

Cena zł 12,00
(VAT 5%)

Indeks 381306
PL ISSN 0043-518X
e-ISSN 2543-8476



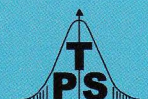
WIADOMOŚCI STATYSTYCZNE

GŁÓWNY
URZĄD
STATYSTYCZNY

POLSKIE
TOWARZYSTWO
STATYSTYCZNE

MIESIĘCZNIK
ROK LXI
WARSZAWA
WRZESIEŃ 2016

Nr **9** (664)



WIADOMOŚCI STATYSTYCZNE

GŁÓWNY
URZĄD
STATYSTYCZNY

POLSKIE
TOWARZYSTWO
STATYSTYCZNE

MIESIĘCZNIK
ROK LXI
WARSZAWA
WRZESIEŃ 2016

Nr **9** (664)

KOLEGIUM REDAKCYJNE

dr Marek Cierpiał-Wolan (redaktor naczelny), dr hab. Andrzej Młodak (zastępca redaktora naczelnego), mgr Renata Bielak, dr Jacek Kowalewski, dr Jan Kubacki, mgr Władysław Wiesław Łagodziński, dr Grażyna Marciniak, dr Stanisław Paradysz, dr hab. Mateusz Pipień, prof. dr hab. Bogdan Stefanowicz, dr Wioletta Wrzaszcz, dr inż. Agnieszka Zgierska

Sekretarz: Alina Świdarska

RADA NAUKOWA

dr Halina Dmochowska (przewodnicząca), dr hab. Bożena Balcerzak-Paradowska, prof. dr hab. Czesław Domański, dr hab. Elżbieta Gołata, prof. dr hab. Semen Matkowski, prof. dr hab. Włodzimierz Okrasa, prof. dr hab. Józef Oleński, prof. dr hab. Tomasz Panek, doc. ing. Iveta Stankovicova, prof. dr hab. Józef Zegar

Sekretarz: Justyna Gustyn

REDAKCJA

al. Niepodległości 208, 00-925 Warszawa, gmach GUS, pok. 353, tel. 22 608 32 25

<http://stat.gov.pl/czasopisma/wiadomosci-statystyczne>

Alina Świdarska (a.swiderska@stat.gov.pl)

Elżbieta Grabowska (e.grabowska@stat.gov.pl)

Wersja internetowa jest wersją pierwotną czasopisma.



ZAKŁAD WYDAWNICTW STATYSTYCZNYCH

al. Niepodległości 208, 00-925 Warszawa, tel. 22 608 31 45.

Informacje w sprawach nabywania czasopism tel. 22 608 32 10, 22 608 38 10.

Zbigniew Karpiński (redaktor techniczny), Ewa Krawczyńska (skład i łamanie),

Wydział Korekty pod kierunkiem Bożeny Gorzycy, mgr Andrzej Kajkowski (wykresy).

Indeks 381306

Prenumerata realizowana przez RUCH S.A.

Zamówienia na prenumeratę w wersji papierowej i na e-wydania można składać bezpośrednio na stronie www.prenumerata.ruch.com.pl.

Ewentualne pytania prosimy kierować na adres e-mail: prenumerata@ruch.com.pl lub kontaktując się z Centrum Obsługi Klienta „RUCH” pod numerami: 22 693 70 00 lub 801 800 803 — czynne w dni robocze w godzinach 7⁰⁰—17⁰⁰.

Koszt połączenia według taryfy operatora.

Szanowni Czytelnicy

We wrześniu 2016 r. mija już 60 lat od ukazania się pierwszego numeru „Wiadomości Statystycznych”. Warto podkreślić, że sam tytuł ma dłuższą historię. Sięga bowiem okresu międzywojennego, kiedy publikacja o takiej nazwie ukazywała się kilka razy w miesiącu i zawierała jedynie wybrane informacje liczbowe o stanie i rozwoju kraju. „Wiadomości Statystyczne” stanowią także swoistego rodzaju kontynuację czasopisma „Kwartalnik Statystyczny”, publikowanego w latach 20. i 30. ub. wieku.

Pierwszym redaktorem naczelnym „Wiadomości Statystycznych” był profesor Władysław Welfe — jeden z najwybitniejszych polskich autorytetów naukowych w dziedzinie statystyki i ekonometrii, znany i ceniony w światowych gremiach naukowych. Kolejnymi redaktorami naczelnymi byli: Michał Szymanowski, Izidor Hrab, Stanisław Gajda, Henryk Białczyński, Stanisław Róg, Tadeusz Walczak i Stanisław Paradysz.

W całym okresie funkcjonowania „Wiadomości Statystycznych” głównym zamierzeniem czasopisma było nie tylko służyć rozwojowi teorii badań statystycznych, ale przede wszystkim praktyczne jej wykorzystanie. Oznaczało to większą integrację środowisk akademickich ze statystyką publiczną. Najlepszym tego przykładem jest wspólne wydawanie od 1989 r. „Wiadomości Statystycznych” przez Główny Urząd Statystyczny oraz Polskie Towarzystwo Statystyczne.

„Wiadomości Statystyczne” odegrały istotną rolę w modernizacji statystyki polskiej w okresie transformacji, szczególnie wtedy stymulowały i wspomagały rozwój poszczególnych dziedzin badań statystycznych, a także przyczyniły się do podnoszenia jakości w wielu działach statystyki publicznej.

Postępujące równocześnie integracja i dezintegracja współczesnego świata i związana z nimi złożoność rozwoju społeczno-gospodarczego powodują, szczególnie dzisiaj, potrzebę doskonalenia badań i analiz statystycznych. Zatem systematyzująca i integrująca rola „Wiadomości Statystycznych” nabiera szczególnego znaczenia.

Realizując nowe wyzwania przed którymi stoi czasopismo, nie możemy zapominać o konieczności zapewnienia odpowiedniego poziomu naukowego „Wiadomości Statystycznych”, które od 60 lat są rozpoznawalne na coraz bardziej konkurencyjnym rynku wydawniczym.

dr Marek Cierpiał-Wolan
Redaktor Naczelny

SPIS TREŚCI

<i>Marek Cierpiał-Wolan</i> — Szanowni Czytelnicy	3
---	---

STUDIA METODOLOGICZNE

<i>Jacek Wesołowski, Jakub Tarczyński</i> — Podstawy matematyczne technik imputacyjnych	7
---	---

<i>Anna Prażmo, Joanna Wójcik, Magdalena Żero</i> — Wyzwania statystyki publicznej w świetle Agendy na rzecz Zrównoważonego Rozwoju 2030	55
--	----

<i>Olga Leszczyńska-Luberek</i> — Zmiany w ujmowaniu zobowiązań emerytalno-rentowych w rachunkach narodowych	69
--	----

<i>Dariusz Kotlewski, Mirosław Błazej</i> — Metodologia rachunku produktywności KLEMS i jego implementacja w warunkach polskich	86
---	----

EDUKACJA STATYSTYCZNA

<i>Czesław Domański</i> — Wyzwania w zakresie edukacji statystycznej społeczeństwa	109
--	-----

INFORMACJE. PRZEGLĄDY. RECENZJE

Marek Hetmański: <i>Świat informacji</i> , 230 stron, Wydawnictwo Dyfin, Warszawa 2015 (oprac. <i>Bogdan Stefanowicz</i>)	116
--	-----

Wydawnictwa GUS — sierpień 2016 r. (oprac. <i>Justyna Gustyn</i>)	119
--	-----

CONTENTS

<i>Marek Cierpiat-Wolan</i> — Dear readers	3
--	----------

METHODOLOGICAL STUDIES

<i>Jacek Wesółowski, Jakub Tarczyński</i> — Basic mathematical imputation techniques	7
--	----------

<i>Anna Prażmo, Joanna Wójcik, Magdalena Żero</i> — The challenges of official statistics to the Agenda for Sustainable Development in 2030 ...	55
---	-----------

<i>Olga Leszczyńska-Luberek</i> — Changes in the treatment of pension liabilities in the national accounts	69
--	-----------

<i>Dariusz Kotlewski, Mirosław Błazej</i> — The methodology of KLEMS productivity accounting and its implementation in Polish conditions	86
--	-----------

STATISTICAL EDUCATION

<i>Czesław Domański</i> — Challenges to statistical education of the society	109
---	------------

INFORMATION. REVIEWS. COMMENTS

Marek Hetmański: <i>World of information</i> , 230 pages, Publiser Difin, Warsaw 2015 (by <i>Bogdan Stefanowicz</i>)	116
---	------------

Publications of the CSO of Poland in August 2016 (by <i>Justyna Gustyn</i>) ...	119
--	------------

СОДЕРЖАНИЕ

<i>Марек Церпял-Волян</i> — Дорогие читатели	3
--	----------

МЕТОДОЛОГИЧЕСКИЕ ИЗУЧЕНИЯ

<i>Яцек Весоловски, Якуб Тарчиньски</i> — Математические основы импутационных методов	7
---	----------

<i>Анна Пражмо, Йоанна Вуйчик, Магдалена Жеро</i> — Вызовы стоящие перед официальной статистикой в свете Агенды для сбалансированного развития 2030	55
---	-----------

<i>Ольга Лециньска-Люберек</i> — Изменения в учете пенсионных обязательств и обязательств по инвалидности в национальных счетах	69
---	-----------

<i>Дарюш Котлевски, Мирослав Блажей</i> — Методология использования счета производительности KLEMS и реализация в польских условиях	86
---	-----------

СТАТИСТИЧЕСКОЕ ОБУЧЕНИЕ

<i>Чеслав Доманьски</i> — Ожидания общества в области статистического обучения	109
--	------------

ИНФОРМАЦИИ. ОБЗОРЫ. РЕЦЕНЗИИ

<i>Марек Хетманьски: Мир информации, 230 страниц, Издательство Дифин, Варшава 2015 (разраб. Богдан Стефанович)</i>	116
--	------------

<i>Публикации ЦСУ — август 2016 г. (разраб. Юстына Густын)</i>	119
--	------------

STUDIA METODOLOGICZNE

Jacek WESOŁOWSKI
Jakub TARCZYŃSKI

Podstawy matematyczne technik imputacyjnych

Streszczenie. *W artykule przedstawiono podstawy metodologii imputacyjnej (w tym metodologii wielokrotnej imputacji), koncentrując się na wyjaśnieniu matematycznej strony zagadnień. Analizowano sytuację, gdy obserwacje tworzące pierwotną próbkę są niezależnymi zmiennymi losowymi o jednakowym rozkładzie, a braki odpowiedzi pojawiają się losowo w sposób niezależny od obserwacji. W szczególności wskazano na problemy pojawiające się, gdy w imputacji wielokrotnej stosowany jest standardowy estymator Rubina wariancji estymatora wielokrotnej imputacji i wskazano na możliwe ulepszenie tego popularnego estymatora. Punktem wyjścia analiz jest sytuacja, gdy za pojawianie się braków odpowiedzi odpowiada mechanizm deterministyczny.*

Słowa kluczowe: imputacja, imputacja wielokrotna, estymator imputacyjny, estymator Rubina, imputacja średnią, imputacja typu hot-deck, imputacja regresyjna.

Słowem imputacja określa się zespół metod uzupełniania braków w próbie w taki sposób, aby zastosowanie estymatorów standardowych, tzn. takich, które byłyby użyte, gdyby w próbie nie było braków obserwacji, prowadziło do efektywnej procedury estymacyjnej.

Literatura dotycząca technik imputacyjnych jest bardzo obszerna. Składają się na nią setki specjalistycznych artykułów w czasopismach statystycznych, artykuły przeglądowe, np. Andridge i Little (2010), Donders, van der Heijden, Stijnen i Moons (2006) czy Norazian Ramli, Yahaya, Ramli i Yusof (2013) oraz kilka większych publikacji monograficznych, takich jak: Rubin (1987), Little i Rubin

(2002), de Waal, Pannekoek i Scholtus (2011) czy Garson (2012) lub van Buuren (2012). Literatura ta najczęściej dotyczy opisu konkretnych metod, ich zastosowań i porównań. Techniki te bywają dostępne w pakietach statystycznych (Horton i Lipsitz, 2001; van Buuren i Groothuis-Oudshoorn, 2011; Misztal, 2012). We wszystkich tych publikacjach stosunkowo mały nacisk kładziony jest na podstawy teoretyczne.

Okazuje się, że w standardowym modelu statystycznym niezależnych obserwacji o jednakowych rozkładach daje się precyzyjnie opisać (w formie twierdzeń matematycznych) podstawy teoretyczne technik imputacyjnych. Przedstawiane opracowanie jest próbą takiego właśnie systematycznego i matematycznego uściślenia uniwersalnych paradygmatów imputacji. W rozdziale I rozważana jest uproszczona sytuacja, gdy zbiór braków obserwacji jest ustalony. Przy tym założeniu wyprowadzono wzory ogólne na estymator imputacyjny średniej i jego parametry oraz zastosowano je w czterech podstawowych schematach imputacyjnych: imputacji średnią, dwóch procedurach typu hot-deck oraz imputacji regresyjnej. W rozdziale II analizowano imputację wielokrotną w modelu z rozdziału I, a następnie opracowany schemat ogólny wykorzystano w sytuacji, gdy w każdym kroku procedury stosowana jest imputacja typu hot-deck. Używany w literaturze termin „imputacja hot-deck”, w tym opracowaniu — nieco rozszerzony — odnosi się do sytuacji, gdy elementy imputowane są zgodnie z pewną procedurą losową, która zależna jest od obserwowanej części próbki.

W rozdziałach III i IV rozważany jest bardziej realistyczny losowy model pojawiania się braków odpowiedzi, tzn. taki, w którym każda obserwacja może być obecna w próbie lub nie, przy czym zdarzenie to ma charakter losowy, czyli można mu przypisać pewne prawdopodobieństwo. Poprzez zastosowanie techniki warunkowania, rezultaty otrzymane w rozdziałach III i IV można stosunkowo prosto wywieść z wyników przedstawionych w rozdziałach I i II. W takim podejściu istotną rolę odgrywa przyjęte w artykule założenie o niezależności statystycznej obserwacji i braków odpowiedzi. W literaturze mówi się w takiej sytuacji o brakach odpowiedzi typu MCAR (*missing completely at random*), np. Little i Rubin (2002). To podejście jest ważne, bo choć często nie jest właściwym modelem dla całej populacji, to okazuje się dobrym przybliżeniem sytuacji rzeczywistej w odpowiednio wybranych podpopulacjach (chodzi o klasy elementów podobnych pod względem skłonności do udzielania odpowiedzi).

Należy zwrócić uwagę na, prezentowane w opracowaniu, ulepszenie klasycznego estymatora Rubina służącego do estymacji wariancji estymatora imputacyjnego w przypadku imputacji wielokrotnej. Zasadnicza część artykułu ma matematyczną strukturę twierdzeń i dowodów. Najciekawsze wyniki matematyczne (w rozdziałach II i IV) zilustrowano eksperymentami numerycznymi. Konkluzje przedstawiono w rozdziale V.

I. DETERMINISTYCZNY ZBIÓR BRAKÓW OBSERWACJI — IMPUTACJA JEDNOKROTNA

I.0. Ogólny schemat imputacyjny

Zmienne losowe $X_1, \dots, X_n$ stanowią pierwotną próbkę, $\mathbf{X} = (X_1, \dots, X_n)$. Niech $\hat{\theta} = g(\mathbf{X})$ oznacza estymator parametru θ , gdzie g jest odpowiednią funkcją n -argumentową. Rozważać będziemy sytuację, gdy obserwowane są tylko zmienne X_i , $i \in R \subsetneq \{1, \dots, n\}$, czyli R jest zbiorem indeksów obserwowanych elementów próbki pierwotnej. W konsekwencji pierwotny wektor obserwacji dzieli się naturalnie na dwa wektory: wektor odpowiedzi obserwowanych $\mathbf{X}_R = (X_i, i \in R)$ oraz wektor odpowiedzi nieobserwowanych $\mathbf{X}_{R^c} = (X_i, i \in R^c = \{1, \dots, n\} \setminus R)$. Taką niepełną próbkę $\mathbf{X}_R$ można uzupełnić przyjmując $\tilde{\mathbf{X}} = (\tilde{X}_1, \dots, \tilde{X}_n)$, gdzie $\tilde{X}_i = X_i$ dla $i \in R$ są zmiennymi obserwowanymi oraz $\tilde{X}_i$, $i \in R^c$ są tzw. zmiennymi imputowanymi, które mogą zależeć (i często zależą) od zmiennych obserwowanych. Często zmienne są imputowane pewnymi funkcjami tych zmiennych, ale zdarzają się też sytuacje, w których nie ma bezpośredniej zależności funkcyjnej, a zmienne obserwowane wpływają jedynie na rozkład prawdopodobieństwa zmiennych imputowanych.

Definicja I.0.1. Imputacyjną wersją estymatora $\hat{\theta} = g(\mathbf{X})$ parametru θ nazywamy estymator postaci

$$\tilde{\theta}_{imp} = g(\tilde{\mathbf{X}}).$$

W skrócie mówimy, że $\tilde{\theta}_{imp}$ jest estymatorem imputacyjnym parametru θ .

Zakładamy, że zmienne $X_1, \dots, X_n$ są niezależne i mają jednakowe rozkłady, takie jak pewna zmienna losowa X o wartości oczekiwanej $E X = \mu$ i wariancji $\text{Var } X = \sigma^2$, które są nieznane.

Do estymacji średniej $\theta = \mu$ standardowo stosowany jest estymator

$$\hat{\mu} = g(\mathbf{X}) = \bar{X} = \frac{1}{n} \sum_{i=1}^n X_i.$$

Jest to estymator nieobciążony, czyli $E \hat{\mu} = \mu$, a jego wariancja wynosi $\text{Var } \hat{\mu} = \frac{\sigma^2}{n}$.

Do estymacji parametru σ^2 standardowo stosowany jest estymator

$$S^2 = S^2(X) = \frac{1}{n-1} \sum_{i=1}^n (X_i - g(X))^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2.$$

Jest to estymator nieobciążony, czyli $E S^2 = \sigma^2$.

Zatem nieobciążony estymator $d^2(\hat{\mu})$ wariancji $\text{Var } \hat{\mu}$ estymatora $\hat{\mu}$ ma postać $d^2(\hat{\mu}) = \frac{1}{n} S^2$.

Z centralnego twierdzenia granicznego (i tzw. twierdzenia Śluckiego) wynika, że $\frac{\hat{\mu} - \mu}{S} \sqrt{n}$ ma asymptotycznie rozkład standardowy normalny, co pozwala na wnioskowanie o przybliżonej precyzji estymatorów, rozumianej jako przedział ufności na zadanym poziomie ufności.

Zadanie imputacji pojawia się, gdy obserwowane są tylko zmienne losowe $X_i, i \in R \subsetneq \{1, \dots, n\}$. Polega ono na uzupełnieniu obserwowanej części próbki o część nieobserwowaną $\tilde{X}_i, i \in R^c$. Niech r oznacza liczbę obserwowanych elementów próbki, czyli $r = \#(R)$. Dla potrzeb rozdziałów I i II zakładamy, że podzbiór $R \subsetneq \{1, \dots, n\}$, a zatem i jego liczebność r jest ustalona. W kolejnych rozdziałach odejmiemy od tego założenia.

Nowa próbka ma postać $\tilde{X} = (\tilde{X}_1, \dots, \tilde{X}_n)$, przy czym zakłada się, że

$$E \tilde{X}_i = \tilde{\mu}_i = \tilde{\mu} \quad \text{oraz} \quad \text{Var } \tilde{X}_i = \tilde{\sigma}_i^2 = \tilde{\sigma}^2, \quad \text{jeśli } i \in R^c$$

oraz

$$\text{Corr}(\tilde{X}_i, \tilde{X}_j) = \tilde{\rho}_{ij} = \tilde{\rho}, \quad \text{jeśli } i, j \in R^c, i \neq j,$$

$$\text{Corr}(X_i, \tilde{X}_j) = \rho_{ij} = \rho, \quad \text{jeśli } i \in R, j \in R^c.$$

Zajmiemy się teraz, zgodnie z Definicją I.0.1, imputacyjnym estymatorem średniej μ .

TWIERDZENIE I.0.1. *Imputacyjny estymator średniej μ ma postać*

$$\hat{\mu}_{\text{Imp}} = \frac{r \bar{X}_R + (n-r) \tilde{\bar{X}}}{n}, \quad (\text{I.0.1})$$

gdzie $\bar{X}_R = \frac{1}{r} \sum_{i \in R} X_i$ oraz $\tilde{\bar{X}} = \frac{1}{n-r} \sum_{i \in R^c} \tilde{X}_i$.

Jego wartość oczekiwana wynosi

$$E \hat{\mu}_{Imp} = \frac{r}{n} \mu + \frac{n-r}{n} \tilde{\mu}, \quad (I.0.2)$$

a jego obciążenie wynosi

$$B \hat{\mu}_{Imp} = \frac{n-r}{n} (\tilde{\mu} - \mu),$$

więc estymator $\hat{\mu}_{Imp}$ jest nieobciążony tylko gdy $\tilde{\mu} = \mu$.

Dowód. Z postaci funkcji g występującej w definicji estymatora $\hat{\mu}$ wynika, że

$$\hat{\mu}_{Imp} = g(\tilde{X}) = \frac{1}{n} \left(\sum_{i \in R} X_i + \sum_{i \in R^c} \tilde{X}_i \right),$$

co jest równoważne z (I.0.1).

Wzór (I.0.1) implikuje

$$E \hat{\mu}_{Imp} = \frac{1}{n} \left(r E \bar{X}_R + (n-r) E \bar{\tilde{X}} \right).$$

Wzór (I.0.2) wynika z tej równości, ponieważ $E \bar{X}_R = \mu$, a $E \bar{\tilde{X}} = \tilde{\mu}$.

Wzór na obciążenie jest natychmiastową konsekwencją definicji obciążenia $B \hat{\mu}_{Imp} = E \hat{\mu}_{Imp} - \mu$ i wzoru (I.0.2). ■

TWIERDZENIE I.0.2. *Wariancja imputacyjnego estymatora średniej $\hat{\mu}_{Imp}$ ma postać*

$$\text{Var } \hat{\mu}_{Imp} = \frac{r\sigma^2 + (n-r)\tilde{\sigma}^2 + (n-r)(n-r-1)\tilde{\rho}\tilde{\sigma}^2 + 2r(n-r)\rho\sigma\tilde{\sigma}}{n^2} \quad (I.0.3)$$

Dowód. Z postaci estymatora $\hat{\mu}_{Imp}$ i podstawowych własności wariancji otrzymujemy

$$\text{Var } \hat{\mu}_{Imp} = \frac{1}{n^2} \left(r^2 \text{Var } \bar{X}_R + (n-r)^2 \text{Var } \bar{\tilde{X}} + 2r(n-r) \text{Cov} \left(\bar{X}_R, \bar{\tilde{X}} \right) \right). \quad (I.0.4)$$

$$\text{Ale } \text{Var } \bar{X}_R = \frac{\sigma^2}{r}.$$

Z kolei

$$\begin{aligned}\text{Var} \bar{\tilde{X}} &= \frac{1}{(n-r)^2} \left(\sum_{i \in R^c} \text{Var} \tilde{X}_i + \sum_{i, j \in R^c, i \neq j} \text{Cov}(\tilde{X}_i, \tilde{X}_j) \right) = \\ &= \frac{1}{(n-r)^2} \left((n-r) \tilde{\sigma}^2 + (n-r)(n-r-1) \tilde{\rho} \tilde{\sigma}^2 \right).\end{aligned}$$

A zatem

$$\text{Var} \bar{\tilde{X}} = \frac{1 + (n-r-1) \tilde{\rho}}{n-r} \tilde{\sigma}^2. \quad (\text{I.0.5})$$

Natomiast

$$\text{Cov}(\bar{X}_R, \bar{\tilde{X}}) = \frac{1}{r(n-r)} \sum_{j \in R} \sum_{i \in R^c} \text{Cov}(X_j, \tilde{X}_i) = \rho \sigma \tilde{\sigma}. \quad (\text{I.0.6})$$

Wstawiając te trzy wyrażenia do wzoru (I.0.4) otrzymujemy

$$\text{Var} \hat{\mu}_{Imp} = \frac{1}{n^2} \left(r \sigma^2 + (n-r) (1 + (n-r-1) \tilde{\rho}) \sigma^2 + 2r(n-r) \rho \sigma \tilde{\sigma} \right)$$

Ten wzór jest równoważny wzorowi (I.0.3). ■

Imputacyjna estymacja wariancji polega na wykorzystaniu standardowego estymatora wariancji do próbki imputowanej $\tilde{X}$, czyli zgodnie z definicją I.0.1

$$S_{Imp}^2 = S^2(\tilde{X}) = \frac{1}{n-1} \left(\sum_{j \in R} (X_j - \hat{\mu}_{Imp})^2 + \sum_{i \in R^c} (\tilde{X}_i - \hat{\mu}_{Imp})^2 \right). \quad (\text{I.0.7})$$

TWIERDZENIE I.0.3. *Imputacyjny estymator wariancji S_{Imp}^2 ma postać*

$$S_{Imp}^2 = \frac{1}{n-1} \left((r-1) S_R^2 + (n-r-1) \tilde{S}^2 + \frac{r(n-r)}{n} (\bar{X}_R - \bar{\tilde{X}})^2 \right), \quad (\text{I.0.8})$$

gdzie $S_R^2 = \frac{1}{r-1} \sum_{j \in R} (X_j - \bar{X}_R)^2$ oraz $\tilde{S}^2 = \frac{1}{n-r-1} \sum_{i \in R^c} (\tilde{X}_i - \bar{\tilde{X}})^2$.

Dowód. Ze wzoru (I.0.7) i definicji $\hat{\mu}_{Imp}$ mamy

$$S_{Imp}^2 = \frac{1}{n-1} \left(\sum_{i \in R} \left(X_i - \frac{1}{n} \left(r \bar{X}_R + (n-r) \tilde{X} \right) \right)^2 + \sum_{j \in R^c} \left(\tilde{X}_j - \tilde{X} + \frac{r}{n} \left(\tilde{X} - \bar{X}_R \right) \right)^2 \right).$$

Wykonując kwadratowanie pod obiema sumami i korzystając z faktu, że $\sum_{i \in R} (X_i - \bar{X}_R) = 0$ oraz $\sum_{i \in R^c} (\tilde{X}_i - \tilde{X}) = 0$ otrzymujemy

$$\begin{aligned} S_{Imp}^2 &= \frac{1}{n-1} \left(\sum_{i \in R} (X_i - \bar{X}_R)^2 + r \frac{(n-r)^2}{n^2} (\bar{X}_R - \tilde{X})^2 + \sum_{i \in R^c} (\tilde{X}_i - \tilde{X})^2 + (n-r) \frac{r^2}{n^2} (\bar{X}_R - \tilde{X})^2 \right) = \\ &= \frac{1}{n-1} \left((r-1) S_R^2 + (n-r-1) \tilde{S}^2 + \left(r \frac{(n-r)^2}{n^2} + (n-r) \frac{r^2}{n^2} \right) (\bar{X}_R - \tilde{X})^2 \right), \end{aligned}$$

co kończy dowód. ■

TWIERDZENIE I.0.4. *Wartość oczekiwana imputacyjnego estymatora wariancji S_{Imp}^2 ma postać*

$$\begin{aligned} ES_{Imp}^2 &= \frac{r}{n} \sigma^2 + \frac{n-r}{n} \tilde{\sigma}^2 - \frac{(n-r)(n-r-1)}{n(n-1)} \tilde{\rho} \tilde{\sigma}^2 - \\ &+ 2 \frac{r(n-r)}{n(n-1)} \rho \sigma \tilde{\sigma} + \frac{r(n-r)}{n(n-1)} (\mu - \tilde{\mu})^2. \end{aligned} \quad (I.0.9)$$

Dowód. Mamy $ES_R^2 = \sigma^2$ oraz $E\tilde{S}^2 = \sigma^2(1 - \tilde{\rho})$. Pierwszy fakt jest powszechnie znany, a drugi wynika z następującego rachunku

$$\begin{aligned} E\tilde{S}^2 &= E \left(\frac{1}{n-r-1} \sum_{i \in R^c} \tilde{X}_i^2 - \frac{1}{(n-r-1)(n-r)} \left(\sum_{i \in R^c} \tilde{X}_i \right)^2 \right) = \frac{1}{n-r-1} \sum_{i \in R^c} E\tilde{X}_i^2 - \\ &+ \frac{1}{(n-r-1)(n-r)} \left(\sum_{i \in R^c} E\tilde{X}_i^2 + \sum_{i, j \in R^c, i \neq j} E\tilde{X}_i \tilde{X}_j \right) = \frac{n-r}{n-r-1} (\tilde{\sigma}^2 + \tilde{\mu}^2) - \\ &+ \frac{1}{(n-r-1)(n-r)} ((n-r)(\tilde{\sigma}^2 + \tilde{\mu}^2) + (n-r)(n-r-1)(\tilde{\rho} \tilde{\sigma}^2 + \tilde{\mu}^2)) = \tilde{\sigma}^2(1 - \tilde{\rho}). \end{aligned}$$

Co więcej

$$\begin{aligned} E\left(\bar{X}_R - \tilde{X}\right)^2 &= E\bar{X}_R^2 + E\tilde{X}^2 - 2E\bar{X}_R\tilde{X} = \text{Var}\bar{X}_R + \left(E\bar{X}_R\right)^2 + \\ &+ \text{Var}\tilde{X} + \left(E\tilde{X}_R\right)^2 - 2\text{Cov}\left(\bar{X}_R, \tilde{X}\right) - 2E\bar{X}_RE\tilde{X}. \end{aligned}$$

Wykorzystując wzory (I.0.5) oraz (I.0.6) otrzymujemy

$$E\left(\bar{X}_R - \tilde{X}\right)^2 = \frac{\sigma^2}{r} + \frac{\tilde{\sigma}^2}{n-r} + \frac{n-r-1}{n-r}\tilde{\rho}\tilde{\sigma}^2 - 2\rho\sigma\tilde{\sigma} + (\mu - \tilde{\mu})^2.$$

Ostatecznie ze wzoru (I.0.8) i powyżej uzyskanych tożsamości wynika, że

$$\begin{aligned} (n-1)ES_{Imp}^2 &= (r-1)\sigma^2 + (n-r-1)\tilde{\sigma}^2(1-\tilde{\rho}) + \\ &+ \frac{r(n-r)}{n}\left(\frac{\sigma^2}{r} + \frac{\tilde{\sigma}^2}{n-r} + \frac{n-r-1}{n-r}\tilde{\rho}\tilde{\sigma}^2 - 2\rho\sigma\tilde{\sigma} + (\mu - \tilde{\mu})^2\right), \end{aligned}$$

co po uproszczeniu daje wzór (I.0.9).■

1.1. Imputacja średnią

Technika ta polega na przypisaniu każdemu nieobserwowanemu elementowi próbki średniej z części obserwowanej próbki, czyli

$$\tilde{X}_i = \frac{1}{r} \sum_{j \in R} X_j = \bar{X}_R, \quad i \in R^c. \quad (\text{I.1.1})$$

TWIERDZENIE I.1.1. *Imputacyjny estymator wartości oczekiwanej w przypadku imputacji średnią ma postać*

$$\hat{\mu}_{Imp} = \bar{X}_R. \quad (\text{I.1.2})$$

Jego wartość oczekiwana oraz wariancja wynoszą, odpowiednio

$$E\hat{\mu}_{Imp} = \mu \quad \text{oraz} \quad \text{Var}\hat{\mu}_{Imp} = \frac{\sigma^2}{r}. \quad (\text{I.1.3})$$

Dowód. Ze wzoru (I.1.1) wynika, że $\tilde{X} = \bar{X}_R$. Zatem wzór (I.0.1) prowadzi natychmiast do (I.1.2). Ponieważ $\tilde{\mu} = E\tilde{X}_i = E\bar{X}_R = \mu$, więc (I.0.2) daje pierwszy ze wzorów (I.1.3).

Dodatkowo z (I.1.1) wynika, że $\tilde{\sigma}^2 = \text{Var } \tilde{X}_i = \text{Var } \bar{X}_R = \frac{\sigma^2}{r}$ oraz

$$\tilde{\rho} = \text{Corr}(\tilde{X}_i, \tilde{X}_j) = \text{Corr}(\bar{X}_R, \bar{X}_R) = 1,$$

$$\rho = \text{Corr}(X_i, \tilde{X}_j) = \text{Corr}(X_i, \bar{X}_R) = \sqrt{r} \frac{\text{Cov}(X_i, \bar{X}_R)}{\sigma^2} = \sqrt{r} \frac{\frac{\sigma^2}{\sqrt{r}}}{\sigma^2} = \frac{1}{\sqrt{r}}.$$

Zatem zgodnie ze wzorem (I.0.3)

$$\text{Var } \hat{\mu}_{Imp} = \frac{r\sigma^2 + (n-r)\frac{\sigma^2}{r} + (n-r)(n-r-1)\frac{\sigma^2}{r} + 2r(n-r)\frac{1}{\sqrt{r}}\sigma\frac{\sigma}{\sqrt{r}}}{n^2}.$$

Po uproszczeniu otrzymujemy drugi ze wzorów (I.1.3).■

Oczywiście wzory (I.1.3) wynikają wprost z postaci (I.1.2) estymatora imputacyjnego. W dowodzie, wyprowadzając je odpowiednio ze wzorów (I.0.2) oraz (I.0.3), chcieliśmy zaznaczyć uniwersalność proponowanego podejścia.

TWIERDZENIE I.1.2. *Imputacyjny estymator wariancji w przypadku imputacji średnią ma postać*

$$S_{Imp}^2 = \frac{r-1}{n-1} S_R^2, \quad (\text{I.1.4})$$

a jego wartość oczekiwana wynosi

$$\text{ES}_{Imp}^2 = \frac{r-1}{n-1} \sigma^2. \quad (\text{I.1.5})$$

Nieobciążony estymator wariancji estymatora $\hat{\mu}_{Imp}$ ma postać

$$v^2(\hat{\mu}_{Imp}) = \frac{n-1}{r(r-1)} S_{Imp}^2. \quad (\text{I.1.6})$$

Dowód. Zgodnie ze wzorem (I.1.1) mamy $\tilde{S}^2 = 0$ oraz $\tilde{X} = \bar{X}_R$. Zatem (I.1.4) wynika bezpośrednio ze wzoru (I.0.8).

Natomiast zgodnie ze wzorem (I.0.9) na wartość oczekiwaną imputacyjnego estymatora wariancji mamy

$$ES_{imp}^2 = \frac{r}{n} \sigma^2 + \frac{n-r}{n} \frac{\sigma^2}{r} - \frac{(n-r)(n-r-1)}{n(n-1)} \frac{\sigma^2}{r} - 2 \frac{r(n-r)}{n(n-1)} \frac{1}{\sqrt{r}} \sigma \frac{\sigma}{\sqrt{r}},$$

skąd po uproszczeniach otrzymuje się wzór (I.1.5).

Wzór (I.1.6) wynika z porównania wzoru (I.1.5) i drugiego ze wzorów (I.1.3).■

Wzór (I.1.5) wynika również bezpośrednio z postaci (I.1.4), ponieważ S_R^2 jest nieobciążonym estymatorem wariancji σ^2 . Wyprowadziliśmy go jednak ze wzoru (I.0.9), aby, podobnie jak w przypadku wzorów (I.0.2) i (I.0.3) (użytych do wyprowadzenia (I.1.3)), wskazać na jego uniwersalność.

Zauważmy, że ze wzoru (I.1.5) wynika, iż estymator S_{imp}^2 jest obciążonym estymatorem wariancji.

I.2. Imputacja typu *hot-deck*

I.2.1. Losowanie spośród respondentów

Technika ta polega na przypisaniu każdemu nieobserwowanemu elementowi próbki elementu wylosowanego spośród elementów obserwowanych. W najprostszej sytuacji jest to losowanie proste ze zwracaniem.

Niech K_i będzie numerem elementu wylosowanym dla i -tego nierespondenta, $i \in R^c$. Wtedy zmienne losowe K_i , $i \in R^c$, są niezależne i mają ten sam rozkład $P(K_i = j) = \frac{1}{r}$, $j \in R$, $i \in R^c$ oraz wektory losowe $\mathbf{K} = (K_i, i \in R^c)$ i $\mathbf{X}_R = (X_j, j \in R)$ są niezależne.

W konsekwencji

$$\tilde{X}_i = X_{K_i}, \quad i \in R^c. \quad (\text{I.2.1})$$

Uwaga. W przypadku obserwacji dwuwymiarowych $((X_j, Y_j), j \in R, Y_i, i \in R^c)$ stosowana bywa metoda wyboru najbliższego sąsiada. Polega ona na ustaleniu dla każdego $i \in R^c$ takiego $K_i \in R$, dla którego odległość $|Y_{K_i} - Y_i|$ jest najmniejsza spośród $|Y_j - Y_i|$, $j \in R$. Przy założeniu, że wektory losowe $\mathbf{w}_i := (X_i, Y_i)$, $i = 1, \dots, n$, są niezależne i mają jednakowe rozkłady (takie jak rozkład pary (X, Y)), zmienne losowe K_i , $i \in R^c$, mają rozkład jednostajny tak jak w opisanym wyżej modelu, natomiast wektory losowe $\mathbf{K}$ i $\mathbf{X}_R$ nie są niezależne. Zwracam

uwagę, że założenie niezależności dotyczy par, tzn. wektory losowe $\mathbf{w}_i$, $i = 1, \dots, n$, są łącznie niezależne, natomiast dla każdego ustalonego $i \in \{1, \dots, n\}$ zmienne losowe X_i oraz Y_i (będące składowymi wektora $\mathbf{w}_i$) nie muszą być niezależne; najlepiej, gdy są np. dobrze skorelowane. Za takim podejściem przemawia następujące uzasadnienie: jeżeli obserwowane wartości Y_j i Y_i , gdzie $i \in R^c$, $j \in R$, są bliskie, to przy znacznej korelacji zmiennych X i Y można oczekiwać, że nieobserwowana wartość X_i jest bliska obserwowanej wartości X_j .

TWIERDZENIE 1.2.1. *Imputacyjny estymator wartości oczekiwanej w przypadku imputacji typu hot-deck polegającej na losowaniu respondentów ma postać*

$$\hat{\mu}_{Imp} = \frac{1}{n} \left(\sum_{j \in R} X_j + \sum_{i \in R^c} X_{K_i} \right). \quad (I.2.2)$$

Jego wartość oczekiwana oraz wariancja wynoszą, odpowiednio

$$E \hat{\mu}_{Imp} = \mu \quad \text{oraz} \quad \text{Var} \hat{\mu}_{Imp} = \frac{\sigma^2}{r} \left(1 + \frac{(n-r)(r-1)}{n^2} \right). \quad (I.2.3)$$

Dowód. Zgodnie z (I.2.1) mamy $\widetilde{X} = \frac{1}{n-r} \sum_{i \in R^c} X_{K_i}$, a zatem wzór (I.0.1) implikuje (I.2.2).

Zauważmy, że dla dowolnego $i \in R^c$

$$E(X_{K_i} | \mathbf{K}) = E X = \mu \quad \text{oraz} \quad \text{Var}(X_{K_i} | \mathbf{K}) = \text{Var}(X) = \sigma^2$$

i w konsekwencji

$$\hat{\mu} = E \widetilde{X}_i = E X_{K_i} = E E(X_{K_i} | \mathbf{K}) = \mu \quad (I.2.4)$$

oraz

$$\widetilde{\sigma}^2 = \text{Var} X_{K_i} = \text{Var} E(X_{K_i} | \mathbf{K}) + E \text{Var}(X_{K_i} | \mathbf{K}) = \sigma^2. \quad (I.2.5)$$

Ze wzoru (I.2.4) wynika, że $E \widetilde{X} = \mu$, zgodnie ze wzorem (I.0.2), mamy więc pierwszy ze wzorów (I.2.3), czyli estymator $\hat{\mu}_{Imp}$ jest nieobciążonym estymatorem średniej.

Podobnie dla dowolnych $i, j \in R^c$ oraz $l \in R$ mamy

$$\text{Cov}(X_{K_i}, X_{K_j} | \mathbf{K}) = \sigma^2 \mathbf{I}(K_i = K_j) \quad \text{oraz} \quad \text{Cov}(X_l, X_{K_j} | \mathbf{K}) = \sigma^2 \mathbf{I}(K_j = l).$$

Stąd wynika, że dla dowolnych $i, j \in R^c$

$$\begin{aligned} \tilde{\rho} &= \text{Corr}(\tilde{X}_i, \tilde{X}_j) = \frac{\text{Cov}(X_{K_i}, X_{K_j})}{\sigma^2} = \\ &= \frac{\text{E Cov}(X_{K_i}, X_{K_j} | \mathbf{K}) + \text{Cov}(\text{E}(X_{K_i} | \mathbf{K}), \text{E}(X_{K_j} | \mathbf{K}))}{\sigma^2}. \end{aligned}$$

Ponieważ $\text{Cov}(\text{E}(X_{K_i} | \mathbf{K}), \text{E}(X_{K_j} | \mathbf{K})) = 0$, z powyższego wzoru otrzymujemy

$$\tilde{\rho} = \text{P}(K_i = K_j) = \frac{1}{r}. \quad (\text{I.2.6})$$

Z kolei dla dowolnych $j \in R$ oraz $i \in R^c$

$$\begin{aligned} \rho &= \text{Corr}(X_j, \tilde{X}_i) = \frac{\text{Cov}(X_j, X_{K_i})}{\sigma^2} = \\ &= \frac{\text{E Cov}(X_j, X_{K_i} | \mathbf{K}) + \text{Cov}(\text{E}(X_j | \mathbf{K}), \text{E}(X_{K_i} | \mathbf{K}))}{\sigma^2}. \end{aligned}$$

Ponieważ $\text{Cov}(\text{E}(X_j | \mathbf{K}), \text{E}(X_{K_i} | \mathbf{K})) = 0$, z powyższego wzoru otrzymujemy

$$\rho = \text{P}(K_i = j) = \frac{1}{r}. \quad (\text{I.2.7})$$

Wykorzystując wzory (I.2.4—I.2.7), zgodnie ze wzorem (I.0.3), mamy

$$\text{Var} \hat{\mu}_{\text{Imp}} = \frac{r\sigma^2 + (n-r)\sigma^2 + (n-r)(n-r-1)\frac{\sigma^2}{r} + 2r(n-r)\frac{\sigma^2}{r}}{n^2}.$$

Po uproszczeniach otrzymujemy drugi ze wzorów (I.2.3). ■

TWIERDZENIE 1.2.2. *Imputacyjny estymator wariancji w przypadku imputacji typu hot-deck polegającej na losowaniu respondentów ma postać*

$$S_{\text{Imp}}^2 = \frac{1}{n-1} \left((r-1)S_R^2 + \sum_{i \in R^c} \left(X_{K_i} - \frac{1}{n-r} \sum_{j \in R^c} X_{K_j} \right)^2 + \frac{r(n-r)}{n} \left(\bar{X}_R - \frac{1}{n-r} \sum_{j \in R^c} X_{K_j} \right)^2 \right),$$

a jego wartość oczekiwana wynosi

$$E S_{imp}^2 = \frac{r-1}{r} \frac{n(n-1)+r}{n(n-1)} \sigma^2. \quad (I.2.8)$$

Nieobciążony imputacyjny estymator wariancji estymatora $\hat{\mu}_{imp}$ ma postać

$$v^2(\hat{\mu}_{imp}) = \frac{n-1}{n(r-1)} \frac{n(n-1)+r+r(n-r)}{n(n-1)+r} S_{imp}^2. \quad (I.2.9)$$

Dowód. Wzór na S_{imp}^2 jest natychmiastową konsekwencją wzoru (I.2.1) oraz ogólnego wzoru na estymator imputacyjny wariancji (I.0.8). Z kolei wstawiając wzory (I.2.4)—(I.2.7) do wzoru (I.0.9) otrzymujemy

$$\begin{aligned} E S_{imp}^2 &= \frac{r}{n} \sigma^2 + \frac{n-r}{n} \sigma^2 - \frac{(n-r)(n-r-1)}{n(n-1)} \frac{\sigma^2}{r} - 2 \frac{r(n-r)}{n(n-1)} \frac{\sigma^2}{r} = \\ &= \left(1 - \frac{(n-r)(n+r-1)}{n(n-1)r} \right) \sigma^2. \end{aligned}$$

Ostatnie wyrażenie jest równoważne ze wzorem (I.2.8).

