
**KONGRES
STATYSTYKI
POLSKIEJ**

POZNAŃ | 18–20 kwietnia 2012

Z OKAZJI JUBILEUSZU 100-LECIA
POLSKIEGO TOWARZYSTWA STATYSTYCZNEGO

HONOROWY PATRONAT
PREZYDENTA RZECZYPOSPOLITEJ POLSKIEJ
BRONISŁAWA KOMOROWSKIEGO

STRESZCZENIA REFERATÓW

POZNAŃ, 18–20 KWIETNIA 2012

ORGANIZATORZY:



Główny
Urząd
Statystyczny



Polskie
Towarzystwo
Statystyczne



Urząd
Statystyczny
w Poznaniu



Uniwersytet
Ekonomiczny
w Poznaniu

Spis treści

Komitety Organizacyjny Kongresu	9
Komitety Naukowy Kongresu	9
Komitety Honorowy Kongresu	10
Program Kongresu Statystyki Polskiej	11
The Programme of the Congress of Polish Statistics	27
Streszczenia referatów	43
Andrzejczak Karol <i>Ewolucja i granica wzrostu wskaźnika nasycenia samochodami osobowymi w skali globalnej i w Polsce</i>	43
Augustyniak Zbigniew <i>Prezentacja danych statystycznych w wersji on-line jako element realizacji projektu SISP</i>	44
Bal-Domańska Beata <i>Modelowanie procesów konwergencji w europejskiej przestrzeni regionalnej</i>	46
Bartniczak Bartosz <i>Moduł wskaźników zrównoważonego rozwoju w ramach Banku Danych Lokalnych</i>	48
Bartosiewicz Stanisława, Stańczyk Elżbieta <i>Trochę o społeczno-ekonomicznej historii Polski w latach 1918-2008</i>	50
Batóg Barbara, Batóg Jacek, Mojsiewicz Magdalena, Wawrzyniak Katarzyna <i>Statystyczny system monitoringu stopnia realizacji strategii rozwoju miasta</i>	51
Bąk Iwona <i>Statystyczna analiza aktywności turystycznej seniorów w Polsce</i>	52
Bednarski Tadeusz <i>Statystyczna analiza przyczyn obciążoności próby w badaniach sondażowych rynku pracy</i>	53
Beręsewicz Maciej <i>Internetowe źródła danych w świetle spisu powszechnego na przykładzie rynku nieruchomości w Poznaniu</i>	55
Berger Jan, Walczak Tadeusz, Witkowski Janusz <i>Statystyka publiczna - rozwój historyczny i aktualne wyzwania</i>	56
Białek Jacek <i>Szczególny przypadek pewnej ogólnej formuły indeksu cen</i>	60
Bielak Renata <i>Rola statystyki w procesie kształtowania polityk rozwoju – priorytety i wyzwania</i>	60
Błaczowska Anna, Grześkowiak Alicja <i>Uwarunkowania demograficzno-gospodarcze procesu kształcenia gimnazjalnego na Dolnym Śląsku i Opolszczyźnie w latach 2003-2010</i>	63
Błażejowski Marcin <i>Algorytmy modelowania i prognozowania na przykładzie rynku turystycznego</i>	65

Błażejowski Marcin, Kufel Paweł, Kufel Tadeusz <i>Budowa banków danych jako podstawa analiz statystyczno-ekonometrycznych w oprogramowaniu GRETL</i>	65
Bonnéry Daniel, Chauvet Guillaume, Deville Jean-Claude <i>Neyman type optimality for marginal quota sampling</i>	66
Borucka Jadwiga, Romaniuk Joanna, Frątczak Ewa <i>Trwałość pierwszych związków (małżeństw i kohabitacji) w kohortach urodzeniowych 1950–1970</i>	66
Borys Tadeusz, Potkański Tomasz <i>Monitorowanie rozwoju regionalnego i lokalnego</i>	67
Brzezińska Justyna <i>Analiza niezależności zmiennych nominalnych z wykorzystaniem modeli logarytmiczno-liniowych w R</i>	68
Buciak Robert, Pieniążek Marek <i>Klasyfikacja przestrzenna obszarów wiejskich w Polsce</i>	70
Caliński Tadeusz <i>Rozwój polskiej myśli statystycznej – Osiągnięcia w biometrii</i>	72
Ceranka Bronisław, Graczyk Małgorzata <i>Estymacja całkowitej masy obiektów w układach wagowych</i>	74
Chambers Ray, Gunky Kim <i>Regression analysis using data obtained by probability linking of multiple data sources</i>	75
Chlebicki Paweł <i>ArcGIS jako potężne narzędzie do wizualizacji i analizy danych statystycznych</i>	77
Chłoń-Domińczak Agnieszka <i>Wykorzystanie wskaźników w obszarze edukacji w kontekście rozwoju polityki opartej na faktach</i>	78
Cmela Piotr, Kornecki Janusz <i>Strategiczne dylematy wsparcia rozwoju mikroprzedsiębiorstw w Polsce</i>	80
Czyżewski Andrzej, Grzelak Aleksander <i>Możliwości wykorzystania statystyki bilansów przepływów międzygałęziowych dla makroekonomicznych ocen w gospodarce</i>	81
Decker Reinhold <i>Model-based Analysis of Online Consumer Reviews. Methods and Applications</i>	82
Dehnel Grażyna <i>Estymacja pośrednia w krótkookresowej statystyce przedsiębiorstw</i>	83
Didkowska Joanna, Wojciechowska Urszula <i>Nowotwory złośliwe w Polsce jako problem zdrowia publicznego. Współczesne narzędzia do mierzenia zjawiska.</i>	84
Doniec Dorota <i>Rachunki regionalne</i>	86
Dudek Hanna, Landmesser Joanna <i>Satysfakcja z osiągniętych dochodów a względna deprywacja</i>	88
Dygaszewicz Janusz, Janczur-Knapiek Magdalena, Wardzińska-Sharif Amelia <i>Spisy powszechne PSR 2010 i NSP 2011 oraz systemy informacji geograficznej w statystyce publicznej</i>	90
Dyzmann-Sroka Agnieszka, Roszak Andrzej, Moczko Jerzy, Trojanowski Maciej <i>Metody podniesienia jakości danych rejestrów nowotworów złośliwych</i>	93
Fattorini Lorenzo <i>Design-Based Inference on Ecological Diversity</i>	94

Filas-Przybył Sylwia, Kaźmierczak Maciej, Stachowiak Dorota <i>Wykorzystanie danych ze źródeł administracyjnych w statystyce miast</i>	97
Franceschi Sara (University of Tuscia), Fattorini Lorenzo (Università di Siena), Maffei Daniela (Università di Firenze) <i>Design-based treatment of unit nonresponse by the calibration approach</i>	98
Frątczak Ewa <i>Tradycja i nowoczesność w modelowaniu zjawisk i procesów demograficznych. Od tablic J. Graunta i siatki Lexisa po modele analiz wzdlużnych z efektem stałym i efektem losowym.</i>	99
Frątczak Ewa, Korczyński Adam <i>Tradycja i współczesność w metodach imputacji danych. Przegląd teorii – wybrane przykłady zastosowań</i>	100
Furmańczyk Konrad, Jaworski Stanisław <i>Wykrywanie zmian poziomu serii niezależnych obserwacji</i>	103
Gamrot Wojciech <i>O empirycznych prawdopodobieństwach inkluzji</i>	103
Ganczarek-Gamrot Alicja <i>Wielowymiarowe modele FIGARCH w ocenie ryzyka na polskim rynku energii elektrycznej</i>	104
Gatnar Eugeniusz <i>Rola NBP w systemie statystyki publicznej w Polsce</i>	105
Gerasymenko Sergiy, Chupryna Olena <i>Optymizacja spisu wskaźników dla oceny porównawczej obiektów socjalno-ekonomicznych</i>	106
Ghosh Malay <i>Finite Population Sampling: A Model-Design Synthesis</i>	108
Głowiński Cezary, Konikiewicz Kamil <i>Wykorzystanie sieci społecznych w modelowaniu zachowań klientów firm telekomunikacyjnych</i>	109
Gołata Elżbieta <i>Spis ludności i prawda</i>	110
Goujon Anne, Bauer Ramon, Samir K.C., Potančoková Michaela <i>The human capital puzzle: Why is it so hard to find good data on educational attainment?</i>	112
Górecki Tomasz, Krzyśko Mirosław <i>Wersja jądrowa funkcjonalnych składowych głównych</i>	113
Groenen Patrick J.F. <i>The support vector machine as a powerful tool for binary classification</i>	114
Gruchociak Hanna <i>Delimitacja lokalnych rynków pracy w Polsce</i>	115
Grzenda Wioletta <i>Bayesowski semiparametryczny model Coxa w badaniu determinant pozostawania bez pracy osób młodych</i>	117
Hanusik Krystyna, Łangowska-Szcześniak Urszula <i>Uwarunkowania różnicowania poziomu i struktury konsumpcji gospodarstw domowych w Polsce</i>	118
Hozer Józef, Hozer-Koćmiel Marta <i>Quantum satis dla firm w Polsce w 2012 roku</i>	120
Jajuga Krzysztof <i>Rozwój polskiej myśli statystycznej w naukach ekonomicznych</i>	121
Jefmański Bartłomiej, Markowska Małgorzata <i>Ocena dynamiki i kierunku zmian rozwoju inteligentnej specjalizacji regionów europejskich z zastosowaniem klasyfikacji rozmytej</i>	122
Józefowski Tomasz <i>Zastosowanie estymatora typu SPREE w szacowaniu liczby osób bezrobotnych w przekroju podregionów</i>	123
Jurečková Jana <i>Regression quantiles and their two-step modifications</i>	125

Kalinowski Marcin, Krzykowski Grzegorz <i>Dobór i selekcja danych jako kluczowy element badań statystycznych na przykładzie rynków finansowych</i>	125
Kamińska-Gawryluk Ewa <i>Zróżnicowania regionalne w Polsce i wybranych krajach Unii Europejskiej</i>	126
Keilman Nico <i>Challenges for statistics on households and families</i>	129
Kisielińska Joanna <i>Dokładna metoda bootstrapowa na przykładzie estymacji średniej i wariancji</i>	130
Klimanek Tomasz <i>Wykorzystanie estymacji pośredniej uwzględniającej korelację przestrzenną w badaniach społecznych w Polsce</i>	131
Kobus Paweł <i>Podejście bayesowskie do estymacji rozkładów cech wybranych jednostek zbiorowości z wykorzystaniem danych zagregowanych</i>	131
Kołodziejczak Barbara, Roszak Magdalena <i>Rapid e-learning w biostatystyce</i>	132
Kończak Grzegorz <i>Wnioskowanie statystyczne dla danych w wielowymiarowych tablicach wielodzielczych</i>	134
Kopczewska Katarzyna <i>Regiony w przestrzeni. Czy lokalizacja i dostępność mają znaczenie dla rozwoju gospodarczo-społecznego?</i>	135
Kordos Jan <i>Współzależność pomiędzy rozwojem teorii i praktyki badań reprezentacyjnych w Polsce</i>	137
Koronacki Jacek <i>Analiza danych o dużym wymiarze w przypadku małej liczności próby</i>	140
Kosiorowski Daniel <i>Głębia położenia-rozrzutu w odpornej analizie strumienia danych ekonomicznych</i>	142
Kot Stanisław Maciej <i>Stochastyczne skale ekwiwalentności dla Polski</i>	144
Kowaleski Jerzy, Majdzińska Anna <i>Starzenie się populacji krajów Unii Europejskiej – nieodległa przeszłość i prognoza</i>	145
Kowalewski Grzegorz <i>Metodologia ankietowych badań koniunktury gospodarczej</i>	147
Kowalewski Jacek, Natkowska Monika <i>Wykorzystanie mapy statystyki krótkookresowej w procesie organizacji badań przedsiębiorstw prowadzonych przez statystykę publiczną</i>	149
Kowerski Mieczysław <i>Badania nastrojów gospodarczych przedsiębiorców i konsumentów na poziomie regionu. Przykład województwa lubelskiego</i>	151
Krzanowski Wojciech <i>Classification: some old principles applied to new problems</i>	153
Krzyśko Mirosław <i>Jan Czekanowski (1882–1965)</i>	154
Krzyśko Mirosław, Waszak Łukasz <i>Jądrowe korelacje kanoniczne</i>	156
Kukło Cezary <i>Rodzina i jej gospodarstwo na ziemiach polskich do końca epoki przedprzemysłowej - mity i prawda</i>	158
Kupiszewski Marek <i>What drives population change?</i>	159
Kwiatkowska-Ciotucha Dorota, Załuska Urszula <i>Ocena wykorzystania środków z europejskiego funduszu społecznego w poszczególnych regionach polski w obecnym okresie programowania</i>	160
Ledwina Teresa <i>Jerzy Neyman (1894–1981)</i>	160
Ledwina Teresa, Wyłupek Grzegorz <i>Test dla dwóch prób przy jednostronnych alternatywach</i>	161

Lemmi Achille, Panek Tomasz <i>The dimensions of poverty: theory, models and new perspectives</i>	164
Lewiński vel Iwański Stanisław, Wilimowska Zofia <i>Struktura finansowa polskich przedsiębiorstw, a modelowanie statystyczne</i>	165
Liberda Barbara, Świeczewska Iwona, Tomaszewicz Łucja <i>Metodologia rachunków satelitarnych na podstawie rachunku satelitarnego sportu dla Polski</i>	167
Liberda Barbara <i>Rachunki generacyjne metodą pomiaru bogactwa kolejnych pokoleń</i>	168
Łagodziński Władysław Wiesław, Krasucki Piotr <i>Zintegrowany system danych statystycznych z ochrony zdrowia - niezbędne minimum</i>	169
Łazowska Bożena <i>Władysław Bortkiewicz (7 VIII 1868 - 15 II 1931)</i>	170
Łuczak Aleksandra, Wysocki Feliks <i>Ocena poziomu rozwoju społeczno-gospodarczego powiatów województwa wielkopolskiego</i>	171
Malaga Krzysztof <i>Dylematy współczesnej statystyki wzrostu gospodarczego</i>	173
Małecka Marta <i>Ocena wariancji estymatorów wartości narażonej na ryzyko (VaR)</i>	174
Markowska Małgorzata, Strahl Dorota <i>Możliwości oceny innowacyjności regionów szczebla NUTS 2 w świetle zasobów informacyjnych Eurostatu</i>	178
Matuszewska-Janica Aleksandra <i>Statystyczna analiza nierówności w płacach kobiet i mężczyzn w Polsce na tle Unii Europejskiej w odniesieniu do wieku, wielkości przedsiębiorstwa i grupy zawodowej</i>	179
Meller Ewa <i>Estymacja dla małych obszarów dla wybranych charakterystyk rynku pracy</i>	180
Mielniczuk Jan <i>Wkład Polaków w rozwój statystyki matematycznej i aplikacyjnej</i>	181
Mielniczuk Jan, Wojtyś Małgorzata <i>Postselekcyjna estymacja gęstości gradacyjnej</i>	182
Migdał-Najman Kamila <i>Konstrukcja hybrydowej, samouczącej się sieci neuronowej typu SOM-GNG</i>	184
Mikulec Artur <i>Empiryczna analiza efektywności kryterium Mojeny i Wisharta w analizie skupień</i>	186
Młodak Andrzej <i>Modernizacja statystyki działalności gospodarczej w wymiarze paneuropejskim</i>	188
Modranka Emilia, Suchecka Jadwiga <i>Statystyka przestrzenna – metody i zastosowania</i>	189
Motoryn Ruslan, Motoryna Iryna, Motoryna Tetiana <i>Porównanie wskaźników wykorzystania oszczędności gospodarstw domowych Ukrainy i Polski</i>	192
Najman Krzysztof <i>Sieci samouczące się w analizie skupień</i>	193
Nowak Lucyna <i>Rozwój badań statystycznych w obszarze demografii i migracji</i>	194
Nowakowski Rafał <i>Wybory samorządowe na przykładzie stolic regionów dolnośląskiego, małopolskiego i wielkopolskiego - 20 lat doświadczeń</i>	194
Okrasa Włodzimierz <i>Statystyka i socjologia: współzależności rozwoju z perspektywy interdyscyplinarnej badań społecznych. Aspekty metodologiczne i instytucjonalne</i>	198

Olbrot-Brzezińska Arleta <i>Spoleczne funkcje informacji statystycznej i odpowiedzialność statystyki publicznej</i>	199
Okupniak Magdalena <i>Zastosowanie analizy wewnątrz- i międzyblokowej w badaniach społecznych</i>	200
Ostasiewicz Walenty <i>Rozwój myśli statystycznej w Polsce</i>	201
Owsiński Jan <i>Optymalny podział rozkładu empirycznego (i kilka problemów z tym związanych)</i>	203
Paradysz Jan <i>Statystyka regionalna: stan, problemy i kierunki rozwoju</i>	205
Parlińska Maria, Babiak Robert <i>Dokumentowanie i ryzyko asymetrii informacji</i>	206
Paszko Elżbieta, Jurek Anna Maria <i>Zastosowanie CQI - ciągłego doskonalenia jakości w procesie poprawy metod nauczania</i>	208
Pawełek Barbara <i>Badanie przydatności wyników testu koniunktury GUS do prognozowania krótkoterminowego zmian aktywności gospodarczej w Polsce na podstawie modelu wektorowej autoregresji</i>	209
Pełka Marcin <i>Podjęcie wielomodelowe w klasyfikacji danych symbolicznych interwałowych</i>	210
Perek-Białas Jolanta <i>Sytuacja starszych generacji w Europie Środkowo-Wschodniej na podstawie danych EU-SILC</i>	211
Pietrzykowski Robert <i>Wykorzystanie regresji kwantylowej w analizach przestrzennych cen ziemi rolniczej</i>	212
Piotrowska-Trybull Marzena, Sirko Stanisław <i>Wpływ jednostki wojskowej na rozwój gminy</i>	213
Pobłocka Agnieszka <i>Polski rynek ubezpieczeń w latach 1991–2010</i>	214
Pociecha Józef <i>Okoliczności powstania polskiego towarzystwa statystycznego w Krakowie i jego pierwszy prezes</i>	215
Podogrodzka Małgorzata <i>Niedostosowanie strukturalne rynku matrymonialnego determinantą spadku procesu zawierania małżeństw w Polsce w latach 1990–2010</i>	217
Ptak-Chmielewska Aneta <i>Uwarunkowania rynku przedsiębiorstw w Polsce. Statystyka przeżywalności firm.</i>	218
Rogalińska Dominika <i>Wyzwania statystyki regionalnej na tle debaty o polityce spójności</i>	219
Roszak Magdalena, Kołodziejczak Barbara <i>E-ewaluacja wiedzy statystycznej studentów medycyny</i>	220
Roszka Wojciech <i>Szacowanie ubóstwa w ujęciu lokalnym na podstawie zintegrowanych źródeł danych</i>	222
Rozkrut Dominik <i>Differentiation of innovation strategies</i>	223
Rozmus Dorota <i>Podjęcie zagregowane w taksonomii</i>	224
Rozpędowska-Matraszek Danuta <i>Ocena skuteczności funkcjonowania opieki zdrowotnej w Polsce - analiza według grup powiatów podobnych</i>	225
Ryś-Jurek Roma <i>Zróżnicowanie regionalne produkcji rolniczej w krajach UE-27</i>	227
Siwiński Włodzimierz <i>Statystyczny obraz kryzysu gospodarczego 2008–2010</i>	229
Sobczak Elżbieta <i>Pracujący wg sektorów intensywności działalności B+R w krajach UE - analiza przestrzenno-strukturalna</i>	231

Sokołowski Andrzej, Basiura Beata <i>Taksonomiczna metoda Warda</i>	234
Stanimir Agnieszka <i>Zróźnicowane techniki prezentacji zmiennych niemetrycz- nych</i>	235
Stańczyk Elżbieta <i>Kwalifikacje zawodowe osób bezrobotnych w świetle danych GUS. Międzywojewódzka analiza porównawcza</i>	237
Stonawski Marcin <i>Kapitał ludzki w warunkach starzenia się ludności w Polsce</i>	238
Strahl Danuta, Markowska Małgorzata <i>Możliwości oceny innowacyjności regionów szczebla NUTS 2 w świetle zasobów informacyjnych Eurostatu .</i>	241
Styrc Marta <i>Czynniki wpływające na stabilność pierwszych małżeństw w Polsce</i>	242
Styrc Marta <i>Gradient edukacyjny rozwodów w Polsce</i>	243
Styrc Marta, Matysiak Anna <i>Status społeczno-ekonomiczny a stabilność mał- żeństw</i>	245
Suchecka Jadwiga, Modranka Emilia <i>Statystyka przestrzenna - metody i za- stosowania</i>	247
Szkutnik Zbigniew <i>Algorytm EM i jego modyfikacje</i>	250
Szmuksta-Zawadzka Maria, Zawadzki Jan <i>O metodach prognozowania bra- kujących danych w szeregach czasowych z wahaniami okresowymi (sezo- nowymi)</i>	251
Sztanderska Urszula <i>Polska statystyka publiczna jako inspiracja i bariera roz- woju badań rynku pracy</i>	252
Szukielój-Bieńkuńska Anna <i>Jakość życia w badaniach GUS</i>	253
Szukielój-Bieńkuńska Anna <i>Pomiar ubóstwa i wykluczenia społecznego w pol- skiej statystyce publicznej</i>	255
Szwarc Krzysztof <i>Poziom życia ludności a zróźnicowanie sytuacji demogra- ficznej w powiatach województwa wielkopolskiego</i>	256
Szymkowiak Marcin <i>Konstrukcja estymatorów kalibracyjnych wartości global- nej dla różnych funkcji odległości</i>	259
Śliwicki Dominik <i>Czynniki determinujące poziom wynagrodzenia. Analiza eko- nometryczna</i>	260
Traat Imbi <i>Domain estimators calibrated on reference survey</i>	262
Trzpiot Grażyna, Orwat-Acedańska Agnieszka <i>Klasyfikacja funduszy in- westycyjnych ze względu na styl zarządzania - podejście regresji kwantylowej</i>	263
Ulman Paweł <i>Wpływ aktywności zawodowej osób niepełnosprawnych na sytu- ację ekonomiczną ich gospodarstw domowych</i>	264
Wawrowski Łukasz <i>Analiza ubóstwa w przekroju powiatów w województwie wielkopolskim z wykorzystaniem metod statystyki małych obszarów</i>	265
Wawrzyniak Katarzyna <i>Ocena stopnia realizacji celów głównych strategii roz- woju kraju według województw z wykorzystaniem analizy korespondencji .</i>	267
Wiatrak Andrzej Piotr <i>Aktualne problemy statystyki rolniczej</i>	268
Wierzbicka Anna <i>Analiza taksonomiczna stanu zdrowia publicznego Polski na tle wybranych krajów europejskich</i>	269
Wilk Justyna <i>Podejście symboliczne w analizach regionalnych</i>	270
Witkowski Janusz, Walczak Tadeusz, Berger Jan <i>Statystyka publiczna - rozwój historyczny i aktualne wyzwania</i>	272

Wróblewska Wiktoria <i>Zmiany maksymalnego trwania życia i długowieczność - wyzwania dla statystyki</i>	276
Wyszkowska Dorota <i>Statystyka regionalna jako instrument wsparcia rozwoju polskich województw</i>	278
Wywiał Janusz <i>Estymacja średniej na podstawie próby prostej porządkowanej za pomocą zmiennej dodatkowej</i>	281
Zdaniewicz Witold <i>Wkład Instytutu Statystyki Kościoła Katolickiego w polską statystykę publiczną</i>	282
Zhang Li-Chun <i>Micro calibration for data integration</i>	283
Zieliński Wojciech <i>Dobór liczności próby w kontrolach prowadzonych przez NIK</i>	284
Żądło Tomasz <i>O predykcji wartości globalnych dla domen skokrelowanych przez strzennie</i>	285
Kontakt	287
Lista uczestników	288

Komitety Organizacyjny Kongresu

Elżbieta Gołata, Uniwersytet Ekonomiczny w Poznaniu - przewodnicząca
Wojciech Adamczewski, Zakład Wydawnictw Statystycznych
Grażyna Dehnel, Uniwersytet Ekonomiczny w Poznaniu
Tomasz Grudziak, Urząd Marszałkowski Województwa Wielkopolskiego
Marek Kalemba, Urząd Miasta Poznania
Tomasz Klimanek, Uniwersytet Ekonomiczny w Poznaniu
Jacek Kowalewski, Urząd Statystyczny w Poznaniu
Władysław Wiesław Łagodziński, Polskie Towarzystwo Statystyczne
Andrzej Plesiński, Wielkopolski Urząd Wojewódzki w Poznaniu
Feliks Wysocki, Uniwersytet Przyrodniczy w Poznaniu

Komitety Naukowy Kongresu

Prof. dr hab. Czesław Domański, Uniwersytet Łódzki - przewodniczący
Prof. dr hab. Andrzej Barczak, Uniwersytet Ekonomiczny w Katowicach
Prof. dr hab. Tadeusz Borys, Uniwersytet Ekonomiczny we Wrocławiu
Prof. dr hab. Tadeusz Caliński, Uniwersytet Przyrodniczy w Poznaniu
Prof. dr hab. Krzysztof Jajuga, Uniwersytet Ekonomiczny we Wrocławiu
Prof. dr hab. Janina Jóźwiak, Szkoła Główna Handlowa w Warszawie
Prof. dr hab. Mirosław Krzyśko, Uniwersytet im. A. Mickiewicza w Poznaniu
Prof. dr hab. Jerzy Moczko, Uniwersytet Medyczny im. K. Marcinkowskiego w Poznaniu
Prof. dr hab. Walenty Ostasiewicz, Uniwersytet Ekonomiczny we Wrocławiu
Prof. dr hab. Tomasz Panek, Szkoła Główna Handlowa w Warszawie
Prof. dr hab. Jan Paradysz, Uniwersytet Ekonomiczny w Poznaniu
Prof. dr hab. Józef Pociecha, Uniwersytet Ekonomiczny w Krakowie
Prof. dr hab. Bogdan Sojkin, Uniwersytet Ekonomiczny w Poznaniu
Prof. dr hab. Mirosław Szreder, Uniwersytet Gdański
Prof. dr hab. Jerzy Wilkin, Uniwersytet Warszawski
Prof. dr hab. Janusz Witkowski, Szkoła Główna Handlowa w Warszawie
Prof. dr hab. Jan Zawadzki, Zachodniopomorski Uniwersytet Technologiczny w Szczecinie

Komitety Honorowy Kongresu

Prof. dr hab. Marek Belka, Prezes Narodowego Banku Polskiego
Dr Andrzej Byrt, Prezes Międzynarodowych Targów Poznańskich
Prof. Jae Chang Lee, Prezydent Międzynarodowego Instytutu Statystycznego
Piotr Florek, Wojewoda Wielkopolski
Prof. dr hab. Marian Gorynia, Rektor Uniwersytetu Ekonomicznego w Poznaniu
Ryszard Grobelny, Prezydent Miasta Poznania
Prof. dr hab. Michał Kleiber, Prezes PAN
Prof. dr hab. Jan Kordos, Prezes PTS w latach 1985–1994
Dr Kazimierz Kruska, Prezes PTS w latach 2006–2009
Prof. dr hab. Bronisław Marciniak, Rektor Uniwersytetu im. A. Mickiewicza w Poznaniu
Walter Radermacher, Dyrektor Generalny Eurostatu
Prof. dr hab. Janusz Witkowski, Prezes GUS
Marek Woźniak, Marszałek Województwa Wielkopolskiego

Program Kongresu Statystyki Polskiej

Sesje plenarne:

Sesja jubileuszowa

18 kwietnia 2012r. godz.9.00–10.30

Aula Uniwersytetu im. A. Mickiewicza

Przewodniczący **Andrzej Sokołowski** (Uniwersytet Ekonomiczny w Krakowie)

1. **Elżbieta Gołata**, Przewodnicząca Komitetu Organizacyjnego, godz. 9.00–9.10
2. **Przedstawiciele organizatorów Kongresu**
 - **Marian Gorynia**, Rektor Uniwersytetu Ekonomicznego w Poznaniu, godz. 9.10–9.15
 - **Janusz Witkowski**, Prezes Głównego Urzędu Statystycznego, godz. 9.15–9.20
 - **Czesław Domański**, Prezes Polskiego Towarzystwa Statystycznego, godz. 9.20–9.25
 - **Jacek Kowalewski**, Dyrektor Urzędu Statystycznego w Poznaniu, godz. 9.25–9.30
3. **Odczytanie listu Bronisława Komorowskiego**, Prezydenta Rzeczypospolitej Polskiej, godz. 9.30–9.40
4. **Marek Ratajczak**, Wiceminister Nauki i Szkolnictwa Wyższego, godz. 9.40–9.45
5. **Roman Słowiński**, Prezes Oddziału Polskiej Akademii Nauk w Poznaniu, godz. 9.45–9.50
6. **Eugeniusz Gatnar**, Członek Zarządu Narodowego Banku Polskiego, godz. 9.50–9.55
7. **Przedstawiciele władz miasta i regionu**
 - **Piotr Florek**, Wojewoda Wielkopolski, godz. 9.55–10.00
 - **Tomasz Bugajski**, Członek Zarządu Województwa Wielkopolskiego, godz. 10.00–10.05
 - **Tomasz Kayser**, Wiceprezydent Miasta Poznania, godz. 10.05–10.10
8. **Przedstawiciele międzynarodowych organizacji naukowych i instytucji statystycznych**
 - **Marie Bohata**, Zastępca Dyrektora Generalnego Eurostatu, godz. 10.10–10.20
 - **Raymond Chambers**, Prezydent International Association of Survey Statisticians, godz. 10.20–10.25
9. **Odczytanie listów gratulacyjnych - Andrzej Sokołowski**, godz. 10.25–10.30

Sesja plenarna 1

18 kwietnia 2012r. godz.11.00–12.50

Rozwój polskiej myśli statystycznej

Aula Uniwersytetu im. A. Mickiewicza

Przewodniczący Mirosław Krzyśko (Uniwersytet im. Adama Mickiewicza w Poznaniu)

1. **Prezentacja osiągnięć polskiej statystyki, organizator M. Krzyśko** (Uniwersytet im. Adama Mickiewicza w Poznaniu)
 - **Walenty Ostasiewicz** (Uniwersytet Ekonomiczny we Wrocławiu), Rozwój myśli statystycznej w Polsce, godz. 11.00–11.15
 - **Jan Mielniczuk** (Instytut Podstaw Informatyki PAN, Politechnika Warszawska), Wkład Polaków w rozwój statystyki matematycznej i aplikacyjnej, godz. 11.15–11.30
 - **Tadeusz Caliński** (Uniwersytet Przyrodniczy w Poznaniu), Rozwój i osiągnięcia w biometrii, godz. 11.30–11.45
 - **Krzysztof Jajuga** (Uniwersytet Ekonomiczny we Wrocławiu), Rozwój polskiej myśli statystycznej w naukach ekonomicznych, godz. 11.45–12.00
 - **Janusz Witkowski, Tadeusz Walczak, Jan Berger** (Główny Urząd Statystyczny), Statystyka publiczna - rozwój historyczny i aktualne wyzwania, godz. 12.00–12.15
2. **Prezentacja wydawnictw okolicznościowych - Ewa Kowalka**, godz. 12.15–12.25
 - **Uroczyste otwarcie wystawy „Polskie Towarzystwo Statystyczne w latach 1912–2012”**
 - **Uroczyste otwarcie wystawy „Statystyka w Wielkopolsce”**
3. **Uroczyste wręczenie nagród laureatom wielkopolskiego konkursu „Statystyka mnie dotyka”**, godz. 12.25–12.50

Sesja plenarna 2

18 kwietnia 2012 r. godz. 14.35–16.05

Rozwój polskiej myśli statystycznej

Aula Uniwersytetu im. A. Mickiewicza

Przewodniczący Czesław Domański (Uniwersytet Łódzki)

1. **Wybitni polscy statystycy**
 - **Teresa Ledwina** (Instytut Matematyczny PAN), Jerzy Neyman (1894–1981) - **referat zaproszony**, godz. 14.35–14.55
 - **Mirosław Krzyśko** (Uniwersytet im. Adama Mickiewicza w Poznaniu), Jan Czekanowski (1882–1965) - **referat zaproszony**, godz. 14.55–15.15
2. **Stanisława Bartosiewicz** (Wyższa Szkoła Bankowa we Wrocławiu), **Elżbieta Stańczyk** (Urząd Statystyczny we Wrocławiu), Trochę o społeczno-ekonomicznej historii Polski w latach 1918–2010 (na podstawie badań GUS-u), godz. 15.15–15.35
3. **Jan Kordos** (Szkoła Główna Handlowa w Warszawie, Wyższa Szkoła Menedżerska w Warszawie), Współzależność pomiędzy rozwojem teorii i praktyki badań reprezentacyjnych w Polsce, godz. 15.35–15.55

4. Dyskusja, godz. 15.55–16.05

Panel dyskusyjny

18 kwietnia 2012 r. godz. 16.30–18.00

Przyszłość statystyki

Aula Uniwersytetu im. A. Mickiewicza

Przewodniczący **Tadeusz Caliński** (Uniwersytet Przyrodniczy w Poznaniu)

Tadeusz Caliński (Uniwersytet Przyrodniczy w Poznaniu), Powołanie zespołu ds. Uchwały Kongresu Statystyki Polskiej, wprowadzenie do dyskusji, godz. 16.30–16.45

Janina Jóźwiak (Szkola Główna Handlowa w Warszawie), Statystyka ludności, godz. 16.45–16.50

Janusz Witkowski (Główny Urząd Statystyczny), Statystyka społeczna, godz. 16.50–16.55

Tomasz Panek (Szkola Główna Handlowa w Warszawie), Statystyka społeczna, godz. 16.55–17.00

Jerzy Wilkin (Uniwersytet Warszawski), Statystyka gospodarcza, godz. 17.00–17.05

Jan Paradysz (Uniwersytet Ekonomiczny w Poznaniu), Statystyka regionalna, godz. 17.05–17.10

Jerzy Andrzej Moczko (Uniwersytet Medyczny w Poznaniu), Statystyka zdrowia, godz. 17.10–17.15

Mirosław Szreder (Uniwersytet Gdański), Dane statystyczne, godz. 17.15–17.20

Jacek Koronacki (Instytut Podstaw Informatyki PAN, Warszawa), Metodologia badań statystycznych, godz. 17.20–17.25

Dyskusja godz. 17.25–18.00

UROCZYSTE POSIEDZENIE RADY GŁÓWNEJ POLSKIEGO TOWARZYSTWA STATYSTYCZNEGO

19 kwietnia 2012 r. godz. 8.30–9.30

Aula Uniwersytetu im. A. Mickiewicza

Statystyka w biznesie - wymiar praktyczny

19 kwietnia 2012 r. godz. 8.30–9.30

sala 407, Uniwersytet Ekonomiczny w Poznaniu

Sesje paralelne:

Historia i rozwój polskiej statystyki

19 kwietnia 2012 r. godz. 9.50–11.20

sala 418, Uniwersytet Ekonomiczny w Poznaniu

1. **Józef Pociecha** (Uniwersytet Ekonomiczny w Krakowie), Okoliczności powstania Polskiego Towarzystwa Statystycznego w Krakowie i jego pierwszy Prezes - **referat zaproszony**

2. **Bożena Łazowska** (Centralna Biblioteka Statystyczna), Władysław Bortkiewicz (7 VIII 1868–15 II 1931)
3. **Cezary Kukło** (Uniwersytet w Białymstoku), Rodzina i jej gospodarstwo na ziemiach polskich do końca epoki przedprzemysłowej - mity i prawda - **referat zaproszony**
4. **Witold Zdaniewicz** (Instytut Statystyki Kościoła Katolickiego w Warszawie), Wkład Instytutu Statystyki Kościoła Katolickiego w polską statystykę publiczną

Statystyka matematyczna 1

19 kwietnia 2012 r. godz. 9.50–11.20

sala 236, Uniwersytet Ekonomiczny w Poznaniu

1. **Tadeusz Bednarski** (Uniwersytet Wrocławski), Statystyczna analiza przyczyn obciążoności próby w badaniach sondażowych rynku pracy - **referat zaproszony**
2. **Wioletta Grzenda** (Szkola Główna Handlowa w Warszawie), Bayesowski semiparametryczny model Coxa w badaniu determinant pozostawania bez pracy osób młodych
3. **Jan W. Owsiniński** (Instytut Badań Systemowych PAN w Warszawie), Optymalny podział rozkładu empirycznego (i kilka problemów z tym związanych)

Statystyka matematyczna 2

19 kwietnia 2012 r. godz. 15.30–17.00

sala 236, Uniwersytet Ekonomiczny w Poznaniu

1. **Jana Jurečková** (Charles University in Prague), Regression quantiles and their two-step modifications
2. **Daniel Kosiorowski** (Uniwersytet Ekonomiczny w Krakowie), Głębia położenia - rozrzutu w odpornej analizie strumienia danych ekonomicznych
3. **Paweł Kobus** (Szkola Główna Gospodarstwa Wiejskiego w Warszawie), Podejście bayesowskie do estymacji rozkładów cech wybranych jednostek zbiorowości z wykorzystaniem danych zagregowanych

Statystyka matematyczna 3

20 kwietnia 2012 r. godz. 9.00–10.30

sala 0011, Uniwersytet Ekonomiczny w Poznaniu

1. **Jacek Koronacki** (Instytut Podstaw Informatyki PAN w Warszawie), Analiza danych o dużym wymiarze w przypadku małej liczności próby - **referat zaproszony**
2. **Tomasz Górecki, Mirosław Krzyśko** (Uniwersytet im. Adama Mickiewicza w Poznaniu), Wersja jądrowa funkcjonalnych składowych głównych
3. **Bronisław Ceranka, Małgorzata Graczyk** (Uniwersytet Przyrodniczy w Poznaniu), Estymacja całkowitej masy obiektów w układach wagowych

Statystyka matematyczna 4

20 kwietnia 2012 r. godz. 14.30–16.00

sala 0011, Uniwersytet Ekonomiczny w Poznaniu

1. **Zbigniew Szkutnik** (Akademia Górniczo-Hutnicza w Krakowie), Algorytm EM i jego modyfikacje - **referat zaproszony**
2. **Jan Mielniczuk** (Instytut Podstaw Informatyki PAN, Politechnika Warszawska), **Małgorzata Wojtyś** (Politechnika Warszawska), Postselekcyjna estymacja gęstości gradacyjnej
3. **Teresa Ledwina, Grzegorz Wyłupek** (Instytut Matematyczny PAN, Oddział Wrocław), Test dla dwóch prób przy jednostronnych alternatywach

Metoda reprezentacyjna i statystyka małych obszarów 1 English language session

19 kwietnia 2012 r. godz. 9.50–11.20 Aula, Uniwersytet Ekonomiczny w Poznaniu

1. **Malay Ghosh** (University of Florida), Finite population sampling: a model-design synthesis - **invited paper**
2. **Ray Chambers, Gunky Kim** (Wollongong University), Regression analysis using data obtained by probability linking of multiple data sources - **invited paper**
3. **Imbi Traat** (University of Tartu), Domain Estimators Calibrated on Reference Survey

Metoda reprezentacyjna i statystyka małych obszarów 2 English language session

19 kwietnia 2012 r. godz. 17.30–19.00

sala 407, Uniwersytet Ekonomiczny

w Poznaniu

1. **Li-Chun Zhang** (Statistics Norway), Micro calibration for data integration - **invited paper**
2. **Janusz L. Wywiał** (Katowice University of Economics), Estimation of population mean on the basis of a simple sample ordered by auxiliary variable
3. **Wojciech Gamrot** (Katowice University of Economics), On empirical inclusion probabilities
4. **Konrad Furmańczyk, Stanisław Jaworski** (Medical University of Warsaw, Warsaw University of Life Sciences), Change point detection in a sequence of independent observations

Metoda reprezentacyjna i statystyka małych obszarów 3 English language session

20 kwietnia 2012 r. godz. 11.00–12.30

sala 407, Uniwersytet Ekonomiczny

w Poznaniu

1. **Lorenzo Fattorini** (University of Siena), Design-Based Inference on Ecological Diversity - **invited paper**

2. **Tomasz Żądło** (Katowice University of Economics), On prediction of totals for spatially correlated domains
3. **Tomasz Klimanek** (Poznan University of Economics, Statistical Office in Poznan), Using indirect estimation with spatial autocorrelation in social surveys in Poland
4. **Tomasz Józefowski** (Statistical Office in Poznan), Using a SPREE estimator to estimate the number of unemployed across subregions

Metoda reprezentacyjna i statystyka małych obszarów 4 English language session
20 kwietnia 2012 r. godz. 14.30–16.00 Aula, Uniwersytet Ekonomiczny w Poznaniu

1. **Jean-Claude Deville, Daniel Bonnéry, Guillaume Chauvet** (ENSAE), Neyman type optimality for marginal quota sampling - **invited paper**
2. **Sara Franceschi** (University of Tuscia), **Lorenzo Fattorini** (Università di Siena), **Daniela Maffei** (Università di Firenze), Design-based treatment of unit nonresponse by the calibration approach
3. **Marcin Szymkowiak** (Poznan University of Economics), Construction of calibration estimators of total for different distance measures

Statystyka ludności 1

19 kwietnia 2012 r. godz. 15.30–17.00
w Poznaniu

English language session

sala418, Uniwersytet Ekonomiczny

1. **Nico Keilman** (University of Oslo), Challenges for statistics on households and families - **invited paper**
2. **Irena E. Kotowska** (Warsaw School of Economics), Evolving population structures and their relevance for demographic and social change - **invited paper**
3. **Marta Styrc, Anna Matysiak** (Warsaw School of Economics), Socioeconomic status and marital stability

Statystyka ludności 2

20 kwietnia 2012 r. godz. 16.30–18.00
w Poznaniu

English language session

sala418, Uniwersytet Ekonomiczny

1. **Francesco Billari** (Bocconi University, Milan), Challenges for new methods in demographic analysis (tentative) - **invited paper**
2. **Ewa Frątczak** (Warsaw School of Economics), Tradition and modernity in the modeling of demographic events and processes. From J. Graunt's life tables, Lexis Diagram to longitudinal models with fixed and random effects.
3. **Marek Kupiszewski** (Central European Forum for Migration and Population Research), What drives population change? - **invited paper**

Statystyka ludności 3

19 kwietnia 2012 r. godz. 17.30–19.00

sala418, Uniwersytet Ekonomiczny w Poznaniu

1. **Elżbieta Gołata** (Uniwersytet Ekonomiczny w Poznaniu), Spis ludności i prawda
2. **Lucyna Nowak** (Główny Urząd Statystyczny), Rozwój badań statystycznych w obszarze demografii i migracji
3. **Wiktoria Wróblewska** (Szkoła Główna Handlowa w Warszawie), Zmiany maksymalnego trwania życia i długowieczność - wyzwania dla statystyki
4. **Jerzy T. Kowaleski, Anna Majdzińska** (Uniwersytet Łódzki), Starzenie się populacji krajów UE - nieodległa przeszłość i prognoza
5. **Jadwiga Borucka, Joanna Romaniuk, Ewa Frątczak** (Szkoła Główna Handlowa w Warszawie) Trwałość pierwszych związków (małżeństw i kohabitacji) w kohortach urodzeniowych 1950–1970

Statystyka społeczna 1

English language session

19 kwietnia 2012 r. godz. 9.50–11.20

sala407, Uniwersytet Ekonomiczny

w Poznaniu

1. **Achille Lemmi** (Siena University, Dipartimento di Economia Politica), **Panek Tomasz** (Warsaw School of Economics) Dimensions of Poverty. Theory, Models and New Perspectives - **invited paper**
2. **Anna Szukielojć-Bieńkuńska** (Central Statistical Office), Measurement of Poverty and Social Exclusion in the Polish Public Statistics
3. **Stanisław Maciej Kot** (Gdansk University of Technology), Stochastic Equivalence Scales for Poland
4. **Krystyna Hanusik, Urszula Łangowska-Szczęśniak** (University of Opole) Determinants of Level and Structure Diversification of Households' Consumption in Poland

Statystyka społeczna 2

19 kwietnia 2012 r. godz. 15.30–17.00

sala407, Uniwersytet Ekonomiczny w Poznaniu

1. **Walenty Ostasiewicz** (Uniwersytet Ekonomiczny we Wrocławiu), Jakość życia jako przedmiot badań statystycznych - **referat zaproszony**
2. **Jolanta Perek-Białas** (Szkoła Główna Handlowa w Warszawie, Uniwersytet Jagielloński w Krakowie), Sytuacja starszych generacji w Europie Środkowo-Wschodniej na podstawie danych EU-SILC
3. **Anna Szukielojć-Bieńkuńska** (Główny Urząd Statystyczny), Jakość życia w badaniach GUS

4. **Krzysztof Szwarc** (Uniwersytet Ekonomiczny w Poznaniu), Poziom życia ludności a zróżnicowanie sytuacji demograficznej w powiatach województwa wielkopolskiego

Statystyka społeczna 3

English language session

20 kwietnia 2012 r. godz. 9.00–10.30
w Poznaniu

sala407, Uniwersytet Ekonomiczny

1. **Anne Valia Goujon, Ramon Bauer, Samir K.C., Michaela Potančoková** (Vienna Institute of Demography (VID) of the Austrian Academy of Sciences), The Human Capital Puzzle: Why It Is so Hard to Find Good Data on Educational Attainment? - **invited paper**
2. **Marcin Stonawski** (Cracow University of Economics, IIASA), Human Capital and Population Ageing in Poland
3. **Agnieszka Chłoń-Domińczak** (Warsaw School of Economics), The use of indicators in education in the context of development of evidence-based social policy
4. **Dominik Rozkrut** (Statistical Office in Szczecin), Differentiation of innovation strategies

Statystyka społeczna 4

20 kwietnia 2012 r. godz. 14.30–16.00

sala407, Uniwersytet Ekonomiczny w Poznaniu

1. **Urszula Sztanderska** (Uniwersytet Warszawski), Polska statystyka publiczna jako inspiracja i bariera rozwoju badań rynku pracy - **referat zaproszony**
2. **Paweł Ulman** (Uniwersytet Ekonomiczny w Krakowie), Wpływ aktywności zawodowej osób niepełnosprawnych na sytuację ekonomiczną ich gospodarstw domowych
3. **Dominik Śliwicki** (Urząd Statystyczny w Bydgoszczy), Czynniki determinujące poziom wynagrodzenia. Analiza ekonometryczna
4. **Aleksandra Matuszewska-Janica** (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie), Statystyczna analiza nierówności w płacach kobiet i mężczyzn w Polsce na tle Unii Europejskiej w odniesieniu do wieku, wielkości przedsiębiorstwa i grupy zawodowej

Statystyka gospodarcza 1

19 kwietnia 2012 r. godz. 9.50–11.20

sala408, Uniwersytet Ekonomiczny w Poznaniu

1. **Włodzimierz Siwiński** (Uniwersytet Warszawski, Akademia L. Koźmińskiego), Statystyczny obraz kryzysu gospodarczego 2008–2010 - **referat zaproszony**
2. **Eugeniusz Gatnar** (Narodowy Bank Polski), Rola NBP w systemie statystyki publicznej w Polsce
3. **Barbara Pawełek** (Uniwersytet Ekonomiczny w Krakowie), Badanie przydatności wyników testu koniunktury GUS do prognozowania krótkoterminowego zmian aktywności gospodarczej w Polsce na podstawie modelu wektorowej autoregresji

4. **Andrzej Czyżewski, Aleksander Grzelak** (Uniwersytet Ekonomiczny w Poznaniu),
Możliwości wykorzystania statystyki bilansów przepływów międzygałęziowych dla makro-
ekonomicznych ocen w gospodarce

Statystyka gospodarcza 2

19 kwietnia 2012 r. godz. 15.30–17.00

sala408, Uniwersytet Ekonomiczny w Poznaniu

1. **Barbara Liberda** (Uniwersytet Warszawski), Rachunki generacyjne metodą pomiaru
bogactwa kolejnych pokoleń - **referat zaproszony**
2. **Krzysztof Malaga** (Uniwersytet Ekonomiczny w Poznaniu), Dylematy współczesnej
statystyki wzrostu gospodarczego
3. **Stanisław Lewiński vel Iwański, Zofia Wilimowska** (Politechnika Wrocławska),
Struktura finansowa polskich przedsiębiorstw a modelowanie statystyczne
4. **Maria Parlińska, Robert Babiak** (Szkola Główna Gospodarstwa Wiejskiego w War-
szawie), Dokumentowanie i ryzyko asymetrii informacji

Statystyka gospodarcza 3

19 kwietnia 2012 r. godz. 17.30–19.00

sala408, Uniwersytet Ekonomiczny w Poznaniu

1. **Elżbieta Mączyńska** (Instytut Nauk Ekonomicznych PAN, Polskie Towarzystwo Eko-
nomiczne), Statystyka jako źródło danych w badaniach naukowych. Dylematy i niedosto-
sowania - **referat zaproszony**
2. **Monika Natkowska, Jacek Kowalewski** (Urząd Statystyczny w Poznaniu), Wyko-
rzystanie mapy statystyki krótkookresowej w procesie organizacji badań przedsiębiorstw
prowadzonych przez statystykę publiczną
3. **Grażyna Dehnel** (Uniwersytet Ekonomiczny w Poznaniu), Estymacja pośrednia w sta-
tystyce przedsiębiorstw
4. **Aneta Ptak-Chmielewska** (Szkola Główna Handlowa w Warszawie), Uwarunkowania
ryнку przedsiębiorstw w Polsce. Statystyka przeżywalności firm

Statystyka gospodarcza 4

20 kwietnia 2012 r. godz. 11.00–12.30

sala408, Uniwersytet Ekonomiczny w Poznaniu

1. **Andrzej P. Wiatrak** (Uniwersytet Warszawski), Aktualne problemy statystyki rolniczej
- **referat zaproszony**
2. **Renata Bielak** (Główny Urząd Statystyczny), Rola statystyki w procesie kształtowania
polityk rozwoju - priorytety i wyzwania
3. **Grzegorz Kowalewski** (Uniwersytet Ekonomiczny we Wrocławiu), Metodologia ankie-
towych badań koniunktury gospodarczej

4. **Sergiy Gerasymenko** (Wyższa Szkoła Bankowa w Poznaniu, Wydział Zamiejscowy w Chorzowie) **Olena Chupryna** (Narodowy Uniwersytet Karazina w Charkowie), Optymalizacja spisu wskaźników dla oceny porównawczej obiektów socjalno-ekonomicznych

Statystyka regionalna 1

20 kwietnia 2012 r. godz. 9.00–10.30

sala408, Uniwersytet Ekonomiczny w Poznaniu

1. **Jan Paradysz** (Centrum Statystyki Regionalnej, Uniwersytet Ekonomiczny w Poznaniu), Statystyka regionalna: stan, problemy i kierunki rozwoju - **referat zaproszony**
2. **Dominika Rogalińska** (Główny Urząd Statystyczny), Wyzwania statystyki regionalnej na tle debaty o polityce spójności
3. **Danuta Strahl, Małgorzata Markowska** (Uniwersytet Ekonomiczny we Wrocławiu), Możliwości oceny innowacyjności regionów szczebla NUTS 2 w świetle zasobów informacyjnych Eurostatu
4. **Ewa Kamińska-Gawryluk** (Urząd Statystyczny w Białymstoku), Zróżnicowania regionalne w Polsce i wybranych krajach Unii Europejskiej

Statystyka regionalna 2

20 kwietnia 2012 r. godz. 11.00–12.30

sala408, Uniwersytet Ekonomiczny w Poznaniu

1. **Tadeusz Borys** (Uniwersytet Ekonomiczny we Wrocławiu), **Tomasz Potkański** (Związek Miast Polskich), Monitorowanie rozwoju regionalnego i lokalnego
2. **Dorota Doniec** (Urząd Statystyczny w Katowicach), Rachunki regionalne
3. **Bartosz Bartniczak** (Urząd Statystyczny we Wrocławiu, Uniwersytet Ekonomiczny we Wrocławiu), Moduł wskaźników zrównoważonego rozwoju w ramach Banku Danych Lokalnych
4. **Dorota Wyszowska** (Uniwersytet w Białymstoku, Urząd Statystyczny w Białymstoku), Statystyka regionalna jako instrument wsparcia rozwoju polskich województw
5. **Jacek Batóg, Barbara Batóg, Magdalena Mojsiewicz** (Uniwersytet Szczeciński), **Katarzyna Wawrzyniak** (Zachodniopomorski Uniwersytet Technologiczny w Szczecinie), Statystyczny system monitoringu stopnia realizacji strategii rozwoju miasta

Statystyka regionalna 3

20 kwietnia 2012 r. godz. 14.30–16.00

sala408, Uniwersytet Ekonomiczny w Poznaniu

1. **Janusz Dygaszewicz, Magdalena Janczur-Knappek, Amelia Wardzińska-Sharif** (Główny Urząd Statystyczny), Spisy powszechne PSR 2010 i NSP 2011 oraz systemy informacji geograficznej w statystyce publicznej
2. **Paweł Chlebicki** (ESRI Polska), ArcGIS jako potężne narzędzie do wizualizacji i analizy danych statystycznych

3. **Sylwia Filas-Przybył, Maciej Kaźmierczak Dorota Stachowiak** (Urząd Statystyczny w Poznaniu), Wykorzystanie danych ze źródeł administracyjnych w statystyce miast
4. **Robert Buciak, Marek Pieniążek** (Główny Urząd Statystyczny, Uniwersytet Warszawski), Klasyfikacja przestrzenna obszarów wiejskich w Polsce
5. **Roma Ryś-Jurek** (Uniwersytet Przyrodniczy w Poznaniu), Zróżnicowanie regionalne produkcji rolniczej w krajach UE-27

Analiza i klasyfikacja danych 1

English language session

19 kwietnia 2012 r. godz. 15.30–17.00 Aula, Uniwersytet Ekonomiczny w Poznaniu

1. **Reinhold Decker** (Universität Bielefeld), Model-based analysis of online consumer reviews - methods and applications - **invited paper**
2. **Patrick Groenen** (Erasmus University Rotterdam), The support vector machine as a powerful tool for binary classification - **invited paper**
3. **Andrzej Sokołowski** (Cracow University of Economics), **Beata Basiura** (AGH University of Science and Technology in Cracow) - Ward's agglomerative method

Analiza i klasyfikacja danych 2

English language session

19 kwietnia 2012 r. godz. 17.30–19.00 Aula, Uniwersytet Ekonomiczny w Poznaniu

1. **Wojciech Krzanowski** (University of Exeter), Classification: some old principles applied to new problems - **invited paper**
2. **Justyna Wilk** (Wrocław University of Economics), Symbolic approach in regional analyses
3. **Marcin Pełka** (Wrocław University of Economics), Ensemble approach for clustering of interval-valued symbolic data
4. **Justyna Brzezińska** (Katowice University of Economics), Independence analysis of nominal data with the use of log-linear models in R

Analiza i klasyfikacja danych 3

20 kwietnia 2012 r. godz. 9.00–10.30

Aula, Uniwersytet Ekonomiczny w Poznaniu

1. **Dorota Rozmus** (Uniwersytet Ekonomiczny w Katowicach), Podejście zagregowane w taksonomii
2. **Aleksandra Łuczak, Feliks Wysocki** (Uniwersytet Przyrodniczy w Poznaniu), Ocena poziomu rozwoju społeczno-gospodarczego powiatów województwa wielkopolskiego
3. **Kamila Migdał-Najman** (Uniwersytet Gdański), Konstrukcja hybrydowej, samouczącej się sieci neuronowej typu SOM-GNG

4. **Krzysztof Najman** (Uniwersytet Gdański), Sieci samouczące się w analizie skupień
5. **Piotr Krasucki** (Instytut Spraw Publicznych), **Władysław Wiesław Łagodziński** (Polskie Towarzystwo Statystyczne), Zintegrowany system danych statystycznych z ochrony zdrowia - niezbędne minimum

Analiza i klasyfikacja danych 4

20 kwietnia 2012 r. godz. 11.00–12.30

Aula, Uniwersytet Ekonomiczny w Poznaniu

1. **Małgorzata Markowska, Bartłomiej Jefmański** (Uniwersytet Ekonomiczny we Wrocławiu), Ocena dynamiki i kierunku zmian rozwoju inteligentnej specjalizacji regionów europejskich z zastosowaniem klasyfikacji rozmytej
2. **Katarzyna Kopczewska** (Uniwersytet Warszawski), Regiony w przestrzeni. Czy lokalizacja i dostępność mają znaczenie dla rozwoju gospodarczo-społecznego?
3. **Robert Pietrzykowski** (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie), Wykorzystanie regresji kwantylowej w analizach przestrzennych cen ziemi rolniczej
4. **Marzena Piotrowska-Trybull, Stanisław Sirko** (Akademia Obrony Narodowej), Wpływ jednostki wojskowej na rozwój gminy
5. **Joanna Zyprych-Walczak, Alicja Szabelska, Idzi Siatkowski** (Uniwersytet Przyrodniczy w Poznaniu), Metody selekcji genów uwzględniające korelację pomiędzy genami

Dane statystyczne 1

19 kwietnia 2012 r. godz. 17.30–19.00

sala 236, Uniwersytet Ekonomiczny w Poznaniu

1. **Jan Zawadzki, Maria Szumksta-Zawadzka** (Zachodniopomorski Uniwersytet Technologiczny w Szczecinie), O metodach prognozowania brakujących danych w szeregach czasowych z wahaniami okresowymi (sezonowymi)
2. **Jadwiga Suchecka, Emilia Modranka** (Uniwersytet Łódzki), Statystyka przestrzenna - metody i zastosowania
3. **Ewa-Zofia Frątczak, Adam Korczyński** (Szkoła Główna Handlowa w Warszawie), Tradycja i współczesność w metodach imputacji danych. Przegląd teorii - wybrane przykłady zastosowań
4. **Artur Mikulec, Aleksandra Kupis-Fijałkowska** (Uniwersytet Łódzki), Empiryczna analiza efektywności kryterium Mojeny i Wisharta w analizie skupień
5. **Cezary Głowiński, Kamil Konikiewicz** (SAS Institute sp. z o.o.), Wykorzystanie sieci społecznych w modelowaniu zachowań klientów firm telekomunikacyjnych

Dane statystyczne 2

20 kwietnia 2012 r. godz. 11.00–12.30

sala 0011, Uniwersytet Ekonomiczny w Poznaniu

1. **Grażyna Trzpiot, Agnieszka Orwat-Acedańska** (Uniwersytet Ekonomiczny w Katowicach), Klasyfikacja funduszy inwestycyjnych ze względu na styl zarządzania - podejście regresji kwantylowej
2. **Tadeusz Kufel** (Uniwersytet Mikołaja Kopernika w Toruniu), **Marcin Błażejowski, Paweł Kufel** (Wyższa Szkoła Bankowa w Toruniu), Budowa banków danych jako podstawa analiz statystyczno-ekonometrycznych w oprogramowaniu GRETL
3. **Zbigniew Augustyniak** (Urząd Statystyczny w Poznaniu), Prezentacja danych statystycznych w wersji on-line jako element realizacji projektu SISP
4. **Arleta Olbrot-Brzezińska** (Urząd Statystyczny w Poznaniu), Społeczne funkcje informacji statystycznej i odpowiedzialność statystyki publicznej
5. **Katarzyna Wawrzyniak** (Zachodniopomorski Uniwersytet Technologiczny w Szczecinie), Ocena stopnia realizacji celów głównych Strategii Rozwoju Kraju według województw z wykorzystaniem analizy korespondencji