Wzór (I.2.9) wynika z porównania wzoru (I.2.8) oraz drugiego ze wzorów (I.2.3).■

Zatem dla dużych wartości r mamy $E S_{imp}^2 \approx \sigma^2$, czyli estymator wariancji jest w przybliżeniu nieobciążony.

I.2.2. Losowanie z rozkładu normalnego

Zakładamy, że próbka pochodzi z rozkładu normalnego, czyli X_i , $i \in R$, są niezależnymi zmiennymi losowymi o rozkładzie normalnym $N(\mu, \sigma^2)$. Zmienne imputowane generowane są w sposób niezależny z rozkładu $N(\bar{X}_R, S_R^2)$, tzn. $\tilde{X}_j | \mathbf{X}_R \sim N(\bar{X}_R, S_R^2)$, $j \in R^c$, oraz $\tilde{X}_j, j \in R^c$, są warunkowo niezależne pod warunkiem $\mathbf{X}_R$. Zatem

$$E \tilde{X}_j = E E(\tilde{X}_j | \mathbf{X}_R) = E \bar{X}_R = \mu$$

oraz

$$\text{Var } \tilde{X}_j = \text{Var} E(\tilde{X}_j | \mathbf{X}_R) + E \text{Var}(\tilde{X}_j | \mathbf{X}_R) = \text{Var } \bar{X}_R + E S_R^2 = \sigma^2 \frac{1+r}{r}.$$

Ponieważ

$$\text{Cov}(X_i, \tilde{X}_j | \mathbf{X}_R) = 0 \quad \text{oraz} \quad \text{Cov}(\tilde{X}_i, \tilde{X}_j | \mathbf{X}_R) = 0.$$

więc

$$\text{Cov}(X_i, \tilde{X}_j) = \text{Cov}(E(X_i | \mathbf{X}_R), E(\tilde{X}_j | \mathbf{X}_R)) = \text{Cov}(X_i, \bar{X}_R) = \frac{\sigma^2}{r}$$

oraz

$$\text{Cov}(\tilde{X}_i, \tilde{X}_j | \mathbf{X}_R) = \text{Cov}(E(\tilde{X}_i | \mathbf{X}_R), E(\tilde{X}_j | \mathbf{X}_R)) = \text{Var } \bar{X}_R = \frac{\sigma^2}{r}.$$

TWIERDZENIE 1.2.3. *Dla imputacyjnego estymatora wartości oczekiwanej w przypadku imputacji typu hot-deck polegającej na losowaniu z rozkładu normalnego wartość oczekiwana oraz wariancja wynoszą odpowiednio*

$$E\hat{\mu}_{Imp} = \mu \quad \text{oraz} \quad \text{Var } \hat{\mu}_{Imp} = \frac{\sigma^2}{r} \left(1 + \frac{(n-r)r}{n^2} \right). \quad (1.2.10)$$

Natomiast

$$ES_{Imp}^2 = \left(1 - \frac{n-r}{n(n-1)} \right) \sigma^2. \quad (1.2.11)$$

Nieobciążony estymator wariancji estymatora $\hat{\mu}_{Imp}$ ma postać

$$\hat{v}^2(\hat{\mu}_{Imp}) = \frac{n-1}{rn} \frac{n^2 + (n-r)r}{n(n-1) - (n-r)} S_{Imp}^2. \quad (1.2.12)$$

Dowód. Wzory (1.2.10) są odpowiednio natychmiastową konsekwencją wzorów (1.0.2) i (1.0.3), natomiast wzór (1.2.11) wynika wprost ze wzoru (1.0.9). Wzór (1.2.12) jest z kolei konsekwencją wzoru (1.2.11) i drugiego ze wzorów (1.2.10). ■

1.3. Imputacja regresyjna

Zakładamy, że pary (X_i, Y_i) , $i = 1, 2, \dots, n$, tworzące próbkę pierwotną mają jednakowe rozkłady, takie jak para (X, Y) i są niezależne. Próbka obserwowana ma postać (X_i, Y_i) , $i \in R$, Y_j , $j \in R^c$ czyli zmienna Y -owa jest obserwowana

w całej próbkę pierwotnej, natomiast zmienna X -owa — jedynie dla respondentów. Jeśli znany jest rozkład warunkowy $X|Y$, to można wartość zmiennej losowej $\tilde{X}_j$ generować z rozkładu warunkowego $X_j|Y_j, j \in R^c$. Wtedy własności statystyczne próbki imputowanej $\tilde{X}$ nie różnią się od własności próbki pierwotnej X . Jednak najczęściej rozkład warunkowy może być jedynie estymowany na podstawie obserwacji $(X_i, Y_i), i \in R$. W takiej sytuacji precyzyjny opis własności estymatorów imputacyjnych jest trudny, w szczególności zależy od metody estymacji rozkładu warunkowego.

Zamiast estymacji rozkładu warunkowego można wykorzystać model regresyjny, który jest konstrukcją teoretyczną znajdującą zastosowanie w takiej sytuacji.

Niech $EX = \mu, EY = v, \text{Var}(X) = \sigma^2, \text{Var}(Y) = \tau^2$. Niech χ będzie współczynnikiem korelacji zmiennych X i Y . Niech λ będzie współczynnikiem regresji X względem Y , tzn. $\lambda = \chi \frac{\sigma}{\tau}$.

W najprostszej sytuacji, gdy znane są średnia v oraz współczynnik regresji λ , technika imputacji regresyjnej polega na przypisaniu nierespondentom wartości predyktora regresyjnego, tzn.

$$\tilde{X}_j = \lambda(Y_j - v) + \bar{X}_R, j \in R^c.$$

TWIERDZENIE 1.3.1. *Imputacyjny estymator wartości oczekiwanej w przypadku imputacji regresyjnej ma postać*

$$\hat{\mu}_{Imp} = \bar{X}_R + \frac{n-r}{n} \lambda (\bar{Y}_{R^c} - v), \quad (I.3.1)$$

gdzie $\bar{Y}_{R^c} = \frac{1}{n-r} \sum_{j \in R^c} Y_j$.

Jego wartość oczekiwana oraz wariancja wynoszą odpowiednio

$$E \hat{\mu}_{Imp} = \mu \quad \text{oraz} \quad \text{Var} \hat{\mu}_{Imp} = \sigma^2 \left(\frac{1}{r} + \chi^2 \frac{n-r}{n^2} \right). \quad (I.3.2)$$

Dowód. Ponieważ

$$\tilde{X} = \frac{1}{n-r} \sum_{j \in R^c} \tilde{X}_j = \bar{X}_R + \lambda (\bar{Y}_{R^c} - v), \quad (I.3.3)$$

więc wzór (I.3.1) wynika wprost ze wzoru (I.0.1).

Z kolei

$$E(\bar{Y}_{R^c} - v) = 0.$$

a zatem pierwszy ze wzorów (I.3.2) wynika natychmiast z (I.3.1).

Ponieważ Y_j oraz $\bar{X}_R$ są niezależne dla $j \in R^c$, to

$$\tilde{\sigma}^2 = \text{Var} \tilde{X}_j = \text{Var}(\lambda(Y_j - v)) + \text{Var}(\bar{X}_R) = \sigma^2 \left(\chi^2 + \frac{1}{r} \right).$$

Dla $i \in R, j \in R^c$

$$\text{Cov}(X_i, \tilde{X}_j) = \text{Cov}(X_i, \bar{X}_R) = \frac{\sigma^2}{r},$$

więc

$$\rho = \frac{\frac{\sigma^2}{r}}{\sigma^2 \sqrt{\chi^2 + \frac{1}{r}}} = \frac{1}{\sqrt{r(1 + r\chi^2)}}.$$

Dla $i, j \in R^c$

$$\text{Cov}(\tilde{X}_i, \tilde{X}_j) = \text{Var}(\bar{X}_R) = \frac{\sigma^2}{r},$$

więc

$$\tilde{\rho} = \frac{\frac{\sigma^2}{r}}{\sigma^2 \left(\chi^2 + \frac{1}{r} \right)} = \frac{1}{1 + r\chi^2}$$

Ponieważ $\tilde{\rho} \tilde{\sigma}^2 = \frac{\sigma^2}{r}$ oraz $\rho \sigma \tilde{\sigma} = \frac{\sigma^2}{r}$, więc wstawiając otrzymane wyżej wyrażenia na $\tilde{\sigma}^2$, ρ oraz $\tilde{\rho}$ do wzoru (I.0.3) mamy

$$\begin{aligned} \text{Var} \hat{\mu}_{Imp} &= \frac{r\sigma^2 + (n-r)\sigma^2 \left(\chi^2 + \frac{1}{r} \right) + (n-r)(n-r-1) \frac{\sigma^2}{r} + 2r(n-r) \frac{\sigma^2}{r}}{n^2} = \\ &= \frac{\sigma^2}{n^2} \left(\frac{n^2}{r} + (n-r)\chi^2 \right), \end{aligned}$$

co jest równoważne z drugim ze wzorów (I.3.2). ■

Uwaga. W odróżnieniu od wcześniej rozważanych przypadków w imputacji regresyjnej wartości imputowane zależą od współczynnika regresji λ oraz od wartości oczekiwanej ν . Najczęściej nie są one znane. Wtedy we wzorach na zmienne imputowane, na estymator imputacyjny średniej oraz na estymator imputacyjny wariancji (poniżej) zamiast λ oraz ν należy wstawić ich estymatory otrzymane na podstawie obserwowanej części próbki $(X_i, Y_i), i \in R, Y_j, j \in R^c$.

TWIERDZENIE I.3.2. *Imputacyjny estymator wariancji w przypadku imputacji regresyjnej ma postać*

$$S_{imp}^2 = \frac{1}{n-1} \left((r-1)S_R^2 + \lambda^2 \left((n-r-1)S_{R^c, Y}^2 + \frac{(n-r)r}{n} (\bar{Y}_{R^c} - \nu)^2 \right) \right), \quad (I.3.4)$$

a jego wartość oczekiwana wynosi

$$E S_{imp}^2 = \sigma^2 \left(\frac{r-1}{n-1} + \left(1 - \frac{r}{n} \right) \chi^2 \right). \quad (I.3.5)$$

Nieobciążony imputacyjny estymator wariancji estymatora $\hat{\mu}_{imp}$ ma postać

$$\hat{\nu}^2(\hat{\mu}_{imp}) = \frac{\frac{1}{r} + \frac{n-r}{n^2} \chi^2}{\frac{r-1}{n-1} + \frac{n-r}{n} \chi^2} S_{imp}^2. \quad (I.3.6)$$

Dowód. Ze wzoru (I.3.1) i definicji $\tilde{S}^2$ mamy

$$\tilde{S}^2 = \lambda^2 S_{R^c, Y}^2, \text{ gdzie } S_{R^c, Y}^2 = \frac{1}{n-r-1} \sum_{j \in R^c} (Y_j - \bar{Y}_{R^c})^2.$$

Ponownie wykorzystując wzór (I.3.1) otrzymujemy

$$\left(\bar{X}_R - \bar{\tilde{X}} \right)^2 = \lambda^2 (\bar{Y}_{R^c} - \nu)^2.$$

Wstawiając oba powyższe wyrażenia na $\tilde{S}^2$ oraz na $\left(\bar{X}_R - \bar{\tilde{X}} \right)^2$ do wzoru (I.0.8) otrzymujemy (I.3.4).

Natomiast wstawiając wyrażenia

$$\tilde{\sigma}^2 = \sigma^2 \left(\chi^2 + \frac{1}{r} \right), \quad \tilde{\varrho} \tilde{\sigma}^2 = \frac{\sigma^2}{r} \quad \text{oraz} \quad \rho \sigma \tilde{\sigma} = \frac{\sigma^2}{r}$$

do wzoru (I.0.9) na wartość oczekiwaną S_{Imp}^2 , po łatwych uproszczeniach, otrzymujemy (I.3.5).

Wzór (I.3.6) wynika z porównania wzoru (I.3.5) oraz drugiego ze wzorów (I.3.2). ■

II. DETERMINISTYCZNY ZBIÓR BRAKÓW OBSERWACJI — IMPUTACJA WIELOKROTNA

II.0. Ogólny schemat imputacji wielokrotnej

Imputacja wielokrotna została wprowadzona przez Rubina w monografii z 1987 r. Poniżej przytaczamy podstawowe zasady tej metody, zwane zasadami Rubina (*Rubin's rules*), przedstawione np. na stronach 39 i 40 monografii Carpenter i Kenward (2013).

Imputacja wielokrotna polega na generowaniu kilku próbek imputacyjnych

$$\tilde{X}^{(l)} = (X_i, i \in R, \tilde{X}_j^{(l)}, j \in R^c), \quad l = 1, \dots, m,$$

na podstawie obserwowanej części próbki $X_R = (X_i, i \in R)$.

Definicja II.0.1. *Estymatorem wielokrotnej imputacji Rubina parametru θ (odpowiadającym estymatorowi $\hat{\theta} = g(X)$) nazywamy średnią z estymatorów imputacyjnych dla próbek $\tilde{X}^{(l)}$, $l = 1, \dots, m$, czyli*

$$\hat{\theta}_{MImp} = \frac{1}{m} \sum_{l=1}^m \hat{\theta}_{Imp}^{(l)}.$$

Estymator Rubina wariancji estymatora $\hat{\theta}_{MImp}$ ma postać

$$\hat{v}_{MImp}^2 = \bar{U}_m + \frac{m+1}{m} B_m,$$

gdzie

$$\bar{U}_m := \frac{1}{m} \sum_{l=1}^m \hat{v}_{Imp}^{(l),2}$$

oznacza średnią estymatorów wariancji estymatorów imputacyjnych, natomiast

$$B_m := \frac{1}{m-1} \sum_{l=1}^m \left(\hat{\theta}_{Imp}^{(l)} - \hat{\theta}_{MImp} \right)^2 \quad (\text{II.0.1})$$

jest wariancją empiryczną estymatorów imputacyjnych.

Poniżej interesować się będziemy estymacją średniej μ w modelu opisanym w poprzednim rozdziale. Zakładać będziemy, że $\tilde{X}^{(l)}$, $l=1, \dots, m$, są warunkowo niezależne pod warunkiem X_R . Oczywiście jest to równoważne temu, że imputowane części próbek $\left(\tilde{X}_j^{(l)}, j \in R^c \right)$, $l=1, \dots, m$, są warunkowo niezależne pod warunkiem X_R . Zakłada się również, że $E \tilde{X}_j^{(l)} = \tilde{\mu}$, $\text{Var } \tilde{X}_j^{(l)} = \tilde{\sigma}^2$ nie zależą od $j \in R^c$ oraz od $l=1, \dots, m$, $\text{Corr}(\tilde{X}_i^{(l)}, \tilde{X}_j^{(l)}) = \tilde{\rho}$ nie zależy od $i, j \in R^c$ oraz od $l=1, \dots, m$, a $\text{Corr}(X_i, \tilde{X}_j^{(l)}) = \rho$ nie zależy od $i \in R, j \in R^c$ oraz $l=1, \dots, m$.

Dodatkowo zakłada się warunek liniowości regresji

$$E(\tilde{X}_j^{(l)} | X_R) = \alpha \bar{X}_R + \beta$$

dla dowolnych $j \in R^c$ oraz $l=1, \dots, m$. Współczynniki α i β można wyliczyć korzystając z zależności

$$\tilde{\mu} = E E(\tilde{X}_j^{(l)} | X_R) = \alpha \mu + \beta$$

oraz

$$\tilde{\mu} \mu + \rho \sigma \tilde{\sigma} = E \tilde{X}_j^{(l)} X_i = E X_i E(\tilde{X}_j^{(l)} | X_R) = E X_i (\alpha \bar{X}_R + \beta) = \alpha \left(\frac{\sigma^2}{r} + \mu^2 \right) + \beta \mu.$$

Rozwiązując powyższy układ równań względem α i β otrzymujemy

$$E(\tilde{X}_j^{(l)} | X_R) = r \rho \frac{\tilde{\sigma}}{\sigma} (\bar{X}_R - \mu) + \tilde{\mu} \quad (\text{II.0.2})$$

dla dowolnych $j \in R^c$ oraz $l=1, \dots, m$.

Na podstawie próbek imputacyjnych konstruowane są estymatory

$$\hat{\mu}_{Imp}^{(l)} = \frac{1}{n} \left[r \bar{X}_R + (n-r) \tilde{X}^{(l)} \right], \quad l=1, \dots, m,$$

a estymator finalny jest ich średnią

$$\hat{\mu}_{MImp} = \frac{1}{m} \sum_{l=1}^m \hat{\mu}_{Imp}^{(l)} = \frac{r}{n} \bar{X}_R + \frac{n-r}{nm} \sum_{l=1}^m \widetilde{X}^{(l)}. \quad (\text{II.0.3})$$

TWIERDZENIE II.0.1. *Wartość oczekiwana estymatora wielokrotnej imputacji dla średniej wynosi*

$$E\hat{\mu}_{MImp} = \frac{r}{n} \mu + \frac{n-r}{n} \widetilde{\mu}. \quad (\text{II.0.4})$$

Jego wariancja ma postać

$$\begin{aligned} \text{Var}(\hat{\mu}_{MImp}) = \\ = \frac{mr\sigma^2 + (n-r)(1 + (n-r-1)\widetilde{\rho})\widetilde{\sigma}^2 + 2mr(n-r)\varrho\sigma\widetilde{\sigma} + (m-1)(n-r)^2r\varrho^2\widetilde{\sigma}^2}{mn^2}. \end{aligned} \quad (\text{II.0.5})$$

Dowód. Wzór (II.0.4) wynika wprost z faktu, że dla każdego z estymatorów $\hat{\mu}_{Imp}^{(l)}$, $l = 1, \dots, m$, obowiązuje wzór (I.0.2).

Znajdziemy teraz wariancję estymatora wielokrotnej imputacji średniej

$$\text{Var} \hat{\mu}_{MImp} = \frac{1}{m^2} \left(\sum_{l=1}^m \text{Var} \hat{\mu}_{Imp}^{(l)} + \sum_{k \neq l}^m \text{Cov} \left(\hat{\mu}_{Imp}^{(k)}, \hat{\mu}_{Imp}^{(l)} \right) \right).$$

Ale zgodnie ze wzorem (I.0.3)

$$\text{Var} \hat{\mu}_{Imp}^{(l)} = \frac{1}{n^2} \left(r\sigma^2 + (n-r)(1 + (n-r-1)\widetilde{\rho})\widetilde{\sigma}^2 + 2r(n-r)\varrho\sigma\widetilde{\sigma} \right).$$

Z kolei

$$\begin{aligned} \text{Cov} \left(\hat{\mu}_{Imp}^{(k)}, \hat{\mu}_{Imp}^{(l)} \right) = \frac{1}{n^2} \left(r^2 \text{Var} \left(\bar{X}_R \right) + (n-r)^2 \text{Cov} \left(\widetilde{X}^{(k)}, \widetilde{X}^{(l)} \right) + \right. \\ \left. + r(n-r) \left(\text{Cov} \left(\bar{X}_R, \widetilde{X}^{(k)} \right) + \text{Cov} \left(\bar{X}_R, \widetilde{X}^{(l)} \right) \right) \right). \end{aligned}$$

Wykorzystując wzór (I.0.6) otrzymujemy

$$\text{Cov} \left(\hat{\mu}_{Imp}^{(k)}, \hat{\mu}_{Imp}^{(l)} \right) = \frac{1}{n^2} \left(r\sigma^2 + (n-r)^2 \text{Cov} \left(\widetilde{X}^{(k)}, \widetilde{X}^{(l)} \right) + 2r(n-r)\varrho\sigma\widetilde{\sigma} \right).$$

Ale

$$\text{Cov}\left(\tilde{X}^{(k)}, \tilde{X}^{(l)}\right) = \text{ECov}\left(\tilde{X}^{(k)}, \tilde{X}^{(l)} | \mathbf{X}_R\right) + \text{Cov}\left(\text{E}\left(\tilde{X}^{(k)} | \mathbf{X}_R\right), \text{E}\left(\tilde{X}^{(l)} | \mathbf{X}_R\right)\right).$$

Z warunkowej niezależności wynika, że pierwszy składnik po prawej stronie jest równy zeru, natomiast

$$\text{Cov}\left(\text{E}\left(\tilde{X}^{(k)} | \mathbf{X}_R\right), \text{E}\left(\tilde{X}^{(l)} | \mathbf{X}_R\right)\right) = \frac{1}{(n-r)^2} \sum_{i,j \in R^c} \text{Cov}\left(\text{E}\left(\tilde{X}_i^{(k)} | \mathbf{X}_R\right), \text{E}\left(\tilde{X}_j^{(l)} | \mathbf{X}_R\right)\right).$$

Z liniowości regresji (II.0.2) otrzymujemy dla dowolnych $i, j \in R^c$

$$\text{Cov}\left(\text{E}\left(\tilde{X}_i^{(k)} | \mathbf{X}_R\right), \text{E}\left(\tilde{X}_j^{(l)} | \mathbf{X}_R\right)\right) = r^2 \varrho^2 \frac{\tilde{\sigma}^2}{\sigma^2} \text{Var}\left(\bar{X}_R\right) = r \varrho^2 \tilde{\sigma}^2.$$

Zatem

$$\text{Cov}\left(\hat{\mu}_{Imp}^{(k)}, \hat{\mu}_{Imp}^{(l)}\right) = \frac{1}{n^2} \left(r \sigma^2 + (n-r)^2 r \varrho^2 \tilde{\sigma}^2 + 2r(n-r) \varrho \sigma \tilde{\sigma} \right). \quad (\text{II.0.6})$$

Ostatecznie

$$\begin{aligned} \text{Var} \hat{\mu}_{MImp} &= \frac{1}{mn^2} \left(r \sigma^2 + (n-r) \left(1 + (n-r-1) \tilde{\rho} \right) \tilde{\sigma}^2 + 2r(n-r) \varrho + \sigma \tilde{\sigma} \right) \\ &\quad + \frac{m-1}{mn^2} \left(r \sigma^2 + (n-r)^2 r \varrho^2 \tilde{\sigma}^2 + 2r(n-r) \varrho \sigma \tilde{\sigma} \right), \end{aligned}$$

co po łatwych uproszczeniach pozwala otrzymać wzór (II.0.5).■

Zgodnie z drugą częścią Definicji II.0.1 estymatorem Rubina wariancji estymatora $\hat{\mu}_{MImp}$ nazywamy statystykę określoną wzorem

$$\hat{v}_{MImp}^2 = \bar{U}_m + \frac{m+1}{m} B_m, \quad (\text{II.0.7})$$

gdzie

$$\bar{U}_m := \frac{1}{mn} \sum_{l=1}^m \left(S_{Imp}^{(l)} \right)^2, \quad (\text{II.0.8})$$

przy czym

$$\left(S_{Imp}^{(l)}\right)^2 = \frac{1}{n-1} \sum_{i=1}^n \left(\tilde{X}_i^{(l)} - \hat{\mu}_{Imp}^{(l)}\right)^2, \quad l=1, \dots, m,$$

natomiast

$$B_m := \frac{1}{m-1} \sum_{l=1}^m \left(\hat{\mu}_{Imp}^{(l)} - \hat{\mu}_{MImp}\right)^2. \quad (\text{II.0.9})$$

Estymator (II.0.7), zwany estymatorem Rubina, na ogół nie jest nieobciążony.

TWIERDZENIE II.0.2. *Wartość oczekiwana estymatora Rubina wynosi*

$$\begin{aligned} E \hat{v}_{MImp}^2 = & \frac{r\sigma^2 + (n-r)\tilde{\sigma}^2 - \frac{(n-r)[(n-r-1)\tilde{\rho}\tilde{\sigma}^2 + 2r\rho\sigma\tilde{\sigma} - r(\mu - \tilde{\mu})^2]}{n-1}}{n^2} + \\ & + \frac{m+1}{m} \frac{(n-r)[1 + (n-r-1)\tilde{\varrho} - (n-r)r\varrho^2]}{n^2} \tilde{\sigma}^2. \end{aligned} \quad (\text{II.0.10})$$

Dowód. Zauważmy najpierw, że

$$E \bar{U}_m = \frac{1}{mn} \sum_{l=1}^m E \left(S_{Imp}^{(l)}\right)^2 = \frac{1}{n} E \left(S_{Imp}^{(1)}\right)^2.$$

A zatem zgodnie ze wzorem (I.0.9) otrzymujemy

$$E \bar{U}_m = \frac{1}{n^2} \left(r\sigma^2 + (n-r)\tilde{\sigma}^2 - \frac{(n-r)[(n-r-1)\tilde{\rho}\tilde{\sigma}^2 + 2r\rho\sigma\tilde{\sigma} - r(\mu - \tilde{\mu})^2]}{n-1} \right), \quad (\text{II.0.11})$$

czyli pierwszy składnik we wzorze (II.0.10).

Z kolei, jeśli $E \hat{\mu}_{Imp}^{(l)} = E \hat{\mu}_{MImp} = v, l=1, \dots, m$, to

$$\begin{aligned} \sum_{l=1}^m E \left(\hat{\mu}_{Imp}^{(l)} - \hat{\mu}_{MImp}\right)^2 &= \sum_{l=1}^m E \left(\hat{\mu}_{Imp}^{(l)} - v\right)^2 - 2E \left(\hat{\mu}_{MImp} - v\right) \sum_{l=1}^m \left(\hat{\mu}_{Imp}^{(l)} - v\right) + \\ &+ mE \left(\hat{\mu}_{MImp} - v\right)^2 = m \left(\text{Var} \hat{\mu}_{Imp}^{(1)} - \text{Var} \hat{\mu}_{MImp} \right). \end{aligned}$$

Ale

$$\begin{aligned}\text{Var } \hat{\mu}_{MImp} &= \frac{1}{m^2} \left(\sum_{l=1}^m \text{Var } \hat{\mu}_{Imp}^{(l)} + \sum_{l,k=1, l \neq k}^m \text{Cov}(\hat{\mu}_{Imp}^{(l)}, \hat{\mu}_{Imp}^{(k)}) \right) = \\ &= \frac{\text{Var } \hat{\mu}_{Imp}^{(1)} + (m-1) \text{Cov}(\hat{\mu}_{Imp}^{(1)}, \hat{\mu}_{Imp}^{(2)})}{m}.\end{aligned}$$

W konsekwencji

$$\begin{aligned}\sum_{l=1}^m \text{E}(\hat{\mu}_{Imp}^{(l)} - \hat{\mu}_{MImp})^2 &= m \left(\text{Var } \hat{\mu}_{Imp}^{(1)} - \frac{\text{Var } \hat{\mu}_{Imp}^{(1)} + (m-1) \text{Cov}(\hat{\mu}_{Imp}^{(1)}, \hat{\mu}_{Imp}^{(2)})}{m} \right) = \\ &= (m-1) (\text{Var } \hat{\mu}_{Imp}^{(1)} + \text{Cov}(\hat{\mu}_{Imp}^{(1)}, \hat{\mu}_{Imp}^{(2)})).\end{aligned}$$

Zatem

$$\text{E} B_m = \frac{1}{m-1} \sum_{l=1}^m \text{E}(\hat{\mu}_{Imp}^{(l)} - \hat{\mu}_{MImp})^2 = \text{Var } \hat{\mu}_{Imp}^{(1)} + \text{Cov}(\hat{\mu}_{Imp}^{(1)}, \hat{\mu}_{Imp}^{(2)}).$$

Stosując w powyższej równości wzory (I.0.3) i (II.0.6) otrzymujemy

$$\text{E} B_m = \frac{(n-r)[1 + (n-r-1)\tilde{\sigma} - (n-r)r\varrho^2]}{n^2} \tilde{\sigma}^2, \quad (\text{II.0.12})$$

czyli drugą część wzoru (II.0.10), bez mnożnika $\frac{m+1}{m}$. ■

II.1. Imputacja wielokrotna typu hot-deck

II.1.1. Losowanie spośród respondentów

Każdą z próbek imputacyjnych $\tilde{X}^{(l)}$, $l = 1, \dots, m$, tworzymy stosując metodę hot-deck opisaną w podrozdziale I.2. Niech $\mathbf{K}^{(l)} = (\mathbf{K}_j^{(l)}, j \in R^c)$ będzie wektorem numerów elementów wylosowanych dla kolejnych nierespondentów przy tworzeniu próbki $\tilde{X}^{(l)}$. Zakładamy, że wektory losowe $\mathbf{K}^{(l)}$, $l = 1, \dots, m$, oraz $\mathbf{X}_R$ są niezależne. Wtedy próbki imputacyjne $\tilde{X}^{(l)}$, $l = 1, \dots, m$, są warunkowo niezależne pod warunkiem $\mathbf{X}_R$.

Zgodnie z teorią rozwiniętą dla jednokrotnej imputacji hot-deck mamy

$$\tilde{\mu} = \mu, \quad \tilde{\sigma}^2 = \sigma^2, \quad \tilde{\rho} = \rho = \frac{1}{r}. \quad (\text{II.1.1})$$

Co więcej, z niezależności $\mathbf{K}^{(l)}$ i $\mathbf{X}_R$ wynika, że dla $j \in R^c$

$$E(\tilde{X}_j^{(l)} | \mathbf{X}_R) = E(X_{K_j^{(l)}} | \mathbf{X}_R) = \bar{X}_R,$$

czyli $\alpha = 1$ oraz $\beta = 0$ w ogólnym wzorze na liniową regresję z podrozdziału II.0.

TWIERDZENIE II.1.1. *Estymator wielokrotnej imputacji ma postać*

$$\hat{\mu}_{MImp} = \frac{r}{n} \bar{X}_R + \frac{1}{nm} \sum_{l=1}^m \sum_{j \in R^c} X_{K_j^{(l)}}.$$

Jego wartość oczekiwana wynosi

$$E\hat{\mu}_{MImp} = \mu,$$

a jego wariancja ma postać

$$\text{Var} \hat{\mu}_{MImp} = \frac{\sigma^2}{r} \left(1 + \frac{(n-r)(r-1)}{mn^2} \right). \quad (\text{II.1.2})$$

Dowód. Wzór na postać estymatora wielokrotnej imputacji wynika wprost z ogólnej postaci (II.0.3) oraz wzoru (I.2.2) z twierdzenia I.2.1, natomiast wzór na jego wartość oczekiwaną jest konsekwencją wzoru (II.0.4) oraz równości $\tilde{\mu} = \mu$.

Zgodnie ze wzorem (II.0.5) mamy

$$\begin{aligned} \text{Var} \hat{\mu}_{MImp} &= \frac{mr\sigma^2 + (n-r) \left(1 + \frac{n-r-1}{r} \right) \sigma^2 + 2m(n-r)\sigma^2 + (m-1) \frac{(n-r)^2}{r} \sigma^2}{mn^2} = \\ &= \frac{\sigma^2}{n^2} \left(r + 2(n-r) + \frac{n^2}{r} - 2n + r \right) + \frac{\sigma^2}{mn^2} \left(n-r + \frac{(n-r)^2}{r} - \frac{n-r}{r} - \frac{(n-r)^2}{r} \right) = \\ &= \frac{\sigma^2}{r} + \frac{\sigma^2(n-r)(r-1)}{mn^2r}, \end{aligned}$$

co daje wzór (II.1.2). ■

TWIERDZENIE II.1.2. *Obciążenie estymatora Rubina wariancji $\hat{v}_{MImp}^2$ wynosi*

$$B\hat{v}_{MImp}^2 = E\hat{v}_{MImp}^2 - \text{Var}\hat{\mu}_{MImp} = -\frac{(n-r)[n(n-r+1)+2(r-1)]}{n^2(n-1)r}\sigma^2. \quad (\text{II.1.3})$$

Estymator nieobciążony wariancji estymatora $\hat{\mu}_{MImp}$ ma postać

$$\hat{v}_*^2 = \frac{n}{r}\bar{U}_m + \left(\frac{1}{m} + \frac{n(n+r-1)}{(n-1)r(r-1)}\right)B_m.$$

Dowód. Wstawiając do (II.0.11) wartości podane w (II.1.1) otrzymujemy

$$E\bar{U}_m = \frac{\sigma^2}{n^2} \left(n - \frac{(n-r)(n+r-1)}{r(n-1)} \right).$$

Z kolei wstawiając do (II.0.12) te same wartości z (II.1.1) otrzymujemy

$$EB_m = \sigma^2 \frac{(n-1)(r-1)}{rn^2}. \quad (\text{II.1.4})$$

Zatem

$$\begin{aligned} E\hat{v}_{MImp}^2 &= E\bar{U}_m + EB_m + \frac{1}{m}EB_m = \frac{\sigma^2(r-1)(2n(n-1)+2r-nr)}{rn^2(n-1)} + \\ &+ \frac{\sigma^2(r-1)(n-r)}{mn^2r}. \end{aligned}$$

Obciążenie wynosi więc

$$E\hat{v}_{MImp}^2 - \text{Var}\hat{\mu}_{MImp} = \frac{\sigma^2(r-1)(2n(n-1)+2r-nr)}{rn^2(n-1)} - \frac{\sigma^2}{r},$$

co daje po uproszczeniach wzór (II.1.3).

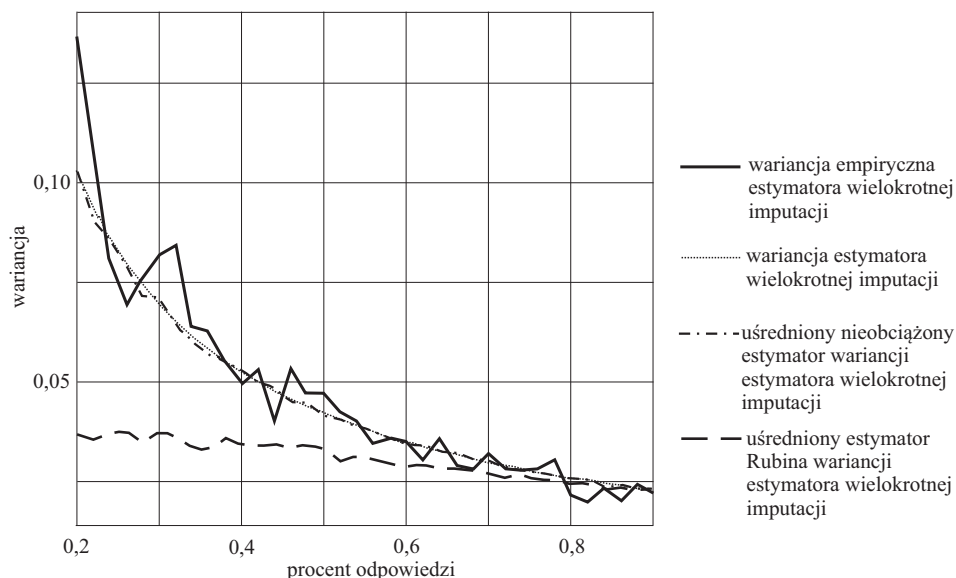
Zauważmy, że

$$\frac{n}{r}E\bar{U}_m + \frac{n(n+r-1)}{(n-1)r(r-1)}EB_m = \frac{\sigma^2}{r}.$$

Nieobciążoność estymatora $\hat{v}_*^2$ wynika z powyższego wzoru oraz (II.1.4) i wzoru (II.1.2) na wariancję estymatora $\hat{\mu}_{MImp}$. ■

Symulacyjna ilustracja wyników zawartych w twierdzeniach II.1.1 i II.1.2 przedstawiona jest na wyk. 1. Szczegółowy opis procedury symulacyjnej znajduje się w aneksie.

**Wykr. 1. WYNIKI ESTYMACJI WARIANCJI ESTYMATORA ŚREDNIEJ
DLA WIELOKROTNEJ IMPUTACJI (HOT-DECK)**



Źródło: obliczenia własne (zobacz aneks).

Wniosek II.1.3. *Asymptotycznie względne obciążenie estymatora Rubina $\hat{v}_{MImp}^2$ wynosi*

$$\lim_{n \rightarrow \infty} \frac{E\hat{v}_{MImp}^2 - \text{Var} \hat{\mu}_{MImp}}{\text{Var} \hat{\mu}_{MImp}} = -1. \quad (\text{II.1.5})$$

Dowód. Ze wzoru (II.1.2) mamy

$$\lim_{n \rightarrow \infty} \text{Var} \hat{\mu}_{MImp} = \frac{\sigma^2}{r}.$$

Z kolei wzór (II.1.3) implikuje

$$\lim_{n \rightarrow \infty} (E\hat{v}_{MImp}^2 - \text{Var} \hat{\mu}_{MImp}) = -\frac{\sigma^2}{r}. \blacksquare$$

Uwaga. Oczywiście istnieją też inne niż $\hat{v}_*^2$ nieobciążone estymatory wariancji estymatora wielokrotnej imputacji $\hat{\mu}_{MImp}$. W szczególności estymatory mające postać kombinacji liniowej statystyk $\bar{U}_m$ i B_m . Analiza tej klasy estymatorów jest równoważna rozważaniu estymatorów wariancji $\hat{\mu}_{MImp}$ postaci

$$\hat{v}_{(\alpha,\beta)}^2 = \alpha \bar{U}_m + \beta B_m + \frac{1}{m} B_m,$$

przy warunku

$$\beta = \frac{1}{n-r} \left(\frac{n^2}{r-1} - \frac{n(n-1)+r}{(n-r)(n-1)} \alpha \right) \quad (\text{II.1.6})$$

zapewniającym nieobciążoność. Istotnym zagadnieniem jest znalezienie optymalnych wartości α , β , tzn. takich, przy których wariancja estymatora $\hat{v}_{(\alpha,\beta)}^2$ osiąga minimum. Rozwiązanie analityczne jest trudne, wymaga znajomości trzecich i czwartych momentów. Możliwe jest natomiast numeryczne badanie wariancji estymatora $\hat{v}_{(\alpha,\beta)}^2$ jako funkcji zmiennej α , przy ograniczeniu (II.1.6).

II.1.2. Losowanie z rozkładu normalnego

Korzystamy ze wzorów ogólnych, przy czym

$$\tilde{\mu} = \mu, \quad \tilde{\sigma}^2 = \sigma^2 \frac{1+r}{r}, \quad \varrho = \frac{1}{\sqrt{r(r+1)}}, \quad \tilde{\varrho} = \frac{1}{r+1}.$$

TWIERDZENIE II.1.4. *Wartość oczekiwana estymatora imputacji wielokrotnej przy losowaniu z rozkładu normalnego wynosi*

$$E \hat{\mu}_{MImp} = \mu,$$

a jego wariancja ma postać

$$\text{Var} \hat{\mu}_{MImp} = \frac{\sigma^2}{r} \left(1 + \frac{(n-r)r}{mn^2} \right). \quad (\text{II.1.7})$$

Dowód. Powyższe wzory wynikają wprost ze wzorów (II.0.4) i (II.0.5) z twierdzenia II.1.1. ■

TWIERDZENIE II.1.5. *Obciążenie estymatora Rubina wariancji $\hat{v}_{MImp}^2$ wynosi*

$$B\hat{v}_{MImp}^2 = E\hat{v}_{MImp}^2 - \text{Var}\hat{\mu}_{MImp} = -\frac{(n-r)[(n-1)(n-r)-1]}{n^2(n-1)r}\sigma^2. \quad (\text{II.1.8})$$

Estymator nieobciążony wariancji estymatora $\hat{\mu}_{MImp}$ ma postać

$$\hat{v}_*^2 = \frac{n}{r}\bar{U}_m + \left(\frac{1}{m} + \frac{n}{(n-1)r}\right)B_m.$$

Dowód. Ze wzorów (II.0.11) oraz (II.0.12) wynika, że

$$E\bar{U}_m = \frac{\sigma^2}{n^2}\left(n - \frac{n-r}{n-1}\right) \quad \text{oraz} \quad EB_m = \frac{\sigma^2(n-r)}{n^2}.$$

Zatem

$$E\hat{v}_{MImp}^2 - \text{Var}\hat{\mu}_{MImp} = \frac{\sigma^2}{n^2}\left(n + \frac{(n-r)(r-2)}{r(n-1)}\right) + \frac{m+1}{m}\frac{\sigma^2(n-r)}{n^2} - \frac{\sigma^2}{r}\left(1 + \frac{(n-r)r}{mn^2}\right)$$

skąd po prostych przekształceniach otrzymujemy wzór (II.1.8).

Ponieważ

$$\text{Var}\hat{\mu}_{MImp} - E\frac{1}{m}B_m = \frac{\sigma^2}{r},$$

więc wystarczy znaleźć α i β takie, że $\alpha E\bar{U}_m + \beta EB_m = \frac{\sigma^2}{r}$. Przyjmujemy

$$\alpha = \frac{n}{r}.$$

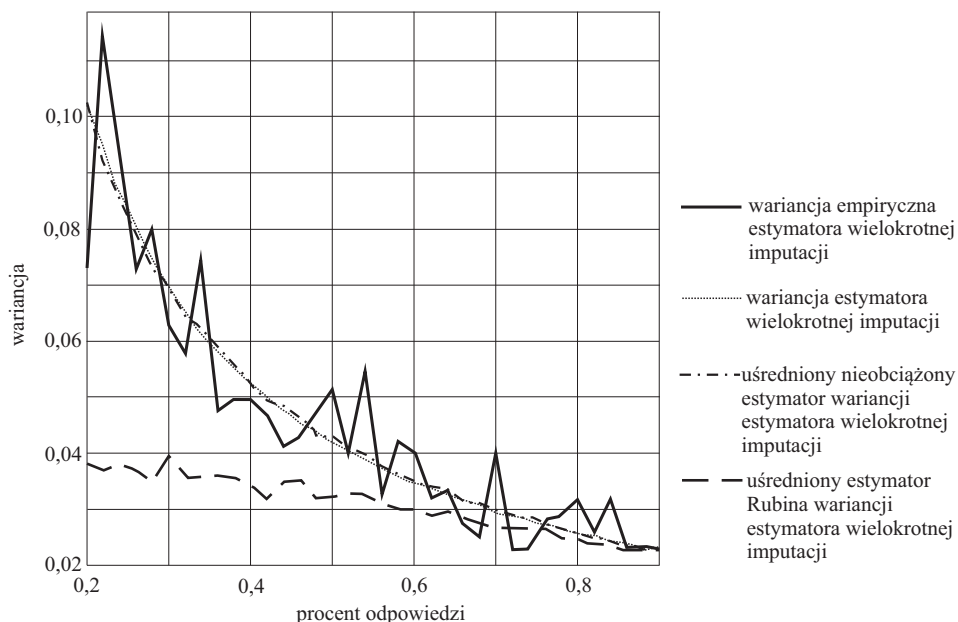
Wtedy nietrudno wyliczyć, że $\beta = -\frac{n(r-2)}{(n-1)r^2}$. ■

Podobnie jak we wniosku II.1.3 i w tym przypadku asymptotyczne obciążenie wynosi

$$\lim_{n \rightarrow \infty} \frac{E\hat{v}_{MImp}^2 - \text{Var}\hat{\mu}_{MImp}}{\text{Var}\hat{\mu}_{MImp}} = -1.$$

Symulacyjna ilustracja wyników zawartych w twierdzeniach II.1.4 i II.1.5 przedstawiona jest na wykr. 2. Szczegółowy opis procedury symulacyjnej znajduje się w aneksie.

**Wykr. 2. WYNIKI ESTYMACJI WARIANCJI ESTYMATORA ŚREDNIEJ
DLA WIELOKROTNEJ IMPUTACJI (HOT-DECK) Z ROZKŁADU NORMALNEGO**



Źródło: jak przy wykr. 1.

III. LOSOWY ZBIÓR BRAKÓW OBSERWACJI — IMPUTACJA JEDNOKROTNA

III.0. Ogólny schemat imputacyjny

Niech wektor losowy $X = (X_1, \dots, X_n)$ oznacza pierwotną próbkę, natomiast $J = (J_1, \dots, J_n)$ niech oznacza wektor losowy, w którym składowa $J_i = 1$, jeśli zmienna X_i jest obserwowana oraz $J_i = 0$, jeśli zmienna X_i nie jest obserwowana. W konsekwencji zbiór $R = \{i \in \{1, \dots, n\} : J_i = 1\}$ jednostek, dla których obserwacja jest dostępna jest losowy.

Przy założeniu, że mechanizm pojawiania się braków odpowiedzi nie jest zależny od wartości badanych zmiennych (MCAR, *missing completely at random*), czyli że wektory losowe X i J są (statystycznie) niezależne, do analizy własności rozkładu warunkowego X pod warunkiem J można wygodnie stosować metodologię rozwiniętą w rozdziałach I i II, czyli w sytuacji, gdy zbiór R jest deterministyczny. Ostateczne wyniki otrzymuje się poprzez uśrednianie

rezultatów dla rozkładu warunkowego względem rozkładu wektora $\mathbf{J}$. Szczegółowo jest to opisane dla imputacji jednokrotnej w bieżącym rozdziale i dla imputacji wielokrotnej w rozdziale kolejnym.

Konstrukcja wektora $\mathbf{J}$ wymaga przeanalizowania pewnej subtelności technicznej. Losowość braków obserwacji modelowana jest pierwotnie przez zmienne losowe (indykatory odpowiedzi) $\delta_1, \dots, \delta_n$, przy czym $\delta_i = 1$, jeśli obserwacja X_i jest dostępna, a $\delta_i = 0$, jeśli obserwacja X_i nie jest dostępna (brak odpowiedzi). Zakładamy, że zmienne $\delta_1, \dots, \delta_n$ są niezależne oraz że wektory $\mathbf{X}$ i $\boldsymbol{\delta} = (\delta_1, \dots, \delta_n)$ są również niezależne. Wprowadzamy oznaczenie p_i na prawdopodobieństwo i -tej odpowiedzi, $P(\delta_i = 1) = p_i$, dla każdego $i = 1, \dots, n$. Niech

$|\boldsymbol{\delta}| = \sum_{i=1}^n \delta_i$ oznacza liczbę elementów obserwowanych w próbce. Jeśli w danej

próbce nie ma żadnej dostępnej obserwacji, tzn. $|\boldsymbol{\delta}| = 0$, to badamy kolejną próbkę i postępujemy tak do momentu, gdy znajdziemy próbkę, dla której $|\boldsymbol{\delta}| > 0$.

Dopiero na zakończenie takiej procedury wprowadzamy wektor ostatecznych braków odpowiedzi $\mathbf{J}$ — jeśli obserwacja X_i jest dostępna, przyjmujemy $J_i = 1$, a jeśli obserwacja X_i nie jest dostępna, przyjmujemy $J_i = 0$, $i = 1, \dots, n$. Czyli zamiast oryginalnego wektora $\boldsymbol{\delta}$ o składowych niezależnych, braki odpowiedzi opisane są za pomocą wektora $\mathbf{J} = (J_1, \dots, J_n)$, którego składowe nie są niezależne. Tym niemniej wektory losowe $\mathbf{X}$ i $\mathbf{J}$ są niezależne. Zauważmy, że

$$P(J_1 = \varepsilon_1, \dots, J_n = \varepsilon_n) = P(\delta_1 = \varepsilon_1, \dots, \delta_n = \varepsilon_n \mid |\boldsymbol{\delta}| > 0) = \frac{\prod_{i=1}^n p_i^{\varepsilon_i} (1-p_i)^{1-\varepsilon_i}}{1 - \prod_{i=1}^n (1-p_i)},$$

dla $\varepsilon_i \in \{0, 1\}$, $i = 1, \dots, n$, takich, że $\sum_{i=1}^n \varepsilon_i > 0$. Wtedy $|\mathbf{J}| = |\boldsymbol{\delta}| > 0$ oraz

$$P(|\mathbf{J}| = k) = \frac{\sum_{1 \leq i_1, \dots, i_k \leq n} \prod_{l=1}^k p_{i_l} \prod_{j \notin \{i_1, \dots, i_k\}} (1-p_j)}{1 - \prod_{i=1}^n (1-p_i)}, \quad k = 1, \dots, n.$$

W szczególnym przypadku, gdy prawdopodobieństwo odpowiedzi jest stałe, tzn. $p_i = p$, $i = 1, \dots, n$,

$$P(J_1 = \varepsilon_1, \dots, J_n = \varepsilon_n) = \frac{p^{\sum_{i=1}^n \varepsilon_i} (1-p)^{n - \sum_{i=1}^n \varepsilon_i}}{1 - (1-p)^n}.$$

Wtedy zmienna losowa $|\delta|$ ma rozkład dwumianowy, $b(n, p)$, tzn.

$$P(|\delta| = k) = \binom{n}{k} p^k (1-p)^{n-k}, \quad k = 0, 1, \dots, n,$$

a zmienna losowa $|J|$ ma rozkład dwumianowy ucięty w zerze, $b_+(n, p)$, tzn.

$$P(|J| = k) = \binom{n}{k} \frac{p^k (1-p)^{n-k}}{q_n}, \quad k = 1, \dots, n,$$

gdzie

$$q_n = 1 - (1-p)^n.$$

Gdy prawdopodobieństwo braków odpowiedzi jest jednakowe i wynosi p , wtedy

$$Eg(|J|) = \frac{1}{q_n} \sum_{i=1}^n g(i) \binom{n}{i} p^i (1-p)^{n-i},$$

w szczególności

$$E|J| = \frac{np}{q_n} \quad \text{oraz} \quad E\frac{1}{|J|} = \frac{1}{q_n} \sum_{i=1}^n \frac{1}{i} \binom{n}{i} p^i (1-p)^{n-i}.$$

Uwaga. W pracy Szablowski, Wesołowski i Wieczorkowski (1996) pokazano, że $E\frac{1}{|J|} \cong \frac{1}{np}$. Dokładniejsze przybliżenie można znaleźć np. w pracach Marciniak i Wesołowski (1999) albo Rempała (2004).

Zmienne imputacyjne oznaczamy symbolem $\tilde{X}_i$, $i = 1, \dots, n$. Estymator imputacyjny średniej ma postać

$$\hat{\mu}_{Imp} = \frac{1}{n} \sum_{i=1}^n (X_i J_i + \tilde{X}_i (1 - J_i)) = \frac{1}{n} \left(\sum_{j \in R} X_i + \sum_{j \in R^c} \tilde{X}_i \right). \quad (\text{III.0.1})$$

Natomiast imputacyjny estymator wariancji ma postać

$$S_{Imp}^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i J_i + \tilde{X}_i (1 - J_i) - \hat{\mu}_{Imp})^2. \quad (\text{III.0.2})$$

III.1. Imputacja średnią

Technika ta polega na przypisaniu każdemu nieobserwowanemu elementowi próbki średniej z obserwowanej części próbki, czyli

$$\tilde{X}_i = \frac{1}{|\mathbf{J}|} \sum_{j=1}^n X_j J_j =: \bar{X}_J. \quad (\text{III.1.1})$$

TWIERDZENIE III.1.1. *Imputacyjny estymator wartości oczekiwanej w przypadku imputacji średnią ma postać*

$$\hat{\mu}_{Imp} = \bar{X}_J. \quad (\text{III.1.2})$$

Jego wartość oczekiwana oraz wariancja wynoszą odpowiednio

$$\mathbb{E} \hat{\mu}_{Imp} = \mu \quad \text{oraz} \quad \text{Var} \hat{\mu}_{Imp} = \sigma^2 \mathbb{E} \frac{1}{|\mathbf{J}|}. \quad (\text{III.1.3})$$

Dowód. Ze wzorów (III.0.1) i (III.0.2) wynika, że

$$\hat{\mu}_{Imp} = \frac{1}{n} \sum_{i=1}^n (J_i X_i + (1 - J_i) \bar{X}_J) = \bar{X}_J + \frac{1}{n} \sum_{i=1}^n J_i X_i - \frac{1}{n} \bar{X}_J \sum_{i=1}^n J_i,$$

co prowadzi bezpośrednio do wzoru (III.1.2).

Z pierwszego ze wzorów (I.1.3) mamy

$$\mathbb{E}(\hat{\mu}_{Imp} | \mathbf{J}) = \mu,$$

a zatem

$$\mathbb{E} \hat{\mu}_{Imp} = \mathbb{E} \mathbb{E}(\hat{\mu}_{Imp} | \mathbf{J}) = \mu.$$

Z kolei

$$\text{Var} \hat{\mu}_{Imp} = \text{Var} \mathbb{E}(\hat{\mu}_{Imp} | \mathbf{J}) + \mathbb{E} \text{Var}(\hat{\mu}_{Imp} | \mathbf{J}) = \mathbb{E} \text{Var}(\hat{\mu}_{Imp} | \mathbf{J}),$$

ponieważ $\mathbb{E}(\hat{\mu}_{Imp} | \mathbf{J})$ nie jest losowa. Z drugiego ze wzorów (I.1.3) mamy

$$\text{Var}(\hat{\mu}_{Imp} | \mathbf{J}) = \frac{\sigma^2}{|\mathbf{J}|}.$$

Czyli zachodzi drugi ze wzorów (III.1.3). ■

TWIERDZENIE III.1.2. *Imputacyjny estymator wariancji w przypadku imputacji średnią ma postać*

$$S_{Imp}^2 = \frac{1}{n-1} \sum_{i=1}^n J_i (X_i - \bar{X}_J)^2, \quad (\text{III.1.4})$$

a jego wartość oczekiwana wynosi

$$ES_{Imp}^2 = \frac{E|\mathbf{J}| - 1}{n-1} \sigma^2. \quad (\text{III.1.5})$$

Nieobciążony imputacyjny estymator wariancji estymatora $\hat{\mu}_{Imp}$ ma postać

$$v^2(\hat{\mu}_{Imp}) = \frac{(n-1)E \frac{1}{|\mathbf{J}|}}{E|\mathbf{J}| - 1} S_{Imp}^2. \quad (\text{III.1.6})$$

W przypadku $p_i = p, i = 1, \dots, n$

$$ES_{Imp}^2 = \frac{np - q_n}{q_n(n-1)} \sigma^2 \quad \text{oraz} \quad v^2(\hat{\mu}_{Imp}) = \frac{q_n(n-1)E \frac{1}{|\mathbf{J}|}}{np - q_n} S_{Imp}^2. \quad (\text{III.1.7})$$

Dowód. Zgodnie ze wzorami (III.0.2) i (III.1.2) mamy

$$S_{Imp}^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i J_i + \bar{X}_J (1 - J_i) - \bar{X}_J)^2 = \frac{1}{n-1} \sum_{i=1}^n J_i (X_i - \bar{X}_J)^2.$$

Zgodnie ze wzorem (I.1.5)

$$E(S_{Imp}^2 | \mathbf{J}) = \frac{|\mathbf{J}| - 1}{n-1} \sigma^2.$$

Zatem równość $ES_{Imp}^2 = EE(S_{Imp}^2 | \mathbf{J})$ pociąga wzór (III.1.5). ■

Zauważmy, że ze wzoru (III.1.5) wynika, iż estymator S_{Imp}^2 jest obciążonym estymatorem wariancji.

III.2. Imputacja typu hot-deck

III.2.1. Losowanie spośród respondentów

Niech K_i będzie numerem elementu wylosowanym dla i -tej jednostki, gdy $i \notin R$. Wtedy zmienne losowe K_i $i \notin R$, są warunkowo niezależne pod warunkiem $\mathbf{J}$ oraz $P(K_i = j | \mathbf{J}) = \frac{1}{|\mathbf{J}|}$, $j \in R$. Co więcej, wektory losowe $\mathbf{K}_{R^c} = (K_i, i \in R^c)$

i $\mathbf{X}_R = (X_j, j \in R)$ są warunkowo niezależne pod warunkiem $\mathbf{J}$.

W konsekwencji

$$\tilde{X}_i = X_{K_i}, \quad i \in \{1, \dots, n\}. \quad (\text{III.2.1})$$

TWIERDZENIE III.2.1. *Imputacyjny estymator wartości oczekiwanej w przypadku imputacji typu hot-deck polegającej na losowaniu respondentów ma postać*

$$\hat{\mu}_{Imp} = \frac{1}{n} \left(\sum_{j \in R} X_j + \sum_{i \in R^c} X_{K_i} \right). \quad (\text{III.2.2})$$

Jego wartość oczekiwana oraz wariancja wynoszą odpowiednio

$$E\hat{\mu}_{Imp} = \mu \text{ oraz } \text{Var}\hat{\mu}_{Imp} = \frac{\sigma^2}{n} \left(1 + \frac{1}{n} + (n-1)E\frac{1}{|\mathbf{J}|} - \frac{E|\mathbf{J}|}{n} \right). \quad (\text{III.2.3})$$

W przypadku $p_i = p$, $i = 1, \dots, n$

$$\text{Var}\hat{\mu}_{Imp} = \frac{\sigma^2}{n} \left(1 + \frac{1}{n} + (n-1)E\frac{1}{|\mathbf{J}|} - \frac{p}{q_n} \right). \quad (\text{III.2.4})$$

Dowód. Z twierdzenia I.2.1 wynika, że $E(\hat{\mu}_{Imp} | \mathbf{J}) = \mu$. W konsekwencji, warunkując względem $\mathbf{J}$ otrzymujemy pierwszy ze wzorów (III.2.3).

Ze wzoru (I.2.3) otrzymujemy

$$\text{Var}(\hat{\mu}_{Imp} | \mathbf{J}) = \frac{\sigma^2}{n} \left(1 + \frac{1}{n} + \frac{n-1}{|\mathbf{J}|} - \frac{|\mathbf{J}|}{n} \right).$$

Ponieważ warunkowa wartość oczekiwana $E(\hat{\mu}_{Imp}|\mathbf{J}) = \mu$ jest nieelosowa, więc

$$\text{Var} \hat{\mu}_{Imp} = E \text{Var}(\hat{\mu}_{Imp}|\mathbf{J}),$$

skąd wynika drugi ze wzorów (III.2.3). Wzór (III.2.4) jest natychmiastową konsekwencją (III.2.3) oraz wzoru na wartość oczekiwaną w rozkładzie $b_+(n, p)$. ■

TWIERDZENIE III.2.2. *Imputacyjny estymator wariancji w przypadku imputacji typu hot-deck polegającej na losowaniu respondentów ma postać*

$$S_{Imp}^2 = \frac{(|\mathbf{J}|-1)S_R^2 + (n-|\mathbf{J}|-1)S_{R^c}^2 + \frac{|\mathbf{J}|(n-|\mathbf{J}|)}{n}(\bar{X}_R - \bar{X}_{R^c})^2}{n-1},$$

gdzie $S_{R^c}^2 = \frac{1}{n-|\mathbf{J}|-1} \sum_{i \in R^c} (X_{K_i} - \bar{X}_{R^c})^2$ oraz $\bar{X}_{R^c} = \frac{1}{n-|\mathbf{J}|} \sum_{j \in R^c} X_{K_j}$.

Jego wartość oczekiwana wynosi

$$E S_{Imp}^2 = \left(1 + \frac{E|\mathbf{J}|-1}{n(n-1)} - E \frac{1}{|\mathbf{J}|} \right) \sigma^2. \quad (\text{III.2.5})$$

Nieobciążony imputacyjny estymator wariancji estymatora $\hat{\mu}_{Imp}$ ma postać

$$v^2(\hat{\mu}_{Imp}) = \frac{n+1+n(n-1)E \frac{1}{|\mathbf{J}|} - E|\mathbf{J}|}{n(n-1)-1+E|\mathbf{J}|-n(n-1)E \frac{1}{|\mathbf{J}|}} \frac{n-1}{n} S_{Imp}^2. \quad (\text{III.2.6})$$

W przypadku gdy $p_i = p$, $i = 1, \dots, n$

$$E S_{Imp}^2 = \left(1 - \frac{1}{n(n-1)} + \frac{p}{q_n(n-1)} - E \frac{1}{|\mathbf{J}|} \right) \sigma^2$$

oraz
$$v^2(\hat{\mu}_{Imp}) = \frac{n+1+n(n-1)E \frac{1}{|\mathbf{J}|} - \frac{np}{q_n}}{n(n-1)-1+\frac{np}{q_n}-n(n-1)E \frac{1}{|\mathbf{J}|}} \frac{n-1}{n} S_{Imp}^2.$$

Dowód. Postać estymatora S_{Imp}^2 wynika wprost z postaci podanej w twierdzeniu I.2.2, natomiast ze wzoru (I.2.8) mamy

$$E(S_{Imp}^2 | \mathbf{J}) = \left(1 - \frac{1}{n(n-1)} + \frac{|\mathbf{J}|}{n(n-1)} - \frac{1}{|\mathbf{J}|} \right) \sigma^2.$$

Więc (III.2.5) wynika z równości $ES_{Imp}^2 = EE(S_{Imp}^2 | \mathbf{J})$. Łącząc wzory (III.2.4) i (III.2.5) otrzymujemy postać (III.2.6) nieobciążonego estymatora wariancji estymatora $\hat{\mu}_{Imp}$. ■

III.2.2. Losowanie z rozkładu normalnego

Stosując wykorzystywaną już wcześniej technikę warunkowania przez obserwowaną część próbki $\mathbf{X}_R$ bezpośrednio z twierdzenia I.2.3 otrzymujemy pełny opis standardowych własności imputacyjnego estymatora w tym przypadku.

TWIERDZENIE III.2.3. *Dla imputacyjnego estymatora wartości oczekiwanej w przypadku imputacji hot-deck polegającej na losowaniu z rozkładu normalnego wartość oczekiwana oraz wariancja wynoszą odpowiednio*

$$E\hat{\mu}_{Imp} = \mu \text{ oraz } \text{Var}\hat{\mu}_{Imp} = \sigma^2 \left(E \frac{1}{|\mathbf{J}|} + \frac{n - E|\mathbf{J}|}{n^2} \right). \quad (\text{III.2.7})$$

Natomiast

$$ES_{Imp}^2 = \left(1 - \frac{n - E|\mathbf{J}|}{n(n-1)} \right) \sigma^2. \quad (\text{III.2.8})$$

Nieobciążony estymator wariancji estymatora $\hat{\mu}_{Imp}$ ma postać

$$v^2(\hat{\mu}_{Imp}) = \frac{n-1}{n} \frac{n^2 E \frac{1}{|\mathbf{J}|} + n - E|\mathbf{J}|}{n(n-1) - n + E|\mathbf{J}|} S_{Imp}^2. \quad (\text{III.2.9})$$

W przypadku gdy $p_i = p$, $i = 1, \dots, n$

$$\text{Var}\hat{\mu}_{Imp} = \sigma^2 \left(E \frac{1}{|\mathbf{J}|} + \frac{q_n - p}{n} \right) \text{ oraz } ES_{Imp}^2 = \left(1 - \frac{q_n - p}{(n-1)q_n} \right) \sigma^2.$$

Ponadto nieobciążony estymator wariancji estymatora $\hat{\mu}_{Imp}$ ma postać

$$v^2(\hat{\mu}_{Imp}) = \frac{n-1}{n} \frac{\left(nE \frac{1}{|\mathbf{J}|} + 1 \right) q_n - p}{(n-2)q_n + p} S_{Imp}^2.$$

Dowód. Wzór (III.2.7) jest bezpośrednią konsekwencją wzoru (I.2.10) oraz tego, że $E(\hat{\mu}_{Imp}|\mathbf{J}) = \mu$, natomiast wzór (III.2.8) wynika wprost ze wzoru (I.2.11) przez warunkowanie względem $\mathbf{J}$. Wzór (III.2.9) na nieobciążony estymator wariancji powstaje z kombinacji wzorów (III.2.7) i (III.2.8). ■

III.3. Imputacja regresyjna

W modelu z podrozdziału I.3 zakładamy, że odpowiedzi i braki odpowiedzi opisane są wektorem $\mathbf{J}$.

TWIERDZENIE III.3.1. *Imputacyjny estymator wartości oczekiwanej w przypadku imputacji regresyjnej ma postać*

$$\hat{\mu}_{Imp} = \bar{X}_R + \frac{n-|\mathbf{J}|}{n} \lambda (\bar{Y}_{R^c} - v), \quad (\text{III.3.1})$$

gdzie $\bar{Y}_{R^c} = \frac{1}{n-|\mathbf{J}|} \sum_{i \in R^c} Y_i$.

Jego wartość oczekiwana oraz wariancja wynoszą odpowiednio

$$E\hat{\mu}_{Imp} = \mu \text{ oraz } \text{Var}\hat{\mu}_{Imp} = \left(E \frac{1}{|\mathbf{J}|} + \chi^2 \frac{n-E|\mathbf{J}|}{n^2} \right) \sigma^2. \quad (\text{III.3.2})$$

W przypadku $p_i = p$, $i = 1, \dots, n$

$$\text{Var}\hat{\mu}_{Imp} = \left(E \frac{1}{|\mathbf{J}|} + \chi^2 \frac{q_n - p}{nq_n} \right) \sigma^2.$$

Dowód. Wzór (III.3.1) wynika wprost ze wzoru (I.3.1).

Z pierwszego ze wzorów (I.3.2) mamy $E(\hat{\mu}_{Imp}|\mathbf{J}) = \mu$, więc pierwszy ze wzorów (III.3.2) jest konsekwencją identyczności $E\hat{\mu}_{Imp} = EE(\hat{\mu}_{Imp}|\mathbf{J})$. Ponieważ $\text{Var}E(\hat{\mu}_{Imp}|\mathbf{J}) = 0$, więc

$$\text{Var}\hat{\mu}_{Imp} = E\text{Var}(\hat{\mu}_{Imp}|\mathbf{J}),$$

a z drugiego ze wzorów (I.3.2) mamy

$$\text{Var}(\hat{\mu}_{Imp}|\mathbf{J}) = \sigma^2 \left(\frac{1}{|\mathbf{J}|} + \chi^2 \frac{n - |\mathbf{J}|}{n^2} \right).$$

Stąd wynika drugi ze wzorów (III.3.2).■

TWIERDZENIE III.3.2. *Imputacyjny estymator wariancji w przypadku imputacji regresyjnej ma postać*

$$S_{Imp}^2 = \frac{(|\mathbf{J}| - 1)S_R^2 + \chi^2 \left((n - |\mathbf{J}| - 1)S_{R^c, Y}^2 + \frac{(n - |\mathbf{J}|)|\mathbf{J}|}{n} (\bar{Y}_{R^c} - \nu)^2 \right)}{n - 1},$$

gdzie $S_{R^c, Y}^2 = \frac{1}{n - |\mathbf{J}| - 1} \sum_{i \in R^c} (Y_i - \bar{Y}_{R^c})^2$.

Jego wartość oczekiwana wynosi

$$ES_{Imp}^2 = \left(\chi^2 - \frac{1}{n-1} + E|\mathbf{J}| \left(\frac{1}{n-1} - \frac{\chi^2}{n} \right) \right) \sigma^2. \quad (\text{III.3.3})$$

Nieobciążony imputacyjny estymator wariancji estymatora $\hat{\mu}_{Imp}$ ma postać

$$v^2(\hat{\mu}_{Imp}) = \frac{E \left(\frac{1}{|\mathbf{J}|} + \frac{n - E|\mathbf{J}|}{n^2} \chi^2 \right)}{\chi^2 - \frac{1}{n-1} + E|\mathbf{J}| \left(\frac{1}{n-1} - \frac{\chi^2}{n} \right)} S_{Imp}^2. \quad (\text{III.3.4})$$

W przypadku $p_i = p$, $i = 1, \dots, n$

$$ES_{Imp}^2 = \left(\chi^2 - \frac{1}{n-1} + \frac{np}{q_n} \left(\frac{1}{n-1} - \frac{\chi^2}{n} \right) \right) \sigma^2$$

oraz

$$v^2(\hat{\mu}_{Imp}) = \frac{E \frac{1}{|\mathbf{J}|} + \frac{q_n - p}{nq_n} \chi^2}{\chi^2 - \frac{1}{n-1} + \frac{np}{q_n} \left(\frac{1}{n-1} - \frac{\chi^2}{n} \right)} S_{Imp}^2.$$

Dowód. Zgodnie ze wzorem (I.3.4) mamy

$$S_{Imp}^2 = \frac{(|\mathbf{J}| - 1)S_R^2 + \lambda^2 \left((n - |\mathbf{J}| - 1)S_{R^c, Y}^2 + \frac{(n - |\mathbf{J}|)|\mathbf{J}|}{n} (\bar{Y}_{R^c} - v)^2 \right)}{n - 1}.$$

Wzór (I.3.5) implikuje

$$E(S_{Imp}^2 | \mathbf{J}) = \sigma^2 \left(\frac{|\mathbf{J}| - 1}{n - 1} + \left(1 - \frac{|\mathbf{J}|}{n} \right) \chi^2 \right),$$

co prowadzi natychmiast do (III.3.3). Wzór (III.3.4) jest natychmiastową konsekwencją wzorów (III.3.2) i (III.3.3). ■

IV. LOSOWY ZBIÓR BRAKÓW OBSERWACJI — IMPUTACJA WIELOKROTNA

IV.1. Losowanie spośród respondentów

Zgodnie ze wzorem (II.0.3) estymator imputacji wielokrotnej jest średnią z estymatorów imputacyjnych

$$\hat{\mu}_{MImp} = \frac{1}{m} \sum_{l=1}^m \hat{\mu}_{Imp}^{(l)},$$

czyli zgodnie z twierdzeniem II.1.1

$$\hat{\mu}_{MImp} = \frac{|\mathbf{J}|}{n} \bar{X}_R + \frac{1}{nm} \sum_{l=1}^m \sum_{j \in R^c} X_{K_j^{(l)}}.$$

TWIERDZENIE IV.1.1. *Wartość oczekiwana estymatora wielokrotnej imputacji wynosi*

$$E\hat{\mu}_{MImp} = \mu,$$

a wariancja ma postać

$$\text{Var}\hat{\mu}_{MImp} = \frac{\sigma^2}{mn^2} \left(n+1 - E|\mathbf{J}| - nE\frac{1}{|\mathbf{J}|} \right) + \sigma^2 E\frac{1}{|\mathbf{J}|}. \quad (\text{IV.1.1})$$

W przypadku $p_i = p, i = 1, \dots, n$

$$\text{Var}\hat{\mu}_{MImp} = \frac{\sigma^2}{mn^2} \left(n+1 - \frac{np}{q_n} - nE\frac{1}{|\mathbf{J}|} \right) + \sigma^2 E\frac{1}{|\mathbf{J}|}.$$

Dowód. Ze wzoru na wartość oczekiwaną estymatora imputacyjnego hot-deck wynika, że

$$E(\hat{\mu}_{MImp}|\mathbf{J}) = \mu.$$

Więc wzór na wartość oczekiwaną estymatora $\hat{\mu}_{Imp}$ jest konsekwencją identyczności

$$E\hat{\mu}_{MImp} = EE(\hat{\mu}_{MImp}|\mathbf{J}).$$

Zgodnie ze wzorem (II.1.2) mamy

$$\text{Var}(\hat{\mu}_{MImp}|\mathbf{J}) = \frac{\sigma^2}{mn} \left(\frac{mn-1}{|\mathbf{J}|} + \frac{n+1}{n} - \frac{|\mathbf{J}|}{n} \right).$$

Ponieważ

$$\text{Var}\hat{\mu}_{MImp} = E\text{Var}(\hat{\mu}_{MImp}|\mathbf{J}) + \text{Var}E(\hat{\mu}_{MImp}|\mathbf{J})$$

i drugi składnik po prawej stronie jest równy zero otrzymujemy wzór (IV.1.1).■

TWIERDZENIE IV.1.2. *Obciążenie estymatora Rubina wariancji $\hat{v}_{MImp}^2$ wynosi*

$$B\hat{v}_{MImp}^2 = E\hat{v}_{MImp}^2 - \text{Var}\hat{\mu}_{MImp} = -\frac{\sigma^2}{n(n-1)} \left(E|\mathbf{J}| + n(n+1)E\frac{1}{|\mathbf{J}|} - 2n+1 \right). \quad (\text{IV.1.2.})$$

Estymator nieobciążony wariancji estymatora $\hat{\mu}_{Mimp}$ ma postać

$$\hat{v}_*^2 = \frac{E \frac{1}{|J|}}{1 - E \frac{1}{|J|}} \left((n-1) \bar{U}_m - B_m \right) + \frac{1}{m} B_m.$$

Dowód. Dla estymatora Rubina wariancji $\hat{v}_{Mimp}^2$, określonego wzorami (II.0.7)—(II.0.9), w przypadku wielokrotnej imputacji hot-deck zgodnie ze wzorem z dowodu twierdzenia II.1.2 mamy

$$E(\bar{U}_m | J) = \frac{\sigma^2 (|J| - 1) (n(n+1) - |J|)}{|J| n^2 (n-1)} \quad \text{oraz} \quad E(B_m | J) = \frac{\sigma^2 (|J| - 1) (n - |J|)}{|J| n^2}.$$

Zatem

$$E(\hat{v}_{Mimp}^2 | J) = \frac{\sigma^2 (|J| - 1) (2n - |J|)}{|J| n(n-1)} + \frac{\sigma^2 (J - 1) (n - J)}{|J| m n^2}.$$

W konsekwencji

$$E \hat{v}_{Mimp}^2 = \frac{\sigma^2}{n(n-1)} \left(2n - 1 - E|J| - 2n E \frac{1}{|J|} \right) + \frac{\sigma^2}{m n^2} \left(n + 1 - E|J| - n E \frac{1}{|J|} \right).$$

Ostatecznie, zgodnie ze wzorem (III.2.6) obciążenie estymatora $\hat{v}_{Mimp}^2$ wynosi

$$E \hat{v}_{Mimp}^2 - \text{Var} \hat{\mu}_{Mimp} = \frac{\sigma^2}{n(n-1)} \left(2n - 1 - E|J| - 2n E \frac{1}{|J|} \right) - \sigma^2 E \frac{1}{|J|},$$

co po uproszczeniach daje wzór (IV.1.2).

Ponieważ $E B_m = \frac{\sigma^2}{n^2} \left(n + 1 - E|J| - n E \frac{1}{|J|} \right)$, więc ze wzoru (IV.1.1) wynika, że wystarczy pokazać równość

$$E((n-1) \bar{U}_m - B_m) = \sigma^2 \left(1 - E \frac{1}{|J|} \right),$$

a to wynika wprost ze wzorów na $\bar{U}_m$ i B_m . ■

Wniosek IV.1.3. *Asymptotycznie względne obciążenie estymatora Rubina $\hat{v}_{MImp}^2$ wynosi*

$$\lim_{n \rightarrow \infty} \frac{E\hat{v}_{MImp}^2 - \text{Var}\hat{\mu}_{MImp}}{\text{Var}\hat{\mu}_{MImp}} = -\frac{(1-p)^2}{1 + \frac{p(1-p)}{m}}. \quad (\text{IV.1.3})$$

Dowód. Ze wzorów (IV.1.1) i (IV.1.2) wynika, że wzór na względne obciążenie można napisać w postaci

$$K_n = -\frac{E|J| + n(n+1)E\frac{1}{|J|} - 2n + 1}{E\frac{n(n-1)}{|J|} + \frac{n-1}{mn}\left(n+1 - E|J| - nE\frac{1}{|J|}\right)}.$$

Ponieważ (Marciniak, Wesołowski, 1999)

$$\lim_{n \rightarrow \infty} nE\frac{1}{|J|} = \frac{1}{p}, \quad \text{czyli} \quad \lim_{n \rightarrow \infty} E\frac{1}{|J|} = 0,$$

więc dzieląc licznik i mianownik przez n i wykorzystując fakt, że

$$\lim_{n \rightarrow \infty} \frac{1}{n} E|J| = p$$

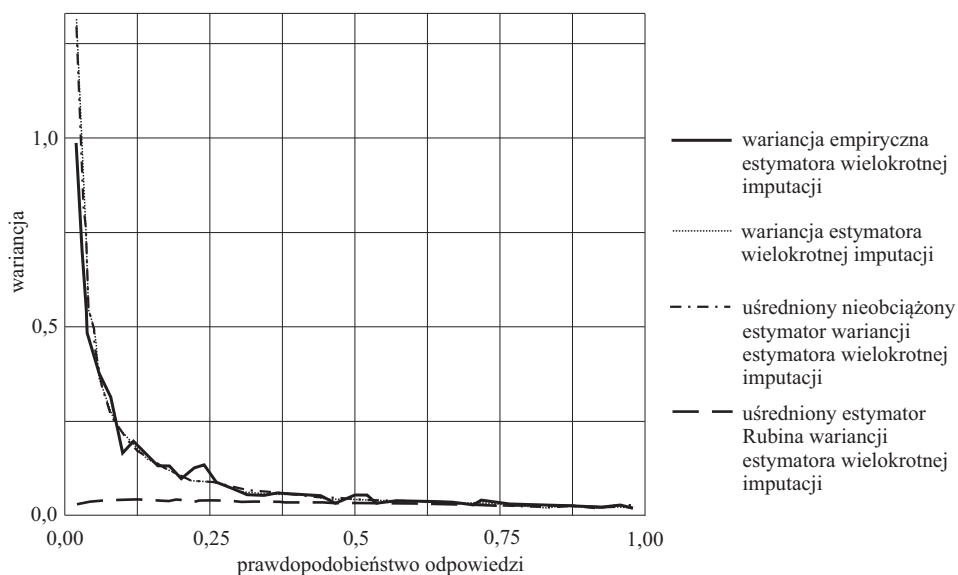
otrzymujemy

$$\begin{aligned} \lim_{n \rightarrow \infty} K_n &= -\frac{\lim_{n \rightarrow \infty} \frac{1}{n} E|J| + \lim_{n \rightarrow \infty} (n+1)E\frac{1}{|J|} - 2 + \lim_{n \rightarrow \infty} \frac{1}{n}}{\lim_{n \rightarrow \infty} (n-1)E\frac{1}{|J|} + \frac{1}{m}\left(\lim_{n \rightarrow \infty} \frac{n-1}{n}\right)\left(1 + \lim_{n \rightarrow \infty} \frac{1}{n} - \lim_{n \rightarrow \infty} \frac{1}{n} E|J| - \lim_{n \rightarrow \infty} E\frac{1}{|J|}\right)} = \\ &= \frac{p + \frac{1}{p} - 2}{\frac{1}{p} + \frac{1}{m}(1-p)}, \end{aligned}$$

co po uproszczeniach daje wzór (IV.1.3).■

Symulacyjna ilustracja wyników zawartych w twierdzeniach IV.1.1 i IV.1.2 przedstawiona jest na wykr. 3. Szczegółowy opis procedury symulacyjnej znajduje się w aneksie.

Wykr. 3. WYNIKI ESTYMACJI WARIANCJI ESTYMATORA ŚREDNIEJ DLA WIELOKROTNEJ IMPUTACJI (HOT-DECK) PRZY LOSOWYM BRAKU OBSERWACJI



Źródło: jak przy wykr. 1.

IV.2. Losowanie z rozkładu normalnego

TWIERDZENIE IV.2.1. *Wartość oczekiwana estymatora imputacji wielokrotnej przy losowaniu z rozkładu normalnego wynosi*

$$E\hat{\mu}_{MImp} = \mu,$$

a jego wariancja ma postać

$$\text{Var}\hat{\mu}_{MImp} = \sigma^2 \left(E \frac{1}{|J|} + \frac{n - E|J|}{mn^2} \right). \quad (\text{IV.2.1})$$

Dowód. Powyższe wzory wynikają wprost ze wzorów z twierdzenia II.1.4. ■

TWIERDZENIE IV.2.2. *Obciążenie estymatora Rubina wariancji $\hat{v}_{MImp}^2$ wynosi*

$$B\hat{v}_{MImp}^2 = E\hat{v}_{MImp}^2 - \text{Var}\hat{\mu}_{MImp} = -\frac{\sigma^2}{n^2} \left(\frac{n - E|\mathbf{J}|}{n-1} - n^2 E \frac{1}{|\mathbf{J}|} - E|\mathbf{J}| + 2n \right). \quad (\text{IV.2.2})$$

Estymator nieobciążony wariancji estymatora $\hat{\mu}_{MImp}$ ma postać

$$\hat{v}_*^2 = \left(n E \frac{1}{|\mathbf{J}|} \right) \bar{U}_m + \left(\frac{1}{m} + \frac{n}{n-1} E \frac{1}{|\mathbf{J}|} \right) B_m.$$

Dowód. Zgodnie ze wzorami (II.1.1) i (II.1.2) mamy

$$E\bar{U}_m = \frac{\sigma^2}{n^2} \left(n - \frac{n - E|\mathbf{J}|}{n-1} \right) \quad \text{oraz} \quad EB_m = \frac{\sigma^2}{n^2} (n - E|\mathbf{J}|).$$

W konsekwencji

$$E\hat{v}_{MImp}^2 - \text{Var}\hat{\mu}_{MImp} = \frac{\sigma^2}{n^2} \left(n - \frac{n - E|\mathbf{J}|}{n-1} + n - E|\mathbf{J}| \right) - \sigma^2 E \frac{1}{|\mathbf{J}|},$$

co prowadzi do wzoru (IV.2.2).

Ze wzoru (IV.2.1) wynika, że wystarczy znaleźć liczby α i β takie, że

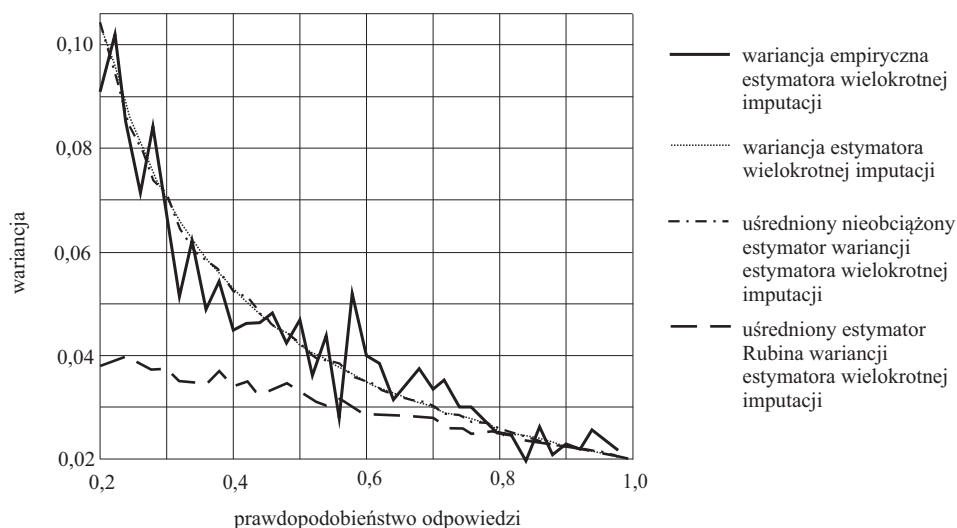
$$\alpha E\bar{U}_m + \beta EB_m = \sigma^2 E \frac{1}{|\mathbf{J}|}.$$

Wtedy $\alpha \bar{U}_m + \left(\beta + \frac{1}{m} \right) B_m$ jest poszukiwanym estymatorem nieobciążonym.

Przyjmując $\alpha = n E \frac{1}{|\mathbf{J}|}$ z powyższej równości otrzymujemy $\beta = \frac{n}{n-1} E \frac{1}{|\mathbf{J}|}$. ■

Symulacyjna ilustracja wyników zawartych w twierdzeniach IV.2.1 i IV.2.2 przedstawiona jest na wykr. 4. Szczegółowy opis procedury symulacyjnej znajduje się w aneksie.

**Wykr. 4. WYNIKI ESTYMACJI WARIANCJI ESTYMATORA ŚREDNIEJ
DLA WIELOKROTNEJ IMPUTACJI (HOT-DECK)
Z ROZKŁADU NORMALNEGO I LOSOWYM BRAKU OBSERWACJI**



Źródło: jak przy wykr. 1.

V. KONKLUZJE

W powyższym opracowaniu przedstawiono podstawowe idee i modele imputacji, w tym imputacji wielokrotnej na podstawie próbki niezależnych obserwacji o jednakowym rozkładzie, w której występują braki. W szczególności rozważono imputację średnią, imputacje typu hot-deck polegające na losowaniu respondentów oraz losowaniu z rozkładu normalnego z parametrami estymowanymi na podstawie obserwowanej części próbki oraz imputację regresyjną. Każda z tych metod jest szczególnym przypadkiem podejścia ogólnego analizowanego w rozdziałach I i III. W przypadku imputacji wielokrotnej wskazano, że popularny estymator wariancji, zwany estymatorem Rubina, nie ma dobrych własności, w szczególności jest obciążony. Zaproponowano wersje nieobciążone estymatorów typu Rubina w przypadku imputacji typu hot-deck metodą losowania respondentów i losowania z rozkładu normalnego. W artykule nie analizowano bayesowskiego modelu obserwacji, który jest ważnym elementem, zaproponowanego przez Rubina, podejścia wykorzystującego imputację wielokrotną. Planujemy przedstawić matematyczne podstawy imputacji w modelu bayesowskim w kolejnym opracowaniu.

LITERATURA

- Andridge, R.R., Little, R.J.A. (2010). A review of hot deck imputation survey non-response. *International Statistical Review*, vol. 78, no. 1, s. 40—64.
- Carpenter, J.R., Kenward, M.G. (2013). *Multiple Imputation and its Application*. Wiley, Chichester.
- de Waal, T., Pannekoek, J., Scholtus, S. (2011). *Handbook of Statistical Data Editing and Imputation*. Wiley, New York.
- Donders, A.R.T., van der Heijden, G.J.M.G., Stijnen, T., Moons, K.G.M. (2006). Review: a gentle introduction to imputation of missing values. *Journal of Clinical Epidemiology*, vol. 59, no. 10, s. 1087—1091.
- Garson, G.D. (2012). Missing Values Analysis & Data Imputation. *Statistical Associates Publishers*, Asheboro.
- Horton, N.J., Lipsitz, S.R. (2001). Multiple imputation in practice: comparison of software packages for regression models with missing data. *American Statistician*, vol. 55, no. 3, s. 244—254.
- Little, R.J.A., Rubin, D.B. (2002). *Statistical Analysis with Missing Data*. Wiley, New York.
- Marciniak, E., Wesołowski, J. (1999). Asymptotic Eulerian expansions for binomial and negative binomial reciprocals. *Proceedings of The American Mathematical Society*, vol. 127, s. 3329—3338.
- Misztal, M. (2012). Imputation of missing data using R package. *Acta Universitatis Lodzensis, Folia Oeconomica*, vol. 269, s. 131—144.
- Norazian Ramli, M.N., Yahaya, A.S., Ramli, N.A., Yusof, N.F.F.M., Abdullah, M.M.A. (2013). Roles of imputation methods for filling the missing values: a review. *Advances in Environmental Biology*, vol. 7, no. 12, s. 3861—3869.
- Rempała, G.A. (2004). Asymptotic factorial powers expansions for binomial and negative binomial reciprocals. *Proceedings of The American Mathematical Society*, vol. 32, no. 1, s. 261—272.
- Rubin, D.B. (1987). *Multiple Imputation for Nonresponse in Surveys*. Wiley, New York.
- van Buuren, S. (2012). *Flexible Imputation of Missing Data*. Chapman&Hall/CRC, London.
- van Buuren, S., Groothuis-Oudshoorn, K. (2011). mice: multivariate imputation by chained equations in R. *Journal of Statistical Software*, vol. 45, no. 3, s. 1—67, <http://www.jstatsoft.org/>.
- Schafer, J.L. (1997). *Analysis of Incomplete Multivariate Data*. Chapman and Hall, Boston.
- Szabłowski, P.J., Wesołowski, J., Wieczorkowski, R. (1996). Estymacja w podpopulacjach. *Wiadomości Statystyczne*, nr 7, s. 1—13.

ANEKS

Badania symulacyjne

Poniżej przedstawiony został opis badań symulacyjnych przeprowadzonych nad rozważanymi w pracy estymatorami wariancji wykorzystującymi imputację wielokrotną metodą typu hot-deck.