Polska i Ukraina w przestrzeni europejskiej 20 kwietnia 2012 r. godz. 14.30–16.00

sala 418, Uniwersytet Ekonomiczny w Poznaniu

1. **Rusłan Motoryn** (Ukraiński Państwowy Uniwersytet Finansów i Handlu Międzynarodowego w Kijowie, Politechnika Lubelska), **Iryna Motoryna** (Akademia Zarządzania Finansami w Kijowie), **Tetiana Motoryna** (Kijowski Narodowy Uniwersytet im. Tarasa Szewczenko), Porównanie wskaźników wykorzystania oszczędności gospodarstw domowych Ukrainy i Polski
2. **Mieczysław Kowerski** (Wyższa Szkoła Zarządzania i Administracji w Zamościu), Badania nastrojów gospodarczych przedsiębiorców i konsumentów na poziomie regionu. Przykład województwa lubelskiego
3. **Andrzej Młodak** (Urząd Statystyczny w Poznaniu, Państwowa Wyższa Szkoła Zawodowa im. Prezydenta Stanisława Wojciechowskiego w Kaliszu), Modernizacja statystyki działalności gospodarczej w wymiarze paneuropejskim
4. **Karol Andrzejczak** (Politechnika Poznańska), Ewolucja i granica wzrostu wskaźnika nasycenia samochodami osobowymi w skali globalnej i w Polsce
5. **Dorota Kwiatkowska-Ciotucha, Urszula Załuska** (Uniwersytet Ekonomiczny we Wrocławiu), Ocena wykorzystania środków z europejskiego funduszu społecznego w poszczególnych regionach Polski w obecnym okresie programowania

Statystyka zdrowia

20 kwietnia 2012 r. godz. 9.00–10.30

sala 418, Uniwersytet Ekonomiczny w Poznaniu

1. **Agnieszka Dyzmann-Sroka, Andrzej Roszak, Maciej Trojanowski** (Wielkopolskie Centrum Onkologii), **Jerzy Moczko** (Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu), Metody podniesienia jakości danych rejestrów nowotworów złośliwych
2. **Urszula Wojciechowska, Joanna Didkowska** (Centrum Onkologii - Instytut im. Marii Skłodowskiej-Curie w Warszawie), Nowotwory złośliwe w Polsce jako problem zdrowia publicznego. Współczesne narzędzia do mierzenia zjawiska
3. **Barbara Kołodziejczak, Magdalena Roszak** (Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu), Rapid e-learning w biostatystyce
4. **Magdalena Roszak, Barbara Kołodziejczak** (Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu), E-ewaluacja wiedzy statystycznej studentów medycyny
5. **Anna Wierzbicka** (Instytut Statystyki i Demografii), Analiza taksonomiczna stanu zdrowia publicznego Polski na tle wybranych krajów europejskich

Statystyka sportu i turystyki

20 kwietnia 2012 r. godz. 16.30–18.00

sala 0011, Uniwersytet Ekonomiczny w Poznaniu

1. **Ryszard Walaszczyk** (Ministerstwo Sportu i Turystyki), System Informacji Sportowej - referat zaproszony
2. **Barbara Liberda** (Uniwersytet Warszawski), **Łucja Tomaszewicz, Iwona Świeczewska** (Uniwersytet Łódzki), Metodologia rachunków satelitarnych na podstawie rachunku satelitarnego sportu dla Polski
3. **Iwona Bąk** (Zachodniopomorski Uniwersytet Technologiczny w Szczecinie), Statystyczna analiza aktywności turystycznej seniorów w Polsce
4. **Marcin Błazejowski** (Wyższa Szkoła Bankowa w Toruniu), Algorytmy modelowania i prognozowania na przykładzie rynku turystycznego

Klasyfikacje i standardy

English language session

20 kwietnia 2012 r. godz. 16.30–18.00
w Poznaniu

sala 418, Uniwersytet Ekonomiczny

1. **Włodzimierz Okrasa** (Central Statistical Office, Cardinal Stefan Wyszyński University in Warsaw), Statistics and sociology: Mutually-supportive development from the perspective of interdisciplinaryisation of social research
2. **Éva Laczka** (Hungarian Central Statistical Office), The role of Population and Housing and Agricultural Censuses in the national statistical systems

3. **Grzegorz Krzykowski, Marcin Kalinowski** (Gdańsk School of Banking), Selection and screening of data as a key element of statistical research on the financial markets

Wnioskowanie statystyczne

20 kwietnia 2012 r. godz. 16.30–18.00

Aula, Uniwersytet Ekonomiczny w Poznaniu

1. **Mirosław Krzyśko, Łukasz Waszak** (Uniwersytet im. Adama Mickiewicza w Poznaniu), Jądrowe korelacje kanoniczne
2. **Grzegorz Kończak** (Uniwersytet Ekonomiczny w Katowicach), Wnioskowanie statystyczne dla danych w wielowymiarowych tablicach wielodzielczych
3. **Joanna Kisielińska** (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie), Dokładna metoda bootstrapowa na przykładzie estymacji średniej i wariancji
4. **Ewa Meller** (Urząd Statystyczny w Poznaniu), Estymacja dla małych obszarów dla wybranych charakterystyk rynku pracy
5. **Łukasz Wawrowski** (Urząd Statystyczny w Poznaniu), Analiza ubóstwa w przekroju powiatów w województwie wielkopolskim z wykorzystaniem metod statystyki małych obszarów

Analiza rozwoju gospodarczego i społecznego
16.30–18.00

20 kwietnia 2012 r. godz.

sala 407, Uniwersytet Ekonomiczny w Poznaniu

1. **Elżbieta Stańczyk** (Urząd Statystyczny we Wrocławiu), Kwalifikacje zawodowe osób bezrobotnych w świetle danych GUS. Międzywojewódzka analiza porównawcza
2. **Józef Hozer, Marta Hozer-Koćmiel** (Uniwersytet Szczeciński), Quantum satis dla firm w Polsce w 2012 roku
3. **Janusz Kornecki, Piotr Cmela** (Urząd Statystyczny w Łodzi), Strategiczne dylematy wsparcia rozwoju mikroprzedsiębiorstw w Polsce
4. **Rafał Nowakowski** (Urząd Statystyczny we Wrocławiu), Wybory samorządowe na przykładzie stolic regionów dolnośląskiego, małopolskiego i wielkopolskiego - 20 lat doświadczeń
5. **Jacek Białek** (Uniwersytet Łódzki), Szczególny przypadek pewnej ogólnej formuły indeksu cen

Sesja plakatowa

19 kwietnia 2012 r. godz. 17.00–17.30

Hol II piętro, Uniwersytet Ekonomiczny w Poznaniu

1. **Beata Bal-Domańska** (Uniwersytet Ekonomiczny we Wrocławiu, WGRiT w Jeleniej Górze), Convergence processes modelling of the European regional space
2. **Maciej Beręsewicz** (Uniwersytet Ekonomiczny w Poznaniu), Internetowe źródła danych w świetle spisu powszechnego na przykładzie rynku nieruchomości w Poznaniu
3. **Anna Błaczowska** (Wyższa Szkoła Bankowa we Wrocławiu), **Alicja Grześkowiak** (Uniwersytet Ekonomiczny we Wrocławiu), Uwarunkowania demograficzno-gospodarcze procesu kształcenia gimnazjalnego na Dolnym Śląsku i Opolszczyźnie w latach 2003–2010
4. **Hanna Dudek, Joanna Landmesser** (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie), Satysfakcja z osiągniętych dochodów a względna deprywacja
5. **Alicja Ganczarek-Gamrot** (Uniwersytet Ekonomiczny w Katowicach), Wielowymiarowe modele FIGARCH w ocenie ryzyka na polskim rynku energii elektrycznej
6. **Hanna Gruchociak** (Uniwersytet Ekonomiczny w Poznaniu), Delimitacja lokalnych rynków pracy w Polsce
7. **Marta Małecka** (Uniwersytet Łódzki), Ocena wariancji estymatorów wartości narażonej na ryzyko (VaR)
8. **Magdalena Okupniak** (Uniwersytet Ekonomiczny w Poznaniu), Zastosowanie analizy wewnątrz- i między-blokowej w statystyce gospodarczej
9. **Dariusz Parys** (Uniwersytet Łódzki), Uwagi o zastosowaniach testowania wielokrotnego w naukach ekonomicznych
10. **Elżbieta Paszko, Anna Maria Jurek** (Uniwersytet Łódzki), Zastosowanie CQI - ciągłego doskonalenia jakości w procesie poprawy metod nauczania
11. **Agnieszka Pobłocka** (Uniwersytet Gdański), Polski rynek ubezpieczeń w latach 1991–2010
12. **Wojciech Roszka** (Uniwersytet Ekonomiczny w Poznaniu), Wykorzystanie metod masowej imputacji do integracji baz danych pochodzących z różnych źródeł
13. **Danuta Rozpędowska-Matraszek** (Państwowa Wyższa Szkoła Zawodowa w Skierniewicach), Ocena skuteczności funkcjonowania opieki zdrowotnej w Polsce - analiza według grup powiatów podobnych
14. **Elżbieta Sobczak** (Uniwersytet Ekonomiczny we Wrocławiu), Pracujący wg sektorów intensywności działalności B+R w krajach UE - analiza przestrzenno-strukturalna
15. **Agnieszka Stanimir** (Uniwersytet Ekonomiczny we Wrocławiu), Zróżnicowane techniki prezentacji zmiennych niemetrycznych
16. **Marta Styrc** (Szkoła Główna Handlowa w Warszawie), Czynniki wpływające na stabilność pierwszych małżeństw w Polsce

The Programme of the Congress of Polish Statistics

Plenary sessions:

Anniversary session

18 April 2012, 9.00–10.30

The Plenary Hall of Adam Mickiewicz University

Chairperson **Andrzej Sokołowski** ((Cracow University of Economics))

1. **Elżbieta Gołata**, Chairperson of the Organizing Committee of the Congress, 9.00–9.10
2. **Representatives of the Organizers of the Congress**
 - **Marian Gorynia**, Rector of the Poznan University of Economics, 9.10–9.15
 - **Janusz Witkowski**, President of the Central Statistical Office, 9.15–9.20
 - **Czesław Domański**, President of the Polish Statistical Association, 9.20–9.25
 - **Jacek Kowalewski**, Director of the Statistical Office in Poznan, 9.25–9.30
3. **Reading Bronisław Komorowski's letter**, The President of the Republic of Poland, 9.30–9.40 4
4. **Marek Ratajczak**, Deputy Minister of Science and Higher Education, 9.40–9.45
5. **Roman Słowiński**, Chairman of the Department of the Polish Academy of Sciences in Poznan, 9.45–9.50
6. **Eugeniusz Gatnar**, Member of the Board of the National Bank of Poland, 9.50–9.55
7. **Representatives of the Authorities of the City and Region**
 - **Piotr Florek**, Voivode of the Wielkopolskie Voivodeship, 9.55–10.00
 - **Tomasz Bugajski**, Member of the Board of the Wielkopolska Region, 10.00–10.05
 - **Tomasz Kayser**, Deputy Major of Poznan, 10.05–10.10
8. **Representatives of International Scientific Organizations and Statistical Institutions**
 - **Marie Bohata**, Deputy Director General of Eurostat, 10.10–10.20
 - **Raymond Chambers**, President of the International Association of Survey Statisticians, 10.20–10.25
9. **Reading letters of congratulations - Andrzej Sokołowski**, 10.25–10.30

Plenary session 1

18 April 2012, 11.00–12.50

Development of Polish statistical research The Plenary Hall of Adam Mickiewicz University

Chairperson Mirosław Krzyśko (Adam Mickiewicz University, Poznań)

1. Presentation of the achievements of Polish statistics

- **Walenty Ostasiewicz** (Wrocław University of Economics), Development of statistical research in Poland, 11.00–11.15
- **Jan Mielniczuk** (Institute of Computer Science, Polish Academy of Sciences, Warsaw University of Technology), Polish Contributions in the development of mathematical and applied statistics, 11.15–11.30
- **Tadeusz Caliński** (Poznań University of Life Sciences), Development and achievements in biometry, 11.30–11.45
- **Krzysztof Jajuga** (Wrocław University of Economics), Development of Polish statistical research in economic sciences, 11.45–12.00
- **Janusz Witkowski, Tadeusz Walczak, Jan Berger** (Central Statistical Office), The Polish official statistics its historical development and current challenges, 12.00–12.15

2. Presentation of anniversary publications - Ewa Kowalka, 12.15–12.25

- **The opening ceremony of the exhibition "Polish Statistical Association in the period between 1912 and 2012"**
- **The opening ceremony of the exhibition "Statistics in Wielkopolska"**

3. The award ceremony of the provincial competition entitled "Statistics concerns me", 12.25–12.50

Plenary session 2

18 April 2012, 14.35–16.05

Development of Polish statistical research The Plenary Hall of Adam Mickiewicz University

Chairperson Czesław Domański (University of Lodz)

1. Eminent Polish statisticians

- **Teresa Ledwina** (Institute of Mathematics, Polish Academy of Sciences, Wrocław), Jerzy Neyman (1894-1981) - invited paper, 14.35–14.55
- **Mirosław Krzyśko** (Adam Mickiewicz University, Poznań), Jan Czekanowski (1882–1965) - invited paper, 14.55–15.15

2. Stanisława Bartosiewicz (Wrocław School of Banking), Elżbieta Stańczyk (Statistical Office in Wrocław), Socio-economic history of Poland in years 1918–2010, 15.15–15.35

3. **Jan Kordos** (Warsaw School of Economics, Warsaw Management Academy), Interplay between sample survey theory and practice in Poland, 15.35–15.55
4. **Discussion**, 15.55–16.05

Discussion Panel

18 April 2012, 16.30–18.00

The Future of Statistics

The Plenary Hall of Adam Mickiewicz University

Chairperson Tadeusz Caliński (Poznań University of Life Sciences)

Tadeusz Caliński (Poznań University of Life Sciences), Appointing a team responsible for preparing a declaration of the Congress of Polish Statistics, introduction to the discussion, 16.30–16.45

Janina Jóźwiak (Warsaw School of Economics), Population statistics, 16.45–16.50

Janusz Witkowski (Central Statistical Office), Social statistics, 16.50–16.55

Tomasz Panek (Warsaw School of Economics), Social statistics, 16.55–17.00

Jerzy Wilkin (University of Warsaw), Economic statistics, 17.00–17.05

Jan Paradysz (Poznań University of Economics), Regional statistics, 17.05–17.10

Jerzy Andrzej Moczko (Poznań University of Medical Sciences), Health statistics, 17.10–17.15

Mirosław Szreder (University of Gdańsk), Statistical data, 17.15–17.20

Jacek Koronacki (Institute of Computer Science, Polish Academy of Sciences, Warsaw), Methodology of statistical research, 17.20–17.25

Discussion, 17.25–18.00

ANNIVERSARY SESSION OF THE MAIN COUNCIL OF THE POLISH STATISTICAL ASSOCIATION

19 April 2012, 8.30–9.30

Aula, Poznań University of Economics

Statistics in business - practical dimension

19 April 2012, 8.30–9.30

room 407, Poznań University of Economics

Parallel Sessions:

History and Development of Polish Statistics

19 April 2012, 9.50–11.20

room 418, Poznań University of Economics

1. **Józef Pociecha** (Cracow University of Economics), Circumstances surrounding the foundation of the Polish Statistical Association and its first President - **invited paper**

2. **Bożena Łazowska** (Central Statistical Library), The legacy of professor Władysław Bortkiewicz (7 VIII 1868–15 II 1931)
3. **Cezary Kukło** (University of Białystok), Myths and the truth on the family and its household on the territory of Poland by the end of the pre-industrial era - **invited paper**
4. **Witold Zdaniewicz** (Institute of Catholic Church Statistics SAC, Warsaw), The contribution of the Institute of Catholic Church Statistics to the Polish public statistics

Mathematical statistics 1

19 April 2012, 9.50–11.20

room 236, Poznan University of Economics

1. **Tadeusz Bednarski** (University of Wrocław), Statistical analysis of nonresponse in unemployment duration studies - **invited paper**
2. **Wioletta Grzenda** (Warsaw School of Economics), Bayesian semiparametric Cox model in the analysis of unemployment duration determinants of youth
3. **Jan W. Owsiniński** (Systems Research Institute, Polish Academy of Sciences, Warsaw), On the optimal division of an empirical distribution (and some related problems)

Mathematical statistics 2

19 April 2012, 15.30–17.00

room 236, Poznan University of Economics

1. **Jana Jurečková** (Charles University in Prague), Regression quantiles and their two-step modifications
2. **Daniel Kosiorowski** (Cracow University of Economics), Location-scale depth in robust analysis of economic data stream
3. **Paweł Kobus** (Warsaw University of Life Sciences), The use of the Bayesian approach to estimate distributions of variables for selected units in a population using aggregate data

Mathematical statistics 3

20 April 2012, 9.00–10.30

room 0011, Poznan University of Economics

1. **Jacek Koronacki** (Institute of Computer Science, Polish Academy of Sciences, Warsaw), Statistical challenges of high dimensional small sample size problems - **invited paper**
2. **Tomasz Górecki, Mirosław Krzyśko** (Adam Mickiewicz University, Poznań), Kernel version of the functional principal components
3. **Bronisław Ceranka, Małgorzata Graczyk** (Poznań University of Life Sciences), Estimation of total weight of objects in weighting design

Mathematical statistics 4

20 April 2012, 14.30–16.00

room 0011, Poznan University of Economics

1. **Zbigniew Szkutnik** (AGH University of Science and Technology in Cracow), EM algorithm and its modifications - **invited paper**
2. **Jan Mielniczuk** (Institute of Computer Science, Polish Academy of Sciences, Warsaw University of Technology), Małgorzata Wojtyś (Warsaw University of Technology), Post-model-selection estimators of grade density
3. **Teresa Ledwina, Grzegorz Wyłupek** (Institute of Mathematics, Polish Academy of Sciences, Wrocław), Two-sample test against one-sided alternatives

Survey sampling and small area statistics 1

English language session

19 April 2012, 9.50–11.20

Aula, Poznan University of Economics

1. **Malay Ghosh** (University of Florida), Finite population sampling: a model-design synthesis - **invited paper**
2. **Ray Chambers, Gunky Kim** (Wollongong University), Regression analysis using data obtained by probability linking of multiple data sources - **invited paper**
3. **Imbi Traat** (University of Tartu), Domain Estimators Calibrated on Reference Survey

Survey sampling and small area statistics 2

English language session

19 April 2012, 17.30–19.00

room 407, Poznan University of Economics

1. **Li-Chun Zhang** (Statistics Norway), Micro calibration for data integration - **invited paper**
2. **Janusz L. Wywiał** (Katowice University of Economics), Estimation of population mean on the basis of a simple sample ordered by auxiliary variable
3. **Wojciech Gamrot** (Katowice University of Economics), On empirical inclusion probabilities
4. **Konrad Furmańczyk, Stanisław Jaworski** (Medical University of Warsaw, Warsaw University of Life Sciences), Change point detection in a sequence of independent observations

Survey sampling and small area statistics 3

English language session

20 April 2012, 11.00–12.30

room 407, Poznan University of Economics

1. **Lorenzo Fattorini** (University of Siena), Design-Based Inference on Ecological Diversity - **invited paper**
2. **Tomasz Żądło** (Katowice University of Economics), On prediction of totals for spatially correlated domains

3. **Tomasz Klimanek** (Poznan University of Economics, Statistical Office in Poznan), Using indirect estimation with spatial autocorrelation in social surveys in Poland
4. **Tomasz Józefowski** (Statistical Office in Poznan), Using a SPREE estimator to estimate the number of unemployed across subregions

Survey sampling and small area statistics 4

English language session

20 April 2012, 14.30 - 16.00

Aula, Poznan University of Economics

1. **Jean-Claude Deville, Daniel Bonn  ry, Guillaume Chauvet** (ENSAE), Neyman type optimality for marginal quota sampling - **invited paper**
2. **Sara Franceschi** (University of Tuscia), **Lorenzo Fattorini** (Universit   di Siena), **Daniela Maffei** (Universita di Firenze), Design-based treatment of unit nonresponse by the calibration approach
3. **Marcin Szymkowiak** (Poznan University of Economics), Construction of calibration estimators of total for different distance measures

Population statistics 1

English language session

19 April 2012, 15.30–17.00

room 418, Poznan University of Economics

1. **Nico Keilman** (University of Oslo), Challenges for statistics on households and families - **invited paper**
2. **Irena E. Kotowska** (Warsaw School of Economics), Evolving population structures and their relevance for demographic and social change - **invited paper**
3. **Marta Styrc, Anna Matysiak** (Warsaw School of Economics), Socioeconomic status and marital stability

Population statistics 2

English language session

20 April 2012, 16.30–18.00

room 418, Poznan University of Economics

1. **Francesco Billari** (Bocconi University, Milan), Challenges for new methods in demographic analysis (tentative) - **invited paper**
2. **Ewa Fr  czak** (Warsaw School of Economics), Tradition and modernity in the modeling of demographic events and processes. From J. Graunt's life tables, Lexis Diagram to longitudinal models with fixed and random effects.
3. **Marek Kupiszewski** (Central European Forum for Migration and Population Research), What drives population change? - **invited paper**

Population statistics 3

19 April 2012, 17.30–19.00

room 418, Poznan University of Economics

1. **El  bieta Go  ata** (Poznan University of Economics), Population census and truth

2. **Lucyna Nowak** (Central Statistical Office), Development of statistical research in the field of demography and migration
3. **Wiktoria Wróblewska** (Warsaw School of Economics), Changes in life expectancy and longevity - challenges for statistics
4. **Jerzy T. Kowaleski, Anna Majdzińska** (University of Lodz), Population ageing of the European Union countries - the past and the perspectives
5. **Jadwiga Borucka, Joanna Romaniuk, Ewa Frątczak** (Warsaw School of Economics), Stability of first relationships (marriages and cohabitations) in birth cohorts 1950–1970

Social Statistics 1

English language session

19 April 2012, 9.50–11.20

room 407, Poznan University of Economics

1. **Achille Lemmi** (Siena University, Dipartimento di Economia Politica), **Panek Tomasz** (Warsaw School of Economics) Dimensions of Poverty. Theory, Models and New Perspectives - **invited paper**
2. **Anna Szukielojć-Bieńkuńska** (Central Statistical Office), Measurement of Poverty and Social Exclusion in the Polish Public Statistics
3. **Stanisław Maciej Kot** (Gdansk University of Technology), Stochastic Equivalence Scales for Poland
4. **Krystyna Hanusik, Urszula Łangowska-Szczęśniak** (University of Opole) Determinants of Level and Structure Diversification of Households' Consumption in Poland

Social Statistics 2

19 April 2012, 15.30–17.00

room 407, Poznan University of Economics

1. **Walenty Ostasiewicz** (Wrocław University of Economics), Life quality as a subject of statistical research - **invited paper**
2. **Jolanta Perek-Białas** (Warsaw School of Economics, Jagiellonian University in Krakow), Situation of older generations based on EU-SILC in selected Central and Eastern European countries
3. **Anna Szukielojć-Bieńkuńska** (Central Statistical Office), Quality of life in the surveys of the Central Statistical Office of Poland
4. **Krzysztof Szwarc** (Poznan University of Economics), The level of life and the diversity of the demographic situation in the districts of the Wielkopolska voivodship

Social Statistics 3

English language session

20 April 2012, 9.00–10.30

room 407, Poznan University of Economics

1. **Anne Valia Goujon, Ramon Bauer, Samir K.C., Michaela Potančoková** (Vienna Institute of Demography (VID) of the Austrian Academy of Sciences), The Human Capital Puzzle: Why It Is so Hard to Find Good Data on Educational Attainment? - **invited paper**
2. **Marcin Stonawski** (Cracow University of Economics, IIASA), Human Capital and Population Ageing in Poland
3. **Agnieszka Chłoń-Domińczak** (Warsaw School of Economics), The use of indicators in education in the context of development of evidence-based social policy
4. **Dominik Rozkrut** (Statistical Office in Szczecin), Differentiation of innovation strategies

Social Statistics 4

20 April 2012, 14.30–16.00

room 407, Poznan University of Economics

1. **Urszula Sztanderska** (University of Warsaw), Polish public statistics as a source of inspiration for and an obstacle to the development of labour market research - **invited paper**
2. **Paweł Ulman** (Cracow University of Economics), The influence of economic activity of the disabled on the economic situation of their households
3. **Dominik Śliwicki** (Statistical Office in Bydgoszcz), Factors determining the level of wages. The econometric analysis
4. **Aleksandra Matuszewska-Janica** (Warsaw University of Life Sciences), Wages inequalities between men and women in the EU: differences analysis by age, size of the enterprise and occupation

Economic statistics 1

19 April 2012, 9.50–11.20

room 408, Poznan University of Economics

1. **Włodzimierz Siwiński** (University of Warsaw, Kozminski University), Statistical image of economic crisis 2008–2010 - **invited paper**
2. **Eugeniusz Gatnar** (National Bank of Poland), The role of the National Bank of Poland in the system of Polish public statistics
3. **Barbara Pawełek** (Cracow University of Economics), Study of suitability of CSO business tendency survey to short-term forecasting of changes in economic activity in Poland based on vector autoregression model
4. **Andrzej Czyżewski, Aleksander Grzelak** (Poznan University of Economics), The possibilities of using the statistic of input-output in macroeconomics evaluation of the economy

Economic statistics 2

19 April 2012, 15.30–17.00

room 408, Poznan University of Economics

1. **Barbara Liberda** (University of Warsaw), Generational accounting - a method of measuring net wealth of each generation - **invited paper**
2. **Krzysztof Malaga** (Poznan University of Economics), The dilemmas of contemporary economic growth statistics
3. **Stanisław Lewiński vel Iwański, Zofia Wilimowska** (Wrocław University of Technology), Financial structure of Polish companies: statistical modelling
4. **Maria Parlińska, Robert Babiak** (Warsaw University of Life Sciences), Screening and risks of information asymmetry

Economic statistics 3

19 April 2012, 17.30–19.00

room 408, Poznan University of Economics

1. **Elżbieta Mączyńska** (Institute of Economics, Polish Academy of Sciences, Polish Economic Society), Statistics as a data source in scientific research. Dilemmas and inadequacy - **invited paper**
2. **Monika Natkowska, Jacek Kowalewski** (Statistical Office in Poznan), Using a map of short-term statistics in the process of organizing business surveys conducted by public statistics
3. **Grażyna Dehnel** (Poznan University of Economics), Indirect micro-estimation in business statistics
4. **Aneta Ptak-Chmielewska** (Warsaw School of Economics), Conditions of enterprises' market in Poland. Enterprises' survival statistics

Economic statistics 4

20 April 2012, 11.00–12.30

room 408, Poznan University of Economics

1. **Andrzej P. Wiatrak** (University of Warsaw), Contemporary problems of agricultural statistics - **invited paper**
2. **Renata Bielak** (Central Statistical Office), The role of statistics in the process of planning development policies - priorities and challenges
3. **Grzegorz Kowalewski** (Wrocław University of Economics), Methodology of business tendency surveys
4. **Sergiy Gerasymenko** (Poznan School of Banking, the Faculty in Chorzow), **Olena Chupryna** (V.N. Karazin Kharkiv National University), Optimization of the set of indexes for the comparative estimation of the social-economic objects

Regional statistics 1

20 April 2012, 9.00–10.30

room 408, Poznan University of Economics

1. **Jan Paradysz** (Centre of Regional Statistics, Poznan University of Economics), Regional statistics: condition, problems and development directions - **invited paper**
2. **Dominika Rogalińska** (Central Statistical Office), Challenges of regional statistics in the context of the debate about policy coherence
3. **Danuta Strahl, Małgorzata Markowska** (Wrocław University of Economics), The possibility of NUTS 2 level regions innovation assessment in the perspective of Eurostat information resources
4. **Ewa Kamińska-Gawryluk** (Statistical Office in Białystok), Regional Differences in Poland and in Selected Countries of the European Union

Regional statistics 2

20 April 2012, 11.00–12.30

room 408, Poznan University of Economics

1. **Tadeusz Borys** (Wrocław University of Economics), **Tomasz Potkański** (Association of Polish Cities), Monitoring of regional and local development
2. **Dorota Doniec** (Statistical Office in Katowice), Regional accounts
3. **Bartosz Bartniczak** (Statistical Office in Wrocław, Wrocław University of Economics), Sustainable development indicators module in local data bank
4. **Dorota Wyszowska** (University of Białystok, Statistical Office in Białystok), Regional statistics as a tool of supporting the development of Polish voivodships
5. **Jacek Batóg, Barbara Batóg, Magdalena Mojsiewicz** (University of Szczecin), **Katarzyna Wawrzyniak** (West Pomeranian University of Technology, Szczecin), Statistical monitoring system of implementation of the city strategy of development

Regional statistics 3

20 April 2012, 14.30–16.00

room 408, Poznan University of Economics

1. **Janusz Dygaszewicz, Magdalena Janczur-Knappek, Amelia Wardzińska-Sharif** (Central Statistical Office), National censuses PSR 2010 and NSP 2011 via geographic information systems in public statistics
2. **Paweł Chlebicki** (ESRI Polska), ArcGIS as a powerful tool for visualization and spatial analysis of statistical data
3. **Sylwia Filas-Przybył, Maciej Kaźmierczak, Dorota Stachowiak** (Statistical Office in Poznan), The use of data from administrative sources in urban statistics
4. **Robert Buciak, Marek Pieniążek** (Central Statistical Office, University of Warsaw), Spatial classification of rural areas in Poland

5. **Roma Ryś-Jurek** (Poznań University of Life Sciences), Regional differences in agricultural production in the EU-27 countries

Data analysis and classification 1

English language session

19 April 2012, 15.30–17.00

Aula, Poznan University of Economics

1. **Reinhold Decker** (Universität Bielefeld), Model-based analysis of online consumer reviews - methods and applications - **invited paper**
2. **Patrick Groenen** (Erasmus University Rotterdam), The support vector machine as a powerful tool for binary classification - **invited paper**
3. **Andrzej Sokołowski** (Cracow University of Economics), **Beata Basiura** (AGH University of Science and Technology in Cracow) - Ward's agglomerative method

Data analysis and classification 2

English language session

19 April 2012, 17.30–19.00

Aula, Poznan University of Economics

1. **Wojciech Krzanowski** (University of Exeter), Classification: some old principles applied to new problems - **invited paper**
2. **Justyna Wilk** (Wrocław University of Economics), Symbolic approach in regional analyses
3. **Marcin Pełka** (Wrocław University of Economics), Ensemble approach for clustering of interval-valued symbolic data
4. **Justyna Brzezińska** (Katowice University of Economics), Independence analysis of nominal data with the use of log-linear models in R

Data analysis and classification 3

20 April 2012, 9.00–10.30

Aula, Poznan University of Economics

1. **Dorota Rozmus** (Katowice University of Economics), Ensemble approach in taxonomy
2. **Aleksandra Łuczak, Feliks Wysocki** (Poznań University of Life Sciences), Evaluation of level of socioeconomic development of powiats in Wielkopolska province
3. **Kamila Migdał-Najman** (University of Gdańsk), The design of the hybrid, self-learning neural network type SOM-GNG
4. **Krzysztof Najman** (University of Gdańsk), Self-organizing networks in cluster analysis
5. **Piotr Krasucki** (The Institute of Public Affairs), **Władysław Wiesław Łagodziński** (Polish Statistical Association), Integrated system of health data and statistics - indispensable minimum

Data analysis and classification 4

20 April 2012, 11.00–12.30

Aula, Poznan University of Economics

1. **Małgorzata Markowska, Bartłomiej Jefmański** (Wrocław University of Economics), The evaluation of smart specialization dynamics and directions of changes development in European regions by applying fuzzy classification
2. **Katarzyna Kopczewska** (University of Warsaw), Regions in space. Is the location and availability are important for economic and social development?
3. **Robert Pietrzykowski** (Warsaw University of Life Sciences), Spatial analysis of agricultural land prices used to spatial quantile regression
4. **Marzena Piotrowska-Trybull, Stanisław Sirko** (National Defense University), Influence of military units on the community's development
5. **Joanna Zyprych-Walczak, Alicja Szabelska, Idzi Siatkowski** (Poznań University of Life Sciences), Gene selection procedures including correlation between genes

Statistical data 1

19 April 2012, 17.30–19.00

room 236, Poznan University of Economics

1. **Jan Zawadzki, Maria Szumksta-Zawadzka** (West Pomeranian University of Technology, Szczecin), On methods of predicting missing data in time series with periodical (seasonal) fluctuations
2. **Jadwiga Suchecka, Emilia Modranka** (University of Lodz), Spatial statistics - methods and applications
3. **Ewa-Zofia Frątczak, Adam Korczyński** (Warsaw School of Economics), Classical and modern approaches in data imputation methods. Review of theories - examples of practical applications
4. **Artur Mikulec, Aleksandra Kupis - Fijałkowska** (University of Lodz), An empirical analysis of the effectiveness of Wishart and Mojena criteria in cluster analysis
5. **Cezary Głowiński, Kamil Konikiewicz** (SAS Institute sp. z o.o.), Using social networks to model the behaviour of customers of telecommunications companies

Statistical data 2

20 April 2012, 11.00–12.30

room 0011, Poznan University of Economics

1. **Grażyna Trzpiot, Agnieszka Orwat-Acedańska** (Katowice University of Economics), The classification of mutual funds based on the management style-the quantile regression approach
2. **Tadeusz Kufel** (Nicolaus Copernicus University), **Marcin Błazejowski, Paweł Kufel** (Torun School of Banking), The creation of databanks as the basis of socio-econometric analyses in GRETL software

3. **Zbigniew Augustyniak** (Statistical Office in Poznan), Online presentation of statistical data as part of the SISP project
4. **Arleta Olbrot-Brzezińska** (Statistical Office in Poznan), Social functions of statistical information and the responsibility of public statistics
5. **Katarzyna Wawrzyniak** (West Pomeranian University of Technology, Szczecin), The evaluation of main goals of the national development strategy according to voivodships with using the correspondence analysis

Poland and Ukraine in the European context

20 April 2012, 14.30–16.00

room 418, Poznan University of Economics

1. **Ruslan Motoryn** (Ukrainian State University of Finance and International Trade, Lublin University of Technology), **Iryna Motoryna** (Academy of Financial Management in Kiev), **Tetiana Motoryna** (Taras Shevchenko National University Of Kyiv), Comparison indexes of the use of gross savings of households Ukraine and Poland
2. **Andrzej Młodak** (Statistical Office in Poznan, Higher Vocational State School in Kalisz), Modernization of economic activity statistics on a pan-European scale
3. **Karol Andrzejczak** (Poznan University of Technology), Changes and the growth limit of the indicator of the automobile market saturation at the global level and in Poland
4. **Dorota Kwiatkowska-Ciotucha**, **Urszula Załuska** (Wrocław University of Economics), Assessing the use of ESF funding in different regions of Poland in the current programming cycle

Health statistics

20 April 2012, 9.00–10.30

room 418, Poznan University of Economics

1. **Agnieszka Dyzmann-Sroka**, **Andrzej Roszak**, **Maciej Trojanowski** (Greater Poland Cancer Center), **Jerzy Moczko** (Poznan University of Medical Sciences), Methods of increasing data quality in malignant tumor registers
2. **Urszula Wojciechowska**, **Joanna Didkowska** (Cancer Center and Institute of Oncology in Warsaw), Cancer in Poland as a public health problem. Modern tools to measure the phenomenon
3. **Barbara Kołodziejczak**, **Magdalena Roszak** (Poznan University of Medical Sciences), Rapid e-learning in biostatistics
4. **Magdalena Roszak**, **Barbara Kołodziejczak** (Poznan University of Medical Sciences), E-evaluation of the statistical knowledge of medical students
5. **Anna Wierzbicka** (University of Lodz), A taxonomic analysis of public health in Poland in comparison with selected European countries

Statistics of sport and tourism

20 April 2012, 16.30–18.00

room 0011, Poznan University of Economics

1. **Ryszard Walaszczyk** (Ministry of Sport and Tourism), Sport Information System - invited paper
2. **Barbara Liberda** (University of Warsaw), **Łucja Tomaszewicz**, **Iwona Świeczewska** (University of Łódź), Methodology of satellite accounts at the example of the sport satellite account for Poland
3. **Iwona Bąk** (West Pomeranian University of Technology, Szczecin), A statistical analysis of tourist activity of senior citizens in Poland
4. **Marcin Błazejowski** (Torun School of Banking), Applying modelling and prediction algorithms to the tourism market

Classifications and standards

English language session

20 April 2012, 16.30–18.00

room 408, Poznan University of Economics

1. **Włodzimierz Okrasa** (Central Statistical Office, Cardinal Stefan Wyszyński University in Warsaw), Statistics and sociology: Mutually-supportive development from the perspective of interdisciplinaryisation of social research
2. **Éva Laczka** (Hungarian Central Statistical Office), The role of Population and Housing and Agricultural Censuses in the national statistical systems
3. **Grzegorz Krzykowski**, **Marcin Kalinowski** (Gdańsk School of Banking), Selection and screening of data as a key element of statistical research on the financial markets

Statistical inference

20 April 2012, 16.30–18.00

Aula, Poznan University of Economics

1. **Mirosław Krzyśko**, **Łukasz Waszak** (Adam Mickiewicz University, Poznań), Kernel canonical correlation
2. **Grzegorz Kończak** (Katowice University of Economics), Statistical inference for multidimensional contingency tables
3. **Joanna Kisielińska** (Warsaw University of Life Sciences), The exact bootstrap method shown on the example of the mean and variance estimation
4. **Ewa Meller** (Statistical Office in Poznan), Small area estimation for some characteristics of labour market
5. **Łukasz Wawrowski** (Statistical Office in Poznan), Poverty analysis across districts of Wielkopolska using small area estimation methods

Economic and social development analysis

20 April 2012, 16.30–18.00

room 407, Poznan University of Economics

1. **Elżbieta Stańczyk** (Statistical Office in Wrocław), Occupational qualifications of the unemployed on the basis of research by CSO. Interviewship Comparative Analysis
2. **Józef Hozer, Marta Hozer-Koćmiel** (University of Szczecin), Quantum satis of Polish companies in 2012
3. **Janusz Kornecki, Piotr Cmela** (Statistical Office in Łódź), Strategic dilemmas of growth support for micro-enterprises in Poland
4. **Rafał Nowakowski** (Statistical Office in Wrocław), The local elections on the example of regional capitals of Lower Silesia, Małopolska and Wielkopolska - 20 years of experience
5. **Jacek Białek** (University of Łódź), Special case of some general price index formula

POSTER SESSION

19 April 2012, 17.00–17.30

The hall, 2nd floor, Poznan University of Economics

1. **Beata Bal-Domańska** (Wrocław University of Economics, Faculty in Jelenia Góra), Convergence processes modelling the European regional space
2. **Maciej Beręsewicz** (Poznan University of Economics), Internet data sources in the light of the national census based on the example of the real estate market in Poznan
3. **Anna Błaczowska** (Wrocław School of Banking), **Alicja Grześkowiak** (Wrocław University of Economics), Demographic and economic conditions of the process of junior high school education in the regions of Dolny Śląsk and Opolszczyzna in the years 2003–2010
4. **Hanna Dudek, Joanna Landmesser** (Warsaw University of Life Sciences), Income satisfaction and relative deprivation
5. **Alicja Ganczarek-Gamrot** (Katowice University of Economics), Multivariate FIGARCH models in risk estimation on the Polish electric energy market
6. **Hanna Gruchociak** (Poznan University of Economics), Delimitation of local labour markets in Poland
7. **Marta Małecka** (University of Lodz), Evaluation of variance of Value-at-Risk estimators
8. **Magdalena Okupniak** (Poznan University of Economics), Applying intra- and inter-block analysis in economic statistics
9. **Dariusz Parys** (University of Lodz), Remarks on the use of multiple testing in economic sciences
10. **Elżbieta Paszko, Anna Maria Jurek** (University of Lodz), Application of TQM - Total Quality Management in process of improvement teaching methods
11. **Agnieszka Pobłocka** (University of Gdańsk), Insurance market in Poland, 1991–2010
12. **Wojciech Roszka** (Poznan University of Economics), Applying mass imputation methods to integrate databases from various sources
13. **Danuta Rozpędowska-Matraszek** (State Higher Vocational School in Skierniewice), The assessment of health care system operation efficiency based on the example of Poland - an analysis using clustering of poviats
14. **Elżbieta Sobczak** (Wrocław University of Economics), Workforce by R&D activity intensity in EU countries - space-structural analysis
15. **Agnieszka Stanimir** (Wrocław University of Economics), Different techniques of presenting non-metric variables
16. **Marta Styrc** (Warsaw School of Economics), Determinants of the marital stability of first marriages in Poland

Streszczenia referatów

Andrzejczak Karol (Politechnika Poznańska)

Ewolucja i granica wzrostu wskaźnika nasycenia samochodami osobowymi w skali globalnej i w Polsce

Celem tej publikacji jest przedstawienie oceny wskaźnika nasycenia zarejestrowanych samochodów osobowych (wskaźnik NSO) w Polsce do roku 2025. Prognoza jest dokonana na podstawie kryterium porównawczego z krajami o dominujących wskaźnikach nasycenia. W opracowaniu wykorzystano dane zawarte w rocznikach GUS. Metoda badawcza oparta jest na paradygmacie ewolucyjnego rozwoju motoryzacji. Pominięte są oceny krótkookresowych zdarzeń dotyczących wpływu rozwoju gospodarczego na ilość rejestrowanych samochodów osobowych. Wzrost wskaźnika NSO jest ściśle związany ze wzrostem mobilności społeczeństw oraz towarzyszącym mu rozwojowi przewozów pasażerskich i motoryzacji indywidualnej. Źródłem tego zjawiska jest oczywiście prowadzenie przez ludzi coraz intensywniejszej działalności, mającej na celu zaspokojenie swoich potrzeb natury biologicznej, społecznej i kulturowej, podejmowanej w miejscach coraz bardziej oddalonych od stałego miejsca pobytu człowieka. Zwiększanie poziomu stopy życiowej i wzrost mobilności ludzi jest sam w sobie ważnym czynnikiem sprawczym zmian w standardach kulturowych i konsumpcyjnych i w istotny sposób wpływa na rozwój społeczeństw. Jednym z istotniejszych wskaźników tego rozwoju jest wskaźnik NSO. Przeglądając wskaźniki NSO dla różnych krajów można zauważyć olbrzymie jego zróżnicowanie. Pojawiają się też naturalne pytania: *Czy istnieją granice wzrostu wskaźnika NSO w skali globalnej, w krajach europejskich, a w szczególności w Polsce? Jeśli tak, to gdzie one są? Jakie są ich graniczne poziomy? w jakim tempie będziemy się do nich zbliżali? Jakie będą ekologiczne skutki rozwoju motoryzacji? Jak normować i projektować produkcję aut?* Próba odpowiedzi na te pytania jest celem tej publikacji.

Słowa kluczowe: wskaźnik nasycenia, samochody osobowe, prognoza, standardy konsumenckie, skutki środowiskowe

The evolution and the limit of growth of the personal cars saturation ratio in global scale and in Poland

The purpose of this paper is to provide an estimate of the saturation ratio of registered passenger cars (NSO ratio) in Poland by 2025. The forecast is made on the basis of comparison with countries with prevailing rates of saturation. Background information is taken from the annals of the Central Statistical Office GUS. Research method is based on the paradigm of evolutionary development of the automotive industry. Short-term assessment of the economic events in the development of the automotive industry and, consequently, the number of registered passenger cars are omitted. Growth of the NSO is closely associated with an increased mobility of

populations and the accompanying development of passenger transport and private cars. The source of this phenomenon are of course people's intensive activities aimed at satisfying their biological, social and cultural needs, undertaken in locations more distant from the permanent residence of man. Improvement of living standards and increase of the mobility of people is in itself an important factor of changes in cultural and consumer standards and significantly affects the development of societies. One of the key indicators of this development is the NSO ratio. Reviewing the NSO rates for different countries, we can see their huge diversity. Today we ask ourselves the questions: Are there limits of growth of the NSO ratios in the global scale and in the European countries, particularly in Poland? If so, where are they? What are their limit levels? How fast are we approaching them? What are the environmental consequences of the automobile development? How to design and regulate the production of cars? An attempt to answer such questions is made in this publication.

Key words: Saturation rate, personal cars, forecast, consumer standards, environmental consequences

BIBLIOGRAPHY

1. Mały Rocznik Statystyczny Polski 2011, GUS
2. http://www.stat.gov.pl/cps/rde/xbcr/gus/PUBL_oz_maly_rocznik_statystyczny_2011.pdf

Augustyniak Zbigniew (Urząd Statystyczny w Poznaniu)

Prezentacja danych statystycznych w wersji on-line jako element realizacji projektu SISP

Prezentacja składać będzie się z następujących elementów:

- Wstęp – krótka informacja na temat SISP oraz dlaczego, moim zdaniem, prezentacja danych statystycznych w wersji on-line jest jednym z jego elementów.
- Prezentacja rozwiązań europejskich na przykładzie rozwiązań stosowanych przez Portugalski Urząd Statystyczny – w tym punkcie zostaną omówione i zaprezentowane poszczególne etapy jakie należy przejść aby uzyskać dostęp do wybranych danych statystycznych.
- Prezentacja obecnie dostępnych polskich rozwiązań na podstawie Banku Danych Lokalnych – w tym punkcie, analogicznie do poprzedniego, zaprezentowane zostanie polskie rozwiązanie prezentacji danych statystycznych w wersji on-line.
- Porównanie obu rozwiązań – w tym punkcie przedstawione zostaną, wady i zalety obu rozwiązań

- Prezentacja aplikacji Banku Danych Makroekonomicznych – punkt ten zostanie zrealizowany jedynie w przypadku ukończenia prac nad aplikacją (przynajmniej do wersji beta).
- Podsumowanie – krótkie podsumowanie prezentacji.

Słowa kluczowe: dane statystyczne, prezentacja danych statystycznych, prezentacja danych statystycznych on-line

Online presentation of statistical data as part of the SISP project

The presentation will consist of the following elements:

- Introduction - brief information on the SISP project and why, in my opinion, online presentation of statistical data is one of its elements.
- Presentation of European solutions based on the solutions used by the Statistical Office of Portugal - at that point will be discussed and presented various stages, to be followed in order to gain access to selected statistical data.
- Presentation of the currently used Polish solutions with Local Data Bank as an example - at this point, analogous to the previous one, Polish solution for online presentation of statistical data will be presented.
- Comparison of both solutions - at this point the advantages and disadvantages of both solutions will be presented.
- Presentation of the application of the Macroeconomic Databank-only if the complete version of the application will be ready (at least the beta version).
- Summary - brief summary of the presentation.

Key words: statistical data, presentation of statistical data presentation, online presentation of statistical data

BIBLIOGRAPHY

1. Rozporządzenie Rady Ministrów z dnia 28 marca 2007 roku w sprawie Planu Informatyzacji Państwa na lata 2007 – 2010, Dziennik Ustaw nr 61 poz. 415,
2. <http://www.ine.pt> (stan na dzień 28 grudnia 2011)

Bal-Domańska Beata (Uniwersytet Ekonomiczny we Wrocławiu)

Modelowanie procesów konwergencji w europejskiej przestrzeni regionalnej

Badania konwergencji regionów Unii Europejskiej pozwalają lepiej zrozumieć procesy wzrostu i czynniki nimi rządzące. W literaturze przedmiotu często pojawiają się odmienne wnioski, co do obecności procesów konwergencji. Uzależnione jest to m.in. od: rodzaju analizowanych procesów, poziomu analizy (kraje, regiony), składu próby wytypowanej do badania, okresu, wykorzystanych narzędzi analitycznych. W artykule przedstawiono podsumowanie prowadzonych przez autorkę badań nad konwergencją regionów Unii Europejskiej (NUTS2) w latach 1999-2007. Przedmiotem rozważań była zarówno sigma, jak i beta konwergencja (absolutna i warunkowa). Badania prowadzono wśród populacji regionów łącznie oraz w ramach klas regionów wyodrębnionych ze względu na poziom rozwoju gospodarczego, innowacyjności lub przynależności do regionów państw starej piętnastki (UE15) lub nowego rozszerzenia (UE12). O beta konwergencji wnioskowano na podstawie konstrukcji wynikającej z neoklasycznego modelu wzrostu Solowa-Swana.

Celem uwiarygodnienia otrzymanych wyników, analizy prowadzone były z wykorzystaniem różnych narzędzi, tj. odchylenia sztywnego logarytmów produktu dla sigma konwergencji oraz modeli przekrojowych i modeli dla danych panelowych (systemowy estymator uogólnionej metody momentów [Arellano i Bover, Blundel i Bond]) dla oceny procesów beta konwergencji absolutnej i warunkowej.

Słowa kluczowe: konwergencja, rozwój regionalny, regiony Unii Europejskiej (NUTS2)

Convergence processes modelling of the European regional space

The research of EU regions convergence allow for better understanding of growth processes and factors influencing them. Professional literature often presents opposing conclusions regarding the presence of convergence processes. It depends, among others, on: the type of analyzed processes, analysis level (countries, regions), composition of the sample chosen for analysis, period of time or used analytical tools. The article presents the summary of research conducted by the author in the area of the European Union regions (NUTS2). The study focused on both sigma and beta convergence (absolute and conditional). The analysis was conducted in the joint population of regions and within the framework of regional classes distinguished with regard to economic development level, innovation or old EU 15 regions membership, as well as the recent accession (EU 12). Conclusions regarding beta convergence were presented based on the construction resulting from a neoclassical Solow-Swan growth model.

In order to justify the obtained results the analyses were conducted by means of diversified tools, i.e. product logarithms standard deviation for sigma convergence and cross-section models, as

well as models for panel data (systemic estimation for generalized moments method [Arellano and Bover, Blundel and Bond]) for the evaluation of absolute and conditional beta convergence processes.

Key words: convergence, regional development, UE regions (NUTS2)

BIBLIOGRAPHY

1. Arellano M., O. Bover, Another Look at the Instrumental Variables Estimation of Error-components Models, "Journal of Econometrics", 1995, vol. 68, pp. 29-51.
2. Bal-Domańska B., Ekonometryczna identyfikacja β konwergencji regionów szczebla NUTS-2 państw Unii Europejskiej, J. Suchecka (red.), Ekonometria Przestrzenna i Regionalne Analizy Ekonomiczne Nr 253, Wyd. UŁ, Folia Oeconomica, Łódź 2011, s. 9-24.
3. Bal-Domańska B., Konwergencja w regionach Unii Europejskiej o różnym poziomie innowacyjności, K. Jajuga, M. Walesiak (red.), Taksonomia 18, Klasyfikacja i analiza danych – teoria i zastosowania, Wyd. UE. 2011, 120-128.
4. Bal-Domańska B., The application of Dynamic Panel Data Models in the Analysis of Conditional Convergence, Pociecha J. (red.), "Data Analysis Methods in Economics Research", Kraków 2010, s. 107-124.
5. Blundell R., S. Bond, Initial Conditions and Moment Restriction in Dynamic Panel Data Models, "Journal of Econometrics" 87 /1998, pp. 115-143.
6. Kubiela S., Innowacje i luka technologiczna w gospodarce globalnej opartej na wiedzy. Strukturalne i makroekonomiczne uwarunkowania [Innovations and technological gap in global, knowledge intensive economy. Structural and macroeconomic determinants]. WUW, Warszawa 2009, 269-279.
7. Mankiw N., D. Romer, D. Weil, A Contribution to the Empirics of Economic Growth. The Quarterly Journal of Economics, Vol. 107 1992, No. 2 (May), pp 407-437.
8. Strahl D., Innowacyjność europejskiej przestrzeni regionalnej a dynamika rozwoju gospodarczego [Innovation of European regional space vs. economic development rate], UE, Wrocław 2010, s. 114-154.
9. Swan T., Economic growth and capital accumulation. "Economic Record" 1956, 32, pp 334-361.

Bartniczak Bartosz (Urząd Statystyczny we Wrocławiu)

Moduł wskaźników zrównoważonego rozwoju w ramach Banku Danych Lokalnych

Zrównoważony rozwój jest jednym z najważniejszych wyzwań współczesnego świata. W związku z tym jego wdrażanie powinno być monitorowane na każdym poziomie zarządzania. Szczególny nacisk powinien być kładziony na monitorowanie na poziomie lokalnym (gminnym i powiatowym) oraz regionalnym (wojewódzkim). Pozwoli to na uzyskanie odpowiedzi na pytanie jaki jest poziom wdrażania zrównoważonego rozwoju na tym poziomie zarządzania. Podstawowym narzędziem oceny tej koncepcji rozwoju są wskaźniki. Umożliwiają one bowiem stworzenie statystycznego obrazu danej jednostki terytorialnej z punktu widzenia wdrażania koncepcji zrównoważonego rozwoju.

Źródłem wiedzy o poziomie wdrażania zrównoważonego rozwoju mogą być zasoby Banku Danych Lokalnych (BDL). W ramach BDL skonstruowany został uniwersalny moduł umożliwiający monitorowanie zrównoważonego rozwoju na poziomie niższym niż kraj. Punktem odniesienia do stworzenia tego modułu były wskaźniki zrównoważonego rozwoju opracowane przez Eurostat. Moduł opracowany w ramach BDL nie ma charakteru zestawu monitorującego konkretną strategię rozwoju na poziomie lokalnym i regionalnym ale ma charakter podstawowego zestawu wskaźników, który może stanowić podstawę dla analiz nad zrównoważonym rozwojem na poziomie niższym niż kraj. Stworzony moduł ma stanowić pewnego rodzaju „rdzeń” umożliwiający ocenę i porównywanie poszczególnych jednostek samorządu terytorialnego. Taki moduł musi się więc składać ze wskaźników, które posiadają pokrycie danymi statystycznymi i dane te są porównywalne dla wszystkich jednostek samorządu terytorialnego. Oprócz tego, ze względu na specyfikę poszczególnych jednostek, moduł będzie mógł być uzupełniany o wskaźniki istotne z punktu widzenia danej jednostki terytorialnej. Wzbogacania modułu powinny dokonywać tzw. grupy wiodące na danym terenie – przedstawiciele sektora samorządowego, przedsiębiorców oraz organizacji społecznych.

Opracowany moduł składa się z tematów, które są istotne z punktu widzenia monitorowania zrównoważonego rozwoju na poziomie regionalnym i lokalnym. W każdym z tematów wydzielone zostały podtematy w ramach których umieszczone zostały konkretne wskaźniki.

W przygotowywanym artykule zaprezentowany zostanie moduł wskaźnikowy w ramach BDL. Omówione zostaną problemy, które pojawiły się w trakcie prac nad modułem. Wskazane zostaną także dalsze kierunki prac nad wskaźnikami zrównoważonego rozwoju na poziomie lokalnym i regionalnym. Zasygnalizowane zostaną możliwości praktyczne wykorzystania opracowanego modułu. W artykule zaprezentowane zostaną także wartości dla wybranych wskaźników.

Sustainable development indicators module in Local Data Bank

Sustainable development is one of the major challenges of the modern world. Therefore, its implementation should be monitored at every level of management. Particular emphasis should be placed on monitoring at the local regional level. This will allow to get answers to the question what is the level of implementation of sustainable development at that management level. The main tool for assessing that concept of development are indicators. They allow for the creation of a statistical picture of the territorial unit in terms of implementation of the concept of sustainable development.

The source of knowledge about the level of implementation of sustainable development can be resources of Local Data Bank (LDB). In the LDB was created a universal module designed to monitor sustainable development at a level lower than the country. The reference point for the creation of this module were sustainable development indicators developed by Eurostat. Module made by the BDL is not a set of indicators developed to monitoring the progress of a specific strategy at local and regional level but it's a core set of indicators which could be the basis for the analysis of sustainable development at a level lower than the country. Created module should be some kind of „core” for the assessment and comparison of individual units of local government. This module must therefore consist of indicators that have statistics coverage and data are comparable for all local government units. In addition, due to the nature of individual units, the module can be supplemented by the indicators relevant to a particular territorial unit. Supplementation of that module should be done by so called leading groups in the area - the representatives of the local government sector, business and social organizations.

The developed module consists of topics that are relevant to the monitoring of sustainable development at regional and local level. In each of the subjects were allocated the sub-themes where specific indicators have been placed.

In a prepared paper will be presented the indicator module of LDB. In paper will be discuss the problems that emerged during the work on the module. The article will be also show further directions of work on indicators of sustainable development at local and regional level. In article will be also mentioned the possibility of practical use developed module. This paper will present the values for selected indicators.

Key words: sustainable development, indicators, Local Data Bank

Bartosiewicz Stanisława (Wyższa Szkoła Bankowa we Wrocławiu), Stańczyk Elżbieta (Urząd Statystyczny we Wrocławiu)

Trochę o społeczno-ekonomicznej historii Polski w latach 1918-2008 (na podstawie badań GUS-u)

W referacie dokonano cząstkowego porównania trzech epok historii Polski w zakresie niektórych zagadnień społecznych i ekonomicznych (okres międzywojenny, PRL, lata od 1990 roku).

Uwzględniono następujące zagadnienia:

1. strukturę ludności ze względu na kilka kryteriów,
2. charakterystyki rynku pracy,
3. wskaźniki dobrobytu ludności,
4. strukturę dochodu narodowego ze względu na wpływ działów gospodarki narodowej na jego wielkość.

Słowa kluczowe: dobrobyt ludności, rynek pracy, dochód narodowy, budżet gospodarstw domowych, spis ludności

Socio-economic history of Poland in years 1918-2010 (on the basis of research by CSO)

The paper presents a partial comparison of three historical periods of Poland in the scope of some social and economic issues (interwar period, the People's Republic of Poland, years since 1990).

The following issues have been included:

1. structure of population on account of several criteria (by sex and age, nationality, educational level),
2. description of labour market,
3. measurement of well-being of people,
4. structure of national income on account of the influence of sectors of national economy on its size.

Key words: well-being of people, labour market, national income, household budget surveys, national population census

Jacek Batóg, Barbara Batóg, Magdalena Mojsiewicz (Uniwersytet Szczeciński),
Katarzyna Wawrzyniak (Zachodniopomorski Uniwersytet Technologiczny w Szczecinie)

S tatystyczny system monitoringu stopnia realizacji strategii rozwoju miasta

W procesie tworzenia strategii rozwoju miasta istotną rolę, oprócz określenia listy celów strategicznych, kierunkowych i szczegółowych oraz sposobów ich osiągnięcia, odgrywa kontrola stopnia realizacji poszczególnych celów. Monitoring realizacji strategii ma na celu bieżącą kontrolę sposobu realizacji i zgodności celów szczegółowych, kierunkowych i strategicznych z zapisami strategii, umożliwienie podjęcia działań eliminujących niedopuszczalne odchylenia wskaźników od przyjętych wartości docelowych (norm) oraz bieżącą ocenę stopnia realizacji przyjętych celów i stwierdzenie, czy ich dalsza realizacja nie jest zagrożona.