Symulacje polegały na $L = 100$ -krotnym powtórzeniu następującego algorytmu:

- 1) generowano dane pochodzące z centralnego rozkładu normalnego z wariancją $\sigma^2 = 4$;
- 2) usuwano (losowo bądź nie, w zależności od rozważanego modelu) obserwacje z wygenerowanych danych;
- 3) generowano $m = 5$ próbek imputacyjnych w sposób określony przez rozważany model imputacyjny;

- 4) wyliczano estymator $\hat{\mu}_{MImp}$ oraz jego estymatory wariancji (Rubina — $\hat{v}_{MImp}^2$ i nieobciążony — $\hat{v}_*^2$) na podstawie wygenerowanych próbek imputacyjnych. Na wszystkich rysunkach przedstawiono wykresy następujących wielkości:
 — $\widehat{\text{Var}}(\hat{\mu}_{MImp})$ — wariancji empirycznej estymatora wielokrotnej imputacji $\hat{\mu}_{MImp}$, obliczanej na podstawie L powtórzeń eksperymentu

$$\widehat{\text{Var}}(\hat{\mu}_{MImp}) = \frac{1}{L-1} \sum_{l=1}^L \left(\hat{\mu}_{MImp}^{(l)} - \frac{1}{L} \sum_{j=1}^L \hat{\mu}_{MImp}^{(j)} \right)^2,$$

gdzie $\hat{\mu}_{MImp}^{(l)}$ oznacza wartość estymatora wielokrotnej imputacji $\hat{\mu}_{MImp}$ w l -tym powtórzeniu eksperymentu, $l=1, \dots, L$;

- $\widehat{\text{Var}}\hat{\mu}_{MImp}$ — wariancji estymatora wielokrotnej estymacji, obliczanej na podstawie wzorów: II.1.2 (wykr. 1), II.1.7 (wykr. 2), IV.1.1 w wersji $p_i = p$ (wykr. 3) oraz IV.2.1 (wykr. 4);
 — $\overline{\hat{v}_*^2}$ — uśrednionej, na podstawie L powtórzeń eksperymentu, wartości $\hat{v}_*^2$ nieobciążonego estymatora wariancji estymatora wielokrotnej estymacji:

$$\overline{\hat{v}_*^2} = \frac{1}{L} \sum_{l=1}^L \hat{v}_{*,l}^2,$$

gdzie $\hat{v}_{*,l}^2$ oznacza wartość estymatora $\hat{v}_*^2$ w l -tym powtórzeniu eksperymentu, $l = 1, \dots, L$;

- $\overline{\hat{v}_{MImp}^2}$ — uśrednionej, na podstawie L powtórzeń eksperymentu, wartości $\hat{v}_{MImp}^2$ estymatora Rubina wariancji estymatora wielokrotnej imputacji:

$$\overline{\hat{v}_{MImp}^2} = \frac{1}{L} \sum_{l=1}^L \hat{v}_{MImp,l}^2,$$

gdzie $\hat{v}_{MImp,l}^2$ oznacza wartość estymatora Rubina $\hat{v}_{MImp}^2$ w l -tym powtórzeniu eksperymentu, $l = 1, \dots, L$.

Summary. *The article presents the basics of imputation methodology (including the methodology of multiple imputation), focusing on understanding its mathematical background. We analyze the situation when observations in the original sample are independent random variables with identical distributions,*

and response or its lack is modeled by a random mechanism which is independent of observations. In particular, we point out to problems that arise when the standard Rubin estimate of the multiple imputation variance estimator is used. A possible improvement of this popular estimator is indicated. The starting point of the analysis is when the appearance of response deficiencies is caused by a deterministic mechanism.

Keywords: imputation, multiple imputation, imputation estimator, Rubin estimator, mean imputation, hot-deck imputation, regression imputation.

Резюме. В статье представлены основы импутационной методологии (в том числе методологии многократной импутации). Внимание в статье сосредоточено на прояснении математической стороны вопросов. Проанализирована ситуация, когда наблюдения формирующие оригинальную выборку являются независимыми случайными величинами с одинаковыми распределениями, а отсутствие ответов появляется случайно независимо от наблюдения. В частности статья указывает на проблемы, которые возникают когда используется стандартная оценка Рубина дисперсии оценки многократной импутации. В статье указано также на возможное улучшение этой популярной оценки. Отправной точкой анализа является ситуация, когда отсутствие ответов объясняет детерминистический механизм.

Ключевые слова: импутация, многократная импутация, импутационная оценка, оценка Рубина, импутация средним, импутация типа hot-deck, регрессионная импутация.

Anna PRAŻMO
Joanna WÓJCIK
Magdalena ŻERO

Wyzwania statystyki publicznej w świetle Agendy na rzecz Zrównoważonego Rozwoju 2030

Streszczenie. *Celem artykułu jest zaprezentowanie nowej Agendy przyjętej przez ONZ w zakresie zrównoważonego rozwoju świata, a także wyzwań stojących przed statystyką publiczną, związanych z monitorowaniem realizacji wyznaczonych celów w ujęciu globalnym, regionalnym oraz krajowym. Międzynarodowe gremia statystyki publicznej uczestniczą w formułowaniu celów, do których świat będzie dążyć. Do monitorowania realizacji tych celów zobowiązano urzędy statystyczne poszczególnych krajów. Sprostanie tym oczekiwaniom będzie wymagało od międzynarodowych i krajowych instytucji statystycznych organizacji sprawnego systemu monitorowania oraz wypełnienia luk informacyjnych. Wyzwania te stoją również przed GUS — koordynatorem monitorowania realizacji celów rozwojowych w Polsce.*

Słowa kluczowe: zrównoważony rozwój, Agenda 2030, monitorowanie, wskaźniki zrównoważonego rozwoju, wyzwania statystyki publicznej.

W dobie szybkiego rozwoju technologicznego oraz postępującej globalizacji współczesny świat staje przed nowymi, wieloaspektowymi wyzwaniami. Wśród tych szczególnych, na których w perspektywie długookresowej powinna być skupiona uwaga społeczeństwa, jest zrównoważony rozwój, obejmujący trzy główne aspekty: społeczny, gospodarczy i środowiskowy. Kierunki planowanych działań o tak szerokim zakresie najczęściej przedstawiane są w formie dokumentów strategicznych. Większość krajów oraz organizacji międzynarodowych posiada strategie zrównoważonego rozwoju, realizowane w perspektywie długoterminowej. Wyznaczane są w nich cele oraz definiowane zadania i narzędzia, dzięki którym idea zrównoważonego rozwoju może być zrealizowana. Realizacja wytyczonych celów jest możliwa dzięki współpracy wielu grup inte-

resariuszy. Nieodłącznym elementem realizacji celów jest ich monitorowanie za pomocą rzetelnych i obiektywnych danych. Wskazują one na stopień realizacji i postępu w uzyskiwaniu założonego efektu. W ostatnim czasie wynikiem takiej współpracy jest przyjęty przez ONZ długofalowy plan rozwoju świata — Agenda na rzecz Zrównoważonego Rozwoju 2030¹. Opracowanie i przyjęcie dokumentu, obejmującego tak szerokie spektrum zagadnień, było możliwe dzięki zaangażowaniu i efektywnemu współdziałaniu różnych organizacji oraz rządów krajów na poziomach — międzynarodowym, regionalnym oraz krajowym. Prowadzenie efektywnych działań na rzecz realizacji Agendy 2030 wymaga zapewnienia globalnego partnerstwa, tj. obejmującego ściśle współpracę wspomnianych jednostek i gremiów. Wśród nich jedną z głównych ról pełnić będzie statystyka publiczna jako gwarant aktualnych, rzetelnych, dostępnych i wysokiej jakości danych. Wskaźniki opracowywane przez statystykę publiczną umożliwiają monitorowanie postępów we wdrażaniu całej Agendy.

IDEA ZRÓWNOWAŻONEGO ROZWOJU ORAZ KAMIENIE MIŁOWE W JEJ IMPLEMENTACJI

U podstaw idei zrównoważonego rozwoju leżą ograniczone możliwości zaspokajania różnorodnych potrzeb ludzkości, co wynika przede wszystkim z szybko uszczuplających się zasobów naturalnych na świecie. Początki tej koncepcji są zwykle wiązane z raportami Klubu Rzymskiego, który w centrum zainteresowań stawiał przyszłość ludzkości oraz planety w perspektywie długookresowej. W pierwszym raporcie tej organizacji pt. *Granice wzrostu*² zwrócono uwagę na fakt, że ludzkość dąży do ciągłego wzrostu populacji, zasiedlania terenów, zwiększania produkcji i konsumpcji, nie zważając na wpływ tej ekspansji na środowisko. Wśród konsekwencji takiego postępowania wskazano zagrożenie załamaniem gospodarczym na świecie i wyczerpywaniem się zasobów naturalnych. Wtedy zrównoważony rozwój oznaczał głównie konieczność skoncentrowania działań na ograniczeniu szkodliwego wpływu działalności produkcyjnej na środowisko naturalne. Obecnie najczęściej stosowana jest definicja zrównoważonego rozwoju opracowana przez Światową Komisję ds. Środowiska i Rozwoju³, która zwróciła uwagę na konieczność integracji działań w zakresie rozwoju społecznego, gospodarczego oraz środowiskowego. W raporcie pt. *Nasza wspólna przyszłość* (1987 r.) Komisja zdefiniowała zrównoważony rozwój jako **proces przemian, który zapewnia zaspokajanie potrzeb obecnego pokolenia bez umniejszania szans rozwojowych przyszłych generacji**⁴.

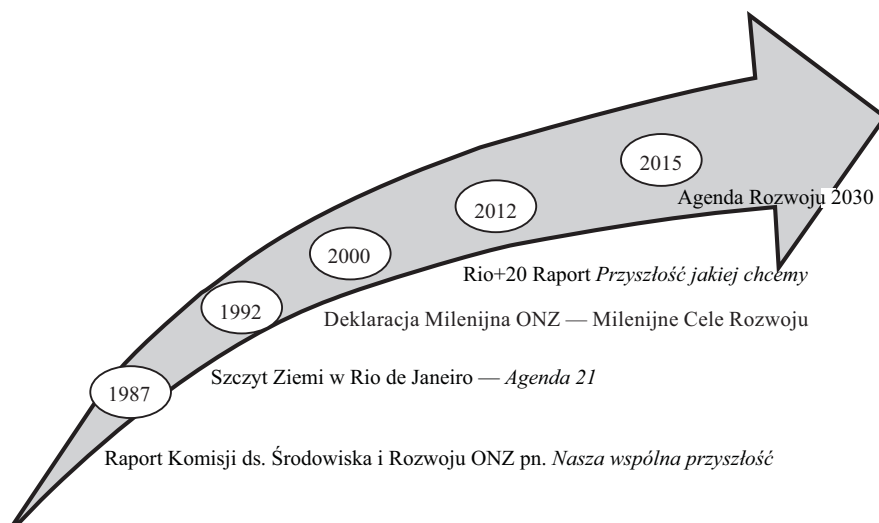
Prace dotyczące obecnych szans dla świata na zapewnienie zrównoważonego rozwoju toczyły się m.in. dzięki różnym inicjatywom międzynarodowym, w tym podejmowanym przez ONZ.

¹ Dokument określany również jako *Agenda post-2015*.

² *Limits to growth* (1972), raport Klubu Rzymskiego, <http://www.donellameadows.org/wp-content/userfiles/Limits-to-Growth-digital-scan-version.pdf>.

³ Komisja powołana w 1983 r. przy ONZ przez *Gro Harlem Brundtland*, zwana Komisją Brundtland.

⁴ Przytoczoną definicję na potrzeby artykułu przyjęto za podstawową.

**SCHEMAT 1. WAŻNIEJSZE INICJATYWY PODEJMOWANE
NA RZECZ ZRÓWNOWAŻONEGO ROZWOJU PRZEZ ONZ**

Źródło: opracowanie własne na podstawie dokumentów dotyczących inicjatyw ONZ w zakresie zrównoważonego rozwoju.

Oprócz ONZ inicjatywy w tym zakresie podejmowała również Komisja Europejska. Do najważniejszych z nich należą:

- przyjęta w 2001 r. (odnowiona w 2006 r.) Strategia Zrównoważonego Rozwoju Unii Europejskiej (UE), mająca na celu wdrażanie działań umożliwiających obywatelom Wspólnoty stałą poprawę jakości życia zarówno obecnego, jak i przyszłych pokoleń, dzięki tworzeniu społeczeństwa rozważnie gospodarującego zasobami, umiejących wykorzystać potencjał społeczny oraz środowiskowy gospodarki;
- ogłoszona w 2007 r. inicjatywa „Wyjść poza PKB”, która zwraca uwagę na potrzebę rozwoju wskaźników przedstawiających zarówno aspekt ekonomiczny rozwoju, jak również uwzględniających czynniki społeczne i środowiskowe wpływające na rozwój;
- przyjęta w 2010 r. strategia *Europa 2020*, która za cel główny stawia zrównoważony wzrost gospodarczy, osiągnięty dzięki niskoemisyjnej gospodarce opartej na wiedzy, promującej przyjazne środowisku technologie, zapewniającej nowe „zielone miejsca pracy”⁵ oraz spójność społeczną.

Do kwestii pomiaru zrównoważonego rozwoju oraz jakości życia odnosił się także, przedstawiony w 2009 r., Raport tzw. Komisji Stiglitz’a. Jednym z celów

⁵ W 1999 r. OECD oraz Eurostat zaproponowały, by w kontekście „zielonych miejsc pracy” mówić o „czynnościach, które służą tworzeniu dóbr oraz usług w celu określenia wielkości, zapobiegania, ograniczenia, zminimalizowania lub naprawy szkód w środowisku związanych z wodą, powietrzem oraz glebą, jak również aktywności związanej z problematyką odpadów, hałasu czy ekosystemów. *Measuring Green Entrepreneurship w Entrepreneurship at a Glance 2011*, OECD.

Raportu była identyfikacja ograniczeń PKB (traktowanego zwykle jako miernik postępu gospodarczego) oraz ocena możliwości zastosowania wskaźników go uzupełniających lub do niego alternatywnych, w celu lepszego przedstawienia jakości życia ludzi. Raport zawierał 12 rekomendacji w zakresie lepszego pomiaru m.in. wyników gospodarczych oraz dobrobytu społecznego. Zrównoważony rozwój m.in. w kontekście Strategii Zielonego Wzrostu promuje również OECD.

Najnowszym i dotychczas najbardziej kompleksowym przedsięwzięciem jest opracowana pod auspicjami ONZ Agenda 2030, której punkt wyjścia stanowiła Deklaracja Milenijna i związane z nią Milenijne Cele Rozwoju (obowiązujące do 2015 r.).

AGENDA NA RZECZ ZRÓWNOWAŻONEGO ROZWOJU 2030

Agenda na rzecz Zrównoważonego Rozwoju 2030 (*The 2030 Agenda for Sustainable Development*) to ramowy plan dla świata, wyznaczający globalne cele w perspektywie do 2030 r. Dokument ten wraz z zestawem 17 celów zrównoważonego rozwoju (*Sustainable Development Goals* — SDGs) stanowiących jego trzon przyjęły 193 państwa członkowskie ONZ podczas specjalnego Szczytu Zrównoważonego Rozwoju w Nowym Jorku (25—27 września 2015 r.). Nowe cele, obowiązujące od 2016 r. i zastępujące Milenijne Cele Rozwoju (*Millennium Development Goals* — MDGs), obejmują wszystkie kraje niezależnie od stopnia ich rozwoju (odmiennie niż w przypadku celów milenijnych, które w praktyce odnosiły się do krajów rozwijających się).

Agenda 2030 ma na celu dążenie do stworzenia sprawiedliwego świata, opartego na poszanowaniu prawa i rządach sprzyjających włączeniu społecznemu. Takie podejście zobowiązuje różne podmioty do współpracy ukierunkowanej na promowanie zachowań pozwalających na wzrost gospodarczy, rozwój społeczny i ochronę środowiska jednocześnie. Jej fundamentalnym przesłaniem jest dążenie do rozwoju, który zagwarantuje godne życie dla wszystkich. Urzeczywistnieniu tego założenia mają służyć przede wszystkim eliminacja ubóstwa we wszystkich jego formach i wymiarach, w tym skrajnego, a także zapewnienie pokoju. Agenda 2030 ma przynieść korzyści zarówno pokoleniom obecnym (włączając różne grupy społeczne: kobiety, dzieci, osoby starsze i osoby niepełnosprawne), jak i przyszłym generacjom. Ma być realizowana w sposób spójny z obecnymi zobowiązaniami państw członkowskich oraz z poszanowaniem prawa międzynarodowego. Uniwersalny charakter Agendy wymaga zintegrowanego podejścia do zrównoważonego rozwoju oraz zbiorowego działania na poziomach: globalnym, regionalnym, krajowym oraz niższych, aby sprostać wyzwaniom naszych czasów, a przede wszystkim aby zapewnić, że nikt nie zostanie pominięty (*leaving no one behind*⁶). Zidentyfikowanie nierówności i dyskryminacji będzie stanowić podstawę w definiowaniu przyszłych działań. Nowy plan rozwojowy wymaga wzmoczonych działań ze strony wielu podmiotów: organizacji międzynarodowych, rządów, środowiska naukowego oraz społeczeństwa

⁶ Zapis ten zawarto w preambule dokumentu Agendy 2030: *As we embark on this collective journey, we pledge that no one will be left behind.*

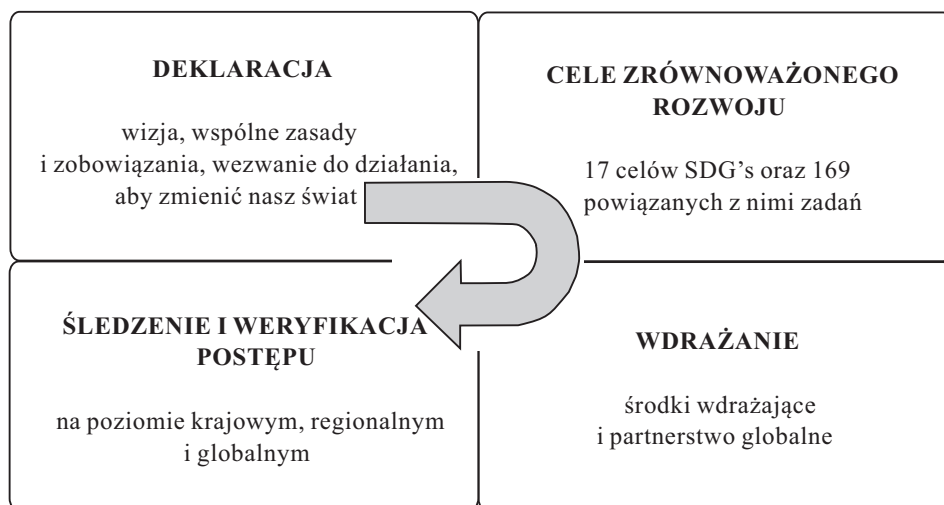
obywatelskiego. Jest również źródłem nowych wyzwań stawianych przed statystyką publiczną, którym będzie musiała sprostać jako podmiot zobowiązany do monitorowania postępów we wdrażaniu Agendy.

STRUKTURA AGENDY 2030

Na strukturę Agendy na rzecz Zrównoważonego Rozwoju 2030 składają się cztery główne elementy:

- 1) wizja i podstawowe zasady zawarte w Deklaracji;
- 2) uzgodnione cele zrównoważonego rozwoju i przypisane im zadania;
- 3) środki wdrażające oraz Globalne Partnerstwo;
- 4) monitorowanie i raportowanie o postępach we wdrażaniu Agendy 2030.

SCHEMAT 2. STRUKTURA AGENDY NA RZECZ ZRÓWNOWAŻONEGO ROZWOJU 2030



Źródło: opracowanie własne na podstawie *Mainstreaming the 2030 Agenda for Sustainable Development. Interim Reference Guide to UN Country Teams* (2015), ONZ.

Zasadniczą część Agendy stanowi 17 celów oraz 169 powiązanych z nimi zadań. Ich ustalanie zainicjowano w 2012 r. w Rio de Janeiro, podczas Konferencji Narodów Zjednoczonych w sprawie Zrównoważonego Rozwoju (Rio+20). W wyniku ponad trzyletnich negocjacji, w których uwzględniono opinie rządów państw, sektora prywatnego, środowiska naukowego oraz społeczeństwa obywatelskiego, przedstawiono katalog celów będący wynikiem kompromisu pomiędzy 193 państwami członkowskimi ONZ.

Cele Agendy stanowią wyraz wspólnych, globalnych priorytetów, których przyjęcie oznacza potwierdzenie gotowości wszystkich państw sygnatariuszy do ich wypełniania. Są one kontynuacją i uzupełnieniem celów milenijnych i będą stymulować działania w ciągu kolejnych 15 lat w dziedzinach o istotnym zna-

czeniu dla rozwoju świata — skupiając się na ludziach, planecie, dobrobycie, pokoju na świecie oraz partnerstwie. W założeniu mają równoważyć trzy wymiary rozwoju: gospodarczy, społeczny i środowiskowy. Chociaż poszczególne cele SDGs są od siebie niezależne, powinny być realizowane w sposób zintegrowany. Aby to osiągnąć państwa członkowskie są zachęcane do wyznaczania celów na szczeblu krajowym, odpowiadających globalnym ambicjom, ale uwzględniających krajowe realia i poziom rozwoju. Rządy będą mogły również zdecydować, w jaki sposób cele te powinny być włączone do krajowego planowania procesów, polityki i strategii.

MONITOROWANIE I RAPORTOWANIE O POSTĘPACH WE WDRAŻANIU CELÓW AGENDY 2030

Monitorowanie i raportowanie rezultatów Agendy na rzecz Zrównoważonego Rozwoju 2030 powinno być oparte na następujących zasadach: uniwersalny charakter, transparentność, wiodąca rola państw członkowskich, skuteczność i efektywność, cykliczność i wielopoziomowość oraz partycypacyjny i włączający różne środowiska charakter działań⁷. Ramy te powinny składać się z następujących priorytetów — monitorowania celów i zadań zrównoważonego rozwoju z wykorzystaniem zestawu wskaźników pozwalających ocenić ich rezultaty oraz oceny jakościowej postępów we wdrażaniu Agendy 2030, w tym wzajemne poznawanie przyczyn sukcesu oraz przeszkód w ich realizacji.

Zgodnie z art. 21 Agendy 2030 założono trójszczeblowy poziom monitorowania wdrażania celów:

- globalny — za koordynację którego odpowiadać będzie Komisja Statystyczna ONZ;
- regionalny⁸ — z wiodącą rolą regionalnych organizacji i agend ONZ;
- krajowy — za który odpowiedzialne będą krajowe urzędy statystyczne.

MONITOROWANIE NA POZIOMIE GLOBALNYM

Monitorowanie ustalonych celów zrównoważonego rozwoju na poziomie globalnym odbywać się będzie przy pomocy zestawu wskaźników globalnych, zgodnych z międzynarodowymi standardami oraz wyliczanych przy wykorzystaniu oficjalnych danych. Zadanie opracowania globalnego systemu monitorowania powierzono Grupie eksperckiej ds. wskaźników monitorujących cele zrównoważonego rozwoju⁹ (*Inter-Agency and Expert Group on SDG Indicators* — IAEG-SDGs). W wyniku jej prac przyjęto 230 wskaźników monitorujących cele Agendy 2030. Listę tych wskaźników przyjęto w marcu 2016 r. podczas 47.

⁷ *Transforming our world: the 2030 Agenda for Sustainable Development* (2015), ONZ, art. 74.

⁸ Poziom regionalny dotyczy regionów ONZ i obejmuje grupy krajów. Polska należy do regionu europejskiego UNECE, tj. Europejskiej Komisji Gospodarczej ONZ.

⁹ Grupa powołana w 2015 r. przez Komisję Statystyczną ONZ. W jej skład weszło 28 państw członkowskich ONZ reprezentujących wszystkie regiony świata. Polska nie jest członkiem IAEG-SDGs. Kraje niebędące członkami grupy, a także organizacje międzynarodowe (w tym Eurostat), zostały uprawnione do uczestniczenia w pracach grupy w charakterze obserwatorów.

sesji plenarnej Komisji Statystycznej ONZ, a następnie skierowano do akceptacji Rady Społecznej i Gospodarczej ONZ oraz Zgromadzenia Ogólnego.

Uzgodniony zestaw mierników jest efektem współpracy szerokiego grona interesariuszy, w tym środowiska naukowego, urzędów statystycznych i organizacji pozarządowych. Opracowana lista ma charakter otwarty i będzie powiększana, a wynika to z braku dostępności niektórych proponowanych wskaźników. Opracowania wymagają przede wszystkim wskaźniki, które nie mają obecnie uzgodnionej międzynarodowej metodologii. Zakłada się, że dane dla monitorowania globalnego będą w największym, możliwym stopniu pobierane z baz międzynarodowych, co pozwoli ograniczyć obciążenia krajowych urzędów statystycznych.

Biorąc pod uwagę kryterium dostępności, przyjęte wskaźniki globalne pogrupowano według trzech poziomów (tzw. *tier-ów*¹⁰):

- I — wskaźniki z wypracowaną metodologią oraz powszechnie dostępnymi danymi (ok. 41% wszystkich wskaźników);
- II — wskaźniki, dla których metodologia została uzgodniona, ale dane nie są łatwo dostępne (ok. 21%);
- III — wskaźniki bez opracowanej metodologii (ok. 32%).

Pozostałe wskaźniki (ok. 6%) nie zostały przypisane do żadnego z *tier-ów* lub zakwalifikowano je do więcej niż jednego poziomu. Podczas 48. plenarnej sesji Komisji Statystycznej w 2017 r. grupa IAEG-SDGs ma przedstawić raport z prac dotyczących uzupełniania globalnej listy wskaźników, w szczególności niezbędnych w zakresie miar z poziomu III.

MONITOROWANIE NA POZIOMIE REGIONALNYM

Monitoring regionalny, jako łącznik pomiędzy szczeblem krajowym i globalnym, stanowi platformę wspierającą wymianę wiedzy i doświadczenia, przegląd partnerski oraz wzajemne uczenie się w poszczególnych regionach. Agendy ONZ i organizacje regionalne powinny zatem stymulować współpracę pomiędzy krajami w zakresie wdrażania i rozwijania systemu monitorowania na rzecz zrównoważonego rozwoju. Wskaźniki regionalne mają obejmować odpowiednie globalne wskaźniki, uzupełnione wskaźnikami krajowymi oraz ukierunkowanymi na konkretne priorytety regionalne.

Roli koordynatora monitorowania regionalnego dla Europy podjęła się Konferencja Statystyków Europejskich (*Conference of European Statisticians* — CES), działająca w strukturach Europejskiej Komisji Gospodarczej (*United Nations Economic Commission for Europe* — UNECE). Obecnie w CES toczą się prace nad przygotowaniem planu dotyczącego rozwoju statystyki publicznej w zakresie monitorowania zrównoważonego rozwoju w regionie UNECE (*Road*

¹⁰ Podział wskaźników według *tier-ów* przedstawiono w dokumencie: *Provisional Proposed Tiers for Global SDG Indicators*, dostępnym pod adresem: <http://unstats.un.org/sdgs/files/meetings/iaeg-sdgs-meeting-03/Provisional-Proposed-Tiers-for-SDG-Indicators-24-03-16.pdf>.

map for setting up the reporting on SDGs in the UNECE region). Do opracowania tego planu powołano Grupę Sterującą, w pracach której uczestniczy przedstawiciel GUS. Przygotowywany dokument ma informować o organizacji monitorowania w ujęciu regionalnym, w tym precyzować kwestie dotyczące raportowania danych, wzmacniania potencjału statystycznego oraz strategii komunikacji i współpracy, a prace nad nim mają się zakończyć do końca bieżącego roku.

MONITOROWANIE NA POZIOMIE KRAJOWYM

Monitorowanie rozwoju w kraju jest przywilejem państw członkowskich ONZ. Każde z państw ma zdefiniować krajowe priorytety, wzięwszy pod uwagę m.in. poziom rozwoju, możliwości, jakimi dysponuje oraz szanse ich osiągnięcia. Na podstawie ustalonych przez rządy priorytetów powstanie zestaw właściwych mierników z punktu widzenia potrzeby ich monitorowania.

W związku z tym, że określanie przez państwa priorytetów krajowych uwzględniających specyfikę tych krajów będzie odbywać się niezależnie, nie wszystkie wskaźniki będą porównywalne. Geograficzne położenie danego kraju, system polityczny i prawny, potencjał gospodarczy, zróżnicowanie społeczne i kulturowe to tylko niektóre kryteria brane pod uwagę przy definiowaniu i implementowaniu priorytetów. Cele te mogą być odmienne dla krajów przynależących do różnych regionów, jak i państw w jednym regionie, skupiających się na rozwijaniu rozmaitych dziedzin życia. Krajowe działania monitorujące mogą wykorzystywać wskaźniki już wcześniej przyjęte na poziomie globalnym, jeśli państwa uznają, że są one odpowiednie do monitorowania ich priorytetów. Zakłada się, że większość danych powinna pochodzić z oficjalnych źródeł.

PROCES RAPORTOWANIA¹¹

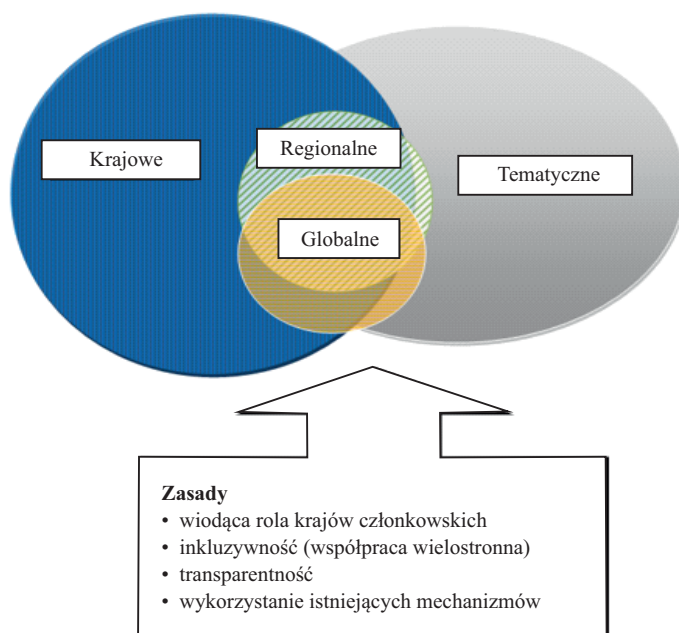
Raportowanie globalne będzie oparte na raportach regionalnych i krajowych oraz raportowaniu tematycznym (schemat 3).

Wiodącą rolę w raportowaniu globalnym ma pełnić, powołane w 2013 r., Forum polityczne wysokiego szczebla (*High Level Political Forum* — HLPF), działające w ścisłej współpracy ze Zgromadzeniem Ogólnym oraz Radą Społeczną i Gospodarczą ONZ. Podstawę informowania HLPF mają stanowić coroczne raporty z postępu we wdrażaniu celów zrównoważonego rozwoju (*Progress Report on the Sustainable Development Goals*). Raporty będą przygotowywane przez Sekretarza Generalnego we współpracy z systemem statystycznym Narodów Zjednoczonych na podstawie wskaźników globalnych, danych otrzymywanych z krajowych systemów statystycznych oraz informacji zebranych w regionach. Dodatkowym źródłem informacji będą Globalne raporty

¹¹ Opracowano na podstawie Raportu Sekretarza Generalnego ONZ *Critical milestones towards coherent, efficient and inclusive follow-up and review at the global level* (2016), raport Sekretarza Generalnego ONZ, http://www.un.org/ga/search/view_doc.asp?symbol=A%20/70/684&Lang=E.

zrównoważonego rozwoju (*Global Sustainable Development Report*), pozwalające na wzmocnienie kontaktów pomiędzy środowiskiem naukowym i politycznym oraz stanowiące instrument wspierający kształtowanie polityki rozwoju. Co cztery lata Forum polityczne wysokiego szczebla będzie spotykało się pod auspicjami Zgromadzenia Ogólnego, w celu dostarczenia wytycznych dotyczących implementacji Agendy, zidentyfikowania pojawiających się wyzwań oraz mobilizacji do działań, które przyspieszą realizację jej zapisów.

SCHEMAT 3. RAPORTOWANIE



Źródło: opracowanie własne na podstawie *Indicators and a Monitoring Framework for Sustainable Development Goals: Launching a data revolution for the SDGs* (2015), Sustainable Development Solutions Network.

Wyniki przeglądów prowadzonych na forach regionalnych będą przekazywane HLPF w formie zagregowanej. Kraje, które poddano równorzędnym lub innym przeglądom na poziomie regionalnym są zachęcane do korzystania z ich wyników podczas przygotowywania raportów krajowych.

Raportowanie krajów członkowskich odbywać się będzie na zasadzie dobrowoli. Zapisy Agendy 2030 nie przedstawiają szczegółów dotyczących okresów przeglądów krajowych, jednak zachęcają do regularnego ich dokonywania. Pozwoli to identyfikować zachodzące trendy i wyciągać z nich wnioski. Każdy kraj członkowski powinien natomiast do 2030 r. rozważyć przeprowadzenie dwóch przeglądów krajowych (w formie prezentacji na szczeblu ministerialnym) i przedstawić je podczas corocznych spotkań HLPF, odbywających się pod au-

spicjami Komisji Społecznej i Gospodarczej ONZ. Raportowanie podczas każdej sesji HLPF ma być tak zorganizowane, aby zapewniało udział przedstawicieli ONZ wszystkich regionów świata oraz obejmowało kraje znajdujące się na różnych etapach rozwoju. Prezentacje mają przedstawiać praktyki krajowe i główne wyzwania, przed którymi stoją państwa podczas realizacji Agendy 2030 oraz dziedziny, w których potrzebują wsparcia. W prezentacji krajowej powinno się również przedstawić w formie pisemnej szczegółowy raport podkreślający wnioski i główne przesłania wynikające z przeprowadzonej weryfikacji postępu. Podczas spotkania HLPF w czerwcu 2016 r. przeprowadzono pierwsze raportowanie krajowe (oparte na zasadzie dobrowolności) — 22 państwa podzieliły się doświadczeniami we wdrażaniu Agendy 2030 (m.in.: Francja, Finlandia, Niemcy, Chiny i Meksyk).

Tematyczne raportowanie z realizacji celów zrównoważonego rozwoju odbywać się będzie corocznie na Forum politycznym wysokiego szczebla i dotyczyć będzie zagadnień przekrojowych odnoszących się do więcej niż jednego celu. Przykłady takich tematów to m.in.: zapewnienie o przeciwdziałaniu wykluczeniu społecznemu, równouprawnienie kobiet i dziewcząt np. w dostępie do edukacji czy wzmocnienie partnerstwa globalnego na rzecz zrównoważonego rozwoju. Raportowanie tematyczne ma odbywać się w cyklu czteroletnim — w każdym roku rozpatrywane będzie jedno zagadnienie przekrojowe. Każdy cel zostanie przypisany do jednego z tematów (z wyjątkiem analizowanego co roku celu 17.). Pozwoli to omówić wszystkie cele zrównoważonego rozwoju w ciągu czterech lat, co będzie stanowić podstawę do ich kompleksowego przeglądu podczas spotkania HLPF pod auspicjami Zgromadzenia Ogólnego.

ROLA STATYSTYKI PUBLICZNEJ I WYZWANIA W ŚWIETLE AGENDY 2030

Osiągnięcie założonych w Agendzie 2030 celów wymaga współpracy i zaangażowania wszystkich interesariuszy, w tym decydentów, środowiska naukowego i społeczeństwa obywatelskiego. Ważnym filarem monitorowania realizacji zakładanych celów jest statystyka publiczna, dysponująca obszernymi zasobami rzetelnych, terminowych i wiarygodnych danych. Wysokiej jakości mierniki stanowią podstawę do dokonywania oceny postępu, ale także służą diagnozie określającej punkt startowy. Dane statystyczne są podstawowym źródłem polityki opartej na dowodach, wspierającym decydentów w ustalaniu właściwych kierunków rozwoju. Przykładem takiego wykorzystania informacji statystycznych jest strategia rozwoju, do której monitorowania organizacje czy też państwa je opracowujące wykorzystują zestawy wskaźników. Mają one uwidocznąć dokonywany postęp w danej dziedzinie oraz umożliwiają szybką modyfikację działań, w przypadku gdy kierunek zmian okazuje się niezgodny z oczekiwanym.

W Agendzie 2030 podkreślono istotność danych, wskazując w Artykule 48, że dobrej jakości, rzetelne, aktualne, wiarygodne dane są niezbędne do pomiaru postępu oraz podejmowania właściwych decyzji.

Polska statystyka publiczna aktywnie uczestniczyła we wszystkich rundach otwartych konsultacji, których celem było przygotowanie wskaźników globalnych. Rolę koordynatora w konsultacji globalnego zestawu wskaźników pełnił GUS, przy współpracy z właściwymi resortami.

Dotychczasowe prace GUS pozwalają zakładać, że w kolejnych etapach prac związanych z realizacją Agendy 2030 znacząca rola Urzędu będzie się pogłębiać. Wynika to m.in. z konieczności zapewnienia narzędzi do monitorowania priorytetów, które w nieodległej przyszłości zostaną zaimplementowane w Polsce.

Coraz większe potrzeby informacyjne, w tym szeroki zakres celów oraz ambitnych zadań sformułowanych w Agendzie 2030 powoduje, że monitorowanie jej wdrażania stanowi duże wyzwanie dla statystyki publicznej. Głównym zadaniem jest zapewnienie szerokiego zakresu informacji umożliwiających analizę retrospektywną, ocenę sytuacji bieżącej oraz śledzenie postępu w rozwoju. Powierzenie statystyce publicznej koordynacyjnej roli w zakresie zapewnienia sprawnego systemu monitorowania wiąże się z wieloma wyzwaniami, którym będą musiały sprostać międzynarodowe i krajowe instytucje statystyczne. Najpilniejsze do rozwiązania kwestie obejmują:

- 1) **identyfikację luk w zasobach informacyjnych** oraz dążenie do ich wypełniania. Szybko rosnące zapotrzebowanie na informacje, a także zmieniający się rodzaj pożądanych danych implikują coraz większą potrzebę skierowania wysiłków na zwiększenie adaptacyjności statystyki w stosunku do potrzeb różnego rodzaju użytkowników (w tym decydentów). Nieodzowne jest poszerzanie analizy o nowe miary postępu oraz wprowadzanie zmian w ich opracowywaniu;
- 2) stałe wysiłki w zakresie **uzupełniania zidentyfikowanych luk informacyjnych**, zarówno w zasobach krajowych jak również w bazach organizacji międzynarodowych (takich jak Eurostat czy OECD), głównie poprzez wprowadzanie nowych badań, rozwijanie istniejących oraz opracowywanie nowych wskaźników;
- 3) **wzmacnianie potencjału statystycznego** krajowych urzędów statystycznych, dzięki któremu możliwe będzie doskonalenie zasobów statystycznych (w tym badań bieżących) i uzupełnienie brakujących danych, jak również modernizacja procesu opracowywania danych;
- 4) **określenie i wdrożenie systemu raportowania** wskaźników w wymiarze globalnym i regionalnym poprzez wskazanie formy oraz określenie ewentualnego harmonogramu przekazywania danych.

Warunkiem realizacji tych wyzwań z punktu widzenia statystyki publicznej jest podjęcie konkretnych zadań, takich jak:

- aktywny udział statystyki publicznej w pracach międzynarodowych instytucji na wszystkich możliwych szczeblach, co zwiększy wpływ na podejmowane w ramach Agendy decyzje;
- ścisłą współpracę statystyki ze środowiskiem naukowo-badawczym, w celu rozwiązywania najważniejszych problemów metodologicznych oraz współpracę z innymi jednostkami, w celu wyeliminowania luk informacyjnych;

- korzystanie z doświadczeń i dorobku międzynarodowego, dzielenie się wiedzą w zakresie monitorowania zrównoważonego rozwoju;
- popularyzację idei zrównoważonego rozwoju poprzez upowszechnianie wyników Agendy 2030.

AKTYWNOŚĆ POLSKIEJ STATYSTYKI PUBLICZNEJ W ZAKRESIE MONITOROWANIA ZRÓWNOWAŻONEGO ROZWOJU

Aktywność polskiego systemu statystyki publicznej w zakresie monitorowania zrównoważonego rozwoju nie ogranicza się do działań dotyczących Agendy 2030. Dotychczas, oprócz bieżących zadań związanych np. z pracami wynikającymi ze współpracy międzynarodowej (m.in. w ramach Europejskiego Systemu Statystycznego), opracowano wskaźniki zrównoważonego rozwoju na poziomach krajowym, regionalnym oraz lokalnym, dla których inspirację stanowiły przedsięwzięcia prowadzone na arenie międzynarodowej (m.in. monitoring strategii zrównoważonego rozwoju UE). Ze względu na brak strategii zrównoważonego rozwoju w Polsce, jako podstawę stworzenia zestawu wskaźników, przyjęto założenia i cele w zakresie zrównoważonego rozwoju zapisane w wielu krajowych dokumentach o charakterze strategicznym (tj. strategię, plany czy programy). W efekcie powstał krajowy zestaw wskaźników liczący 101 mierników. Wskaźniki pogrupowano według czterech łańdów: społecznego, gospodarczego, środowiskowego oraz instytucjonalno-politycznego. Wyodrębniono 25 dziedzin, które odzwierciedlają zarówno cele, jak i priorytety zrównoważonego rozwoju. Wyniki prac zamieszczono w publikacji pt. *Wskaźniki Zrównoważonego Rozwoju Polski* (wydania z lat 2011 i 2015)¹², gdzie oprócz informacji metodologicznych zawarto również część analityczną. Następnie, ze względu na konieczność dostosowania zakresu informacyjnego do potrzeb jednostek różnych szczebli terytorialnych, stworzono moduły wskaźników na poziomach regionalnym i lokalnym o strukturze spójnej z modułem krajowym. Zestaw regionalny obejmuje ok. 73 wskaźniki, w którym najniższym poziomem agregacji są województwa. W zestawie lokalnym znalazło się natomiast ok. 56 wskaźników, gdzie najniższy dostępny przekrój stanowią powiaty.

Najnowszą pracą wpisującą się w działania monitorujące zrównoważony rozwój jest zestaw 92 wskaźników dotyczących „zielonej” gospodarki. Pogrupowano je według 5 tematów: kapitał naturalny, środowiskowa efektywność produkcji, środowiskowa jakość życia ludności, reakcje polityczne i możliwości gospodarcze oraz uwarunkowania społeczno-gospodarcze.

W celu popularyzacji, jak również w związku z potrzebą zebrania i zaprezentowania w jednym miejscu działań GUS i podejmowanych inicjatyw dotyczących zrównoważonego rozwoju zaprojektowano i udostępniono na portalu informacyjnym GUS specjalną zakładkę pn. *Zrównoważony Rozwój*. Odnaleźć tam można informacje o podejmowanych inicjatywach krajowych i międzynaro-

¹² Publikacja opracowana przez GUS i Urząd Statystyczny w Katowicach.

dowych oraz dane statystyczne dotyczące pomiaru tego rozwoju. Zakładka obejmuje takie elementy, jak:

- opis idei zrównoważonego rozwoju;
- proces wypracowywania Agendy Rozwoju 2030;
- monitorowanie zrównoważonego rozwoju Polski (Aplikacja WZR);
- opis koncepcji „zielonej” gospodarki;
- zestaw publikacji nt. zrównoważonego rozwoju;
- zbiór aktów prawnych i dokumentów strategicznych;
- listę najważniejszych inicjatyw międzynarodowych.

Częścią zasobów zakładki jest aplikacja *Wskaźniki Zrównoważonego Rozwoju*¹³, udostępniająca wskaźniki monitorujące zrównoważony rozwój w kraju, regionach, województwach i powiatach. Wskaźniki pogrupowano według dziedzin i czterech łańdów: społecznego, gospodarczego, środowiskowego oraz instytucjonalno-politycznego. Aplikacja oferuje obszerny zbiór metadanych oraz pozwala na dokonywanie porównań i przeprowadzanie analiz przy użyciu narzędzi do wizualizacji danych w formie wykresów i map. W przyszłości ma zostać rozbudowana o dodatkowe moduły pozwalające na umieszczenie wskaźników służących monitorowaniu krajowych celów zrównoważonego rozwoju wynikających z Agendy 2030 oraz stanu „zielonej” gospodarki.

Podsumowanie

Przed polską statystyką publiczną stoją ambitne wyzwania związane z planem rozwojowym, jakim jest Agenda na rzecz Zrównoważonego Rozwoju 2030. GUS, jako instytucja odpowiedzialna w kraju za koordynację monitorowania jej celów, ma sprostać potrzebom raportowania na poziomie globalnym i regionalnym. Na szczeblu krajowym natomiast będzie zobowiązany do przygotowania i utrzymania systemu monitorowania celów przyjętych dla Polski. Rola statystyki jest wspierająca, rozpoczęcie zatem aktywności w tym zakresie wymaga podjęcia decyzji i zdefiniowania priorytetów rozwojowych na szczeblu politycznym. Jest to kolejne przedsięwzięcie na rzecz planowania polityki rozwoju, które wymaga bliskiej współpracy decydentów ze statystykami. Ale doświadczenia polskiej statystyki pozwalają z optymizmem patrzeć na możliwość sprawnego wywiązania się z tych zadań.

Osiągnięcie celów zrównoważonego rozwoju — obecnie priorytetowego wyzwania dla świata — wymaga współdziałania szerokiego grona interesariuszy na wszystkich szczeblach zarządzania rozwojem. Zadania, wobec których stoją organizacje międzynarodowe, rządy krajowe i samorządy będą sprawdzianem zdolności do efektywnej współpracy i wykorzystania potencjału swoich partnerów, w tym statystyki publicznej. Polskie doświadczenia w zakresie instytucjonalnej współpracy oraz dotychczasowy dorobek statystyki dotyczący zrównowa-

¹³ Aplikacja dostępna bezpośrednio pod adresem <http://wskaznikizrp.stat.gov.pl/index.jsf>.

żonego rozwoju stanowią solidną podstawę do utworzenia kompleksowego planu działań ukierunkowanego na osiągnięcie celów Agendy 2030 w naszym kraju.

mgr Anna Prażmo, mgr inż. Joanna Wójcik, mgr Magdalena Żero — GUS

Summary. *The aim of the article is to present a new development Agenda for the world and the challenges for official statistics, related to monitoring progress in achieving its goals — at the global, regional and national level. Official statistics, from the beginning participating in the process of goals agreeing for which the world will endeavour in the next 15 years, was indicated as the authority responsible for coordinating and ensuring the continuity of monitoring their implementation. Meeting these expectations will require from international and national statistical institutions an increased efforts aimed at ensuring an efficient monitoring system and filling information gaps. These challenges also facing the Central Statistical Office of Poland — coordinator of the development goals monitoring at the national level.*

Keywords: sustainable development, 2030 Agenda, monitoring, sustainable development indicators, challenges of official statistics.

Резюме. *Целью статьи является представление новой Агенты принятой ООН в области сбалансированного развития мира, а также вызовов стоящих перед официальной статистикой, связанных с мониторингом реализации поставленных целей в глобальном, региональном и национальном подходе. Международная официальная статистика формирует цели, к которым стремится мир. Статистические управления отдельных стран обязаны наблюдать за реализацией этих целей. Удовлетворение этих ожиданий потребует у международных и национальных статистических учреждений организации эффективной системы мониторинга и заполнения пробелов в информации. Эти проблемы касаются также ЦСУ — координатора мониторинга реализации целей в области развития в Польше.*

Ключевые слова: сбалансированное развитие, Агенда 2030, мониторинг, показатели сбалансированного развития, вызовы перед официальной статистикой.

Olga LESZCZYŃSKA-LUBEREK

Zmiany w ujmowaniu zobowiązań emerytalno-rentowych w rachunkach narodowych

Streszczenie. *W artykule opisano zmiany rozwiązań w ujmowaniu nabytych uprawnień emerytalno-rentowych w ramach ubezpieczeń społecznych w rachunkach narodowych. Omówiono zasady obowiązujące do września 2014 r. oraz nowe wytyczne zawarte w Systemie Rachunków Narodowych (SNA 2008) i Europejskim Systemie Rachunków Narodowych i Regionalnych (ESA 2010). Wskazano różnice między systemami oraz przedstawiono rozwiązania europejskie. Podano argumenty, które przesądziły o nie włączeniu zobowiązań emerytalnych niekapitałowych publicznych systemów emerytalnych do statystyki opracowywanej na potrzeby procedury nadmiernego deficytu.*

Słowa kluczowe: rachunki narodowe, dług sektora instytucji rządowych i samorządowych, ubezpieczenia społeczne.

W wielu krajach, nawet najwyżej rozwiniętych, coraz bardziej dają o sobie znać trudności z finansowaniem potrzeb emerytalnych i ze zbilansowaniem funduszy tworzonych z wpływów ze składek emerytalnych z potrzebami związanymi z wypłatami emerytur (Barr i Diamond, 2014). Wśród głównych ich przyczyn wymienia się trendy demograficzne: wydłużanie się przeciętnego trwania życia, spadek umieralności i dzietności czy też spadek aktywności zawodowej mężczyzn w wieku emerytalnym.

W starzejących się społeczeństwach zwiększa się liczba osób w wieku poprodukcyjnym, a maleje liczba osób w wieku produkcyjnym, powodując wzrost współczynnika obciążenia demograficznego osobami starszymi. W 2013 r. wartość tego współczynnika wynosiła w przypadku miast 22 osoby oraz wsi 19 osób. W opublikowanej przez GUS w 2014 r. prognozie demograficznej na lata 2014—2050 czytamy, że wskaźniki starzenia (takie jak mediana czy udział ludności w wieku 65+ w populacji ogółem) pokazują, że deformacja struktur wyni-

kająca ze starzenia się ludności w większości krajów „starej” Unii Europejskiej (UE) jest bardziej zaawansowana niż w Polsce (GUS, 2014). Jednakże, zdaniem demografów, sytuacja zmieni się diametralnie już w ciągu najbliższej dekady, zaś w 2050 r. Polska stanie się w Europie jednym z krajów o najbardziej zaawansowanym w starzeniu się populacji. Prognozy mówią o potrojeniu liczby osób starszych przypadających na 100 osób w wieku 15—64 lata do 2050 r.

Wydatki emerytalne rządów rosną, rośnie bowiem liczba osób pobierających świadczenia, jak i ich wysokość. Skalę, perspektywę średniookresową czy też presję na budżet państwa można oceniać poprzez wysokość składki, która zapewniłaby zbilansowanie pieniężne systemu (w przypadku systemu o zdefiniowanym świadczeniu), analizę ogólnej kwoty wydatków na świadczenia emerytalno-rentowe czy też prezentowanie ich jako procentu PKB (Franco, 1995).

Państwa dysponują różnymi sposobami korekty sytuacji finansowej systemów emerytalnych. Do podstawowych należą: podniesienie składek¹, obniżenie świadczeń (wprowadzenie zmian w formułach służących do ich wyliczania), zastępowanie systemu o zdefiniowanym świadczeniu systemem o zdefiniowanej składce. Brakuje jednak metod do oceny długoterminowych skutków planowanych bądź wdrażanych reform, skali ich wpływu na wartość przyszłych pojedynczych świadczeń oraz skumulowanych zobowiązań. Nie istnieją dane porównywalne między krajami mówiące o tym, ile kosztowałoby zrealizowanie wszystkich nabytych w ramach rządowych systemów repartycyjnych² zobowiązań emerytalno-rentowych.

WYBRANE POJĘCIA ZWIĄZANE Z UBEZPIECZENIAMI SPOŁECZNYMI WYSTĘPUJĄCE W ESA 2010

Punktem wyjścia do omówienia problematyki ujmowania zobowiązań emerytalnych w systemie rachunków narodowych jest pojęcie systemu ubezpieczeń społecznych (*social insurance*), w ramach którego wyróżniamy zabezpieczenie społeczne (*social security*) oraz ubezpieczenia społeczne związane z zatrudnieniem inne niż zabezpieczenie społeczne (*employment related social insurance other than social security*).

Systemy ubezpieczeń społecznych definiowane są jako systemy, w których uczestnicy (osoby ubezpieczone) są zobowiązani lub zachęcani przez stronę trzecią do ubezpieczenia się od pewnego ryzyka społecznego lub okoliczności, które mogą mieć znaczący wpływ na ich sytuację lub sytuację osób na ich utrzymaniu. Nie ma tu znaczenia, kto (pracownik, inna osoba, pracodawca

¹ Zwiększenie wysokości składki może być sposobem na bilansowanie jedynie systemów o zdefiniowanym świadczeniu. W systemach o zdefiniowanej składce, wyższa składka wpływa na wyższą wartość emerytur do wypłaty.

² Systemami repartycyjnymi nazywane są niekapitałowe systemy ubezpieczeń społecznych. W systemach tych świadczenia wypłacane w danym okresie są finansowane, przede wszystkim, ze składek zapłaconych w tym samym okresie.

w imieniu pracownika) opłaca składkę na ubezpieczenie społeczne. Opłacanie składki zapewnia uprawnienie do świadczeń z ubezpieczeń społecznych w okresie bieżącym lub w przyszłości. Uprawnienia nabywa pracownik lub inna osoba opłacająca składki, osoby na ich utrzymaniu lub uprawnione do świadczeń z powodu śmierci bliskiego. Systemem ubezpieczeń społecznych mogą być też objęte osoby pracujące na własny rachunek lub niezatrudnione.

W rachunkach narodowych systemy zabezpieczenia społecznego (nazywane także funduszami zabezpieczenia społecznego) definiowane są jako obejmujące całe społeczeństwo lub dużą jego część, kontrolowane i finansowane³ przez jednostki sektora instytucji rządowych i samorządowych.

Systemy związane z zatrudnieniem wynikają z relacji pomiędzy osobą zatrudnioną i pracodawcą. Odpowiedzialność za wypłatę świadczeń ponosi pracodawca lub instytucja zarządzająca systemem w jego imieniu.

Najważniejszym świadczeniem, do którego nabywa się uprawnienie w ramach ubezpieczenia społecznego jest świadczenie emerytalno-rentowe. Pod tym pojęciem należy rozumieć nie tylko emerytury i renty wypłacane pracownikom lub osobom, które zakończyły karierę zawodową, ale także świadczenia wypłacane wdowom i wdowcom, sierotom czy tym, którzy utracili zdolność do pracy.

W ESA 2010 podzielono systemy emerytalno-rentowe na systemy o zdefiniowanej składce i systemy o zdefiniowanym świadczeniu⁴. W przypadku pierwszym, świadczenie jest uzależnione od wartości składek wpłaconych do systemu, dochodów z ich inwestowania oraz zysków lub strat z tytułu posiadania aktywów. W drugim — świadczenia wyliczane są według określonej formuły. Te dwa rodzaje systemów odróżnia także kwestia ponoszenia ryzyka związanego z inwestowaniem składek. W systemie o zdefiniowanej składce pozostaje ono po stronie pracownika, z kolei w systemie o zdefiniowanym świadczeniu — ponosi je pracodawca lub jednostka zarządzająca systemem w jego imieniu.

Dodatkowo w ESA 2010 wyróżniono systemy hybrydowe, łączące elementy systemu o zdefiniowanej składce i o zdefiniowanym świadczeniu oraz umowne systemy o zdefiniowanej składce, w których informacja o wysokości wpłaconych składek jest zapisywana na indywidualnych umownych rachunkach. Ich umowny charakter wynika z tego, że składki wpłacane do systemu nie są w rzeczywistości gromadzone tylko wykorzystywane do wypłaty bieżących świadczeń. W przypadku tych systemów ryzyko zapewnienia odpowiedniego świadczenia emerytalnego jest dzielone między pracownika i pracodawcę lub jednostkę zarządzającą systemem w jego imieniu. Aby system emerytalny został uznany w rachunkach narodowych za hybrydowy, musi zawierać gwarancję minimalnej kwoty świadczenia.

³ Finansowanie należy rozumieć jako ponoszenie odpowiedzialności za wypłatę świadczeń.

⁴ Rozporządzenie Parlamentu Europejskiego i Rady (UE) nr 549/2013 z 21 maja 2013 r. w sprawie Europejskiego Systemu Rachunków Narodowych i Regionalnych w Unii Europejskiej (ESA 2010), Dz. U. L 169 z 21.06.2013 r.

W zestawieniu przedstawiono podsumowanie zawartych w ESA 2010 wytycznych dotyczących typów, charakterystyki, formy organizacji oraz klasyfikacji do właściwego sektora instytucjonalnego systemów ubezpieczeń społecznych.

**ZESTAWIENIE INFORMACJI DOTYCZĄCYCH
SYSTEMU UBEZPIECZEŃ SPOŁECZNYCH ESA 2010**

Typ ubezpieczenia	Charakterystyka	Forma organizacji	Sektor/Podsektor
Zabezpieczenie społeczne	sektor instytucji rządowych i samorządowych zobowiązuje uczestników systemu do ubezpieczenia od określonych rodzajów ryzyka	organizowane przez sektor instytucji rządowych i samorządowych za pośrednictwem funduszy zabezpieczeń społecznych	sektor instytucji rządowych i samorządowych — fundusze zabezpieczenia społecznego
Systemy ubezpieczeń społecznych związane z zatrudnieniem, inne niż zabezpieczenie społeczne	pracodawcy mogą uzależnić zatrudnienie od tego, czy pracownik ubezpieczy się od określonych rodzajów ryzyka	organizowane przez pracodawców w imieniu ich pracowników oraz osób na ich utrzymaniu lub przez inne podmioty w imieniu określonej grupy	sektor lub podsektor pracodawcy, instytucji ubezpieczeniowych, funduszy emerytalno-rentowych lub instytucji niekomercyjnych działających na rzecz gospodarstw domowych

Źródło: opracowanie na podstawie ESA 2010.

Przedstawiona typologia jest kluczowa dla rozważań dotyczących zobowiązań emerytalno-rentowych w systemie rachunków narodowych. Uprawnienia powstające w ramach zabezpieczenia społecznego nie są wykazywane w zasadniczych rachunkach narodowych⁵, a jedynie w tablicy uzupełniającej (omówiona w dalszej części artykułu). W przypadku pozostałych systemów ubezpieczeń społecznych związanych z zatrudnieniem uprawnienia emerytalno-rentowe są rejestrowane w zasadniczych rachunkach narodowych w miarę ich powstawania.

POJĘCIE ZOBOWIĄZAŃ EMERYTALNO-RENTOWYCH

W literaturze przedmiotu (Europejski Bank Centralny i Eurostat, 2011) wprowadzono trzy podstawowe koncepcje zobowiązań emerytalno-rentowych:

- zobowiązania narosłe do roku bazowego (*accrued-to-date liabilities* — ADL),
- zobowiązania osób pracujących i pobierających świadczenia (*current workers' and pensioners' liabilities* — CWL),
- narosłe zobowiązania emerytalne dla systemu otwartego (*open-system gross liabilities* — OSL).

⁵ Zasadnicze rachunki narodowe to rachunki, dla których opracowywane są standardowe tablice, a zawarte w nich informacje mają potencjalny wpływ na kluczowe agregaty, takie jak PKB. Rachunki satelitarne i tablice uzupełniające przenoszą zasady rachunków narodowych na określone dziedziny (np. środowisko, zdrowie, edukacja) i nie mają wpływu na kluczowe agregaty.

Uprawnienia i zobowiązania w podejściu ADL obejmują teraźniejszą wartość świadczeń, które mają być zapłacone w przyszłości, a wynikają z już zapłaconych składek przez osoby pracujące oraz pozostałych do zrealizowania uprawnień osób pobierających świadczenia. To ujęcie pokazuje zatem zobowiązania w ograniczonej perspektywie czasowej. Nie uwzględnia uprawnień nabytych po okresie badanym zarówno przez osoby już pracujące, jak i te, które dopiero podejmą zatrudnienie. Podejście takie jest uznawane za istotne dla oceny wpływu efektów fiskalnych na decyzje dotyczące oszczędności i konsumpcji gospodarstw domowych. Informuje również o koszcie likwidacji systemu przy jednoczesnej wypłacie świadczeń, do których jego uczestnicy nabyli uprawnienia do momentu likwidacji systemu i zgodnie z obowiązującymi zasadami.

Gdyby w danym momencie podjęto decyzję o zaprzestaniu pobierania składek, gdyż byłyby wpłacane do innego systemu, ADL powie nam, jakie środki finansowe trzeba będzie uzyskać z innych źródeł (np. podatkowych), aby zrealizować uprawnienia nabyte przez uczestników systemu. W opracowaniach dotyczących tematyki ubezpieczeń społecznych powtarza się stwierdzenie, że ujęcie ADL nie może być wykorzystywane do oceny kondycji systemu emerytalnego lub szerzej do kondycji fiskalnej państwa. Wszystko, co można powiedzieć, to tyle, że im wyższy jest wskaźnik odnoszący wartość uprawnień emerytalnych ADL do PKB, tym wyższy udział przyszłych zasobów publicznych, które będą musiały być przeznaczone na wypłaty emerytur i tym wyższe ryzyko, że jeśli wzrost gospodarczy nie będzie odpowiedni, konieczne będą korekty (np. podniesienie składek, obniżenie uprawnień).

Szacując ADL można przyjąć jedno z dwóch zobowiązań z tytułu przewidywanych świadczeń (*projected benefit obligations approach* — PBO) lub z tytułu nabytych uprawnień do świadczeń (*accumulated benefit obligations approach* — ABO). Różnica między nimi odnosi się do odpowiedzi na pytanie o to, jak (z uwzględnieniem wzrostu wynagrodzeń w przyszłości — PBO czy zakładając stały ich poziom — ABO) oszacować uprawnienia osób, które jeszcze nie przeszły na emeryturę. Może ona wynikać z ogólnego wzrostu wynagrodzeń, który w większości przypadków jest związany ze wzrostem gospodarczym lub ze wzrostem wynagrodzenia podczas kariery zawodowej. Oznacza to jednocześnie, że uprawnienia osób, które otrzymują świadczenia w roku bazowym — nabyły już pełne uprawnienia emerytalne — nie różnią się w podejściu ABO i PBO (Heidler, Müller i Weddige, 2009).

W podejściu CWL dokonuje się doszacowania tak, by pokazać system emerytalno-rentowy do momentu, w którym umiera ostatnia z osób już płacących składki. Nie uwzględnia się natomiast nowych członków. Podejście to stanowi sumę zobowiązań oszacowanych w podejściu ADL i szacunku zobowiązań emerytalno-rentowych, które powstaną w przyszłości wobec osób już płacących składki.

Ujęcie OSL, oprócz zobowiązań/uprawnień naliczonych w podejściu CWL, obejmuje szacunek teraźniejszej wartości świadczeń nowych pracowników przystępujących do systemu. Przyjmuje się tu założenie, że system emerytalno-

-rentowy będzie funkcjonował na niezmiennych zasadach w relatywnie długim okresie. Bieżąca wartość OSL może być szacowana dla nieograniczonego horyzontu czasowego. Nie dostarcza ona informacji o ewentualnym niezbilansowaniu systemu, jego presji na budżet, a zatem nie może stanowić podstawy do wnioskowania o potrzebie zmian. OSL uznawane jest jednakże za właściwe do oceny długoterminowej stabilności finansów publicznych oraz systemu emerytalnego. Po uwzględnieniu przyszłych dochodów (mówimy wtedy o narosłych zobowiązaniach emerytalnych netto w systemie otwartym, *open-system net liabilities* — OSNL) pozwala na skonfrontowanie z nimi uprawnień poprzez pokazanie dłuższej perspektywy obejmującej również uprawnienia, które zostaną nabyte w przyszłości.

Problemy dotyczące metod szacowania zobowiązań systemów emerytalno-rentowych i ich zastosowań zarówno w odniesieniu do systemów kapitałowych, jak i niekapitałowych analizowane były przez księgowych i aktuariuszy, jak również ekspertów pracujących nad nowymi (światowymi i europejskimi) systemami rachunków narodowych (SNA 2008/ESA 2010). W ramach tych prac opracowano koncepcję rachunków pokoleniowych. Pojawiło się także podejście mające na celu opracowanie syntetycznych wskaźników fiskalnych służących do oceny kwot, jakie są niezbędne w celu zbilansowania bieżącej wartości długoterminowych ograniczeń budżetowych.

Inne podejście mówiło o rozszerzeniu rachunków sektora instytucji rządowych i samorządowych o aktywa i pasywa tak, by podawać jego wartość netto, przy czym pasywa powinny obejmować również zobowiązania uznawane za warunkowe, a zatem także emerytalno-rentowe. Jednakże, zgodnie z zasadami SNA 1993 oraz obowiązującego do września 2015 r. ESA 95, w statystyce makroekonomicznej podawano jedynie zobowiązania systemów kapitałowych, ponieważ zobowiązania systemów niekapitałowych nie były ujmowane w rachunkach narodowych. Wymóg opracowywania szacunków w tym zakresie nie był ujęty w żadnym akcie prawnym dotyczącym statystyki makroekonomicznej.

ZOBOWIĄZANIA EMERYTALNO-RENTOWE W SNA 1993 i ESA 95

W latach 80. XX w. ekonomiści zajmujący się problematyką finansów publicznych wskazywali na wady konwencjonalnego pomiaru deficytu i długu sektora instytucji rządowych i samorządowych oraz jego nieprzydatność do oceny równowagi fiskalnej. Uważali oni, że informacja o wielkości i zmianach długu nie dostarcza kompleksowego obrazu finansów państw. Pierwsze szacunki zobowiązań repartycyjnych rządowych systemów emerytalnych pokazały, że w większości krajów „dług emerytalny”, nazywany „ukrytym”, jest znacznie większy niż tradycyjnie definiowany dług publiczny. To utwierdziło zwolenników poszerzenia definicji długu publicznego o zobowiązania rządowych systemów emerytalno-rentowych w przekonaniu, że dostępna statystyka jest niedoskonała, gdyż nie obrazuje przyszłych efektów bieżącej polityki rządów. Jedno-

częściej podkreślali, że stosowane metody i dostępne informacje mogą nie być wystarczające dla właściwego monitorowania i kontroli sytuacji fiskalnej. W tym kontekście statystycy zaczęli formułować propozycje dotyczące włączenia emerytur i rent do rachunków sektora instytucji rządowych i samorządowych w momencie, kiedy narastają zobowiązania, a nie wtedy, kiedy ponoszone są wydatki.

Równocześnie postawiono szereg pytań, bo pojęcie ukrytego długu emerytalnego, choć powszechnie używane, nie zostało jednoznacznie zdefiniowane. Dwie propozycje były szczególnie popularne. Według pierwszej, ukryty dług emerytalny miała stanowić różnica między zdyskontowaną wartością wydatków i dochodów określoną dla wybranego okresu. Według drugiej — wartość nabytych zobowiązań nieuwzględniająca przyszłych dochodów i zgromadzonych aktywów. Drugie z podejść nazywane jest ujęciem brutto, gdyż pomija się w nim przyszłe dochody oraz zgromadzone aktywa. W obu przypadkach istotny wpływ na wyniki ma przyjęta do szacunków perspektywa czasowa oraz założenia dotyczące np. przewidywanego trwania życia.

Eksperci zajmujący się statystyką sektora instytucji rządowych i samorządowych podjęli również kwestię ujęcia w rachunkach narodowych jednorazowych dużych transferów zobowiązań. Występują one np. w kontekście planowanej prywatyzacji — rząd przejmuje zobowiązania emerytalne przedsiębiorstwa wobec jego pracowników w zamian otrzymując jednorazowy transfer aktywów. Podejście obowiązujące w SNA 93 i ESA 95 odzwierciedlało tylko jedną stronę takiej transakcji — transfer aktywów. Powodowało to jednorazowe poprawienie deficytu/nadwyżki sektora instytucji rządowych i samorządowych, podczas gdy nie rejestrowano w żaden sposób, w rachunku finansowym lub niefinansowym, przejścia zobowiązań.

Jak wspomniano, wycena zobowiązań emerytalnych dotyczy także systemów prywatnych. Korzystne trendy na rynkach finansowych w Stanach Zjednoczonych czy Wielkiej Brytanii przynosiły nadwyżki systemom emerytalnym, składając korporacje do podnoszenia wysokości świadczeń (Semeraro, 2007). Pogorszenie się jednak sytuacji oraz brak właściwego oszacowania zobowiązań emerytalnych negatywnie wpłynęły na wypłacalność korporacji, gdyż konieczne stało się finansowanie deficytowych programów emerytalnych. Zjawisko było szczególnie istotne w tych krajach, gdyż pracownicze systemy emerytalne o zdefiniowanym świadczeniu były tam powszechne.

Wdrożone na przełomie XX w. oraz XXI w. międzynarodowe standardy rachunkowości FRS17 i IAS19 zakładały ujednolicone metody określania zobowiązań (sposób wyceny i ustalanie ich w miarę nabywania, a nie w momencie wypłaty) systemów o zdefiniowanym świadczeniu i o zdefiniowanej składce. Wprowadzeniu wytycznych towarzyszyło przekonanie, że umożliwią one pracodawcom i inwestorom dokonywanie bardziej realistycznych ocen, które będą mniej zależne od okresowej poprawy przepływów finansowych. Wdrożenie wspomnianych standardów było bardzo istotne, gdyż kwestia zharmonizowania tak dalece, jak to tylko możliwe, zasad rachunków narodowych i standardów

rachunkowości biznesowej jest stałym elementem dyskusji w pracach nad SNA/ESA.

Eksperti pracujący nad nowymi wytycznymi w zakresie rachunków narodowych rozpatrywali przydatność zasad rachunkowości dotyczących zobowiązań emerytalnych systemów zakładowych w sytuacji, w której sektor instytucji rządowych i samorządowych zobowiązany jest do wypłaty świadczeń. Za prawdopodobne uznano, że deficyt sektora instytucji rządowych jako pracodawcy lub gwaranta systemu ubezpieczeń społecznych może być niedoszacowany. Analogia występująca z FRS i IAS oraz błędy w oszacowaniach zobowiązań pracodawców stały się jednym z kluczowych argumentów w dyskusji dotyczącej zmiany ujęcia emerytur w rachunkach narodowych (Semeraro, 2011). Ponadto zwracano uwagę, iż odmienne ujęcie w statystyce systemów kapitałowych i niekapitałowych daje różny wpływ na takie zmienne, jak dochody i oszczędności gospodarstw domowych czy udziały netto gospodarstw domowych w rezerwach funduszy emerytalnych i ubezpieczeń na życie. Ze względu na różnorodność systemów emerytalnych, różną skalę programów repartycyjnych i kapitałowych wartości te nie mogły być porównywane między krajami. Jednocześnie, wobec zmian demograficznych i zwiększającego się zakresu systemów emerytalnych, pojawiło się zainteresowanie dostępem do szerszego spektrum informacji statystycznych dotyczących przyszłych zobowiązań rządów oraz wpływu, jaki mają na ich wysokość wdrożone lub planowane reformy systemów emerytalnych.

Zgodnie z SNA 93/ESA 95 zobowiązania emerytalne były przedstawione jako element aktywów gospodarstw domowych w pozycji udziały netto gospodarstw domowych w rezerwach funduszy emerytalnych (składowa technicznych rezerw ubezpieczeniowych) oraz pasywów funduszy emerytalnych. W paragrafie 7.59 ESA 95 czytamy, że natura zobowiązań funduszu (i aktywów gospodarstw domowych) zależy od rodzaju systemu emerytalnego. Dla systemu o zdefiniowanym świadczeniu równa się ona bieżącej wartości obiecanych (i gwarantowanych) świadczeń. Dla systemów kapitałowych jest to bieżąca wartość rynkowa aktywów funduszu. Dodatkowo w rozdziale 5 ESA 95 znajduje się stwierdzenie, że techniczne rezerwy ubezpieczeniowe stanowią zobowiązania autonomicznych funduszy emerytalnych sklasyfikowanych w podsektorze instytucji ubezpieczeniowych i funduszy emerytalnych oraz nieautonomicznych funduszy emerytalnych sklasyfikowanych w tym sektorze, do którego zalicza się podmiot założycielski. Równocześnie zgodnie z ESA 95 ujmowano w rachunkach narodowych rezerwy tworzone przez pracodawców w celu wypłaty świadczeń pracownikom tylko, jeśli były wyliczane tak, jak ma to miejsce w przypadku zakładów ubezpieczeń lub autonomicznych funduszy emerytalnych. Zobowiązaniom ujętym w rachunku finansowym i w rachunku niefinansowym odpowiadała pozycja udziały netto gospodarstw domowych w rezerwach funduszy emerytalnych. W paragrafie 5.102 ESA 95 czytamy, że nie obejmuje się rezerw ustanowionych przez jednostki sklasyfikowane w podsektorze ubezpieczeń społecznych, gdyż system nie uznaje ich za zobowiązania podsektora.

ZOBOWIĄZANIA EMERYTALNO-RENTOWE W SNA 2008 i ESA 2010

Pierwsze, wypracowane na forum międzynarodowym, propozycje dotyczące nowego SNA mówiły o ujmowaniu w zasadniczych rachunkach narodowych zobowiązań systemów pracowniczych bez względu na to, czy są kapitałowe czy niekapitałowe. Do ich wyceny proponowano wykorzystanie metod aktuarialnych oraz zaleceń zawartych w IAS19. Jednocześnie propozycja nie mówiła nic o ewentualnej zmianie ujęcia zobowiązań rządowych systemów repartycyjnych. Równolegle trwały dyskusje na forum europejskim. Od początku zakładano spójność z nowym SNA, ale szczególną uwagę zwracano na konsekwencje ewentualnych zmian dla statystyki wykorzystywanej w procedurze nadmiernego deficytu (EDP). Podczas gdy system światowy mógł pozostawiać pewną dowolność w prezentacji zobowiązań nie było to możliwe w przypadku wytycznych europejskich. Zobowiązanie wprowadzane rozporządzeniem Parlamentu Europejskiego i Rady UE stanowi bowiem wiążącą wykładnię dla opracowywania danych, które wykorzystywane są w celach administracyjnych i politycznych w Unii.

Efektom prac nad SNA 2008 był kompromis pozwalający na częściowe uwzględnienie pewnych zobowiązań w zasadniczych rachunkach, a innych w pozycjach uzupełniających. W paragrafach 17.191—17.193 SNA 2008 zawarto zapisy mówiące o tym, że ponieważ ubezpieczenia społeczne mają charakter repartycyjny, to zobowiązania nabyte w ich ramach nie są zazwyczaj podawane w systemie (*are not normally shown in the SNA*). Uzasadnieniem jest to, że wiarygodne szacunki ich wartości nie są dostępne w przeciwieństwie do systemów prywatnych oraz dodaje się, że nawet gdyby istniały takie obliczenia, to ich użyteczność należy uznać za ograniczoną, gdyż rząd ma możliwość zmiany zasad funkcjonowania systemu.

W SNA czytamy także, że konsekwencją akceptacji ujmowania w rachunkach narodowych jednych uprawnień, a innych nie, będzie to, że niektóre kraje (mające przewagę systemów prywatnych) będą podawać w statystyce większą część zobowiązań, podczas gdy inne (z przewagą systemów rządowych) nie będą ich podawać. Mając na względzie różnice między krajami w zakresie organizacji systemu ubezpieczeń społecznych SNA pozostawia statystykom swobodę rejestracji zobowiązań niekapitałowych systemów emerytalnych sponsorowanych przez rząd. Część z nich może być ujęta w zasadniczych rachunkach, podczas gdy przybliżone szacunki pozostałych muszą być ujmowane w dodatkowej tablicy uzupełniającej. Jej wzór znajdziemy w rozdziale 17 SNA 2008. W rozdziale tym zawarto także kryteria brane pod uwagę przy podejmowaniu decyzji, które z zobowiązań należy prezentować w zasadniczych rachunkach, a które tylko w tablicy uzupełniającej. Prace metodyczne na ten temat będą kontynuowane⁶.

⁶ SNA może zawierać zapisy pozwalające na nie do końca jednolite opracowywanie danych przez kraje stosujące standard, gdyż dane nie są wykorzystywane dla celów administracyjnych i politycznych tak, jak ma to miejsce w UE.

OECD proponowała równe traktowanie niekapitałowych systemów pracowniczych i systemów ubezpieczeń społecznych oraz włączenie wszystkich zobowiązań emerytalnych do systemu rachunków narodowych, obok zasadniczych standardowych rachunków. Proponowano równoległe opracowywanie rachunków — z uwzględnieniem lub nie zobowiązań i transakcji związanych z niekapitałowymi systemami emerytalnymi.

Ostateczną propozycję wytycznych dla krajów UE opracowano w 2005 r. Opierała się ona na przedstawieniu wszystkich zobowiązań emerytalnych (w tym w ramach ubezpieczeń społecznych) w dodatkowej tabeli uzupełniającej, nieujętej w zasadniczych rachunkach narodowych. Przy czym uprawnienia nabyte w ramach systemu ubezpieczeń społecznych oraz systemów o zdefiniowanym świadczeniu dla pracowników sektora instytucji rządowych i samorządowych ESA 2010 nakazuje ujmować jedynie w tablicy uzupełniającej, a nie w zasadniczych rachunkach narodowych. Nadal w zasadniczych rachunkach narodowych podawane będą transakcje takie, jak płatności składek i wypłaty świadczeń. Decyzja była odzwierciedleniem występującego w wielu krajach UE podobieństwa pomiędzy takimi rodzajami systemów. Pozwoliła ona na zapewnienie porównywalności danych pomiędzy krajami UE.

Tablicę z nabytymi uprawnieniami emerytalno-rentowymi w ramach ubezpieczeń społecznych włączono do rozporządzenia wdrażającego ESA 2010. Jest ona elementem załącznika B do rozporządzenia określającego program transmisji danych ESA 2010. Zagadnieniu temu poświęcono osobny rozdział zatytułowany *Ubezpieczenia społeczne wraz z emeryturami i rentami*.

Pierwsze obowiązkowe przekazanie danych nastąpi w grudniu 2017 r. i obejmować będzie szacunki dla roku 2015. Tablica uzupełniająca rejestruje stany i przepływy zobowiązań emerytalno-rentowych w odniesieniu do wszystkich systemów ubezpieczeń społecznych. Pokazuje także drugą stronę, czyli uprawnienia emerytalno-rentowe gospodarstw domowych odpowiadające tym zobowiązaniom. Informacje przedstawione poniżej w tablicy mają obrazować perspektywę dłużnika (systemu emerytalno-rentowego) — pokazują zobowiązania emerytalno-rentowe, jak i wierzyciela (gospodarstwa domowego) oraz przedstawiają uprawnienia. Stanom i przepływowi zobowiązań zawsze odpowiadają stany i przepływy uprawnień emerytalno-rentowych.

Tablicę przygotowano tak, by odczytać z niej nie tylko informację o wartości nabytych uprawnień emerytalno-rentowych, ale także o transakcjach i pozostałych przepływach zaobserwowanych w danym okresie powodujących ich zmianę. Dowiemy się z niej, które ze zobowiązań zostały ujęte w zasadniczych rachunkach narodowych, a które tylko w tablicy uzupełniającej. Dane będą przedstawiane osobno dla systemów emerytalno-rentowych zarządzanych przez sektor instytucji rządowych i samorządowych oraz przez pozostałe sektory instytucjonalne. Dowiemy się także, jaka jest wartość zobowiązań nabytych w systemach o zdefiniowanej składce, a jaka w systemach o zdefiniowanym świadczeniu.

NABYTE UPRAWNIENIA EMERYTALNO-RENTOWE W RAMACH UBEZPIECZEŃ SPOŁECZNYCH (dok.)

Kod transakcji według ESA 2010	Rejestrowanie	Zasadnicze rachunki narodowe				Poza zasadniczymi rachunkami narodowymi				Systemy emerytalno-rentowe w ramach systemu ubezpieczeń społecznych	Uprawnienia emerytalno-rentowe gospodarstw domowych niezysdentów		
		sektory inne niż sektor instytucji rządowych i samorządowych		sektor instytucji rządowych i samorządowych		sektor instytucji rządowych i samorządowych		systemy emerytalno-rentowe w ramach systemu ubezpieczeń społecznych					
		systemy o zdefiniowanej składce	systemy o zdefiniowanym świadczeniu oraz pozostałe systemy o niezdefiniowanej składce	ogółem	systemy o zdefiniowanej składce	zaliczane do instytucji finansowych	zaliczane do sektora instytucji rządowych i samorządowych	zaliczane do sektora instytucji rządowych i samorządowych	systemy emerytalno-rentowe w ramach systemu ubezpieczeń społecznych				
		XPCIW	XPBIW	XPCBIW	XPCG	XPBG12	XPBG13	XPBOUT13	XPI314	XPTOT	XPTOTNRH		
		numer kolumny										I	J
Zmiany w uprawnieniach emerytalno-rentowych z tytułu transakcji (dok.)													
XD81	6	0	0	0	0	0	0	0	0	0	0		
XD82	7	0	0	0	0	0	0	0	0	0	0		
Zmiany w uprawnieniach emerytalno-rentowych z tytułu pozostałych przepływów													
XK7	8	0	0	0	0	0	0	0	0	0	0		
XK5	9	0	0	0	0	0	0	0	0	0	0		
Bilans zamknięcia													
XAF63LE	10	0	0	0	0	0	0	0	0	0	0		
Powiązane wskaźniki													
XP1	11	0	0	0	0	0	0	0	0	0	0		

Źródło: opracowanie własne na podstawie ESA 2010.

WYTYCZNE DO SZACUNKÓW ZOBOWIĄZAŃ EMERYTALNO-RENTOWYCH W ESA 2010

Pojawiły się w dyskusjach problemy dotyczące niepewności co do danych wykorzystywanych do wyliczeń. Zwracano uwagę na zależność i bardzo dużą wrażliwość wyników na przyjętą stopę dyskontową. W przypadku wybrania stałej wartości mówi się o arbitralności wyboru, w przypadku wykorzystywania zmiennej stopy — o wahaniach wyników. Kolejnym, trudnym zagadnieniem była kwestia dotycząca opracowania prognoz demograficznych, jak również projekcji przyszłych wynagrodzeń. Statystycy byli jednak niejednomysłni co do sposobu ujęcia w nowej metodzie składek na ubezpieczenia społeczne. Kluczowe pytanie dotyczyło tego, czy w szacunkach zobowiązań do wypłaty świadczeń w przyszłości powinno ujmować się (ujęcie netto) czy pomijać (ujęcie brutto) przyszłe składki. Zwolennicy podejścia netto wskazywali, że kwestię wypłaty świadczeń i płatności składek regulują zazwyczaj te same akty prawne. Ich oponenci argumentowali, że składki płacone w danym okresie są podstawą nabycia przez osoby je płacące uprawnień do świadczeń w przyszłości, więc nie powinny być odejmowane od zobowiązań nabytych do okresu bieżącego.

Szczegółowe zalecenia, które zapewnią porównywalność obliczeń pomiędzy krajami UE, zawarto w wydanym wspólnie przez Eurostat i Europejski Bank Centralny (EBC) w 2011 r. podręczniku *Technical Compilation Guide for Pension Data in National Accounts*⁷. Ustalono, że dane do tablicy należy opracować w ujęciu ADL z zastosowaniem podejścia PBO. Inne najważniejsze wytyczne można podsumować w następujących punktach:

- realna stopa dyskontowa wynosząca na stałym poziomie — 3%, nominalna — 5% (stopa realna plus stała inflacja — 2% zgodna ze średniookresowym celem inflacyjnym EBC);
- indeksacja świadczeń zgodna z zasadami obowiązującymi w poszczególnych krajach (np. indeksacja wskaźnikiem zmiany cen konsumpcyjnych, wzrostem wynagrodzeń lub inflacją w połączeniu ze wzrostem wynagrodzeń);
- zgodność założeń dotyczących realnego wzrostu wynagrodzeń z długookresowymi założeniami Grupy roboczej ds. starzenia się społeczeństw (*Ageing Working Group* — AWG) dotyczącymi produktywności krajów UE;
- nominalny wzrost wynagrodzeń równy wzrostowi realnemu powiększonemu o inflację;
- spójność założeń demograficznych (ludność, urodzenia, zgony, migracje) z opracowywanymi przez Eurostat prognozami Europop dla krajów UE.

⁷ Podręcznik dostępny jest na stronie internetowej Eurostatu, <http://ec.europa.eu/eurostat/ca/web/products-manuals-and-guidelines/-/KS-RA-11-027>.

*PIERWSZE UPROSZCZONE SZACUNKI ZOBOWIĄZAŃ
EMERYTALNO-RENTOWYCH*

Pierwsze, uproszczone szacunki oparte na europejskich wytycznych opracowano na zlecenie EBC w 2009 r. na Uniwersytecie we Freiburgu. Obliczenia dotyczyły 2006 r. i obejmowały one 19 krajów UE. Zdaniem ich autorów, nowe zasady rejestrowania zobowiązań emerytalno-rentowych otwierają obszerne pole dyskusji dla polityków i naukowców, w szczególności zajmujących się kwestią zasobności gospodarstw domowych. Będzie to jednakże możliwe tylko wtedy, jeśli zagwarantowane zostanie spójne podejście do obliczeń między krajami.

Przeprowadzone analizy wrażliwości wyników na przyjęte założenia wykazały znaczący wpływ na obliczenia zmian w stopie dyskontowej oraz wzroście wynagrodzeń. Z tej przyczyny tak ważne, zdaniem autorów, dla porównań międzynarodowych jest przyjęcie zharmonizowanych założeń.

Spośród analizowanych czynników najmniej istotny wpływ na wyniki miała umieralność, przeciwnie było w przypadku zasad indeksacji świadczeń. Zmiana z indeksacji wyłącznie wskaźnikiem wzrostu cen na indeksację łączoną — po połowie indeks cen i wzrost wynagrodzeń — powoduje znaczący wzrost wartości zobowiązań. Uprawnienia są tym większe, im większa skala indeksacji wzrostem wynagrodzeń (Kaier i Müller, 2013). Prowadzący szacunki Kaier i Müller stwierdzili zależność między wysokością wydatków emerytalnych, wyrażonych jako % PKB, a wysokością zobowiązań emerytalnych — im wyższe wydatki, tym wyższe zobowiązania.

Na wysokość wydatków wpływa m.in. liczba osób pobierających świadczenia, średni wiek przejścia na emeryturę, hojność systemu emerytalnego wyrażona stopą zastąpienia. Z przeprowadzonych analiz wynika, że zmiana przyjętych założeń ze średnich do optymistycznych albo pesymistycznych prowadzi do zmiany wyników na skalę większą od długu krajowego pozostającego w rękach prywatnych. Jednakże zdaniem przedstawicieli Uniwersytetu we Freiburgu, to zmiany demograficzne stanowią główne, przyszłe ryzyko publicznych systemów emerytalnych w większości krajów rozwiniętych (Heidler i in., 2009).

Chociaż opracowania międzynarodowe są użyteczne w promowaniu zainteresowania tematyką zobowiązań emerytalnych, to szacunków nie można uznać za satysfakcjonujące. Wyniki mogą znacząco różnić się od bardziej dokładnych opracowań krajowych. Zakres informacji wymaganych do dobrego oszacowania zobowiązań jest tak duży, że nie można tego zrobić bez wsparcia ekspertów krajowych. To stwierdzenie wydaje się potwierdzać pierwsze wstępne opracowania przygotowane i udostępnione przez statystyków brytyjskich i holenderskich. W obu wypadkach wskazuje się na konieczność przyjęcia uproszczeń (np. wyłączenie z obliczeń mniej istotnych elementów systemu ubezpieczeń społecznych, przyjęcie bez wprowadzania dostosowań do europejskich wytycznych szacunków opracowanych przez podmioty prywatne prowadzące pracownicze programy emerytalne) i trudności w dostępie do szczegółowych danych źródłowych.

ZOBOWIĄZANIA EMERYTALNO-RENTOWE I DŁUG ORAZ DEFICYT SEKTORA INSTYTUCJI RZĄDOWYCH I SAMORZĄDOWYCH

W toku prac nad SNA 2008/ESA 2010 oraz w dyskusjach nad sposobem ujęcia zobowiązań emerytalnych w systemie podnoszono także kwestię zmiany definicji deficytu i długu sektora instytucji rządowych i samorządowych. Choć nazywane długiem ukrytym zobowiązania emerytalne różnią się od ujmowanych w długu konwencjonalnym. Jako podstawowe różnice wskazuje się:

- brak ustalonych z góry terminów i wysokości wypłat, które zależą od czynników takich, jak wiek przejścia na emeryturę, długość życia, zmiany elementów formuły stosowanej do wyliczania świadczenia;
- możliwość jednostronnej modyfikacji przez rząd zarówno czasu, jak i wysokości płatności bez poniesienia konsekwencji prawnych;
- brak swobody uczestnictwa w systemie, a zatem i nabycia uprawnień emerytalnych.