Ważnym elementem systemu monitoringu jest zestaw wskaźników oceniających, który wraz z przyporządkowanymi im wartościami normatywnymi prowadzi do wypracowania metody oceny stopnia realizacji celów szczegółowych strategii. Zaproponowane wskaźniki powinny być stosunkowo proste w rozumieniu i interpretacji, politycznie niezależne, koherentne oraz możliwe do zmierzenia w założonych jednostkach czasu. Ponadto powinny pochodzić z wiarygodnych źródeł, które gwarantują jednolitą metodologię pomiaru oraz nie powinny stanowić dużego obciążenia dla uczestników procesu monitorowania.

Celem artykułu jest konstrukcja systemu monitoringu stopnia realizacji strategii rozwoju miasta wykorzystującego wyselekcjonowany zestaw wskaźników. W zależności od rodzaju monitorowanych celów zastosowane zostaną poniższe metody oceny:

- dla prostych celów szczegółowych zaproponowany zostanie pojedynczy wskaźnik oceniający, bądź wskaźnik budowany przy pomocy ankietowej metody ekspertów,
- dla złożonych (wieloaspektowych) celów szczegółowych wykorzystany będzie wskaźnik syntetyczny ze wskaźników pokrywających zakres zdefiniowany w celu lub/i z wykorzystaniem badań ankietowych skierowanych do mieszkańców miasta lub ekspertów,
- dla złożonych celów szczegółowych, w przypadku niemożności budowy wskaźnika syntetycznego z różnorodnych wskaźników pokrywających zakres zdefiniowany w celu, dla których nie można jednoznacznie ustalić ani normy, ani poziomu wskaźników na dany rok, zostanie zastosowany wskaźnik agregatowy, obliczany na podstawie wszystkich indywidualnych wskaźników dla danego celu szczegółowego przy zastosowaniu ocen na skali dychotomicznej 0-1.

Ocena stopnia realizacji strategii wymaga nie tylko analizy poziomów poszczególnych wskaźników oceniających, lecz również ich porównywania z odpowiadającymi im wartościami normatywnymi wyznaczającymi docelowy poziom realizowanych celów szczegółowych.

W artykule wskazane zostaną także ograniczenia procesu oceny stopnia realizacji celów szczegółowych strategii rozwoju miasta związane na przykład z brakiem możliwości pomiaru ilościowego czy przenikaniem się części celów.

Słowa kluczowe: strategia rozwoju, system monitoringu, miary statystyczne

Statistical monitoring system of implementation of the city strategy of development

During the process of creating the city strategy of development it is important not only to define the strategic, directional and detailed goals and ways of their realization but also to control their implementation levels. The monitoring system should allow for present evaluation of the implementation degree as well as for initiation of activities which eliminate unadmittable deviations of indices values from accepted norms. The set of evaluating indices together with their norms is important part of monitoring system. These indices should have quite simple interpretation, should be coherent and should come from reliable sources. The aim of the paper is to present the statistical monitoring system of realization of the city strategy of development along with the set of appropriate indices and their norms. According to the kind of goal the following methods of evaluating will be applied:

- for simple detailed goals - simple index or index formulated on the base on expert survey,
- for complex detailed goals - aggregate measure based on several simple indices with using surveys directed to experts or citizens,
- for complex detailed goals in case of impossibility of creating synthetic index because of lack of norms - aggregate index based on dichotomic indices.

The barriers of above kinds of evaluation will be also discussed.

Key words: strategy of development, monitoring system, statistical measures

Bąk Iwona (Zachodniopomorski Uniwersytet Technologiczny w Szczecinie)

Statystyczna analiza aktywności turystycznej seniorów w Polsce

Jednym z najważniejszych problemów społecznych współczesnego świata jest starzenie się społeczeństwa. Szczególnie Europa, a w niej Polska staje się społeczeństwem szybko starzejącym, a określenie „Stary Kontynent” nabiera nowego znaczenia. Starzenie się społeczeństw jest procesem demograficznym, definiowanym najogólniej jako zmiany w stanie i strukturze według wieku ludności kraju (regionu, świata) polegające na wzroście w ogólnej liczbie ludności liczby i udziału ludności starszej. We wszystkich krajach występuje stałe przedłużanie się przeciętnego trwania życia, głównie na skutek systematycznego spadku natężenia zgonów niemowląt,

jak i na skutek wydłużającego się życia ludzkiego. Wzrastająca liczba osób starszych w strukturze populacji to nie tylko wyzwanie dla polityki gospodarczej i społecznej, ale także potężne wyzwanie dla rynku dóbr i usług konsumpcyjnych. Coraz bardziej rośnie siła nabywcza tego segmentu rynku i to zarówno w odniesieniu do usług życia codziennego, jak i towarów i usług luksusowych, szczególnie w branży zdrowotnej, motoryzacyjnej, gastronomicznej i turystycznej.

Zmiany demograficzne już dziś wywołują wiele implikacji w różnych sferach życia, wpływają na model produkcji, konsumpcji, strukturę wydatków budżetowych, rynek pracy, system ubezpieczeń i świadczeń socjalnych, inwestycje, oszczędzanie, stosunki społeczne, rodzinne, styl życia. Rosnąca liczba osób starszych może stać się też potężnym źródłem koniunktury dla szeroko pojętego przemysłu turystycznego. Ludzi w starszym wieku należy zachęcać do uprawiania turystyki o każdej porze roku i dążyć do wyrobienia trwałego nawyku aktywności. Pozytywny wpływ aktywności turystycznej na zdrowie i samopoczucie starszych osób jest niezaprzeczalny i potwierdzony wieloma wynikami badań. Turystyka jest jednym z ważnych czynników zaspakajających podstawowe potrzeby biologiczne, psychiczne i społeczne ludzi starszych.

Celem artykułu jest analiza zależności między zmienną określającą wystąpienie wyjazdu turystycznego a zmiennymi ekonomicznymi i społeczno-demograficznymi charakteryzującymi gospodarstwa domowe seniorów w Polsce. Podstawę informacyjną badań będą stanowiły dane statystyczne dotyczące wyjazdów turystycznych gospodarstw domowych, w których głowa gospodarstwa jest w wieku 60 lat i więcej. Zostały one zaczerpnięte z przeprowadzonego w 2009 roku przez Główny Urząd Statystyczny badania ankietowego nt. „Turystyka i wypoczynek w gospodarstwach domowych”. Podstawą wnioskowania o istotności wpływu czynników determinujących wystąpienie wyjazdu turystycznego będzie test niezależności χ^2 oraz oszacowanie modelu logitowego.

Bednarski Tadeusz (Uniwersytet Wrocławski)

Statystyczna analiza przyczyn obciążoności próby w badaniach sondażowych rynku pracy

Skuteczność przeciwdziałania bezrobociu jest uwarunkowana wiarygodną analizą czynników wpływających na stan i czas jego trwania. Jakość tej analizy, w przypadku badań sondażowych, jest z kolei zależna nieobciążoności próby losowej. W statystycznych badaniach społecznych warunek ten praktycznie nigdy nie jest spełniony, a w sondażowych badaniach rynku pracy poziom obciążoności niejednokrotnie przekracza 40%. Badania takie, jak na przykład BAEL są dodatkowo obciążone „wyczerpywaniem się” danych – część jednostek rezygnuje z udziału w trakcie trwania cyklu badawczego. Dla danych BAEL z roku 2002 obciążoność, w perspektywie statystycznej analizy czasu poszukiwania pracy przez bezrobotnego, przekracza 60%.

Instytucje powołane do systematycznego prowadzenia badań społecznych próbują ten problem zredukować poprzez stosowne zachęty do udziału w projektach sondażowych, jednak nie są w stanie go całkowicie rozwiązać. Dlatego też, część prac naukowych skierowano na ocenę

wpływu obciążoności próby na efekt wnioskowania. Wyniki takich analiz dają podstawę dla ewentualnej korekty wniosków końcowych, tak jak na przykład w pracach Pyy-Martikainena i Rendtela (2008) oraz van den Berga, Lindeboom i Doltona (2006) w odniesieniu do rynku pracy. Polegają one w części na analizie porównawczej wyników badań sondażowych oraz tychże badań sondażowych uzupełnionych o administracyjnie zebrane dane tych jednostek, które wcześniej zrezygnowały z udziału w sondażu (w warunkach polskich byłoby to na przykład zestawienie danych BAELu i urzędów pracy). Tak „uzupełnione” archiwalne dane sondażowe pozwalają także na badanie mechanizmu odmowy w sondażach. W kontekście analizy wpływu czynników bezrobocia wyróżnia się dwa podstawowe mechanizmy odmowy: selektywny, będący skutkiem wpływu tylko tych czynników, które determinują rozkład czasu poszukiwania pracy, oraz przyczynowy wynikający z rezygnacji z udziału w badaniach z powodu znalezienia pracy. Rozróżnienie tych dwóch kategorii jest o tyle istotne, że czynnik selektywny, który wydaje się być dominujący w generowaniu odmowy udziału w sondażach nie prowadzi do obciążoności estymatorów określających strukturę zależności pomiędzy czasem poszukiwania pracy i jego determinantami.

W wykładzie przedstawi się szczegółowy opis powyższych mechanizmów generujących odmowy w badaniach sondażowych. Głównym celem jest prezentacja metody testowania obecności mechanizmu przyczynowego przy założeniu, że relacje pomiędzy czasem poszukiwania pracy i determinantami jego rozkładu są opisane, jak często postuluje się w badaniach empirycznych, za pomocą modelu Coxa.

Słowa kluczowe: model Coxa, analiza danych bezrobocia, obciążoność próby, analiza przyczynowości

Statistical analysis of nonresponse in unemployment duration studies

High survey nonresponse in unemployment duration studies, depending on nonresponse type - causal or selective, may result in different estimation bias. A method of testing the nonresponse causality for time to event social studies is proposed for data sets where survey information can be combined with complete administrative records. The method is based on a simple application of the Cox regression model which is frequently used in the unemployment duration studies.

Key words: Cox model, sampling bias, causality, unemployment duration studies

BIBLIOGRAPHY

1. T. Bedmarski and F. Borowicz (2010). Analysis of non-response causality in labor market surveys. *Acta Universitatis Lodzensis. Folia Oeconomica* 253, s. 217–224,
2. T. Bednarski (2011). On Testing Causality Non-response in Unemployment Duration Studies under the Cox Model. Submitted for publication.
3. G. J. van den Berg, M. Lindeboom and P. J. Dolton (2006). Survey non-response and the duration of unemployment. *J. R. Statist. Soc. A* (2006) 169, Part 3, pp. 585–604

4. M. Pyy-Martikainen and U. Rendtel (2008) Assessing the impact of initial nonresponse and attrition in the analysis of unemployment duration with panel surveys. *Advances in Statistical Analysis* 92: 297–318.

Beręsewicz Maciej (Uniwersytet Ekonomiczny w Poznaniu)

Internetowe źródła danych w świetle spisu powszechnego na przykładzie rynku nieruchomości w Poznaniu

W obecnych czasach Internet jest głównym źródłem informacji dotyczącej otaczającej nas rzeczywistości. Główną zaletą korzystania z tych źródeł jest ich aktualność, jednak ogromna ilość danych znajdująca się w sieci najczęściej nie współgra z ich jakością. Bardzo trudno mówić w przypadku takich danych o reprezentatywności, tym bardziej, że często powielają informacje i nie podają źródeł. Jednym z przykładów jest rynek nieruchomości, gdzie ogłoszenia umieszczane są zarówno przez pośredników jak i prywatne osoby.

Celem artykułu jest próba oceny internetowych źródeł danych dotyczących rynku nieruchomości, w szczególności mieszkań w kontekście spisu powszechnego. Wykorzystane zostaną dane pozyskane z www.otodom.pl, www.poznan.gumtree.pl, www.nieruchomosci.zumi.pl i innych portali internetowych w celu próby weryfikacji hipotezy o reprezentatywności tych źródeł, co może przyczynić się do wykorzystania danych pochodzących z sieci w kontekście badań reprezentacyjnych.

Słowa kluczowe: rynek nieruchomości, internetowe źródła danych, badania reprezentacyjne

Internet data sources in context of Census on the example of real estate market in Poznań

Nowadays Internet is main source of information about our surroundings. Main advantage of using internet data sources is topicality, however the vast amount of data which is available, typically does not adequately reflect their quality. It is very difficult to speak in the case of such data of representativeness, the more that often duplicate information and do not give sources or references. One example is the housing market, where advertisements are placed by both intermediaries and private persons.

The aim of this paper is an attempt to evaluate the Web data sources of real estate market, in particular housing market in Census context. Author will use data from main Polish websites like www.otodom.pl, www.poznan.gumtree.pl, www.nieruchomosci.zumi.pl and others in attempt to verify hypothesis of representativeness of these sources, which may help to use Web data in context of sample surveys.

Keywords: real estate market, internet data sources, sample surveys

Berger Jan, Walczak Tadeusz, Witkowski Janusz (Główny Urząd Statystyczny)

S

tatystyka publiczna - rozwój historyczny i aktualne wyzwania

W referacie postanowiliśmy przypomnieć korzenie polskiej statystyki publicznej oraz zwrócić uwagę na jej niełatwe drogi rozwoju, którego celem nadrzędnym była zawsze służba dla kraju oraz jego obywateli. Z tego względu wyodrębniono kilka najważniejszych okresów historii statystyki publicznej, akcentując te wydarzenia, które wywarły największy wpływ na realizację jej zadań oraz na jej przyszły rozwój.

Punktem wyjścia rozważań jest charakterystyka warunków, w jakich powstawały pierwsze założenia organizacji i zadań służb statystycznych po utworzeniu niepodległego Państwa Polskiego po 123 latach zaborów. Podkreślono rolę polskich statystyków, w tym działaczy Polskiego Towarzystwa Statystycznego, tworzących pierwsze opracowania i publikacje statystyczne charakteryzujące sytuację społeczno-gospodarczą na ziemiach polskich pod panowaniem mocarstw zaborczych. Opracowania te stały się jakby załącznikiem przyszłych założeń organizacji służb statystyki publicznej po odzyskaniu przez Polskę niepodległości w 1918 r.

Ważny etap rozwoju statystyki obejmuje okres międzywojenny, w którym ukształtowały się jej organizacyjne podstawy. Ogromne znaczenie miało tu utworzenie Głównego Urzędu Statystycznego podległego bezpośrednio Prezydium Rady Ministrów. Przyjęto wówczas tzw. scentralizowany model organizacji badań statystycznych realizowanych w oparciu o program badań zatwierdzany przez Radę Ministrów. GUS był nie tylko głównym realizatorem tych badań, ale także otrzymał uprawnienia do koordynacji badań statystycznych prowadzonych w całym kraju. Zainicjowano wówczas wiele badań statystycznych, podjęto szereg prac metodologicznych i w zakresie standardów klasyfikacyjnych. Dzięki temu stworzono podstawę dla lepszej spójności informacji wynikowych.

O uzyskaniu wysokiego profesjonalnego poziomu pracy służb statystycznych w okresie międzywojennym świadczy przygotowanie, przeprowadzenie i opracowanie wyników dwóch kolejnych powszechnych spisów ludności w latach 1921 i 1931. Spotkało się to z uznaniem środowiska międzynarodowego czego wyrazem było przyznanie Polsce przez Międzynarodowy Instytut Statystyczny funkcji organizatora 18. Kongresu tej organizacji. Kongres pod patronatem prezydenta Rzeczypospolitej Ignacego Mościckiego odbył się w dniach 21–24 sierpnia 1929 r.

W referacie omówiono także krótko tragiczne losy polskiej statystyki w latach okupacji hitlerowskiej 1939–1945, a następnie rozwój statystyki w okresie po II wojnie światowej, podzielony na trzy okresy historyczne: lata 1945 do 1989, 1989 do 2004 oraz po 2004 roku.

Mimo niezwykle trudnych warunków lokalowych i kadrowych tuż po II wojnie światowej, służby statystyczne przystąpiły do organizacji pierwszych badań szczególnie użytecznych do oceny stanu ludności i majątku narodowego pozostałego po zniszczeniach wojennych. Najpilniejszą sprawą było oszacowanie liczby ludności zamieszkałej w kraju i jej terytorialnego rozmieszcze-

nia. Pierwszy, tzw. sumaryczny spis według stanu z dnia 14 lutego 1946 r. wykazał, że ogólna liczba ludności Polski w jej nowych granicach wyniosła 23930 tys. osób. Kolejne, pełne spisy ludności przeprowadzono z częstotliwością generalnie co dziesięć lat (1950, 1960, 1970, 1978, 1988, 2002 oraz ostatni - w 2011 roku).

W okresie powojennym przystąpiono także do opracowania systemu badań statystycznych odpowiadających najpilniejszym potrzebom organów zarządzania gospodarką w nowych warunkach ustrojowych, co oznaczało m.in. konieczność dostosowania tematyki badań do wymagań wieloletnich, rocznych i kwartalnych planów gospodarczych, zgodnie z przepisami dekretu z dnia 31 lipca 1946.

Od roku 1956 rozwinęła się działalność publikacyjna, wznowiono także wydawanie roczników statystycznych, których publikowanie zaprzestano w 1951 r. Znacznie ożywiono współpracę międzynarodową. Podjęto wówczas dwu i wielostronne prace metodologiczne dotyczące porównań ważniejszych kategorii społeczno-ekonomicznych, w tym w dziedzinie rachunków narodowych oraz spożycia i cen. Ważną rolę w tych pracach odgrywał Zakład Badań Statystyczno-Ekonomicznych GUS. Wyrazem uznania międzynarodowej społeczności statystyków dla aktywności GUS była decyzja Międzynarodowego Instytutu Statystycznego przyznająca Polsce zadanie zorganizowania kolejnego - 40 Kongresu tej organizacji w Warszawie w dniach 1-9 września 1975 r.

W drugiej połowie lat 1980. przystąpiono do wprowadzania bardziej wszechstronnych zmian w statystyce publicznej. Zmiany te objęły niemal wszystkie aspekty funkcjonowania systemu informacji statystycznej. Bardziej precyzyjnie określono rolę statystyki w kraju zmierzającym do wprowadzenia gospodarki rynkowej oraz przestrzegania zasad społeczeństwa demokratycznego. Szczególny akcent położono na zapewnienie warunków uzyskiwania obiektywnych i rzetelnych informacji bazujących na społecznym zaufaniu do statystyki i partnerskich stosunkach z respondentami.

Prace nad wprowadzeniem zasadniczych zmian w statystyce nasiliły się zwłaszcza po 1990 roku w ramach przygotowania Polski do wstąpienia do UE. Oprócz głębokich zmian metodologicznych i prawnych dostosowano zakres badań i opracowań do obowiązujących programów badań i terminów dostarczania informacji w Unii Europejskiej. Prawnym fundamentem tych zmian było przyjęcie ustawy z dnia 29 czerwca 1995 r. o statystyce publicznej.

Od 1 maja 2004 r., po wstąpieniu Polski do Unii Europejskiej polska statystyka publiczna funkcjonuje jako pełnoprawny członek Europejskiego Systemu Statystycznego (ESS). Z tego faktu wynika cały szereg nowych obowiązków, ale także uprawnień i możliwości rozwojowych. W obecnych warunkach do zadań tych należą w szczególności:

- ciągłe doskonalenie zakresu i metodologii badań w celu zaspokojenia zmieniających się potrzeb informacyjnych użytkowników,
- dalsza modernizacja procesu badań statystycznych z uwzględnieniem różnicowania źródeł danych, w tym szerszego wykorzystania źródeł administracyjnych oraz zastosowania osiągnięć technologii informatycznych,
- wzmocnienie pozycji polskiej statystyki na arenie międzynarodowej,

- zagwarantowanie odpowiedniej jakości statystyki zgodnie z wymogami Europejskiego Kodeksu Praktyk Statystycznych,
- doskonalenie pomiaru postępu społecznego, dobrobytu oraz zrównoważonego rozwoju, wychodząc poza tradycyjny pomiar wzrostu gospodarczego mierzonego PKB.

Słowa kluczowe: statystyka oficjalna, historia statystyki, metodologia badań, Europejski System Statystyczny, jakość statystyki

The Polish official statistics its historical development and current challenges

The paper presents a historical overview of the Polish Official Statistics and focuses on challenges encountered in its development, which has always been primarily treated as service for the country and its citizens. Most significant periods in the history of Official Statistics are identified, each highlighting events with the biggest influence on the performance of its duties and its future development.

The paper begins with a description of conditions that shaped the underlying assumptions about the organization and tasks of statistical services in the newly formed independent state of Poland after 123 years of partitions. It stresses the role of Polish statisticians, including members of the Polish Statistical Association, who prepared the first reports and statistical materials about the socio-political situation in the ethnically Polish regions under foreign rule. This early work constituted the basis of the future organizing principles of statistical services after Poland regained independence in 1918.

Another important period comprises the years before World War II, when the organizational structure of Official Statistics was established. This refers, in particular, to the creation of the Central Statistical Office (CSO), reporting directly to the Prime Minister. It was then that the so-called 'centralized model of organizing statistical surveys' was adopted, which was to be conducted according to a survey programs approved by the Council of Ministers. CSO was not only charged with the task of conducting surveys but also coordinating surveys across the whole country. Many survey programs were started at that time, accompanied by methodological studies and development of classification standards. All this helped to lay the foundations for the production of more coherent statistical output.

The high professional standard of statistical services in the period of 1918-1939 was demonstrated during the preparation, administration and analysis of two censuses in 1921 and 1931. This achievement received international recognition, which was marked by International Statistical Institute's decision to grant Poland the task of organizing its 18th Congress. The Congress, enjoying the patronage of the President of Poland, Ignacy Mościcki, was held between 21 and 24 of August 1929.

The paper continues with a brief look at the tragic circumstances in the Nazi-occupied Poland and progress made in the field of statistics in the years following the war, focusing on three periods: 1945-1989, 1989-2004 and after 2004.

Despite extremely difficult operating conditions and limited staff after World War II, statistical services started to conduct their first surveys in order to assess the state of population and national wealth and resources in the war-ravished country. The most urgent need was to estimate the population size and its territorial distribution. The first, so-called 'summary' census at 14 February 1946 put the country's population at 23,930,000 people. Consecutive censuses were conducted more or less every decade (1950, 1960, 1970, 1978, 1988, 2002 and the last one in 2011).

During the postwar period work was started to develop a system of statistical surveys that would meet the needs of authorities responsible for managing the country's economy; this meant adapting the survey program in accordance with longterm, annual and quarterly economic plans.

In 1956 publishing activity was restarted, including the publication of statistical yearbooks, which was discontinued in 1951. Several bilateral and multilateral comparative methodological studies were initiated, for example in the field of national accounts, consumption and prices. In recognition of the work conducted by CSO, International Statistical Institute selected Poland as the organizer of its 40th congress, which was held in Warsaw, between 1–9 September 1975.

In the late 1980s Official Statistics underwent sweeping changes. The changes were intended to redefine the role of statistics in a country, which was about to switch to free market economy and become a democratic society. The matter of particular significance was to create conditions for obtaining objective and reliable information and to build up trust towards statistics and develop partner relationships with respondents.

Following 1990, the reforming efforts within the field of Official Statistics were intensified in preparation for Poland's accession to the EU. In addition to far-reaching methodological and legal changes, survey programs and information provision dead-lines were adapted to those existing in the EU.

Since the moment of joining the EU on 1 May, 2004, Polish Official Statistics has been a full member of the European Statistical System (ESS). This entails a number of duties as well as privileges and development prospects. Under new circumstances, these tasks include:

- ongoing development in terms of survey scope and methodology to meet the changing needs of information users;
- further modernization of the statistical survey process and making use of various data sources, particularly more reliance on administrative data and advances in information technology;
- strengthening the position of Polish statistics internationally;
- ensuring the required quality of statistical information in accordance with the European Statistics Code of Practice
- improving ways of measuring social progress, welfare and sustainable development, which go beyond the traditional measure of economic progress by means of GDP.

Key words: Official Statistics, history of statistics, survey methodology, European Statistical System, statistical quality

Białek Jacek (Uniwersytet Łódzki)

Szczególny przypadek pewnej ogólnej formuły indeksu cen

W ramach aksjomatycznej teorii indeksów postuluje się, aby właściwie skonstruowany indeks statystyczny spełniał określoną grupę postulatów. W literaturze spotkać można propozycje systemów minimalnych wymagań wobec indeksów statystycznych (Martini (1992), Eichhorn i Voeller (1976), Olt (1996)). Wymóg współmierności i test odwróconego czasu należą do grona restrykcyjnych testów i wiele, nawet powszechnie używanych indeksów, nie spełnia tych postulatów.

W pierwszej części artykułu przedstawiono i omówiono autorską, ogólną formułę indeksu cen. W pracy pokazano, iż proponowana formuła spełnia przynajmniej te aksjomaty, które zawarte są w systemach minimalnych wymagań. Przy pewnym dodatkowym założeniu, formuła ta spełnia również test odwróconego czasu. W drugiej części pracy zaproponowany zostanie indeks cen, który stanowi szczególny przypadek formuły ogólnej. A zatem z aksjomatycznego punktu widzenia indeks ten jest skonstruowany właściwie. Dokonane zostanie porównanie proponowanego indeksu ze znanymi z literatury indeksami cen.

Bielak Renata (Główny Urząd Statystyczny)

Rola statystyki w procesie kształtowania polityk rozwoju – priorytety i wyzwania

Zmiany zachodzące we współczesnym świecie powodują konieczność kompleksowego podejścia do zarządzania gospodarczego. Podstawowe problemy i wyzwania stojące przed decydentami – m.in. potrzeba zapewnienia zrównoważonego rozwoju, czy sprawiedliwości i spójności społecznej – są aktualne zarówno w wymiarze lokalnym, krajowym, jak i ponadnarodowym. Warunkiem powodzenia w sprostaniu tym wyzwaniom jest wzajemna współpraca i koordynacja działań na różnych poziomach zarządzania.

Programowanie polityki rozwoju bazuje na odpowiednich systemach informacji, stanowiących podstawę do podejmowania decyzji. Zapewnienie dostępu do rzetelnych, aktualnych i użytecznych informacji jest zadaniem statystyki publicznej. Informacje statystyczne stanowią narzędzie

kształtowania świadomości w zakresie kierunków i tempa zmian zachodzących w życiu społecznym, gospodarczym, czy w środowisku naturalnym.

Celem wystąpienia jest zaprezentowanie działań realizowanych przez polską statystykę publiczną, które wspierają efektywność zarządzania strategicznego na poziomie krajowym i europejskim. Instrumentem usprawnienia koordynacji procesów rozwojowych w Unii Europejskiej jest strategia Europa 2020. Jest to strategia na najbliższą dekadę, która zakłada zapewnienie wzrostu inteligentnego, zrównoważonego i sprzyjającego włączeniu społecznemu. Nadrzędne cele strategii mają być transponowane na cele krajowe, a postęp w zakresie ich osiągnięcia będzie monitorowany w oparciu o ustalony zestaw wskaźników. W ślad za przyjęciem strategii ustanowiono europejski okres oceny (tzw. *europejski semestr*) oraz zobowiązano kraje członkowskie do opracowywania i wdrażania *Krajowych Programów Reform*, które mają zawierać działania wskazujące na sposób dojścia do celów strategii Europa 2020.

Krajowy Program Reform dla Polski zakłada nowe, ponadsektorowe podejście do zarządzania rozwojem. Zrewidowane zasady prowadzenia polityki rozwoju, uwzględnione m.in. w *Długookresowej strategii rozwoju kraju* oraz w *Koncepcji przestrzennego zagospodarowania kraju*, zakładają zwiększenie przejrzystości procesu programowania. Efektem tych prac jest zredukowanie liczby dokumentów strategicznych i opracowanie dziewięciu zintegrowanych strategii rozwoju kraju, które mają służyć realizacji celów średnio- i długookresowej strategii. Integralną częścią wszystkich tych dokumentów są wskaźniki monitorowania realizacji założonych celów.

Wzmacnianie koordynacji polityki gospodarczej na poziomie Unii Europejskiej odbywa się również poprzez wprowadzanie dodatkowych mechanizmów nadzoru – w ramach funkcjonującej procedury nadmiernego deficytu, czy wdrażanie nowych elementów – jak procedura nadmiernych nierównowag makroekonomicznych. Przyjęty we wrześniu 2011 r. pakiet pięciu regulacji i dyrektywy (tzw. *six pack*) ma podstawowe znaczenie dla zapewnienia wzmoczonej dyscypliny fiskalnej i uniknięcia nadmiernych zakłóceń równowagi makroekonomicznej. Ocena poziomu tych zagrożeń oraz podjęcie ewentualnych interwencji ma się odbywać w oparciu o zestaw wskaźników (tzw. *Scoreboard indicators*), pochodzących z krajowych systemów statystycznych i bankowych.

Biorąc pod uwagę wielość obszarów podlegających monitorowaniu, wyzwaniem dla statystyki jest zapewnienie wysokiej jakości informacji, dostępnych również na poziomie regionalnym i lokalnym. Co więcej, kompleksowy charakter opisywanych procesów powoduje konieczność dokonywania szczegółowych analiz prezentowanych informacji i właściwej interpretacji zachodzących zmian. Jest to zadanie ambitne i odpowiedzialne, zwłaszcza mając świadomość szczególnej roli omawianych mechanizmów w kształtowaniu trwałego rozwoju.

Słowa kluczowe: Europa 2020, europejski semestr, strategie rozwoju kraju, scoreboard indicators

Statistics for policymaking – priorities and challenges

Changes in the contemporary world engender the necessity of a complex approach to economic

governance. Basic problems and challenges such as the need for sustainable development, or justice and social coherence are still present at local, national and also supranational level. The necessary condition for facing these challenges is mutual cooperation and coordination of activities at all levels of governance.

Growth's policy planning relies on the proper information systems that constitute the basis for decision-making process. The main task of official statistics is to ensure the access to reliable, current and useful information. Statistical information is used to raise awareness of directions and rate of changes occurring in both socio-economic life and natural environment.

The main purpose of this speech is to present activities performed by the Polish official statistics, which support the efficiency of strategic management at national and European level. The Europe 2020 Strategy is a key instrument for the coordination of the European Union's development. It is the strategy for the nearest decade, which aims to ensure smart, sustainable and inclusive growth. The primary targets of the Europe 2020 Strategy are to be translated into national targets, and the progress of its implementation will be monitored by means of a given set of indicators. The approval of the Europe 2020 Strategy was followed by the establishment of a European semester and Member States were obliged to develop and implement National Reform Programmes, indicating the ways of achieving the Europe 2020 Strategy targets.

National Reform Programme for Poland takes a new cross-sectoral approach to development governance. The revised rules of development policy, enclosed in "Long-term national development strategy" and "National Spatial Management Concept", imply transparency improvement of the planning process. The reduction of the number of strategic documents and the compilation of nine integrated national development strategies, aiming at implementing targets of mid-term and long-term strategies are the direct result of those actions. A set of indicators used to monitor their realization is an integral part of the above-mentioned documents.

The economic policy at European Union level is also reinforced through the implementation of additional supervisory mechanisms – within the excessive deficit procedure and through the implementation of new elements such as an excessive imbalance procedure. The Package of five regulations and a directive (Six Pack), approved in September 2011 is crucial for maintenance of fiscal discipline and avoidance of excessive macroeconomic balance disruptions. The evaluation of dangers and performing possible interventions is to be done by means of a set of indicators (Scoreboard indicators), coming from statistical and banking systems.

Taking into account the amount of areas to be monitored, it is a challenge for statistics to ensure information of high quality, available also at regional and local level. Moreover, the complexity of described processes engenders the necessity of detailed analysis and the proper interpretation of changes. It is an ambitious and responsible task, especially when we realize the special role of discussed mechanisms in shaping stable development.

Key words: Europe 2020, European semester, country's growth strategies, scoreboard indicators

Błackowska Anna (Wyższa Szkoła Bankowa we Wrocławiu), Grześkowiak Alicja (Uniwersytet Ekonomiczny we Wrocławiu)

Uwarunkowania demograficzno-gospodarcze procesu kształcenia gimnazjalnego na Dolnym Śląsku i Opolszczyźnie w latach 2003-2010

Celem pracy jest próba zastosowania wybranych metod analizy wielowymiarowej do oceny zróżnicowania wiedzy i umiejętności uczniów gimnazjów na Dolnym Śląsku i Opolszczyźnie, na tle sytuacji demograficzno-gospodarczej w tych regionach.

Trzy lata po powołaniu do życia gimnazjów (w 1999 roku) wprowadzono porównywalne testy sprawdzające oraz oceniające wiedzę i umiejętności uczniów. Testy te mające charakter egzaminu zewnętrznego, oceniane przez zewnętrznych egzaminatorów, dały podstawę do podjęcia szczegółowych badań porównawczych na temat wiedzy i umiejętności uczniów według różnych kryteriów. Testy obejmują dwie części egzaminu: humanistyczną i matematyczno-przyrodniczą.

W części humanistycznej weryfikowano kompetencje uczniów w dwóch zakresach:

1. czytanie i odbiór tekstów kultury;
2. tworzenie własnego tekstu.

W części matematyczno-przyrodniczej egzaminu wyróżniono cztery obszary:

1. umiejętne stosowanie terminów, pojęć i procedur z zakresu przedmiotów matematyczno-przyrodniczych niezbędnych w praktyce życiowej i dalszym kształceniu;
2. wyszukiwanie i stosowanie informacji;
3. wskazywanie i opisywanie faktów, związków i zależności, w szczególności przyczynowo-skutkowych, funkcjonalnych, przestrzennych i czasowych;
4. stosowanie zintegrowanej wiedzy i umiejętności do rozwiązywania problemów.

W prowadzonych analizach próbowano wykazać wpływ zróżnicowania regionalnego, związanego z przeszłością historyczną, tradycjami kulturowymi czy cywilizacyjnymi na wyniki egzaminu. Wskazano również na różnice przestrzenne na niższym stopniu podziału terytorialnego – na poziomie gmin i powiatów, w których istnieją dosyć istotne rozpiętości pomiędzy wynikami uzyskiwanymi przez gimnazjalistów. Badano także wpływ płci gimnazjalistów na osiągane rezultaty oraz czynniki gospodarcze, status majątkowy gospodarstw domowych i ogółu ludności, w gminach i powiatach w badanych województwach. Wykazano, że wyniki egzaminów gimnazjalnych wskazują na dużą rozbieżność pomiędzy umiejętnościami w zakresie części humanistycznej i matematyczno-przyrodniczej. W rozkładach rezultatów występują także różnice ze względu na płeć oraz lokalizację szkoły.

Słowa kluczowe: gimnazjum, umiejętności, statystyczna analiza wielowymiarowa, demografia, gospodarka

Economic and demographic determinants of gymnasium educational process in lower Silesia and Opole region in the period 2003-2010

The paper presents an attempt to apply selected multivariate statistical methods to assess the diversity of knowledge and skills of gymnasium students in Lower Silesia and Opole region against the background of demographic and economic situation in this area. Comparable tests evaluating students' knowledge and skills were introduced three years after the establishment of gymnasiums (in 1999). These tests are treated as external examination and are assessed by external examiners, what gave the possibility to conduct an extensive comparative study on the knowledge and skills of students according to various criteria. The considered tests comprise two parts of the examination: the humanities and the science.

In the humanities part students' competences were verified in two areas:

1. reading and the interpretation of texts;
2. creation of own text.

In the science part four areas were regarded:

1. application of terms and procedures in the area of exact and natural sciences, essential in life and further education;
2. finding and using information;
3. identifying and describing facts, relationships and dependences in particular cause-effect, functional, spatial and temporal;
4. application of integrated knowledge and skills to solve problems.

The analyses were oriented to identify the impact of regional diversity associated with the historical past, cultural tradition and civilization on the results of the examination. Spatial differences at a lower level of territorial division (communities and districts) were indicated according to the results obtained by gymnasium students. The study includes also an evaluation of the influence of gender, economic factors, financial status of households on gymnasium students achievements in communities and districts in the regarded voivodships. The gymnasium examination results show that there is a substantial difference between the skills and competences concerning the two parts - humanities and science. The distributions of results vary when gender and school location are taken into consideration.

Key words: gymnasium, competences, multivariate statistical analysis, demography, economy

Błażejowski Marcin (Wyższa Szkoła Bankowa w Toruniu)

Algorytmy modelowania i prognozowania na przykładzie rynku turystycznego

W referacie zostanie przedstawiony aspekt metodologii automatycznej specyfikacji modelu (ang. model selection) i prognozowania zjawisk na przykładzie rynku usług turystycznych dla danych: miesięcznych, kwartalnych i rocznych. Przedstawiony zostanie procedura automatycznej specyfikacji modelu wg koncepcji modelowania zgodnego CongruentSpecification dostępna w programie GRETL i opracowana przez zespół: Marcin Błażejowski, Paweł Kufel oraz Tadeusz Kufel. Omówiona zostanie metodologia modelowania zgodnego jak również założenia i przyjęte rozwiązania w procedurze CongruentSpecification. Całość zostanie zilustrowana wieloma przykładami empirycznymi dotyczącymi rynku turystycznego.

Błażejowski Marcin (Wyższa Szkoła Bankowa w Toruniu), Kufel Paweł (Wyższa Szkoła Bankowa w Toruniu), Kufel Tadeusz (Uniwersytet Mikołaja Kopernika w Toruniu)

Budowa banków danych jako podstawa analiz statystyczno-ekonometrycznych w oprogramowaniu GRETL

Celem artykułu jest przedstawienie procedury budowy baz danych w oprogramowaniu GRETL (GNU Regression, Econometric and Time-series Library) dla danych importowanych z dziedzinowych banków danych GUS, Eurostatu i innych instytucji. Utworzone zostaną dziedzinowe banki danych dla oprogramowania GRETL w układzie danych przekrojowych (w podziale terytorialnym) i szeregów czasowych, a także danych panelowych. Przedstawiona zostanie propozycja analiz statystyczno – ekonometrycznych wykonanych za pomocą funkcji oprogramowania GRETL.

Bonnéry Daniel, Chauvet Guillaume, Deville Jean-Claude (ENSAE)

Neyman type optimality for marginal quota sampling

A long time ago, in his seminal paper on representativeness, Jerzy Neyman had shown how to optimize a stratified sampling. By the way he proved that representativeness, seen as equivalent to precision, is not achieved by equivalent repartitions in the sample and in the population.

Balanced sampling generalize stratification, in particular when it is used to fit the margins of a contingency table: it is marginal quota sampling. We seek for a similar optimization procedure leading to inclusion probabilities specific for each cell of the table. A simple model based method can be used in a quite straightforward way and is, in some sense, a natural generalization. However the hypotheses of the model are generally too strong. In fact, when the sampling is near from maximal entropy, there exists a quite simple approximation of the variance which is model free. We simply try to optimize this approximation.

A natural iterative procedure is used. It is shown to converge to this optimum.

BIBLIOGRAPHY

1. J.C.Deville, Y. Tillé, Efficient balanced sampling: the cube method, *Biometrika*, 2004.
2. J.C.Deville, Y. Tillé, Variance approximation under balanced sampling, *Journal of statistical planning and Inference*, 2005.
3. J.Neyman, On the two different aspects of the representative method : the method of stratified sampling and the method of purposive selection. *Journal of the Royal Statistical Society*, 1934.
4. Y.Tillé, C. Favre, Optimal allocation in balanced sampling. 2005.

Borucka Jadwiga, Romaniuk Joanna, Frątczak Ewa (Szkoła Główna Handlowa w Warszawie)

Trwałość pierwszych związków (małżeństw i kohabitacji) w kohortach urodzeniowych 1950–1970

Rodzina jako podstawowa komórka społeczna stanowi przedmiot badań statystyków i demografów od lat. Na przestrzeni czasu instytucja małżeństwa zmieniała się, zarówno pod względem

wieku osób, które zakładają rodzinę, jak i liczby osób wchodzących w jej skład oraz trwałości związku małżeńskiego, będącego podstawą każdej rodziny w tradycyjnym rozumieniu tego słowa. Równolegle zaczęły rozwijać się alternatywne formy rodziny, przede wszystkim kohabitacje oraz urodzenia pozamałżeńskie. Takie zróżnicowanie znacznie komplikuje analizę funkcjonowania i trwałości rodziny rozumianej tym razem szerzej, a nie tylko w tradycyjnej formie. Dodatkowo utrudnia ją fakt, iż bardzo często kohabitacja poprzedza zawarcie związku małżeńskiego, stąd powstaje pytanie o definicje konieczne w analizowaniu trwałości związków, nie jest bowiem jednoznaczne, czy taka kohabitacja powinna być traktowana zupełnie niezależnie, czy też analizowana w odniesieniu do następującego po niej związku małżeńskiego. Niezależnie jednak od przyjętej definicji, rosnąca w ostatnich latach liczba rozwodów oraz coraz częściej spotykane kohabitacje, skłaniają do stawiania pytań o trwałość małżeństw i związków nieformalnych oraz zmiany jakie dają się zaobserwować w tym zakresie w ostatnich latach. Niniejszy artykuł poświęcony jest analizie trwałości pierwszych związków, zarówno małżeństw, jak i kohabitacji przeprowadzonej w oparciu o wyniki modeli czasu przeżycia w estymacji zarówno klasycznej, jak i bayesowskiej. Analiza kohortowa umożliwia przedstawienie zmian, jakie zachodziły w funkcjonowaniu związków na przestrzeni lat. Badania przeprowadzono z wykorzystaniem danych z formularza D „Dietność kobiet” z formularzem a „Narodowy Spis Powszechny Ludności i Mieszkań w 2002r.” pochodzących z Narodowego Spisu Powszechnego Ludności i Mieszkań w roku 2002.

Borys Tadeusz (Uniwersytet Ekonomiczny we Wrocławiu), Potkański Tomasz (Związek Miast Polskich)

Monitorowanie rozwoju regionalnego i lokalnego

W Polsce istnieje wiele dokumentów strategicznych określających cele społeczne, gospodarcze i środowiskowe perspektywie długookresowej oraz kierunki działań zgodne z zasadami rozwoju zrównoważonego, uwzględniającymi wymóg spójności sfery ekonomicznej, społecznej i ekologicznej. Dokumentami strategicznymi, w których zawarto cele związane z polityką zrównoważonego rozwoju są między innymi: Strategia Rozwoju Kraju 2007–2015, Narodowe Strategiczne Ramy Odniesienia 2007–2013, strategie sektorowe (np. Strategia zmian wzorców produkcji i konsumpcji na sprzyjające realizacji zasad zrównoważonego rozwoju). Konieczność stworzenia w Polsce – na wzór innych krajów Unii Europejskiej – systemu wskaźników monitorujących realizację celów tych strategii jest oczywista. W referacie przedstawiono główne problemy tworzenia w Polsce zestawów wskaźników do monitorowania zrównoważonego rozwoju na poziomie lokalnym, regionalnym i krajowym.

Słowa kluczowe: monitoring, wskaźniki, rozwój regionalny, rozwój lokalny

Monitoring of regional and local development

There are many strategic documents in Poland which define social, economic and environmental objectives in a long-term perspective, as well as operational directions in line with sustainable development principles, considering the requirements of economic, social and ecological sphere cohesion. Strategic documents which include goals related to sustainable development policy are, among others, the following ones: National Development Strategy 2007–2015, National Strategic Reference Framework 2007–2013, sector strategies (e.g. Strategy of changes in production and consumption patterns into these facilitating sustainable development principles implementation). The need to establish in Poland – following other European Union countries – the system of indicators monitoring the implementation of these strategies objectives, has become obvious. The paper presents major problems with building in Poland sets of indicators to monitor sustainable development on local, regional and national level.

Key words: sustainable development, indicators, regional development, local development

Brzezińska Justyna (Uniwersytet Ekonomiczny w Katowicach)

Analiza niezależności zmiennych nominalnych z wykorzystaniem modeli logarytmiczno-liniowych w R

Modele logarytmiczno-liniowe są techniką analizy dowolnej liczby zmiennych niemetrycznych tj. zmiennych nominalnych bądź porządkowych. Termin model logarytmiczno-liniowy pochodzi z równania modelu, w którym modelowane są logarytmy liczebności empirycznych w tablicy kontyngencji. W modelu pojawiają się efekty wpływu poszczególnych oraz interakcje pomiędzy nimi w dowolnych konfiguracjach. Modele te pozwalają na szczegółową analizę niezależności poprzez wyróżnienie wielu modeli pomiędzy modelem niezależności a zależności jak np. modelu zależności warunkowej, modelu zależności całkowitej, modelu zależności homogenicznej lub też modelu zerowy. Do oceny modelu wykorzystywane są dwie statystyki chi kwadrat: tradycyjne chi-kwadrat oraz iloraz największej wiarygodności [Fisher 1922]. Interpretacja wyników obu statystyk jest taka sama i mówi o tym, czy wartości empiryczne w tablicy kontyngencji różnią się od teoretycznych. Jeśli tak, wówczas badany model zostaje odrzucony. Innymi miernikami oceny modelu są m. in. kryteria informacyjne AIC [Akaike 1973] i BIC [Raftery 1986].

Wartości teoretyczne dla danego modelu wyznaczone są poprzez algorytm dopasowania iteracyjno-proporcjonalnego IPF (*iterative proportional fitting algorithm*) [Deminga, Stephena 1940]. Dla wyznaczonych w ten sposób liczebności empirycznych danego modelu wyznaczone są parametry modelu odpowiadające zmiennym oraz interakcjom występującym w modelu. Głównym narzędziem opisu tych parametrów są szanse rozumiane jako stosunek liczby obiektów przynależących do danej kategorii zmiennej do liczby obiektów przynależących do pozostałych kategorii tej zmiennej.

Modele logarytmiczno-liniowe dostępne są w programie R dzięki funkcji `loglm` w pakiecie MASS oraz dzięki funkcji `glm` w pakiecie stats. Funkcje te różnią się algorytmem iteracyjnym a także postacią danych wejściowych. W niniejszym artykule zaprezentowana zostanie analiza niezależności zmiennych nominalnych z wykorzystaniem modeli logarytmiczno-liniowych w programie R.

Słowa kluczowe: modele logarytmiczno-liniowe, zmienne niemetryczne, tablica kontyngencji, analiza niezależności zmiennych nominalnych

Independence analysis of nominal data with the use of R software

Log-linear models are used to analyze the relationship between two or more categorical (e.g. nominal or ordinal) variables. The term log-linear derives from the fact that one can, through logarithmic transformations, restate the problem of analyzing multi-way frequency tables in terms that are very similar to ANOVA. Specifically, one may think of the multi-way frequency table to reflect various main effects and interaction effects that add together in a linear fashion to bring about the observed table of frequencies. There are several types of models between dependence and independence like: homogenous association, partial association, conditional association and null model. There are two types of Chi-squares: the traditional Pearson Chi-square statistic and the maximum likelihood ratio Chi-square statistic [Fisher, 1922]. In practice, the interpretation and magnitude of those two Chi-square statistics are essentially identical. Both tests evaluate whether the expected cell frequencies under the respective model are significantly different from the observed cell frequencies. If so, the respective model for the table is rejected. There are also information criteria AIC and BIC to test the goodness of fit.

Expected cell frequencies are obtained with the use of iterative proportional fitting algorithm (IPF) [Deming, Stephen 1940]. The next step is to derive model coefficients for single variables as well as for interaction parameter. The most useful tool for interpreting model parameter is odds and odds ratio.

Log-linear models are available in R software with the use of `loglm` function in MASS library and `glm` function in stats library. These functions are different in terms of fitting algorithm and data.

In this paper log-linear analysis will be presented.

Key words: log-linear models, cross-tabulation, qualitative data, independence analysis of nominal data

BIBLIOGRAPHY

1. Akaike H. (1973), Information theory and an extension of the maximum likelihood principle, w: Proceedings of the 2nd International Symposium on Information, Petrow B. N., Czaki F., Budapest: Akademiai Kiado.
2. Deming W., Stephan F., (1940), On a least squares adjustment of a sampled frequency table when the expected marginal totals are known, Annals of Mathematical Statistics,

11, 427-444.

3. Fisher R. A. (1922), On the interpretation of chi-square from contingency tables, and the calculations of P, J. Roy. Statist. Soc. 85, 87-94.
4. Raftery A.E. (1986), Choosing models for cross-classification, Amer. Sociol. Rev. 51, 145-146.

Buciak Robert, Pieniążek Marek (Główny Urząd Statystyczny, Uniwersytet Warszawski)

Klasyfikacja przestrzenna obszarów wiejskich w Polsce

Celem prezentacji jest przedstawienie procedury, dzięki której opracowano klasyfikację przestrzenną obszarów wiejskich w Polsce. Oficjalne statystyki rozróżniają obszary miejskie i wiejskie w Polsce na podstawie kryterium administracyjnego. Obecnie obszary wiejskie nie mogą być traktowane jedynie jako obszary rolne i leśne. W wielu przypadkach wzrosło znacznie funkcji turystycznych lub mieszkalnych. W związku z zachodzącymi zmianami Główny Urząd Statystyczny aktywnie dąży do sporządzenia nowego podziału obszarów wiejskich. Celem pracy metodologicznej „Typologia obszarów wiejskich” jest stworzenie standardu podziału gmin w oparciu o oficjalne dane statystyczne, z uwzględnieniem rozbieżności gmin miejsko-wiejskich na części miejskie i wiejskie. Założono, że klasyfikacja powinna być oparta na prostych i łatwych do interpretacji wskaźnikach, co umożliwiłoby jej szerokie wykorzystanie.

Typologie przestrzenne obszarów wiejskich często opierają się na gęstości zaludnienia, które może być zastąpione przez użytkowanie gruntów. Z przeprowadzonej analizy wynika, że korelacja pomiędzy udziałem obszarów zabudowanych i zurbanizowanych a gęstością zaludnienia na poziomie gmin wynosi 0,94. Użytkowanie gruntów pozwala na wyróżnienie miast i zróżnicowanie obszarów wiejskich na rolnicze, leśne i częściowo zurbanizowane. W pracy wykorzystano dane nt. użytkowania gruntów otrzymanych z Głównego Urzędu Geodezji i Kartografii, pochodzących z Ewidencji Gruntów i Budynków. Grunty zabudowane i zurbanizowane obejmują grunty: mieszkaniowe, przemysłowe, zurbanizowane niezabudowane, rekreacji i wypoczynku, komunikacyjne, użytki kopalne i inne zabudowane. Założono, że udział tych terenów powinien być przedstawiany jako procent łącznej powierzchni użytkowanej, a nie jako udział powierzchni całkowitej jednostki. Oznacza to, że od całkowitej powierzchni jednostki odjęto m. in. grunty pod wodami, nieużytki, tereny różne. Otrzymaną w ten sposób powierzchnię podzielono następnie na trzy grupy: rolnicze, leśne, zabudowane i zurbanizowane. Efekt tego podziału można przedstawić dla każdej jednostki na wykresie trójkątnym.

Zasadniczym problemem do rozważenia była kwestia wyznaczenia granic klas obszarów wiejskich i miejskich. W trakcie szczegółowej analizy danych okazało się, że udział gruntów zabudowanych i zurbanizowanych przekracza 6% w tych jednostkach wiejskich, w których obserwuje

się procesy suburbanizacji lub rozległe inwestycje przemysłowe. Z drugiej strony okazało się, że udział gruntów zabudowanych i zurbanizowanych jest niższy od 15% tylko w tych miastach, które w swoich granicach posiadają znaczne powierzchnie gruntów rolnych lub leśnych. Żaden taki przypadek nie został odnotowany w województwie warmińsko-mazurskim.

W klasyfikacji obszarów wiejskich w Polsce wyróżniono 7 klas: a) zurbanizowane; b) leśne, częściowo zurbanizowane; c) rolnicze, częściowo zurbanizowane; d) przeważająco rolnicze, częściowo zurbanizowane; e) leśne; f) rolnicze; g) przeważająco rolnicze. Przy tak przyjętych kryteriach do ostatnich trzech klas zaliczono 32 miasta (3,5% ogólnej ich liczby). Natomiast do klasy „zurbanizowane” przypisano 18 jednostek wiejskich (0,8% ogólnej ich liczby). W grupie tej znalazły się jednostki graniczące z miastami, położone w strefie podmiejskiej lub związane z górnictwem odkrywkowym.

Słowa kluczowe: obszary wiejskie, klasyfikacja, użytkowanie gruntów

Spatial classification of rural areas in Poland

The presentation demonstrate the procedure which led us to spatial classification of rural areas in Poland. Official statistics in Poland traditionally distinguish urban and rural areas on the basis of administrative criteria. Nowadays, rural areas cannot be considered strictly as agricultural and forest areas. In many cases have increased their touristic and residential functions. Central Statistical Office of Poland actively aims to classify rural areas. The aim of the methodological work „The typology of rural areas” is to create a standard of division based on the official statistic data for gminas with partition of urban-rural gminas into urban and rural part. We assume that our classification should contain simple and easily interpretable indicators what provides its intelligibility.

Spatial typologies of rural areas are commonly based on population density what we found could be substituted by land use. Population density and share of build-up and urbanized areas in total used land correlates highly (0.94) at the local territorial unit level. Land use allows one to distinguish cities and diversify rural areas to agricultural, forest and partly urbanized classes. We use the share of build-up and urbanized areas determined by Central Geodesy and Cartography Office. This administrative data are based on the legal status of the plots. Build-up and urbanized areas include residential, industrial, urbanized non-built-up, recreational, transport, mining and other build-up areas. We recognize that shares should be counted not as a percentage of total area, but as a percentage of total use area. It means that we subtract from total area the land which could not be use like lakes, rivers, sea, wastelands and others. Whole use area could be trisect into agricultural, forest and build-up plus urbanized areas. This division allows us to show each unit on a trilinear-diagram.

The question is, if there exist any natural borders between rural and urban areas considered as a share of build-up and urbanized areas. During detailed data analysis we found out that there is a share of build-up and urbanized areas higher than 6% in the rural units mainly in case of a unit where suburbanization or industrial investment take place. On the other hand we found out that there is a share of build-up and urbanized areas lower than 15% in cities only if there

are vast surrounding agricultural or forest areas incorporated into the cities' boundaries. None of these situations took place intensively in Warmińsko-Mazurskie voivodeship.

Spatial classification of rural areas in Poland distinguishes 7 classes: a) urbanized; b) forest, partly urbanized; c) agricultural, partly urbanized; d) predominantly agricultural, partly urbanized; e) forest; f) agricultural; g) predominantly agricultural. We recognized and analyzed 32 cities (3.5% of the total number of cities) which are classified in non urbanized classes and 18 rural units (0.8% of the total number of rural units) that belong to urbanized class. The urbanized rural units either lie near city in suburban area or there is a sizeable open pit mine within unit's border.

Key words: rural areas, classification, land use

Caliński Tadeusz (Uniwersytet Przyrodniczy w Poznaniu)

Rozwój polskiej myśli statystycznej – Osiągnięcia w biometrii

Mówiąc biometria, mamy na myśli dyscyplinę naukową zajmującą się zastosowaniami metod matematycznych i statystycznych w rozwiązywaniu problemów biologicznych, w szerokim tego słowa znaczeniu, a zwłaszcza w planowaniu i analizie eksperymentów. W Polsce pionierami biometrii było dwóch wybitnych uczonych. Antropolog, Jan Czekanowski (1882–1965), oraz chemik, agrotechnik i hodowca roślin, Edmund Załęski (1863–1932).

Jednym z uczniów Załęskiego był Stefan Barbacki (1903–1979), od 1945 roku związany z Poznaniem, uczony o ogromnych osiągnięciach w zakresie rozwoju metodyki doświadczalnictwa rolniczego i biometrii. Jego działalność naukowa i pedagogiczna przyczyniła się do stworzenia poznańskiej szkoły statystyki matematycznej i biometrii (patrz [1] i [2]).

Podobnie rodziły się zainteresowania biometrią w kilku innych polskich ośrodkach naukowych. Należy tu wymienić zwłaszcza następujące trzy osoby i ich szkoły.

Mikołaj Olekiewicz (1896–1971), inicjator lubelskiej szkoły biometrycznej, rozwiniętej na Akademii Rolniczej (Uniwersytecie Przyrodniczym) w Lublinie przez jego przyjaciela i ucznia Wiktora Oktabę (1920–2009) i nadal rozwijanej przez uczniów W. Oktaby (patrz [3] i [4]).

Zygmunt Nawrocki (1910–1978), kontynuator myśli Jerzego Neymana, współtwórca warszawskiej szkoły statystyki matematycznej i biometrii, rozwijanej nadal przez jego następców i uczniów w SGGW (patrz [5] i [6]).

Julian Perkal (1913–1965), z wrocławskiej szkoły zastosowań matematyki, który skupił wokół siebie grono wybitnych matematyków i przyrodników różnych specjalności, inspirując ich do stosowania metod statystycznych w badaniach eksperymentalnych (patrz np. [7] i [8]).

Słowa kluczowe: biometria, planowanie i analiza doświadczeń

Development of Polish statistical research – Achievements in biometry

By biometry we mean the branch of science which deals with applications of mathematical and statistical methods to biological problems, in a broad sense, particularly to the design and analysis of experiments. Two prominent scientists are considered as pioneers of biometry in Poland. An anthropologist, Jan Czekanowski (1882–1965), and a chemist, agricultural researcher and plant breeder, Edmund Załęski (1863–1932).

One of Załęski's followers was Stefan Barbacki (1903–1979), from 1945 in Poznań, a scientist of great achievements in the development of agricultural research methodology and biometry. His scientific and educational activity contributed essentially to the formation of the Poznań school of mathematical statistics and biometry (see [1] i [2]).

The interest in biometry in some other Polish scientific centers was initiated in a similar way. Particularly, the following three persons and their schools should be mentioned.

Mikołaj Olekiewicz (1896–1971), the initiator of the Lublin biometric school, developed further by his friend and follower Wiktor Oktaba (1920–2009) at the University of Life Sciences in Lublin, and still developing due to the followers of Professor Oktaba (see [3] i [4]).

Zygmunt Nawrocki (1910–1978), a continuator of Jerzy Neyman's ideas, cofounder of the Warsaw school of mathematical statistics and biometry, developed further by his successors and followers at the Warsaw University of Life Sciences (SGGW) (see [5] i [6]).

Julian Perkal (1913–1965), from the Wrocław school of applied mathematics, who was able to assemble a group of eminent mathematicians and scientists from various branches, inspiring them to applications of statistical methods in experimental research (see, e.g., [7] i [8]).