Najważniejszą różnicą pozostaje fakt, że uprawnienia emerytalne mogą być ograniczone poprzez zmianę formuły określania świadczeń. W wielu krajach dokonano zmian indeksacji, wieku emerytalnego, zasad nabywania uprawnień i innych elementów systemów. Konsekwentnie błędem byłoby uznanie, że programy emerytalne stanowią wiążące prawnie lub moralnie zobowiązanie. Ujęcie zobowiązań emerytalnych w definicji długu publicznego wymagałoby analogicznej zmiany w definicji deficytu. Składki stałyby się pożyczkami dla sektora publicznego, emerytury w części byłyby uznawane za spłatę pożyczki, a w części jako zapłacone odsetki.

Fakt, że do opracowania szacunków konieczne jest przyjęcie szeregu założeń może powodować konieczność ich ciągłej rewizji, zapewne na dużą skalę. W efekcie zmienne fiskalne, takie jak dług i deficyt sektora instytucji rządowych i samorządowych byłyby obciążone dużą niepewnością wynikającą z wrażliwości szacunków na przyjmowane założenia.

Włączenie zobowiązań niekapitałowych publicznych systemów emerytalnych do statystyk deficytu i długu publicznego wpłynęłoby również na porównywalność międzynarodową tych wskaźników. Dług ogółem zależałby od struktury systemu emerytalnego. Kraje z dominującym publicznym systemem repartycyjnym miałyby większy dług niż te, w których dominują prywatne systemy kapitałowe, których zobowiązania nie byłyby uwzględniane w danych dotyczących długu publicznego. Zobowiązania emerytalne są zazwyczaj większe niż dług konwencjonalny, dlatego miałyby to znaczący wpływ na lokatę krajów w porównaniach międzynarodowych ze względu na relację długu sektora instytucji rządowych i samorządowych do PKB.

Istnieją zatem praktyczne i teoretyczne powody nieujmowania zobowiązań emerytalnych w długu i deficycie sektora instytucji rządowych i samorządowych wykorzystywanych do oceny bieżącej polityki fiskalnej krajów UE.

Podsumowanie

W artykule opisano prace nad metodyką ujmowania nabytych uprawnień emerytalno-rentowych w ramach ubezpieczeń społecznych w rachunkach narodowych. Przedstawiono wytyczne zawarte w ESA 95 oraz rozwiązania przyjęte w SNA 2008 i ESA 2010. Wskazano argumenty merytoryczne przemawiające przeciwko włączeniu zobowiązań emerytalnych rządów do statystyki deficytu i długu sektora instytucji rządowych i samorządowych. Zwrócono uwagę na negatywny wpływ, jaki miałyby to dla porównywalności danych między krajami.

Wśród wniosków płynących z pierwszych, uproszczonych szacunków należy podkreślić wrażliwość wyników na przyjęte założenia, w szczególności dotyczącej stopy dyskontowej i wzrostu wynagrodzeń.

Należy podkreślić, że po raz pierwszy obowiązkowo dane o nabytych uprawnieniach emerytalno-rentowych w ramach ubezpieczeń społecznych państwa członkowskie UE przekażą do Komisji Europejskiej (Eurostatu) w grudniu 2017 r. Szacunki powinny dotyczyć 2015 r. To całkiem nowa tablica i nowa informacja. Jej opracowanie będzie wymagać od statystyków realizacji zadań, które są właściwe aktuariuszom. Dla żadnej innej tablicy programu transmisji danych ESA 2010 nie ustalono tak odległego terminu pierwszej transmisji. Nie można wykluczyć, iż metodyka wykorzystana do obliczeń w 2017 r. będzie podlegała doskonaleniu.

Informacje dotyczyć będą wartości uprawnień nabytych w ramach całego systemu ubezpieczeń społecznych. Będą opracowane na podstawie jednolitych wytycznych i z przyjęciem spójnych założeń. Podawane jako procent PKB mogą stać się przydatne do analiz i porównań w skali UE. Możliwa stanie się analiza zmian, gdyż szacunki będą opracowywane regularnie. Dowiemy się, ile kosztowałoby zrealizowanie wszystkich nabytych w ramach rządowych systemów repartycyjnych zobowiązań emerytalno-rentowych w poszczególnych krajach. Obliczenia mogą także pomóc w ocenie długoterminowych skutków planowanych bądź wdrażanych reform, oceny skali ich wpływu na wartość przyszłych pojedynczych świadczeń oraz skumulowanych zobowiązań.

Ważne, żeby pamiętać, że wartość zobowiązań podawana przez krajowe urzędy statystyczne nie będzie wskaźnikiem świadczącym o kondycji systemu ubezpieczeń społecznych lub o kondycji fiskalnej państwa. Szacunki w ujęciu ADL obejmują bieżącą wartość świadczeń, które mają być wypłacone w przeszłości, a wynikają z już zapłaconych przez osoby pracujące składek oraz pozostałych do zrealizowania uprawnień osób pobierających świadczenia. Będą zatem źródłem informacji o koszcie likwidacji systemu przy jednoczesnej wypłacie świadczeń, do których jego uczestnicy nabyli uprawnienia, do momentu likwidacji systemu i zgodnie z obowiązującymi zasadami.

LITERATURA

- Barr, N., Diamond P. (2014). *Reformy systemu emerytalnego*. PTE, Warszawa.
- Europejski Bank Centralny, Eurostat (2011). *Technical Compilation Guide for Pension Data in National Accounts*. Publications Office of the European Union, Luksemburg.
- Franco, D. (1995). *Pension liabilities — their use and misuse in the assessment of fiscal policies*. European Commission Directorate-General for Economic and Financial Affairs, Bruksela.
- GUS (2014). *Prognoza ludności na lata 2014—2050*.
- Heidler, M., Müller, Ch., Weddige, O. (2009). *Measuring Accrued-to-date Liabilities of Public Pension Systems — Method, Data and Limitations*. Uniwersytet we Freiburgu.
- Kaier, K., Müller, Ch. (2013). *New Figures on Unfunded Public Pension Entitlements Across Europe — Concept, Results and Applications*. Uniwersytet we Freiburgu.
- Semeraro, G. (2007). *Should financial accounts include future pension liabilities*, rozdział *Proceedings of the IFC Conference on Measuring the financial position of the household sector*. Bank for International Settlements, Bazylea.
- Semeraro, G. T. (2011). *Should household wealth and government liabilities include future pension rights?*, rozdział *The Financial Systems of Industrial Countries*. Springer.

Summary. *The article describes the changes solutions in the recognition of acquired pension rights in the framework of social security in national accounts. The principles applicable to September 2014 and the new guidelines contained in the System of National Accounts (SNA 2008) and the European System of Accounts (ESA 2010) are discussed. The author indicates differences between the systems and presents European solutions. Arguments are given that determined the non-inclusion of unfunded pension liabilities of public pension systems to statistics produced for the purposes of the excessive deficit procedure.*

Keywords: national accounts, debt of the government, social security.

Резюме. *В статье обсуждались изменения решений в признании приобретенных пенсионных прав и прав по инвалидности в рамках социального обеспечения в национальных счетах. Были обсуждены правила обязывающие до сентября 2014 г., а также новые рекомендации содержащиеся в Системе национальных счетов (SNA 2008) и в Европейской системе национальных и региональных счетов (ESA 2010). Статья показывает различия между системами и представляет европейские решения. Авторы назвали причины, которые решили не включать неденежные пенсионные обязательства государственных пенсионных систем в статистику разрабатываемую для целей процедуры чрезмерного дефицита.*

Ключевые слова: национальные счета, долг сектора правительственных учреждений и учреждений местного самоуправления, социальное обеспечение.

Dariusz KOTLEWSKI
Mirosław BŁĄŻEJ

Metodologia rachunku produktywności KLEMS i jego implementacja w warunkach polskich

Streszczenie. *W artykule omówiono pochodzenie i status rachunku produktywności KLEMS. Podano szczegóły metodologiczne rachunku, przedstawiono jego międzynarodowy kontekst oraz aspekty implementacyjne w warunkach polskich. Wskazano, że rachunek ten jest prowadzony w GUS. Opisano dostępność danych wykorzystywanych do rachunku w Polsce oraz zaprezentowano metodę uwzględnienia w tych rachunkach pracy i kapitału.*

Słowa kluczowe: rachunek produktywności, KLEMS, czynniki produkcji, czynniki pierwotne, czynnik praca, czynnik kapitał, przyrost produktywności, dekompozycja, kompozycja pracy, godziny przepracowane, pracownicy, godziny na pracownika.

Rachunkowość wzrostu gospodarczego, która jest podstawą rachunku produktywności gospodarki KLEMS, jest rozwijana od wielu lat i cieszy się dużym zainteresowaniem. Nazwa rachunku produktywności gospodarki — KLEMS — pochodzi od pierwszych liter słów w języku angielskim (K — *Capital*, L — *Labour*, E — *Energy*, M — *Materials*, S — *Services*). Wskazuje ona na czynniki, do których ten rachunek się odwołuje. Są to tzw. czynniki pierwotne, czyli kapitał oraz praca i czynniki wtórne, które są składowymi zużycia pośredniego, czyli energia, materiały oraz usługi rozumiane jako wkłady uzyskiwane przez przedsiębiorstwa z zewnątrz. KLEMS jest rachunkiem statystycznym pokazującym procesy gospodarcze *ex post* zasadniczo od strony podażowej. Wywodzi się z neoklasycznej teorii wzrostu gospodarczego, a ściślej z tzw. dekompozycji Solowa z lat pięćdziesiątych XX w. Ta teoria stanowi źródło dwóch głównych nurtów metodologicznych rachunku produktywności obecnie realizowanych na świecie. Jednym z nich jest metodologia OECD, skonstruowana tak, aby zapew-

nić jak najdalej idącą porównywalność międzynarodową, nawet za cenę istotnych kompromisów teoretycznych¹. Drugim jest metodologia EU KLEMS, w większym stopniu przystająca do rygoru zgodnego z głównym nurtem teorii, na której jest oparty rachunek. EU KLEMS obejmuje obecnie 10 krajów europejskich oraz Stany Zjednoczone i Japonię. Ma on swoją odmianę w postaci WORLD KLEMS opartą na tych samych podstawach metodologicznych, która jeszcze się nie upowszechniła, choć w założeniu miała być platformą o zasięgu globalnym. W artykule zaprezentowano podstawy teoretyczne rachunku KLEMS, jego realizację na świecie oraz elementy implementacji w warunkach polskich.

PODSTAWY TEORETYCZNE RACHUNKU PRODUKTYWNOŚCI KLEMS

Inicjalnie tempo wzrostu gospodarczego wiązano tylko z jednym czynnikiem produkcji, tj. kapitałem albo pracą. W latach dwudziestych XX w. przetestowano na danych statystycznych zastosowanie funkcji Cobba-Douglasa jako funkcji produkcji dwóch czynników, tj. kapitału — K i pracy — L w postaci:

$$Y = AK^{\alpha}L^{\beta} \quad (1)$$

gdzie, jeśli się przyjmie, że występują stałe przychody względem skali oraz warunki doskonałej konkurencji, to α stanowi udział wynagrodzenia kapitału w łącznym wynagrodzeniu czynników produkcji, natomiast β — taki udział pracy². Z tego wynika, że PKB (GDP) widniejący po lewej stronie powyższego równania wiąże się z DNB (GDI)³, czyli dochodami czynników produkcji po prawej stronie tego równania. Założenie o stałych przychodach skali znajduje swoje odzwierciedlenie w tym, że parametry α oraz β sumują się do jedności⁴, czyli że $\beta = 1 - \alpha$.

Przy tym założeniu Solow (1956) wyprowadził wzór wiążący wzrost gospodarczy $\Delta Y/Y$, rozumiany jako względny (lub inaczej procentowy) przyrost PKB,

¹ W metodologii OECD dopuszcza się występowanie niestałych przychodów skali oraz niedoskonałej konkurencji. Stosuje się w niej także szereg innych uproszczeń, dzięki którym uzyskuje się lepszą porównywalność międzynarodową — korzystanie z łatwiej dostępnych danych pozwala przeprowadzić rachunek dla większej liczby krajów.

² W ogólnym przypadku jest to elastyczność produkcji względem nakładów czynników produkcji.

³ GDP — *Gross Domestic Product*, GDI — *Gross Domestic Income*.

⁴ Wynika to z tzw. zasady powtarzania, polegającej na tym, że w skali całej gospodarki danego kraju nie występują efekty skali, gdyż wzrost produkcji jest związany wyłącznie ze zwiększeniem ilości obiektów wytwórczych, a nie ze zwiększeniem ich wielkości, z którą związane są te efekty. Jest to jednak kontrowersyjne i dlatego w niektórych metodologiach przyjmuje się możliwość występowania (w ograniczonym zakresie) makroekonomicznych efektów skali, a zatem zakłada się, że $\alpha + \beta > 1$.

ze względnym przyrostem wynagrodzeń dwóch pierwotnych czynników produkcji, znany jako „dekompozycja Solowa”, w postaci:

$$\frac{\Delta Y}{Y} = \frac{\Delta A}{A} + \alpha \frac{\Delta K}{K} + (1 - \alpha) \frac{\Delta L}{L} \quad (2)$$

przy czym:

$$\alpha = r \frac{K}{Y} \quad (3)$$

$$\beta = w \frac{L}{Y} \quad (4)$$

gdzie:

r — stopa procentowa,

w — średnia godzinowa stawka płac.

Podstawiając (3) i (4) do (2) można wykazać, że w tym ostatnim równaniu (2) wyrażenia związane z czynnikami produkcji kapitałem K i pracą L stanowią przyrosty ich wynagrodzeń względnych, ujęte jako ich wkłady (inaczej kontrybucje) we wzrost gospodarczy $\Delta Y/Y$. K należy przy tym rozumieć jako wartość zgromadzonego kapitału, zaś L jako nakład pracy w godzinach⁵. Równania (3) i (4) definiują α i β także jako udziały wynagrodzeń czynników produkcji w gospodarce, czyli $\alpha = WK/Y$, gdzie WK to całkowite wynagrodzenie kapitału w gospodarce oraz $\beta = WP/Y$, gdzie WP to całkowite wynagrodzenie pracy.

We wzorze (2) pojawia się rezydualna wartość $\Delta A/A$, zwana „resztą Solowa” (*Solow's residual*), która reprezentuje nieznany czynnik przyczyniający się również do wzrostu gospodarczego, inny niż wymienione dwa podstawowe⁶. Solow uważał, że owa „reszta” reprezentuje postęp techniczny, ujmowany w tym modelu jako egzogeniczny, który towarzyszy równolegle nakładom kapitału i pracy. Można to również interpretować w ten sposób, że skojarzone nakłady kapitału i pracy dają większy przyrost produktywności w gospodarce niż gdyby wykorzystywano każdy z tych nakładów oddzielnie. Współcześnie najczęściej uważa się, że chodzi głównie o tzw. nieucieleśniony (w kapitale lub pracy) postęp techniczny i/lub organizacyjny. Reszta Solowa jest siłą rzeczy wyliczana rezydualnie⁷, może zawierać wszystkie pozostałe znane i nieznane czynniki przyczynia-

⁵ Znaczenie symbolu L zmienia się, o czym dalej.

⁶ Ważone udziałem w gospodarce przyrosty czynników mają inną łączną wielkość niż przyrost produktu — Hulten (2009), s. 17.

⁷ Choć zawsze spełnione, równanie (2) jest przybliżonym odpowiednikiem równania (1), gdyż wyprowadzone zostało z zastosowaniem uproszczeń. Dla małych zmian (oznaczonych w równaniu symbolem Δ) jest ono zbliżone z równaniem (1). Jeżeli zmiany te są duże, to rezydualnie wyliczona wartość Δ jest inna w obu równaniach.

jące się do wzrostu gospodarczego inne niż kapitał i praca. Funkcja produkcji Cobba-Douglasa i dekompozycja Solowa stały się punktem wyjścia dla rozwoju rachunku KLEMS.

Wątpliwości związane z interpretacją reszty Solowa (często określanej mianem *Total Factor Productivity* — TFP, które to określenie sugeruje, że jest to dodatkowa rezydualna produktywność łączna wszystkich czynników) doprowadziły do rozwinięcia innych wersji tej pierwszej „wieloczynnikowej” teorii wzrostu produktywności gospodarki, w których m.in. uwzględniano dodatkowe czynniki produkcji lub dzielono czynniki kapitał i praca na podczynniki składowe. Podczynniki te nie zawsze były w stosunku do siebie całkowicie rozdzielne, czyli zachodziły na siebie, co powodowało, że część wartości tych czynników była mnożona przez więcej niż jeden udział (z szeregu udziałów oznaczanych zwykle jako: $\alpha, \beta, \gamma, \dots$ itd. lub — jak w podręczniku EU KLEMS — przez v z odpowiednimi indeksami (Timmer i in., 2007a). Dodawanie takich zachodzących na siebie czynników powodowało nieuzasadnione zmniejszanie udziału reszty Solowa w funkcji produkcji. W istocie wymienione w funkcji produkcji czynniki powinny być całkowicie rozdzielne, czyli nie powinny się pokrywać nawet częściowo, aby pomiar reszty Solowa, czyli TFP, był właściwy.

W rozwoju teorii istotną zmianą było podejmowanie prób wykonania dekompozycji przyrostu produkcji globalnej. Wymagało to wprowadzenia czynnika zużycie pośrednie do funkcji produkcji wraz z dekompozycją przyrostu wartości dodanej brutto, które wchodzi w skład produkcji globalnej, zamiast dotychczasowej dekompozycji przyrostu PKB, jak w modelu Solowa. Zmianom tym i dalszym modyfikacjom⁸ towarzyszyło powiązanie koncepcji dekompozycji ze statystyką w oparciu o system rachunków narodowych (SNA). W rezultacie rachunki produktywności obejmują obecnie poddekompozycję na sektory⁹ w ujęciu statystycznym. Wprowadzono nowe określenie dla reszty Solowa — TFP, mianowicie *Multifactor Productivity* — MFP (wieloczynnikowa produktywność gospodarki), choć w innych publikacjach niż publikacje metodologiczne OECD i EU KLEMS starsze określenia są zwykle nadal używane (Timmer i in., 2007a; OECD, 2001)¹⁰.

Według obecnych ustaleń teoretycznych rachunek EU KLEMS zasadniczo powinien opierać się na wieloczynnikowej dekompozycji przyrostu produkcji globalnej, która według teoretycznej podstawy tego rachunku jest najodpowiedniejszym sposobem rezydualnego pomiaru przyrostu wartości wieloczynnikowej

⁸ Wprowadzonym głównie przez autorów Jorgenson (1963), Jorgenson i Griliches (1967), Jorgenson, Gollop i Fraumeni (1987), Jorgenson (1989), Jorgenson, Ho i Stiroh (2005).

⁹ Określenie „sektor” będzie tutaj oznaczać dowolną agregację podmiotów gospodarczych, pośrednią pomiędzy agregatem całej gospodarki a indywidualnymi firmami. Przyjęto to określenie, gdyż w rachunkach typu KLEMS występuje podział na „przemysł” (ang. *Industries*), które odpowiadają różnym poziomom agregacji, tj. grupom sekcji, sekcjom, grupom działów, działom NACE i ewentualnie głębiej.

¹⁰ W niniejszym artykule tradycyjna reszta Solowa oraz MFP będą traktowane jako dwie możliwe odmiany TFP.

produktywności gospodarki (czyli MFP). Teoretycznie MFP odpowiada wówczas najściślejszej idei nieucieleśnionego w czynnikach postępu technicznego i/lub organizacyjnego. Jednak dekompozycja ta jest obciążona istotną wadą polegającą na zaburzającym wpływie stopnia pionowej integracji firm w gospodarce. Im bardziej firmy są w danej gospodarce „pionowo” zintegrowane, tym większa część zużycia pośredniego staje się niewidzialna statystycznie. Wpływa to szczególnie mocno na porównywalność wyników według różnych krajów świata, które jak wiadomo bardzo się różnią w zakresie tej pionowej integracji firm. Aby zminimalizować ten problem, wybiera się odpowiedni podział gospodarki na sektory, w którym zakłada się, że przy pewnym poziomie agregacji pionowa integracja pomiędzy firmami występuje tylko wewnątrz wybranych sektorów, zaś nie występuje pomiędzy sektorami.

Sposobem na uniknięcie tego problemu jest także dekompozycja przyrostu wartości dodanej brutto zamiast przyrostu produkcji globalnej¹¹. Dodatkową przy tym korzyścią jest to, że wartość ta jest zbliżona do PKB, a nawet przyjmowana przez teoretyków za identyczną, jeśli nie brać pod uwagę pozycji „podatki pośrednie i subsydia”. Dekompozycja przyrostu wartości dodanej brutto daje taki sam wynik przyrostu wieloczynnikowej produktywności MFP¹², jak dla dekompozycji przyrostu produkcji globalnej, tylko wtedy jeśli się założy, że poza MFP postęp techniczny i/lub organizacyjny jest ucieleśniony tylko w dwóch czynnikach pierwotnych, czyli w kapitale i pracy, i nie dotyczy zużycia pośredniego. Założenie o braku zmian produktywności zużycia pośredniego jest jednak trudne do utrzymania w tzw. dłuższym okresie.

Przykładowo, może wystąpić spadek materiałochłonności (czynnik M) czy energochłonności gospodarki (czynnik E), które *ceteris paribus* prowadzą do kurczenia się zużycia pośredniego na rzecz MFP. Może także wystąpić zmiana zużycia pośredniego w odwrotnym kierunku, czyli jego powiększanie się spowodowane utrzymującym się wzrostem cen niektórych surowców, szczególnie energetycznych oraz częstszym wykorzystywaniem outsourcingu w gospodarce (czynnik S). Te zmiany niekoniecznie przenoszą się w całości na wielkość produkcji globalnej, a zatem mogą wpływać na wielkość rezydualnego MFP. W rezultacie, gdy występuje zauważalnie wysokie tempo zmian zużycia pośredniego, to może także wystąpić rozbieżność pomiędzy tempem zmian MFP wyliczanym rezydualnie z dekompozycji produkcji globalnej a tempem zmian MFP wyliczanym rezydualnie z dekompozycji wartości dodanej brutto¹³. Dekompozycja wartości dodanej brutto, która nie zawiera zużycia pośredniego, pozwala uniknąć tego problemu.

¹¹ Rozumowanie oparte na opracowaniu — Hulten (2009), s. 25—28.

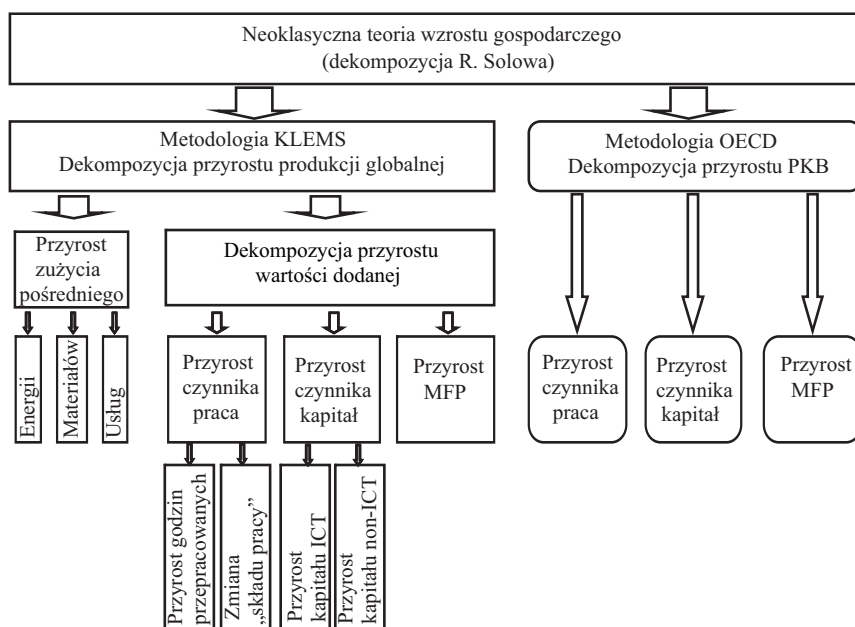
¹² W ujęciu procentowym pomnożony przez stosunek wartości produkcji globalnej do wartości dodanej.

¹³ Różnica ta nie występowałaby, gdyby stosowano czas ciągły, a nie dyskretny, a to dlatego, że wartość dodana brutto jest również zwykle obliczana rezydualnie jako różnica pomiędzy produkcją globalną a zużyciem pośrednim, które są danymi otrzymywanymi wprost od przedsiębiorstw.

Ostatecznie rachunek produktywności gospodarki realizowany dla większości krajów na platformie EU KLEMS oparto na pomiarach przyrostu różnych odmian nakładów kapitału (K) i pracy (L) dla dekompozycji przyrostu wartości dodanej brutto. W przypadku ewentualnej dekompozycji przyrostu produkcji globalnej, dokonywanej fakultatywnie dla poszczególnych krajów przy założeniu braku możliwości jej wykorzystywania do porównań międzynarodowych, stosuje się także dekompozycję wykorzystującą przyrosty składowych zużycia pośredniego w postaci energii (E), materiałów (M) i usług (S) ujmowanych wartościowo¹⁴.

Pomiary wspomnianych wielkości — w rozbiciu na agregaty pogrupowane według linii podziału na odpowiednie grupy sekcji, sekcje, grupy działów i działy NACE (dla krajów pozaeuropejskich jej odpowiednika ISIC), zgodnie z wymaganiami danego rachunku — są podstawą dla tworzenia baz danych, które należy wykorzystać w algorytmach tu zaprezentowanych.

SCHEMAT DWÓCH DOMINUJĄCYCH METODOLOGII I RACHUNKU PRODUKTYWNOŚCI GOSPODARKI



Źródło: opracowanie własne.

Na tej samej teorii wzrostu gospodarczego oparto metodologię stosowaną w rachunku produktywności OECD. Jednak, w założeniu, ma on służyć porównywalności gospodarki możliwie największej grupy krajów. Dlatego realizowa-

¹⁴ Pozwala to ukazać, jaka część zmian MFP jest związana ze zmianami składowych zużycia pośredniego.

na jest obecnie tylko dekompozycja PKB bez dekompozycji produkcji globalnej lub wartości dodanej brutto (OECD, 2001; OECD, 2009; a szczególnie OECD, 2013, s. 66–70). Zrezygnowano całkowicie z wszelkich odniesień do zużycia pośredniego. Dekompozycja PKB w metodologii OECD uwzględnia także „poluzowanie” niektórych bardzo sztywnych założeń, takich jak stałe przychody skali, które traktuje się jako obowiązujące w przybliżeniu. Ponadto w metodologii OECD nie praktykuje się dalszej dekompozycji pierwotnych czynników produkcji, jak pokazano na schemacie, na którym zestawiono ze sobą te dwie dominujące metodologie dekompozycji wzrostu gospodarczego.

Jedną z ważniejszych zmian w stosunku do tradycyjnej dekompozycji Solowa jest także rozszerzenie dekompozycji czynnika praca. L w przytoczonych wyżej wzorach oznacza czynnik praca pojmowany jako godziny pracy. W obecnie prowadzonym w Polsce rachunku KLEMS¹⁵ symbol ten jednak oznacza całkowite wynagrodzenie pracy (WP), czyli $L = WP$. Tempo przyrostu wynagrodzenia całego czynnika praca jest inne niż tempo przyrostu ilości godzin przepracowanych¹⁶. Na skutek takiej zmiany w rachunku KLEMS dekomponuje się czynnik praca na dwie składowe według równania:

$$\frac{\Delta L}{L} = \frac{\Delta LC}{LC} + \frac{\Delta H}{H} \quad (5)$$

gdzie symbol H oznacza godziny przepracowane, zaś rezydualnie obliczona wartość LC (*labor composition*) to „skład pracy” lub „kompozycja pracy” (dwa proponowane określenia)¹⁷. Zmiana kompozycji pracy teoretycznie polegać ma na zmianie udziału różnych rodzajów pracy, czyli wzrost (spadek) udziału pracy lepiej opłacanej w stosunku do pracy gorzej opłacanej jest odpowiedzialny za wyższe (niższe) tempo przyrostu wynagrodzenia całego czynnika pracy ($\Delta L/L$) w stosunku do przyrostu ilości godzin przepracowanych ($\Delta H/H$). Oczywiście, na skutek tej zmiany, zmiana produktywności czynnika praca zostaje „wyjęta” z „tradycyjnej” reszty Solowa, czyli MFP jest od tej tradycyjnej reszty Solowa zwykle odpowiednio mniejsze i ściślej odpowiada jego koncepcji postępu technicznego nieucieleśnionego w pracy i kapitale. To wszystko oczywiście przy założeniu, że przyrost wynagrodzenia czynnika praca jest równy krańcowemu (inaczej marginalnemu) przyrostowi jego produktywności¹⁸.

¹⁵ Jako skutek zmian wprowadzonych przez Jorgenson i Griliches, 1967; Jorgenson i in., 1987; Jorgenson i in., 2005.

¹⁶ Według teorii rachunku KLEMS pojęcie „godziny przepracowane” jest najodpowiedniejsze dla tego rachunku.

¹⁷ Kompozycja pracy może być szeroko rozumiana, jak tutaj, i wąsko rozumiana, jak na platformie EU KLEMS. Wąsko rozumiana kompozycja pracy zawiera tylko wpływ zmian grup wiekowych, płci i wykształcenia pracowników na ich wynagrodzenia. Z kolei szeroko rozumiana kompozycja pracy obejmuje także pozostałe zmiany wynagrodzeń.

¹⁸ Jak przedstawia schemat, zaproponowano także dekompozycję czynnika kapitał na kapitał ICT i kapitał Non-ICT, lecz nie wszyscy uczestnicy rachunku produktywności KLEMS ją stosują.

Inną ważną zmianą o charakterze technicznym jest stosowanie wyrażeń logarytmicznych zamiast zwykłych wzorów na przyrosty względne. Powodem takiej zmiany jest zalecana przez teorię statystyki procedura Tornqvista stosowana przy agregacji wybranych sektorów gospodarki, w której stosowanie wyrażeń logarytmicznych jest wymagane. Ponadto addytywne człony równań logarytmicznych są ściśle zgodne z ich multiplikatywnymi odpowiednikami, jak w funkcji produkcji o postaci typu (1).

METODOLOGIA RACHUNKU KLEMS I CZYNNIK PRACA

Podręcznik EU KLEMS (Timmer i in., 2007a), wykorzystujący zarysowane założenia teoretyczne, wychodzi z następującej formuły na wzrost produkcji globalnej w danym sektorze j w okresie t :

$$\Delta \ln Y_{jt} = \bar{v}_{jt}^X \Delta \ln X_{jt} + \bar{v}_{jt}^K \Delta \ln K_{jt} + \bar{v}_{jt}^L \Delta \ln L_{jt} + \Delta \ln A_{jt}^Y \quad (6)$$

gdzie:

Y — produkcja globalna,
 X — zużycie pośrednie,
 K — kapitał,
 L — praca,
 A^Y — MFP.

Wartości te są subskrybowane, dotyczą sektorów j i okresów t . Δ oznacza dla wartości pod tym znakiem ich zmianę pomiędzy okresem uprzednim $t-1$, a obecnym t , które zwykle identyfikowane są jako okresy jednoroczne. Generalnie, jeżeli zmiany są niewielkie, jak zwykle w okresach jednorocznych, zachodzi w przybliżeniu $\Delta \ln x = \Delta x/x$. W tym przypadku chodzi o zmianę względną, np. wyrażoną w sposób procentowy. Oprócz tego uzasadnienia stosowanie wyrażeń logarytmicznych jest także bardziej właściwe dla zmian większych, obserwowanych w okresach wieloletnich, gdyż nie amplifikuje rozbieżności i narastania błędów w wyniku kolejnych obliczeń (Schreyer, 2004; Milana, 2009). Z kolei $\bar{v}$ z odpowiednimi indeksami oznacza średni udział, w ujęciu wartościowym, danego czynnika (określonego w superskrypcie jako X , K i L) pomiędzy okresami $t-1$ i t , który wylicza się według wzoru $\bar{v} = (v_{t-1} + v_t)/2$ (dla prostoty pominięto subskrypty obecne we wzorze (6)). Dokonuje się zatem interpolacji liniowej dla udziałów, której celem jest zmniejszenie nieścisłości, jakie byłyby obecne w rachunku, gdyby przyjęto udziały bieżące. Przyrost wartości A^Y , czyli MFP, jest rezydualnie wyliczany z tego wzoru, tak iż wzór (6) jest *ex post* fundamentalnie zawsze spełniony.

Podobną funkcję stosuje się do dekompozycji wartości dodanej brutto, ale nie występuje w niej zużycie pośrednie X jako czynnik składowy:

$$\Delta \ln V_{jt} = \bar{w}_{jt}^K \Delta \ln K_{jt} + \bar{w}_{jt}^L \Delta \ln L_{jt} + \Delta \ln A_{jt}^V \quad (7)$$

gdzie V to wartość dodana, zaś pozostałe symbole (z odpowiednimi indeksami) mają takie samo znaczenie, jak we wzorze (6), ale, oprócz kapitału K i pracy L , oznaczają inne wartości. Analogiczne średnie udziały $\bar{w}$ nie są identyczne ze średnimi udziałami $\bar{v}$ (są one w ujęciu procentowym oraz wyliczane podobnie jak średnie udziały $\bar{v}$), jak również wkład MFP przy dekompozycji wartości dodanej A^V nie jest identyczny w ujęciu względnym z wkładem MFP przy produkcji globalnej A^Y , choć jej przyrost absolutny w idealnym przypadku, gdy nie występuje zmiana produktywności zużycia pośredniego, powinien być identyczny.

We wzorach tych przyrost czynnika praca¹⁹ w sektorze j definiuje się według wzoru:

$$\Delta \ln L_t = \sum_l \bar{v}_{l,t} \Delta \ln H_{l,t} \quad (8)$$

gdzie:

L — wartość czynnika praca (obejmująca kompozycję pracy),

l — rodzaj czynnika praca,

$\bar{v}$ — z odpowiednimi indeksami oznacza średnie udziały poszczególnych rodzajów czynnika praca l pomiędzy okresami $t-1$ i t (wyliczane analogicznie do tych udziałów jako średnia arytmetyczna),

H — z odpowiednimi indeksami oznacza ilość godzin przepracowanych dla danego rodzaju czynnika praca l w okresie t .

We wzorze (8) (w ślad za podręcznikiem EU KLEMS) pominięto subskrypt j określający sektor, który „normalnie” powinien widnieć obok wszystkich innych subskryptów tego równania.

Zakłada się, że tzw. usługi czynnika praca każdego rodzaju l wyrażone wartościowo są proporcjonalne do ilości godzin przepracowanych w tym rodzaju pracy (czyli ich przyrosty względne powinny być identyczne, jak dla godzin przepracowanych $\Delta \ln H_{l,t}$), zaś pracownicy danego rodzaju pracy są opłacani według ich krańcowych produktywności. Jest to uwzględnione w udziale wynagrodzenia danego rodzaju pracy $\bar{v}_{l,t}$ w łącznym wynagrodzeniu wszystkich rodzajów pracy. Te rodzaje czynnika praca l są wyróżnione w metodologii rachunku EU KLEMS

¹⁹ Podobnie postępuje się z pozostałymi czynnikami, tj. zużyciem pośrednim X i kapitałem K . Dyskusja na temat tych czynników wykracza jednak poza ramy niniejszego artykułu.

według trzech poziomów wykształcenia, płci i trzech grup wiekowych, a ich udziały w sektorze j są wyliczane wartościowo.

Wzór (8) uwzględnia niejednorodność czynnika praca pomiędzy sektorami z punktu widzenia jego wynagrodzenia. Jednocześnie fizyczny przyrost czynnika praca (roboczogodziny) jest ujmowany względnie, czyli w procentach w postaci przyrostu wyrażenia logarytmicznego po prawej stronie. Gdy tempo przyrostu godzin przepracowanych w różnych rodzajach pracy l jest identyczne, wówczas nie ma różnicy pomiędzy tempem przyrostu godzin przepracowanych $\Delta H/H$ a tempem przyrostu wynagrodzenia całego czynnika praca $\Delta L/L$, czyli $\Delta LC/LC = 0$ (wzór 5). Wartość ta nie jest równa 0, gdy tempo przyrostu godzin przepracowanych w rodzajach pracy l , które są lepiej wynagradzane, jest inne od tempa przyrostu godzin przepracowanych w rodzajach pracy l , które są gorzej wynagradzane. Może to wynikać z różnego tempa przyrostu godzin przepracowanych w rozmaitych sektorach j o różnym wynagradzaniu za pracę, jak również z innego tempa przyrostu godzin przepracowanych dla różnych rodzajów pracy wewnątrz poszczególnych sektorów. Wprowadzenie subskryptu j do wzoru (8) dekomponuje tę analizę na sektory, jednak nie wpływa na końcowy wynik rachunku dla zagregowanego przyrostu LC . Konsekwentnie jednak w rachunku KLEMS zaleca się stosowanie wyrażen logarytmicznych zamiast zwykłych przyrostów względnych, jak we wzorze (5), gdy dokonuje się dekompozycji na sektory j , które agreguje się w procedurze Tornqvista. Dekompozycja czynnika praca dla sektora j na jego podczynniki LC i H przyjmuje zatem następującą postać:

$$\Delta \ln L_t = \Delta \ln LC_{jt} + \Delta \ln H_{jt} \quad (9)$$

przy czym wyrażenie związane z LC oblicza się rezydualnie, gdyż statystycy dysponują tylko danymi dotyczącymi wielkości L (obejmującą kompozycję pracy) oraz liczby godzin przepracowanych H .

We wdrażaniu rozwiązań teoretycznych problemem podstawowym w wielu krajach i w wielu przypadkach jest to, że czynnik praca jest rejestrowany w nieodpowiednich jednostkach (w osobach zamiast w roboczogodzinach). Kolejnym problemem jest to, że gdy pracę ujmuje się w roboczogodzinach, są one traktowane jako opłacane, a nie przepracowane roboczogodziny. Problemem jest też uwzględnienie zjawiska samozatrudnienia. Ze względu na liczne ograniczenia w zakresie dostępności danych stosuje się różne metody zastępczych estymacji.

Ze względu na to, że rachunek produktywności OECD jest adresowany do większej grupy krajów i ma na celu zapewnienie porównywalności pomiędzy nimi, czynnik praca nie jest w nim dekomponowany na różne rodzaje pracy l , tj. według wieku, płci i wykształcenia, gdyż dane tego rodzaju nie są dostępne dla większości krajów. Czynnik praca jest zatem tutaj dekomponowany jedynie według sektorów. Podział natomiast na te sektory jest w metodologii OECD

bardziej szczegółowy, co częściowo niweluje brak podziału na rodzaje pracy²⁰. W przypadku braku danych dotyczących godzin przepracowanych według sektorów wartość zagregowana dla całej gospodarki jest rozszacowywana na sektory według (kolejność według malejącej priorytetowości) pełnych etatów „ekwiwalentnych”, pełnozatrudnionych pracowników „ekwiwalentnych”, wielkości zatrudnienia lub liczby pracowników (Arnaud, Dupont, Seung-Hee i Schreyer, 2011). W przypadku Polski, OECD wykorzystuje tu trzecią z tych czterech możliwości rozszacowania, tj. wielkość zatrudnienia (czyli godziny przepracowane). Tę metodę stosuje także większość krajów uczestniczących w EU KLEMS, w celu rozszacowania zagregowanych w sektory wartości na 3 grupy wiekowe, według płci i 3 poziomów wykształcenia, czyli według ww. rodzajów pracy, których jest 18 ($3 \times 2 \times 3$).

STATUS RACHUNKU KLEMS NA ŚWIECIE

Rachunek KLEMS dzieli się zasadniczo na dwie części — WORLD KLEMS oraz EU KLEMS, które w zasadzie funkcjonują tylko jako odrębne platformy internetowe służące do prezentowania uzyskanych wyników przez kraje wykonujące ten rachunek. W 2013 r. na stronie internetowej WORLD KLEMS umieszczono 5 plików tablic z naliczonymi danymi w formacie Excel. Reprezentują one: Rosję, Japonię, Stany Zjednoczone i Kanadę, przy czym Stany Zjednoczone publikują 2 pliki tablic nieznacznie różniących się metodologicznie. Ostatnio (2015—2016) rozpoczęto umieszczanie na tej platformie danych dotyczących Korei Południowej i Indii. Poza tym podano 3 linki: do EU KLEMS i do indywidualnie realizowanych przez Chiny i Argentynę rachunków typu KLEMS²¹. Tworzone są również platformy regionalne Asia KLEMS i Latin America KLEMS. Podano również linki do stron internetowych krajowych instytucji statystycznych (tzw. NSI)²² dla: Australii, Danii, Finlandii, Holandii, Meksyku, Stanów Zjednoczonych, Szwecji, Wielkiej Brytanii i Włoch (według stanu z początku 2016 r.).

²⁰ Wkład zmiany „struktury subtelnej” według wynagrodzeń (czyli zmiany udziału poszczególnych rodzajów pracy, jak w EU KLEMS) oraz według podziału na sektory (uproszczonego do sekcji i grup sekcji w EU KLEMS, ale bardziej szczegółowego w metodologii OECD) ma decydujący wpływ na zmniejszenie wielkości reszty Solowa. Jeżeli kompozycja pracy (czyli wielkość i struktura wynagrodzeń) się zmienia, to MFP może być dużo mniejsza od klasycznej reszty Solowa, przy której nie wyliczano wkładu zmiany struktury wynagrodzeń (OECD, 2001, s. 47). Bardziej subtelny podział na sektory zastosowany w metodologii OECD częściowo przechwytuje efekty różnic w wynagrodzeniach związanych z jakością pracy (wąsko rozumiane sektory wysokopłatne i niskopłatne), co częściowo rozwiązuje problem zmian w jakości pracy (OECD, 2001, s. 48). W rachunku KLEMS jednak uwzględnione są także zmiany w jakości pracy wewnątrz sektorów.

²¹ Każdorazowo, gdy użyto dalej określenia KLEMS, dotyczy to zarówno EU KLEMS, jak i WORLD KLEMS, a przede wszystkim koncepcji metodologicznej, na której te platformy internetowe zbudowano.

²² NSI, czyli *National Statistical Institute*.

Wspomniane 4 kraje wykonują rachunek produktywności na platformie WORLD KLEMS nieco inaczej z punktu widzenia doboru zmiennych KLEMS oraz okresu czasu, dla którego przeliczono dane. Rosja np. za okres 1995—2009²³ stosuje klasyfikację NACE 1 w podziale na 34 sektory. Szczególna sytuacja występuje w przypadku Japonii, która wykorzystuje w kraju własną oryginalną klasyfikację JIP (*Japan Industrial Productivity Database*) z podziałem gospodarki na 108 sektorów. Dla potrzeb rachunku KLEMS, ten układ danych „przełożono” na wszystkie sektory międzynarodowego systemu klasyfikacyjnego ISIC 3 (stanowiącego odpowiednik europejskiego NACE 1 stosowanego w okresie 1973—2009). W Japonii tylko nominalne wartości są dostępne dla wartości dodanej brutto, dlatego jej realny przyrost został wyliczony rezydualnie poprzez odjęcie realnego przyrostu zużycia pośredniego od realnego przyrostu produkcji globalnej dla poszczególnych sektorów. W Stanach Zjednoczonych stosowany jest system klasyfikacyjny NAICS (*North American Industry Classification System*) z podziałem na 70 sektorów, który dla potrzeb rachunku KLEMS należało „przełożyć” na 31 sektorów systemu klasyfikacyjnego ISIC 3. Zamiast zastosować odpowiedni deflator dla wartości dodanej, wykorzystano tę samą co dla Japonii metodę wyliczenia jej przyrostu dla indywidualnych sektorów. System NAICS jest także stosowany w Kanadzie, dlatego dane w tym systemie zostały także przełożone na 31 sektorów systemu ISIC 3. Rachunki WORLD KLEMS Kanady są oparte na największej liczbie zmiennych spośród wszystkich krajów prowadzących ten rachunek na obu wymienionych platformach (EU i WORLD KLEMS), są zatem najbardziej szczegółowe. Różnice tego rodzaju (a nawet większe) dotyczą również pozostałych krajów niedawno umieszczonych na stronie WORLD KLEMS. Kraje te prezentują dane w nieco innych układach.

Z punktu widzenia liczby krajów realizujących rachunek KLEMS, tych które realizują jego odmianę jako rachunek EU KLEMS było do niedawna dużo więcej. W systemie klasyfikacyjnym NACE 1 jest ich ok. 30 i wykonywały one rachunek EU KLEMS w podziale na 72 sektory gospodarki. Jednak z uwagi na to, że rachunek EU KLEMS w tej klasyfikacji nie jest już kontynuowany, w diagramie 1 zaprezentowano tylko te kraje, które uczestniczą w realizowanym obecnie rachunku EU KLEMS w systemie klasyfikacyjnym NACE 2 z podziałem na 34 sektory, przy czym dane są prawie dla wszystkich krajów (oprócz Finlandii) równolegle prezentowane na platformie internetowej EU KLEMS w obu tych systemach klasyfikacyjnych (są one ciągle przeliczane z jednego systemu na drugi). W diagramie 1 zaprezentowano również tylko 4 kraje inicjalnie obecne na platformie WORLD KLEMS, gdyż pozostałe kraje, które niedawno dołączyły, prezentują nadal dane w nieco innych układach.

²³ W EU KLEMS stosuje się w NACE 1 podział na 72 sektory, zaś w NACE 2 na 34 sektory, czyli dla Rosji występuje łączenie starszej i nowszej metodologii EU KLEMS. Stosowanie zróżnicowanej w pewnym zakresie metodologii nie jest jednak przypadkiem odosobnionym wśród krajów wykonujących rachunek produktywności typu KLEMS.

**DIAGRAM 1. SPÓJNE METODOLOGICZNIE KLUBY^a KRAJÓW KLEMS
WEDŁUG STANU Z 2014 R.**

WORLD KLEMS					EU KLEMS												
Rosja	Japonia I	Stany Zjednoczone I	Stany Zjednoczone II	Kanada	Japonia II	Stany Zjednoczone III	Szwecja	Finlandia	Belgia	Holandia	Wielka Brytania	Niemcy	Włochy	Austria	Hiszpania	Francja	

^a Przez klub należy rozumieć zdefiniowaną spójną metodologicznie grupę krajów.

Źródło: opracowanie własne.

Pomiędzy 12 krajami, które obecnie uczestniczą w EU KLEMS w systemie klasyfikacyjnym NACE 2 także występują różnice metodologiczne. I tak statystyka szwedzka wykorzystuje dla potrzeb krajowych własny podział na 52 sektory, realizowany jednak według linii podziału zgodnych z NACE 2, które dla lat 1993—2011 przełożono na 35 sektorów²⁴. Finlandia pracuje wyłącznie według klasyfikacji NACE 2 w podziale również na 35 sektorów²⁵, które prezentuje dla okresu 1975—2012. W przypadku Japonii, która odrębnie uczestniczy zarówno w EU KLEMS²⁶, jak i w WORLD KLEMS, wymagane były identyczne zabiegi, jak dla WORLD KLEMS, z tym że zastosowano docelowo podział na 35 wybranych sektorów według linii podziału w klasyfikacji NACE 2 (odpowiednika międzynarodowego systemu klasyfikacyjnego ISIC 4, a nie ISIC 3). Stany Zjednoczone są drugim oprócz Japonii krajem prezentującym dane na obu platformach WORLD KLEMS i EU KLEMS, postąpiono zatem podobnie, „przełożono” wybrane 63 sektory z systemu NAICS na 35 sektorów według linii podziału w NACE 2, ale tylko dla okresu 1998—2010. W związku z tym kraj ten prezentuje dane w trzech układach — dwa na WORLD KLEMS i jeden na EU KLEMS.

²⁴ Nie wiadomo dlaczego dokumenty źródłowe podają liczbę 35 zamiast prawidłowo 34! Być może jest to tzw. powielona literówka — Gouma i Timmer (2013a); Gouma i Timmer (2013b), natomiast starsze opisy wykorzystywanych źródeł — Timmer i in. (2007b).

²⁵ Uwaga jak wyżej.

²⁶ Najwyraźniej nie ma przeciwskażeń do uczestnictwa w tym projekcie krajów pozaeuropejskich.

Diagram 1 w symboliczny i uproszczony sposób przedstawia strukturę rachunku produktywności KLEMS na świecie. Komórki poniżej nagłówka symbolicznie przedstawiają bloki naliczanych danych. Oprócz już wymienionych wyżej, pozostałe 8 krajów uczestniczących w EU KLEMS tworzy „ścisły blok metodologiczny” charakteryzujący się identycznym doбором zmiennych podstawowych (patrz: pliki Excel w kolumnie *Basic file* na stronie internetowej EU KLEMS, <http://www.euklems.net/>). W tym bloku przede wszystkim zrezygnowano z odrębnego przeliczania i wykorzystywania zmiennych dotyczących produkcji globalnej oraz zużycia pośredniego (pierwszy rząd komórek pod nagłówkiem) i skupieniu się na dekompozycji wartości dodanej, co oznaczono szarym polem w trzech odcieniach. Wśród nich 7 krajów (bez Belgii) oznaczono szarym polem w dwóch ciemniejszych odcieniach, gdyż korzystają z bazy danych, która jest zgodna z systemem STAN (*OECD Structural Analysis Database*) wykorzystywanym w rachunku produktywności OECD, z wyjątkiem danych dotyczących kapitału. 5 krajów spośród owych 7 stosuje identyczne zmienne, co oznaczono najciemniejszym szarym odcieniem. 7 krajów ścisłego bloku metodologicznego stosuje także uproszczony system konwersji z NACE 1 na NACE 2 dla czynnika praca, polegający na przypisaniu do siebie w przybliżeniu odpowiadających sobie 14 grup sekcji i sekcji z obu systemów klasyfikacyjnych. Finlandia dokładnie rozpisuje czynnik praca na wszystkie 34 sektory EU KLEMS dzięki temu, że posługuje się tylko klasyfikacją NACE 2²⁷. Oprócz owych 7 krajów wraz z Finlandią, inne kraje EU KLEMS nie prezentują danych dotyczących czynnika praca w kolumnie *Labour input file* na stronie internetowej EU KLEMS.

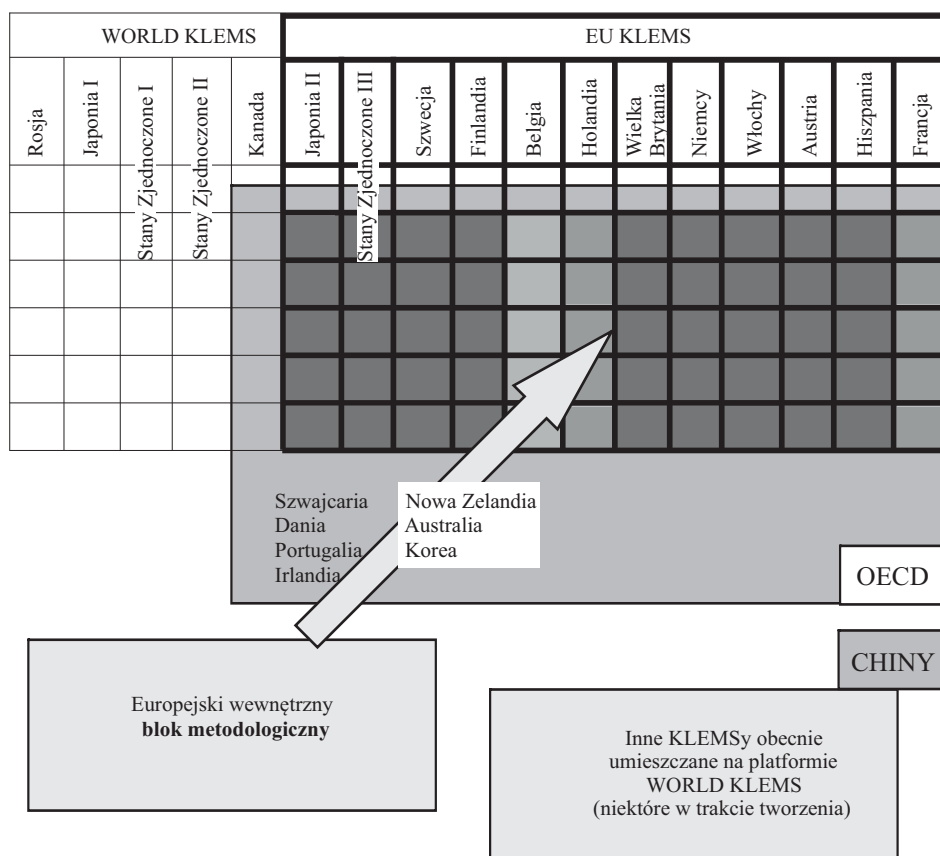
Wymieniony ścisły blok metodologiczny składa się z 7 krajów Europy Zachodniej: Austria, Francja, Hiszpania, Holandia, Niemcy, Włochy i Wielka Brytania, do których można jeszcze doliczyć ósmy kraj — Belgię, nieco luźniej związaną z tym blokiem. Można uznać, że stosowanie metodologii tej grupy krajów jest najbardziej uzasadnione z punktu widzenia porównywalności międzynarodowej, szczególnie w przypadku krajów europejskich, które mogłyby dołączyć do rachunku EU KLEMS, takich jak np. Polska. Wynika to także z przesłanek geograficznych oraz związanych z tym podobieństw gospodarczych i ujednoliconych systemów statystycznych.

Dekompozycje podobne do rachunku produktywności WORLD KLEMS i EU KLEMS są prowadzone też dla innych krajów, co przedstawiono na diagramie 2. Charakterystyczny jest przy tym blok 20 krajów OECD, dla których dekompozycja wzrostu gospodarczego (produktywności gospodarki od strony podaźowej) jest wykonywana według nieco innej metodologii niż dla dwóch platform KLEMS.

²⁷ Wynika to także z faktu, że wykorzystuje znacznie krótsze szeregi czasowe w porównaniu z innymi krajami prowadzącymi rachunek produktywności KLEMS.

12 krajów EU KLEMS oraz 1 kraj WORLD KLEMS (Kanada) uczestniczą także w rachunku produktywności OECD, który jest realizowany jeszcze dla 7 krajów (diagram 2). Według specyficznej metodologii Chiny, prowadzą rachunek produktywności z podziałem na prowincje, a nie na sektory (Kang i Peng, 2013). Ponadto zakłada się, że do systemu WORLD KLEMS będą dołączać inne bloki regionalne, takie jak EU KLEMS. Jednak WORLD KLEMS pozostaje jeszcze w stadium załączkowym, a EU KLEMS jest samodzielną platformą (*de facto* ponadregionalną skoro dołączyły do niej Stany Zjednoczone i Japonia), której dorównuje jedynie platforma OECD. Łącznie na świecie jest ok. 40 krajów, które prowadzą lub podejmują próbę tworzenia rachunku KLEMS.

DIAGRAM 2. KLUBY KLEMS A KLUB OECD WEDŁUG STANU Z 2014 R.



Źródło: jak przy diagramie 1.

Status metodologii wykonywania dekompozycji wzrostu gospodarczego można także ujmować w jeszcze szerszym kontekście, co pokazano na diagramie 3.

Od pewnego czasu Komisja Europejska²⁸ i Stany Zjednoczone wdrażają metodologię badania wzrostu potencjalnego i tzw. luki produktowej (*potential growth & output gap*) z wykorzystaniem funkcji produkcji wywiedzionej z teorii ekonomii. Ta funkcja to nic innego, jak tylko funkcja produkcji identyczna z tą, która posłużyła do wyprowadzenia dekompozycji Solowa. Zamiast wykorzystywać do tego celu metody czysto statystyczne polegające na badaniu własności szeregów czasowych i ich korelacji z cyklem koniunkturalnym, wprowadza się do funkcji czynniki produkcji oraz TFP (*Total Factor Productivity*), oczyszczone z elementów cyklicznych za pomocą dodatkowych technik ich obliczania²⁹.

To pozwala wyznaczyć nieobserwowany tzw. produkt potencjalny (rozumiany jako PKB). Z kolei porównanie owego produktu potencjalnego z faktycznym produktem (obserwowalnym) umożliwia wyznaczenie tzw. luki produktowej. Technika tą wykorzystuje się do wykonywania prognoz do 10 lat od roku wyjściowego, przy założeniu braku zmian w polityce gospodarczej i braku nieoczekiwanych wydarzeń o pochodzeniu egzogenicznym. W tej metodologii TFP jest klasyczną resztą Solowa, czyli dekompozycja w zakresie czynnika praca obejmuje tylko godziny przepracowane bez zmiany tzw. składu lub kompozycji pracy, czyli wkładu zmiany struktury i wielkości wynagrodzeń, jak to jest w rachunku KLEMS.

PRZYGOTOWANIE DANYCH W RACHUNKU KLEMS DLA POLSKI³⁰

W świetle powyższych ustaleń uzasadnione jest przeprowadzanie rachunku KLEMS dla Polski na podstawie założeń zbliżonych do stosowanych przez kraje EU KLEMS, dla których podstawowym podejściem jest dekompozycja wartości dodanej brutto³¹.

Obliczenia w zakresie czynnika praca

W przypadku czynnika praca dane GUS w ramach badania reprezentacyjnego *Badanie struktury wynagrodzeń według zawodów* (formularz Z-12) są dostępne za lata parzyste: 2004, 2006, 2008, 2010 i 2012, czyli do wykonania rachunku KLEMS dla Polski od 2005 r. są wystarczające³². Za lata nieparzyste należało

²⁸ Dokument: ECFIN Economic Paper No. 535 (2014) stanowi kontynuację poprzedniego: ECFIN Economic Paper No. 420 (2010), który z kolei jest kontynuacją: ECFIN Economic Paper No. 247 (2006) oraz ECFIN Economic Paper No. 176 (2002).

²⁹ Przyjmuje się w tym celu, że $Y = (U_L L E_L)^{\alpha} (U_K K E_K)^{1-\alpha} = L^{\alpha} K^{1-\alpha} * TFP$, czyli że $TFP = (E_L^{\alpha} E_K^{1-\alpha}) (U_L^{\alpha} U_K^{1-\alpha})$, patrz: ECFIN ... 535 (2014), w przypisie powyższym.

³⁰ Autorzy, którzy są jednocześnie wykonawcami rachunku produktywności KLEMS dla gospodarki polskiej, pragną podziękować pracownikom GUS za życzliwą współpracę i przekazanie odpowiednich danych do rachunku KLEMS.

³¹ W celu uzyskania informacji o starszych naliczeniach, także dla Polski dokonywanych tylko częściowo i w systemie NACE 1 (Timmer i in., 2007b).

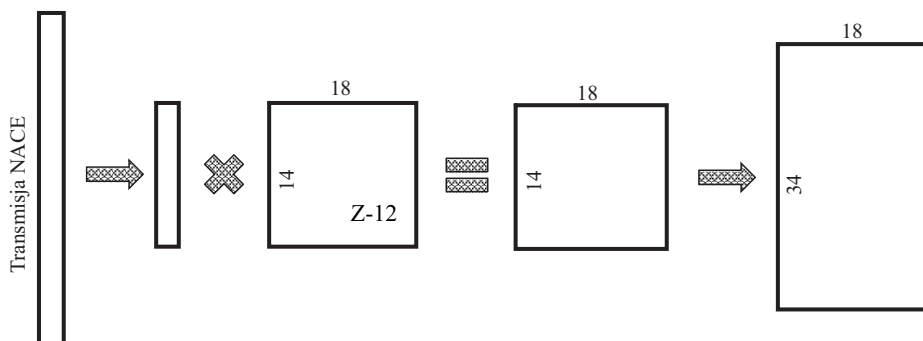
³² Przed 2004 r. badania w zakresie czynnika praca były prowadzone w sposób niesystematyczny, mogą zatem wystąpić trudności w rozszerzeniu wstecz okresu dla rachunku KLEMS dla Polski.

dokonać interpolacji liniowej. Dane w zakresie czynnika praca za 2004 r. dotyczą liczby pracowników pełnozatrudnionych, przeciętnego godzinowego wynagrodzenia brutto w zł za godzinę faktycznie przepracowaną w czasie nominalnym i nadliczbowym w całym roku przez pracowników pełnozatrudnionych w złotych oraz liczby godzin faktycznie przepracowanych przez pracowników pełnozatrudnionych. Od 2006 r. dane dotyczą również pracowników zatrudnionych, a nie jak wcześniej tylko pełnozatrudnionych. Dane są doszacowywane strukturą obejmującą cały rynek pracy, przy czym ewentualny błąd stał się tutaj marginalny.

Dane za lata 2004 i 2006 są dostępne według PKD 2004, natomiast od 2008 r. według PKD 2007, jednak dane za rok 2008 przeliczono na system PKD 2004. Uproszczona korespondencja pomiędzy tymi dwoma systemami klasyfikacyjnymi, stosowana przez kraje ścisłego bloku metodologicznego, zastosowana tylko dla czynnika praca (w podziale na 14 grup sekcji i sekcje) pozwala ominąć niespójność danych z tego wynikającą, ale tylko w sytuacji, gdy przyrosty pomiędzy latami 2007 i 2008 są naliczane w PKD 2004, natomiast przyrosty pomiędzy latami 2008 i 2009 są naliczane w PKD 2007. Oznaczało to wykorzystywanie danych za rok 2008 w zależności od potrzeb albo w jednej albo w drugiej klasyfikacji.

Dane są dostępne zgodnie z wymaganiami rachunku EU KLEMS w podziale na płeć, trzy grupy wiekowe (15—29 lat, 30—49, 50 lat i więcej) i trzy poziomy wykształcenia, czyli w podziale na 18 rodzajów czynnika praca ($2 \times 3 \times 3$). Powstała „macierz” o przyjętej rozpiętości „poziomej”, sięgającej 18 rodzajów pracy oraz o rozpiętości „pionowej”, sięgającej 14 wspólnych dla obu klasyfikacji PKD grup sekcji i sekcji. Te dane dotyczyły ok. 7—8 mln zatrudnionych, a zatem nie objęły całego rynku pracy.

DIAGRAM 4. PRZYGOTOWANIE DANYCH DOTYCZĄCYCH RYNKU PRACY



Źródło: jak przy diagramie 1.

Ten układ danych umożliwiał ich wykorzystanie jedynie jako struktury, którą należało doszacować pełniejszymi danymi dotyczącymi całego rynku pracy, które pochodzą z *Badania Aktywności Ekonomicznej Ludności* (BAEL). „Pionową” strukturę danych z tego źródła w podziale na grupy działów i działów

zgodną z wymogami transmisyjnymi dotyczącymi tablic do Eurostatu (64 agregacje) konwertowano na „pionowy” wektor o rozpiętości 14 sekcji i grup sekcji. Wektorem tym doszacowano 14 wierszy macierzy 18×14 z danymi z badania reprezentacyjnego na formularzu Z-12. Następnie macierz tę rozszacowano na macierz 18×34 , czyli na 34 agregacje wymagane w rachunku EU KLEMS w systemie NACE 2. Tak przygotowane wstępnie dane wykorzystano do przeliczeń zgodnie z metodologią rachunku KLEMS (Timmer i in., 2007a). Diagram 4 symbolicznie przedstawia te procedury.

Czynnik praca według tej metody został obliczony z uwzględnieniem samoza-trudnienia dla całkowitej liczby godzin przepracowanych oraz całkowitej wielkości wynagrodzenia tego czynnika. W metodologii EU KLEMS dopuszcza się we wszystkich rozszacowaniach stosowanie struktur całkowitej liczby pracujących lub, alternatywnie, całkowitej liczby godzin przepracowanych według 34 sektorów i 18 rodzajów pracy dla każdego przeliczanego roku (z metodologicznie dopuszczalnymi uproszczeniami do 14 grup sekcji i sekcji). W obliczeniach dla Polski wykorzystano godziny przepracowane.

Obliczenia w zakresie czynnika kapitał

Rachunek KLEMS wyróżnia następujące kategorie kapitału³³:

- 1) mieszkania,
- 2) pozostałe budowle i budynki,
- 3) sprzęt transportowy,
- 4) pozostałe maszyny i urządzenia,
- 5) sprzęt komputerowy,
- 6) urządzenia telekomunikacyjne,
- 7) aktywa kultywowane,
- 8) wartości niematerialne i prawne,
- 9) oprogramowanie komputerowe.

W praktycznej implementacji tzw. aktywa kultywowane oraz wartości niematerialne i prawne łączy się w kategorię „pozostałe aktywa”, dlatego formularze na stronie internetowej EU KLEMS zawierają tylko 8 kategorii kapitału. Mieszkania mogłyby zostać zaliczone do kategorii „kapitału produkcyjnego”, gdyby w Polsce większość majątku rezydencjalnego była w posiadaniu podmiotów zajmujących się wynajmem usług mieszkaniowych, od których odprowadzałyby podatek i które można by rejestrować jak każdą inną działalność gospodarczą³⁴. Ale tak nie jest i w rachunku KLEMS dla Polski wykorzystano 7 kategorii

³³ Określenia stosowane w rachunkach narodowych są w niektórych przypadkach inne, ale chodzi tu o odpowiedniki terminów angielskich w rachunku KLEMS, dla których określenia „maszyny biurowe i sprzęt komputerowy” oraz „urządzenia radiowe, telewizyjne i komunikacyjne” byłyby nieodpowiednie.

³⁴ W metodologii KLEMS analizowane są także możliwości powiększenia usług kapitału produkcyjnego, poprzez ich doszacowanie usługami kapitału, występującego jako obiekty mieszkalne.

kapitału a nie 8 lub 9³⁵. Ponadto w Polsce kategorie sprzęt komputerowy oraz urządzenia telekomunikacyjne nie są wydzielone z kategorii pozostałych maszyn i urządzeń, a także oprogramowanie komputerowe nie jest wydzielona z kategorii wartości niematerialnych i prawnych. Te trzy niewydzielone kategorie agreguje się w rachunku KLEMS w jedną kategorię kapitału ICT, natomiast pozostałe uwzględnione w rachunku (czyli pozostałe 4) w jedną kategorię kapitału non-ICT.

Dla czynnika kapitał podstawową operacją było zatem podjęcie próby wydzielenia tych trzech rodzajów kapitału ICT przed ich zagregowaniem we wspólną kategorię kapitału ICT. Dokonano tego na podstawie tablic podaży i wykorzystania, które zawierają pozycje w kolumnie „nakłady” dla każdej z tych trzech kategorii kapitału ICT. Pozycje te rozszacowano następnie poziomą strukturą usług związanych z oprogramowaniem w podziale na działy z tych samych tablic podaży i wykorzystania, którą wcześniej transponowano i zagregowano w 34 pionowo ułożone agregacje KLEMS, wnosząc, że w agregatach sektorowych usługi związane z oprogramowaniem są w przybliżeniu proporcjonalne do tych trzech kategorii kapitału ICT³⁶. Następnie przyjęto założenie, że skoro kapitał ICT starzeje się szybko, to nie występuje potrzeba wydzielania z istniejących szerszych agregatów kapitału starszych części kapitału ICT, z uwagi na ich niewielką wartość. W związku z tym w celu ustalenia stanu środków trwałych w tym zakresie uznano, że nie ma potrzeby uwzględniania kapitału ICT sprzed 2005 r., którego łączna wartość wynikająca z wysokich stóp jego zużycia jest dużo mniejsza niż 10% jego wartości początkowej. Począwszy od 2005 r. do 2010 r. zsumowano nakłady, które indywidualnie dla każdego roku zamortyzowano według podanych w podręczniku EU KLEMS (Timmer i in., 2007a) stóp deprecjacji o geometrycznym przebiegu.

Tak obliczone stany środków trwałych w 2010 r. dla trzech kategorii kapitału ICT wydzielono z agregatów stanów środków trwałych. Uznano, że stany środków trwałych dla kapitału ICT za lata ubiegłe są w tej samej proporcji do stanu pozostałych środków trwałych, jak w 2010 r. Lata 2011 i 2012 naliczono w prosty sposób, traktując je jako bazowe³⁷. Tablice podaży i wykorzystania są sporządzone w dwóch odrębnych klasyfikacjach PKD 2004 oraz PKD 2007 i nie będą w przyszłości przeliczane. Siłą rzeczy skorzystano tam gdzie to okazało się niezbędne z uproszczonej korespondencji pomiędzy tymi klasyfikacjami w podziale na 14 grup sekcji i sekcji stosowanej dla czynnika praca. Dzięki temu ewentualny błąd wynikający z tego uproszczenia dla tego czynnika powinien pojawiać się w tym samym miejscu, tylko działać w przeciwną stronę, co błąd

³⁵ Może się to zmienić, gdy Polska dołączy do systemu EU KLEMS.

³⁶ W 34 sektorach wielkość zakupów sprzętu komputerowego oraz tzw. *software'u* jest w przybliżeniu proporcjonalna do zapotrzebowania na usługi związane z oprogramowaniem. To założenie można rozszerzyć na sprzęt telekomunikacyjny ze względu na informatyzację tego sprzętu i jego niewielkie znaczenie w stosunku do pozostałego kapitału ICT.

³⁷ Czyli tradycyjnie poprzez odjęcie wartości amortyzacji od wartości stanu środków trwałych w roku uprzednim i dodanie bieżącej wartości nakładów na środki trwałe.

dla czynnika kapitał, dlatego te czynniki lepiej się bilansują do wartości dodanej na poziomie wybranych sektorów. Dla lat 2011 i 2012 przyjęto dane z tablic podaży i wykorzystania dla 2010 r., gdyż tablice te za 2011 r. nie były dostępne w chwili wykonywania obliczeń³⁸.

Opisane działania miały swoją konsekwencję polegającą na tym, że rezydualnie wyliczone zagregowane wynagrodzenie czynnika kapitał, jako różnica pomiędzy zagregowaną wartością dodaną brutto i zagregowanym wynagrodzeniem czynnika praca, trzeba było rozszacować na wybrane 34 sektory w EU KLEMS strukturą stanów środków trwałych, aby była jednocześnie możliwość rozdzielania wynagrodzenia kapitału na jego 7 kategorii środków trwałych.

Oczekiwanym problemem w rachunku KLEMS było przejście z systemu ESA 95 na ESA 2010, gdyż nie wszystkie dane zostały przeliczone z jednego systemu na drugi. Są także dane, które z założenia nigdy nie będą przeliczane (np. tablice podaży i wykorzystania sprzed 2010 r.). Stąd, choć rzadko (dotyczy to tylko problemu rozszacowania usług kapitału stanami środków w celu wydzielenia kapitału ICT), występuje niekiedy konieczność mieszanego wykorzystania danych. Aby sprawdzić, czy jest to dopuszczalne, przeprowadzono analizę błędów poprzez porównanie różnic przyrostów stanu środków trwałych. Wyniki tej analizy wskazują, że błędy, które mogą wyniknąć z mieszanego wykorzystania danych pochodzących z oby systemów, tj. ESA 95 i ESA 2010, są zaniebawiane z punktu widzenia potrzeb rachunku produktywności KLEMS.

Choć przeliczenie danych w celu wydzielenia kapitału ICT jest możliwe, to kontrowersyjna okazała się sama idea rozdzielania w rachunku produktywności KLEMS usług kapitału na kategorie ICT i non-ICT. W Stanach Zjednoczonych, Szwecji i Rosji nie dokonuje się tego rozdzielania. Znaczenie wkładu kapitału ICT w przyrost wartości dodanej brutto jest niewielkie i stale maleje w okresie od 2005 r. prawie dla wszystkich krajów EU KLEMS, przy czym Polska jest krajem o szczególnie małym jego znaczeniu.

Uwagi końcowe

Wstępnie przygotowane dane przetworzono zgodnie z metodologią EU KLEMS, z tym że dla porównania zastosowano cztery alternatywne warianty przeliczenia. Wynikają one z dwóch dychotomii:

- a) przeliczanie danych na poziomie wszystkich wybranych w rachunku EU KLEMS agregacji lub alternatywnie przeliczanie danych tylko na najniższych wybranych agregacjach i końcowe sumowanie ich do wyższych;
- b) stosowanie dwóch zapisów matematycznych dla przyrostów wartości.

Pierwsza z tych dychotomii polega na tym, że można przeliczyć wszystko dla założonych agregatów, czyli 34 wybranych sektorów EU KLEMS, sześciu czę-

³⁸ Stosowanie „starych” tablic podaży i wykorzystania w sytuacji braku aktualnych jest częstą praktyką statystyczną. W chwili składania tego artykułu obliczenia zostały uzupełnione do roku 2014, a tablice podaży i wykorzystania za 2011 r. stały się już dostępne.

ściowych subagregatów EU KLEMS oraz jednego agregatu dla całej gospodarki, czyli naliczyć wszystko odrębnie dla 41 pozycji w pełnym pionowym wektorze agregacji KLEMS. Można również przeliczyć wszystko dla 34 sektorów EU KLEMS i tak obliczone dane końcowe zsumować do wyższych agregacji. Pierwszą metodę należy uznać za generalnie prostszą w odniesieniu do wyższych agregacji, drugą zaś za generalnie bardziej złożoną, przy czym technik pośrednich przeliczania jest wiele (przy ich wyborze kierowano się zarówno argumentacją merytoryczną, jak i dostępnością danych GUS). Drugą dychotomią jest to, że można przyrosty wartości liczyć jako $\Delta x/x$ lub $\Delta \ln x$. W celu porównania wykonano rachunki według wszystkich tych sposobów, choć teoria EU KLEMS zaleca stosowanie wyrażenia logarytmicznego, co jest związane m.in. z wymaganą jako preferencyjną procedurą Tornqvista przy agregacji. Szezej na ten temat napisali Schreyer (2004) i Milana (2009).

Pomimo bardzo dużej liczby operacji różnice w wynikach pomiędzy różnymi wybranymi metodami przeliczania danych okazały się niewielkie.

dr Dariusz Kotlewski — Uniwersytet Warszawski — Wydział Geografii i Studiów Regionalnych; GUS
mgr Mirosław Błażej — GUS

LITERATURA

- Arnaud, B., Dupont, J., Seung-Hee, K., Schreyer, P. (2011). *Measuring Multi-Factor Productivity by Industry: Methodology and First Results from the OECD Productivity Database*. OECD.
- Gouma, R., Timmer, M. (2013a). *EU KLEMS Growth and Productivity Accounts — 2013 update*. Groningen Growth and Development Centre.
- Gouma, R., Timmer, M. (2013b). *WORLD KLEMS Growth and Productivity Accounts — 2013 update*. Groningen Growth and Development Centre.
- Hulten, C. R. (2009). *Growth Accounting*. NBER Working Paper Series 15341.
- Jorgenson, D. W. (1963). Capital Theory and Investment Behavior. *American Economic Review*, no. 53(2), s. 247—259.
- Jorgenson, D. W. (1989). *Productivity and Economic Growth*. W: Berndt E.R., Triplett J.E. (eds.), *Fifty Years of Economic Measurement*. University of Chicago Press.
- Jorgenson, D. W., Gollop, F. M., Fraumeni, B. M. (1987). *Productivity and US Economic Growth*. Cambridge MA: Harvard University Press.
- Jorgenson, D. W., Griliches Z. (1967). The explanation of Productivity Change. *Review of Economic Studies*, no. 34, s. 249—283.
- Jorgenson, D. W., Ho M., Stiroh, K. (2005). *Information Technology and the American Growth Resurgence*. MIT.
- Kang L., Peng, F. (2013). *Growth Accounting in China 1978—2009*. MPRA Paper, no. 50827.
- Milana, C. (2009). *Solving the Index-Number Problem in a Historical perspective*. EU KLEMS Working Paper Series 43.
- OECD (2001). *Measuring Productivity*. OECD Manual.
- OECD (2009). *Measuring Capital*. OECD Manual.
- OECD (2013). *OECD Compendium of Productivity Indicators 2013*. OECD Publishing.
- Schreyer, P. (2004). *Chain Index Number Formulae in the National Accounts*. 8th OECD — NBS Workshop on National Accounts.

- Solow, R. M. (1956). A Contribution to the Theory of Economic Growth. *The Quarterly Journal of Economics*, no. 70, vol. 1, s. 65—70.
- Timmer, M., van Moergastel, T., Stuivenwold, E., Ypma, G., O'Mahony, M., Kangasniemi, M. (2007a). *EU KLEMS Growth and Productivity Accounts — Metodology*. EU KLEMS Consortium.
- Timmer, M., van Moergastel, T., Stuivenwold, E., Ypma, G., O'Mahony, M., Kangasniemi, M. (2007b). *EU KLEMS Growth and Productivity Accounts — Sources by country*. EU KLEMS Consortium.

Summary. *The article presents the origins of the KLEMS Growth and Productivity Accounts, that is descendant from the Solow's well known economic growth decomposition. They discuss further the methodological details of the Accounts, present the international context and the implementation aspects of the KLEMS Growth and Productivity Accounts in Poland, that is actually being performed in the Central Statistical Office of Poland (CSO). They focus on the specific data environment for this productivity accounts in Poland. Specific techniques of data pre-calculation are presented for the two primary factors — labour and capital.*

Keywords: KLEMS Growth and Productivity Accounts, production factors, the primary factors, labor factor, capital factor, productivity increase, decomposition, work composition, worked hours, employees, hours per worker.

Резюме. *В статье рассматривается происхождение и состояние счета производительности KLEMS. Были сообщены методологические детали счета, был представлен также его международный контекст и аспекты использования в польских условиях. Было отмечено, что учетная запись проводится в ЦСУ. Авторы представили доступность данных, используемых для счета в Польше, а также представили метод включения в эти счета труда и капитала.*

Ключевые слова: счет производительности, KLEMS, факторы производства, основные факторы, фактор труда, фактор капитала, рост производительности труда, разложение, состав труда, количество отработанных часов, сотрудники, часы на одного сотрудника.

EDUKACJA STATYSTYCZNA

Czesław DOMAŃSKI

Wyzwania w zakresie edukacji statystycznej społeczeństwa¹

Kto mówi prawdę, w prawdzie ma obronę
Sofokles

Streszczenie. *W artykule omówiono najnowsze wyzwania badawcze i dydaktyczne związane z edukacją statystyczną społeczeństwa. Rewolucja technologiczna powoduje radykalne zmiany życia gospodarczego i społecznego. Wymagają one przekształcenia sposobu myślenia i działania. Społeczeństwu potrzebna jest statystyka, która dostarcza informacji o kraju i jego sąsiadach, ponieważ dzięki niej obywatele mogą aktywnie uczestniczyć w procesach demokratycznych. Te nowe warunki określają zadania stawiane przed statystykami — rozwijanie umiejętności statystycznych oraz dostosowanie przekazu informacji statystycznej do potrzeb obywateli.*

Słowa kluczowe: edukacja statystyczna, rewolucja technologiczna, umiejętności statystyczne.

Możemy wymienić trzy czynniki kształtujące rolę statystyki i kierunki jej rozwoju na początku XXI w. Pierwszym z nich jest wprowadzenie modelu otwartej i konkurencyjnej gospodarki rynkowej. Drugim czynnikiem jest rewolucja technologiczna, która nastąpiła w sektorze informacji i która radykalnie zmieniła społec-

¹ Artykuł opracowano na podstawie referatu pt. *Istota edukacji statystycznej społeczeństwa* przedstawionego podczas ogólnopolskiej konferencji naukowej z okazji obchodów Dnia Statystyki, Gdynia 17 i 18 marca 2016 r.

czeństwo. Trzeci czynnik częściowo wynika z dwóch poprzednich — jest to powstanie globalnej wioski gospodarczej. Proces ten jest konieczny w gospodarce rynkowej i jest możliwy właśnie dzięki rewolucji technologicznej.

Wymienione czynniki wywołują znaczące zmiany w życiu politycznym, gospodarczym i społecznym oraz przyczyniają się do zacieśnienia współzależności między państwami, w których lokalne władze z coraz większym trudem podejmują działania całkowicie niezależne od współpartnerów. Jednocześnie mamy do czynienia z radykalną zmianą życia gospodarczego i społecznego, wymagającą powszechnej zmiany stylu myślenia i działania. Konieczne staje się informowanie, wyjaśnianie i kształcenie ludzi, by mogli zrozumieć skalę zmian i aktywnie w nich uczestniczyć.

Oficjalna statystyka musi umieć przewidywać kierunki ewolucji społeczeństwa, określać warunki podejmowania decyzji poprzez dostarczanie decydentom wiarygodnych i aktualnych informacji o sytuacji bieżącej i jej zmianach we wszystkich sektorach ekonomicznych i grupach społecznych.

Spółeczeństwu potrzebna jest statystyka, która dostarcza informacji o ojczyźnie i jej sąsiadach, ponieważ dzięki niej ludzie mogą aktywnie uczestniczyć w procesach demokratycznych. W działalności gospodarczej statystyka umożliwia określenie własnej pozycji ekonomicznej oraz pomaga wyznaczyć przyszłą strategię.

Doświadczenia wskazują, że gospodarka rynkowa nie może sama dostarczać danych statystycznych potrzebnych do jej działania. Oficjalna statystyka stanowi więc rodzaj infrastruktury gospodarki rynkowej, niemniej jednak uzyskiwanie danych jest procesem długotrwałym i kosztownym, w którym podmioty gospodarki rynkowej nie chcą partycypować. W tej sytuacji organizacja i utrzymanie systemu statystyki spoczywają na barkach władz.

Okresom szybkich zmian w społeczeństwie towarzyszy rosnąca niepewność. W efekcie decyzji, przedsiębiorcy i obywatele potrzebują coraz więcej danych statystycznych, których otrzymywanie jest coraz trudniejsze.

W okresie transformacji systemowej instytucje statystyczne muszą przekształcić się z administratora w dostawcę informacji statystycznych. Jest to ogromne wyzwanie z punktu widzenia rozwiązań prawnych i zachowania wiarygodności. Wymaga ono:

- zmiany mentalności pracowników statystyki, szczególnie ich roli i motywacji,
- zmiany systemu przekazu informacji, szczególnie w odniesieniu do zasad i terminologii dotyczących procesów ekonomicznych oraz społecznych,
- wprowadzenia nowych zasad prawnych, organizacyjnych oraz metodologicznych, w szczególności zasad etyki, zachowania poufności informacji, organizacji gromadzenia i udostępniania danych.

Trendy w rodzajach wyzwań stających przed statystykami wskazują perspektywy rozwoju statystyki, także strategii Polskiego Towarzystwa Statystycznego:

- istotne jest rozwiązywanie małych problemów. W przemyśle, handlu i usługach naciski konkurencji są tak silne, okresy sprawozdawcze tak krótkie, zaś

- zbieranie danych tak kosztowne, że przygotowanie planów i decyzji musi opierać się na coraz mniejszej ilości danych,
- faktyczne wyzwanie wiąże się z rozwiązywaniem coraz większych problemów. Zarówno statystycy, jak i dobrze wykształceni przedstawiciele innych dyscyplin, posiadający zdolności do kreowania racjonalnych zmian, mają tendencję do stymulowania tego kierunku,
 - w wielu dziedzinach można obserwować nadmierne ilości informacji, np. gigabajty czy terabajty danych.

SPOŁECZNY WYMIAR STATYSTYKI

Obecny etap rozwoju społeczeństw charakteryzuje się gwałtownym rozpowszechnieniem technologii informacyjnych, które obejmuje informatyka statystyczna (Domański i Jędrzejczak, 2015). Rozwój gospodarczy zależy w decydującym stopniu od umiejętności zarządzania informacją oraz od umiejętności wykorzystania zdobyczy nowoczesnych technologii. Podstawową rolę w zbieraniu, analizach, interpretacji i udostępnianiu informacji odgrywa statystyka. Fakt ten nie jest jednak w wystarczającym stopniu doceniany ani przez sfery rządzące, ani przez prowadzących działalność gospodarczą. Nie docenia się również statystyki jako przedmiotu studiów i badań naukowych zarówno ze względu na jej niedostateczną atrakcyjność dla studentów, jak i ograniczone możliwości kariery zawodowej dla absolwentów.

W wieloletnim procesie rozwoju historycznego statystyka np. w Kanadzie wypracowała wiele narzędzi, które można wykorzystać do rozwiązywania licznych zadań związanych z doskonaleniem działalności przedsiębiorstw, wiele z tych narzędzi w praktyce się jednak nie wykorzystuje. W przedsiębiorstwach statystyka powinna wspomagać kierownictwo w realizacji zadań na poziomie strategicznym i operacyjnym. Na poziomie strategicznym najważniejszą rolę odgrywa myślenie statystyczne polegające na: umiejętności kojarzenia różnych zjawisk, podejmowaniu decyzji opartych na podstawie informacji, zrozumieniu pojęcia zmienności oraz systematyczności w działaniu.

Ciekawego przeglądu definicji myślenia statystycznego dokonał Mallows, mówiąc: *Myślenie statystyczne dotyczy relacji danych ilościowych do problemu świata rzeczywistego, któremu często towarzyszy zmienność i niepewność. Próbuje ono dokonać precyzyjnego i jasnego wyjaśnienia, co dane mają do powiedzenia odnośnie interesującego nas problemu.* Do tej definicji zgłaszanych jest wiele zastrzeżeń, których nie będziemy omawiali. Wydaje mi się, że bardziej odpowiednia jest ta podana przez Amerykańskie Stowarzyszenie Jakości (Kordos, 1999, 2001). Dotyczy ona rozpoznawania otaczającego nas świata, prób jego zrozumienia oraz działań mających na celu podejmowanie decyzji. Badania i metody statystyczne odgrywają tu podstawową, ale nie jedyną, rolę. Ogólnie można to wyrazić następująco — myślenie statystyczne

jest filozofią uczenia się i działania opartego na trzech fundamentalnych zasadach:

- działanie odbywa się w systemie procesów współzależnych,
- zmienność i niepewność występują we wszystkich procesach,
- zrozumienie zmienności i zredukowanie niepewności jest kluczem do sukcesu w podejmowaniu decyzji i zrozumieniu otaczającego nas świata.

Statystyka i informatyka statystyczna jest obecna dzisiaj w niemal wszystkich dziedzinach działalności człowieka i społeczeństwa. Obecnie państwa nie mogłyby funkcjonować bez stosowania metod statystycznych, ale jednocześnie wiedza statystyczna pozostaje niemal całkowicie poza świadomością znacznej części obywateli lub postrzegana jest jako specjalistyczna dyscyplina matematyczna zastrzeżona dla garstki uczonych bujających w obłokach. Równocześnie media przekazują niemal codziennie wiele informacji o różnych przejawach aktywności człowieka: o wynikach działalności gospodarczej, sytuacji na giełdach, bezpieczeństwie w ruchu drogowym i powietrznym.

Informacje te nie zawsze znajdują zrozumienie u szerokiego grona odbiorców. Świadczy to o potrzebie zasadniczego przemyślenia roli statystyki, zarówno w szkolnictwie, biznesie, administracji publicznej jak i w społeczeństwie. Zadanie to staje się szczególnie aktualne na obecnym etapie rozwoju społeczeństw.

ROZWIJANIE UMIEJĘTNOŚCI STATYSTYCZNYCH

W ostatnich latach zostały podjęte próby w celu zdefiniowania i rozróżnienia konstrukcji powiązanych i związanych ze zdolnością ludzi z różnych środowisk do działania jako świadomi odbiorcy informacji statystycznej (Steen, 2001).

Autorzy są zgodni co do potrzeby prowadzenia zajęć z przedmiotu statystyka na wszystkich kierunkach, w celu rozwijania zdolności rozumienia, interpretacji i krytycznej oceny wiadomości z elementami statystycznymi lub argumentów przekazywanych przez media i inne źródła. Ta umiejętność nazywana jest „alfabetyzacją statystyczną” lub „umiejętnością statystyczną” (Wallman, 1993).

Według Gala (2002) zrozumienie, interpretacja i reakcja na codzienne i rzeczywiste wiadomości, które zawierają informacje i dane statystyczne wymaga znacznie więcej umiejętności niż posiadanie wiedzy statystycznej. Takie działania są oparte na wzajemnym oddziaływaniu na siebie kilku dziedzin wiedzy. Umiejętności te muszą być uaktywnione wspólnie: wiedza statystyczna, wiedza matematyczna oraz łącznie ogólna wiedza o świecie.