Key words: biometry, design and analysis of experiments

BIBLIOGRAPHY

1. T. Caliński, Statystyka matematyczna i biometria w ośrodku poznańskim, Z. Palka (red.), Poznańska szkoła matematyczna, Wyd. Nauk. UAM, Poznań 1995, 31–40.
2. T. Caliński, A. Dobek, Z. Kaczmarek, Stefan Barbacki i jego szkoła biometryczna w Poznaniu, P. Krajewski, Z. Zwierzykowski, P. Kachlicki (red.), Genetyka w ulepszaniu roślin użytkowych, Instytut Genetyki Roślin PAN, Poznań 2004, 197–205.
3. W. Oktaba, Agro-biometria w Polsce, Listy Biometryczne Nr 51–54 (1976), 10–28.
4. M. Wesołowska-Janczarek, 50 lat Katedry Zastosowań Matematyki Akademii Rolniczej w Lublinie (1952–2002, Wyd. Akademii Rolniczej w Lublinie, Lublin 2002.
5. Prekursorzy statystyki i biometrii w SGGW, Agricola, Pismo SGGW w Warszawie, Nr 70 (2008), 9–10.
6. Katedra Biometrii, J. Łabętowicz (red.), 100-lecie Wydziału Rolnictwa i Biologii 1906–2006, Tom 1, Wyd. SGGW, Warszawa 2006, 195–222.

7. T. Bogdanik, Przegląd dorobku biometrii w naukach medycznych w Polsce w ostatnim 30-leciu, Listy Biometryczne Nr 51-54 (1976), 29–47.
8. Z. Welon, Biometria a antropologia w Polsce, Listy Biometryczne Nr 51-54 (1976), 1–9.

Ceranka Bronisław, Graczyk Małgorzata (Uniwersytet Przyrodniczy w Poznaniu)

Estymacja całkowitej masy obiektów w układach wagowych

Rozważamy eksperyment, w którym chcemy oszacować nieznane miary p obiektów w oparciu o n operacji pomiarowych. Doświadczenie to można opisać modelem liniowym Gaussa – Markowa $\mathbf{y} = \mathbf{X}\mathbf{w} + \mathbf{e}$, gdzie $\mathbf{y}$ jest $n \times 1$ wymiarowym wektorem obserwacji, $\mathbf{X}$ jest $n \times p$ wymiarową macierzą układu wagowego, $\mathbf{w}$ jest $p \times 1$ wymiarowym wektorem nieznanymi miar obiektów, $\mathbf{e}$ jest $n \times 1$ wymiarowym wektorem błędów losowych. Jeżeli macierz układu jest macierzą pełnego rzędu kolumnowego, to możemy estymować wszystkie nieznane miary obiektów. W przeciwnym przypadku, który właśnie rozważamy, możemy estymować całkowitą masę obiektów, albo pewne funkcje nieznanymi parametrów. W przedstawionej pracy zajmujemy się estymacją całkowitej masy obiektów przy dodatkowych założeniach dotyczących wektora błędów. Przyjmujemy, że w modelu nie występują błędy systematyczne, błędy są jednakowo skorelowane i mają jednako- kowe wariancje, tzn. $\mathbf{E}(\mathbf{e}) = \mathbf{0}_n$, $\mathbf{E}(\mathbf{e}\mathbf{e}') = \sigma^2 \mathbf{G}$, gdzie $\mathbf{G} = g((1 - \rho)\mathbf{I}_n + \rho \mathbf{1}_n \mathbf{1}_n')$ jest macierzą dodatnio określoną.

Pod pojęciem macierzy układu wagowego rozumie się macierz o elementach 0 i 1, tzn. macierz sprężynowego układu wagowego oraz macierz o elementach -1 , 0 i 1, tzn. macierz chemicznego układu wagowego.

Dla obu typów układów podano warunki konieczne i dostateczne dla optymalnej estymacji całkowitej masy. Rozważania teoretyczne zilustrowano przykładami konstrukcji odpowiednich układów.

Słowa kluczowe: całkowita masa, odporność, układ wagowy

Estimation of total weight of objects in weighing design

We consider the experiment in which we estimate unknown parameters of p objects using n measurement operations. Such experiment we can describe using Gauss – Markoff linear model $\mathbf{y} = \mathbf{X}\mathbf{w} + \mathbf{e}$ where $\mathbf{y}$ is an $n \times 1$ vector of observations, $\mathbf{X}$ is the $n \times p$ design matrix, $\mathbf{w}$ is a $p \times 1$ vector of unknown measurements of objects, $\mathbf{e}$ is an $n \times 1$ vector of random errors. If the design matrix is of full column rank, then we are able to estimate all unknown measurements of

objects. In the other case which is considered here, we can estimate the total weight of objects or some functions of unknown parameters, only.

In present paper we study the estimation of total weight under special assumptions connected with the vector of errors. We assume that there are not systematic errors, moreover the errors are equal correlated and they have the same variances, i.e. $\mathbf{E}(\mathbf{e}) = \mathbf{0}_n$, $\mathbf{E}(\mathbf{e}\mathbf{e}') = \sigma^2\mathbf{G}$, where $\mathbf{G} = g((1 - \rho)\mathbf{I}_n + \rho\mathbf{1}_n\mathbf{1}_n')$ positive definite matrix.

Weighing design there is design in which the design matrix has elements equal to 0 and 1 and then is called spring balance weighing design or there is the design in which the design matrix has elements equal to -1 , 0 and 1 and then is called chemical balance weighing design.

For both design matrices we give the necessary and sufficient conditions for the optimal estimation of the total weight. Moreover, we give the examples of construction of optimal designs.

Key words: robustness, total weight, weighing design

BIBLIOGRAPHY

1. Ceranka B., Graczyk M.: *Robustness optimal spring balance weighing designs for estimation total weight*. "Kybernetika", 2011, 47, 902-908
2. Ceranka B., Graczyk M.: *Robustness of optimal chemical balance weighing designs for estimation of total weight*. "Communications in Statistics-Theory and Methods", 2012 (praca przyjęta do druku)
3. Ceranka B., Katulska K.: *Optimum singular spring balance weighing designs with non-homogeneity of the variance of errors for estimating the total weight*. "Australian Journal of Statistics", 1986, 28, 200-205.
4. Ceranka B., Katulska K.: *On the estimation of total weight in chemical balance weighing designs under the covariance matrix of errors $\sigma^2\mathbf{G}$* . "SoftStat'95, Advances in Statistical Software 5, F. Faulbaum and W. Bandila Eds., Stuttgart, Lucius and Lucius", 1996, 461-468.
5. Katulska K.: *On the estimation of total weight in singular spring balance weighing designs under the covariance matrix of errors $\sigma^2\mathbf{G}$* . "Australian Journal of Statistics", 1989, 31, 277-286.

Chambers Ray, Gunky Kim (Wollongong University)

Regression analysis using data obtained by probability linking of multiple data sources

Data linkage is now an important research tool in many areas of scientific research. For example, Wilkins, Shields and Rotermann (2009) describe an analysis that models the relationship

between an individual's probability of hospitalization and length of time spent subsequently in hospital and his/her smoking status using a linked data set obtained by merging data collected in the Canadian Community Health Survey and data held in Statistics Canada's Hospital Person-Oriented Information database. In Australia, Brook, Rosman and Holman (2008) claim that linked health record data sets produced by the Western Australia Data Linkage Unit over the period 1995 - 2003 were used in 708 research outputs, comprising journal articles, reports, presentations, conference proceedings and theses. In this context, Moorin and Holman (2008) use merged data sets from the Western Australia mortality register and the hospital morbidity data system to explore patterns of health expenditure for in-patient care in the last 3 years of life, while Zhao, Connors, Wright and Guthridge (2008) use data obtained by linking a primary care chronic disease register with hospital inpatient databases to determine the 2005 prevalence rates of chronic diseases for the remote indigenous population of the Northern Territory. In all of these linkage applications, different data sets relating to the same individuals at different points in time are linked to provide a longitudinal data record for each individual, thus permitting longitudinal analysis.

However, the use of probabilistic linkage raises issues about potential biases induced by linkage errors. In particular, there is always the possibility that linkage errors in the merged data could lead to a longitudinal record ostensibly relating to a single individual being actually made up of a composite of data items from different individuals. Optimal probability-based linkage, based on maximising the probability of a declared link being correct, is described in Fellegi and Sunter (1969). Unfortunately, most practical implementations of their approach require one to trade off the number of links made against their accuracy. As a consequence, any implementation of probabilistic linkage will result in unmade linkages or non-linkages, as well as linkage errors where linkages are actually made. Furthermore, even a small amount of mismatching can lead to significant response errors, as demonstrated by Neter, Maynes and Ramanathan (1965). In this context, Scheuren and Winkler (1993, 1997) and Lahiri and Larsen (2005) investigate methods for eliminating the resulting bias in the context of regression analysis based on data from two probabilistically linked data sets. In this talk we will describe how the framework for inference using probability-linked data described in Chambers (2009) can be extended to define methods for efficient bias-corrected regression analysis when three registers are linked, or when a sample is linked to two separate registers, enabling longitudinal analysis. Our development allows for correlated linkage errors as well as errors arising when not all records can be linked. It can also be extended to when more than three data sets are linked.

BIBLIOGRAPHY

1. Brook, E. L., Rosman, D. L., & Holman, C. D. (2008). Public good through data linkage: measuring research output from the Western Australian data linkage system. *Australian and New Zealand Journal of Public Health*, 32 (1), 19–23.
2. Chambers, R. (2009). Regression analysis of probability-linked data. *Statisphere*, 4, Official Statistics Research Series, Statistics New Zealand.
3. Fellegi, I. P. & Sunter, A. B. (1969). A theory for record linkage. *Journal of the American Statistical Association*, 64 (328), 1183–1210.

4. Lahiri, P. & Larsen, M. D. (2005). Regression analysis with linked data. *Journal of the American Statistical Association*, 100 (469), 222–230.
5. Moorin, R. E. & Holman, C. D. (2008). The cost of in-patient care in Western Australia in the last years of life: a population-based data linkage study. *Health Policy*, 85, 380–390.
6. Neter, J., Maynes, E. S. & Ramanathan, R. (1965). The effect of mismatching on the measurement of response error. *Journal of the American Statistical Association*, 60 (312), 1005–1027.
7. Scheuren, F. & Winkler, W. E. (1993). Regression analysis of data files that are computer matched. *Survey Methodology*, 19, 39–58.
8. Scheuren, F. & Winkler, W. E. (1997). Regression analysis of data files that are computer matched - Part II. *Survey Methodology*, 23, 157–165.
9. Wilkins, K., Shields, M. & Rothermann, M. (2009). Smoker's use of acute care hospitals - a prospective study. *Health Reports*, 20 (4), pp. 75–83.
10. Zhao, Y., Connors, C., Wright, J. & Guthridge, S. (2008). Estimating chronic disease prevalence among the remote aboriginal population of the northern territory using multiple data sources. *Australian and New Zealand Journal of Public Health*, 32 (4), 307–313.

Chlebicki Paweł (ESRI Polska)

ArcGIS jako potężne narzędzie do wizualizacji i analizy danych statystycznych

Wszystkie dane statystyczne w pewnym sensie odnoszą się do przestrzeni geograficznej, a co z tym związane analiza tych danych powinna uwzględniać również aspekt przestrzenny. Zwykle informacje publikowane w formie tabelarycznej, bez odwzorowania ich na mapach, nie pozwalają na dostrzeżenie wielu zależności pomiędzy tymi danymi. Dlatego coraz częściej w statystyce wykorzystywane jest oprogramowanie GIS (System Informacji Geograficznej), które umożliwia wizualizację danych statystycznych na mapach oraz zaawansowane analizy przestrzenne tych danych. W prezentacji zostaną pokazane korzyści, jakie niesie ze sobą wykorzystanie do tego celu oprogramowania ArcGIS firmy Esri – światowego lidera w dziedzinie oprogramowania GIS

ArcGIS as a powerful tool for visualization and spatial analysis of statistical data

All of statistical data in some way refer to geographical space and consequently analysis of this data should also take into consideration geographical aspect. Usually simple tabular data, without visualization on the maps, don't allow to notice many relationships between this data.

For that reason in statistics GIS (Geographic Information System), which allows to visualize statistical data on the maps and also advanced analysis of this data, is used more and more frequently. During this presentation will be shown advantages of using ArcGIS by Esri – the world leader in GIS software.

Chłoń-Domińczak Agnieszka (Szkola Główna Handlowa w Warszawie)

Wykorzystanie wskaźników w obszarze edukacji w kontekście rozwoju polityki opartej na faktach

Artykuł ma na celu prezentację i analizę wskaźników które mogą być używane do monitorowania oraz ewaluacji polityki edukacyjnej w Polsce oraz efektów aktywności edukacyjnej. Analiza będzie się skupiać na istniejących źródłach danych, w tym danych administracyjnych (System Informacji Oświatowej) oraz dane z badań reprezentacyjnych, takich jak Badanie Aktywności Ekonomicznej Ludności, Badanie Dochodów i Warunków Życia (EU-SILC), badania międzynarodowe (badanie PISA, badanie Generation and Gender Survey) oraz badania krajowe (np. Diagnoza Społeczna).

Edukacja jest obszarem polityki publicznej, który ma wpływ na wiele obszarów rozwoju społecznego. W artykule chciała wskazać na określone grupy wskaźników z wybranych obszarów, zarówno z perspektywy indywidualnego uczestnictwa w edukacji i efektów tego uczestnictwa, jak również międzypokoleniowe aspekty aktywności edukacyjnej gospodarstw domowych.

Pierwszym obszarem jest uczestnictwo w uczeniu się przez całe życie, w tym udział w edukacji dzieci i młodzieży, ale także uczenie się dorosłych w różnych formach uczenia się formalnego, pozaformalnego oraz nieformalnego. W tym celu ważne jest wskazanie uwagi na potrzebę precyzyjnego zdefiniowania różnych form uczenia się tak, aby możliwe było ich właściwe monitorowanie. Uwzględnione zostanie również międzypokoleniowy kontekst poziomu wykształcanie, uwzględniając wpływ wykształcenia rodziców na wykształcenie dzieci. Analiza uwzględni również charakterystyki demograficzne (wiek, płeć, niepełnosprawność, stan cywilny).

Drugi obszar uwzględnia powiązanie pomiędzy edukacją i rynkiem pracy. W szczególności, dotyczy to statusu na rynku pracy osób w zależności od wykształcenia, ale też zależności pomiędzy aktywnością edukacyjną dzieci i statusem rodziców na rynku pracy. Uwzględniony zostanie również rola pracodawców w organizowaniu uczenia się przez całe życie pracowników, jak również przechodzenie pomiędzy edukacją a rynkiem pracy.

Trzeci obszar uwzględnia zasoby gospodarstw domowych i ich wykorzystanie na cele edukacyjne, w tym również deprywacja gospodarstw domowych dotycząca możliwości uczestnictwa w edukacji. Zwrócona zostanie uwaga zarówno na zasoby finansowe, jak i czasowe gospodarstw domowych przeznaczone na edukację osób dorosłych oraz dzieci.

Czwarty obszar dotyczy rozwoju kapitału ludzkiego osób, w tym ich umiejętności i kompetencji (językowych, ICT). Piąty obszar dotyczy terytorialnych aspektów polityki edukacyjnej, takich

jak zróżnicowanie wyników edukacyjnych lub dostęp do wybranych usług edukacyjnych.

W podsumowaniu wskazane są rekomendacje dotyczące dalszego rozwoju statystyki edukacyjnej, w tym rekomendacje badań mających na celu wypełnienie istniejących luk informacyjnych.

Słowa kluczowe: edukacja, statystyka społeczna, wskaźniki

The use of indicators in education in the context of development of evidence-based social policy

The article will focus on the presentation and analysis of indicators that can be used to monitor and evaluate the development of education policy in Poland as well as outcomes of education activities. The analysis will concentrate on the existing sources of information covering administrative data (such as System of Education Information (SIO)) or survey data, including surveys within public statistics such as Labour Force Survey or EU SILC), international surveys (for example PISA or Generations and Gender Survey) and national surveys (for example Diagnoza Społeczna).

Education is an area of public policy that has an impact on many spheres of social developments. Thus, in the article, I will focus on specific groups of indicators covering selected areas, both from the perspective of individual participation in education and the outcomes of such participation as well as inter-generational aspects of educational activities within households.

The first area is the participation in the life-long learning, including initial education and participation in the education of children and youth, but also further education of adults and their participation in various forms of life-long learning activities in formal and non-formal education as well as informal learning. In this context, I will also discuss the need of defining various forms of learning in the way that would allow for their monitoring. Special focus will be also put on the inter-generational context of educational attainment, including the influence of parents' education on the education of their children. The analysis will also take into account characteristics of individuals (age, gender, disability, marital status).

The second area will be the links between education and labour market. Analysis will cover both the labour market status of individuals depending on their educational attainment, but also the relation between the education activities of children and parents' labour market participation. It will also take into account the role of employers' in the development of life-long learning as well as the transition between education and labour market.

The third area will focus on the resources of the individuals and households, including those used for education purposes, as well as educational deprivation of households. This will cover not only material and financial resources but also time resources of households' members devoted to their own education or to the education of other households' members (in particular children).

The fourth area covers the development of human capital of individuals and households, including in particular development of core skills and competences (language ability, ICT).

The fifth area will cover territorial aspects of educational policy developments, such as school performance and access to educational services.

As a result, the proposal for future development of statistics in education will be formulated, including among others, the development of new sources of information that should fill the existing information gaps.

Key words: education, social statistics, indicators

BIBLIOGRAPHY

1. Czapiński J., Panek T. (2009) Diagnoza społeczna 2009 - warunki i jakość życia Polaków (oraz wcześniejsze raporty)
2. GUS (2011), Dochody i warunki życia ludności Polski (raport z badania EU-SILC 2009),
3. GUS, (2009), Kształcenie dorosłych, Warszawa
4. Kotowska, I. E., U. Sztanderska, I. Wóycicka (2007), Aktywność zawodowa i edukacyjna a obowiązki rodzinne kobiet w Polsce w świetle badań empirycznych, Scholar, Warsaw
5. OECD (2009), Jobs for Youth. Poland, Paris
6. OECD (2011), Education at a Glance, Paris
7. Szapiro, T. (red.), (2004), Mechanizmy kształtujące decyzje edukacyjne, Wydawnictwo SGH, Warsaw
8. Witkowski, J. (red.), 2008, Badanie aktywności zawodowej absolwentów w kontekście realizacji programu „Pierwsza praca”, Warszawa

Cmela Piotr, Kornecki Janusz (Urząd Statystyczny w Łodzi)

Strategiczne dylematy wsparcia rozwoju mikroprzedsiębiorstw w Polsce

Mikroprzedsiębiorstwa są uznawane za kluczowy czynnik rozwoju gospodarki Polski. Ta kategoria podmiotów gospodarczych generuje około jednej piątej wartości dodanej ogółu przedsiębiorstw, stanowiąc przy tym ponad dziewięć na dziesięć wszystkich firm, dając zatrudnienie jednej piątej ogółu zatrudnionych w przedsiębiorstwach. Jak pokazują dostępne dane statystyczne, cechą firm tego sektora w Polsce jest relatywnie niski poziom rozwoju w porównaniu do innych krajów UE. Polskie mikroprzedsiębiorstwa są mniejsze niż ich odpowiednicy za granicą, zatrudniają mniejszą liczbę pracowników i osiągają niższe obroty i wartość bilansową. Doceniając ich istotną rolę społeczno-gospodarczą, firmy mikro są ważnym adresatem polityki wsparcia. W referacie zostanie podjęta próba oceny dotychczasowej polityki z uwzględnieniem stopnia jej dopasowania do cech tej kategorii firm. Jednocześnie postawiono kwestię charakteru tego wsparcia, który z jednej strony, pozwoli zdynamizować rozwój tego sektora przedsiębiorstw, a z drugiej strony, nie zakłóci wolnorynkowej konkurencji.

Słowa kluczowe: mikroprzedsiębiorstwa, polityka wsparcia

Strategic dilemmas of growth support for micro-enterprises in Poland

Micro enterprises are regarded to be a key factor in the development of Polish economy. Business entities belonging to this category generate about one fifth of GVA of all enterprises, constitute nine out of ten of all businesses, offer employment to every fifth worker employed in enterprises. As statistical figures show, this category of enterprises is underdeveloped as compared with their EU counterparts. Polish micro enterprises employ fewer workers, generate lower turnover and show lower balance sheet value. Appreciating their important socio-economic role, micro firms are an important target of support policy. The paper attempts to evaluate current policies, including the degree of their match to the characteristics of this category of enterprises. Simultaneously, a question of the nature of support was raised so that it boosts the development of micro enterprises, on one hand, and does not distort market competition on the other hand.

Key words: micro enterprises, support policy

Czyżewski Andrzej, Grzelak Aleksander (Uniwersytet Ekonomiczny w Poznaniu)

Możliwości wykorzystania statystyki bilansów przepływów międzygałęziowych dla makroekonomicznych ocen w gospodarce

W artykule zaprezentowano możliwości wykorzystania modelu przepływów międzygałęziowych w zakresie ocen makroekonomicznych gospodarki. Zasadniczym przesłaniem bilansu przepływów międzygałęziowych jest ukazanie wzajemnych zależności w gospodarce, które decydują o jej rozwoju jako całości, jak i jej części. Stwierdzono, że oceny dokonywane za pośrednictwem bilansów przepływów międzygałęziowych umożliwiają i poszerzają perspektywę badawczą uwzględniając pozycję badanych sektorów (grup produktów) w gospodarce, a także ich efektywność makroekonomiczną i współzależności w procesie rozwoju. Dostrzeżono, że model przepływów międzygałęziowych wykazuje znaczące walory dla pełniejszego zrozumienia funkcjonowania mechanizmu rynkowego i budżetowego, schematu przepływu strumieni dochodów i wydatków w gospodarce oraz kształtowania się podstawowych parametrów makroekonomicznych. Istnieje także potrzeba poszerzenia zapisu i interpretacji tabeli przepływów o ćwiartkę czwartą (podział dochodów), szersze uwzględnienie powiązań badanych podmiotów z sektorem bankowym i tym samym pokazanie skali i struktur kreacji pieniądza emisyjnego i kredytowego w gospodarce. Do rozważenia pozostaje także wykorzystanie tabeli input-output do określenia zakresu nierównowagi podażowej w gospodarce poprzez zestawienie popytu potencjalnego i efektywnego oraz określenie w ten sposób przybliżonego poziomu nierównowagi.

Słowa kluczowe: model przepływów międzygałęziowych, gospodarka, makroekonomia

The possibilities of using the statistic of input-output in macroeconomics evaluation of the economy

The possibilities of using of the input-output table for macroeconomics evaluation of economy was presented in the article. The principal message of input-output model is appearance mutual interdependence in economy, which decide about her development as the whole, and as her part. It has stated that the evaluation conducted by use of this model makes possible extension of investigative perspective regarding such questions as: position of a examined sectors in the economy, macroeconomics efficiency, and interdependences in developmental processes. One had noticed that the input-output model points out considerable values in for fuller understanding of functioning of the market and budget mechanism, the roundabout pattern of streams of incomes and the expenses in economy, and the shaping of the basic macroeconomics parameters. There is a need extension of record and the interpretation of input-output model about fourth quarter (the division of incomes), the wider regard of connections of studied subjects with bank sector and the same show the scale and the structures of creation of issue and credit money in the economy. It stays for considering also the use of this model for evaluation the range of the supply non-equilibrium in the economy by composition of potential demand and effective and qualification in this way an approximate level the non-equilibrium.

Key words: the input-output model, economy, macroeconomics

Decker Reinhold (Universität Bielefeld)

Model-based Analysis of Online Consumer Reviews. Methods and Applications

The statistical analysis of online consumer reviews has attracted increasing attention in quantitative marketing research in the recent past. Typically, consumer reviews, as posted on websites like amazon.com, epinions.com, or airlinequality.com, consist of different components which are assumed to be interrelated, particularly product ("star") ratings, pros and cons, as well as free format texts. Analyzing such data can create a comprehensive picture of consumers' opinions and preferences regarding the products or services of interest. Against this background the talk deals with three topics, demonstrating the usefulness of statistical online consumer review analysis and inevitable limitations due to the specificity of this data.

The first part of the talk is about data quality and completeness, with a special focus on the general reliability of („star”) ratings as the key element of most online consumer reviews. We suggest a new framework, based on a beta binominal model, which accommodates the uncertainty usually inhering posted ratings and allows inferences about their informative value. Beyond, it is shown how the calibrated model can used to estimate product recommendation probabilities. The latter play an important role in online word-of-mouth and therefore are of special interest for marketing managers who want to know whether the consumers would tend to recommend the product of interest, given an observed distribution of ratings. The usefulness

and limitations of this approach are demonstrated using synthetic as well as real data from a popular airline review site.

Another important issue of online review analysis is opinion heterogeneity. Similar to the concept of consumer heterogeneity in buying behavior research, a methodology based on negative binomial regression is presented. The suggested approach enables the estimation of parameters, which allow inferences on the relative effect of product attributes and brand names on product evaluations. It is shown how the parameter estimates can be used to identify those attributes of a product which feature potentials for product improvement and redesign. In order to illustrate this approach we make use of the pros and cons collected from thousands of online consumer reviews on a widespread electronic device.

The third part of the talk is devoted to differences in consumer opinions across countries. By means of ordered probit and logistic models we show that the product attributes being considered important may diverge significantly across borders. As an empirical basis for this research, a large data set is used, which contains several thousands of online consumer reviews posted in different countries. The available results clearly underline the advisability of considering both the usual („star”) ratings and the willingness to recommend a product and, there-with, complement recent findings in this field of marketing data analysis.

Dehnel Grażyna (Uniwersytet Ekonomiczny w Poznaniu)

Estymacja pośrednia w krótkookresowej statystyce przedsiębiorstw

Zachodzące obecnie przekształcenia systemu informacji z zakresu statystyki gospodarczej zmierzają w kierunku szerokiego wykorzystania przez statystykę publiczną danych administracyjnych. Informacje zawarte w rejestrach administracyjnych mają wpłynąć na poprawę efektywności badań statystycznych przedsiębiorstw. Chodzi tu między innymi o zmniejszenie obciążeń podmiotów gospodarczych wynikających ze sprawozdawczości statystycznej, zmniejszenie kosztów pozyskania danych.

Wszystkie korzyści wynikające z wykorzystania informacji zawartych w rejestrach administracyjnych nabierają szczególnego znaczenia w odniesieniu do statystyki krótkookresowej, gdzie konieczne jest wypracowanie nowych metod szacunku oraz systematyczności i odpowiedniej częstotliwości w przekazywaniu informacji przez gestorów. Stąd konieczność prowadzenia prac mających na celu wypracowanie rozwiązań metodologicznych pozwalających na podniesienie efektywności systemu badań statystycznych statystyki krótkookresowej.

Jednym z głównych celów referatu jest wskazanie możliwości wykorzystania zasobów rejestrów administracyjnych dla poprawy krótkookresowej statystyki przedsiębiorstw. Wśród szczegółowych celów należy wskazać: ocenę jakości danych pozyskiwanych w ramach badania DG1, na podstawie źródeł administracyjnych oraz próbę wykorzystania metodologii statystyki małych

obszarów do szacunku podstawowych charakterystyk przedsiębiorstw w odniesieniu do statystyki krótkookresowej.

Słowa kluczowe: estymacja pośrednia, statystyka krótkookresowa, statystyka gospodarcza

Indirect estimation in short-term statistics

Changes in the information system of business statistics are geared towards a greater reliance on administrative data. Information stored in administrative register is expected to improve the effectiveness of business statistics. The changes are designed to reduce the response burden for companies owing to statistical reporting, reduce the cost of obtaining data.

All the benefits of using administrative data are especially significant when it comes to short-term statistics, which requires adequate estimation methods as well as systematic and timely provision of data by its administrators. This explains the need for continued efforts to find methodological solutions that will improve the effectiveness of the short-term statistics research system.

The main objective of the paper is to highlight the possibility of using administrative registers to improve short term business statistics, in particular: to assess the quality of data collected by means of the DG1 survey, based on administrative registers and to apply small area estimation methodology to estimate basic characteristics of enterprises in terms of short-term statistics.

Key words: indirect estimation, short-term statistics, business statistics

Didkowska Joanna, Wojciechowska Urszula (Centrum Onkologii - Instytut im. Marii Skłodowskiej-Curie w Warszawie)

Nowotwory złośliwe w Polsce jako problem zdrowia publicznego. Współczesne narzędzia do mierzenia zjawiska.

Nowotwory złośliwe w Polsce stanowią istotny problem społeczny nie tylko u osób w starszych grupach wieku, ale są główną przyczyną przedwczesnej umieralności przed 65. rokiem życia. Na tym tle Polska negatywnie wyróżnia się wśród krajów europejskich. Zjawisko to jest szczególnie widoczne w populacji kobiet - nowotwory są już od kilku lat najczęstszą przyczyną zgonów przed 65. rokiem życia (stanowią około 30% zgonów w grupie młodych kobiet i prawie 50% zgonów wśród kobiet w średnim wieku). Do końca tej dekady nowotwory prawdopodobnie staną się także najczęstszą przyczyną przedwczesnej umieralności mężczyzn.

Przewidywany wzrost liczby zachorowań i zgonów z powodu nowotworów spowodowany jest głównie zmianami demograficznymi w polskiej populacji (wzrost liczby i frakcji osób po 65.

roku życia). Według prognozy Głównego Urzędu Statystycznego liczba osób w tej grupie wiekowej wzrośnie do 2025 roku o prawie 75%. Zwiększenie populacji starszych mężczyzn przekłada się na wzrost liczby przypadków nowotworów występujących u osób w starszym wieku (gruczołu krokowego, jelita grubego i płuca). Podobnie wśród kobiet należy liczyć się ze znacznym wzrostem liczby zachorowań (o czym decydują rosnące trendy zachorowalności na raka płuca, piersi, trzonu macicy).

Zachodzące zmiany w stanie zdrowia populacji polskiej wymagają zbudowania silnego narzędzia do monitorowania i analizy zachodzących procesów. W przypadku nowotworów takim narzędziem jest rejestr nowotworów złośliwych.

Obecna struktura systemu rejestracji nowotworów w Polsce (16 niezależnych wojewódzkich baz danych), w dobie procesów cyfryzacji informacji medycznej, wymaga zmian, które pozwolą na korzystanie z nowoczesnych metod wymiany informacji pomiędzy instytucjami włączonymi do systemu informacji w ochronie zdrowia wdrażanemu w Polsce (Dz. U. Nr 113, poz. 657). W odpowiedzi na taką potrzebę Krajowy Rejestr Nowotworów rozpoczął projekt pt. „Utworzenie pierwszej w Polsce informatycznej platformy naukowej do wymiany wiedzy o zagrożeniu nowotworami złośliwymi w Polsce” współfinansowany ze środków EFRR w ramach Programu Operacyjnego Innowacyjna Gospodarka. Celem projektu jest stworzenie centralnej bazy zachorowań na nowotwory złośliwe, dającej podstawy do prowadzenia nowoczesnych badań epidemiologicznych z wykorzystaniem technologii społeczeństwa informacyjnego.

Obecna infrastruktura ochrony zdrowia w Polsce jest tylko w niewielkim stopniu przygotowana na „eksplozję” liczby nowych zachorowań i chorych na nowotwory złośliwe. Prognozy zachorowalności na nowotwory, wskazują na konieczność przygotowania się na gwałtownie rosnącą liczbę pacjentów z chorobą nowotworową pod względem edukacyjnym, diagnostycznym, leczniczym, ekonomicznym i logistycznym. Skala tego zjawiska wymaga udziału państwowych struktur ochrony zdrowia. Oprócz konieczności diagnozy i leczenia przeciwnowotworowego chorzy będą wymagali długoletniej opieki lekarskiej i pielęgniarstwa, a w wielu przypadkach również hospicyjnej. Żadne z wymienionych działań nie może być planowane bez rzetelnej, bieżącej informacji o zjawisku, którą może zaoferować jedynie dobrze funkcjonujący rejestr nowotworów złośliwych.

Słowa kluczowe: nowotwory złośliwe, zachorowalność, umieralność, prognozy, przedwczesna umieralność, Krajowy Rejestr Nowotworów

Cancer in Poland as a public health problem. Modern tools to measure the phenomenon.

Malignant neoplasms are in Poland a major social problem not only in the older age groups, but they are the main cause of premature mortality (before age 65) and, therefore, Poland negatively stands out among the European countries. This phenomenon is observed particularly in women - for several years cancers have been leading cause of death before age 65 (about 30% of deaths among young women and nearly 50% of deaths among middle-aged women). By the end of the XX century, cancers are likely to become the main cause of premature mortality also in male population.

The projected increase in the number of cancer cases and deaths is attributable to demographic changes of the Polish population (increase of the number and the fraction of people over 65 years of age. Central Statistical Office in Poland estimates that the number of people in this age group will increase by almost 75% by 2025. Increasing in the population of older men translates into an increase in the number of new cases of the most common cancers in elderly (prostate, colon and lung). Similarly, among women it should be expected a significant increase in incidence (growing trends in the incidence of lung cancer, breast, endometrial).

The health phenomena taking place in the Polish population require building a strong tool to monitor and analyze the processes involved. In case of cancer such tool is the cancer registry.

The current structure of the system of cancer registration in Poland (the network of 16 regional databases), in the time of digitization of medical information, requires changes that will allow using modern methods of exchanging information among medical centers in Poland. In response to this need the National Cancer Registry launched a project: „The creation of the first Polish scientific information platform to exchange knowledge on the risks of cancer in Poland”, supported by the ERDF in the Operational Programme Innovative Economy. The project aims to create a central database of cancers, establishing social baseline of conducting modern epidemiological research using information society technologies.

The current health care infrastructure in Poland is only marginally prepared for the explosion of new cancer cases. Projections of cancer trends indicate the need to prepare for the rapidly growing number of cancer patients in educational, diagnostic, therapeutic, economic and logistics terms. The scale of this phenomenon requires the participation of public health structures. Patients diagnosed with cancer will require treatment, long-term medical care, social services and in many cases, hospice care. None of these actions can be planned without a reliable, current information on the phenomenon, which can offer only a well-functioning cancer registry.

Key words: cancer, incidence, mortality, forecast, premature mortality, National Cancer Registry

Doniec Dorota (Urząd Statystyczny w Katowicach)

Rachunki regionalne

W statystyce publicznej priorytetową rolę odgrywają badania prowadzone w ramach systemu rachunków narodowych oraz badania regionalne. Wśród nich szczególną wagę mają rachunki regionalne, które zawierają istotne cechy zarówno systemu rachunków narodowych, jak również badań regionalnych. Rachunki regionalne są bowiem przestrzenną specyfikacją rachunków narodowych – podsystemem rachunków narodowych, w którym działalność ekonomiczna wszystkich podmiotów gospodarki narodowej jest grupowana w przekroju terytorialnym.

Rachunki regionalne dostarczają spójnych i kompleksowych informacji do prowadzenia analiz

ekonomicznych niezbędnych do ocen rozwoju na szczeblu regionalnym. Jest to ważne źródło informacji o zróżnicowaniu poziomu rozwoju gospodarczego regionów, województw i podregionów oraz o zmianach zachodzących w strukturze gospodarki w ujęciu przestrzennym. Podstawowym wskaźnikiem makroekonomicznym charakteryzującym w sposób syntetyczny sytuację gospodarki narodowej ogółem oraz w przekroju terytorialnym jest produkt krajowy brutto (PKB).

Na podstawie danych z rachunków regionalnych możliwe jest międzynarodowe porównanie produktu krajowego brutto dla poszczególnych jednostek podziału terytorialnego Polski do średniego poziomu krajów UE oraz innych regionów państw Wspólnoty.

Wyniki badania z rachunków regionalnych wspierają procesy podejmowania decyzji w obszarze polityki regionalnej. Dane dotyczące regionalnego PKB są wykorzystywane dla potrzeb funduszy strukturalnych UE do selekcji i wyboru regionów kwalifikujących się do wsparcia finansowego oraz do monitorowania rezultatów udzielonej pomocy, a także do określania części regionalnej subwencji ogólnej dla województw. Służą one również do opracowywania i monitorowania wielu strategii rozwoju.

W ostatnich latach w systemie rachunków regionalnych wprowadzone zostały istotne zmiany. Artykuł ma na celu przedstawienie aktualnego stanu rachunków regionalnych – zakresu obliczanych kategorii makroekonomicznych i dostępnych danych wynikowych oraz dalszych zamierzeń i kierunków prac metodologicznych.

Regional accounts

In public statistics surveys conducted within national accounts system as well as regional surveys play a priority role. Among them, regional accounts, which include essential features of both national accounts system and regional surveys, are of a great importance. Regional accounts are regional specification of national accounts – the subsystem of national accounts, in which economic activities of all entities of national economy are grouped by territorial division.

Regional accounts provide cohesive and comprehensive information for carrying out economic analysis indispensable in development evaluating on the regional level. They constitute an important source of information on the diversity of economic development level of regions, voivodships and subregions as well as on changes in the structure of economy in spatial approach. Gross domestic product (GDP) is a basic macroeconomic indicator characterizing in the synthetic way situation of national economy as a whole as well as in the territorial breakdown.

International comparisons of gross domestic product for particular units of territorial division of Poland with the average level of EU countries as well as with other regions of Community countries are possible on the basis of data from regional accounts.

Results from regional accounts support decision-making processes in the field of regional policy. Data concerning regional GDP are used for EU structural funds purposes, for selection and choosing regions eligible for financial support, for monitoring results of granted aid as well as

for defining regional part of total subsidy for voivodships. They serve also for preparing and monitoring many development strategies.

In the last years essential changes were introduced in the regional accounts system. The aim of this article is to present current state of regional accounts – range of calculated macroeconomic categories and available result data as well as further intentions and directions of methodological works.

Dudek Hanna, Landmesser Joanna (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie)

Satysfakcja z osiągniętych dochodów a względna deprivacja

Analiza subiektywnego postrzegania własnej sytuacji dochodowej stanowi stosunkowo nowy, lecz dynamicznie rozwijający się dział statystyki społecznej. Badania dobrobytu gospodarstw domowych w wielu rozwiniętych krajach nie ograniczają się jedynie do analizy konsumpcji i obiektywnej kondycji dochodowej, lecz uwzględniają także zagadnienie deprivacji obejmującej wiele sfer życia, w tym deprivacji subiektywnej. Poznanie uwarunkowań satysfakcji z osiągniętych dochodów może pomóc w kształtowaniu polityki socjalnej mającej na celu złagodzenie skutków subiektywnego ubóstwa.

W pracy podjęto temat określenia determinant subiektywnej oceny dochodów przez polskie gospodarstwa domowe. W tym celu poddano analizie sytuację dochodową gospodarstw, zarówno w ujęciu bezwzględnym jak i względnym. W literaturze przedmiotu sugeruje się bowiem, że ta druga może mieć silniejszy związek z subiektywną percepcją swojego bytu niż pierwsza [Easterlin 1995]. Dlatego też dla każdego z gospodarstw domowych wyznaczono wartość funkcji relatywnej deprivacji dochodowej w formie zaproponowanej w pracy [D'Ambrosio, Frick 2007]: $d(y_{(i)}) = \frac{1}{n} \sum_{j=i+1}^n (y_{(j)} - y_{(i)})$, gdzie $y_{(j)}$, $y_{(i)}$ – dochody ekwiwalentne i-tego oraz j-tego gospodarstwa, przy czym: $y_{(1)} \leq y_{(2)} \leq \dots \leq y_{(n)}$. Ponadto, w charakterze determinant subiektywnej oceny dochodów rozpatrzono wiele cech socjodemograficznych takich jak wiek, płeć oraz charakterystyki odnoszące się do statusu członków gospodarstw domowych na rynku pracy. W artykule przedstawiono wyniki weryfikacji hipotezy, czy efekt wpływu relatywnej deprivacji na subiektywną ocenę dochodów jest taki sam w różnych grupach określonych przez w/w cechy.

W pracy zastosowano metodę częściowo uogólnionych uporządkowanych modeli logitowych. Analizę przeprowadzono na podstawie danych z badań budżetów gospodarstw domowych zrealizowanych przez GUS w 2009 roku. Proponowana metodyka jest modyfikacją podejścia zaprezentowanego w publikacjach [D'Ambrosio, Frick 2007] i [Ferrer-i-Carbonell 2005], gdzie do estymacji parametrów modelu objaśniającego satysfakcję z osiągniętych dochodów wykorzystano uporządkowane modele probitowe i logitowe nie weryfikując dość silnych założeń nakładanych na te modele.

Słowa kluczowe: satysfakcja z osiągniętych dochodów, względna deprywacja

Literatura:

- D'Ambrosio C., Frick J., Income satisfaction and relative deprivation: An empirical link, "Social Indicators Research", 2007, vol. 81, nr 3, pp. 497–519.
- Easterlin R.A., Will raising the incomes of all increase the happiness of all?, "Journal of Economic Behavior & Organization", 1995, vol. 27, pp. 35-47.
- Ferrer-i-Carbonell A., Income and well-being: An empirical analysis of the comparison income effect, "Journal of Public Economics", 2005, vol. 89, pp. 997-1019.

Income satisfaction and relative deprivation

Researches on income satisfaction are sparse but growing rapidly. Many studies of household wealth in developed countries are not limited to analysis of consumption and objective income condition, but also take into account the issue of deprivation covering many areas of life, including subjective deprivation. Understanding the determinants of satisfaction with incomes may help in creation of social policy aimed at mitigating the effects of subjective poverty.

The study's main objectives is the identification of determinants of income satisfaction in Poland. For this purpose, income situation of households, both in absolute and relative terms, is analysed, because some authors suggest that the natural relationship is more between subjective well-being and relative deprivation rather than between subjective well-being and income itself [Easterlin 1995]. For each household with income $y(i)$ value of the relative deprivation function is computed as follows: $d(y(i)) = \frac{1}{n} \sum_{j=i+1}^n (y(j) - y(i))$, where $y(j)$, $y(i)$ – equivalent incomes of i -th and j -th household, $y_{(1)} \leq y_{(2)} \leq \dots \leq y_{(n)}$ [D'Ambrosio, Frick 2007]. Moreover, regressors in the econometric model include socio-demographic characteristics such as gender, age, marital status, and dummies for employment status. The effects of relative deprivation on income satisfaction in various socio-demographic groups of households is examined.

The method of partial generalized ordered logit models is used. The empirical investigation is based on data from household budget surveys carried out by GUS in 2009. The proposed methodology is a modification of the approach presented in [D'Ambrosio, Frick 2007] and [Ferrer-i-Carbonell 2005].

Key words: income satisfaction, relative deprivation

Dygaszewicz Janusz, Janczur-Knapiek Magdalena, Wardzińska-Sharif Amelia (Główny Urząd Statystyczny)

Spisys powszechne PSR 2010 i NSP 2011 oraz systemy informacji geograficznej w statystyce publicznej

Lata 2010–2011 były dla statystyki publicznej okresem bardzo wyťažonej pracy. W 2010 roku - od 1 września do 31 października - został przeprowadzony Powszechny Spis Rolny (PSR 2010). W 2011 roku, w okresie od 1 kwietnia do 30 czerwca, został przeprowadzony Narodowy Spis Powszechny Ludności i Mieszkań (NSP 2011). Mając na względzie konieczność oszczędnego gospodarowania środkami finansowymi, w realizacji spisów powszechnych wprowadzono nowoczesne i tańsze rozwiązania niż były stosowane wcześniej. Opierają się one na pozyskaniu danych od 16 gestorów - z 25 systemów informacyjnych, w tym administracyjnych, oraz wykorzystaniu narzędzi komunikacji elektronicznej. Całkowicie wyeliminowane zostały formularze papierowe. Pozwoliło to na zmniejszenie obciążenia respondentów, jak i na ograniczenie kosztów druku materiałów spisowych. W spisach powszechnych zastosowane zostały następujące kanały pozyskiwania danych:

- źródła administracyjne,
- Internet (CAII - *Computer Assisted Internet Interview*, samospis internetowy),
- wywiad telefoniczny (CATI - *Computer Assisted Telephone Interview*),
- spis za pośrednictwem rachmistrza (CAPI - *Computer Assisted Personal Interview*) - wyposażonego w terminal przenośny typu hand-held.

Innowacją wprowadzoną w obydwu spisach powszechnych było wykorzystanie na każdym etapie prac technologii GIS (Geographic Information System). Mapy cyfrowe były niezbędnym narzędziem pracy dla rachmistrzów spisowych (w zakresie poruszania się w terenie, weryfikacji operatu itd.), liderów gminnych oraz dyspozytorów wojewódzkich i centralnych, którzy na mapie mogli weryfikować postęp spisu oraz np. marszrutę lub położenie rachmistrza. Aplikacje GIS wykorzystywały materiały pozyskane z PZGiK (ortofotmapę, granice województw, powiatów i gmin, nazwy miejscowości), granice rejonów statystycznych i obwodów spisowych oraz statystyczne punkty adresowe przygotowane przez służby statystyki publicznej, a także warstwę działek ewidencyjnych (ARiMR) oraz drogi i ulice.

Mapy cyfrowe w technologii GIS miały zastosowanie w obydwu spisach podczas następujących etapów prac:

1. aktualizacji gminnej, dzięki której urzędy gmin miały możliwość edycji punktów adresowych on-line,
2. obchodu przedspisowego:

- w aplikacji rachmistrza spisowego do wprowadzania nowych punktów, edycji istniejących oraz usuwania nieistniejących w terenie,
- w aplikacji lidera gminnego AGMIS oraz aplikacji dyspozytorskiej ADYS, dzięki którym można było na bieżąco monitorować przebieg spisu i reagować na zaistniałe sytuacje,

3. samego spisu:

- w aplikacji rachmistrza spisowego do przeprowadzania wywiadu oraz
- w aplikacji lidera gminnego AGMIS oraz aplikacji dyspozytorskiej ADYS, pozwalając na monitorowanie i raportowanie przebiegu obchodu oraz spisu, a także planowanie i zarządzanie pracą rachmistrzów na podległych obszarach.

Słowa kluczowe: spis powszechny, system informacji geograficznej, wyniki spisów powszechnych, prezentacja danych statystycznych, GIS

National censuses PSR 2010 and NSP 2011 via geographic information systems in public statistic

The years 2010–2011 were a period of very hard work for public statistics. In 2010, from 1st of September to 31st of October Agricultural Census was carried out (PSR 2010). This year, in the period from 1st of April to 30th of June, the National Census of Population and Housing was conducted (Census 2011). With a view to necessity of efficient financial management, the implementation of the census introduced a modern and relatively cheaper solutions than were used previously. It focused on obtaining data from 16 holders - from 25 information systems (including the administration systems), and the use of electronic communication tools. Paper forms were completely eliminated. This allowed to reduce the burden on respondents, and to reduce printing costs of census materials. In connection with this method, following channels were used for data acquisition in the censuses:

- administrative sources,
- Internet interview (CAII - *Computer Assisted Internet Interview, self-enumeration On-line*),
- Telephone interview (CATI - *Computer Assisted Telephone Interview*),
- inventory through Enumerator (CAPI - *Computer Assisted Personal Interview*) - equipped with a portable terminal - hand-held

For the purpose of preparing and conducting national censuses, CSO implemented a Computer – based Census System (ISS). This system, comprising of over 10 components designed by various contractors, provided a computer assistance in the census process. ISS integrated various technologies (applications installed on mobile terminals, applications for administration and support of telephone interviews, specialized databases, data warehouses and analytical-report tools).

Furthermore, the solutions applied in ISS ensured the high level of security of the processed census data and implemented set of organizational measures obliging the census participants to comply with statistical confidentiality and to protect personal data.

An innovation introduced in both censuses was use of the GIS technology (Geographic Information System) at every stage. Digital maps were an essential tool for census enumerators (in the range of their movement in the field, verification of sampling frame etc.), municipal leaders and voivodship and central dispatchers, who could verify the progress of the census and for example a route or location of enumerator. GIS applications contained materials derived from PZGiK (orthophotmaps, the borders of voivodships, counties (powiats) and municipalities (gminas), and the names of towns), the borders of statistical regions and census enumeration areas and statistical address points, made by the official statistics services, as well as the parcels layer (ARiMR) and the roads and streets.

Digital maps in GIS technology were applied in both censuses during the following stages:

1. update in municipalities (gminas) - thanks to which municipalities (gminas) offices had the opportunity to edit the address points on-line;
2. pre-census field check:
 - in census enumerator application- to set new, edit existing and remove non- existing address points in the field,
 - in the gmina leaders application AGMIS and the dispatchers application ADYS through which it was possible to monitor the census progress up to date and respond to occurred situations;
3. census itself:
 - census enumerators application to carry out the interview and
 - in the gmina leaders application AGMIS and the dispatchers application ADYS, allowing to monitor and report progress of census and the pre-census field check, as well as to plan and manage of census enumerators work on subordinated areas.

Key words: census, geographic information system, census results, presentation of statistical data, GIS

Dyzmann-Sroka Agnieszka (Wielkopolskie Centrum Onkologii), Roszka Andrzej (Wielkopolskie Centrum Onkologii), Moczko Jerzy (Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu), Trojanowski Maciej (Wielkopolskie Centrum Onkologii)

Metody podniesienia jakości danych rejestrów nowotworów złośliwych

Nowotwory złośliwe będące trzecią na świecie, a w krajach rozwiniętych drugą przyczyną zgonów - stanowią obecnie istotny problem zdrowotny społeczeństw XXI wieku. Biorąc pod uwagę starzenie się populacji - problem ten będzie narastał.

Rejestracja nowotworów złośliwych odbywa się na Świecie od ponad stu-lat, rozpoczęła się w roku 1900 w Niemczech (w Polsce w 1951).

Jako główną bazę danych niezbędną dla realizacji wyznaczonych celów pracy przyjęto informacje zawarte w Wielkopolskim Rejestrze Nowotworów (WRN). W celu uzupełnienia danych zawartych w rejestrze korzystano z dodatkowych źródeł informacji, takich jak baza: Głównego Urzędu Statystycznego, Wydziału Meldunkowego Wielkopolskiego Urzędu Wojewódzkiego, Zakładów i Pracowni histopatologicznych, Dokumentacji medycznej Wielkopolskiego Centrum Onkologii. Analizę jakości bazy danych Wielkopolskiego Rejestru Nowotworów oparto o trzy podstawowe wskaźniki zaproponowane przez WHO: kompletność (K_r), jakość (HV%), przypadki Death Certificate Only (DCO).

W wyniku podjętych działań udało się opracować metody pozyskiwania z WRN wiarygodnych danych o zachorowalności, leczeniu i losach chorych na nowotwory złośliwe w Wielkopolsce.

Uzyskana kompletność (100% w roku 2007) i jakość rejestracji (odpowiednio 82%) przy równoczesnym obniżeniu odsetka przypadków zgłaszanych wyłącznie na podstawie aktu zgonu (DCO z 12% w roku 2000 do 7% w 2007) potwierdzają wiarygodność danych Wielkopolskiego Rejestru Nowotworów. Uzupełniona baza danych WRN pozwala na ocenę skuteczności leczenia chorych na podstawie 5-letnich przeżyć.

Słowa kluczowe: rejestracja nowotworów, jakość rejestracji, epidemiologia nowotworów

Income satisfaction and relative deprivation

Cancer, being the world's third and developed countries' second death cause, represents a major health issue for societies of the 21st century. The problem is likely to be even more severe with the population becoming older.

Cancer registration has been conducted for over a hundred years. It was first started in Germany

in 1900 (in Poland since 1951).

The main source of data used in pursuing the aims of the study was the database held by the Wielkopolska Cancer Registry (WCR). Complementary data were obtained from additional sources, such as: the Central Statistical Office, Population Registry Department of the Wielkopolska Voivodeship Office, histopathology laboratories and units, medical records of the Greater Poland Cancer Centre.

The quality assessment of the WCR data was based on three basic indicators as proposed by WHO: completeness (K_r), quality (HV%), and Death Certificate Only (DCO) cases.

Effective methods were developed to obtain reliable data from the WCR on incidence, treatment and follow-up of cancer cases in the region of Wielkopolska.

The level of completeness achieved (100% in 2007), registration quality (82% in the same year), combined with lower proportion of DCO cases (12% in 2000 vs. 7% in 2007), confirm high reliability of data collected by the Wielkopolska Cancer Registry. The completed WCR database allows to evaluate the treatment efficacy based on five-year survival rates.

Key words: cancer registration, registration quality, cancer epidemiology

BIBLIOGRAPHY

1. Boyle p., Levin B. World Cancer Report 2008, WHO, Lyon 2008,
2. Dyzmann-Sroka A. „Metody uzyskiwania i opracowania wiarygodnych danych pozwalających na ocenę skuteczności leczenia chorych w Wielkopolskim Centrum Onkologii”, Zeszyty Naukowe WCO Tom 8, Nr 2/2011,
3. GLOBOCAN 2008 v1.2, Cancer Incidence and Mortality Worldwide: IARC CancerBase No. 10 [Internet],
4. Wojciechowska U., Didkowska J., Zatoński W. Nowotwory złośliwe w Polsce w 2008 roku. Centrum Onkologii - Instytut im. M. Skłodowskiej-Curie, Warszawa 2010.

Fattorini Lorenzo (University of Siena)

Design-Based Inference on Ecological Diversity

Three main issues of ecological diversity analysis are considered: i) the problem of measuring ecological diversity by means of suitable indexes; ii) the problem of estimating these indexes when biological populations are sampled by means of sampling schemes actually adopted by biologists rather than by unrealistic schemes often supposed by statisticians just for their convenience; iii) the problem of comparing and ordering biological communities with respect to their diversity.

Even if a formal definition is still lacking, the concept of diversity traditionally relies on the apportionment of abundance into the animal or plant species forming the ecological community under study. In this contest it is of basic importance to define the term community. Indeed, it would be unfeasible to consider every living form in the area, from molecules to species. The abundance of nematodes cannot be compared with the abundance of oaks, pines or other tree species. Accordingly, a taxonomic group at a higher level than species (for examples families or classes) is usually chosen as the community of interest and the ecological diversity refers to all members of the groups which is usually called taxocene.

For a long time the primary aim of statisticians faced with ecological diversity has been its quantification by means of indexes which may take into account some important aspects of diversity such as evenness, dominance and rarity of species. To this purpose, a very effective and unifying approach for measuring diversity was offered by Patil and Taille (1982) on the basis of the intuition that a community is diverse when there is a large number of rare species. The authors propose measuring the rarity of each species and adopting the average community rarity as a diversity index. The most familiar diversity indexes such as Shannon's and Simpson's index can be viewed as rarity measures. Contemporary to this proposal, Rao (1982) furnished an axiomatization of diversity measures on the basis of their capability to split diversity between and within the sub-populations determined by a hierarchical classification of the community under study. The author defined a diversity measure to be perfect if the diversity splitting can be carried out for hierarchical classification of any order. These indexes are referred to as Rao's quadratic entropy and surprisingly they had a limited impact on diversity studies. Some anomalies of quadratic entropy have been pointed out in literature, showing that the problem of measuring diversity still remains very open.

But, irrespective of the index adopted, any diversity index H is a function of the relative abundance vector p which, in turn, is a function of the abundance vector N . However, knowledge of these quantities obviously requires a census of the ecological community under study, which is unfeasible in most cases. Accordingly, abundance is actually unknown and must be estimated by means of a sample survey in order to subsequently estimate H . A variety of results are available in literature regarding the estimation of diversity indexes when individuals are selected from the ecological community using simple random sampling with replacement. The reason is that under this scheme the species frequencies in the sample are optimal estimators of the relative abundances. Then, the subsequent inference on diversity indexes is straightforward. Unfortunately, these works are of few practical utility since animals and plants cannot be selected from the ecological community just like balls from an urn.

Few papers recognize this fact in the huge literature regarding the estimation of diversity indexes. Most of these papers are based on simulation studies or field works, while theoretical results of general validity are lacking. Barabesi and Fattorini (1998) attempt to give general results on diversity index estimation when N is estimated by plot, point or transect sampling. Indeed, in sampling ecological communities it is important to recognize that an ecological community on a study area constitutes a without-frame population of organisms spread over the area. Owing to the lack of frame, the most effective schemes for sampling ecological populations differ from the traditional schemes and their choice is mainly determined by practical considerations on the nature of the community to be sampled. Usually, in environmental schemes of sampling, the

selected units are those encountered from points or within plots or along transects randomly thrown onto the study area. For example when dealing with communities of trees, floating plot sampling are usually adopted while when dealing with shrub communities line intercept sampling is suitable. On the other hand, when dealing with animal communities, line transect sampling or point transect sampling should be performed to handle problems related to the elusive behaviour of animals. Unbiased estimators of N can be achieved with this scheme by means of the Horvitz-Thompson criterion.

Since a study area cannot be adequately sampled using one plot or one line or one point only, it is suitable to replicate the sampling procedure n times, in such a way that n plots or n lines or n points are randomly and independently thrown onto the study area, in order to obtain an adequate coverage of this area. In this framework, the arithmetic mean of the resulting estimates achieved in each replication obviously constitutes an unbiased, asymptotically normal and consistent estimator of N . Moreover, once the N estimator is obtained, the p estimator and subsequently the Δ estimator are asymptotically normal and consistent owing to delta method.

Even if asymptotically unbiased, bias in Δ estimators occurs for a finite number of replications. It is a stated result that most diversity indexes, when evaluated from sample surveys, heavily underestimate the population counterparts since most of the rare species are lost. Quoting from the standard results on jackknife, conditions are given under which the bias of the jackknifed estimators of diversity indexes are consistent, asymptotically normal with a bias of order smaller than n^{-1} . These results ensure a reliable inference on diversity indexes under sampling schemes actually adopted in the field.

However, single diversity indexes are not suitable for comparing ecological communities in that different indexes may lead to different rankings. In order to avoid inconsistent rankings, Patil and Talle (1982) introduced the concept of intrinsic diversity ordering, defining a community to be intrinsically more diverse than another if the second is obtained from the first through a finite sequence of the following operations: transferring abundance from more to less abundant species without reversing the rank-order of the species; transferring abundance to a new species; re-labelling the species. Subsequently, the authors prove that any intrinsic diversity ordering, if it exists, can be determined by means of the so-called intrinsic diversity profiles: if a community is intrinsically less diverse than another, its diversity profile is everywhere below the profile of the other community. On the other hand, if the two profiles intersect one or more times, no intrinsic ordering of the two communities is possible.

But once again an intrinsic diversity profile is a multivariate function of N , so once again the statistical problem lies in estimating the diversity profiles as functions of the estimated abundance. When the abundance estimates are obtained by means of independent replications of an environmental scheme, Fattorini and Marcheselli (1999) derive the asymptotic properties of the sample diversity profiles, considering conditions under which the sample profile is consistent and asymptotically normal. Joint confidence bands are also constructed around sample diversity profiles, and an asymptotically conservative procedure for comparing couples of diversity profiles by means of some methods previously adopted to make inference on Lorenz curves are considered. The procedure is able to distinguish between the three possible outcomes of profile comparison: dominance of a profile over the other, equivalence of the two profiles and profile crossing.

Applications to some case studies are considered while the possibilities of improvements and further developments are discussed.