Krytyczna ocena wiadomości statystycznych zależy również od umiejętności zadawania odpowiednich pytań i samego usposobienia, charakteru i skłonności osób np. do przyjęcia stanowiska krytycznego.

Analiza najnowszej literatury dotyczącej edukacji statystycznej pokazuje dwie opcje rozwoju zdolności do działania w sposób statystyczny — poprzez transfer generalnej wiedzy statystycznej lub bezpośrednią praktykę w krytycznych pyta-

niach i komunikatach modelowych. Umiejętności statystyczne są traktowane jako kluczowy cel kształcenia.

Biorąc pod uwagę ograniczenia każdej z opisanych metod oraz nieuwagę i niedbałość związaną z oceną umiejętności czytania informacji statystycznych, wykładowcy i statystycy powinni ponownie przeanalizować swoje oczekiwania w zakresie kształcenia w stosunku do efektywności oraz wydajności obecnych działań na rzecz opanowania umiejętności statystycznych przez studentów. Wydają się również zasadne rozważania Moore'a (2001) odnośnie „zmniejszenia oczekiwań”, jako propozycja — przesłanie dla nauczycieli podstaw statystyki, aby nie nauczać o „wszystkich tematach”. Nauczanie powinno uwzględniać odpowiednie tempo i jakość, pozwalając studentom konstruować własne spostrzeżenia i zrozumienie.

Nauczyciele akademicki nauczający statystyki powinni znaleźć sposób, aby:

- a) zajmować się podstawowymi zagadnieniami i problemami alfabetyzacji statystycznej w tym samym czasie oraz
- b) oceniać rzeczywiste postępy studentów w zdobywaniu umiejętności wykorzystywania zagadnień statystycznych.

PRZEKAZ INFORMACJI STATYSTYCZNEJ W SPOŁECZEŃSTWIE

Dążenie do rozwinięcia umiejętności statystycznych wszystkich obywateli może obrać różną drogę, w zależności od grupy docelowej i samych zainteresowanych. Jednocześnie istnieją tysiące informacji, wydawnictw prasowych, streszczeń i podsumowań, spośród których nauczyciele mogą wybierać. Istnieje jednak potrzeba wyeksponowania tych szczególnie wartościowych na zajęciach, w celu skutecznej promocji statystyki i umiejętności statystycznych. Ponadto występuje konieczność opracowania wytycznych i sugestii, aby umiejętnie wybierać materiały do odpowiednich etapów nauczania czy tematów.

Potrzeba selekcji i przygotowania materiałów i sugestii nie może spoczywać jedynie na barkach nauczycieli, którzy są uwikłani w starania o poprawę nauczania podstawowych tematów statystycznych lub są ograniczeni logistycznie. Potrzebny jest w tym wypadku udział i zaangażowanie zainteresowanych organizacji, w tym przede wszystkim PTS na łamach czasopism „Wiadomości Statystyczne” i „Kwartalnik Statystyczny” (Stefanowicz, 1999, 2001).

Społeczność statystyczna poszukuje sposobów na ulepszenie przekazywania wiedzy statystycznej. Ważne jest przy tym, że żadne studia nie testowały skuteczności różnych metod nauczania statystyki. Ponadto nie istnieją modele oceny umiejętności statystycznych w grupach studenckich. Nie jest to zaskakujące, że funkcjonuje mnóstwo różnorodnych i możliwych bodźców, które nauczyciele mogą wykorzystywać w stosunku do reakcji studentów na ich poznawcze podstawy wiedzy i umiejętności, jak również na ich przekonania, postawy i skłonność do działania.

Dane i informacje są generowane w zatrważającym tempie: dane satelitarne, dane o ochronie zdrowia w tysiącu form, statystyki rządowe, dane środowiskowe, meteorologiczne, informacje genetyczne, dane eksperymentalne generowane przez

komputery itp. Interpretowanie niektórych z nich będzie obejmowało ponowne rozwijanie znanych technik, wymuszone przez sam rozmiar danego problemu. Mówiąc ogólniej, problem polega na tym, że nie wiemy, jak sporządzać dane lub nawet jak je zdobyć, nie mówiąc o tym, jak je interpretować (Szreder, 2016).

Gdy będziemy starali się sprostać tym wyzwaniom, to rzeczą podstawową jest, by pamiętać, dlaczego to robimy. Jeśli o to chodzi, to mamy klienta, którym jest socjolog, biolog, inżynier, agencja rządowa, firma farmaceutyczna, nauczyciel lub inny niestatystyk. Jesteśmy odpowiedzialni za pielęgnowanie użytecznej i właściwej analizy statystycznej po to, by rzucić światło na ważne problemy naukowe. Mowa statystyka nie powinna być dalej niejasnym, tajemniczym symbolem matematycznym i skróconymi zawiłościami potoku słów zawierającymi przeskoki przy wnioskowaniu.

Potrzebujemy całkowicie nowego podejścia, takiego które wpaja studentowi docenienie statystyki, tak jak poezji, sztuki, muzyki i filozofii. Odnosi się to zwłaszcza do tych, dla których statystyka jest przedmiotem kierunkowym. Bardziej niż wprowadzenia do kilku specyficznych metod potrzebujemy „statystyki w społeczeństwie”. Kierunek pamiętany i bardzo dobrze oceniany przez studentów będzie miał o wiele bardziej trwały i pozytywny wpływ na to jak postrzegany jest nasz przedmiot. Jesteśmy jedną z niewielu wyspecjalizowanych dyscyplin mających przywilej edukowania poprzez program studiów przyszłych przywódców narodu i pracodawców.

Musimy zapytać, czy odpowiadamy na potrzeby społeczeństwa?

Pytamy, czy jesteśmy maszynami do liczenia lub czy umiemy liczyć? Jednak pytanie nie powinno brzmieć, czy umiemy liczyć, tylko raczej, czy liczymy? Czy ludzie służący społeczeństwu, urzędnicy rządowi, menadżerowie przemysłowi, naukowcy różnych profesji natychmiast przyswajają metodę statystyczną jako integralną nić w tkaninie ich własnych przedsięwzięć czy też izolują się?

Tak więc musimy wyteńczyć nasze zdolności i energię intelektualną, od nowa ukierunkować naszą świadomość, gdy poszukujemy swego przeznaczenia i przyszłości, która będzie produktywna i pasująca do niechybnie czekających nas problemów społecznych. Należy znaleźć nowe wzorce nie tylko dla naszego przedmiotu, ale również dla naszych związków z rządem, przemysłem, światem nauki i współpracy z innymi dyscyplinami. Do nas należy wybór, czy zharmonizujemy się, jak w pięknej symfonii czy też będziemy monotonicznie brzęczeć. Nasza przeszłość jest rozświetlana dokonaniem wielkich ludzi i wspaniałymi osiągnięciami. Teraźniejszość i przyszłość są bardzo obiecujące. Jest to ekscytujący okres życia. Do nas, statystyków, należy wytyczenie kursu, który skupiałby się na unikalnych możliwościach zawartych w statystyce i nieograniczonych okazjach odegrania kluczowej i niezbędnej roli w rozwiązywaniu współczesnych problemów. Kursu, który gwarantowałby sukces naszej profesji. Czy liczymy? Lubimy wierzyć, że tak. Pytanie podstawowe to, czy inni myślą, że my liczymy? Odpowiedź na to pytanie i nasze podejście do niej będzie kształtowało naszą przyszłość.

LITERATURA

- Domański, Cz., Jędrzejczak, A. (2015). Statistical computing in information society. *Folia Oeconomica Stetinensia*, De Gruyter Open, s. 145—152.
- Gal, I. (2002). Adult Statistical Literacy: Meanings, Components, Responsibilities. *International Statistical Review*, vol. 70, s. 1—25.
- Gal, I. (2003). Teaching for Statistical Literacy and Services of Statistics Agencies. *The American Statistician*, vol. 57, no. 2.
- Kordos, J. (1999). Prace badawcze w zakresie edukacji statystycznej. *Kwartalnik Statystyczny*, nr 4, s. 45—47.
- Kordos, J. (2001). Czterowymiarowa struktura myślenia statystycznego w badaniu empirycznym. *Kwartalnik Statystyczny*, R. III, nr 1, s. 55—57.
- Moore, D. (Ed) (2001). *Les représentations des langues et de leur apprentissage. Références, modèles, données et méthode*. Paris, Coolection Crédif — Essais, Didier.
- Steen, L.A. (2001). *Mathematics and Democracy: The Case for Quantitative Literacy*. USA, Washington, DC, National Council on Education and the Disciplines.
- Stefanowicz, B. (1999). Funkcja informacji statystycznej. *Kwartalnik Statystyczny*, R. III, nr 1, s. 25 i 26.
- Stefanowicz, B. (2001). Edukacja statystyczna. *Kwartalnik Statystyczny*, R. III, nr 1, s. 57 i 58.
- Szreder, M. (2016). O niektórych nowych wyzwaniach i oczekiwaniach wobec statystyki. *Wiadomości Statystyczne*, nr 6, s. 1—9.
- Wallman, K.K. (1993). Enhancing Statistical Literacy: Enriching Our Society. *Journal of the American Statistical Association*, vol. 88, s. 1—8.

Summary. *The article discusses the latest challenges of research and teaching related to statistical education of the society. The technological revolution change the economic and social life radically. They require a transformation of the way of thinking and acting. Society need statistics that provides information about the country and its neighbors. Through statistics citizens can actively participate in democratic processes. These new conditions define the tasks set for the statisticians — to develop statistics skills and adapt statistical information to the needs of citizens.*

Keywords: statistical education, technological revolution, statistics skills.

Резюме. *В статье рассматриваются последние проблемы исследования и преподавания в области образования статистического общества. Научно-техническая революция влияет на изменения экономической и социальной жизни. Изменения требуют трансформации мышления и действия. Общество нуждается в статистике, которая предоставляет информации о стране и ее соседях. Благодаря ее граждане могут активно участвовать в демократических процессах. Новые условия определяют задачи, поставленные статистике — развивать статистические знания и регулировать поток статистических информации согласно потребностям граждан.*

Ключевые слова: статистическое обучение, научно-техническая революция, статистические знания.

INFORMACJE. PRZEGLĄDY. RECENZJE

Marek Hetmański: *Świat informacji*,
230 stron, Wydawnictwo Difin, Warszawa 2015

Jest to tytuł książki profesora Marka Hetmańskiego, filozofa i badacza teorii komunikacji, informacji i wiedzy. Jej treść, w dużym skrócie, można ująć jako poszukiwanie odpowiedzi na pytania: co oznacza (opisuje) termin *informacja* oraz czy termin ten trafnie określa to, co jesteśmy skłonni uznać za *informację*. Pytanie pierwsze zmierza do ustalenia treści, jaką można przypisać temu terminowi, zaś pytanie drugie wiąże się z poszukiwaniem terminu odpowiadającego treści, którą wiążemy z pojęciem zwanym informacją. Oba pytania wzajemnie się uzupełniają.

Autor rozwija analizę w dużej mierze na podstawie podejścia filozoficznego z wykorzystaniem znajomości problematyki komunikacji (w znaczeniu powiadamiania, przekazywania wiadomości). Prof. Hetmański rozpoczyna swoją książkę od analizy etymologicznej tytułowego terminu. Odwołując się do średniowiecznej łaciny i pochodzącego z niej słowa *informatio*, tłumaczenie opiera na słowniku łacińsko-polskim i przedstawia ewolucję pojęciową tego słowa. Autor analizuje różne formy gramatyczne i wskazuje kilka znaczeń: *odzworowywanie*, *odkształcanie* (kształtowanie). Podkreśla, że współcześnie termin ten występuje w wielu dziedzinach: w fizyce, informatyce i psychologii. Przedstawia też historię ewolucji zakresu znaczeniowego *informacji* — od etapu inżyniersko-fizycznego i cybernetycznego po okres zastosowań w naukach przyrodniczych i społecznych. Ten proces trwa i trwają uściślenia oraz uogólnienia znaczenia tego terminu w celu budowania *filozofii informacji*. M. Hetmański dodaje, że (s. 181) *informacja nie ma jednej ojczyzny (dziedziny — B. S.), lecz wiele krajów (dziedzin — B. S.), w których zamieszkuje*. Zaznacza, że pojęcie informacji zajmuje podobną pozycję, jak pojęcia masy i energii w fizyce. Zarzuca przy tym, że nie jest ona (...) *przedmiotem refleksji ze strony specjalistów od nauk, które mają go (tzn. termin informacja) w nazwie lub się nim bezpośrednio zajmują*. Dodaje (s. 13), że *bez precyzji terminologicznej oraz kategorialno-pojęciowych analiz nie można zrozumieć w pełni zjawisk informacyjnych*.

Tej tematyce poświęca rozdział 1.

Interesujący jest rozdział 2, w którym znajdujemy bogaty opis pojęcia informacji w metaforach. Autor rozwija ten wątek jako skuteczny sposób ujmowania wie-

lu złożonych zjawisk i pojęć. Prof. Hetmański nazywa je metaforami informacyjnymi. Rozwija ich analizę przede wszystkim w kontekście komunikowania się — przesyłania informacji za pośrednictwem *przewodu*. M. Hetmański przywołuje metaforyczne zwroty Claud’a E. Shannona (1948) — autora matematycznej teorii komunikacji (uznawanej obecnie za teorię informacji), w której można znaleźć opisy zjawisk komunikacyjnych w sposób metaforyczny. Przykładem jest porównanie łączności z telegrafistką (metafora *teoria komunikacji/informacji jest jak telegrafistka* — nie interesuje ją treść telegramu, lecz jak go sprawnie nadać). Hetmański przytacza wiele metafor, a wśród nich dwie autorstwa Stanisława Lema:

- *Informacja jako bomba megabitowa*. Lem w ten sposób opisywał zjawiska informacyjne w skali globalnej, obejmującej nie tylko cywilizację ludzką, lecz cały Wszechświat, łącznie ze światami możliwymi. Metafora podkreśla, że przyrost informacji i wiedzy, rozpatrywany w skali trwania Wszechświata, podlega uniwersalnej prawidłowości — swoistej eksplozji. Metafora bomby podkreśla dynamiczny charakter informacji i zjawisk informacyjnych — pozytywnych i negatywnych;
- *Informacja jako odpadki*. Stanisław Lem w ten sposób charakteryzował nadmiar informacji krążących w Internecie. Podkreślał „zaśmiecenie” środowiska informacyjnego ogromną ilością głupstw i kłamstw, bylejakości i szkodliwości.

Prof. Hetmański podkreśla, że ze względu na swobodną grę skojarzeń i oryginalnych porównań informacji z bardzo odległymi zjawiskami i obiektami, metafory informacyjne uzyskują przewagę nad tradycyjnymi definicjami i opisami, stają się kluczem do budowania bogatych sieci skojarzeń między informacją i występującymi w metaforach pojęciami i terminami.

Autor formułuje interesującą tezę w sprawie wartości poznawczej metafor (s. 78): *Im zjawisko, proces czy zdarzenie, o którym mówi dana metafora, jest mniej prawdopodobne, tym większą wartość poznawczą ma ona wówczas, można ją określić jako informacyjną zawartość (moc) metafory*. Stwierdzenie to wynika z kluczowej tezy Shannona: *Im mniejsze jest prawdopodobieństwo pojawienia się jakiegoś komunikatu, tym więcej informacji wnosi jego opublikowanie*.

Kolejny temat to relacja między pojęciami *informacja* i *komunikacja* (rozdział 3). Autor podkreśla, że powiązania między nimi są oczywiste i naturalne. Analizę swoją przeprowadza na podstawie teorii Shannona o przesyłaniu sygnałów w systemach komunikacyjnych, w której wprowadzono pojęcie entropii jako miary nieuporządkowania zbiorowości komunikatów oraz informacji jako miary porządku w tej zbiorowości. Tę analizę prof. M. Hetmański uzupełnił przeglądem prac innych autorów, w szczególności Norberta Wienera — uznawanego za twórcę teoretycznych podstaw cybernetyki i autora koncepcji sprzężenia zwrotnego oraz Umberto Eco — włoskiego filozofa, bibliofila i powieściopisarza, także badacza w zakresie semiotyki.

Kontynuując podjętą analizę, Hetmański bada skutki dynamiki procesów informacyjnych i komunikowania się we współczesnym społeczeństwie — pojawienie się tzw. *zwrotu cywilizacyjnego*, jaki daje się obecnie zaobserwować. Tematowi temu poświęca rozdział 4 — „Cywilizacja informacji”. Autor zaznacza, że terminu

zwrot używa się zazwyczaj do określenia zjawisk mających zasadniczy, radykalny i nieodwracalny wpływ na kształtowanie się zjawisk i procesów we wszystkich sferach życia danej zbiorowości i jednostek, które tę zbiorowość tworzą. Warunki te spełniają przemiany społeczne, gospodarcze i kulturowe zachodzące współcześnie pod wpływem informacji i technologii informatycznych. Daje to podstawę do formułowania opinii w sprawie *zwrotu informacyjnego o niespotykanej dotychczas teoretycznej i praktycznej mocy oraz oddziaływaniu daleko głębszym i rozleglej- szym niż inne wcześniejsze zwroty cywilizacyjne* — pisze autor (s. 154).

Z tą tematyką wiążą się także inne pojęcia: *postawa informacyjna* jako nieod- łączny wyznacznik kultury informacyjnej, *światopogląd informacyjny* oraz *in- formatologia* jako nauka o informacji, kształtujące się pod wpływem swoistego pola informacyjnego podczas intensywnej wymiany informacji i procesów in- formacyjnych w społeczeństwie.

Czytelnik znajdzie w książce Marka Hetmańskiego dużo oryginalnych myśli o informacji i jej miejscu w naszej rzeczywistości. Główny nurt rozważań wy- wodzi się z teorii C. E. Shannona, w której istotą jest analiza „upakowania” sy- gnałów w kanale komunikacyjnym. James Gleick (2012), autor obszernej mono- grafii na temat informacji, tak określił główny wątek badań Schannona: *jak zmaksymalizować przepływ przez rozgałęzioną sieć*.

Lektura książki może sprawić wrażenie, że świat informacji według Shannona jest dość ubogi i nie obejmuje wielu dziedzin (i problemów), z jakimi stykamy się w codzienności wynikającej z zanurzenia się w przestrzeni informacyjnej. Sam autor dostrzegał w swej koncepcji pewne ograniczenie: *Taka interpretacja pojęcia informacji ogranicza je do procesów przesyłania wiadomości bez uwzględnienia treści — jedynie ze względu na jej „upakowanie” w kanale, w którym jest przesyłana*.

Warto przytoczyć jeszcze opinię Jamesa Gleicka (2012, s. 386): *Narodziny teorii informacji były związane z bezlitosnym poświęceniem znaczenia — tej cechy, która nadaje informacji wartość i celowość*. We wprowadzeniu do *A Mathematical Theory of Communication* Shannon musiał być szczery. Oświadczył po prostu, że *znaczenie nie jest powiązane z problemem inżyneryjnym. Zapomnijmy o psychologii; odrzućmy subiektywizm*. Nieco dalej Gleick dodaje: *Dla niektórych było to zbyt bezduszne, zimne*. Na jednej z pierwszych konferencji cybernetyków Heinz von Foerster (austriacko-amerykański filozof, cybernetyk — B. S.) uważał się, że teoria informacji dotyczy tylko *piskliwych sygnałów*, i stwierdził, iż dopiero gdy dochodzi do ich zrozumienia w ludzkim mózgu, *wtedy rodzi się informacja — a nie w piskach*.

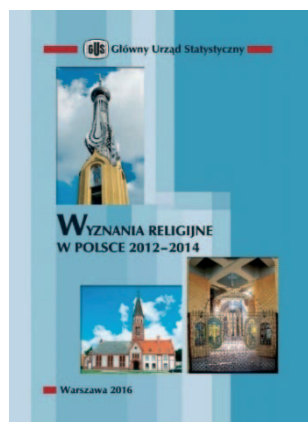
Oprac. prof. dr hab. Bogdan Stefanowicz

LITERATURA

- Gleick, J. (2012). *Informacja: Bit, wszechświat, rewolucja*. Tłumaczył Grzegorz Siwek. Kraków: Wydawnictwo Znak.
- Hetmański, M. (2015). *Świat informacji*. Warszawa: Wydawnictwo Difin.
- Shannon, C. E. (1948). A Mathematical Theory of Communication. Bell System. *Technical Journal*, vol. 27, no. 3—4.

Wydawnictwa GUS — sierpień 2016 r.

Z sierpniowej oferty wydawniczej GUS warto zwrócić uwagę na publikację cykliczną **„Wyznania religijne w Polsce 2012—2014”** oraz zeszyt metodologiczny **„Zasady metodyczne badań statystycznych z zakresu energii ze źródeł odnawialnych”**.



Pierwsza z nich jest kolejnym wydaniem, ukazującego się co 3 lata, informatora o wyznaniach religijnych, a zawarte w nim informacje uzyskano z ankiet statystycznych, kierowanych do władz kościołów i związków wyznaniowych.

Opracowanie zawiera opis struktur organizacyjnych wspólnot wyznaniowych, który poprzedzono szczegółowym omówieniem problemów metodologicznych badań nad wyznaniem religijnym w Polsce. Zamieszczono w nim m.in. źródła danych o wyznaniach oraz krótki rys historyczny omawianego tematu. Dalszą część opracowania stanowi prezentacja poszczególnych wyznań, obejmująca indeksy oraz karty informacyjne.

Indeksy, zawierające pełny wykaz uwzględnionych w opracowaniu kościołów i związków wyznaniowych, porządkują materiał alfabetycznie i demograficznie, co umożliwia wyszukiwanie danych i opisów. Zasadniczą część publikacji stanowią karty informacyjne, w których znajduje się charakterystyka wyznań. Przy każdym opisywanym wyznaniu Czytelnik znajdzie podstawowe dane adresowe (rok powstania i rejestracji), charakterystykę opisową wyznania (m.in. historię, doktrynę, zasady życia religijnego wyznawców, zasady przyjmowania członków, strukturę i działalność związku wyznaniowego) oraz dane statystyczne (liczbę wyznawców, duchownych, jednostek kościelnych, obiektów sprawowania kultu, udzielonych chrztów czy ślubów). Blok kart informacyjnych zamyka zestawienie krótkich notek o wyznaniach nowo zarejestrowanych oraz niesklasyfikowanych, a także tych które formalnie zaprzestały działalności, uległy rozwiązaniu bądź przyłączyły się do innych wspólnot. Zamieszczony materiał wzbogacono tablicami, wykresami oraz mapami, ukazującymi struktury administracyjno-terytorialne niektórych kościołów.

Publikacja wydana w języku polskim, dostępna również na płycie CD oraz na stronie internetowej GUS.



Zeszyt metodologiczny natomiast zawiera ogólne zasady metodyczne sporządzania sprawozdań z produkcji energii ze źródeł odnawialnych (OZE) oraz opracowywania danych statystycznych z tego zakresu. Jego celem jest prezentacja jednolitych zasad zbierania i publikowania danych o energii ze źródeł odnawialnych i uzupełnienie tą problematyką zeszytu metodycznego pt. „Zasady metodyczne sprawozdawczości statystycznej z zakresu gospodarki paliwami i energią oraz definicje stosowanych pojęć”, wydane go w 2006 r.

Tegoroczna edycja opracowania zawiera regulacje prawne dotyczące energii ze źródeł odnawialnych, a także szczegółową charakterystykę odnawialnych nośników energii oraz obiektów technicznych, w których zachodzą przemiany energetyczne i procesy technologiczne, w wyniku których uzyskiwana jest energia użytkowa. Zweryfikowano również przeliczniki wagowo-objętościowe niektórych paliw zgodnie z obowiązującymi rozporządzeniami, a także uzupełniono wykaz podstawowych materiałów źródłowych i algorytmy obliczania kategorii dotyczących OZE.

Zeszyt jest adresowany do osób sporządzających sprawozdania statystyczne z zakresu OZE oraz osób korzystających z publikowanych przez statystykę publiczną wyników na ten temat.

Opracowanie w wersji polskiej, dostępne jest również na stronie internetowej Urzędu.

W sierpniu br. ukazały się także: „Aktywność ekonomiczna ludności Polski, I kwartał 2016 r.”, „Biuletyn Statystyczny Nr 7/2016”, „Ceny robót budowlano-montażowych i obiektów budowlanych. Czerwiec 2016 r.”, „Ceny w gospodarce narodowej w 2015 r.”, „Informacja o sytuacji społeczno-gospodarczej kraju w lipcu 2016 r.”, „Handel zagraniczny. Styczeń—Grudzień 2015 r.”, „Obroty towarowe handlu zagranicznego w 2015 r.”, „Produkcja ważniejszych wyrobów przemysłowych w lipcu 2016 r.”, „Rachunki narodowe według sektorów i podsektorów instytucjonalnych w latach 2011—2014”, „Regiony Polski 2016” (folder), „Transport — wyniki działalności w 2015 r.”, „Uniwersytety Trzeciego Wieku w roku akademickim 2014/2015” (folder), „Zmiany strukturalne grup podmiotów gospodarki narodowej w rejestrze REGON, I półrocze 2016 r.” oraz „Wiadomości Statystyczne nr 8/2016 (663)”.

Do Autorów

Szanowni Państwo!

- W „Wiadomościach Statystycznych” publikowane są artykuły o charakterze naukowym poświęcone teorii i praktyce statystycznej, prezentujące wyniki oryginalnych badań teoretycznych lub analitycznych wykorzystujących metody statystyki matematycznej, opisowej lub ekonometrii. W miesięczniku zamieszczane są również artykuły przeglądowe, popularnonaukowe, recenzje publikacji naukowych oraz inne opracowania informacyjne. Prezentowany w artykule naukowym problem badawczy powinien być jednoznacznie zdefiniowany oraz istotny dla oceny zjawisk społecznych lub gospodarczych. Wyniki studiów przeprowadzanych w artykułach winny oddziaływać na rozwój myśli statystycznej oraz edukacji, wnosząc oryginalny wkład do tej dziedziny.

Zasopismo publikuje także artykuły i opracowania prezentujące informacje o teorii i praktyce statystycznej, jak również o problemach edukacji statystycznej. Dotyczą one: programów badań statystycznych statystyki publicznej, systemu zbierania i udostępniania informacji statystycznych, zastosowań informatyki w statystyce, informacji o konferencjach naukowych, działalności organów doradczych prezesa GUS oraz edukacji statystycznej.

- Artykuły kierowane do opublikowania w „Wiadomościach Statystycznych” powinny zawierać precyzyjny opis badanych zjawisk i stosowanych metod oraz autorskie wnioski i sugestie dotyczące rozwoju badań i analiz statystycznych. Autorzy winni wyraźnie określić cel artykułu oraz jasno przedstawić uzyskane wyniki przeprowadzonej analizy. W przypadku prezentacji badań prowadzonych przez Autorów należy opisać zastosowaną w nich metodę. Przy prezentacji nowatorskich metod analizy pożądane jest podanie przykładu pokazującego ich zastosowanie w praktyce statystycznej.
- Artykuły zamieszczane w „Wiadomościach Statystycznych” powinny wyrażać opinie własne Autorów. Autorzy ponoszą odpowiedzialność za treści prezentowane w artykułach. W razie zgłaszania przez czytelników zastrzeżeń odnoszących się do tych treści, Autorzy są zobligowani do udzielenia odpowiedzi na łamach miesięcznika.
- Po wstępnej ocenie przez Redakcję „Wiadomości Statystycznych” tematyki artykułu pod względem zgodności z profilem czasopisma, artykuły mające charakter naukowy przekazywane są do oceny osobom specjalizującym się w poszczególnych dziedzinach, które kierują się kryterium oryginalności i jakości opracowania, w tym treści i formy, a także potencjalnego zainteresowania czytelników.
- Autorzy artykułów, które otrzymały pozytywne recenzje, wprowadzają zasugerowane przez recenzentów poprawki i dostarczają Redakcji zaktualizowaną wersję opracowania. Autorzy poświadczają w przysłanym piśmie uwzględnienie wszystkich poprawek. Jeśli pojawi się różnica zdań co do zasadności proponowanych zmian, należy wyjaśnić, które poprawki zostały uwzględnione, a w przypadku ich nieuwzględnienia przedstawić motywy swojego stanowiska.

- Kontroli poprawności stosowanych przez Autorów metod statystycznych dokonują redaktorzy statystyczni.
- Decyzję o publikacji artykułu podejmuje Kolegium Redakcyjne „Wiadomości Statystycznych”. Podstawą tej decyzji jest wynik dyskusji dotyczącej zgłoszonego artykułu, w której uwzględniane są opinie przedstawione w recenzjach wraz z rekomendacją ich opublikowania.
- Redakcja „Wiadomości Statystycznych” przestrzega zasady nietolerowania przejawów nierzetelności naukowej autorów artykułów polegającej na:
 - nieujawnianiu współautorów, mimo że wnieśli oni istotny wkład w powstanie artykułu, określanemu w języku angielskim terminem „ghostwriting”;
 - podawaniu jako współautorów osób o znikomym udziale lub niebiorących udziału w opracowaniu artykułu, określanemu w języku angielskim terminem „guest authorship”.

Stwierdzone przypadki nierzetelności naukowej w tym zakresie mogą być ujawniane. W celu przeciwdziałania zjawiskom „ghostwriting” i „guest authorship” należy dołączyć do przesłanego artykułu oświadczenie, którego wzór zamieszczono na stronie internetowej czasopisma (link do załącznika znajduje się w zakładce „Do Autorów”).

Główną odpowiedzialność za rzetelność przekazanych informacji, łącznie z informacją na temat wkładu poszczególnych współautorów w powstanie artykułu, ponosi zgłaszający artykuł.

Redakcja zastrzega sobie prawo dokonywania w artykułach zmian tytułów, skrótów i przeredagowania tekstu i tablic bez naruszenia zasadniczej myśli Autora.

Informacje dotyczące przysyłania artykułów do „Wiadomości Statystycznych”

- Artykuły należy dostarczać pocztą elektroniczną na adres:

a.swiderska@stat.gov.pl lub e.grabowska@stat.gov.pl

Redakcja „Wiadomości Statystyczne”

Główny Urząd Statystyczny

al. Niepodległości 208, 00-925 Warszawa

- Konieczne jest dołączenie do artykułu skróconej informacji (streszczenia) o jego treści (ok. 10 wierszy) w języku polskim i, jeżeli jest to możliwe, także w językach angielskim i rosyjskim. Streszczenie powinno być utrzymane w formie bezosobowej i zawierać: zwięźle sprecyzowany cel badania, przybliżony jego zakres i przyjętą metodologię badania oraz ważniejsze wnioski.
- Prosimy również o podawanie słów kluczowych, w języku polskim i angielskim, przybliżających zagadnienia w artykule.
- Redakcja rozpoczyna postępowanie kwalifikujące artykuł do opublikowania po spełnieniu warunku przesłania przez Autora oświadczenia.
- Pytania dotyczące przesłanego artykułu, co do jego aktualnego statusu itp., należy kierować do redakcji na adres: **a.swiderska@stat.gov.pl** lub **e.grabowska@stat.gov.pl** lub tel. 22 608-32-25.

Wymogi czasopisma dotyczące przygotowania artykułu

Artykuł powinien mieć optymalną objętość (łącznie z wykresami, tablicami i literaturą) 10—20 stron przygotowanych zgodnie z poniższymi wytycznymi:

1. Edytor tekstu — Microsoft Word, format *.doc lub *.docx.
2. Czcionka:
 - o autor — Arial, wersalik, wyrównanie do lewej, 12 pkt.,
 - o tytuł opracowania — Arial, wyśrodkowany, 16 pkt.,
 - o tytuły rozdziałów i podrozdziałów — Times New Roman, wyśrodkowany, kursywa, 14 pkt.,
 - o tekst główny — Times New Roman, normalny, wyjustowany, 12 pkt.,
 - o przypisy — Times New Roman, 10 pkt.
3. Marginesy przy formacie strony A4 — 2,5 cm z każdej strony.
4. Odstęp między wierszami półtorej linii oraz interlinia przed tytułami rozdziałów.
5. Pierwszy wiersz akapitu wcięty o 0,4 cm, enter na końcu akapitu.
6. Wyszczególnianie rozmaitych kategorii należy zacząć od kropek, a numerowanie od cyfr arabskich.
7. Strony powinny być ponumerowane automatycznie.
8. Wykresy powinny być załączone w osobnym pliku w oryginalnej formie (Excel lub Corel), tak aby można było je modyfikować przy opracowaniu edytorskim tekstu. W tekście należy zaznaczyć miejsce ich włączenia. Należy także przekazać dane, na podstawie których powstały wykresy.
9. Tablice należy zamieszczać w tekście, zgodnie z treścią artykułu. W tablicach nie należy stosować rastrów, cieniowania, pogrubiania czy też podwójnych linii itp.
10. Pod wykresami i tablicami należy podać informacje dotyczące źródła opracowania.
11. Oznaczenia literowe należy wyróżniać następująco: macierze — wersalik, proste, pogrubione (np. **P**, **N_{ij}**); wektory — małe litery, kursywa, pogrubione (np. **w**, **x_i**); pozostałe zmienne — małe litery, kursywa, bez pogrubienia (np. *w*, *x_i*).
12. Stosowane są skróty: tablica — tabl., wykres — wykr.
13. Przypisy do tekstu należy umieszczać na dole strony.
14. Przytaczane w treści artykułu pozycje literatury przedmiotu należy wykonać według stylu APA.

Zasady przywoływania pracy w tekście:

- a. Jeden autor: bez względu na to ile razy przywoływana jest praca, zawsze należy podać nazwisko autora i datę publikacji pracy, w przypadku więcej niż jednej pracy danego autora opublikowanej w tym samym roku należy dodać kolejne litery alfabety przy dacie (np., 2001a), zasada ta obowiązuje także w przypadku większej liczby autorów danej pracy.

Przykład zapisu:

Jak stwierdza Iksiński (2001)...

Badania wskazują, iż ... (Iksiński, 2001).

- b. Dwóch autorów: bez względu na to ile razy przywoływana jest praca, zawsze należy podać nazwiska obu autorów i datę publikacji pracy, w przypadku więcej niż jednej pracy tych autorów opublikowanej w tym samym roku należy dodać kolejne litery alfabetu przy dacie. Nazwiska autorów zawsze należy łączyć spójnikiem „i”, nawet w przypadku przywoływania publikacji obcojęzycznej.

Przykład zapisu:

Jak sugerują Iksiński i Nowak (1999)...

Badania wskazują, iż ... (Iksiński i Nowak, 1999).

- c. 3—5 autorów: przywołanie po raz pierwszy — należy wymienić nazwiska wszystkich autorów, rozdzielając je przecinkami i stawiając spójnik „i” pomiędzy dwoma ostatnimi nazwiskami. Przy kolejnych wskazaniach tej samej pracy można zastosować określenie „i współpracownicy” (w przypadku umieszczenia przywołania nazwisk w strukturze zdania) lub „i in.” (w przypadku, gdy nazwiska autorów nie stanowią części struktury zdania).

Przykład zapisu:

Przywołanie po raz pierwszy:

Jak sugerują Nowak, Iksiński i Jankiewicz (2003) ...

Badania (Nowak, Iksiński i Jankiewicz, 2003) wskazują, iż ...

Kolejne przywołania:

Badania Nowaka i współpracowników (2003) wskazują, iż ... Badania te (Nowak i in., 2003) ...

- d. 6 i więcej autorów: wymienić należy tylko nazwisko pierwszego autora, zarówno gdy praca przywoływana jest po raz pierwszy, jak i w późniejszych przywołaniach, natomiast pozostałych autorów należy zastąpić skrótem „i in.” (gdy nazwiska nie stanowią części struktury zdania). W literaturze cytowanej należy umieścić nazwiska wszystkich autorów pracy.

Przykład zapisu:

Nowakowski i współpracownicy twierdzą, iż ... (1997).

Pierwsze badania na ten temat (Nowakowski i in., 1997) sugerują

- e. Przywoływanie jednocześnie kilku prac: należy wymienić je alfabetycznie, według nazwiska pierwszego autora. Przywołania kolejnych prac muszą być oddzielone średnikiem i umieszczone w nawiasie. Lata wydania prac tego samego autora/autorów muszą być oddzielone przecinkiem.

Przykład zapisu:

(Iksiński, 2001; Nowak i Iksiński, 1999)

(Iksiński, 1997, 1999, 2004a, 2004b; Nowak i Iksiński, 1999)

- f. Przywoływanie pracy za innym autorem: stosujemy w tekście, natomiast w literaturze cytowanej umieszczamy jedynie pracę czytaną.

Przykład zapisu:

Jak wykazał Nowakowski (1990; za: Zieniecka, 2007) ...

Badania sugerują, iż ... (Nowakowski, 1990; za: Zieniecka, 2007).

15. Wykaz literatury należy zamieszczać na końcu opracowania.

Prace zapisujemy przy zachowaniu kolejności alfabetycznej cytowanych dzieł, przy czym decyduje pierwsza litera nazwiska autora.

Każdą nową pracę zaczynamy bez wcięcia, wyrównanie do lewego marginesu, a kolejne wiersze danego adresu bibliograficznego powinny być zapisane z wcięciem 0,4 cm.

Zasady zapisu literatury załącznikowej:

Poniżej znajdują się schematy zapisów bibliograficznych podstawowych źródeł (artykułów i książek). Sposoby zapisu innych, rzadziej powoływanych źródeł są szczegółowo opisane w szóstym wydaniu „Publication Manual of the American Psychological Association”.

- a. artykuł w czasopiśmie, w którym każdy kolejny numer/zeszyt (*issue*) w ramach jednego rocznika ma osobną numerację stron (w każdym zeszycie pierwsza strona opatrzona jest numerem 1):
Nazwisko, X., Nazwisko2, X. Y., Nazwisko3, Z. (rok). Tytuł artykułu. *Tytuł Czasopisma, nr rocznika* (nr zeszytu), strona początku–strona końca.
 - b. artykuł w czasopiśmie, w którym kolejne numery/zeszyty (*issues*) w ramach jednego rocznika nie mają osobnej numeracji stron (pierwsza strona w kolejnym zeszycie opatrzona jest numerem kolejnym, po ostatniej stronie w zeszycie poprzednim):
Nazwisko, X., Nazwisko2, X. Y., Nazwisko3, Z. (rok). Tytuł artykułu. *Tytuł Czasopisma, nr rocznika*, strona początku–strona końca.
 - c. jeśli artykuł ma numer DOI (*Digital Object Identifier*), należy podać go na końcu zapisu bibliograficznego: Nazwisko, X., Nazwisko2, X. Y. (rok). Tytuł artykułu. *Tytuł Czasopisma, nr rocznika*, strona początku–strona końca. DOI: xxxxx.
 - d. książka:
Nazwisko, X., Nazwisko, X. Y. (rok). *Tytuł książki*. Miejsce wydania: Wydawnictwo.
 - e. książka napisana pod redakcją:
Nazwisko, X. (red.). (rok). *Tytuł książki*. Miejsce wydania: Wydawnictwo.
 - f. rozdział w pracy zbiorowej:
Nazwisko, X. (rok). Tytuł rozdziału. W: Y. Nazwisko, B. Nazwisko (red.), *Tytuł książki* (s. strona początku–strona końca). Miejsce wydania: Wydawnictwo.
W stylu APA proponuje się zapis bibliograficzny bez użycia dwukropka po przymku W (*In*), pisany wielką literą. W polskim zapisie jednak przyjmujemy zasadę pisania dwukropka po W:
 - g. jeśli dany tekst znajduje się na stronie internetowej i nie jest artykułem w czasopiśmie, książką ani rozdziałem w książce, należy podać autora, datę publikacji (jeśli jest znana), tytuł, a następnie zamieścić informacje o stronie, skąd został pobrany tekst:
Nazwisko, X. (rok). *Tytuł tekstu*. Pobrane z: adres strony internetowej.
16. W wykazie literatury należy zamieścić wyłącznie pozycje przytoczone w artykule.
17. Opracowanie przygotowane w sposób niezgodny z powyższymi wskazówkami będzie odesłane z prośbą o dostosowanie jego formy do wymagań redakcji.

Charakterystyka zakresu tematycznego poszczególnych działów „Wiadomości Statystycznych”

STUDIA METODOLOGICZNE

W dziale tym zamieszczane są artykuły naukowe zawierające prezentacje teoretycznych rozwiązań metodologicznych, ze wskazaniem ich praktycznej użyteczności, w tym prace o charakterze przeglądowym i porównawczym oraz dotyczące zagadnień etyki statystycznej. Poruszane tu zagadnienia mogą obejmować różnorodne dziedziny statystyki, ekonomii matematycznej i ekonometrii, a prezentowane rezultaty badawcze stwarzają możliwość efektywnego zastosowania w empirycznych badaniach i analizach statystycznych, umożliwiając doskonalenie ich jakości i zasobu informacyjnego.

STATYSTYKA W PRAKTYCE

Dział ten dotyczy prac naukowych poświęconych nowatorskim zastosowaniom znanych narzędzi i modeli statystycznych w praktyce, analizy i statystycznej oceny określonych zjawisk społeczno-ekonomicznych i innych, a w szczególności artykułów wykorzystujących dane pochodzące z zasobów statystyki publicznej. Publikowane są tutaj także teksty sygnalizujące praktyczne problemy związane z: projektowaniem badań statystycznych, uzyskiwaniem, integracją i przetwarzaniem danych oraz generowaniem wyników informacji statystycznych i kontrolą ich ujawniania wraz z propozycjami efektywnych metod rozwiązywania owych problemów.

EDUKACJA STATYSTYCZNA

Artykuły publikowane w tym dziale dotyczą metod i efektów nauczania statystyki oraz popularyzacji myślenia statystycznego. W szczególności odnosi się to do problemów związanych z kształceniem w zakresie stosowania statystyki na wszystkich poziomach edukacji, a także wykorzystywania nowoczesnych idei i metod dydaktycznych (w tym eksperymentów i pokazów) oraz pomocy naukowych (np. komputerów, Internetu i innych urządzeń) w nauczaniu statystyki. Szczególną uwagę koncentruje się tutaj na rozumieniu prawdopodobieństwa i statystyki, badaniach z zakresu nauczania statystyki, postaw i zachowań społecznych w odniesieniu do statystyki, jak również rozumieniu informacji statystycznych. Ponadto ukazywane są problemy związane z prezentacją danych statystycznych oraz ich interpretacją w powszechnym obiegu informacyjnym (np. w środkach społecznego przekazu).

STATYSTYKA W SPOŁECZEŃSTWIE INFORMACYJNYM

Jest to blok tematyczny zawierający artykuły z zakresu wykorzystania narzędzi informatycznych do uzyskiwania i przetwarzania informacji statystycznych, naliczania danych wyników, ich prezentacji i rozpoznawania oraz dotyczące nowoczesnych technik programistycznych, interaktywnych i komunikacyjnych umożliwiających potencjalnym użytkownikom danych statystycznych ich wykorzystanie w oczekiwanym przez siebie zakresie i formie. W dziale tym przedstawiane mogą być również artykuły dotyczące: wykorzystania technologii informacyjnych i komunikacyjnych (ICT), gospodarki opartej na wiedzy, problematyki innowacyjności, zagadnień dotyczących przepływu informacji we współczesnym społeczeństwie (w tym z użyciem Internetu) oraz przetwarzania i analizy zagadnień związanych z Big Data.

Z DZIEJÓW STATYSTYKI

Prace należące do tego działu tematycznego poświęcone są historii prowadzenia obserwacji statystycznych, rozwoju i doskonalenia ich metodologii narzędzi. Ponadto zamieszczane są opisy wartościowych faktów dotyczących życia i osiągnięć zawodowych wybitnych statystyków, jak również wiodących instytucji i organizacji statystycznych w Polsce i za granicą.

INFORMACJE. PRZEGLĄDY. RECENZJE

Dział ten obejmuje informacje o najważniejszych wydarzeniach w życiu statystyki polskiej i międzynarodowej, działalności Rady Statystyki oraz z życia Polskiego Towarzystwa Statystycznego, a także sprawozdania z prestiżowych konferencji naukowych, recenzje książek naukowych i popularnonaukowych z zakresu statystyki i ekonometrii, jak również rekomendacje nowych, istotnych i ciekawych pozycji wydawniczych dotyczących tego obszaru wiedzy. Jest to jedyna część czasopisma zawierająca teksty niemające charakteru artykułów naukowych.