BIBLIOGRAPHY

1. Barabesi L and Fattorini L (1998) The use of replicated plot, line and point sampling for estimating species abundances and ecological diversity, *Environmental and Ecological Statistics*, 5, 353-370.
2. Fattorini L. and Marcheselli M. (1999) Inference on intrinsic diversity profiles of biological populations, *Environmetrics*, 10, 589-599.
3. Patil G.P. and Taille C. (1982) Diversity as a concept and its measurement, *Journal of the American Statistical Association*, 77, 548-567.
4. Rao C.R. (1982) Diversity and dissimilarity coefficients: a unified approach, *Theoretical Population Biology*, 21, 24-43.

Filas-Przybył Sylwia, Kaźmierczak Maciej, Stachowiak Dorota (Urząd Statystyczny w Poznaniu)

W wykorzystanie danych ze źródeł administracyjnych w statystyce miast

Rosnące zapotrzebowanie na szczegółowe a zarazem kompleksowe informacje o mieście - z jednej strony- i kłopoty z jego zaspokojeniem z drugiej strony powodują, że statystyka publiczna poszukuje źródeł ich pozyskania. Badanie zapotrzebowania informacyjnego wskazuje również na konieczność badania miasta postrzeganego w oderwaniu od granic administracyjnych (obszar funkcjonalny, szersza strefa miejska, strefy wewnątrzmięskie). Konieczność innego spojrzenia na „miasto” wynika także z dokumentów strategicznych takich jak Krajowa Strategia Rozwoju Regionalnego, Koncepcja Przestrzennego Zagospodarowania Kraju oraz z działań Komisji Europejskiej na rzecz wspierania rozwoju miast. Te przesłanki to główne powody inspirujące do zwiększenia wykorzystania danych zawartych w źródłach administracyjnych w uzupełnieniu badań statystycznych. Wyniki prowadzonych badań oraz prac metodologicznych pokazują jakie możliwości daje dostęp do danych punktowych w połączeniu z narzędziami GIS. W referacie przedstawione zostaną wyniki pracy metodologicznej dającej podstawy do delimitacji stref wewnątrzmięskich o specyficznych cechach i znaczeniu. Prace mają na celu umożliwienie prowadzenia analiz zróżnicowania wewnątrzmięskiego na potrzeby polityki miejskiej.

Słowa kluczowe: statystyka miejska, źródła administracyjne, GIS

The use of data from administrative sources in urban statistics

Faced with a growing demand for detailed and complex information about cities on the one hand, and problems involved in meeting this demand on the other, public statistics is looking for new data sources. Studies into information needs indicate that urban statistics should focus on cities perceived quite apart from their administrative borders (i.e. functional areas, grey urban zones, inner-city zones). The need for adopting another research perspective with regard to the notion of the „city” can also be seen in strategic documents, such as the National Strategy for Regional Development, Poland’s Spatial Management Policy as well as the European Commission’s actions taken to support city development. All of these facts call for a more reliance on data from administrative sources to supplement statistical surveys. Results of research and methodological studies reveal possibilities offered by using point data combined with GIS tools. The paper presents results of a methodological study, which provides the basis for delimiting inner-city zones with specific characteristics and status. The aim of the study is to enable analyses of inner-city differentiation for purposes of urban policy-making.

Key words: urban statistics, administrative sources, GIS

Franceschi Sara (University of Tuscia), Fattorini Lorenzo (Università di Siena), Maffei Daniela (Università di Firenze)

Design-based treatment of unit nonresponse by the calibration approach

Unit nonresponse is often a problem in survey sampling, arising when the values of the interest variable cannot be recorded for some sampled units. Widely applied methods to account for unit nonresponse are the so-called nonresponse propensity weighting methods. Nevertheless these procedures require that every population unit had its own invariably positive response probabilities, while this requirement is often unrealistic because most populations contain units that do not respond under any circumstances. Moreover there are situations, e.g. forest inventories, in which the response pattern is fixed, in the sense that the population is strictly partitioned into respondent and nonrespondent units and responses are fixed characteristics of the units, just like the values of the interest variable.

When nonresponse is a fixed characteristic imputation or nonresponse calibration weighting should be used. Nevertheless it should be pointed out that, since knowledge about response behaviour is usually limited, it is difficult to defend any proposed method/model of imputation as being more realistic than others. As an alternative to imputation, the nonresponse calibration weighting (NCW) may be adopted.

In this work the use of NCW is considered in a complete design-based framework. Approximate expressions of design-based bias and variance of the calibration estimator are derived, design-

based consistency is investigated and some estimators of the sampling variance are proposed. Furthermore an extensive simulation study was performed. The results demonstrate how the reliability of the procedure is mainly determined by the capability of selecting auxiliary variables in such a way that their relationship with the interest variable is similar for both the respondent and nonrespondent sub-populations.

Frątczak Ewa (Szkoła Główna Handlowa w Warszawie)

Tradycja i nowoczesność w modelowaniu zjawisk i procesów demograficznych. Od tablic J. Graunta i siatki Lexisa po modele analiz wzdlużnych z efektem stałym i efektem losowym.

Metody stosowane w demografii ewoluowały w kontekście rozwoju demografii jako nauki, w interakcji z innymi dziedzinami nauki oraz w związku z nowymi możliwościami, które zostały przedstawione w postaci nowych źródeł danych i nowych technologii komputerowych. Nowe technologie i nowe paradygmaty w naukach społecznych, w tym w demografii, warunkują rozwój nauki, ale wymuszają postęp zarówno w obszarze badań jak i metod analizy. Ale,... nowe badania i nowe metody analizy mają swoje korzenie w szeroko pojętej tradycji.

Dla przykładu: tablice trwania życia, podstawowe narzędzie tradycyjnej analizy demograficznej, pierwsze wprowadzone przez J.Graunt'a w 1662 roku, są z powodzeniem stosowane współcześnie. Wraz z upływem czasu idea tablic umieralności jako modelu matematycznego została zastosowana do innych zdarzeń demograficznych jak: płodność, małżeństwo, rozwody oraz do zdarzeń niedemograficznych jak zdarzenia z rynku pracy, zdarzenia z procesu edukacji, przeżywalności urzędów czy przedsiębiorstw. O tych ostatnich zastosowaniach coraz więcej mówi się w Polsce, ale nie w środowisku demografów.

Przykład tablic trwania życia jest dobrym przykładem do przedstawienia ciągłości metod analizy demograficznej od tradycji po nowoczesność. Ale tablice umieralności to nie jedyna metoda analiz demograficznych. Artykuł przedstawia w sposób możliwie wszechstronny rozwój metod analizy demograficznej, poczynając od pierwszych tablic trwania życia J.Graunt'a po nowoczesne metody XXI wieku, w dobie współcześnie dostępnej technologii, wskazując na związek pomiędzy tradycją a nowoczesnością.

Tradition and modernity in the modeling of demographic events and processes. From J. Graunt's life tables, Lexis Diagram to longitudinal models with fixed and random effects

The methods used in demography evolved in the context of demography as a science, in interaction with other fields of science and in connection with the new possibilities that are presented in the form: of new data sources and new computer technology. New technologies and new paradigms in the social sciences, including demography, determine the development of science, but forcing progress both in research and analysis methods.

But, A new research and new methods of analysis have their roots in the broad sense of tradition. For example, life expectancy tables, the basic tool of traditional demographic analysis, first introduced by J. Graunt 'in 1662, are successfully used today. With time the idea of mortality tables as a mathematical model has been applied to other demographic events such as fertility, nuptiality, divorce and to non-demographic events as events of the labor market, the events of the educational process, equipment or enterprise (business) survival. The latter uses more and more say in Poland, but not among demographers.

An example of the life table is a good example to provide continuity of demographic analysis methods from tradition to modernity. But the mortality tables is not the only method of demographic analysis. The paper presents a comprehensive manner possible methods of demographic analysis development, from the first life tables J. Graunt 'and the modern methods of the twenty-first century, in the era of available modern technology, indicating the relationship between tradition and modernity.

Frątczak Ewa, Korczyński Adam (Szkoła Główna Handlowa w Warszawie)

Tradycja i współczesność w metodach imputacji danych. Przegląd teorii – wybrane przykłady zastosowań

Braki danych są nieuniknionym problemem dla badaczy wielu zjawisk i procesów będących przedmiotem analiz i modelowania. Stają się one szczególnie uciążliwe w przypadku badań panelowych oraz badań wzdlużnych. Występowanie braków danych w tego typu badaniach może uniemożliwiać wykorzystanie gromadzonych zbiorów do modelowania zmian w czasie a co za tym idzie braku modelowania mechanizmu przyczynowości (ang. *causality mechanism*). Ze względu na wzrost zainteresowania badaniami wzdlużnymi, poważnym wyzwaniem postawionym statystykom staje się problem braków danych.

Zagadnienie analizy niekompletnych zbiorów danych nie jest nowością w analizach statystycznych, niemniej w ciągu ostatnich 20 lat nastąpił znaczący rozwój dostępnych w tym obszarze metod.

W przypadku analiz przekrojowych i całkowicie losowych braków danych zastosowanie metody tradycyjnej tj. pomijania obserwacji z brakami danych pozwala zachować pożądane własności estymatorów parametrów szacowanych modeli. Stosowanie tej metody prowadzi jednak do zwiększenia błędów standardowych ocen parametrów. W przypadku analiz wzdlużnych jej zastosowanie prowadzi do utraty dużej liczby informacji, które potencjalnie mogłyby zostać wykorzystane w analizie. Podejście tradycyjne coraz częściej zastępowane jest skomplikowanymi metodami imputacji, bazującymi na modelach ekonometrycznych, które pozwalają uniknąć pomijania części zgromadzonych danych.

W grupie metod imputacji wyróżniamy metody deterministyczne oraz metody stochastyczne. Najbardziej podstawową z deterministycznych metod imputacji pozostaje zastąpienie braków

pewną stałą wartością, za którą najczęściej przyjmuje się średnią lub medianę. Stosowanie tego typu metod pozwala zachować spójność analizy niemniej niesie ze sobą zagrożenie zniekształcenia rozkładu zmiennych, które prowadzić może do zniekształcenia obrazu badanego zjawiska lub procesu.

Znacznie efektywniejsze okazują się być metody imputacji oparte na modelach ekonometrycznych, które pozwalają uzupełniać braki danych uwzględniając mechanizm według, którego - zgodnie z przyjętym założeniem - braki te powstają.

W grupie zaawansowanych metod imputacji danych wyróżnia się metody wykorzystujące algorytm EM (ang. *Expectation Maximization*) oraz metody wielokrotnych uzupełnień. Zastosowanie tego typu metod wymaga przyjęcia założenia, iż braki danych powstają w sposób losowy. Rozumiemy przez to, że braki te nie są zależne od wartości zmiennej dla której występują, zależą zaś od wartości pozostałych zmiennych.

Algorytm EM jest techniką optymalizacji funkcji wiarygodności otrzymanej po imputacji danych za pomocą oszacowanego modelu regresji liniowej. Algorytm ten cechuje się prostotą i numeryczną stabilnością. Wadą tego algorytmu jest to, iż nie zawsze zbiega on do maksimum globalnego.

Coraz szersze zastosowanie znajdują dziś metody wielokrotnych uzupełnień należące do grupy metod stochastycznych. Wśród najczęściej stosowanych w tym przypadku metod wyróżniamy algorytm Lavioriego, Dawsona i Shera (LDS) oraz algorytm Schafera. Algorytm LDS bazuje na modelu regresji logistycznej, za pomocą którego ocenia się prawdopodobieństwa wystąpienia braków danych. Algorytm Schafera oparty jest na modelu regresji liniowej z dodatkowym czynnikiem losowym, za pomocą którego generuje się dane, imputowane następnie do analizowanego zbioru. Jak się okazuje metody wielokrotnych uzupełnień pozwalają otrzymać nieobciążone estymatory parametrów szacowanych modeli nawet w przypadku znaczącej liczby braków danych (przekraczających 25% zbioru wejściowego).

Słowa kluczowe: zbiory niekompletne, mechanizmy braków danych, imputacja danych, deterministyczne i stochastyczne metody imputacji

Classical and modern approaches in data imputation methods. Review of theories – examples of practical applications.

The problem of missing data is unavoidable in many examples of qualitative researches focused on analyzing and modelling of the phenomena and processes. The incompleteness of the data set becomes problematic especially in panel or longitudinal research. In this type of research the missing date can hinder modelling of the causality mechanism. The considerable increase in application of the panel and longitudinal research and analysis, issues a challenge to modern statistics.

In recent 20 years the methods helping to handle the missing data problem have significantly developed. However in practice listwise deletion is still the most commonly used method. The traditional approach is acceptable in cross-sectional analysis assuming that the values are missing completely at random. The only drawback in this case is that it increases the standard errors of the parameter estimates. In case of the longitudinal analysis the use of the listwise

deletion method leads to loss of the large part of the data. That is why the traditional approach is increasingly replaced by more complicated methods based on econometric models helping to use all of the accessible data.

The imputation methods can be divided into deterministic and stochastic. The basic deterministic method consists in substituting the missing values for a constant, usually mean or median value. This method allows to use the whole potential of the gathered data but it can distort the distribution of the variables which can lead to wrong conclusions concerning the analyzed phenomenon or process.

The imputation methods based on the econometric models are much more efficient than those based on a deterministic algorithms. These methods allows us to fill in the missing values according to a predefine mechanism which produces these missing values.

The method based on the Expectation Maximization (EM) algorithm and the multiple imputation methods can be distinguished in the group of advanced imputation methods. These methods requires making the assumption that the variables are missing at random meaning that the probability of missing data on a variable can depend on other observed variables, but not on this variable itself .

The EM algorithm is a technique of the likelihood function optimisation. The optimized likelihood function is obtained after the imputation based on the linear regression model. The EM algorithm is simple and numerically stable. However it do not always reach the global maximum of the likelihood function.

The alternative stochastic method is based on multiple imputation. The two basic algorithms can be distinguished in this group of methods, namely the Lavori, Dawson and Shera (LDS) algorithm and Schafer algorithm. The LDS algorithm is based on the logistic regression model which calculate a predicted probability of being missing. The Schafer algorithm is based on linear regression model with random component. This model calculates the values to be imputed in the date set. The multiple imputation methods produce unbiased estimates of the parameters even in data sets with large number of the missing values (over 25

Key words: incomplete data sets, missing data mechanisms, imputation for missing data, deterministic and stochastic imputation methods

BIBLIOGRAPHY

1. Allison P. D., *Multiple Imputation for Missing Data – a Cautionary Tale*, “Sociological Methods & Research February”, 2000, Vol. 28 No. 3, p. 301–309.
2. Lavori P., Dawson R. and Shera D., *A Multiple Imputation Strategy for Clinical Trials with Truncation of Patient Data*, “Statistics in Medicine”, 1995, Vol. 14, p. 1913–1925.
3. McLachlan G. J., Krishnan T., *The EM Algorithm and Extentions*, Wiley, New York 1997.
4. Rubin D. B., *Inference and Missing Data*, “Biometrika”, 1976, Vol. 63 No. 3, p. 581–592.
5. Rubin D. B., *Multiple Imputation for Nonresponse in Surveys*, J. Wiley & Sons, New York 1987.

6. Schafer J. L., *Analysis of Incomplete Multivariate Data*, Chapman & Hall, London 1997.

Furmańczyk Konrad, Jaworski Stanisław (Warszawski Uniwersytet Medyczny, Szkoła Główna Gospodarstwa Wiejskiego w Warszawie)

Wykrywanie zmian poziomu serii niezależnych obserwacji

W pracy zaproponowano procedurę służącą do wykrycia zmiany średniego poziomu obserwacji w ciągu niezależnych obserwacji. Przy czym zmiana jest identyfikowana po przekroczeniu pewnej z góry zadanej wartości. Decyzja o tym czy zmiana miała miejsce oparta jest na regule minimaksu.

Słowa kluczowe: reguła minimaksu, jednomodalne rozkłady, rozkłady oparte na wielomianach

Change point detection in a sequence of independent observations

We present a new method of detection a change in a sequence of independent observations. The change is identified when some given threshold is exceeded. Our method is based on minimax decision rule.

Key words: change-point, minimax decision, unimodal distributions, polynomial based distributions

Gamrot Wojciech (Uniwersytet Ekonomiczny w Katowicach)

O empirycznych prawdopodobieństwach inkluzji

Wyznaczenie wartości estymatora Horvitz-Thompsona wartości globalnej w populacji skończonej i ustalonej w ramach podejścia randomizacyjnego wymaga znajomości prawdopodobieństw inkluzji pierwszego rzędu charakteryzujących schemat losowania, według którego przeprowadzone zostało losowanie próby. Do estymacji lub wyznaczania na drodze analitycznej jego wariancji niezbędne są również prawdopodobieństwa inkluzji rzędu drugiego. Tymczasem, wiele stosowanych w praktyce schematów losowania charakteryzuje się znaczną złożonością. W rezultacie, zjawisko eksplozji kombinatorycznej uniemożliwia dokładne wyznaczenie

tychże prawdopodobieństw inkluzji. Dzieje się tak często między innymi dla rozmaitych sekwencyjnych schematów losowania przestrzennego. W takiej sytuacji, w miejsce oryginalnej statystyki Horvitz-Thompsona postuluje się wykorzystanie tak zwanego empirycznego estymatora Horvitz-Thompsona, który konstruowany jest poprzez zastąpienie nieznanymi prawdopodobieństw inkluzji ich oszacowaniami uzyskanymi poprzez wielokrotne symulowanie losowania próby tym samym schematem losowania. Oszacowania te nazywane są często empirycznymi prawdopodobieństwami inkluzji. W tej sytuacji nasuwa się pytanie: jak duża powinna być liczba replikacji dla zapewnienia wymaganej dokładności oszacowania poszczególnych wag i samej statystyki Horvitz-Thompsona. W referacie podjęta zostanie próba odpowiedzi na tak postawione pytanie.

Słowa kluczowe: populacja skończona, wartość globalna, prawdopodobieństwa inkluzji

On empirical inclusion probabilities

Computation of the Horvitz-Thompson statistic for the finite population total requires the knowledge of first-order inclusion probabilities that characterize sampling scheme which was employed to draw a sample. To estimate its variance second order inclusion probabilities are also needed. Meanwhile many sampling schemes are characterized by high complexity. As a result, combinatorial explosion prevents the exact computation of inclusion probabilities. This is particularly true for sequential sampling schemes. In such a case it is reasonable to replace them with estimates obtained in a simulation study and called empirical inclusion probabilities. A question arises, how many sample replications should be drawn to ensure desired accuracy of population total estimates. In this paper an attempt is made to review known solutions to this problem and to improve upon them for a certain sampling scheme.

Key words: finite population, population total, inclusion probabilities

Ganczarek-Gamrot Alicja (Uniwersytet Ekonomiczny w Katowicach)

Wielowymiarowe modele FIGARCH w ocenie ryzyka na polskim rynku energii elektrycznej

W pracy przeprowadzono analizę porównawczą ryzyka estymowanego w oparciu o klasyczny model GARCH oraz modelu FIGARCH na przykładzie modelu warunkowej korelacji DCC. Analiza porównawcza została przeprowadzona na bazie portfeli zbudowanych w oparciu o 24 szeregi czasowe z Rynku Dnia Następnego (RDN) Towarowej Giełdy Energii (TGE) w okresie od 02.03.2009 do 27.02.2011 roku.

Słowa kluczowe: SARIMA, GARCH, VaR, portfel, rynek energii elektrycznej

Multivariate FIGARCH models in risk estimation on Polish electric energy market

Based on Dynamic Conditional Correlation model (DCC) comparative risk analysis was made. The risk estimated by classical GARCH model was compared to risk estimated by FIGARCH model. Analysis was made on the basis of portfolio, which was built from 24 contracts form Day Ahead Market of Polish Power Exchange in the period from 02.03.2009 to 27.02.2011.

Key words: SARIMA, GARCH, VaR, portfolio, electric energy market

Gatnar Eugeniusz (Narodowy Bank Polski)

Rola NBP w systemie statystyki publicznej w Polsce

W referacie zostanie omówiona rola Narodowego Banku Polskiego, który jest bankiem centralnym Polski oraz członkiem Europejskiego Systemu Banków Centralnych (ESBC) w systemie statystyki publicznej w Polsce.

NBP realizuje zadania statystyczne przede wszystkim publikując opracowania w zakresie statystyki monetarnej, statystyki bilansu płatniczego oraz prowadząc kwartalne rachunki finansowe w ramach systemu rachunków narodowych, a następnie przekazując te dane do GUS.

Słowa kluczowe: Statystyka publiczna, statystyka finansowa, bank centralny

National Bank of Poland in the system of official statistics in Poland

The aim of the paper is to discuss the role of the National Bank of Poland as Poland's central bank, and member of the European System of Central Banks (ESCB), in the system of official statistics in Poland.

The NBP publishes monthly, quarterly and annually reports in the field of monetary statistics, balance of payments, international investment position and keeps quarterly financial accounts of the system of National Accounts in Poland.

Key words: official statistics, financial statistics, central bank

Gerasymenko Sergiy (Wyższa Szkoła Bankowa w Poznaniu, Wydział Zamiejscowy w Chorzowie), Chupryna Olena (Narodowy Uniwersytet Karazina w Charkowie)

Optymalizacja spisu wskaźników dla oceny porównawczej obiektów socjalno-ekonomicznych

Podjęcie decyzji co do inwestowania potrzebuje informacji wyczerpującej, operatywnej i pewnej pro rankingi oddzielnych krajów, terytorium albo biznes-partnerów. Ustalenie rankingów potrzebuje uwzględnienia wielu charakterystyk. Ilościową miarą tych charakterystyk jest wskaźniki statystyczne. Jak wiadomo, dla ich uogólnienia mogą być wykorzystane: obliczenie wielowymiarowej średniej (w większości przypadków); ocena różnic między obiektami przy pomocy metod analizy wielowymiarowej, mianowicie - klasterowej analizy. Na odmianę od metody obliczenia wielowymiarowej średniej metody analizy wielowymiarowej oceniają realne różnice między poziomami wskaźników. To znaczy że wnioski według wyników analizy wielowymiarowej są więcej pewny.

W obu przypadkach formowanie spisu wskaźników zwykle odbędzie się przy pomocy metody ekspertowej bez następnego sprawdzania jego adekwatności. To powoduje albo wykorzystanie spisu ograniczonego który nie może scharakteryzować obiekty wszechstronne, albo do obecności w spisu wskaźników który dublują jeden drugiego. Ostatnie uwarunkuje rozrastanie się spisu co zwiększa wartość i czas badania bez zwiększenia go jakości. Więcej tego dublowanie przekreśli wielowymiarową ocenę obiektu przesadzając wkład do niej składnicy dla charakterystyki której wykorzystano kilka wskaźników związanych między sobą.

Udowodnieniem wyżej wymienionego są wyniki oceny porównawczej rozwoju 25 krajów UE (plus Ukraina). Właśnie klasterowa analiza wykorzystana do spisu z 15 wskaźników rozwoju socjalno-ekonomicznego w latach 2005-2009 wyznaczyła pewną tendencję podziału krajów za klasterami. W dalszym ciągu wskaźniki były podzielone o 2 części - socjalne i ekonomiczne. Klasterowa analiza zastosowana do każdego bloku wskaźników ujawniła dla niektórych krajów tendencje które bardzo różną się od wyników analizy za 15 wskaźnikami.

W celu optymalizacji spisu klasterowa analiza była zastosowana do spisu 15 wskaźników jako elementów całokształtu wskutek czego wskaźniki zostali podzielone o kilka grup. Parzysty współczynniki korelacji obliczone dla wskaźników każdego klastera pozwolili ujawnić wskaźnik jaki ma największą korelację co do innych. Ten wskaźnik był zostawiony w spisu, a inne wyłączone z niego. Optymalizowany w taki sposób spis wskaźników naszym zdaniem musi powiększyć jakość oceny porównawczej wielowymiarowych obiektów.

Jako miara charakterystyki zmiany ocen w tym przypadku może być wykorzystana odległość między klasterami: przy wykorzystaniu optymalizowanego spisu wskaźników odległość musi być mniejszą w związku z likwidacją zwiększonego wpływu oddzielnych charakterystyk. Jak wiadomo odległość Ewklidowa która jest wykorzystywana dla tego w klasterowej analizie nie uwzględnia skalę całokształtu. W związku z tym dla oceny porównawczej odległości między klasterami w całokształtach różnego rozmiaru była zaproponowana do wykorzystania odległość

Ewklidowa przerobiona za wzorem odchylenia standardowego:

$$\overline{C}_m = \left[\frac{\sum_{i=1}^m \sum_{j=1}^k (z_{ij} - z_{ik})^2}{m(k-1)} \right]^{\frac{1}{2}}$$

W naszym przykładzie codo 25 krajów UE (plus Ukraina) poziom takiego odchylenia dla optymizowanego spisu wskaźników okazał się mniejszym w porównaniu z go poziomem dla spisu z 15 wskaźników. Z tego można wywnioskować że po wyeliminowaniu ze spisu wskaźników związanych między sobą była otrzymana ocena która wskazuje na większe podobieństwo rozwoju socjalno-ekonomicznego krajów europejskich.

Słowa kluczowe: ocena porównawcza, obiekt wielowymiarowy, spis wskaźników, analiza klasterowa, odchylenie standardowe

Optimization of the set of indexes for the comparative estimation of the social-economic objects

Taking a decision about the investment needs a thorough, efficient and reliable information about the ratings of definite countries, regions and business partners. While defining the ratings it is necessary to take in account a great deal of characteristics. Statistic indexes act as the quantitative expression of the characteristics. As we know for their generalization we can use two approaches: the calculation of multidimensional mean (in most cases); the estimation of the differences between the objects with the help of the methods of multidimensional analysis, in particular – cluster analysis. In contrast to the methods of building the multidimensional mean, the methods of multidimensional analysis estimate the really existing differences in the levels of indexes. So, the conclusions according to the results of the multidimensional analysis are more reliable.

The results of the estimation of the 25 countries of European Union (+ Ukraine) are a proof of the above mentioned. Thus, the cluster analysis for the set of 15 indexes of the social-economics development for 2005-2009 defined the definite tendency of distribution of the countries according to the cluster. Then the indexes were divided into 2 groups – social and economic. Cluster analysis which was used for each group of indexes, revealed for some countries the tendencies which differ from the results of the analysis according to 15 indexes.

Cluster analysis which was used for the set of 15 indexes as the elements of the aggregate allotted them to the groups. Twin indexes of the correlation, calculated for each group, allowed to choose the index which has the highest correlation with all the others and to keep it in the set excluding all the others. Optimized in this way the set of indexes must raise the quality of the comparative estimation of the multidimensional objects.

The distance between the clusters can serve as the measure of the characteristics of the estimation change: while using the optimized set of indexes it must be lesser because of elimination of the exaggerated influence of the definite characteristics. As it is known, the used for this aim in cluster analysis Euclidian distance doesn't take into account the scale of the aggregate. That's why for the comparative estimation of the distances between the clusters in the aggregates of

different volume the following formula of Euclidian distance transformed upon the model of standard deviation was used.

For our example the value of such transformed Euclidian distance for the optimized set of indexes turned out to be lesser as compared with its value calculated for the full set of 15 indexes. Thus, excluding interrelated indexes from the set under consideration we got the estimation indicating the more semblance of the social-economic development of the countries of Europe.

Key words: comparative estimation, multidimensional object, set of indexes, cluster analysis, standard deviation

Ghosh Malay (University of Florida)

Finite Population Sampling: A Model-Design Synthesis

There are primarily two basic approaches towards inference from sample survey data. The first, the design-based approach, finds estimators of population quantities based on probability distributions generated by a given selection mechanism. In contrast, a model-based approach assumes the population units to be generated from some super- population, and the assumed superpopulation model governs any subsequent inference.

There has been a long-standing debate among survey statisticians regarding which one of the two is the preferred inferential approach. The advocates of design-based methods often criticize model-based inference regarding its failure to guard against any possible model misspecification. On the other hand, those advocating the use of models, question the ability of design-based methods to provide inference with sufficient accuracy in the face of small sample sizes, and often in the needed justification of approximations for small or moderate samples.

While the basic conceptual disagreement between the two approaches cannot be resolved, from an operational point of view, it is often possible to find an agreement between the two. The objective of this article is to produce some general results which provide this agreement, exact or asymptotic. We will illustrate our general results with several examples.

Głowiński Cezary, Konikiewicz Kamil (SAS Institute sp. z o.o.)

W wykorzystanie sieci społecznych w modelowaniu zachowań klientów firm telekomunikacyjnych

Praktycznie wszystkie duże firmy telekomunikacyjne wykorzystują statystykę i modele analityczne do wspomagania podejmowania decyzji w procesach biznesowych. Dotychczas dominowały modele, które bazowały na zmiennych opisujących historyczne zachowanie klientów. Takie podejście pozwalało budować całkiem skuteczne oceny punktowe (modele i scoringi). Przy tym podejściu modele koncentrowały się na kliencie jako pojedynczej jednostce i starały się wychwycić zmiany zachodzące w zachowaniu pojedynczego klienta w czasie. Z zależności od problemu był to okres od 3 do 12 ostatnich miesięcy i wymagało to odpowiedniego przygotowania zmiennych wejściowych tak aby model (drzewa decyzyjne, regresje, sieci neuronowe itd.) były w stanie zależności takie wykrywać i wiązać je z faktem przyszłej rezygnacji z telefonu bądź dokupienia kolejnej usługi.

W ostatnim czasie zaczęliśmy wzbogacać dotychczasowe podejście, a dokładnie listy zmiennych do modelowania, o wyniki analizy sieci społecznych (social network analysis). Z naszych doświadczeń wynika, że wyniki uzyskiwane dzięki połączeniu informacji behawioralnych z informacjami o otoczeniu społecznym klientów daje w wielu przypadkach lepsze wyniki. Takie podejście pozwala spojrzeć na klienta nie tylko jako jednostkę, ale również członka pewnej społeczności, która może mieć wpływ na jego zachowanie czy też decyzje. Natomiast z drugiej strony zaobserwowaliśmy, że nie wszystkie informacje o otoczeniu klientów okazywały się użyteczne, jak również istnieją jednostki, które można uznać, że są liderami swoich społeczności i takie które podążają za liderem.

W niniejszym referacie chcielibyśmy podsumować nasze doświadczenia z wykorzystania analizy sieci społecznych w połączeniu z klasycznym modelowaniem predykcyjnym na przykładzie modelowania zachowań klientów firm telekomunikacyjnych. W szczególności chcielibyśmy opisać, które rodzaje statystyk powiązań społecznych pozwalają uzyskać lepsze wyniki modelowania w stosunku do tradycyjnego podejścia behawioralnego, a także odpowiedzieć na pytanie jak dalekie powiązania pomiędzy klientami można czy też warto uznawać za istotne.

Słowa kluczowe: modele behawioralne, analiza sieci społecznych, telekomunikacja

Gołata Elżbieta (Uniwersytet Ekonomiczny w Poznaniu)

Spis ludności i prawda

W artykule podjęto dyskusję jakości informacji pozyskiwanych w najstarszych, największych oraz najbardziej znanych na świecie badaniach statystycznych jakimi są spisy ludności. Ocena jakości spisów ludności dotyczy warunków polskich. Celem artykułu jest przedstawienie propozycji zintegrowanego podejścia do spisów opartych na rejestrach administracyjnych i badaniach specjalnych w odniesieniu do warunków polskich. Pomimo długiej tradycji jak również dobrze wypracowanej metodologii badań spisowych, nie są one badaniami dostarczającymi 'idealnych' rezultatów. Przede wszystkim, spisy ludności jako badanie wyczerpujące obarczone mogą być błędami także o innym aniżeli losowy charakterze. Biorąc pod uwagę różnorodne metody przeprowadzania spisów, w tym metodę tradycyjną, wykorzystującą dane rejestrów administracyjnych i badań reprezentacyjnych, oraz metodę mieszaną, w rozważaniach uwzględniono różne źródła błędów.

W badaniu podjęto próbę identyfikacji spisywanej populacji w świetle standardów międzynarodowych: ludności rezydującej, stale i faktycznie zamieszkałej oraz wynikających z tego tytułu konsekwencji. Rozważając źródła błędów w spisach ludności, przedstawiono sposoby ich identyfikacji oraz eliminacji. Szczególną uwagę zwrócono na błędy pokrycia, które zilustrowano przykładami z polskich spisów odwołując się do wyników badań J. Paradysza (2002a, 2002b), M. Kędelskiego (1990), B. Sakson (2002). Uwzględniono wyniki dotychczasowych badań oceniających jakość rejestrów administracyjnych, które przeprowadzono w ramach przygotowań do spisu ludności 2011 r. (por. B. Pietrzak-Rynarzewska, T. Józefowski, 2010 oraz J. Paradysz, 2010).

Korzystanie z „pozastatystycznych” źródeł informacji w spisach ludności wiąże się z rachunkiem korzyści i zagrożeń. W badaniu uwzględniono przede wszystkim konsekwencje merytoryczne. Do często podnoszonych w tym zakresie problemów należą między innymi możliwości wnioskowania statystycznego w przypadku korzystania ze zintegrowanych źródeł danych (np. rejestr administracyjny i badanie specjalne), czy zagadnienia synchronizacji danych z rejestrów (określenie tzw. momentu krytycznego). Zebranie informacji o warunkach życia ludności w gospodarstwach domowych, to kolejny często podnoszony problem merytoryczny (por. Zhang Li-Chun, 2011). Przedstawiane w tym względzie propozycje definiowania gospodarstw domowych poprzez uwzględnienie ludności zamieszkującej w mieszkaniu, spotykają się z wieloma zastrzeżeniami. Przedstawione zostaną także konsekwencje dla miar bieżącej koniunktury demograficznej jak zniekształcenia współczynników demograficznych w skutek błędów szacunku. Innym problemem jest ocena sytuacji na rynku pracy czy charakterystyka społeczno-ekonomiczna ludności.

Ocena spisów ludności 'tradycyjnie' przeprowadzana jest w oparciu o wyniki badań popisowych. W Wielkiej Brytanii wyniki badania Census Coverage Survey są podstawą szacunku stopnia niedoszacowania ludności spisowej. Polskie doświadczenia w tym zakresie są niewielkie (por. J.

Kordos 2011). Na zakończenie rozważań uwzględniono także problem harmonizacji, a w szczególności niebezpieczeństwa niejednoznacznych wyników i rozbieżnych szacunków w warunkach korzystania z wielu źródeł. W tym względzie przedstawiono brytyjski projekt określany mianem One Number Census (ONC, por. np. Brown i in. 1999). Idea tego projektu sprowadzała się do takiego oszacowania liczby ludności, które byłoby spójne w przekroju różnych źródeł informacji, kategorii ludności i jednostek terytorialnych w ujęciu regionalnym i na poziomie całego kraju.

Population Census and truth

The article discuss quality of population census data as information obtained in the oldest, largest and most famous statistical surveys. Quality assessment of census data refers to Polish experiences. The article aims to propose an integrated approach to census based on administrative records and sample surveys. Despite long tradition as well as a well-developed research methodology, censuses do not provide 'perfect' results. First of all, the census as a comprehensive study may be burdened with the errors of another than random character. Taking into account all the different methods for conducting censuses, including the traditional method, using data from administrative records with sample surveys, and the mixed method, different sources of error are analyzed.

The study attempts to identify census population in the light of international standards: residents, people registered for permanent residence and actually living, as well as the resulting consequences. Considering sources of errors, methods of their identification and elimination are discussed. Particular attention is paid to coverage errors, which are illustrated by examples from Polish censuses referring to the analysis conducted by J. Paradysz (2002a, 2002b), M. Kędelski (1990), B. Sakson (2002). The results of previous studies assessing the quality of administrative records, which were conducted in preparation for Population Census 2011 are taken into account (see B. Pietrzak-Rynarzewska, T. Józefowski, 2010 and J. Paradysz, 2010).

Use of the 'non statistical' sources of information in the population census is associated with the analysis of advantages and disadvantages. The frequently raised issues in this respect include, inter alia, the possibility of statistical inference when using integrated data sources (e.g. administrative record and the sample survey), or synchronization of data from registers. Information about living conditions for households, is another often raised issue of substantive importance (cf. Zhang Li-Chun, 2011). The proposal to define households by taking into account the population living in an apartment face many restrictions. The consequences due to estimation errors will also be measured in respect to the current demographic situation and demographic coefficients. Another problem is the assessment of labor market situation and socio-economic characteristics of the population.

Evaluation of population census is 'traditionally' carried out based on the results of post-census survey. In the UK, results of Census Coverage Survey are the basis for estimating the degree of underenumeration in 2001 Census. Polish experiences in this field are rather small (cf. J. Kordos 2011). Discussion includes also harmonization problems, particularly the danger of divergent results and estimates, in terms of several data sources. In this regard, the British project known as One Number Census is presented (see Brown et al., 1999). The idea of the project was to

reduce the population underestimate and provide outputs which would be consistent in cross-section of territorial units in the regional and on a national levels.

BIBLIOGRAPHY

1. Brown J.J., Diamond I.D., Chambers R.L., Buckner, L. J., Teague, A. D., 1999, a methodological strategy for one-number census in the UK, „Journal of the Royal Statistical Association” No. 162 Part 2 s. 247 - 267
2. Józefowski T., Rynarzewska-Pietrzak B., 2010, Ocena możliwości wykorzystania rejestru PESEL w spisie ludności, [w:] Pomiar i informacja w gospodarce, Zeszyt Naukowy WIGE, E. Gołata (red.), Wydawnictwo UEP, Poznań
3. Kordos J., 2011, Methods Of Quality Assessment Of Population Census Data, „Studia Demograficzne”, Nr 1/2011
4. Paradysz J., 2002 a, Badanie małżeńskości i dzietności kobiet w narodowych spisach powszechnych [w:] Spisy ludności Rzeczypospolitej Polskiej 1921 - 2002. Wybór pism demografów, red. Z. Strzelecki, T. Toczyński, Polskie Towarzystwo Demograficzne, Główny Urząd Statystyczny, s.726-735
5. Paradysz J., 2002 b, O błędach nielosowych w badaniu dzietności kobiet w ramach Narodowego Spisu Powszechnego 1970. [w:] Spisy ludności Rzeczypospolitej Polskiej 1921 - 2002. Wybór pism demografów, red. Z. Strzelecki, T. Toczyński, Polskie Towarzystwo Demograficzne, Główny Urząd Statystyczny, s.479-482
6. Paradysz J., 2010, Konieczność estymacji pośredniej na użytek spisów powszechnych , [w:] Pomiar i informacja w gospodarce, Zeszyt Naukowy WIGE, E. Gołata (red.), Wydawnictwo UEP Poznań
7. Sakson B., 2002, Wpływ „niewidzialnych” migracji zagranicznych lat osiemdziesiątych na struktury demograficzne Polski, seria Monografie i Opracowania nr 481, Szkoła Główna Handlowa, Warszawa
8. Zhang Li-Chun, 2011, a Unit-Error Theory for Register-Based Household Statistics, „Journal of Official Statistics”, Vol.27, No.3, 2011 s. 415-432

Goujon Anne, Bauer Ramon, Samir K.C., Potančoková Michaela (Vienna Institute of Demography (VID) of the Austrian Academy of Sciences)

The human capital puzzle: Why is it so hard to find good data on educational attainment?

Data series on levels of educational attainment of the adult population consistent across time and space cannot be found as such, readily available, not in an aggregated form and not in

by age and sex. This is a pity because levels of educational attainment of the working age population are the main component of human capital that is used in many models, mostly related to economics, IT and health models. IIASA/VID have developed a methodology to reconstruct levels of educational attainment (see Lutz et al. 2007) based on the information contained in the best source for the most recent year. An improved and increased version will become available in 2012. We are showing in this paper that contrary to what is preconized by Cohen and Soto (2001), it does not really make sense to “keep the data as close as possible to those directly available from national censuses (in the spirit of the work of for OECD countries)” or other datasets since a large majority of those suffer from severe flaws, hampering any trend and regression analysis on levels of educational attainment. Besides showing the deficiencies of several human capital datasets and the advantage of our methodology (1st part), we show how picking the right dataset for the starting year can be a real hassle (2nd part). In a 3rd part, we orient the discussion towards the necessity to invest in harmonizing and mapping levels of educations to facilitate academic research for the benefit of societies.

Górecki Tomasz, Krzyśko Mirosław (Uniwersytet im. Adama Mickiewicza w Poznaniu)

Wersja jądrowa funkcjonalnych składowych głównych

W tej prezentacji proponowana jest nowa metoda konstrukcji funkcjonalnych składowych głównych, oparta o metodę składowych głównych dla danych wektorowych. Zakłada się, że dane funkcjonalne $\{x_i(t), i = 1, 2, \dots, N, t \in [0, T]\}$ są realizacjami procesu $X(t) \in L_2([0, T])$, gdzie $L_2([0, T])$ jest przestrzenią funkcji całkowalnych z kwadratem na przedziale $[0, T]$ wyposażoną w iloczyn skalarny $\langle u, v \rangle = \int u(s)v(s)ds$ oraz, że proces $X(t)$ ma następującą reprezentację: $X(t) = \mathbf{c}'\phi(t)$, gdzie $\phi(t) = (\phi_0(t), \phi_1(t), \dots, \phi_{N-1}(t))'$, $\mathbf{c} = (c_0, c_1, \dots, c_{N-1})'$, $0 < N < \infty$ oraz $E(\mathbf{c}) = \mathbf{0}$ i $Var(\mathbf{c}) = \Sigma$. Funkcje $\phi_0(t), \phi_1(t), \dots, \phi_{N-1}(t)$ są ortonormalnymi funkcjami bazowymi. Funkcja wagowa $u_k(t)$ wyznaczająca k -tą funkcjonalną składową główną $U_k = \langle u_k(t), X(t) \rangle$ ma postać $u_k(t) = \mathbf{u}'_k \phi(t)$, gdzie u_k jest wektorem własnym macierzy kowariancji Σ . Po zastąpieniu macierzy kowariancji Σ macierzą jądrową $\mathbf{K}$ otrzymujemy wersję jądrową funkcjonalnych składowych głównych.

Słowa kluczowe: analiza składowych głównych, funkcjonalna analiza składowych głównych, dane funkcjonalne, funkcja jądrowa, jądro reprodukujące

Kernel version of the functional principal components

In this presentation a new construction of functional principal components is proposed, based on principal components for vector data. It is assumed that functional data

$$\{x_i(t), i = 1, 2, \dots, N, t \in [0, T]\}$$

are the realization of the process $X(t) \in L_2([0, T])$, gdzie $L_2([0, T])$ is the space of square integrable functions on an interval $[0, T]$ equipped with the scalar product $\langle u, v \rangle = \int u(s)v(s)ds$ and that the process $X(t)$ can be represented as follows: $X(t) = \mathbf{c}'\phi(t)$, where $\phi(t) = (\phi_0(t), \phi_1(t), \dots, \phi_{N-1}(t))'$, $\mathbf{c} = (c_0, c_1, \dots, c_{N-1})$, $0 < N < \infty$ with $E(\mathbf{c}) = \mathbf{0}$ and $Var(\mathbf{c}) = \Sigma$. The functions $\phi_0(t), \phi_1(t), \dots, \phi_{N-1}(t)$ are orthonormal basis functions. The weight function $u_k(t)$ that determines the k th functional principal components $U_k = \langle u_k(t), X(t) \rangle$ has the form $u_k(t) = \mathbf{u}'_k \phi(t)$, where u_k is an eigenvector of the covariance matrix Σ . After replacing the covariance matrix Σ by the kernel matrix $\mathbf{K}$ we obtain a kernel version of the functional principal components.

Key words: PCA, FPCA, functional data, kernel function, reproducing kernel

Groenen Patrick J.F. (Erasmus University Rotterdam)

The support vector machine as a powerful tool for binary classification

Support vector machines (SVMs) have become one of the main stream methods for two-group classification, particularly so in computer science. Although it is a nonparametric method, it performs remarkably well in comparison with competing techniques such as linear discriminant analysis and logistic regression. The good performance of the SVM is due to several characteristics. First, it is robust against outlying observations. Second, it is regularized through a penalty parameter thereby avoiding overfitting of the data. Third, it can handle nonlinear predictions.

The formulation of the SVM in computer science is often done by the specification of a dual problem. In this presentation, we stick to a regression approach and show how the features of SVM can be explained from that angle. In addition, we show how the SVM loss function can be minimized through a majorization algorithm, SVM-Maj. A big advantage of majorization is that the SVM-Maj algorithm is guaranteed to decrease the loss in each iteration until the global minimum is reached. Nonlinearity was reached by replacing the predictor variables by either their monotone spline bases followed by a linear SVM or by applying kernels. An application is given through our R-package SVMMAj.

BIBLIOGRAPHY

1. Groenen, P.J.F., Nalbantov, G., Bioch, J.C. (2008). SVM-Maj: a majorization approach to linear support vector machines with different hinge errors. *Advances in Data Analysis and Classification*, 2, 17–43.

Gruchociak Hanna (Uniwersytet Ekonomiczny w Poznaniu)

Delimitacja lokalnych rynków pracy w Polsce

Stosowane w Europie i na świecie algorytmy służące do delimitacji lokalnych rynków pracy wymagają znajomości danych na temat dojazdów do pracy, w szczególności macierzy migracji związanej z zatrudnieniem. W związku z powyższym do niedawna przeprowadzenie delimitacji lokalnych rynków pracy w Polsce nie było możliwe ze względu na brak odpowiednich danych statystycznych. W 2010 roku Urząd Statystyczny w Poznaniu opublikował wyniki unikatowego Badania Przepływów Związanych z Zatrudnieniem za rok 2006, opartego na danych pozyskanych z zasobów rejestrów podatkowych Ministerstwa Finansów [por. Filas-Przybył i Stachowiak, 2010]. Do szacunków włączono po raz pierwszy informacje zgromadzone w rejestrach podatkowych (POLTAX). Udostępnienie danych przez Ministerstwo Finansów poprzedzone było kilkuletnimi pracami studialnymi w Urzędzie Statystycznym w Poznaniu. Są to pierwsze od 1988 roku dane dotyczące dojazdów do pracy. Ponadto, w roku 2011, wyniki badania dojazdów do pracy zostały uzupełnione poprzez udostępnienie macierzy migracji związanej z zatrudnieniem. Udostępnione dane umożliwiają przeprowadzenie pierwszej od wielu lat kompleksowej oraz niezależnej od subiektywnej opinii badacza delimitacji lokalnych rynków pracy w Polsce.

Delimitacja lokalnych rynków pracy dokonana zostanie przy pomocy algorytmu opracowanego dla Wielkiej Brytanii przez M. Coombes'a, A. Greena i S. Openshaw'a [por. Coombes, Green i Openshaw 1986], który rekomendowany jest przez Eurostat jako standardowe podejście do definiowania lokalnych rynków pracy w krajach europejskich [por. Eurostat 1992]. Oprócz Wielkiej Brytanii jest on stosowany w wielu krajach, między innymi Hiszpanii, Danii, Nowej Zelandii oraz Australii [por. Casado-Diaz 2000, Andersen 2002, Newell i Papps 2001, 2002, Bamber i Walter 2009]. Z wstępnych badań wynika, że może być on zastosowany również w Polsce, po dostosowaniu kilku parametrów zależnych od wielkości populacji oraz powierzchni jednostek wejściowych, które w badaniu zostały zdefiniowane jako gminy. Celem opracowania jest zweryfikowanie tej tezy.

Słowa kluczowe: delimitacja lokalnych rynków pracy, dojazdy do pracy, macierz przepływów związanych z zatrudnieniem

Delimitation of local labor markets in Poland

Used in Europe and worldwide algorithms for the delimitation of local labor markets require knowledge of data on commuter routes, in particular the migration matrix associated with employment. Therefore, the conduct of delimitation of local labor markets in Poland was not possible until recently due to lack of relevant statistical data. In 2010, the Statistical Office in Poznań published the results of a unique research concerning employment-related population flows for the year 2006, based on data obtained from tax records of the resources of the Ministry

of Finance [see Filas-Przybył and Stachowiak, 2010]. It was first time in Poland when the information collected in the tax registers (POLTAX) was included to statistical research. Sharing of data by the Finance Ministry was preceded several years of study work in the Statistical Office in Poznań. This has been the first study on the commuter routes since 1988. In addition, in 2011, the result of commuter routes was complemented by providing a matrix of migration associated with employment. Available data allow to perform a comprehensive and independent researcher on the subjective opinion of delimitation of local labor markets in Poland for first from many years.

Delimitation of local labor markets will be conducted using an algorithm developed for the UK, by M. Coombes'a, A. Green and S. Openshaw'a [see Coombes, Green, and Openshaw, 1986], which is recommended by Eurostat as a standard approach to defining local labor markets in European countries [see Eurostat 1992]. In addition to the UK it is used in many countries, including Spain, Denmark, New Zealand and Australia [see Casado-Diaz 2000, Andersen 2002, Newell and Papps 2001, 2002, Bamber and Walter 2009]. Preliminary research shows that it can be used also in Poland, after the adjustment of several parameters that depend on population size and area of input units, which in the study were defined as municipalities. The purpose of this paper is to verify this hypothesis.

Key words: delimitation of local labor markets, commuter routes, the matrix of employment-related flows

BIBLIOGRAPHY

1. Andersen, A.K. (2002) 'Are Commuting Areas Relevant for the Delimitation of Administrative Regions in Denmark?', *Regional Studies*, 36(8), 833–844.
2. Bamber, A. and R. Walter (2009). Implementing the Australian statistical geography standard (ASGS). In: Ostendorf B., Baldock, P., Bruce, D., Burdett, M. and P. Corcoran (eds.), *Proceedings of the Surveying & Spatial Sciences Institute Biennial International Conference, Adelaide 2009*, Surveying & Spatial Sciences Institute, pp. 605–616.
3. Casado-Diaz, J.M. (2000) 'Local Labour Market Areas in Spain', *Regional Studies*, 34(9), 843–856.
4. Coombes, M.G., Green, A.E. and Openshaw, S. (1986) 'An Efficient Algorithm to Generate Official Statistical Reporting Areas', *Journal of the Operational Research Society*, 37(10), 943–953.
5. *Dojazdy do pracy w Polsce. Terytorialna identyfikacja przepływów ludności związanych z zatrudnieniem*, 2010, red. K. Kruszka, GUS, Warszawa.
6. Eurostat (1992) *Study on Employment Zones*, Eurostat (E/LOC/20), Luxembourg
7. Filas-Przybył S., Stachowiak D., 2010, *Badanie przepływów ludności związanych z zatrudnieniem w: Procesy metropolizacyjne w teorii naukowej i praktyce konferencja naukowa* Łódź, 12-14 października 2009, GUS, Warszawa.

8. Newell, J. and Papps, K., 2001, Identifying Functional Labour Market Areas in New Zealand: A Reconnaissance Study using Travel-to-Work Data, New Zealand Department of Labour Occasional Paper Series, 2001/6.
9. Papps, K.L. and Newell, J.O. (2002) 'Identifying Functional Labour Market Areas in New Zealand: A Reconnaissance Study Using Travel-to-Work Data', IZA Discussion Paper 443.

Grzenda Wioletta (Szkola Główna Handlowa w Warszawie)

Bayesowski semiparametryczny model Coxa w badaniu determinant pozostawiania bez pracy osób młodych

Obecnie wśród osób rozpoczynających karierę zawodową obserwuje się szczególnie dużą wartość wskaźnika bezrobocia. Celem niniejszego opracowania jest identyfikacja czynników demograficznych oraz społeczno-ekonomicznych wpływających na długość czasu pozostawiania bez pracy tych osób. W badaniu wykorzystano bayesowski semiparametryczny model Coxa dla danych indywidualnych. Wykorzystanie modelu przeżycia daje możliwość analizy jednoczesnego wpływu wybranych zmiennych objaśniających na czas pozostawiania bez pracy. Natomiast podejście bayesowskie za pomocą rozkładów a priori umożliwia uwzględnienie w badaniu dodatkowej informacji spoza próby. Łącząc wiedzę z próby i spoza próby można uzyskać większą precyzję i wiarygodność otrzymanych oszacowań. Estymację modeli przeprowadzono z wykorzystaniem metod Monte Carlo opartych na łańcuchach Markowa, w szczególności próbnika Gibbsa.

Słowa kluczowe: bezrobocie, semiparametryczny model Coxa, wnioskowanie bayesowskie, metody MCMC

Bayesian semiparametric Cox model in the analysis of unemployment duration determinants of youth

High unemployment rates are observed among people beginning job careers nowadays. The aim of the work is to identify demographic and socio-economic factors influencing the unemployment duration in this age group. In this research, Bayesian semiparametric Cox model for individual data has been used. The advantage of survival model is the possibility of the analysis of the impact of selected independent variables on unemployment duration. The Bayesian approach with a priori distribution makes the use of out of the sample knowledge possible. By combining the information from the sample and out of the sample we can receive more precise and more credible estimations. The model has been estimated using Markov chain Monte Carlo method with Gibbs sampling.

Key words: unemployment; semiparametric Cox model; Bayesian inference; MCMC method

BIBLIOGRAPHY

1. Biggeri L., Bini M., Grilli L. (2001), The transition from university to work: a multilevel approach to the analysis of the time to obtain the first job, "Journal of the Royal Statistical Society", 164, pp. 293–305.
2. Blossfeld H.P., Hamerle A., Mayer K. (1989), Event history analysis, Statistical theory and application in the social sciences, Hillsdale, NJ: L. Erlbaum.
3. Blossfeld H.P., Rohwer G. (1995), Techniques of event history modeling, New approaches to causal analysis, Hillsdale, NJ: L. Erlbaum.
4. Collier W. (2003), The impact of demographic and individual heterogeneity on unemployment Duration: a regional study, "Studies in Economics", 0302.
5. Congdon P. (2007), Bayesian Statistical Modelling, Wiley & Sons, New York.
6. Ibrahim J.G., Chen M-H, Sinha D. (2001), Bayesian survival analysis, Springer-Verlag, New York.
7. Lancaster T. (1979) Econometric methods for the duration of unemployment, "Econometrica", 47, pp. 939–956.
8. Sinha D., Dey D.K. (1997) Semiparametric Bayesian Analysis of Survival Data, "Journal of the American Statistical Association", 92, pp. 1195–1212
9. Szkutnik W., Balcerowicz-Szkutnik M., Dyduch M, (2010), Wybrane modele i analizy rynku pracy: uwarunkowania rynku pracy i wzrostu gospodarczego, Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach, Katowice.

Hanusik Krystyna, Łangowska-Szczęśniak Urszula (Uniwersytet Opolski)

Uwarunkowania zróżnicowania poziomu i struktury konsumpcji gospodarstw domowych w Polsce

W rozważaniach na temat makroekonomicznych celów działalności gospodarczej ważne miejsce zajmuje dyskusja kategorii dobrobytu. Powiększanie dobrobytu indywidualnych jednostek, poprawa rozkładu dobrobytu w społeczeństwie traktowane są jako najważniejszy cel, który powinien być realizowany w procesie rozwoju społeczno- gospodarczego. Wychodząc z powyższych przesłanek przyjęto, iż badania dobrobytu są bardzo ważne poznawczo i mogą stanowić podstawę do oceny skuteczności procesu transformacji w naszym kraju.

Dobrobyt jest jednak kategorią bardzo złożoną, którą trudno zdefiniować i jeszcze trudniej zmierzyć. Na ogół przyjmuje się że dobrobyt indywidualnej jednostki jest sumą użyteczności

jaką czerpie ona z osiąganego całkowitego dochodu. Dobrobyt społeczny natomiast jest traktowany jako funkcja poziomu i rozkładu dobrobytu indywidualnych konsumentów. Ze względu na trudności z pomiarem użyteczności, zarówno z pozycji teorii użyteczności kardynalnej, jak i teorii użyteczności porządkowej, jako dobre, mierzalne odzwierciedlenie dobrobytu przyjmuje się szeroko rozumianą kategorię dochodu.

Pełny dochód jednostki konsumującej tworzy strumień usług pochodzących od całego jej indywidualnego majątku, który składa się z dochodu pieniężnego płacowego i pozapłacowego oraz dochodu niepieniężnego, obejmującego zadowolenie z pracy, strumień usług od majątku rzeczowego, wartość produkcji własnej i zadowolenie z korzystania z czasu wolnego. Z zaprezentowanej definicji wynika, że pełny dochód stanowi podstawę całkowitych możliwości konsumpcyjnych jednostki w określonym czasie.

Przedmiotem badań przedstawionych w niniejszym artykule jest dobrobyt gospodarstw domowych w Polsce osiągnięty w wyniku przemian ustrojowych, a celem ustalenie uwarunkowań i zróżnicowania dobrobytu tych podmiotów. Jako podstawę przeprowadzonych badań i analiz przyjęto następujące tezy:

- o dobrobycie gospodarstw domowych można pośrednio wnioskować na podstawie realizowanego poziomu i struktury konsumpcji,
- bieżące dochody gospodarstw domowych są podstawową determinantą ich możliwości konsumpcyjnych,
- typ społeczno-ekonomiczny gospodarstwa domowego, miejsce zamieszkania i wykształcenie głowy domu determinuje preferencje konsumpcyjne,
- modele ekonometryczne stanowią dobre narzędzie badania zależności konsumpcji od dochodów rozporządzalnych i tym samym pozwalają na identyfikację wybranych preferencji konsumpcyjnych w różnych klasach gospodarstw domowych.

Podstawowym źródłem wykorzystanych w badaniach informacji o dochodach i wyposażeniu gospodarstw domowych były dane z reprezentacyjnych badań budżetów gospodarstw domowych w Polsce organizowanych przez Główny Urząd Statystyczny. Obserwacje dotyczyły 2010 r.

Słowa kluczowe: konsumpcja, gospodarstwo domowe, dobrobyt

Determinants of level and structure diversification of households' consumption in Poland

The discussion about the economic welfare is important in the consideration of macroeconomic goals of economic activity. Increasing individual entities' welfare and improving its distribution in society are treated as the most important goals, which should be realized in process of socioeconomic development. Starting from these premises it was assumed that welfare studies are cognitively very important and may be a basis for evaluating the effectiveness of the transformation process in our country.

Welfare is a very complex category, which is very difficult to define and even harder to measure. It is generally accepted that the individual entity's welfare is a sum of the utility which it draws from the achieved total income. Social welfare is then treated as a function of the level and distribution of welfare of individual consumers. Due to the difficulties with measuring the utility, from the standpoint of utility's theories (both the cardinal and the ordinal), as a measurable presentation of welfare can be taken a widely defined category of income.

The full income of consumer creates a stream of services coming from the whole of its individual property, which consists of financial salary income, other financial income and non-financial income, such as job satisfaction, the stream of services from material assets, the value of domestic production and satisfaction from leisure time. From the presented definition it appears that the full income is the basis of the total consumption capacity of consumer in a given time.

The subject of presented in this article studies is welfare of households in Poland achieved through the political changes. The goal of the studies is to determine the conditions and differentiation of the welfare of considered entities. As a basis for the studies and analysis, the following theses have been taken:

- economic welfare of households can be indirectly evaluated by the realized level and the structure of consumption
- current incomes of households are the primary determinant of their consumption capacity,
- socio-economic type of household, place of living and education level of household's leader determine the consumer preferences,
- econometric models are a good tool to study the relation between consumption and disposable income and thus allow to identify selected consumer preferences in different types of households.

The primary source of information used in the studies were data from the representative households' budget surveys in Poland, organized by the Central Statistical Office. Observations concerned the year 2010.

Key words: consumption, household, economic welfare

Hozer Józef, Hozer-Koćmiel Marta (Uniwersytet Szczeciński)

Quantum satis dla firm w Polsce w 2012 roku

Liczba firm w gospodarce związana jest z liczbą gospodarstw domowych. Analizie tej ważnej proporcji przeznacza się zbyt mało uwagi w badaniach ekonomicznych. Sporo niejasności wiąże się z samą liczbą firm - inaczej określa ją GUS i inną liczbę podaje ZUS. Różnica jest tak

duża (ok. 2 mln), że wymaga to wyjaśnienia. Wyjaśnienia wymaga też to czy powinno się budować strategie wzrostu przedsiębiorczości czy wzrostu konsolidacji. W artykule zostaną również poruszone problemy gender - kulturowej tożsamości płci; podjęta zostanie m.in. próba odpowiedzi na pytanie czy ilość firm założonych przez kobiety jest na właściwym poziomie.

Słowa kluczowe: liczba firm, liczba gospodarstw domowych, przedsiębiorczość, gender

Quantum satis of Polish companies in 2012

The number of companies operating in an economy is connected with the number of households. Economic researchers do not pay due attention to the analysis of this crucial proportion. What is more, the very number of companies is not clearly defined - the data published by the Central Statistical Office and by the Social Insurance Institution differ significantly. The discrepancy is so vast (2 millions) that it should be clarified. It must also be elucidated what is more important to build - the strategies of entrepreneurship growth or of consolidation growth. The article discusses as well the gender issues (the gender-related cultural identity) - the authors will attempt to answer the question if the number of firms set up by women is satisfactory.

Key words: the number of companies, the number of households, entrepreneurship, gender

Jajuga Krzysztof (Uniwersytet Ekonomiczny we Wrocławiu)

Rozwój polskiej myśli statystycznej w naukach ekonomicznych

W referacie omówiony jest dorobek polskich statystyków w dziedzinie nauk ekonomicznych. Przedstawiony jest dorobek ponad 30 uczonych, poczynając od Fryderyka Moszyńskiego i Stanisława Staszica, kończąc na współcześnie tworzących statystykach. Przedstawione są zarówno osiągnięcia teoretyczne, jak i dokonania w zakresie organizacji sprawozdawczości statystycznej.

Jefmański Bartłomiej, Markowska Małgorzata (Uniwersytet Ekonomiczny we Wrocławiu)

Ocena dynamiki i kierunku zmian rozwoju inteligentnej specjalizacji regionów europejskich z zastosowaniem klasyfikacji rozmytej

Inteligentna specjalizacja regionów europejskich jest istotnym elementem ich inteligentnego rozwoju będącego jednym z trzech priorytetów strategii *Europa 2020 — Strategia na rzecz inteligentnego i zrównoważonego rozwoju sprzyjającego włączeniu społecznemu*. Przez pojęcie inteligentnej specjalizacji rozumie się m.in. współpracę przedsiębiorstw, ośrodków badawczych i szkół wyższych w celu zidentyfikowania najbardziej obiecujących obszarów specjalizacji w danym regionie. Niestety monitorowanie poziomu i kierunków zmian inteligentnej specjalizacji celem prowadzenia np. analiz benchmarkingowych utrudnia brak sprecyzowanych w literaturze przedmiotu parametrów identyfikujących inteligentną specjalizację regionu oraz dostępność danych na poziomie klasyfikacji NUTS2. Dlatego autorzy opracowania przyjęli założenie, że dobrym wskaźnikiem poziomu tego zjawiska może być serwicyzacja czyli ustalenie znaczenia sektora usług na tle innych sektorów w gospodarkach regionów. Ponadto należy mieć na uwadze fakt, że w przypadku analiz regionalnych określenie wyraźnych granic między regionami jest często utrudnione, co podważa sensowność stosowania dychotomicznej regionalizacji. Dlatego w referacie zaprezentowane zostaną wyniki analiz korzystających z dorobku teorii zbiorów rozmytych, a w szczególności z algorytmów rozmytych metod klasyfikacji. Ich zastosowanie umożliwi podział przestrzeni europejskiej na klasy regionów o określonych profilach inteligentnej specjalizacji. Dla każdego regionu zostanie oszacowany w poszczególnych latach przyjętego okresu badawczego 1997-2008 stopień przynależności do wyodrębnionych klas co pozwoli na analizę dynamiki i kierunków zmian ich usytuowania w podzielonej na klasy przestrzeni europejskiej.

Słowa kluczowe: klasyfikacja rozmyta, inteligentna specjalizacja, regiony europejskie

The evaluation of smart specialization dynamics and directions of changes development in European regions by applying fuzzy classification

Smart specialization in European regions presents the crucial component of intelligent development as one of three priorities in the strategy: *Europe 2020 – The strategy for smart and sustainable development facilitating social inclusion*. The concept of smart specialization is understood, among others, as the cooperation of enterprises, research centres and universities in order to identify the most promising areas of specialization in a given region. Unfortunately, monitoring smart specialization level and directions of changes, in order to carry out e.g. benchmarking analyses, turns out difficult due to the absence of parameters, in professional literature, which identify smart specialization in regions and also the availability of data at the level of NUTS 2 classification. Therefore, the authors of the study assumed that service orientation, i.e. defining the significance of service sector at the background of other sectors in regional eco-

nomies, may become a good indicator of this phenomenon level. Additionally, attention should be paid to the fact that in case of regional analyses specifying clear borders between regions is often a difficult task to perform, which undermines the reasons behind the application of dichotomous regionalization. Therefore, the paper presents analyses results taking advantage of fuzzy sets theories output and mainly of fuzzy classification methods algorithms. Their application will allow for the division of European space into regional classes of particular smart specialization profiles. For each region the level of membership in the distinguished classes will be estimated in particular years of the accepted research period, 1999-2010, which will facilitate the analysis of dynamics and directions of changes in their positioning in the European space divided into classes.

Key words: fuzzy classification, smart specialization, NUTS 2 regions

BIBLIOGRAPHY

1. Bezdek J.C., *Pattern Recognition with Fuzzy Objective Function Algorithms*, Plenum Press, New York 1981.
2. EUROPA 2020. *Strategia na rzecz inteligentnego i zrównoważonego rozwoju sprzyjającego włączeniu społecznemu*. Komisja Europejska, Komunikat Komisji, KOM(2010) 2020 wersja ostateczna, Bruksela 2010.
3. Höppner F., *Fuzzy cluster analysis: methods for classification, data analysis, and image recognition*, John Wiley & Sons, Chichester 1999.
4. *Polityka regionalna jako czynnik przyczyniający się do inteligentnego rozwoju w ramach strategii Europa 2020*, Komunikat Komisji do Parlamentu Europejskiego, Rady, Europejskiego Komitetu Ekonomiczno-Społecznego i Komitetu Regionów, KOM(2010) 553, Bruksela 2010.

Józefowski Tomasz (Urząd Statystyczny w Poznaniu)

Zastosowanie estymatora typu SPREE w szacowaniu liczby osób bezrobotnych w przekroju podregionów

Metody statystyki małych obszarów odgrywają istotną rolę w kształtowaniu nowoczesnych technik pozyskiwania informacji, które ukierunkowane są na obniżenie kosztów badań przy jednoczesnym zmniejszeniu obciążeń respondentów. Dzięki swoim własnościom umożliwiają uzyskiwanie wiarygodnych szacunków na niższych poziomach agregacji przestrzennej oraz bardziej szczegółowych domen, dla których klasyczne metody estymacji charakteryzują się zbyt dużą wariancją estymatorów. Mają one tą przewagę nad estymatorami znanymi z klasycznej metody reprezentacyjnej, iż umożliwiają dostarczenie potrzebnych informacji w sytuacji niewielkiej liczebności lub nawet braku obserwacji w próbie dla danego przekroju. Uzyskane w ten

sposób oszacowania dla niższych poziomów przestrzennych bądź subpopulacji różnią się często po zsumowaniu od szacunków uzyskanych za pomocą metody reprezentacyjnej dla wyższego poziomu, który jest możliwy ze względu na wystarczającą liczebność próby. Jednym ze sposobów poradzenia sobie z powyżej opisaną niezgodnością jest zastosowanie estymatora typu SPREE. Umożliwia on dostosowanie wartości w poszczególnych komórkach szacowanej tabeli kontyngencji do wartości brzegowych otrzymanych przy użyciu metody reprezentacyjnej. Komórki wewnętrzne tabeli początkowo wypełniane mogą być danymi z poprzednich spisów, bądź też z bieżących rejestrów administracyjnych. Metoda ta wydaje się być szczególnie atrakcyjna w kontekście estymacji parametrów charakteryzujących rynek pracy, gdyż techniki użyte w Badaniu Aktywności Ekonomicznej Ludności pozwalają na publikowanie danych jedynie na poziomie województwa. Odbiorcy danych statystycznych oczekują jednak informacji dla bardziej szczegółowych przekrojów geograficznych. W związku z powyższym głównym celem referatu będzie zaprezentowanie możliwości jakie oferuje estymator typu SPREE do oszacowania liczby osób bezrobotnych na poziomie podregionów województwa wielkopolskiego przy wykorzystaniu danych pochodzących z rejestru bezrobotnych oraz Badania Aktywności Ekonomicznej Ludności.

Słowa kluczowe: statystyka małych obszarów, rynek pracy, estymator SPREE

Using a SPREE estimator to estimate the number of unemployed across subregions.

The methodology of small area estimation (SAE) plays an important role in field of modern information provision, which aims to cut survey costs while lowering the respondent burden. Thanks to their properties, the SAE methods enable reliable estimation at lower levels of spatial aggregation and with more specific domains, where direct estimation techniques display too much estimator variance. Another advantage over estimators used in the representative approach is that small area estimation can be used to handle cases with few or no observations for a given domain in the sample. However, estimators of cell totals obtained for lower levels of spatial aggregation or subpopulations tend to differ from marginal totals calculated by means of the direct method for a higher level, which is possible considering the adequate sample size. One way of coping with this inconvenience is by applying a SPREE estimator. It is used to adjust values in the cells of an estimated contingency table to the totals obtained by means of the representative method. Internal cells can initially be filled with data from previous censuses, or current administrative registers. The method seems to be particularly useful for estimating parameters of the labour market, since the methodology used in the LFS can only yield data at the level of a province. Users of statistical data, however, expect information which is more geographically disaggregated. Considering the above, the aim of the present paper is to demonstrate the potential of the SPREE estimator for estimating the number of unemployed at the level of subregions of the Wielkopolska province using data about registered unemployment and LFS.

Key words: small area estimation, labour market, SPREE estimator

BIBLIOGRAPHY

1. Rao, J.N.K, (2003) Small area estimation, John Wiley & Sons, Hoboken, New Jersey.

2. EURAREA Consortium (2004), SPREE, GLSM and GLSMM: SAS Programs and Documentation
3. PURCELL, N.J AND KISH, L (1980). Postcensal estimates for local areas (domains). International Statistical Review 48, 3—18.

Jurečková Jana (Charles University in Prague)

Regression quantiles and their two-step modifications

Two-step α - regression quantile in the linear regression model $Y_i = \beta_0 + \mathbf{x}_i^T \beta + e_i$, $i = 1, 2, \dots, n$ first estimates the slope components β with a suitable (rank) R-estimate $\hat{\beta}_R(\alpha)$ and then determines the $[n\alpha]$ -quantile $\hat{e}_{[n\alpha]}$ of the residuals $\{e_i - \mathbf{x}_i^T \hat{\beta}_R(\alpha)\}$, $i = 1, 2, \dots, n$. It is asymptotically equivalent to the regression quantile of Koenker and Bassett, very close numerically even for small sample sizes, and the extreme regression quantiles of both types coincide exactly. We study their relations under finite samples and show under which conditions their values exactly coincide, and illustrate it numerically.

Kalinowski Marcin, Krzykowski Grzegorz (Wyższa Szkoła Bankowa w Gdańsku)

Dobór i selekcja danych jako kluczowy element badań statystycznych na przykładzie rynków finansowych

Celem artykułu jest analiza problemów pojawiających się podczas doboru i selekcji danych do badań statystycznych na przykładzie rynków finansowych.

Podstawowym problemem spotykanym w analizie statystycznej danych z rynków finansowych jest niezgodność ich rozkładów z rozkładem normalnym. Bezpośrednie zastosowanie standardowych miar lokalizacji i rozrzutu nie oddaje natury statystycznej prezentowanych wielkości. W artykule zaproponowano szczegółową procedurę transformacji i normalizacji klasycznych wskaźników finansowych umożliwiającą zastosowanie statystyk opisowych i wnioskowań statystycznych. Naturalnym zagadnieniem występującym podczas analiz tego rodzaju danych jest występowanie danych znacznie odbiegających od ich głównej lokalizacji. Zjawisko to, znane jako rozkłady o ciężkich ogonach, wymaga specyficznych procedur badawczych danych surowych. Dążą one do odzwierciedlenia miary wpływu klasycznych charakterystyk rynku finansowego na poziom rozwoju giełd akcji.

W pracy wprowadzono szereg zmiennych wskaźnikowych i przeprowadzono procedury kalibracji i normalizacji pozwalające na porównanie wpływu tych zmiennych na strukturę giełd akcji.

Słowa kluczowe: kalibracja, normalizacja, kwantyl, rozwój giełd

Selection and screening of data as a key element of statistical research on the financial markets

The aim of the article is the analysis of issues that arise while selecting and screening the data used in financial markets' statistical research.

The main problem appearing in the statistical analysis of the financial markets data is their divergence from normal distribution. Direct application of standard measures of location and dispersion does not reflect the statistical character of the analyzed data. The article contains propositions of a detailed analysis of the transformation and standardization of traditional financial ratios. This analysis enables us to apply descriptive and decisive statistical procedures.

A typical problem occurring during the analysis of financial markets is the appearance of the data significantly deviating from their main location. This statistical phenomenon known as heavy tails distributions requires specific research procedures of a raw data. It seeks to reflect an influence of typical financial characteristics on stock exchange level of development.

This article introduces a number of indicator variables and presents calibration and standardization procedures. This approach allows us to compare the impact of these variables on the structure of stock exchanges.

Key words: calibration, normalization, quantile, development of exchanges

Kamińska-Gawryluk Ewa (Urząd Statystyczny w Białymstoku)

Zróżnicowania regionalne w Polsce i wybranych krajach Unii Europejskiej

Mimo, że Unia Europejska należy do najbogatszych regionów świata, to w niej samej istnieją zasadnicze różnice w poziomie rozwoju i warunkach życia mieszkańców. Najzamożniejsze państwo - Luksemburg - jest ponad siedmiokrotnie bogatsze niż Rumunia czy Bułgaria, najbiedniejsi i najmłodszy członkowie UE. Dysproporcje pomiędzy regionami państw Wspólnoty są jeszcze bardziej widoczne.

Unia dostrzega problem tych różnicowań i podejmuje działania zmierzające do ich ograniczania poprzez realizację wspólnotowej polityki spójności. W latach 2007 – 2013 polityka ta nabrała szczególnego znaczenia, co uwidoczniło się w wielkości środków budżetowych przeznaczonych na jej realizację. Po raz pierwszy w dotychczasowej historii Unii fundusze na jej działania przekroczyły 1/3 budżetu, spychając na drugie miejsce Wspólną Politykę Rolną.

Wzmocnienie wymiaru terytorialnego oddziaływania Wspólnoty wynika z braku dostatecznych efektów dotychczasowych działań na rzecz poprawy stopnia konwergencji wewnętrznej i zewnętrznej Unii. Oceny efektów wsparcia Komisja Europejska dokonuje na podstawie przygotowywanych przez Eurostat cyklicznych raportów na temat spójności. Ostatni taki raport przyjęty przez Komisję w roku 2010 wskazuje na dalszą potrzebę wspierania i koordynacji działań zmierzających do harmonijnego rozwoju terytorium całej Wspólnoty.

Nadmierne dysproporcje rozwojowe utrudniają bowiem podejmowanie współpracy oraz prowadzą do napięć społecznych. Stąd w zasadzie wszystkie państwa członkowskie dążą do zmniejszania nierówności regionalnych. W największym zakresie dotyczy to krajów biedniejszych, do których należy zaliczyć również Polskę. Dlatego też celem strategicznym polityki regionalnej naszego kraju jest pobudzanie endogenicznych czynników wzrostu w układzie terytorialnym za pośrednictwem działań służących podnoszeniu konkurencyjności województw i całego kraju oraz konwergencji ¹.

Celem artykułu będzie prezentacja poziomu zróżnicowań regionalnych Polski na tle zróżnicowań występujących w wybranych krajach Unii Europejskiej. W porównaniu stopnia konwergencji analizowanych państw, wykorzystane zostaną regionalne dane statystyczne pochodzące ze źródeł statystyki publicznej Polski (GUS) oraz Unii Europejskiej (Eurostat). Analiza obejmować będzie zarówno dane dotyczące sfery gospodarczej, jak i społecznej.

Słowa kluczowe: poziom rozwoju, sfery zróżnicowań, konwergencja regionalna

Regional Differences in Poland and in Selected Countries of the European Union

Although the European Union belongs to the richest regions in the world, in the European Union itself there are significant differences in the level of development and living conditions of its inhabitants. The richest country – Luxembourg – is more than seven times richer than Romania or Bulgaria – the states that are the poorest and with the shortest time span of belonging to the EU. The disparities between the regions of the Union member states are even more visible.

The Union takes notice of the disparities and undertakes activities to reduce them by carrying out the European cohesion policy. In the years 2007 – 2013 this policy has taken on a greater meaning, which is reflected by the amount of budget funds for this aim. It was the first time in history of the Union that funds for the cohesion policy aims exceeded 1/3 of the budget, making the Common Agricultural Policy not the first but the second most important aim.

A lack of sufficient results of activities to improve the level of internal and external convergence of the Union results in putting emphasis on the territorial aspect of the Union's activities. The evaluation of the aid's results is done by the European Commission on the basis of recurrent Eurostat reports on cohesion. The latest report of this kind was accepted by the Commission in 2010 and shows further need to support and coordinate activities aiming to the cohesive development in the entire territory of the Union.

¹Ministerstwo Rozwoju Regionalnego, *Identyfikacja i delimitacja obszarów problemowych i strategicznej interwencji w Polsce*, Warszawa 2009, s. 18.

Excessive developmental disparities hinder undertaking the cooperation and lead to social tensions, therefore almost all member states aim to reduce regional disparities. This situation is mostly visible among poorer countries, including Poland. Bearing all this in mind, the strategic aim of the regional policy of our country is to stimulate endogenous development factors in the territorial aspect by undertaking activities leading to the increase in both the level of the country and its voivodships' competitiveness and the convergence².

This article presents the level of regional disparities of Poland against the background of the disparities noticeable in selected countries of the European Union. To compare the level of convergence of the analysed states, regional statistical data coming from Polish public statistics (CSO – GUS) as well as the European Union's ones (Eurostat) are used. The analysis is based on data describing economic and social sphere.

BIBLIOGRAPHY

1. Eurostat, epp.eurostat.ec.europa.eu/portal/page/portal/eurostat/home
2. Główny Urząd Statystyczny, *Rocznik Statystyczny Województw za lata 2005-2010*, Warszawa
3. Kokocińska K., *Polityka regionalna w Polsce i w Unii Europejskiej*, Wydawnictwo Naukowe Uniwersytetu im. Adama Mickiewicza, Poznań 2010
4. Komisja Europejska, *Inwestowanie w przyszłość Europy. Piąty raport na temat spójności gospodarczej, społecznej i terytorialnej*, Wydawnictwo Eric von Breska, Komisja Europejska, Dyrekcja Generalna ds. Polityki Regionalnej, Luksemburg 2010
5. Ministerstwo Rozwoju Regionalnego, *Krajowa Strategia Rozwoju Regionalnego 2010 - 2020. Regiony – Miasta – Obszary wiejskie*, Warszawa 2010
6. Ministerstwo Rozwoju Regionalnego, *Raport o rozwoju i polityce regionalnej*, Warszawa 2007
7. Ministerstwo Rozwoju Regionalnego, *Identyfikacja i delimitacja obszarów problemowych i strategicznej interwencji w Polsce*, Warszawa 2009
8. Pietrzak I., *Polityka regionalna Unii Europejskiej i regiony w państwach członkowskich*, Wydawnictwo Naukowe PWN, Warszawa 2007

²Ministerstwo Rozwoju Regionalnego, *Identyfikacja i delimitacja obszarów problemowych i strategicznej interwencji w Polsce*, Warsaw 2009, p. 18.

Keilman Nico (University of Oslo)

Challenges for statistics on households and families

Households and families are of basic importance for characterizing the life of most persons, and they are important units of statistical analysis in many areas, for instance, in studies of expenditure and income distribution, of the supply and demand of dwellings, of demographic behaviour, of labour market participation, to name a few. At the same time, rapid transformations in living arrangements and the emergence of new household types have been noted in many industrialized countries in the recent past. Prominent trends were, for instance, later start of family life, increased cohabitation, larger numbers of one-parent families as a result of divorce, more reconstituted families, and increased proportions living alone, in younger ages in particular.

Given these rapid changes in household and family behaviour, it is important that statistics in the field map the developments accurately, and that statistics are comparable across countries. Important data sources for family and household statistics are population censuses and sample surveys. Also, a number of countries are increasingly giving attention to administrative registers as sources of household and family data. Household and family statistics derived from censuses may be compiled approximately every tenth year, while information from surveys and registers may be used to cover inter-census periods.

The Conference of European Statisticians (CES) at the UN Economic Commission of Europe (UN EC) in Geneva has been keen to monitor new household and family structures, and the consequences of these new trends for household and family statistics. In this presentation I will summarize the work done for the CES and the UNECE during the past decade. The following emerging types of households and families will be considered:

1. Reconstituted families
2. Same-sex couples
3. LAT couples
4. Living apart but within a network
5. Commuters between households
6. Consensual unions and de facto marital status
7. Private versus institutional households

Issues about definitions, measurement, and data collection will be discussed.

Kisielińska Joanna (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie)

Dokładna metoda bootstrapowa na przykładzie estymacji średniej i wariancji

Metoda bootstrapowa polega na wtórnym próbkowaniu pierwotnej próby losowej pobranej z populacji o nieznanym rozkładzie. W artykule pokazano, że ze względu na postęp w technice komputerowej wtórne próbkowanie nie jest aktualnie konieczne, jeśli rozmiar próby nie jest zbyt duży. Możliwe jest wówczas automatyczne wygenerowanie wszystkich możliwych prób wtórnych i obliczenie wszystkich realizacji wybranej statystyki. Uzyskany rozkład można wykorzystać do punktowej lub przedziałowej estymacji parametrów populacji, bądź testowania hipotez. Metodę zastosowano do szacowania średniej i wariancji. Porównanie uzyskanych rozkładów z rozkładami granicznymi potwierdziło poprawność dokładnej metody bootstrapowej. Pobieranie próby losowej może być interpretowane jako dyskretyzacja ciągłej zmiennej. Ze względu na postęp w technice komputerowej, można mieć nadzieję, że znaczenie zmiennych dyskretnych w statystyce będzie wzrastało.

Słowa kluczowe: metoda bootstrapowa, nieparametryczna estymacja średniej i wariancji

The exact bootstrap method shown on the example of the mean and variance estimation

The bootstrap method is based on resampling of the original random sample drawn from a population with an unknown distribution. In the article it was shown that because of the progress in computer technology resampling is actually unnecessary if the sample size is not too large. It is possible to automatically generate all possible resamples and calculate all realizations of the required statistic. The obtained distribution can be used in point or interval estimation of population parameters or in testing hypotheses. The method was used to estimate mean and variance. The comparison of the obtained distributions with the limit distributions confirmed the accuracy of the exact bootstrap method. Random sampling may be interpreted as discretization of a continuous variable. Because of the progress in computer technology we may hope that the role of discrete variables in statistics will increase.

Key words: bootstrap, nonparametric mean and variance estimation

Klimanek Tomasz (Uniwersytet Ekonomiczny w Poznaniu, Urząd Statystyczny w Poznaniu)

W wykorzystanie estymacji pośredniej uwzględniającej korelację przestrzenną w badaniach społecznych w Polsce

Artykuł przedstawia propozycję wykorzystania metod estymacji pośredniej (w tym także tej metody, która uwzględnia korelację przestrzenną) do oszacowania pewnych charakterystyk rynku pracy w populacji osób w wieku 15 lat i więcej w przekroju powiatów w województwie wielkopolskim w 2008 roku. Jest to bardziej szczegółowy poziom agregacji przestrzennej niż ten prezentowany w publikacjach Głównego Urzędu Statystycznego opartych na wynikach Badania Aktywności Ekonomicznej Ludności. Drugim celem jest porównanie miar precyzji estymatora bezpośredniego z precyzją estymatora typu EBLUP (*empirical best linear unbiased predictor*) oraz estymatora typu SEBLUP (uwzględniającego korelację przestrzenną).

Słowa kluczowe: statystyka małych obszarów, autokorelacja przestrzenna, bezrobocie, Badanie Aktywności Ekonomicznej Ludności

U sing indirect estimation with spatial autocorrelation in social surveys in Poland

The article presents one possible application of indirect estimation methods (including the method accounting for spatial correlation) to estimate some characteristics of labor market in the population of people aged 15 and over at the level of NUTS4 in Wielkopolska region in 2008. This is a more detailed spatial aggregation of data compared with that found in publications of the Central Statistical Office based on Labour Force Survey results. The second aim of the article is to compare the precision measures of the direct estimator with those of the EBLUP estimator (*empirical best linear unbiased predictor*) and the SEBLUP estimator (which takes into account spatial correlation).

Key words: small area statistics, spatial autocorrelation, unemployment, Labour Force Survey

Kobus Paweł (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie)

Podejście bayesowskie do estymacji rozkładów cech wybranych jednostek zbiorowości z wykorzystaniem danych zagregowanych

W pracy analizowane jest wykorzystanie danych zagregowanych do poprawienia jakości estymacji rozkładów obserwowanych cech dla wybranych jednostek zbiorowości.

W praktyce częstym przypadkiem jest dostępność stosunkowo długich szeregów czasowych dla danych zagregowanych, rzędu kilkudziesięciu obserwacji. Jednocześnie w przypadku wybranych jednostek populacji będących przedmiotem zainteresowania brak jest jakichkolwiek obserwacji lub w najlepszym razie dysponujemy bardzo krótkimi szeregami czasowymi rzędu kilku obserwacji.

Przyjmijmy, że dysponujemy n_i obserwacjami zmiennych X_i o rozkładzie F_i dla $i = 1, \dots, k$ i m obserwacjami zmiennej Z będącej średnią zmiennych $X_1, \dots, X_k$. Uwzględnienie informacji na temat zmiennej X_i zawartej w średniej Z , szczególnie w przypadku $m > n$, powinno poprawić jakość estymacji rozkładu zmiennej X_i . W pracy rozważane jest podejście bayesowskie do estymacji rozkładu zmiennej X_i .

Słowa kluczowe: wnioskowanie bayesowskie, dane zagregowane, estymacja, Monte Carlo

The exact bootstrap method shown on the example of the mean and variance estimation

Bayesian approach to the estimation of individual units characteristics distribution with the use of aggregated data.

The paper deals with the problem of using the aggregated data for the improvement of the distribution estimation quality, in case of individual units characteristics.

In practice often the case is the availability of relatively long time series for the aggregated data. However, in case of individual population units the data could be unavailable or in the best circumstances it would be only a few observations.

Assume that we have n_i observations of a variable X_i following the distribution F_i for $i = 1, \dots, k$, and m observations of variable Z , which is the average of the variables $X_1, \dots, X_k$. The use of information contained in the variable Z , especially in the case $m > n$, should improve the quality of the variable X_i distribution estimation. The Bayesian approach to the distribution estimation of the variable X_i is considered.

Key words: Bayesian inference, aggregated data, estimation, Monte Carlo

Kołodziejczak Barbara, Roszak Magdalena (Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu)

Rapid e-learning w biostatystyce

Codziennym i czasochłonnym obowiązkiem każdego nauczyciela akademickiego jest przygotowanie materiałów kursowych do wykładanego przedmiotu. Aby podnieść efektywność oraz jakość

kształcenia studentów coraz częściej korzystamy z materiałów elektronicznych, które uatrakcyjniamy i wzbogacają zagadnienia prezentowane w sposób tradycyjny. Prawie każdy student posiada dziś komputer z dostępem do Internetu, tak więc zastosowanie w procesie kształcenia materiałów multimedialnych jest naturalne i pożądane. Przygotowanie tego typu zasobów musi być w miarę proste, aby nauczyciel akademicki mógł je wykonać samodzielnie. Brak czasu oraz umiejętności obsługi specjalistycznych programów, to najczęstsza przyczyna rezygnacji z opracowania nowych, ciekawych pomocy dydaktycznych. Zatem dostępność i łatwość obsługi aplikacji wykorzystywanych do tworzenia materiałów edukacyjnych wydaje się tutaj kwestią kluczową.

Referat będzie zawierał przegląd ogólnodostępnych programów wspomagających tworzenie szybkich i efektywnych materiałów edukacyjnych w postaci multimedialnej, w tym prezentacji i filmów instruktażowych. Zaprezentowane zostaną przykładowe materiały wspomagające nauczanie statystyki na kierunku lekarskim. Zostały one opracowane z myślą o wdrożeniu ich w kursowe zajęcia z przedmiotu biostatystyka, która oferowana jest studentom uczelni medycznej. Przedstawione zostaną także opinie studentów z wstępnych badań dotyczących testowego zastosowania tych materiałów na zajęciach z przedmiotu. Badania te, stanowią przykład wdrożenia nowatorskich metod kształcenia w proces nauczania biostatystyki oraz wskazują kierunki dalszych działań w zakresie wspomagania nauczania statystyki na uczelni medycznej.

Słowa kluczowe: biostatystyka, rapid e-learning, e-learning, kształcenie medyczne

Rapid e-learning in biostatistics

Daily and time consuming duty of every university teacher is to prepare the course material for the lined subject. To improve the efficiency and quality of education students are increasingly using electronic resources, which improve the visual and enrich issues presented in the traditional manner. Almost every student now has a computer with internet access, so the use of multimedia education process is natural and desirable. Preparation of such resources must be fairly simple to make a university teacher do it them self. Lack of time and specialized skills programs, is the most common cause of resignation from the development of new and interesting teaching aids. Thus, the availability and ease of application used to create educational material here seems to be a key issue.

The paper will include a review of public programs to help create high-speed and effective educational materials in the form of multimedia, including demonstration and instructional videos. We will present examples of materials that support the teaching of statistics in the Faculty of Medicine. They have been developed to implement them in exchange classes in biostatistics, which is offered to medical school students. Presented are the opinions of students from pre-test study on the use of these materials in the class of the subject. These studies are an example of implementation of innovative methods of teaching and learning in biostatistics and suggest directions for further action in support of teaching statistics in medical school.

Key words: biostatistics, rapid e-learning, e-learning, medical education

BIBLIOGRAPHY

1. Kołodziejczak B. (2011), „Zastosowanie programu PowerPoint w e-learningu”, e mentor nr 5 (42), 51-55.
2. Lenkiewicz J.(2010), „Rapid e-learning przełamał barierę czasu i finansów oraz pomaga tworzyć w pełni interaktywne kursy”, <http://www.eid.edu.pl/>
3. „Rapid E-Learning Guides”, <http://kineo.com/rapid-elearning/rapid-e-learning-guides.html>.
4. Unneberg L., Helle L. E. (2009), „Rapid Advances”, Upfront Publishing, UK.

Kończak Grzegorz (Uniwersytet Ekonomiczny w Katowicach)

Wnioskowanie statystyczne dla danych w wielowymiarowych tablicach wielodzielczych

W referacie zostanie poruszona problematyka wnioskowania statystycznego dla danych zawartych w wielowymiarowych tablicach wielodzielczych. Klasyczne metody wymagają spełnienia założenia odnośnie minimalnych liczebności oczekiwanych w komórkach tablicy wielodzielczej. Dotyczy to zarówno testowania niezależności zmiennych jak i testowania jednorodności danych. Dla prób niespełniających tego kryterium możliwe jest przeprowadzenie testów dokładnych. Takie analizy są przeprowadzane dla tablic o niewielkich wymiarach, w szczególności dla tablic o 2 wierszach i 2 kolumnach. W przypadku wielowymiarowej tablicy wymóg dotyczący minimalnych liczebności oczekiwanych w komórkach tablicy prowadzi do konieczności dysponowania bardzo dużymi próbami. Problem jest jeszcze poważniejszy jeżeli dla zmiennych klasyfikujących wyróżniono dużą liczbę kategorii. W referacie przedstawione zostaną możliwości wnioskowania w oparciu o rozkłady dokładne i ich przybliżenia. Rozwiązania dokładne ze względu na złożoność obliczeń nie zawsze są możliwe do zastosowania w praktyce. Rozwiązania przybliżone można otrzymać na drodze symulacji komputerowych. Przedstawiona propozycja rozwiązania pozwala znacznie osłabić założenia dotyczące minimalnej liczebności próby.

Słowa kluczowe: tablica wielodzielcza, test wielowymiarowy, testy permutacyjne

Statistical inference for multidimensional contingency tables

The problem of statistical inference for data in multidimensional contingency tables is analyzed. The multidimensional contingency table is one which contains information on three or more categorical variables. The chi-square test for independence or homogeneity of proportions is used in the two-way case contingency tables. These tests can be generalized for use with multidimensional tables. The procedure for conducting the chi-square analysis is identical to that

employed for a two-dimensional table. The proposal of the method for testing multi-directional hypothesis will be presented. This method is based on the permutation test. The properties of this test are analyzed in the Monte Carlo study.

Key words: contingency table, multi-dimensional test, permutation test

BIBLIOGRAPHY

1. Agresti A.: (1996) An Introduction to Categorical Data Analysis. John Wiley & Sons, New York.
2. Good P. (2005) Permutation, Parametric and Bootstrap Tests of Hypotheses, Springer Science Business Media, Inc., New York.
3. Good P. (2006) Resampling Methods. A Practical Guide to Data Analysis. Birkhauser Boston.
4. Sheskin D.J. (2004) Handbook of Parametric and Nonparametric Statistical Procedures, Chapman & Hall/CRC, Boca Raton.

Kopczewska Katarzyna (Uniwersytet Warszawski)

Regiony w przestrzeni. Czy lokalizacja i dostępność mają znaczenie dla rozwoju gospodarczo-społecznego?

Dążenie do podnoszenia efektywności gospodarczo-społecznej, przez pryzmat polityki spójności, w tym terytorialnej, przełożyło się na nową terytorialną klasyfikację regionów. W środowisku naukowym trwa dyskusja na ile lokalizacja ma znaczenie dla osiągniętych przez region wyników. O ile lokalizacja w sensie absolutnym, tj. położenie geograficzne, wiązana jest z korzyściami m.in. Z turystyki, kopalin etc., o tyle lokalizacja w sensie relatywnym, tj. względem głównych miast ma znaczenie dla osiągniętych ogólnych wskaźników rozwojowych. Nakłada się na to kwestia dostępności komunikacyjnej w transporcie wielomodalnym, gdzie stawia się hipotezę, że regiony słabo skomunikowane i/lub zlokalizowane peryferyjnie są predysponowane do osiągania gorszych wyników niż regiony zlokalizowane korzystnie. W niniejszym badaniu na poziomie NTS3 dla Polski oraz UE wykorzystane zostaną dane z badania ESPON dotyczące dostępności komunikacyjnej oraz z EUROSTAT obejmujące wskaźniki rozwoju gospodarczo-społecznego. Przy wykorzystaniu tradycyjnych oraz przestrzennych metod ilościowych ocenie poddane zostanie znaczenie dostępności i lokalizacji absolutnej i relatywnej w budowaniu konkurencyjności regionalnej, przy uwzględnieniu struktury gospodarczo-społecznej badanych obszarów.

Słowa kluczowe: dostępność komunikacyjna, spójność terytorialna, lokalizacja relatywna

Regions in space. Is the location and availability are important for economic and social development?

Tendency to improve the efficiency of economic and social cohesion, with regard to territorial cohesion policy, was reflected in a new territorial classification of regions. The scientific community leads the discussion on the issue for how much a location is important for the results achieved by the region. While the location in an absolute sense, i.e., geographical location, is tied to the benefits from tourism, minerals, etc., so the location in a relative sense, i.e. in relation to the major cities is important for overall performance and development indicators. Additionally, the issue of accessibility with a multi-modal transportation is important. There exists the hypothesis that regions poorly communicated and / or located peripherally are predisposed to underperform compared with better localized regions. In this study, the NTS3 level data for Poland and EU regions from the ESPON study regarding accessibility and EUROSTAT indicators covering economic and social development will be used. Traditional and spatial quantitative methods will be used to test accessibility and an importance of absolute and relative location in the process of the regional competitiveness building, taking into account the economic and social structure of the investigated areas.

Key words: accessibility, territorial cohesion, relative location

BIBLIOGRAPHY

1. Beuthe, M.: Transport evaluation methods: from cost-benefit analysis to multicriteria analysis and the decision framework. In: Giorgi, L., Pearman, A. (eds) Project and Policy evaluation in Transport. Ashgate, Aldershot (2002)
2. Briggs, R.: The impact of interstate highway system on non-metropolitan growth, US Department of Transportation, Office of University Research, Washington D.C. (1980)
3. Cohen, J.P.: The broader effects of transportation infrastructure: Spatial econometrics and productivity approaches. *Transportation Research Part E* 46 317-326 (2010)
4. ESPON, Scenarios on the territorial future of Europe. ESPON Project 3.2 Final Report, Luxemburg (2007)
5. Gutierrez, J., Condeco-Melhorado, A., Martin, J.C.: Using accessibility indicators and GIS to assess spatial spillovers of transport infrastructure investment, *Journal of Transport Geography* 18 141-152 (2010)
6. Humphrey, C.R., Sell, R.R.: The impact of controlled access highways on population growth in Pennsylvania nonmetropolitan communities, 1940-1970, *Rural Sociology*, vol.42, 332-343 (1975)
7. Munell, A.H.: How does public infrastructure affect regional economic performance, *New England Economics Review* September-October: 11-32 (1990)
8. Rephann, T., Isserman, A.: New highways as economic development tools: An evaluation using quasi-experimental matching methods. *Regional Science and Urban Economics*, vol.24 723-751 (1994)

9. Schaltegger, Ch.A., Zemp, S.: Spatial Spillovers in Metropolitan Areas: Evidence from Swiss Communes. CREMA, Basel, Working Paper No.2003-06 (2003)
10. Tobler, W.: Speculations On The Geometry Of Geography. National Center For Geographic Information And Analysis, Technical Report 93-1 (1993)

Kordos Jan (Szkoła Główna Handlowa, Wyższa Szkoła Menedżerska w Warszawie)

Współzależność pomiędzy rozwojem teorii i praktyki badań reprezentacyjnych w Polsce

Praktyczne problemy napotkane w czasie projektowania i analizy badań reprezentacyjnych wpłynęły w znacznej mierze na rozwój teorii tych badań. Z drugiej strony, teoria badań reprezentacyjnych wpływała na praktykę tych badań, co prowadziło często do ich udoskonalenia (Rao, 2005).

Autor rozważa współzależności między teorią i praktyką badań reprezentacyjnych w Polsce w czasie ostatnich ponad 70 lat. Rozpoczyna od klasycznej pracy Neymana, (Neyman, 1934), która dała teoretyczne podstawy probabilistycznego podejścia wyboru próbki. Główne idee tej pracy były najpierw opublikowane po polsku w 1933 r. (Neyman, 1933) i miały istotny wpływ na praktykę badań reprezentacyjnych Polsce przed i po Drugiej Wojnie Światowej. Badania reprezentacyjne prowadzone w GUS w latach 1950-tych i 1960-tych były konsultowane z J. Neymanem w czasie jego wiz w Polsce w latach 1950 i 1958 (Fisz, 1950; Zasepa, 1958).

Praktyczne problemy występujące przy planowaniu i analizie badań reprezentacyjnych Polsce były częściowo rozwiązywane przez Komisję Matematyczną GUS powołaną w 1949 r. (Kordos, 1975). Komisja ta działała do 1992 r. i wpływała w istotny sposób na praktykę badań reprezentacyjnych Polsce.

Polscy statystycy rozwiązywali praktyczne problemy związane z badaniami reprezentacyjnymi w czasie (w badaniach budżetów gospodarstw domowych, w badaniach rolniczych) i wiele uwagi poświęcali metodzie rotacyjnej i badaniom panelowym (Greń, 1969; Kordos, 1967, 1982). Zajmowali się także optymalną lokalizacją próby dla kilku parametrów oraz problemami estymacji złożonej (Bracha, 1996; Greń, 1964, 1966, 1970; Kordos, 1982; Zasepa, 1962, 1972, 1993). Występowały znaczne problemy związane z jakością danych (Kordos, 1987, 1988). Podejmowali próby związane z imputacją i kalibrowaniem wyników, aby poprawić oceny i zapewnić zgodność między dodatkowymi informacjami dla danych brzegowych.

Poważne problemy występowały przy estymacji dla małych obszarów, które wymagają specjalnego omówienia (Domański i Pruska, 1996; Kalton, Kordos i Platek, 1993; Kordos, 1999, 2006; Paradysz, 1996). Intensywne prace badawcze w tej dziedzinie są dalej kontynuowane.

Ważny i trudny problem stanowią operaty losowania w badaniach ekonomicznych. Szczególnie dotyczy to małych przedsiębiorstw, dla których trudno uzyskać aktualny operat (Dehnel,

2003). Występują również wysokie odsetki braku odpowiedzi, co stanowi poważny problem przy uogólnianiu wyników.

Specjalną uwagę poświęcono ostatnim pracom polskich statystyków w zakresie teorii i praktyki badań reprezentacyjnych w Polsce, uwzględniając ogólny rozwój teorii tych badań.

Słowa kluczowe: badanie reprezentacyjne, jakość danych, metoda rotacyjna, statystyka małych obszarów

Interplay between sample survey theory and practice in Poland

A large part of sample survey theory has been directly motivated by practical problems encountered in the design and analysis of sample surveys. On the other hand, sample survey theory has influenced practice, often leading to significant improvements (Rao, 2005).

The author examines interplay between sample survey theory and practice in Poland over the past 70 years or so. He begins with the Neyman's (1934) classic landmark paper which laid theoretical foundations to the probability sampling (or design-based) approach to inference from survey samples. Main ideas of that paper were first published in Polish in 1933 (Neyman, 1933) and had a significant impact on sampling practice in Poland before and after the World War II. Sample surveys conducted in 1950s and 1960s were consulted with J. Neyman during his visits in Poland in 1950 and 1958 (Fisz, 1950; Zasepa, 1958).

Some practical problems encountered in the design and analysis of sample surveys were partly solved by the Mathematical Commission of the CSO which was established in 1949 (Kordos, 1975). The Commission had a significant impact on sampling practice in Poland and was active till 1992.

Polish sampling statisticians had problems with sampling in time (household budget surveys, agricultural sample surveys) and some research works were devoted to rotation methods and panel surveys (Greń, 1969; Kordos, 1967, 1982). They studied issues connected with size of sample determination, sample allocation and estimation methods (Bracha, 1996; Greń, 1964, 1966, 1969, 1970; Kordos, 1967, 1987; Zasepa, 1962, 1972, 1993) and with data quality (Kordos, 1987, 1988). There were problems with imputation and calibration estimation that ensures consistency with user specific totals of auxiliary variables, unequal probability sampling without replacement, analysis of survey data.

Very important and difficult problems for sampling practice are data for small areas. That is why Polish statisticians are interested in the issues related to estimation methods for small areas (Kalton, Kordos and Platek, 1993; Domański & Pruska, 2001; Paradysz, 1998). Activities in this field are continued (Dehnel, 2003).

Another important and difficult problem is obtaining adequate sampling frames in economic surveys. In that period many small-size enterprises came into existence for which there were no up-to-date registers which could be used as sampling frames (Dehnel, 2003). There are considerable rates of non-response which make generalisation of the results for the whole universe difficult.

Special attention is devoted to last development in sample survey theory and practice in Poland taking into account general development of sampling theory.

Key words: Sampling survey, rotation sampling, data quality, small area statistics

BIBLIOGRAPHY

1. Dehnel, G., *Statystyka małych obszarów jako narzędzie oceny rozwoju ekonomicznego regionów*. „Wydawnictwo Akademii Ekonomicznej w Poznaniu”, Poznań, 2003.
2. Domański, Cz., Pruska, K., *Metody statystyki małych obszarów*, „Wydawnictwo Uniwersytetu Łódzkiego”, Łódź, 2001.
3. Greń, J., *O pewnych metodach wyznaczania lokalizacji próby w losowaniu warstwowym wieloparametrycznym*, „Przegląd Statystyczny”, , 1964, z. 3, s. 361-369.
4. Greń, J., *O pewnym zastosowaniu programowania nieliniowego do metody reprezentacyjnej*, „Przegląd Statystyczny”, 1966, z. 3, s. 361-369.
5. Greń, J., *Efektywność powtarzalnych badań metodą reprezentacyjną*, „BSW”, 1969, t. 7, s. 97-104.
6. Greń, J., *Wielowymiarowy estymator regresyjny średniej dla skończonej populacji*, „Przegląd Statystyczny”, 1970, z. 1, s. 73-78.
7. Kordos, J., *25 lat działalności Komisji Matematycznej GUS*, „Przegląd Statystyczny”, 1975, z. 1, s. 171-173.
8. Kordos, J., *Metoda rotacyjna w badaniach reprezentacyjnych*, „Przegląd Statystyczny”, 1967, z. 4, s. 373-394.
9. Kordos, J., *Metoda rotacyjna w badaniach budżetów rodzinnych w Polsce*, „Wiadomości Statystyczne”, 1982, nr 9.
10. Kordos, J., *Dokładność danych w badaniach społecznych*, BWS, t. 35, 1987, GUS, Warszawa.
11. Kordos, J., *Jakość danych statystycznych*, „Polskie Wydawnictwa Ekonomiczne”, Warszawa, 1988.
12. Neyman, J., *Zarys teorii i praktyki badania struktury ludności metodą reprezentacyjną*, „Instytut Spraw Społecznych”, Warszawa, 1933.
13. Neyman, J., *On the Two Different Aspects of the Representative Method: The Method of Stratified Sampling and the Method of Purposive Selection*. “Journal of the Royal Statistical Society”, 1934, vol. 97, s. 558-625.
14. Paradysz, J., *Small Area Statistics in Poland - First Experiences and Application Possibilities*, Statistics in Transition, 1998, Vol. 3, Nr 5, s. 1003-1015.

15. Rao, J.N.K. , *Interplay Between Sample Survey Theory and Practice: An Appraisal*, "Survey Methodology", 2005, vol. 31, No. 2, pp.117–138.
16. Zasepa, R., *Problematyka badań reprezentacyjnych GUS w świetle konsultacji z prof. J. Neymanem*, „Wiadomości Statystyczne” 1958, nr 6, s. 7–12.
17. Zasepa, R., *Badania statystyczne metodą reprezentacyjną*, PWN, Warszawa, 1962.
18. Zasepa, R., *Metoda reprezentacyjna*, PWE, Warszawa, 1972.
19. Zasepa, R., *Use of Sampling Methods in Population Censuses in Poland*, „Statistics in Transition”, 1993, vol. 1, No. 1, s. 69–78.

Koronacki Jacek (Instytut Podstaw Informatyki PAN, Warszawa)

Analiza danych o dużym wymiarze w przypadku małej liczności próby

Miniony wiek bywa nazywany wiekiem informacji. Moce obliczeniowe komputerów oraz pojemności ich pamięci rosły w ostatnich dziesięcioleciach nieomal z dnia na dzień. Dziś nierzadko słyszymy, iż weszliśmy w wiek biologii a nawet szerzej, wiek nauk o życiu. Stało się tak między innymi dlatego, że niezwykle rozwinęły się biotechnologie pozwalające na zbieranie masowych danych o żywych komórkach.

Typowe dane w takich zastosowaniach biologicznych charakteryzują się małą liczbą obserwacji (obiektów) – rzędu dziesiątków lub setek – z których każda opisana jest tysiącami lub większą liczbą atrybutów (cech). Dobrymi i ważnymi przykładami problemów tego typu są dane mikromacierzowe, dotyczące ekspresji genów, lub dane proteomiczne. Zadania uczenia pod nadzorem w przypadku takich zbiorów danych nie są typowymi zadaniami statystycznej analizy wielowymiarowej („data miningu”) i wymagają specjalnego potraktowania. Ze względu na ogrom wymiaru danych (liczby atrybutów), pierwszoplanowego znaczenia nabiera (pod)zadanie wyboru z całego wektora atrybutów podzbioru atrybutów interesujących (ważnych) z punktu widzenia zadania uczenia (analizy regresji lub analizy dyskryminacyjnej). Praktyka podpowiada, że np. W przypadku zadania analizy dyskryminacyjnej właściwie zawsze o podziale na klasy decyduje tylko część – i to niewielka – wszystkich obserwowanych atrybutów. Zwykle większość, i to ogromna, atrybutów nie ma nic wspólnego z właściwym zadaniem klasyfikacji i w sposób niejako sztuczny „przenosi” to zadanie do przestrzeni o wielkim wymiarze.

W ramach wykładu przedstawiony zostanie stan badań w obszarze uczenia pod nadzorem w opisaney sytuacji. Jakkolwiek nacisk będzie położony na zadanie analizy dyskryminacyjnej, omówione zostaną także zadania z liczbową zmienną odpowiedzi. Zarysowane zostaną zagadnienia doboru modelu oraz regularyzacji metod. Specjalną uwagę poświęcimy zagadnieniu wyboru (selekcji) cech.

Literatura przedmiotu jest bardzo bogata. Niżej wymieniamy tylko jedną, za to doskonałą, monografię na temat statystycznych metod uczenia się oraz kilka niedawno temu opublikowanych prac z polskim współautorstwem (bibliografie zawarte w tych pracach dają szerszy pogląd na interesujący nas tu obszar badań).

Słowa kluczowe: uczenie pod nadzorem, problemy wielowymiarowe, selekcja cech, problemy biologii obliczeniowej

Statistical Challenges of High Dimensional Small Sample Size Problems

The past century, sometimes dubbed the “age of information,” saw incredible growth in computers’ processing speed and memory capacity. Now that computational power is helping us forge ahead into a new “age of biology”. Indeed, we have entered into what might be described as an era of the Life Sciences. Biotechnologies now enable massive amounts of data to be collected about living cells and organisms.

A major challenge in the analysis of such data sets is their size: a very small number of observations (records), of the order of tens or hundreds, versus thousands (or more) of attributes or features per each observation (such problems are now dubbed small- n -large- p -problems, where n stands for the sample size and p denotes the observation vector’s dimension). Typical examples include microarray gene expression experiments (where the features are gene expression levels) or data coming from genome-wide association studies (where the features are the so-called genetic markers, e.g., single nucleotide polymorphisms); in the former example, a typical dimension of the observation vector is of the order of thousands or tens of thousands, while in the latter it is of the order of hundreds of thousands.

In all these tasks supervised learning is quite different from that within a typical statistical inference or data mining framework when one is faced with small- p -large- n -problems. In particular, in the latter context and when, e.g., dealing with a problem of discriminant analysis, the main task is to propose a classifier of the highest possible quality of classification. In class prediction for typical gene expression data it is not the classifier per se that is crucial; rather, selection of informative (discriminative) features and the discovered interdependencies among them are to give the Life Scientists a much desired possibility of the interpretation of the classification results – since it is rather a rule than an exception that most features in the data are not informative, it is indeed of utmost interest to select the few ones that are informative and that may form the bases for class prediction. Equally interesting is a discovery of interdependencies between the informative features.

In the talk, an up-to-date, albeit brief, overview of the area of supervised learning for small- n -large- p -problems will be provided. While the emphasis will be on discriminant analysis, problems with quantitative responses will also be dealt with. Issues of model selection and regularization of methods will be presented. Special focus will be placed on the issue of feature selection.

The literature on the subject is enormous. In concluding, we refer the Reader to just one, however brilliant, monograph of general interest and a few recent papers with Polish co-authors (the bibliographies there provide a broad view of the area).

Key words: supervised learning, small n large p problems, feature selection, genomic and proteomic data

BIBLIOGRAPHY

1. Bogdan, M., Chakrabarti, A., Frommlet, F., Ghosh, J.K.: *Asymptotic Bayes-optimality under sparsity of some multiple testing procedures*, "Annals of Statistics", 2011, Vol. 39, No. 3, 1551–1579.
2. Bogdan, M., Frommlet, F., Biecek, P., Cheng, R., Ghosh, J.K., Doerge, R.: *Extending the Modified Bayesian Information Criterion (mBIC) to dense markers and multiple interval mapping*, "Biometrics", 2008, Vol. 64, No. 8, 1162–1169.
3. Damiński, M., Rada-Iglesias, A., Enroth, S., Wadelius, C., Koronacki, J., Komorowski, J.: *Monte Carlo Feature Selection for Supervised Classification*, "Bioinformatics", 2008, Vol. 24, No. 1, 110–117.
4. Damiński, M., Kierczak, M., Koronacki, J., Komorowski, J.: *Monte Carlo feature selection and interdependency discovery in supervised classification*, w: "Advances in Machine Learning", vol. II, J. Koronacki, Z.W. Ras, S.T. Wierchoń, J. Kacprzyk (red.), Springer, 2009, 371–385.
5. Frommlet, F., Ruhhaltinger, F., Twaróg, P., Bogdan, M.: *Modified versions of the Bayesian Information Criterion for genome-wide association studies*, "Computational Statist. and Data Anal.", 2012, w druku.
6. Hastie, T., Tibshirani, R., Friedman, J.: *Elements of Statistical Learning: Data Mining, Inference and Prediction*, Springer, 2009, wydanie II.
7. Żak-Szatkowska, M., Bogdan, M.: *Modified versions of the Bayesian Information Criterion for sparse Generalized Linear Models*, "Computational Statist. and Data Anal.", 2011, Vol. 55, No. 11, 2908–2924.

Kosiorowski Daniel (Uniwersytet Ekonomiczny w Krakowie)

Głęboka analiza rozrzutu w odpornej analizie strumienia danych ekonomicznych

Z praktycznego punktu widzenia priorytetowym celem analizy ekonomicznego szeregu czasowego jest uzyskanie wglądu w krótkookresowe właściwości probabilistyczne modelu generującego dane na podstawie stale uaktualnianej próby umiarkowanej długości. Na podstawie takiej w ogólności nieprecyzyjnej wiedzy dokonywanych jest szereg decyzji ekonomicznych oraz prognoz. W praktyce bardzo ważną kwestią jest odpowiedź na pytanie co mówi nam większość danych o przyszłym zachowaniu większości uczestników pewnego rynku. Szczególnie trudno

odpowiedzieć na takie pytanie w przypadku wielkich zbiorów danych generowanych przez zmieniający się wielowymiarowy model.

Powszechnie wiadomo, że ekonomiczne szeregi czasowe niemal zawsze zawierają obserwacje odstające. W związku z faktem, że prognozowanie każdego szeregu czasowego opiera się o ekstrapolację wzorca który ujawnił się w przeszłości – jeżeli estymatory parametrów modelu mocno zależą od kilku nietypowych obserwacji związanych z izolowanymi lub niepowtarzalnymi zdarzeniami, wówczas możemy oczekiwać słabej jakości prognoz. Jeżeli wspomniane parametry posiadają ekonomiczne interpretacje – obecność niewykrytych wpływowych obserwacji może prowadzić do błędnych decyzji ekonomicznych.

W statystyce jednowymiarowej znanych jest wiele odpornych procedur statystycznych wykorzystujących funkcję kwantylową. W statystyce wielowymiarowej wobec braku naturalnego porządku nie istnieje uniwersalna definicja funkcji kwantylowej. Statystyczne funkcje głębi stanowią szczególnie udane przykłady wielowymiarowych funkcji kwantylowych.

W artykule badamy własności głębi położenia – rozrzutu wprowadzonej przez Mizere i Muller (2004). W szczególności zastanawiamy się nad informacją dotyczącą modelu generującego szereg czasowy dostarczaną przez zastosowanie procedur indukowanych przez tę statystyczną funkcję głębi. Zwracamy uwagę na krótkookresowy opis szeregu czasowego wykonany za pomocą tej głębi – badamy odporność i ekonomiczną użyteczność takiego opisu. W szczególności naszą uwagę zwraca wykorzystanie ruchomej mediany Studenta (dwuwymiarowa mediana Tukey'a w zagadnieniu położenia-rozrzutu) i uogólnionej głębi pasma (funkcja głębi dla danych będących funkcjami) w badaniach strumieni danych typu ARMA, GARCH in mean i SV. Rozważania teoretyczne ilustrujemy przykładami empirycznymi dotyczącymi kryzysu finansowego.

Słowa kluczowe: strumień danych, odporność, statystyczna funkcja głębi

Location–scale depth in robust analysis of economic data stream

Very often in practice a primary aim of an economic time series analysis would be to provide an insight into the short term probabilistic features of the underlying model on base of constantly updated sample of a moderate length. On base of such a generally imprecise knowledge are made various economic decisions or predictions. In practice it is very important to answer a question what tells us a majority of data about future behavior of a majority of participants of a certain market. The answer is especially difficult in case of huge amounts of data generated by changing, multivariate model.

It is well known that economic time series almost always consist of atypical observations. Since the forecast from any time series model are based on the extrapolation of the historical patterns, if the parameters of the model estimates are very dependent on a few atypical observations resulting from isolated or nonrepeatable events, then the quality of the forecasts can be expected to be poor. Moreover, when these parameters have economic interpretations, the presence of undetected influential observations can lead the economist to wrong decisions.

In one dimension, the univariate quantile function uniquely characterizes the underlying distri-

bution. One could point out many quantile induced procedures having very good properties in terms of robustness. In high dimensions due to the lack of the natural ordering there is no universally preferred definition of multivariate quantiles. Statistical depth functions provide very successful representatives of the multivariate quantiles.

In this paper we study the properties of the location-scale depth procedures introduced by Mizera & Müller – we look into the probabilistic information of the underlying time series model carried by them. We focus our attention on short term multivariate quantile based description of the possible time series model. We study robustness and utility of such the description in a decision making process. In particular we investigate properties of the moving Student median (two dimensional Tukey median in a location–scale problem) and generalized band depth (depth function for functional data) for ARMA, GARCH in mean and SV models analysis. We illustrate theoretical considerations by means of empirical examples concerning financial crisis.

Key words: data stream, robustness, statistical depth function

BIBLIOGRAPHY

1. Kong L., Zuo Y. (2010), *Smooth Depth Contours Characterize the Underlying Distribution*, “Journal of the Multivariate Analysis” 101, 2222–2226
2. Kosiorowski D. (2010). *Depth Based Procedures for Estimation ARMA and GARCH Models*, Y. Lechevallier, G. Saporta (Red.) “Proceedings of COMPSTAT’2010”, Physica – Verlag, 1207 – 1214
3. Kosiorowski, D. (2011), *Local Characterization of a Time Series Model Using Generalized Tukey Depth*, “Proceedings of 58th Congress of the ISI – Dublin 2011”, ISI
4. Lopez–Pintado, S., Romo J. (2009). *On the Concept of Depth for Functional Data*. “Journal of the American Statistical Association” 104(486), 718 – 734.
5. Mizera I., C.H. Müller (2004), *Location-scale Depth (with discussion)*, “Journal of the American Statistical Association” 99, 949–966.

Kot Stanisław Maciej (Politechnika Gdańska)

Stochastyczne skale ekwiwalentności dla Polski

W pracy zaproponowano *stochastyczne skale ekwiwalentności* (SES), jako nowe narzędzia porównywania dobrobytu, nierówności i ubóstwa heterogenicznych grup gospodarstw domowych. SES jest to każda funkcja, która przekształca rozkład wydatków określonej grupy gospodarstw domowych w taki sposób, że wynikowy rozkład jest stochastycznie indyferentny z rozkładem wydatków gospodarstw grupy odniesienia. Kryterium stochastycznej indyferencji stanowi też

podstawę opracowania metody estymacji SES. Na podstawie mikro-danych z Budżetów Gospodarstw Domowych oszacowano wybrane parametryczne pragmatyczne SES dla Polski dla lat 2005-2010.

Słowa kluczowe: skale ekwiwalentności, indyferencja stochastyczna, estymacja, rozkłady wydatków

Stochastic equivalence scale for Poland

The paper presents the concept of the *stochastic equivalence scale* (SES) as a new tool of comparisons of welfare, inequality and poverty of heterogenic households' groups. The SES is any function that transforms the expenditure distribution of a specific group of households in such a way that the resulting distribution is stochastically indifferent from the expenditure distribution of a reference group of households. The stochastic indifference criteria are also used in developing the method of the estimation of the SES. Some selected pragmatic parametric SESs are estimated using micro-data from the Polish Household Budget Survey 2005-2010.

Key words: equivalence scale, stochastic indifference, estimation, expenditure distribution

Kowaleski Jerzy, Majdzińska Anna (Uniwersytet Łódzki)

Starzenie się populacji krajów Unii Europejskiej – nieodległa przeszłość i prognoza

Celem opracowania jest przedstawienie procesu starzenia się społeczeństw krajów dzisiejszej Unii Europejskiej oraz pokazanie dynamiki tego zjawiska w latach 1991–2030. Sytuacja w Polsce w tym względzie jest przedmiotem szczególnej uwagi. Jako cel cząstkowy traktujemy również zasygnalizowanie wybranych społecznych następstw i wyzwań wynikających z powiększania się obszaru starości demograficznej.

Wprowadzeniem do zasadniczej części tekstu jest omówienie pojęcia starzenia się ludności (w kontekście statystycznym) oraz przedstawienie wykorzystywanych w dalszej analizie miar zaawansowania i dynamiki starości demograficznej.

Część analityczna opracowania składa się z trzech wątków badawczych. Pierwszy z nich pokazuje powiększające się udziały osób starszych i sędziwych (czyli w wieku odpowiednio - 65 lat i więcej oraz 80 lat i więcej) w latach 1991-2030. Podejście takie pozwoliło na wyodrębnienie grup krajów Unii Europejskiej o odmiennym zaawansowaniu starości demograficznej i różnym tempie zmian mierników starości. W kolejnym wątku przedstawiona została ocena zaawansowania przemian demograficznych w rozpatrywanej zbiorowości krajów w oparciu o mierniki wynikające z relacji międzygrupowych, tj. indeksy starości i współczynniki potencjalnego wsparcia. Zakres czasowy

analizy w tym przypadku był taki sam jak w części poprzedniej. Dodatkowym przyczynkiem analitycznym jest przedstawienie poziomu zaawansowania i postępu starości demograficznej przy wykorzystaniu jednej z miar pozycyjnych (kwintyla IV).

Podjęta została również próba objaśnienia wewnętrznych różnic regionalnych w zaawansowaniu starości demograficznej poszczególnych krajów. Ograniczenia analityczne w tym względzie wynikają z dostępności danych statystycznych.

W zakończeniu opracowania zwrócono uwagę na wybrane społeczne następstwa procesu starzenia się ludności dla krajów UE. Wśród nich wskazano na spodziewany wzrost liczby osób niepełnosprawnych wywodzących się z populacji seniorów. Innym poważnym następstwem i wyzwaniem wynikającym z opisywanego procesu jest powiększająca się liczba i udział jednoosobowych gospodarstw domowych tworzonych przez osoby starsze, które w miarę przechodzenia do coraz bardziej zaawansowanych grup wieku będą wymagały asysty i opieki, zarówno rodzinnej jak i pozarodzinnej.

Słowa kluczowe: starzenie się populacji, miary starości demograficznej, kraje Unii Europejskiej

Population ageing of the European Union countries – the past and the perspectives

The aim of this paper is to present the population ageing of the European Union member countries and its dynamics in the years 1991–2030. In this respect, situation in Poland seems of prime attention. Furthermore, our partial aim is to indicate selected social consequences and challenges arising from the progress of demographic ageing.

The introductory part of the text discusses the definition of the population ageing (in the statistical aspect) and presents measures used later in the paper. The analytical part consists of three research threads. The first one treats about the growing number of older people (at the age of 65 and over, as well as 80 and over) in the years 1991–2030; the next one comprises an evaluation of the progress of demographic changes in the considered countries, which is based on measures resulting from the intergroup relations — that is the aging index (population 65 and over to population 0–14), and the potential support ratio (population 15 to 64 years to population 65 and over). In the analysis, there were also used positional measures. The last research part attempts at explaining regional differences in the progress of demographic ageing of particular countries.

At the end, the paper focuses on chosen social consequences of the ageing process of the European Union member countries' population. One of these consequences seems an expected increase in the number of the disabled persons originating in the old people population. Another serious result, and at the same time challenge, seems a growing number of one-person households of the old people, who reaching more advanced age groups need familial and extra-familial care.

Key words: population ageing, old age people, the European Union member countries

BIBLIOGRAPHY

1. Cieślak M. (2004), Pomiar procesu starzenia się, „Studia Demograficzne ” nr 2/146.
2. Długosz Z. (1998), Próba określenia zmian starości demograficznej Polski w ujęciu przestrzennym, „Wiadomości Statystyczne” nr 3, GUS-PTS, Warszawa.
3. Kot M., Kurkiewicz J. (2004), The new measures of the population ageing, “Studia Demograficzne” nr 2/146.
4. Kowaleski J.T. (2008), Struktura demograficzna starszego odłamu ludności (rozważania metodologiczne i elementy obrazu sytuacji w województwach i powiatach na przełomie stuleci) w: Kowaleski J.T., Szukalski P., red. (2008), Starzenie się ludności Polski. Między demografią a gerontologią społeczną, Wydawnictwo UŁ, Łódź.
5. Kowaleski J.T, Majdzińska A. (2011), Starzenie się ludności – aspekt krajowy i regionalny [w]: Organiściak-Krzykowska A. (red.), Regionalne aspekty rynku pracy, IPISS, Uniwersytet Warmińsko-Mazurski w Olsztynie, Warszawa-Olsztyn.
6. Kurek S. (2008), Typologia starzenia się ludności Polski w ujęciu przestrzennym, Wydawnictwo Naukowe Akademii Pedagogicznej, Kraków.
7. Rosset E. (1959), Proces starzenia się ludności, PWG, Warszawa.
8. Szukalski P. (2005), Następcy Matuzalema, ”Prace Instytutu Ekonometrii i Statystyki Uniwersytetu Łódzkiego” nr 140, Łódź.
9. UN, Population Division (2009), World Population Prospects. The 2008 Revision, Volume II, Sex and Age Distribution of the World Population, UN, New York.

Kowalewski Grzegorz (Uniwersytet Ekonomiczny we Wrocławiu)

Metodologia ankietowych badań koniunktury gospodarczej

Wśród metod służących do bieżącej diagnozy i prognozy koniunktury gospodarczej bardzo popularne są testy koniunktury. Podstawą tej metody są ankiety zawierające pytania dotyczące podstawowych wielkości ekonomicznych, symptomatycznych dla koniunktury. Badaniem mogą być objęte różne dziedziny gospodarcze i społeczne charakteryzujące się znacznym udziałem podmiotów gospodarczych podejmujących wolne decyzje ekonomiczne, działających w ramach gospodarki rynkowej (przemysł przetwórczy, budownictwo, handel, różne rodzaje usługi, gospodarstwa domowe).

Zbierane informacje w ankietach mają przeważnie charakter jakościowy, tzn. nie wymagają udzielenia odpowiedzi ilościowych. Najczęściej pytania są w formie zamkniętej, jednokrotnego

wyboru z trzema wariantami odpowiedzi: pozytywnym, neutralnym i negatywnym. Dla każdego takiego pytania oblicza się strukturę odpowiedzi sumującej się do 100%. Na podstawie rozkładu odpowiedzi oblicza się wskaźnik prosty jako różnicę między procentowym udziałem odpowiedzi pozytywnych i negatywnych, co tworzy tzw. saldo odpowiedzi na dane pytanie. Przy obliczaniu wskaźnika prostego nie są brane pod uwagę odpowiedzi neutralne, co może być uznane za wadę takiego sposobu obliczania wskaźników prostych. Bowiemy wtedy „sklejane” są różne struktury odpowiedzi - saldo obliczane dla dwóch różnych struktur odpowiedzi może dać taką samą wartość wskaźnika prostego.

W referacie zostaną więc przedstawione inne metody agregacji struktury odpowiedzi:

1. indeks dyfuzji;
2. suma odpowiedzi pozytywnych i negatywnych jako uzupełnienie salda odpowiedzi;
3. metoda probabilistyczna lub regresyjna. Aby można było je zastosować konieczne jest skorzystanie nie tylko ze struktury odpowiedzi pytania prognostycznego danego zjawiska ale także z rzeczywistych wartości badanego zjawiska. Zatem metody te można stosować jeśli pytanie w ankiecie dotyczy mierzalnego zjawiska. Większość pytań ankiety nie ma odpowiedników ilościowych. Z tego powodu metody te stosuje się w zasadzie tylko do kwantyfikacji oczekiwanej zmiany cen.

W referacie będą też przedstawione wady (m.in. subiektywizm, brak kwantyfikacji, wycinkowy zasięg badań, kosztowność, wybiórczość informacji, wąski zakres informacji, reprezentacyjność) i zalety (aktualność, regularność, możliwość pozyskiwania wiedzy niedostępnej w inny sposób, predykcja i inne) badań ankietowych w porównaniu z ilościowym opisem koniunktury gospodarczej.

Rozważania dotyczące sposobów korzystania z wyników ankietowych będą ilustrowane przykładami z badań koniunktury gospodarczej przedsiębiorstw prowadzonych przez Główny Urząd Statystyczny.

Słowa kluczowe: testy koniunktury, saldo koniunktury, indeks dyfuzji

Methodology of business tendency surveys

Among methods serving as current diagnosis and business cycle forecast - business tendency surveys are very popular. The basis of this method includes surveys containing questions related to basic economic amounts symptomatic to business tendency. The research can involve various business and social areas characterized by major contribution of business entity making free economic decisions and functioning within market economy (processing industry, construction, trade, services, ménage).

Gathering information in the surveys is usually qualitative and there is no requirement of quantitative answers. The questions are mainly closed-end, single-choice questions with three versions of responses (positive, neutral or negative).

For each such question one calculates distribution of answers that adds up to 100%.

On the basis of such distribution of answers one calculates a simple index as a difference between percentage contribution of positive and negative answers, what forms so-called balance of answers to a given question.

In the calculation of a simple index neutral answers are not taken into account. It can be seen as a fault of this kind of calculation, because in this case different answer structures are combine. The balance calculated for two different answer structures can give the same value of a simple index.

The paper is going to present other methods of answer structure aggregation:

1. diffusion indices
2. sum of positive and negative responses as a supplement of balance
3. probabilistic and regressing method. In order to apply them there is a need to use not only the answer structure of prognostic question of the phenomena but also the real value of the analysed phenomena. Therefore, these methods may be used if the question in the survey is related to a measurable phenomena. The majority of the survey question does not have quantitative equivalents. That is why they are used only in the quantification of inflation expectations.

The paper will also present the disadvantages (subjectivism, the lack of quantification, fragmentary survey range, costliness, information selectiveness, narrow information range, representativeness) and advantages (topicality, frequency, the ability to source the inapproachable knowledge, prediction and others) of the survey in comparison to a quantitative description of business cycle.

Consideration related to the ways of usage of the survey results will be illustrated with the examples of business cycle research conducted by Central Statistical Office.

Key words: business tendency surveys, balance of answers, diffusion indices

Kowalewski Jacek, Natkowska Monika (Urząd Statystyczny w Poznaniu)

Wykorzystanie mapy statystyki krótkookresowej w procesie organizacji badań przedsiębiorstw prowadzonych przez statystykę publiczną

Statystyka krótkookresowa obejmuje zbieranie, przetwarzanie i dostarczanie szybkiej i wielodziedzinowej informacji w sferze gospodarczej. W ramach organizacji pracy w statystyce publicznej została ona przydzielona jako specjalizacja Urzędowi Statystycznemu w Poznaniu. W Polsce statystyka krótkookresowa realizowana jest głównie poprzez badanie DG-1, które charakteryzuje

się tym, że jest dużym badaniem (ponad 30 tysięcy podmiotów) przeprowadzanym co miesiąc w bardzo krótkim okresie czasu, z rozbudowanym systemem przetwarzania (łączenie danych z innych badań, uogólnianie oraz przeliczenia na ceny stałe). Konieczność redukcji obciążeń respondentów oraz rozwój nowych technik informatycznych i statystycznych wymusza podejmowanie prób usprawnienia organizacji przeprowadzenia badania. Jednym z takich działań jest opracowanie mapy statystyki krótkookresowej, która pozwala na określenie w jaki sposób pozyskiwane dane są przekształcane w informację wynikową. Pozwala ona również na identyfikację odbiorców oraz czasu i kanałów dystrybucji poszczególnych informacji. Artykuł stanowi próbę ukazania możliwości wykorzystania mapy w ramach optymalizacji badań statystycznych.

Słowa kluczowe: statystyka krótkookresowa, statystyka przedsiębiorstw, badanie działalności gospodarczej, optymalizacja badań, termin publikacji danych statystycznych

Using a map of short term statistics in the process of organizing business surveys conducted by public statistics

Short-Term Statistics (STS) deals with collecting, compiling and providing quick and multi-domain information about economic activity. Within the organizational system of surveys in Polish official statistics, conducting STS surveys is the task of the Statistical Office in Poznań, which specializes in this field. The main source of STS data is the survey of economic activity, code-named DG-1, which covers over 30 thousands reporting units every month. The system of data processing in DG-1 is very complex and involves combining data from other surveys, generalizing and converting data to constant prices. Faced with the necessity to reduce the response burden and given modern advances in computer and statistical techniques, we are obliged to continue our efforts at improving the organization of surveys. One step in this process is compiling „*A Map of Short-Term Statistics*”, which enables us to determine how input data are processed into output information and to identify users of statistical information, distribution channels and data release deadlines. This article provides an example of how the Map can be used in the process of optimizing statistical surveys.

Key words: short-term statistics, business statistics, survey of economic activity, users of statistical information, survey optimization, data release deadlines

BIBLIOGRAPHY

1. GUS 2010. Założenia i wytyczne do system informacji o jednostkach prowadzących działalność gospodarczą na podstawie danych zawartych w sprawozdaniu miesięcznym DG-1 w 2010 r.
2. Eurostat 2006. Methodology of short-term business statistics. Interpretation and guidelines. European Communities. Luxembourg: Office for Official Publications of the European Communities.
3. Regulation No 1158/2005 of the European Parliament and of the Council amending Council Regulation No 1165/98 concerning short-term statistics.
4. GUS 2010. Plan publikacji GUS na 2010 rok.

5. GUS 2010. Plan publikacji US na 2010 rok.
6. GUS 2009. Plan opracowań statystycznych Głównego Urzędu Statystycznego na 2010 rok.
7. GUS 2009. Program badań statystycznych statystyki publicznej na rok 2010. Załącznik do rozporządzenia Rady Ministrów z dnia 8 grudnia 2009 r. (Dz. U. Z 2010 r. Nr 3, poz. 14).
8. Kowalewski J., Model optymalizacji badań statystycznych // W: Prace Statystyczne i Demograficzne/ red. I.Roeske-Słomka /Wydawnictwo Uniwersytetu Ekonomicznego w Poznaniu/ Zeszyt naukowy 133/2010

Kowerski Mieczysław (Wyższa Szkoła Zarządzania i Administracji w Zamościu)

Badania nastrojów gospodarczych przedsiębiorców i konsumentów na poziomie regionu. Przykład województwa lubelskiego

Wyższa Szkoła Zarządzania i Administracji w Zamościu od II kwartału 2001 roku prowadzi kwartalne badania ankietowe umożliwiające obliczenie sektorowych wskaźników nastrojów gospodarczych w barometrach nastrojów gospodarczych w województwie lubelskim. Analiza przebiegu barometrów nastrojów gospodarczych skłania do wniosku, że gospodarka województwa lubelskiego w badanym okresie „przeszła” pełny średniookresowy cykl koniunkturalny, który rozpoczął się recesją w 2001 roku, i poprzez fazę ożywienia z lat 2004–2006 wszedł w fazę wzrostu w okresie 2007 – pierwsza połowa 2008 roku ze szczytem w IV kwartale 2007 roku, by od drugiej połowy 2008 roku wejść w fazę spowolnienia gospodarczego, które zakończyło się recesją lat 2009–2010 roku z „dołkiem” zanotowanym w II kwartale 2009 roku. Natomiast w dalszym ciągu pozostaje otwartym problem sposobu i tempa wychodzenia z recesji. Ostatnie badania wskazują iż jest możliwa sytuacja ponownego spowolnienia gospodarczego, po którym dopiero gospodarka regionu wejdzie w fazę ożywienia.

Wnioski te potwierdzają również uzyskane metodą wygładzania ekonometrycznego za pomocą filtru Hodricka – Prescottta długookresowe trendy barometrów i sektorowych wskaźników nastrojów gospodarczych.

W pracy zaproponowano też sposób wykorzystania regionalny barometru nastrojów gospodarczych do szacowania poziomu rozwoju gospodarczego województwa, w tym tempa zmian produktu krajowego brutto. Jest to szczególnie ważne w sytuacji, kiedy wartości PKP w województwach obliczany jest z ponad dwuletnim opóźnieniem.

Słowa kluczowe: Barometr nastrojów gospodarczych, Cykl koniunkturalny, Filtr Hodricka – Prescottta, Województwo lubelskie

Economic sentiment surveys of entrepreneurs and consumers on the region level. The case of Lubelskie voivodeship

Since the II quarter of 2001 the University of Management and Administration in Zamość has been realizing the economic sentiment surveys which help to calculate sector confidence indicators and economic sentiment indexes for Lubelskie region. The analysis of the trends of economic sentiment indexes permit to reach conclusion that in the investigated period the economy of Lubelskie region has gone through the full medium term business cycle, which began with 2001 recession and through the recovering phase of 2004–2006 entered the boom phase of 2007 – the first half of 2008 with the pick in the IV quarter of 2007 and then the slow down in the economy which ended in 2009–2010 recession with the “doldrums” registered in the II quarter of 2009. Whereas, the way and pace of coming out of 2009–2010 recession constitute a problem. The results of the last surveys that coming out of recession can be W-shaped where economy and economic sentiment reached “the second doldrums”.

These conclusions can be confirmed by econometric smoothing with the Hodrick-Prescott filter long-term trends of general and entrepreneurs’ economic sentiment indexes and sector confidence indicators.

In the paper the method of application of economic sentiment index to estimation of regional GDP rate of growth was proposed. It is very important because the regional GDP in Poland is calculated with two years delay.

Key words: economic sentiment surveys, Business cycle, Lubelskie voivodeship, Hodrick-Prescott filter

BIBLIOGRAPHY

1. Adamowicz E., *Polityka gospodarcza w województwie lubelskim – rekomendacje*, „Barometr Regionalny. Analizy i Prognozy”, nr 2(24), 2011, s. 7–14
2. Bracha Cz., *Metoda reprezentacyjna w badaniu opinii publicznej i marketingu*, Efekt, Warszawa 1998
3. *Business Tendency Surveys: a Handbook*, Source OECD, Transition Economies, March 2003, vol. 2003, no. 8
4. Drozdowicz-Bieć M., *Regionalne cykle koniunkturalne. Doświadczenia światowe – implikacje dla Polskie*, „Barometr Regionalny. Analizy i Prognozy”, nr 3(13), 2008, s. 5–15
5. Drozdowicz-Bieć M., *Wzrost w cieniu światowego kryzysu*, „Barometr Regionalny. Analizy i Prognozy”, nr 3(17), 2009, s. 7–13
6. Hodrick R., E. C. Prescott, Postwar U.S. *Business Cycles: An Empirical Investigation*, “Journal of Money, Credit and Banking”, 29/1997, s. 1–19
7. Kowalewski G., *Wartość informacyjna wskaźników koniunktury dla danych jakościowych* [w:] J. Garczarczyk (red.), *Rynek usług finansowych a koniunktura gospodarcza*, CE-DEWU.PL, Warszawa 2009, s. 341–349

8. Kowerski M., *Wpływ przystąpienia Polski do Unii Europejskiej na nastroje gospodarcze w województwie lubelskim*, „Gospodarka Narodowa”, nr 7 – 8/2005, s. 101–111
9. Kowerski M., *Wartość informacyjna odpowiedzi „bez zmian” w badaniach nastrojów gospodarczych*, „Barometr Regionalny. Analizy i Prognozy”, nr 4(14), 2008, s. 47–61
10. Kowerski M., *Wpływ wprowadzenia przepisów układu z Schengen na gospodarcze nastroje mieszkańców województwa lubelskiego*, „Biblioteka Wiadomości Statystycznych”, Tom 60, GUS 2008c, s. 282–293
11. Kowerski M., *Metodyka badania nastrojów gospodarczych w województwie lubelskim na tle badań Komisji Europejskiej*, „Barometr Regionalny. Analizy i Prognozy”, nr 3(17)/2009a, s. 15–28
12. Kowerski M., *Wahania koniunkturalne w województwie lubelskim*, „Gospodarka Narodowa”, nr 7–8/2009b, s. 93 – 106
13. Kufel T., *Ekonometria. Rozwiązywanie problemów z wykorzystaniem programu GRETL*, PWN, Warszawa, 2004
14. Maddala G. S., *Ekonometria*, Wydawnictwo Naukowe PWN, Warszawa 2006
15. *The Joint Harmonised EU Programme of Business and Consumer Surveys User Guide*, European Commission, Directorate General Economic and Financial Affairs, Brussels 4 July 2007
16. Welfe A., *Ekonometria*, PWE, Warszawa 1998

Krzanowski Wojciech (University of Exeter)

Classification: some old principles applied to new problems

Modern technology and instrumentation frequently generate very large data sets which give rise to new problems of analysis. Classification, whether unsupervised (i.e. cluster analysis) or supervised (i.e. discriminant analysis), is often of central interest. However, many seemingly complicated problems do not necessarily need commensurately complicated methods for solution, as they can be tackled using old and well-established principles.

This talk will consider three such problems that have arisen in practical applications of classification. When handling biological or medical microarray data, variables may need to be screened in order to select a small subset of the most effective ones for clustering. In financial credit scoring, a test may be required for deciding whether or not two scorecards are significantly different in performance. In risk assessment, a bank may wish to predict whether or not those customers who apply for a loan are likely to default on repayment. The first problem involves unsupervised

classification, the second involves supervised classification, and the third has a mixture of the two types.

All of these applications can be tackled using relatively simple adaptations of well-known classification and computational methods. In each case I will provide a motivation for the analysis, give a brief outline of the technical details, and summarize the results to demonstrate the efficacy of the methods.

Key words: Clustering, discrimination, error rates, ROC curves, sums of squares

Krzyżko Mirosław (Uniwersytet im. Adama Mickiewicza w Poznaniu)

Jan Czekanowski (1882–1965)

Czekanowski Jan -, ur. 6 X 1882 w Głuchowie, zm. 20 VII 1965 w Szczecinie; antropolog, etnolog, slawista. Odbił studia uniwersyteckie w l. 1902-06 w Zurychu w zakresie antropologii i matematyki. Od 1906 został asystentem w Konigliches Museum fur Volkerkunde w Berlinie, w dziale Afryki i Oceanii. W l. 1907-09 był uczestnikiem naukowej ekspedycji do Afryki zorganizowanej przez ks. meklemberskiego Adolfa Fryderyka. Ekspedycja badała rejon międzyrzecza Nilu i Konga a owocem jej było obszerne 5-tomowe dzieło napisane przez Czekanowskiego. Za oryginalny wkład w studia afrykanistyczne Czekanowski uzyskał liczne odznaczenia m. in. Order Korony Belgijskiej, Gryfa Meklemberskiego oraz medal pamiątkowy ekspedycji. Od 1911 C. był kustoszem Muzeum Etnograficzno-Antropologicznego Cesarskiej Akademii Nauk w Petersburgu a w l. 1913-1941 profesorem i kierownikiem katedry antropologii oraz dziekanem i rektorem uniwersytetu we Lwowie. Po II wojnie światowej objął przejściowo katedrę antropologii i etnografii na Katolickim Uniwersytecie Lubelskim, a w l. 1946-1960 był kierownikiem katedry antropologii na Uniwersytecie im. Adama Mickiewicza w Poznaniu.

Wprowadzenie przez Czekanowskiego do antropologii zasad statystyki matematycznej stanowiło przełom w metodyce tej nauki i dało początek odrębnej oryginalnej szkole antropologicznej znanej na świecie jako polska szkoła Czekanowskiego (lwowska szkoła antropologiczna). Oprócz zastosowania metod ilościowych przyjęto hipotezę genetyczną występowania typów czystych i rozszczepiania się mieszańców przy krzyżowaniu. Diagramiczna metoda Czekanowskiego umożliwiła analizę materiałów antropologicznych i znalazła szersze zastosowanie także w innych dyscyplinach, zwłaszcza w językoznawstwie. Czekanowski – przyjmując założenie weryfikowania danych antropologii w świecie etnologii, historii, archeologii itp. – wychodził daleko poza ramy swej dyscypliny. Dokonał syntetycznego ujęcia etnogenezy Słowian, uzasadniając tezę, że ich prapokolebką były ziemie dzisiejszej Polski zachodniej. Za prace slawistyczne odznaczony został Krzyżem Komandorskim św. Sawy. Był członkiem rzeczywistym PAN, doktorem honoris causa Uniwersytetu Wrocławskiego (1959). Międzynarodowa pozycja naukowa Czekanowskiego znalazła wyraz w nadaniu członkostw honorowych: *Societa Italiana di Antropologia e Etnografia*, Cesarskiego Rosyjskiego Towarzystwa Geograficznego, British Royal Anthropological Institute.

Był także członkiem korespondentem *Suomalais Ugrilainen Seura* i Instytutu Słowiańskiego w Pradze, członkiem towarzystw antropologicznych w Paryżu i w Wiedniu.

Słowa kluczowe: Czekanowski Jan, Antropologia, Zasady statystyki, Metoda diafrazowa Czekanowskiego

Jan Czekanowski (1882 - 1965)

Czekanowski Jan – (born October 6 1882 in Gluchow – died July 20 1965 in Szczecin) anthropologist, ethnologist, Slavist. He attended university in the years 1902 to 1906 in Zurich in the field of anthropology and mathematics. From 1906 he was appointed assistant at the Royal Museum of Folk Culture in Berlin, in the “Africa and Oceania” section. In the years 1907-1909 he was member of a scientific expedition to Africa, organized by the Duke Adolf Friedrich of Mecklenburg. The expedition examined the region between Nile and Congo and it resulted in an extensive 5-volume work written by Czekanowski. In recognition of his accomplishments and contribution in African studies the scientist was presented with the Order of the Belgian Crown, the Order of Mecklenburg Griffon and a commemorative medal of the expedition among others. In 1911 he was appointed custodian of the Museum of Ethnography and Anthropology of the Imperial Academy of Sciences in St. Petersburg. In the years 1913-1941, Czekanowski was head of the Anthropology Department, Dean and President of the University in Lviv. After World War II, he temporarily became professor at the Catholic University of Lublin and from 1946 till 1960 he was head of the Anthropology Department Adam Mickiewicz University in Poznan.

By introducing principles of mathematical statistics to anthropology, Professor Czekanowski marked a breakthrough in the methodology of the study and became the founder of an original anthropological school known in the world as the Lviv School of Anthropology, which for many years played the lead in all anthropological research in Poland. In addition to the application of quantitative methods the school accepted the hypothesis of pure genetic types and incidence of splitting in hybrids when crossing. The Czekanowski’s diaphrasic method enabled the analysis of anthropological materials and found a wider application in other disciplines, especially in linguistics. Assuming the verification of anthropology data using other sciences like ethnology, history, archeology, etc., Czekanowski went far beyond the borders of his discipline. Using a statistical approach and gathering anthropometric materials, he synthesized the ethnogenesis of the Slavs, justifying the idea that the cradle of the Slavonic tribes lies in the Doab of the Vistula River and the Oder River. For his works on the Slavs he was awarded the Order of St. Sava. He was Full Member of the Academy of Sciences, doctor honoris causa at the University of Wrocław (1959). Czekanowski’s international scientific position found expression in the granting of honorary membership of the following organizations: *Societa Italiana di Antropologia e Etnografia*, the Imperial Russian Geographical Society, British Royal Anthropological Institute. He was also a corresponding member of *Suomalais Ugrilainen Seura*, the Slavic Institute in Prague and a member of the anthropological societies of Paris and Vienna.

Key words: Czekanowski Jan, Anthropology, Statistical rules, Czekanowski’s diaphrasic method

BIBLIOGRAPHY

1. Archiwum Uniwersytetu im. Adama Mickiewicza w Poznaniu.
2. Chodźdło T., *Wspomnienie pośmiertne. Śp. prof. Jan Czekanowski (1882-1965)*, "Zeszyty Naukowe KUL" 9(3), 1966, 91-94.
3. Ćwirko-Godycki M., *Profesor Jan Czekanowski*, "Przegląd Antropologiczny" XXXI(2), 1965, 211-233.
4. Gajek J., *Jan Czekanowski. Sylwetka uczonego*, "Nauka Polska" 6(2), 1958, 118-127.
5. Malinowski A., *Życie i działalność Jana Czekanowskiego*, w: J. Piontek i inni (red.), *Teoria i empiria w polskiej szkole antropologicznej. W 100-lecie urodzin Jana Czekanowskiego*, Poznań 1985, 71-77.
6. Oderfeld J., *Historia życia cz. II-Petroniusz i Ursus*. <http://www.absolwenci.sieniu.czest.pl/index.php?show=art&which=1468>
7. Perkal J., *Jan Czekanowski (1882-1965)*, "Listy Biometryczne" 9-11(1965), III-IV.
8. Szelaąg Z., *Grójeckie we wspomnieniach*, Seria IV, Towarzystwo Literackie im. Adama Mickiewicza, Oddział w Grójcu, 2004.
9. Wanke A., *Sześćdziesiąt lat pracy naukowej Jana Czekanowskiego*, "Materiały i Prace Antropologiczne" 70(1964), 7-27.
10. Wokroj F., *50-lecie promocji doktorskiej prof. dra Jana Czekanowskiego, 1906-1956*, "Życie Szkoły Wyższej" 7/8(1957), 93-96.

Krzyśko Mirosław, Waszak Łukasz (Uniwersytet im. Adama Mickiewicza w Poznaniu)

Jądrowe korelacje kanoniczne

Założmy, że $\mathbf{X}=(\mathbf{X}_1, \dots, \mathbf{X}_N)'$ i $\mathbf{Y}=(\mathbf{Y}_1, \dots, \mathbf{Y}_N)'$ są dwoma scentrowanymi N -elementowymi zbiorami zmiennych losowych oraz $\mathbf{X}_i \in \mathbb{R}^p$, $\mathbf{Y}_i \in \mathbb{R}^q$ dla $i = 1, \dots, N$. Klasycznym zadaniem analizy korelacji kanonicznych jest znalezienie zależności pomiędzy tymi zbiorami, tzn. znalezienie liniowej kombinacji $U = \mathbf{a}'\mathbf{X}$ i $V = \mathbf{b}'\mathbf{Y}$ oryginalnych zmiennych, mających maksymalną korelację, gdzie $\mathbf{a}, \mathbf{b} \in \mathbb{R}^n$, przy założeniu że $E(U^2) = 1 = E(V^2)$. Naszym zadaniem jest więc maksymalizacja korelacji $\rho = \rho(U, V) = E(UV)$. Problem ten jest równoważny rozwiązaniu uogólnionego zagadnienia własnego (*) $\mathbf{S}_0 \mathbf{c} = \rho \mathbf{S}_D \mathbf{c}$, gdzie $\mathbf{S}_D = \text{diag}(\mathbf{S}_{XX}, \mathbf{S}_{YY})$, $\mathbf{S}_0$ jest

macierzą, która poza główną przekątną ma elementy $\mathbf{S}_{XY}$ i $\mathbf{S}_{YX}$, a na przekątnej blokowe macierze złożone z samych zer, natomiast $\mathbf{S}_{XX}=\mathbf{X}'\mathbf{X}$, $\mathbf{S}_{YY}=\mathbf{Y}'\mathbf{Y}$, $\mathbf{S}_{XY}=\mathbf{X}'\mathbf{Y}=\mathbf{S}_{YX}'$. Największe ρ (będące rozwiązaniem tego zagadnienia) jest pierwszym współczynnikiem korelacji kanonicznej.

W celu skonstruowania nieliniowych (jądrowych) współczynników korelacji kanonicznych należy odwzorować przestrzeń $\mathbb{R}^p$ do przestrzeni Hilberta z jądrami reprodukującymi $H(k_x)$: $\varphi: \mathbb{R}^p \rightarrow H(k_x)$, a przestrzeń $\mathbb{R}^q$ do przestrzeni Hilberta z jądrami reprodukującymi $H(k_y)$: $\psi: \mathbb{R}^q \rightarrow H(k_y)$. W nowych przestrzeniach otrzymamy nieliniowe współczynniki korelacji kanonicznych, rozwiązując uogólnione zagadnienie własne $(**)$ $\mathbf{K}_O\tau=\rho\mathbf{K}_D\tau$, gdzie $\mathbf{K}_D=\text{diag}(\mathbf{K}_x^2, \mathbf{K}_y^2)$, $\mathbf{K}_O$ jest macierzą, która poza główną przekątną ma elementy $\mathbf{K}_x\mathbf{K}_y$ oraz $\mathbf{K}_y\mathbf{K}_x$, a na przekątnej blokowe macierze złożone z samych zer, natomiast $\mathbf{K}_x=[k_x(\mathbf{X}_i, \mathbf{X}_j)]$ i $\mathbf{K}_y=[k_y(\mathbf{Y}_i, \mathbf{Y}_j)]$ są macierzami jądrowymi, których elementami są funkcje jądrowe, gdzie $i, j = 1, \dots, N$. Nowe, jądrowe uogólnione zagadnienie własne ma zawsze rozwiązanie $\rho = 1$ i jest to typowy problem nadmiernego dopasowania. Przedstawimy metody rozwiązania tego problemu oraz porównamy wyniki otrzymane za pomocą klasycznej i jądrowej analizy korelacji kanonicznej.

Słowa kluczowe: analiza korelacji kanonicznych, uogólnione zagadnienie własne, przestrzeń Hilberta z jądrem reprodukującym

Kernel canonical correlation

Suppose that $\mathbf{X}=(\mathbf{X}_1, \dots, \mathbf{X}_N)'$ and $\mathbf{Y}=(\mathbf{Y}_1, \dots, \mathbf{Y}_N)'$ are two centered N -elements sets of random variables and $\mathbf{X}_i \in \mathbb{R}^p$, $\mathbf{Y}_i \in \mathbb{R}^q$ for $i = 1, \dots, N$. Classical canonical correlation analysis seeks the associations between these sets, i.e. it searches for a linear combination $U = \mathbf{a}'\mathbf{X}$ and $V = \mathbf{b}'\mathbf{Y}$ of the original variables having maximal correlation, where $\mathbf{a}, \mathbf{b} \in \mathbb{R}^n$, assuming $E(U^2) = 1 = E(V^2)$. Our task is to maximize the correlation $\rho = \rho(U, V) = E(UV)$. This problem is equivalent to solving a generalized eigenvalue problem $(*)$ $\mathbf{S}_O\mathbf{c}=\rho\mathbf{S}_D\mathbf{c}$, where $\mathbf{S}_D=\text{diag}(\mathbf{S}_{XX}, \mathbf{S}_{YY})$, $\mathbf{S}_O$ is an off-diagonal matrix with $\mathbf{S}_{XY}$ and $\mathbf{S}_{YX}$ on the off diagonal and $\mathbf{S}_{XX}=\mathbf{X}'\mathbf{X}$, $\mathbf{S}_{YY}=\mathbf{Y}'\mathbf{Y}$, $\mathbf{S}_{XY}=\mathbf{X}'\mathbf{Y}=\mathbf{S}_{YX}'$. The maximal ρ (being a solution of this problem) is the first canonical correlation.

To construct nonlinear (kernel) canonical correlations, a space $\mathbb{R}^p$ must be mapped to reproducing a kernel Hilbert space $H(k_x)$: $\varphi: \mathbb{R}^p \rightarrow H(k_x)$ and a space $\mathbb{R}^q$ must be mapped to reproducing kernel Hilbert space $H(k_y)$: $\psi: \mathbb{R}^q \rightarrow H(k_y)$. In the new spaces we will obtain a nonlinear canonical correlation by solving the generalized eigenvalue problem $(**)$ $\mathbf{K}_O\tau=\rho\mathbf{K}_D\tau$, where $\mathbf{K}_D=\text{diag}(\mathbf{K}_x^2, \mathbf{K}_y^2)$, $\mathbf{K}_O$ is a off-diagonal matrix with $\mathbf{K}_x\mathbf{K}_y$ and $\mathbf{K}_y\mathbf{K}_x$ on the off diagonal and $\mathbf{K}_x=[k_x(\mathbf{X}_i, \mathbf{X}_j)]$ and $\mathbf{K}_y=[k_y(\mathbf{Y}_i, \mathbf{Y}_j)]$ are kernel matrices, whose elements are kernel functions, where $i, j = 1, \dots, N$. The new kernel generalized eigenvalue problem has always the solution $\rho = 1$ and this is a typical case of over-fitting. We present methods to solve this problem and we compare results obtained by classical and kernel canonical correlation analysis.

Key words: canonical correlation analysis, generalized eigenvalue problem, reproducing kernel Hilbert space

Kukło Cezary (Uniwersytet w Białymstoku)

Rodzina i jej gospodarstwo na ziemiach polskich do końca epoki przedprzemysłowej - mity i prawda

Obraz rodziny i gospodarstwa domowego do połowy XIX stulecia jeszcze do niedawna w polskiej historiografii był kształtowany przez nieliczne pamiętniki epoki, bardziej dotyczące elit niż grup średniozamożnych (nie mówiąc już o mieszczaństwie czy chłopach) z jednej strony, z drugiej przez głośny artykuł angielskiego demografa Johna Hajnala z 1965 r., poświęcony europejskiemu modelowi małżeństwa. W opinii J. Hajnala małżeństwo i rodzina na ziemiach polskich sytuowała się w modelu wschodnioeuropejskim, którego cechami były m.in. wczesne zawieranie małżeństw, niski odsetek osób pozostających w celibacie czy sporadyczne występowanie w gospodarstwach czeladzi złożonej z bezżennych młodych ludzi. Znaczący udział wśród gospodarstw domowych miały mieć gospodarstwa rodzin rozszerzonych o krewnych zarówno w linii bocznej, jak i wspólnie zamieszkujących rodziców. Z punktu widzenia liczby generacji dominowały rodziny trzy- i czteropokoleniowe.

Prezentowany referat stanowi polemikę z podstawowymi tezami Johna Hajnala i przedstawia syntetycznie wyniki najnowszych polskich badań z ostatniego ćwierćwiecza w zakresie struktur rodzinnych i gospodarstw domowych. Autor odwołuje się przy tym do prac które powstały z zastosowaniem dwóch światowych technik badawczych stosowanych w demografii historycznej, tj. metody francuskiej Louis Henry i klasyfikacji gospodarstw domowych autorstwa Peter Laslett.

Słowa kluczowe: małżeństwo, rodzina, gospodarstwo domowe, demografia historyczna

MYths and the truth on the family and its household on the territory of Poland by the end of the pre-industrial era

A picture of the family and its household up to the mid-nineteenth century was shaped even recently in the Polish historiography, by a few diaries of the era, which were referring more to elites than a group of middle wealthy class (not to mention the burghers or peasants) on one hand, on the other - by the popular article from 1965 of an English demographer - John Hajnal, that was dedicated to the European model of marriage. In the opinion of J. Hajnal marriage and the family on the territory of Poland was part of the Eastern European model, whose characteristic features were i.e.: early marriages, low percentage of people living in celibacy or occasional occurrence of servants comprised of unmarried young people in households.

A significant share of households was those extended by the relatives in both the lateral line, and parents living together. In terms of the number of generations the three- and four-generational families dominated.

The presented paper is polemics with the basic thesis of John Hajnal. It presents synthetically

the results of the latest Polish researches referring to family structures and households from the last quarter of the century. The author bases on the papers which were arisen from an application of two global research techniques used in historical demography (the French method of Louis Henry and the classification of households used by Peter Laslett).

Key words: marriage, family, household, historical demography

Kupiszewski Marek (Central European Forum for Migration and Population Research)

What drives population change?

The aim of this overview paper is to look at the main drivers of population change. We will examine the evolution of the populations magnitude and structures in the last half of the century and then we will discuss the role each of the components of change (fertility, mortality and migration) in this evolution.

Both, for population in general and for the components of change we will look at the following aspect:

- long term developments,
- recent changes,
- geographical diversification of patterns in Europe,
- the grand theories explaining evolution of these components of change,
- the main dilemmas concerning future developments.

As the result of the above considerations possible divergent trajectories of future developments of migration, fertility and mortality. In particular, for fertility a return to sub-replacement fertility versus "low fertility trap" are considered. For migration a concept and consequences of replacement migration will be looked into.

Kwiatkowska-Ciotucha Dorota, Załuska Urszula (Uniwersytet Ekonomiczny we Wrocławiu)

Ocena wykorzystania środków z europejskiego funduszu społecznego w poszczególnych regionach polski w obecnym okresie programowania

W referacie przedstawione zostaną zarówno możliwości pozyskiwania środków z Europejskiego Funduszu Społecznego w latach 2007 - 2013, jak i wyniki absorpcji środków w poszczególnych województwach. Analizie poddana będzie realizacja różnych form wsparcia w ramach komponentów regionalnych i komponentów centralnych Programu Operacyjnego Kapitał Ludzki. Do oceny wykorzystania środków w poszczególnych regionach Polski zastosowane zostaną metody wielowymiarowej analizy porównawczej.

Słowa kluczowe: Europejski Fundusz Społeczny, wielowymiarowa analiza porównawcza, Program Operacyjny Kapitał Ludzki

Evaluation of the implementation of funds from the european social fund in different regions of Poland within the current programming period

The paper presents both the possibilities of obtaining funds from the European Social Fund in the years 2007 - 2013, and the results of their absorption in each voivodship. The implementation of different forms of support within the regional and centralized components of the Human Capital Operational Programme will be the subject to the analysis. To evaluate the use of funds in various regions of Poland, the multivariate comparative analysis methods will be implemented.

Key words: European Social Fund, multivariate comparative analysis, Human Capital Operational Programme

Ledwina Teresa (Instytut Matematyczny PAN)

Jerzy Neyman (1894–1981)

W referacie przedstawione zostaną podstawowe fakty z życia i działalności Jerzego Neymana. Referat oparty będzie głównie na źródła podanych poniżej.

Słowa kluczowe: estymacja przedziałowa, metody asymptotyczne, metoda reprezentacyjna, testowanie hipotez

Jerzy Neyman (1894-1981)

In the talk we shall present main facts from the life and scientific work of Jerzy Neyman. The presentation shall mainly be based on the references listed below.

Key words: confidence intervals, asymptotic methods, representative method, testing statistical hypotheses

BIBLIOGRAPHY

1. Kendall, D.G., Bartlett, M.S., Page, T.L. „Jerzy Neyman 1894-1981”, Biographical Memoirs of Fellows of the Royal Society, (1982), 28, 378-412.
2. Klonecki, W. „Jerzy Neyman (1894 - 1981)”, Probability and Mathematical Statistics, 15 (1995), 7-14.
3. Klonecki, W., Zonn, W. „Jerzy Sława-Neyman”, Wiadomości Matematyczne XVI (1973), 55-70.
4. LeCam, L., Lehmann, E.L. „J. Neyman. On the occasion of his 80th birthday”, Annals of Statistics, 2 (1974), vii-xiii.
5. LeCam, L. „Neyman and stochastic models”, Probability and Mathematical Statistics, 15 (1995), 37-45.
6. Lehmann, E.L. „Jerzy Neyman 1894-1981”, Biographical Memoir, National Academy of Sciences, (1994), 395-420, Washington, D.C.
7. Łazowska, B. „Bibliografia prac prof. dr Jerzego Neymana (1894 - 1981)”, Zestawienia Bibliograficzne, Centralna Biblioteka Statystyczna im. Stefana Szulca, 22 (1995).
8. Reid, C. „Neyman - from life”, (1982), Springer, New York.
9. Scott, E.L. „Neyman Jerzy”, Encyclopedia of Statistical Sciences, 8 (2006), 5479-5487, Wiley- Interscience.

Ledwina Teresa, Wyłupek Grzegorz (Instytut Matematyczny PAN, Oddział Wrocław)

Test dla dwóch prób przy jednostronnych alternatywach

Referat będzie oparty na pracy Ledwiny i Wyłupka (2012). Rozważamy problem dwóch prób bez żadnych dodatkowych założeń dotyczących dystrybuant z wyjątkiem, że są ciągłe. Zatem, obserwacje składające się z dwóch niezależnych prób $X_1, \dots, X_m$ oraz $Y_1, \dots, Y_n$ pochodzących

z populacji o ciągłych dystrybuantach F i G , odpowiednio. G opisuje grupę zabiegową, podczas gdy F odpowiada grupie kontrolnej.

Pytanie brzmi, czy zabieg ma pozytywny wpływ, w tym sensie, że zwiększa odpowiedź. Formalny opis pozytywnej reakcji jest następujący

$$F(x) \geq G(x) \quad \text{dla wszystkich } x \text{ z ostrą nierównością przynajmniej dla jednego } x. \quad (1)$$

Relacja $F(x) \geq G(x)$ dla wszystkich x znana jest jako stochastyczny porządek lub stochastyczna dominacja pierwszego rzędu.

Problem testowania formułujemy w następujący sposób:

$\mathcal{H}$: brak stochastycznej dominacji (1)

przeciwko

$\mathcal{A}_+$: obecność stochastycznej dominacji (1).

Nasze niestandardowe postawienie problemu motywowane jest istniejącymi przykładami, które pokazują, że standardowe postawienie problemu i odpowiadające mu procedury mogą prowadzić do poważnych błędów we wnioskowaniu.

Rozwiązanie tak postawionego problemu, przedstawione w pracy Ledwiny i Wyłupek (2012), jest intuicyjnym i naturalnym rozwinięciem rozwiązań i idei przedstawionych w pracach Janic-Wróblewskiej i Ledwiny (2000), Inglota i Ledwiny (2006a) oraz Wyłupek (2010). Oczywiście, bieżący problem jest dużo trudniejszy ponieważ wspomniane prace rozważają klasyczny problem testowania jednorodności dystrybuant oraz testowanie zgodności. Dlatego też, kontrola błędu pierwszego rodzaju była łatwa. Poza tym, w przeciwieństwie do tych prac, teraz alternatywa jest ograniczona. Wymaga to innej implementacji i zastosowania bardziej wyrafinowanych narzędzi probabilistycznych do badania nowego rozwiązania.

Procedura, którą przedstawiamy łączy testowanie i selekcję modelu. Dokładniej, reparametryzujemy problem testowania w języku współczynników Fouriera, dobrze znanej gęstości porównawczej. Ponieważ alternatywa $\mathcal{A}_+$ jest jednostronna, rozważamy układ funkcji, który zachowuje stochastyczny porządek. Następnie, estymowane współczynniki Fouriera zostają użyte do zdefiniowania pewnego rodzaju znakowanej gładkiej statystyki rangowej. W takim ujęciu, liczba współczynników Fouriera włączanych do statystyki jest parametrem gładkości. Parametr ten wybieramy z użyciem elastycznej reguły selekcji. To prowadzi do testu adaptacyjnego dla $\mathcal{H}$ przeciwko $\mathcal{A}_+$. Test ten kontroluje asymptotyczne błędy obu rodzajów w satysfakcjonującym stopniu. Rezultaty te są potwierdzone przez obszerne symulacje.

Słowa kluczowe: selekcja modelu, stochastyczna dominacja, stochastyczny porządek, test adaptacyjny, test nieparametryczny

T wo-sample test against one-sided alternatives

Our contribution is based on the paper Ledwina & Wyłupek (2012). We consider the two-sample problem without any restrictions concerning the form of the underlying distributions, except the

assumption that they are continuous. The observations then consist of two independent samples $X_1, \dots, X_m$ and $Y_1, \dots, Y_n$ from two populations with continuous distribution functions F and G , respectively. G describes the treatment effect while F pertains to the control group.

The question is whether the treatment has a beneficial effect in the sense that it increases the response. Precisely, we formalize it as follows:

$$F(x) \geq G(x) \quad \text{for all } x \text{ with strict inequality for at least one } x. \quad (1)$$

The relation $F(x) \geq G(x)$ for all x is known as stochastic order or stochastic dominance of the first degree.

We formulate the testing problem as follows:

$\mathcal{H}$: lack of stochastic dominance (1)

against

$\mathcal{A}_+$: presence of stochastic dominance (1).

Our nonstandard formulation of the problem is motivated by existing evidence showing that standard formulations and pertaining procedures may lead to serious errors in inference.

The solution of a problem such posed, introduced in Ledwina & Wyłupek (2012), is an intuitive and natural development of the solutions and ideas presented in Janic-Wróblewska & Ledwina (2000), Inglot & Ledwina (2006a), and Wyłupek (2010). Though, evidently, the present problem is much more difficult to handle, as the above papers considered classical homogeneity testing and goodness-of-fit problem. Therefore, for example, control of the error of the first kind was easy. Moreover, in contrast to these contributions, now the alternative is a restricted one. This requires a different implementation and the application of more refined probabilistic tools to investigate the new solution.

The procedure that we introduce matches testing and model selection. More precisely, we reparametrize the testing problem in terms of Fourier coefficients of well known comparison densities. Since the alternative $\mathcal{A}_+$ is one-sided, we consider a system which preserves the stochastic order. Next, the estimated Fourier coefficients are used to form a kind of signed smooth rank statistic. In such a setting, the number of Fourier coefficients incorporated into the statistic is a smoothing parameter. We determine this parameter via some flexible and fully automatic selection rule. This leads to a data driven test for $\mathcal{H}$ against $\mathcal{A}_+$. The test asymptotically controls errors of both kinds to a satisfactory extent. This result is further supported by extensive simulations.

Key words: data driven test, model selection, nonparametric test, stochastic dominance, stochastic order

BIBLIOGRAPHY

1. Inglot T., Ledwina T. (2006a), "Towards Data Driven Selection of a Penalty Function for Data Driven Neyman Tests", Linear Algebra and its Applications, Vol. 417, 124–133.

2. Janic-Wróblewska A., Ledwina T. (2000), "Data Driven Rank Test for Two-Sample Problem", Scandinavian Journal of Statistics, Vol. 27, 281–297.
3. Ledwina T., Wyłupek G. (2012), "Two-Sample Test Against One-Sided Alternatives", accepted to Scandinavian Journal of Statistics.
4. Wyłupek G. (2010), "Data-Driven k -Sample Tests", Technometrics, Vol. 52, 107–123.

Lemmi Achille (Siena University, Dipartimento di Economia Politica), Panek Tomasz (Warsaw School of Economics)

The dimensions of poverty: theory, models and new perspectives

Poverty analysis literature is generally concerned with identifying the characteristic aspects of financial deprivation. This approach does not take into account other aspects that are specific of well-being and living conditions. In this paper the most promising extensions of this field are identified, and effective measures in order to capture fully the complexity of well-being are proposed.

First one must realise that deprivation is a matter of degree, hence fuzzy set theory needs to be adopted. Second, although income can be exchanged for many other goods and services, it does not fully represent many fundamental aspects of well-being such as stock investments (e.g. debt and durable goods) or surrounding context (e.g. crime and pollution). Finally, given a generic level of deprivation the impact that such level has on the well-being can change depending on if it is temporary or permanent.

This time evaluation through different time periods allows to identify very different scenarios, even if the underlying statistics or micro-level distributions are equal, hence would not be captured by static measures. These extensions allow a complete overview of what is commonly intended for economic well being, can be adopted simultaneously, and are compatible with traditional poverty analysis, which can be described in this context through a specific membership functions.

Moreover, indicators are most useful when they are comparable across time and space, as this allows determining trends and making meaningful comparisons between different situations. Policy research and applications require disaggregated statistics to increasingly lower levels and smaller subpopulations in order to take informed decisions, as national statistics are insufficiently precise.

After having discussed these extensions to traditional poverty analysis, the evolution of fuzzy measures is presented terminating in the Integrated Fuzzy and Relative methodology, which stems directly from previous proposals, integrating aspects from the Totally Fuzzy and Relative (TFR) and the Fuzzy Monetary Indicator (FMI). Brief adaptation to discrete variables is

discussed and a relative proposed integration of multiple dimensions is discussed by identifying the most important steps. To obtain micro-data of specific deprivation measures one must abstract the most relevant information from a set of available variables, which is computed through standard confirmatory and explanatory factor analysis.

On these theoretical themes, in the last five years, the Research Centre on Income Distribution “C. Dagum” (Cridire) of the University of Siena and its research group have been awarded of several research grants funded by International bodies such as the European Union, the World Bank, and the Eurostat. In particular, the project SAMPLE funded by the EU under the 7th FP (<http://www.sample-project.eu/>) has represented a splendid occasion for scientific co-operation and common research between authoritative international academic and governmental institutions and agencies; in particular the above mentioned Cridire centre and the Warsaw School of Economics (Wse) have shared and implemented estimation methods and statistical measures for the comparative analysis among Poland and Italy using Eu-Silc data. Multidimensional fuzzy poverty indexes of diffusion and of intensity have been established according to the well known traditional measures called Head-Count Ratio and Average Poverty Gap.

The comparison, conducted at regional level (italian „regioni” and polish „voivodship”), shows the implementation in the analysis introduced by the multidimensional approach.

Lewiński vel Iwański Stanisław, Wilimowska Zofia (Politechnika Wrocławska)

Struktura finansowa polskich przedsiębiorstw, a modelowanie statystyczne

Rozpoczęcie, a także kontynuowanie działalności gospodarczej nie jest możliwe bez wystarczających zasobów finansowych, pochodzących z różnych źródeł, zarówno własnych jak i obcych. Decyzja o wyborze odpowiedniej struktury finansowania może decydować o właściwym rozwoju organizacji. Z drugiej strony zła struktura z kolei może przesądzać o dużych problemach przedsiębiorstwa w przyszłości. Przy czym decyzja ta jest podejmowana w toku prowadzonej działalności, przy ścisłym związku z otoczeniem gospodarczym przedsiębiorstwa. W tym miejscu też pojawia się pytanie, czy można ustalić optymalną strukturę finansową dla danego (każdego) przedsiębiorstwa, która wpływałaby na wzrost jego wartości rynkowej.

Wielu ekonomistów (zarówno badaczy, teoretyków jak i praktyków) stara się odpowiedzieć na powyższe postawione pytanie. Pomimo przeprowadzonych wielu analiz kwestia kształtowania struktury finansowej stanowi nadal otwarty problem badawczy, a staje się szczególnie istotna w dobie (obecnego) kryzysu gospodarczego na świecie, jak i w Polsce.

Ze względu na powyższe zostały przeprowadzone analizy statystyczne dla próby danych finansowych 100 polskich spółek giełdowych, z lat 2001 - 2010. Badania dotyczyły struktury finansowej oraz wartości rynkowej przedsiębiorstw. W analizie brano pod uwagę zarówno rozkład badanych

cech (reprezentujących strukturę oraz wartość rynkową), a także wzajemną ich zależność, pod kątem istotności statystycznej. W pracy badawczej wykorzystano zaawansowane oprogramowanie komputerowe analizy statystycznej, symulacyjnej i niepewności o nazwie @RISK firmy Palisade. Dzięki zastosowanemu narzędziu można było podejść wielowymiarowo do zidentyfikowanych zależności, analizowanych pod kątem statystycznym oraz prognostycznym. W tej pracy szczegółowo zostaną przedstawiono metody analizy oraz jej wyniki, wraz z obszernym komentarzem. Metodologia prowadzonych badań oraz otrzymane rezultaty będą stanowiły podstawę do dalszej pracy nad nowym modelem zarządzania strukturą finansową, dedykowanym dla polskich przedsiębiorstw.

Słowa kluczowe: analiza statystyczna w finansach, struktura finansowa, modelowanie dynamiczne

Statistical modeling of the financial structure of Polish enterprises

Running and continuing a business is not possible without suitable resources of financing, which comes from different sources, both own and foreign. Decision of choosing the right financial structure is essential for company, it is a matter of fundamental importance to subsequent operation of organization. On the other hand bad structure can cause serious problems for company in future. Moreover, the very decision is reached, in the course of business activity, with close relation to context of economical environment of company. In this point appears a question, whether it is possible to create an optimal financial structure for certain (every) company, which has an influence on its market value.

Many economists (both researchers, theorist and practices) try to answer on above mentioned question. Despite the fact that there have been carried out analyses, the issue of managing the financial structure remains to be resolved, and becomes very essential in time of (current) world (and also polish) economical crisis.

Because of above mentioned issues, there has been performed statistical analysis for a group financial data of 100 polish joint stock companies, from years 2001 - 2010. Research was related with financial structure and company's market value. In study there was taken to concern a distribution of analyzed measures (representing financial structure and market value), and their common interdependence, in sense of statistical essentiality. In research there was used a sophisticated statistical, simulating and uncertainty analysis software @RISK (from Palisade company). Thanks to this tool there was possible to use a multidimensional approach to identified relations, which were analyzed with use of statistical and forecasting methods. In this paper methods of research and its results will be presented in detail, with comprehensive comment. Methodology of this study and its outcome will be basis for future work on new model of financial structure, dedicated for polish companies.

Key words: statistical analysis in finances, financial structure, dynamical modelling

BIBLIOGRAPHY

1. DeAngelo H., DeAngelo L. i Whited T. M., Capital structure dynamics and transitory debt, "Journal of Financial Economics", nr 99 ,2011,

2. Durand D., Cost of Debt and Equity Funds for Business. Trends and Problems of Measurement. Conference on Research in Business Finance, National Beureau of Economic Research, 1952.
3. Fischer E., Heinkel R., Zechner J., Dynamic capital structure choice: theory and tests, "Journal of Finance" nr 44, 1989,
4. Goldstein R., Ju N., Leland H., An EBIT - based model of dynamic capital structure, "Journal of Business" nr 74, 4, 2001,
5. Jerzemowska M., Kształtowanie struktury kapitału w spółkach akcyjnych, Warszawa, PWN, 1999,
6. Kryszicki W., Bartos J., Dyczka W., Królikowska K., Wasilewski M., Rachunek prawdopodobieństwa i statystyka matematyczna w zadaniach. Część II. Statystyka matematyczna, wydanie II poprawione, Warszawa, Wydawnictwo Naukowe PWN 1994,
7. Ostaszewski J., Źródła pozyskiwania kapitału przez spółkę akcyjną, Warszawa, Difin, 2000,
8. Maliński M., Weryfikacja hipotez statystycznych wspomagana komputerowo, Gliwice, Wydawnictwo Politechniki Śląskiej, 2004,
9. Modigliani F., Miller M., Koszt kapitału, finanse przedsiębiorstw i teoria inwestycji, "The American Economic Review" nr 48, 1958,
10. Pod redakcją Suszyńskiego C., Przedsiębiorstwo, Wartość, Zarządzanie, Warszawa, Polskie Wydawnictwo Ekonomiczne, 2007,
11. Shyam-Sunder L., Myers S.C., Testing static tradeoff against pecking order models of capital structure, "Journal of Financial Economics" nr 51, 1999.

Liberda Barbara (Uniwersytet Warszawski), Świczewska Iwona, Tomaszewicz Łucja (Uniwersytet Łódzki)

Metodologia rachunków satelitarnych na podstawie rachunku satelitarnego sportu dla Polski

Celem artykułu jest przedstawienie metodologii rachunków satelitarnych zastosowanej do stworzenia rachunku satelitarnego sportu dla Polski. Rachunek satelitarny sportu jest rozszerzeniem i uzupełnieniem Systemu Rachunków Narodowych, pozwalającym na identyfikację przepływu dóbr i usług związanych ze sportem. Rachunek satelitarny sportu dla Polski bazował głównie na tablicach wykorzystania i podaży, w podziale na 58 rodzajów działalności oraz 465 grup produktów, wśród których produkty sportowe zidentyfikowano w 43 grupach bilansowych. Przedstawiamy obliczenia udziału sportu w PKB oraz w zatrudnieniu.

Słowa kluczowe: rachunek satelitarny, sport, rachunki narodowe, tablice wykorzystania i podaży

Methodology of satellite accounts at the example of the sport satellite account for Poland

The aim of the paper is to present the methodology of satellite accounts at the example of sport satellite account for Poland. The sport satellite account acts as an extension and a supplement to the System of National Accounts and allows identifying the sport related flows in the economy. The calculation of sport satellite account for Poland was based mainly on the supply and use tables for the year 2006 for 58 types of activities and 465 product groups from which 43 product groups were significantly related to sport. We estimated the share of GDP generated by sport as well as the sport's share in overall employment.

Key words: satellite account, sport, national accounts, supply and use tables

BIBLIOGRAPHY

1. European Commission, 2007, White Paper on Sport, COM (2007) 391final.
2. Sport Satellite Account for Poland - Rachunek satelitarny sportu dla Polski, Institute of Official Statistics, Central Statistical Office, December 2010, Warsaw, Poland, pp. 82.
3. Sport Satellite Account for the UK, 2004, Sport Industry Research Centre, Sheffield Hallam University, 2010.

Liberda Barbara (Uniwersytet Warszawski)

Rachunki generacyjne metodą pomiaru bogactwa kolejnych pokoleń

Celem artykułu jest przedstawienie metodologii rachunków generacyjnych oraz ich zastosowań w polityce fiskalnej oraz w reformach systemów ubezpieczeń społecznych, zdrowia oraz edukacji. Rachunki generacyjne służą mierzeniu przyrostu majątku netto kolejnych pokoleń, w tym wielkości długu implicite. Dyskusji poddany zostanie zakres danych niezbędnych do stworzenia rachunków generacyjnych oraz metody obliczeń sumy podatków netto w skali życia poszczególnych pokoleń. Rachunki generacyjne stanowią również narzędzie symulacji warunków równowagi i stabilności systemów fiskalnych. Przedstawimy możliwości i ograniczenia takich symulacji.

Słowa kluczowe: rachunki generacyjne, pokolenie, majątek, podatki netto, dług implicite

Generational accounting - a method of measuring net wealth of each generation

The aim of this paper is to discuss the methodology of generational accounting and its application in fiscal policy and in reforms of social security system as well as health care and education system. Generational accounting is a method of measuring the net increase of wealth of each generation, including the implicit debt. We will discuss the data requirements of generational accounts and methods of calculating the sum of net taxes during the life span of a generation. Generational accounting is a tool to simulate the conditions of equilibrium and sustainability of fiscal systems. We will discuss the conditions of such simulations.

Key words: generational accounting, generation, wealth, net taxes, implicit debt

Łagodziński Władysław Wiesław (Polskie Towarzystwo Statystyczne), Krasucki Piotr (Instytut Spraw Publicznych)

Zintegrowany system danych statystycznych z ochrony zdrowia - niezbędne minimum

Zgodnie z Decyzją 1400/97 Parlamentu Europejskiego i Rady Europy, musimy zbierać, gromadzić, opracowywać, przechowywać i publikować dane na temat stanu zdrowia ludności Polski oraz funkcjonowania systemu ochrony zdrowia w naszym kraju. Teoretycznie wszystkie te dane powinny być zbierane i przetwarzane całościowo w jednym systemie informacyjnym przez instytucję ponad resortową, wyposażoną w kompetencje prawne i odpowiednie środki celowe, wolne od doraźnych, koniunkturalnych zabiegów oszczędnościowych. Musi to być instytucja niezależna ” od Ministerstwa Zdrowia. W polskim systemie społeczno-prawnym optymalnym rozwiązaniem było by powierzenie tych obowiązków Głównemu Urzędowi Statystycznemu. Przyjęcie takiego rozwiązania pozwoliłoby zatrzymać trwającą od kilkunastu lat dezintegrację systemu informacji w zakresie ochrony zdrowia oraz zapobiec hellenizacji statystycznych ocen stanu zdrowia i potrzeb zdrowotnych ludności Polski. W naszym wystąpieniu prezentujemy ramową propozycję dotyczącą zakresu i zasad zbierania niezbędnych informacji z uwzględnieniem podziału na różne grupy ryzyka zdrowotnego, związane z dostępem populacji do różnego poziomu ” warunków zdrowotnych” rekomendowanych przez Światową Organizację Zdrowia. Bardzo ważnym elementem będzie analiza relacji pomiędzy kosztami a efektami zdrowotnymi, powszechnie stosowanych procedur profilaktycznych, terapeutycznych i rehabilitacyjnych, a także ocena zatrważających opóźnień w przygotowywaniu takiego systemu informacyjnego. Tylko w oparciu o taki system i dane w nim zgromadzone możliwe będzie kreowanie prawidłowej polityki zdrowotnej. Koszty społeczne i zdrowotne takich błędów i zaniechań ilustruje po raz kolejny konflikt w obszarze zdrowia, jakim w grudniu 2011 r. i styczniu 2012 r. była afera lekowa.

Słowa kluczowe: informacja z ochrony zdrowia, system informacji, dane zintegrowane, niezbędne minimum informacji

Integrated system of health data and statistics - indispensable minimum

According to the Decision No 1400/97/EC of the European Parliament and of the Council we are obliged to gather, accumulate, elaborate, store and publish data on the health condition of Polish population and the functioning of public health system in Poland. Theoretically, all data should be gathered and proceeded integrally by an overdepartmental institution in the one information system. The institution should be provided with legal competence and appropriate funds, independent of ad hoc saving plans and decisions made by the Minister of Health. In Polish sociolegal system the optimal solution would be to consign those obligations to Central Statistical Office. In this way the continued disintegration of information system in the field of health care could be stopped and could avoid false assessment of health status and health needs of Polish population. In our paper we present frame proposal of the scope and principles of gathering indispensable information regarding division on certain health risk groups related to the access of population to different levels of "sanitary conditions" recommended by WHO. An important part of the presentation will be analysis of relation between costs and health effects of commonly applied preventive, therapeutic and remedial procedures. Another significant point in the paper will be evaluation of horrifying delays in preparation of such an information system. Without such a system and data gathered inside the creation of proper health policy would not be possible. The social and health costs of such mistakes and omissions illustrates once again the conflict in health care system, which occurred in December 2011 and January 2012 as a „medicine affair”.

Key words: information from health system, information system, integrated data, the minimum information necessary

Łazowska Bożena (Centralna Biblioteka Statystyczna)

Władysław Bortkiewicz (7 VIII 1868 - 15 II 1931)

Władysław Bortkiewicz urodził się w polskiej rodzinie w Petersburgu, gdzie ukończył studia na Wydziale Prawa w 1890 r. Był ekonomistą i statystykiem, który większość zawodowego życia spędził w Niemczech, gdzie wykładał na Uniwersytecie w Strasburgu (docent, 1895-1897) i na Uniwersytecie w Berlinie (profesor statystyki i ekonomii politycznej, 1901-1931). W swoim pokoleniu Bortkiewicz był najbardziej reprezentatywnym przedstawicielem statystyki matematycznej w Europie. Jego dzieła dotyczyły ekonomii politycznej, rachunku prawdopodobieństwa, demografii, studiów aktuarialnych, statystyki ekonomicznej i fizyki statystycznej. Publikował swoje prace głównie w języku niemieckim, a także w języku rosyjskim i po polsku. W 1898 r. opublikował pracę o rozkładzie Poissona Prawo małych liczb. W pracy tej jako pierwszy wykazał, że szereg zjawisk rzadkich ulega z wielką dokładnością prawu wyrażonemu wzorem Poissona. W ekonomii politycznej jest znany z analizy krytycznej schematu reprodukcji Karola Marksa zawartej w ostatnich dwóch tomach Kapitału. Jego największe prace to : *Die mittlere*

Lebensdauern (1893), *Das Gesetz der kleinen Zahlen* (1898), *Die Radioaktive Strahlung als Gegenstand wahrscheinlichkeitstheoretischer Untersuchung* (1913), *Die Iterationen* (1917).

Słowa kluczowe: statystyka, ekonomia polityczna, rachunek prawdopodobieństwa, demografia

Ladislaus Bortkiewicz (7 VIII 1868-15 II 1931)

Ladislaus Bortkiewicz was born in a Polish family in St. Petersburg, where he graduated from the Law Faculty in 1890. He was an economist and statistician, who spent most of his professional life in Germany, where he taught at Strasburg University (associate professor, 1895-1897) and Berlin University (professor in statistics and political economy, 1901-1931). In his generation Bortkiewicz was the main representative of mathematical statistics in Europe. His works range from political economy through probability, demography, actuarial studies, economic statistics to statistical physics. His published works are mainly in German, but also in Russian and in Polish. In 1898 he published a book about the Poisson distribution, titled *The Law of Small Numbers*. In this book he first noted that events with low frequency in a large population follow a Poisson distribution even when the probabilities of the events varied. In political economy Bortkiewicz is important for his analysis of Karl Marx's reproduction schema in the last two volumes of *Capital*. His important works are: *Die mittlere Lebensdauern* (1893), *Das Gesetz der kleinen Zahlen* (1898), *Die Radioaktive Strahlung als Gegenstand wahrscheinlichkeitstheoretischer Untersuchung* (1913), *Die Iterationen* (1917).

Key words: statistics, political economy, probability, demography

Łuczak Aleksandra, Wysocki Feliks (Uniwersytet Przyrodniczy w Poznaniu)

Ocena poziomu rozwoju społeczno-gospodarczego powiatów województwa wielkopolskiego

Celem pracy jest rozpoznanie poziomu rozwoju społeczno-gospodarczego powiatów ziemskich województwa wielkopolskiego. Do jego określenia wykorzystano cechę syntetyczną, która jest funkcją cech prostych – wyznaczników cząstkowych poziomu rozwoju społeczno-gospodarczego. W procedurze tworzenia cechy syntetycznej zastosowano do mierzenia odległości od wzorca i antywzorca rozwoju uogólnioną miarę odległości GDM (Walesiak 1993, 2006), a do agregacji cech metodę TOPSIS (Technique for Order Preference by Similarity to an Ideal Solution) (Wysocki 2010). Miara Walesiaka (1993) pozwala na obliczanie odległości pomiędzy cechami opisywanymi na różnych typach skal tj. porządkowych, przedziałowych i ilorazowych. Natomiast metoda TOPSIS jest metodą wzorcową służącą do agregacji cech. Zaproponowane podejście pozwala dokonać oceny poziomu rozwoju społeczno-gospodarczego powiatów opisywanych przez cechy proste, zarówno metryczne, jak i porządkowe.

W procesie oceny poziomu społeczno-gospodarczego powiatów województwa wielkopolskiego wykorzystano dane statystyczne z ankiety przeprowadzonej w starostwach powiatowych w województwie wielkopolskim nt. Stanu i możliwości rozwojowych powiatów województwa wielkopolskiego (2000) oraz Banku Danych Lokalnych Głównego Urzędu Statystycznego za rok 2010.

Słowa kluczowe: poziom rozwoju społeczno-gospodarczego, porządkowanie liniowe, cechy metryczne, cechy porządkowe, uogólniona miara odległości GDM, metoda TOPSIS

Evaluation of level of socioeconomic development of powiats in Wielkopolska province

The paper aims to evaluate the level of socioeconomic development of powiats (i.e. second level administrative units) in Wielkopolska province. With this purpose in mind a synthetic feature was utilized that aggregates simple features: partial indicators of socioeconomic development. The procedure of designing this feature used the generalized distance measure GDM to compute the distance from the ideal and anti-ideal points (Walesiak 1993, 2006), and TOPSIS method to aggregate the features (Wysocki 2010).

The measure of Walesiak (1993) facilitates calculating distances between features described on different scale types: ordinal, interval, or ratio, while TOPSIS is the standard method for aggregation of features. Proposed approach allows evaluating the level of socioeconomic development of powiats when expressed through simple features, both metric and ordinal.

The evaluation process made use of the data gathered from the survey conducted in the powiat offices of the Wielkopolska province: "The present state and development potential of the powiats of Wielkopolska" (2000) and from the Local Data Bank of the Central Statistical Office report of 2010.

Key words: level of socioeconomic development, linear ordering, metric features, ordinal features, generalized distance measure (GDM), Technique for Order Preference by Similarity to an Ideal Solution (TOPSIS)

BIBLIOGRAPHY

1. Stan i możliwości rozwojowe powiatów województwa wielkopolskiego (2000). Ankieta przeprowadzona w starostwach powiatowych w województwie wielkopolskim.
2. Walesiak M. (1993): Statystyczna analiza wielowymiarowa w badaniach marketingowych. Prace Naukowe AE we Wrocławiu nr 654. Seria: Monografie i opracowania nr 101, Wrocław.
3. Walesiak M. (2006): Uogólniona miara odległości w statystycznej analizie wielowymiarowej. Wyd. Akademii Ekonomicznej we Wrocławiu, Wrocław.
4. Wysocki F. (2010): Metody taksonomiczne w rozpoznawaniu typów ekonomicznych rolnictwa i obszarów wiejskich. Wyd. UP w Poznaniu, Poznań 2010.

Malaga Krzysztof (Uniwersytet Ekonomiczny w Poznaniu)

Dylematy współczesnej statystyki wzrostu gospodarczego

Celem referatu jest wyodrębnienie i przedstawienie zasadniczych dylematów statystyki wzrostu gospodarczego, pozostających w ścisłym związku ze współczesną teorią wzrostu gospodarczego.

W referacie omówione zostaną następujące zagadnienia: alternatywne wobec PKB miary wzrostu gospodarczego, metodologiczne problemy wyboru miar i statystycznego pomiaru długookresowego lub trwałego wzrostu gospodarczego, problem częstotliwości pomiaru wzrostu gospodarczego w kontekście dychotomii pomiędzy realną i nominalną sferą gospodarki, rola statystyki w identyfikacji stylizowanych faktów wzrostu gospodarczego, wybrane problemy metodologiczne statystycznej rejestracji PKB i alternatywnych miar wzrostu gospodarczego, (nie)adekwatność statystycznej rejestracji parametrów i zmiennych występujących we współczesnych modelach egzogenicznego i endogenicznego wzrostu gospodarczego, a także (nie)zgodność metodologii statystyki wzrostu gospodarczego ze współczesną teorią wzrostu gospodarczego.

We wnioskach końcowych zostaną sformułowane postulaty, których spełnienie powinno sprzyjać dostosowaniu zasad statystyki wzrostu gospodarczego do reguł konstrukcji współczesnych modeli egzogenicznego i endogenicznego wzrostu gospodarczego.

Słowa kluczowe: wzrost gospodarczy, statystyka wzrostu, metodologia statystyki wzrostu, miary wzrostu długookresowego i trwałego, stylizowane fakty wzrostu gospodarczego, współczesna teoria egzogenicznego i endogenicznego wzrostu gospodarczego

The dilemmas of contemporary economic growth statistics

The main goal of this paper is to identify and to present the principal dilemmas of economic growth statistics, which remain to be in an exact relationship with contemporary economic growth theory.

The following issues will be discussed in the paper: alternative (to GDP) measurements of economic growth; methodological problems with choosing the proper measures and statistical measurement of long-run or sustained economic growth; the question of frequency of measuring the economic growth in the context of dichotomy between the real and the nominal spheres of economy; the role of statistics in identification of stylized facts of economic growth; chosen methodological problems of both statistical GDP registration and alternative measures of economic growth; (in)adequacy of statistical registration of parameters and variables found in contemporary exogenous and endogenous growth models; (in)compatibility of methodology used in statistics of economic growth with the contemporary growth theories.

In the summary adequate postulates will be risen, the fulfilment of which should support adjusting the principles of economic growth statistics to the rules of constructing the contemporary endogenous and exogenous economic growth models.

Key words: economic growth, economic growth statistics, methodology of economic growth statistics, measurements of long-run and sustainable growth, stylized facts of economic growth, contemporary exogenous and endogenous theories of economic growth

Źródła bibliograficzne: Eurostat, OECD, GUS, statystyczne bazy danych makroekonomicznych, ważniejsze pozycje literaturowe z zakresu współczesnej teorii wzrostu gospodarczego (artykuły, monografie).

Małecka Marta (Uniwersytet Łódzki)

O cena wariancji estymatorów wartości narażonej na ryzyko (VaR)

Wartość narażona na ryzyko (VaR), będąc jedną z najbardziej popularnych miar ryzyka na rynku kapitałowym, jest jednocześnie miarą poddawaną poważniej krytyce w literaturze przedmiotu, zwłaszcza w świetle kryzysu finansowego 2008/2009. Problemem podnoszonym przez badaczy rynku finansowego jest istnienie wielu elementów uznaniowych związanych z kalkulacją tej miary, w tym wybór rzędu kwantyla rozkładu straty z portfela oraz wybór metody estymacji. Podczas gdy zagadnienie rzędu kwantyla jest rozstrzygnięte w świetle zaleceń instytucji nadzoru bankowego, wzorujących się na opracowaniach Komitetu Bazylejskiego, nie istnieją jednoznaczne rekomendacje dotyczące metody estymacji VaR. Skutkuje to brakiem porównywalności ocen VaR raportowanych przez różne instytucje.

Wybór metody oceny estymatorów VaR jest kwestią budzącą kolejne kontrowersje. Ze względu na brak obserwowalności realizacji VaR, nie istnieje bezpośrednia możliwość oceny średniego błędu estymatora. Stosowanie analitycznych wzorów pozwalających na wyznaczenie wariancji estymatora jest sprzeczne z założeniem o braku stacjonarności zmiennych obserwowanych na rynku finansowym. Popularnie stosowana metoda porównywania udziału przekroczeń VaR z przyjętym poziomem tolerancji prowadzi do sprzecznych między sobą wyników, ze względu na fakt, że zależy ona od wyboru okresu w historii szeregu czasowego. Skuteczną techniką oszacowania błędu estymatorów VaR wydaje się zatem wykorzystanie metod symulacyjnych.

Celem prezentowanego badania była ocena wariancji związanej z estymacją VaR za pomocą eksperymentu symulacyjnego. Przy założeniu, że stopy zwrotu na rynku finansowym są realizacjami procesu GARCH, przeprowadzono analizę porównawczą obciążenia i wariancji szacunków VaR uzyskiwanych za pomocą metod kwantylowych, metody wariancji kowariancji oraz metod szeregów czasowych. Wyniki pokazały wyraźną przewagę podejścia opartego na wariancji z próby nad techniką kwantylową. Dla różnych długości szeregów czasowych, obciążenie

estymatorów przy drugiej z metod było od kilku do kilkunastu razy większe. Nieco mniejsze różnice zaobserwowano przy porównaniu wariancji estymatorów. Otrzymane rezultaty wskazały ponadto na zacieranie się różnic między alternatywnymi podejściami do szacowania VaR przy wydłużaniu szeregów czasowych.

Wyniki badania pokazały zależność jakości otrzymanych estymatorów od długości szeregu czasowego stanowiącego podstawę estymacji. Zaobserwowano systematyczny spadek obciążenia i wariancji estymatorów przy skracaniu szeregu od 1000 do 100 obserwacji. W konsekwencji wysunięto hipotezę o wpływie braku stacjonarności procesu na pogarszanie jakości estymatorów przy dużej liczbie obserwacji.

Dla najdłuższych rozważanych szeregów czasowych, wynoszących 1000 obserwacji, przeprowadzono porównanie metod opartych na momentach oraz kwantylach z próby do metod szeregów czasowych, które wymagają dużej liczby obserwacji. Mniejsze błędy estymacji uzyskano przy ocenie VaR z zastosowaniem tego ostatniego podejścia – pokazana została kilkukrotna przewaga modeli szeregów czasowych w stosunku do metod bazujących na szacowaniu wariancji czy kwantyla z próby, zarówno w zakresie obciążenia jak i wariancji estymatorów.

Badanie pokazało wpływ wyboru rzędu kwantyla na jakość otrzymanych oszacowań. Przy założeniu 1% poziomu tolerancji otrzymano wielokrotnie większą wariancję estymatorów VaR niż dla poziomu tolerancji 5%. Ze względu na bardziej regularne zachowanie szeregów stóp zwrotu na poziomie kwantyla rzędu 5%, w stosunku do kwantyla 1%, dla kwantyla wyższego rzędu pokazano brak dużych różnic w oszacowanych wartościach wariancji estymatorów. Wniosek ten wyciągnięto zarówno z porównania jakości estymatorów pod kątem metody estymacji jak i długości szeregu czasowego.

Słowa kluczowe: VaR, wariancja estymatorów, symulacja

Evaluation of variance of Value-at-Risk estimators

Being one of the most popular risk measures in capital market, Value-at-Risk (VaR) has also been subject to serious criticism in the literature, boosted by the financial crisis of 2008/2009. The arbitrary factors connected with VaR calculation rise many controversies among academic researchers, with the quantile order and estimation method choice being most frequently cited problems. While the quantile order issue is regulated by the banking supervision, with the view on Basel Committee recommendations, there are no specific regulations referring to the estimation method, with the result being lack of comparability of VaR numbers reported by different institutions.

Further controversies around VaR are related to the estimators assessment. The mean estimation error can not be calculated directly in the light of the lack of possibility to observe realized VaR. Application of analytical formulas for estimation error is contradictory to the lack of stationarity of variables observed in financial market. The popular method of comparing the sample fraction of VaR exceptions to the assumed level of tolerance, being dependant on the period in time series history, has pointed to contrary conclusions in different studies. In result, the simulation techniques appear to be the effective way to assess error in VaR estimation.

The aim of the study was evaluation of variance of VaR estimators by means of the simulation experiment. We presented the comparative analysis of bias and variance connected with VaR estimation by sigma-based, quantile-based and time series methods, obtained under the assumption that returns observed in capital markets are generated by the GARCH process. The results showed a major advantage of the sigma-based approach over methods that directly estimate the sample quantile. With the use of the latter approach, for various time series lengths, the bias of the estimators was from several to more than ten times larger. Smaller differences between methods were observed in comparative analysis of the variance of the estimators. The study showed also that lengthening time series translated into smaller differences between alternative VaR estimation techniques.

The study indicated the dependence of the quality of estimators on the length of the series, on which the estimation process was based. A systematic decline in bias and variance of estimators was achieved by shortening the time series from 1000 to 100 observations. In consequence, we posed the hypothesis of the influence of lack of stationarity in returns process on low quality of estimators obtained for large number of observations.

For longest considered time series, we presented a comparative analysis of sigma- and quantile-based methods to time series models, which require large number of observations. Lower estimation errors were achieved with the use of the latter approach – several times smaller bias and variance of estimators was observed, indicating the superiority of time series models over methods based on a direct estimation of sample variance or quantile.

Further enquiry showed the dependence of the quality of estimators on the quantile order. Under the assumption of 1% tolerance level, the variance of estimators was several times higher than for 5% level. Higher regularity observed in times series at 5% level of significance, in comparison to 1% level, translated into smaller differences between examined methods of estimation in terms of estimation bias and variance. The differences obtained in relation to various series lengths were also smaller in case of 5% VaR estimation.

Keywords: VaR, variance of estimators, simulation

BIBLIOGRAPHY

1. Amendment of the Capital Accord to Incorporate Market Risk [1996], Bank for International Settlements, online: www.bis.org.
2. Bałamut T. [2002], Metody estymacji Value at Risk, „Materiały i studia NBP”, zeszyt 147, Warszawa, 1-107.
3. Best P. [2000], Wartość narażona na ryzyko, Dom wydawniczy ABC, Kraków.
4. Butler C. [2001], Tajniki Value at Risk, Praktyczny podręcznik zastosowań metody VaR, LIBER, Warszawa.
5. Cheng W., Hung J. [2010], Skewness and leptokurtosis in GARCH-typed VaR estimation of petroleum and metal asset returns, “Journal of Empirical Finance” 18, ELSEVIER, 160-173.

6. Doman M., Doman R. [2004], Ekonometryczne modelowanie dynamiki polskiego rynku finansowego, Wydawnictwo Akademii Ekonomicznej w Poznaniu, Poznań.
7. Fiszeder P. [2009], Modele klasy GARCH w empirycznych badaniach finansowych, Wydawnictwo naukowe uniwersytetu Mikołaja Kopernika, Toruń.
8. Ganczarek A. [2007], Analiza niezależności przekroczeń VaR na wybranym segmencie rynku energii, „Dynamiczne modele ekonometryczne”, Uniwersytet Mikołaja Kopernika w Toruniu, Toruń, 315-320.
9. Grabowska A. [2000], Metody kalkulacji wartości narażonej na ryzyko (VaR), „Bank i Kredyt”, 29-36.
10. Jajuga K. [1999], Miary ryzyka rynkowego – cz. I, „Rynek Terminowy” nr 4, 67-69.
11. Jajuga K. [2000], Miary ryzyka rynkowego – część trzecia, „Rynek Terminowy” nr 8, 112-117.
12. Jajuga K. [2000], Value at Risk, „Rynek Terminowy” nr 9, 18-20.
13. Jajuga K., Ronka-Chmielowiec W. [red.] [2003], Inwestycje finansowe i ubezpieczenia – tendencje światowe a polski rynek, Wydawnictwo Akademii Ekonomicznej im. Oskara Langego we Wrocławiu, Wrocław.
14. Jorion P. [1996], Risk2: Measuring the Risk in Value at Risk, „Financial Analysis Journal” 52,6, ABI/INFORM Global, 47-56.
15. Kuziak K. [2003], Koncepcja wartości zagrożonej VaR (Value at Risk), StatSoft Polska, 29-38.
16. Łach M., Weron A. [2000], Skuteczność wybranych metod obliczania VaR dla danych finansowych z polskiego rynku, „Rynek Terminowy” nr 9, 133-137.
17. Markowski K. [2000], O weryfikacji historycznej, „Rynek Terminowy” nr 9, 24-26.
18. Mentel G. [2011], Value at Risk w warunkach polskiego rynku kapitałowego, CeDuWu.pl Wydawnictwa fachowe, Warszawa.
19. Osińska M. [2006], Ekonometria finansowa, Polskie Wydawnictwo Ekonomiczne, Warszawa.
20. Osińska M., Fałdziński M. [2007], Modele GARCH i SV z zastosowaniem teorii wartości ekstremalnych, „Dynamiczne Modele Ekonometryczne, Uniwersytet Mikołaja Kopernika w Toruniu”, Toruń, 27-34.
21. Piontek K. [2002], Pomiar ryzyka metodą VaR a modele AR-GARCH ze składnikiem losowym o warunkowym rozkładzie z „grubymi ogonami”, „Rynek kapitałowy Skuteczne inwestowanie”, 467-483.
22. Pipień M. [2006], Wnioskowanie bayesowskie w ekonometrii finansowej, Wydawnictwo Akademii Ekonomicznej w Krakowie, Kraków.

23. Stawczyk A. [1999], Wprowadzenie do metodologii pomiaru ryzyka – Value at Risk, „Rynek Terminowy” nr 4, 132-137.

Markowska Małgorzata, Strahl Dorota (Uniwersytet Ekonomiczny we Wrocławiu)

O Możliwości oceny innowacyjności regionów szczebla NUTS 2 w świetle zasobów informacyjnych Eurostatu

Referat przedstawia dorobek metodologiczny Eurostatu dotyczący pomiaru innowacyjności regionów państw Unii Europejskiej na szczeblu NUTS 2. Omówiona zostanie ewolucja podejść do pomiaru innowacyjności a także zmiany w zasobach informacyjnych ilustrujących dynamikę i poziom charakterystyk innowacyjności.

Dorobek Eurostatu zostanie oceniony na tle podejść stosowanych do oceny innowacyjności w świetle badań prowadzonych w literaturze przedmiotu.

Oceniona zostanie możliwość prowadzenia statystycznych analiz porównawczych w ujęciu dynamicznym w świetle zasobów informacyjnych dostępnych w bazach danych Eurostatu.

Słowa kluczowe: innowacyjność regionów, bazy danych, Eurostat, pomiar innowacyjności, regiony szczebla NUTS 2

The possibility of NUTS 2 level regions innovation assessment in the perspective of Eurostat information resources

The paper presents Eurostat methodological output referring to innovation measurement of NUTS 2 level regions representing EU member states. It also discusses the evolution of approaches towards innovation measurement and changes in information resources illustrating the dynamics and level of innovation characteristics.

Eurostat output will be evaluated at the background of approaches applied for the assessment of innovation in the perspective of research conducted in professional literature.

The assessment will refer to the possibility of performing statistical comparative analyses in dynamic perspective based on information resources available in Eurostat data bases.

Key words: innovation in regions, data bases, Eurostat, innovation measurement, NUTS 2 level regions

Matuszewska-Janica Aleksandra (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie)

Statystyczna analiza nierówności w płacach kobiet i mężczyzn w Polsce na tle Unii Europejskiej w odniesieniu do wieku, wielkości przedsiębiorstwa i grupy zawodowej

Celem badania jest ocena stopnia zróżnicowania nierówności płacowych pomiędzy kobietami i mężczyznami przy uwzględnieniu takich czynników jak: wiek, wielkość przedsiębiorstwa i zaszeregowanie do grupy zawodowej, przy wykorzystaniu narzędzi statystyki opisowej i matematycznej. Analiza zostanie przeprowadzona na bazie danych pozyskanych przez Eurostat w ramach projektu *Structure of earnings survey* (SES) w latach 2002 i 2006. Przeprowadzona analiza zostanie wykorzystana do wstępnej zjawiska jakim jest segregacja wertykalna na rynku pracy w poszczególnych krajach Unii Europejskiej z uwzględnieniem takich czynników jak wiek zatrudnionych i wielkość przedsiębiorstwa.

Słowa kluczowe: nierówności płacowe między kobietami i mężczyznami, luka płacowa, Unia Europejska

Wages inequalities between men and women in the EU: differences analysis by age, size of the enterprise and occupation

"Equal pay for equal work" is one of the fundamental principles of the European Union. However, statistical surveys show that in EU countries the wages differences between men and women (so-called Gender Pay Gap - GPG) are 18% on average. Therefore, one of the priorities of the European Commission is to reduce these differences. According to press reports, the new European Union strategy on gender equality and the elimination of differences between the earnings of men and women, should be presented in the second half of 2010 and it should be one of EU main priorities.

The most popular measure that describing wages differences is Gender Pay Gap (GPG). It represents the difference between average gross hourly earnings of male paid employees and of female paid employees as a percentage of average gross hourly earnings of male paid employees.

Presented statistical analysis based on the Eurostat data collected within the framework of *Structure of Earnings Survey* (SES) in years 2002 and 2006 and it is provided for the EU member states. The aim of the research is to analyze the differences between women and men wages in different groups of employees. The groups of employees are defined by occupation, size of the enterprise and age. Provided analysis is the introductory review of gender pay gap similarities across European Union countries referring to vertical segregation.

Key words: wages inequalities between men and women, gender pay gap, European Union

BIBLIOGRAPHY

1. Gawrycka M., Wasilczuk J., Zwiech P.: *Szklany sufit i ruchome schody – kobiety na rynku pracy*, CeDeWu.pl, Warszawa 2010.
2. Lisowska E., *Równouprawnienie kobiet i mężczyzn w społeczeństwie*, Oficyna Wydawnicza SGH, Warszawa 2008.
3. Preston J.A.: *Occupational gender segregation Trends and explanations*, "The Quarterly Review of Economics and Finance", 1999, 39, s. 611–624.

Meller Ewa (Urząd Statystyczny w Poznaniu)

Estymacja dla małych obszarów dla wybranych charakterystyk rynku pracy

Wykorzystanie tradycyjnych estymatorów znanych z klasycznej metody reprezentacyjnej w badaniach prowadzonych przez urzędy statystyczne pozwala uzyskać wiarygodne szacunki dla domen odpowiednio reprezentowanych w próbie. W przypadku przekrojów, których liczebność jest niewystarczająca bądź w skrajnych przypadkach zerowa ich stosowanie obarczone jest bardzo dużym błędem. W takiej sytuacji oszacowanie wybranych charakterystyk jest możliwe z wykorzystaniem technik jakie oferuje statystyka małych obszarów. Jest to szczególnie ważne z gospodarczego punktu widzenia w odniesieniu do rynku pracy. Stale rosnące potrzeby informacyjne w tym zakresie, na niskich poziomach agregacji danych, wymuszają konieczność stosowania metod estymacji pośredniej. Dzięki temu możliwe jest oszacowanie wybranych charakterystyk rynku pracy z akceptowalną przez odbiorcę precyzją szacunków.

Głównym celem artykułu będzie przedstawienie możliwości wykorzystania wybranych estymatorów statystyki małych obszarów na potrzeby rynku pracy. W części teoretycznej omówiona zostanie idea statystyki małych obszarów wraz ze wskazaniem wykorzystywanych w pracy estymatorów, ich wadami i zaletami. Część empiryczna dotyczyć będzie zastosowania metod estymacji pośredniej dla wybranych charakterystyk rynku pracy na poziomach niepublikowanych przez Główny Urząd Statystyczny.

Słowa kluczowe: statystyka małych obszarów, estymacja pośrednia, rynek pracy

Wages inequalities between men and women in the EU: differences analysis by age, size of the enterprise and occupation

The use of traditional estimators derived from the representative method provides reliable estimates for domains adequately represented in the sample in surveys conducted by statistical offices. In the case of domains whose size is insufficient or in some extreme cases equal to zero, the use of such estimators is burdened with a high degree of error. In this situation one can estimate selected characteristics using techniques offered by small area estimation. It is

particularly important from the economic point of view in relation to the labour market. The growing demand for information in this field at low levels of data aggregation calls for the use of indirect methods of estimation, which produce estimates of selected characteristics of the labour market with an acceptable precision.

The main aim of this paper is to present the possibility of using some small area estimators that meet the needs of the labour market. The theoretical part describes the basic principles of small area estimation as well as advantages and disadvantages of selected estimators. The empirical part is a report on the application of indirect estimation methods to estimate some labour market characteristics at lower levels of aggregation than those published by the Central Statistical Office.

Key words: small area estimation, indirect estimation, labour market

BIBLIOGRAPHY

1. Rao J. N. K., *Small Area Estimation*, Wiley, New York 2003
2. Longford N.T., *Missing data and small-area*, Modern Analytical Equipment for the Survey Statistician, Springer, New York 2005
3. Gołata E., *Estymacja pośrednia bezrobocia na lokalnym rynku pracy*, Wydawnictwo Akademii Ekonomicznej w Poznaniu, Prace habilitacyjne nr 11, Poznań 2004
4. Gołata E., *Metody i źródła pozyskiwania informacji w statystyce publicznej*, Wyd. UE Poznań, Poznań 2009

Mielniczuk Jan (Instytut Podstaw Informatyki PAN, Politechnika Warszawska)

Wkład Polaków w rozwój statystyki matematycznej i aplikacyjnej

W referacie omówię wkład w rozwój statystyki matematycznej i aplikacyjnej Polaków pracujących na uniwersytetach, politechnikach i w instytutach PAN po drugiej wojnie światowej. Zachowany zostanie porządek chronologiczny; najpierw przedstawiona zostanie działalność statystyków ukształtowanych przed wojną, przede wszystkim Hugona Steinhausa, Juliana Perkała i Jana Oderfelda, a następnie ich uczniów i uczniów ich uczniów. Nie będę mówił o ogromnym wkładzie Jerzego Neymana, któremu będzie poświęcony odrębny referat.

Polish contribution to development of mathematical and applied statistics

Main achievements in mathematical and applied statistics of Polish researchers working at universities, technical universities and at Polish Academy of Sciences after the second world

war will be discussed. The subject will be treated in chronological order, starting from the research of statisticians active before the war, mainly Hugo Steinhaus, Julian Perkal and Jan Oderfeld and then proceeding to their students and their students' students. We will not discuss the fundamental work of Jerzy Neyman, which will be the subject of a separate lecture.

BIBLIOGRAPHY

1. R. Duda, A. Weron, Wrocławska szkoła matematyczna, Wiadomości Matematyczne XLII (2006), 73-101
2. B. Kopociński i J. Łukaszewicz, Stefan Zubrzycki (26.III.1927-18.XII.1968), Wiadomości Matematyczne XIII (1971), 57-65
3. Koronacki, Robert Bartoszyński (1933-1998), Wiadomości Matematyczne XXXV (1999), 167-179
4. J. Łukaszewicz, Julian Perkal (1913-1965), Wiadomości Matematyczne X (1967), 29-36
5. J. Łukaszewicz, Rola Hugona Steinhausa w rozwoju zastosowań matematyki, Wiadomości Matematyczne XVIII (1973), 51-63
6. H. Steinhaus, Wspomnienia i zapiski, Aneks Publishers, 1992
7. Strona Zakładu Statystyki Matematycznej IM PAN <http://www.impan.pl/Zaklady/stat.html>
8. S. Zubrzycki, O niektórych pracach seminarium z zastosowań matematyki, Zastosow. Mat. 8 (1966), 267-288

Mielniczuk Jan (Instytut Podstaw Informatyki PAN, Politechnika Warszawska),
Wojtyś Małgorzata (Politechnika Warszawska)

Postselekcyjna estymacja gęstości gradacyjnej

Rozważone zostanie zagadnienie estymacji gęstości gradacyjnej przy zastosowaniu metod selekcji modelu z listy wykładniczych rodzin rozkładów postaci:

$$\mathcal{M}_{\underline{i}} = \left\{ f_{\theta}(x) = \exp \left\{ \sum_{j \in \underline{i}} \theta_j b_j(x) - \psi(\theta) \right\} I(x \in [0, 1]) : \theta_j \in \mathbb{R} \text{ dla } j \in \underline{i} \right\},$$

dla wszystkich $\underline{i} \subset \{1, \dots, m_n\}$ i pewnego m_n zależnego od liczności n próby losowej. Układ $(b_j)_{j=1}^{\infty}$ jest tu systemem ortonormalnych funkcji w przestrzeni $L^2([0, 1], \lambda)$, gdzie λ to miara Lebesgue'a na $[0, 1]$.

Przedstawione zostaną wyniki mówiące o zgodności postselekcyjnych estymatorów gęstości gradacyjnej, w których nie zakłada się, że gęstość ta należy do jakiegokolwiek rodziny wykładniczej

o skończonym wymiarze. Twierdzenia te podają oszacowanie tempa zbieżności estymatorów w zależności od własności reguł wyboru modelu, na których estymatory te są oparte.

Zaprezentowane też zostaną przykłady zgodnych reguł wyboru modelu z listy wykładniczych rodzin rozkładów.

Słowa kluczowe: selekcja modelu, gęstość gradacyjna, wykładnicza rodzina rozkładów

Post-model-selection estimators of grade density

We consider the problem of grade density estimation with the use of model selection methods for choosing a model from a list of exponential families of densities of the form

$$\mathcal{M}_{\underline{i}} = \left\{ f_{\theta}(x) = \exp \left\{ \sum_{j \in \underline{i}} \theta_j b_j(x) - \psi(\theta) \right\} I(x \in [0, 1]) : \theta_j \in \mathbb{R} \text{ for } j \in \underline{i} \right\},$$

for all $\underline{i} \subset \{1, \dots, m_n\}$ and some m_n which depends on the sample size n . Here $(b_j)_{j=1}^{\infty}$ is a system of orthonormal functions in $L^2([0, 1], \lambda)$ where λ is the Lebesgue measure on $[0, 1]$.

The results on consistency of post-model-selection estimators of grade density of which we do not assume that it belongs to any exponential family of densities of finite dimension will be presented. The results give bounds for the rate of convergence of the estimators which depend on properties of model selection rule used to construct such estimators.

Moreover, some examples of consistent model selection criteria which use exponential families of densities will be considered.

Key words: model selection, grade density, exponential family of densities

BIBLIOGRAPHY

1. Barron, A. R., Sheu, C. H. (1991). Approximation of density functions by sequences of exponential families, *Annals of Statistics*, **19**, 1347–1369.
2. Pokarowski, P., Mielniczuk, J. (2011). P-values of likelihood ratio statistic for consistent testing and model selection, *in preparation*.
3. Wojtyś M. (2011) Post-model- selection method for density estimation. *Comm. Statist. Th. Meth.* **40**, 3082–3098.

Migdał-Najman Kamila (Uniwersytet Gdański)

Konstrukcja hybrydowej, samouczącej się sieci neuronowej typu SOM-GNG

Samouczące się sieci neuronowe znajdują liczne zastosowania w grupowaniu i klasyfikacji danych. Do sieci tego typu zalicza się sieć SOM (Self Organizing Map) zaproponowaną przez Kohonen'a [Kohonen 1995, 1997, 2001] i sieć GNG (Growing Neural Gas) zaproponowaną przez Fritzke'go [Fritzke 1994]. Sieć SOM nazywana jest w literaturze mapą samoorganizującą, siecią lub mapą Kohonen'a, samoorganizującym się odwzorowaniem, mapą cech (Self Organizing Feature Map – SOFM) a w okresie powstania znana była pod nazwą topologii preserving map (por. [Berthold, Hand 1999; Kohonen 1995, 1997, 2001; Fort, Pages 1995; Yin 2002; Deboeck, Kohonen 1998]). Zaproponowana i rozwinięta około 1982 roku przez fińskiego profesora Teuvo Kohonen'a, jest obecnie jedną z bardziej znanych nienadzorowanych modeli sztucznych sieci neuronowych. Jest stosowana głównie w problemach klasyfikacji, grupowania, redukcji wymiarowości, wyszukiwania anomalii i odchyłeń od wartości typowych, wizualizacji wielowymiarowych zbiorów danych i badania dynamiki zjawisk (por. [Papadimitriou i inni 2002; Fessant, Midenet 2002; Migdał Najman, Najman 2004]). Poza wieloma zaletami posiada także jedną zasadniczą wadę. Gdy zastosujemy ją do poszukiwania skupień w zbiorze jednostek wyznaczenie granic skupień odbywa się w sposób subiektywny przez wizualną ocenę macierzy ujednoliconych odległości (U-matrix). Kohonen i Deboeck proponowali rozwiązanie tego problemu wprowadzając dodatkowy etap procesu grupowania. Na tym etapie następowało grupowanie metodą k-średnich nauczonych neuronów sieci SOM. Rozwiązanie to należy jednak uznać za połowiczne. Z jednej strony został rozwiązany problem subiektywnego określenia granic skupień jednak z drugiej pojawia się szereg nowych problemów związanych z samą metodą k-średnich. Wymaga ona między innymi ustalenia liczby istniejących klas, która zwykle nie jest znana (lub może być ustalona ponownie w sposób subiektywny w oparciu o wizualizację macierzy U-matrix) a także zakłada istnienie jedynie sferycznych skupień.

Z rozważań teoretycznych a także badań symulacyjnych (por. [Migdał Najman 2009; Najman 2009, 2010]) wynika, że większy potencjał w grupowaniu danych wielowymiarowych prawdopodobnie ma sieć GNG. Mimo licznych zalet sieci tego typu, do których zaliczyć należy ich szybką zbieżność, oszczędną strukturę sieci, automatyczne rozpoznawanie liczby skupień separowalnych, sieć ta posiada także wady. Jedną z poważniejszych wad jest trudność, z jaką sieć GNG rozpoznaje skupienia słabo separowalne. Jeżeli odległość między najbliższymi obiektami z różnych skupień jest tego samego rzędu, co przeciętne odległości między obiektami w tych skupieniach sieć GNG ich nie rozdzieli.

Wydaje się, że połączenie zalet obu powyższych typów sieci może wywołać efekt synergii i znacząco poprawić jakość grupowania jednocześnie eliminując główne wady obu metod. Jeżeli na pierwszym stopniu zastosujemy sieć SOM uzyskamy pewną liczbę nowych jednostek (neuronów) odpowiadających za najważniejsze własności grupowanych jednostek. Uzyskamy znaczące uproszczenie badanego zjawiska ponieważ neuronów jest niewiele w stosunku do liczby

jednostek. Jeżeli na drugim stopniu dokonamy grupowania uzyskanych w sieci SOM neuronów w oparciu o sieć GNG rozwiązany zostanie problem subiektywizmu w ocenie granic i liczby skupień. Jednocześnie jest mało prawdopodobne, aby na granicy między skupieniami istniały neurony sieci SOM, sieć GNG powinna więc istniejące skupienia rozdzielić i to niezależnie od ich struktury przestrzennej. W ten sposób problem sieci GNG z rozdzieleniem skupień słabo separowalnych powinien stracić na znaczeniu. Celem prezentowanych badań jest weryfikacja hipotezy o wysokim potencjale hybrydowej, samouczącej się sieci neuronowej SOM-GNG w analizie skupień.

Słowa kluczowe: analiza skupień, nienadzorowane sieci neuronowe, Self Organizing Map, Growing Neural Gas

The design of the hybrid, self-learning neural network type SOM-GNG

The self-learning neural networks with variable structure are numerous applications in clustering and data classification. The neural network proposed by T. Kohonen SOM (Self Organizing Map) [Kohonen 1995, 1997, 2001] and B. Fritzke GNG (Growing Neural Gas) belong to this type of network [Fritzke 1994]. In this paper the hybrid SOM-GNG neural network are presented and discussed. The combination should give a synergy effect increasing the quality of the clustering. The connection should also minimize the disadvantages of both methods. The aim of this study is to verify the above hypothesis.

Key words: cluster analysis, unsupervised neural networks, self organizing map, growing neural gas

BIBLIOGRAPHY

1. Berthold M., Hand D.J., *Intelligent data analysis*, Springer-Verlag, Berlin Heidelberg, 1999.
2. Deboeck G., Kohonen T., *Visual explorations in finance with Self-Organizing Maps*, Springer-Verlag, London, 1998.
3. Fessant F., Midenet S., *Self-Organizing Map for data imputation and correction in surveys*, „Neural Computing&Applications”, Springer-Verlag, 2002, 10, 304.
4. Fort J.C., Pages G., *About the Kohonen algorithm: strong or weak self-organization?* „Neural Networks”, Pergamon, Elsevier Science, 1996, 9, 5, 773.
5. Fritzke B., *Growing cell structures - a self-organizing network for unsupervised and supervised learning*, „Neural Networks”, 1994, 7, 9, 1441-1460.
6. Kohonen T., *Self-Organizing Maps*, Springer-Verlag, Berlin Heidelberg, 1995, 1997, 2001.
7. Migdał Najman K., *Analiza porównawcza własności nienadzorowanych sieci neuronowych typu Self Organizing Map i Growing Neural Gas w analizie skupień*, Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu nr 47, Taksonomia 16, UE, Wrocław, 2009, 205-213.

8. Migdał Najman K., Najman K., *Diagnozowanie kondycji finansowej spółek notowanych na GPW w Warszawie w oparciu o sieć SOM*, Zeszyty Naukowe nr 389, Rynek Kapitałowy, Skuteczne inwestowanie, część I, Wydawnictwo Naukowe Uniwersytetu Szczecińskiego, Szczecin 2004, 507-519.
9. Najman K., *Zastosowanie nienadzorowanych sieci neuronowych typu Growing Neural Gas w analizie skupień*, Prace Naukowe UE we Wrocławiu nr 47, UE, Wrocław, 2009, 196 – 204.
10. Najman K., *Ocena wpływu parametrów sterujących procesem samouczenia się sieci GNG na ich zdolność do separowania skupień*, Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu nr 107, Taksonomia 17, UE, Wrocław, 2010, 296-304.
11. Papadimitriou S., Mavroudi S., Vladutu L., Pavlides G., Bezerianos A., *The supervised network Self-Organizing Map for classification of large data sets*, „Applied Intelligence”, Kluwer Academic Publishers, 2002, 16, 187.
12. Yin H., *Data visualization and manifold mapping using the ViSOM*, „Neural Network”, Pergamon, 2002, 15, 1005

Mikulec Artur (Uniwersytet Łódzki)

Empiryczna analiza efektywności kryterium Mojeny i Wisharta w analizie skupień

Kryteria Mojeny i Wisharta to metody wyboru optymalnego wyniku grupowania stosowane w przypadku metod hierarchicznych (np. aglomeracyjnych) analizy skupień. Dwa kryteria zaproponowane przez Mojenę w latach 70-tych XX w.: reguła górnego obszaru odrzucenia oraz reguła średniej ruchomej, oba bazują na analizie poziomów połączeń obiektów na wykresie drzewa, celem wyboru miejsca jego przycięcia, tj. wyboru optymalnego wyniku grupowania. Trzecie kryterium: sprawdzania drzewa opracowane przez Wisharta polega na ocenie losowości podziału obiektów na wykresie drzewa.

Celem referatu jest prezentacja wyników empirycznej analizy efektywności kryteriów Mojeny i Wisharta wyboru liczby skupień na tle innych stosowanych w tym zakresie kryteriów, zaproponowanych m.in. przez: Bakera i Huberta, Calińskiego i Harabasza, Daviesa i Bouldina, Hartigana, Huberta i Levine’a, Krzanowskiego i Lai. Analiza empiryczna przeprowadzona zostanie z wykorzystaniem programu ClustanGraphics 8 oraz wybranych pakietów środowiska R na podstawie sztucznie wygenerowanych zbiorów danych oraz zbiorów wybranych z *UCI Machine Learning Repository*.

Słowa kluczowe: reguła górnego obszaru odrzucenia, reguła średniej ruchomej, kryteria Mojeny, kryterium Wisharta (sprawdzania drzewa), ClustanGraphics

An empirical analysis of the effectiveness of Wishart and Mojena criteria in cluster analysis

Mojena and Wishart criteria are methods of selecting the optimal grouping result of hierarchical cluster analysis methods (e.g. agglomerative). Two criteria were proposed by Mojena in the 70's of the 20th century: the upper tail rule and moving average quality control rule, both are based on an analysis of the fusion levels of objects in the dendrogram and aim to determine the cut-off point of it, i.e. to choose the optimal clustering result. The third criterion: tree validation was created by Wishart and evaluates the randomness of the objects clustering in the dendrogram.

The purpose of this paper is to present the results of the empirical analysis of the effectiveness of Mojena and Wishart criteria for number of clusters selections, in comparison to other applicable criteria in this area, including those proposed by: Baker and Hubert, Calinski and Harabasz, Davies and Bouldin, Hartigan, Hubert and Levine, Krzanowski and Lai. The empirical analysis will be done in ClustanGraphics 8 Program and selected packages in R-environment for artificially generated data sets and chosen sets from *UCI Machine Learning Repository*.

Key words: upper tail rule, moving average quality control rule, Mojena criteria, Wishart criterion (tree validation), ClustanGraphics

BIBLIOGRAPHY

1. Balicki A., *Statystyczna analiza wielowymiarowa i jej zastosowania społeczno-ekonomiczne*, Wydawnictwo Uniwersytetu Gdańskiego, Gdańsk 2009.
2. Cormack R., *A review of classification*, „Journal of the Royal Statistical Society”, Series a 1971, vol. 134(3), p. 321-367.
3. Domański Cz., Pruska K., Wagner W., *Wnioskowanie statystyczne przy nieklasycznych założeniach*, Wydawnictwo Uniwersytetu Łódzkiego, Łódź 1998.
4. Gan G., Ma C., Wu J., *Data clustering: theory, algorithms, and applications*, SIAM, Philadelphia 2007.
5. Gatnar E., Walesiak M. (red.), *Statystyczna analiza danych z wykorzystaniem programu R*, Wydawnictwo PWN, Warszawa 2009.
6. Gordon A.D., *Hierarchical classification*, [w:] Arabie P., Hubert L.J., De Soete G. (editors), *Clustering and Classification*, World Scientific, Singapore 1996.
7. Kaufman L., Rousseeuw P.J., *Finding groups in data. An introduction to cluster analysis*, (reprint 2005), John Wiley & Sons, New York 2005.
8. Milligan G.W., Cooper M.C., *An examination of procedures for determining the number of clusters in a data set*, „Psychometrika” 1985, vol. 50(2), p. 159-179.
9. Mojena R., *Hierarchical grouping methods and stopping rules: an evaluation*, „Computer Journal” 1977, vol. 20(4), p. 359-363.

10. Wishart D., *Clustangraphics primer: a guide to cluster analysis*, (4th edition), Edinburgh 2006.

Młodak Andrzej (Urząd Statystyczny w Poznaniu, Państwowa Wyższa Szkoła Zawodowa w Kaliszu)

Modernizacja statystyki działalności gospodarczej w wymiarze paneuropejskim

W referacie zaprezentowane zostaną podstawowe cele i założenia realizacyjne projektu modernizacji metodologii prowadzenia badań w zakresie statystyki gospodarczej i oceny jakości ich wyników realizowanego w ramach programu grantowego Europejskiego Systemu Statystycznego „Metodologia nowoczesnej produkcji statystyki przedsiębiorstw” (ESSnet MeMoBuSt). Zadaaniem tego przewidzianego na lata 2011 – 2013 przedsięwzięcia ośmiu krajów: Holandii, Węgier, Polski, Szwecji, Norwegii, Szwajcarii, Włoch i Grecji jest opracowanie nowej wersji podręcznika metodologicznego z tego zakresu uwzględniającej najnowsze osiągnięcia naukowe w tym zakresie oraz wypracowane kierunki ich efektywnego wykorzystania. Omówiona będzie zawartość dotychczasowego podręcznika napisanego jeszcze w 1997 r. – ze wskazaniem na ułomności i niedostatki zawartych tam informacji odniesieniu do współczesnych realiów i wyzwań. Kolejny element wystąpienia stanowi charakterystyka założeń obecnego podręcznika. Obejmuje to także zakres zgodności jego treści z powszechnie stosowanym w statystyce europejskiej i światowej schematem organizacji badań statystycznych GSBPM (Ogólny Model Przeprowadzania Badań Statystycznych), przyjętym przez ONZ, OECD i Eurostat. Następnie w prezentacji znajdzie się opis aktualnych ustaleń i rozbieżności odnośnie szczegółowej wizji zawartości poszczególnych rozdziałów oraz omówienie tych fragmentów podręcznika, które zostały już opracowane. Całość zwięźliwie analiza roli strony polskiej w tym przedsięwzięciu oraz wkład, który ona doń wniosła i jeszcze może wnieść.

Słowa kluczowe: Metodologia badań statystycznych, badania działalności gospodarczej, model GSBPM, metody analizy danych

Modernization of business statistics in the paneuropean context

In the paper basic purposes and executive assumptions of the project leading to the modernization of methodology of statistical surveys related to economic statistical and an assessment of quality of their results realized within the framework of the European Statistical System „Methodology for Modern Business Statistics” programme (ESSnet MeMoBuSt) will be presented. This venture of eight countries: the Netherlands, Hungary, Poland, Sweden, Switzerland, Norway, Italy and Greece planned to be conducted in years 2011 – 2013 is aimed at an elaboration of a new version of the methodological handbook concerning these issues taking the newest scientific achievements in this field and worked out directions of their effective application into

account. The content of the current version of the handbook written even in 1997 will be described, focusing on disadvantages and gaps existing in the information contained there in relation to modern realities and challenges. Next component of this presentation is a characteristics of assumptions of the current version of the handbook. It includes also a degree of its consistency with commonly used in European and world statistics scheme of organization of statistical surveys GSBPM (Generic Statistical Business Process), adopted by UN, OECD and Eurostat. Next, in the presentation a description of current statements and divergences concerning detailed vision of the content of particular chapters and review of these fragments which were already written is planned. The wholeness of the paper is supplemented with an analysis of the role of the Polish part in this venture and the issues which our country has contributed to this project and which it can contribute in the future.

Key words: Methodology of statistical surveys, surveys of economic activity, GSBPM model, method of statistical data analysis

BIBLIOGRAPHY

1. Willeboordse A. (ed.) (1997), *Handbook on the Design and Implementation of Business Surveys*, Eurostat, Lux-embourg.

Modranka Emilia, Suchecka Jadwiga (Uniwersytet Łódzki)

Statystyka przestrzenna – metody i zastosowania

Zasadniczym celem referatu jest przedstawienie najważniejszych problemów stojących do rozwiązania przed statystyką przestrzenną i wskazanie jej roli w rozwoju metodologii badań statystycznych.

Znaczący wzrost zainteresowania danymi zlokalizowanymi (przestrzennymi i przestrzenno-czasowymi) w badaniach statystycznych obserwujemy dopiero od połowy lat 70. XX wieku. Był to istotny impuls w rozwoju metod ilościowych umożliwiających ich prawidłową analizę. Szczególnie dotyczy to metod, które umożliwiają opis struktur przestrzennych oraz przestrzenną predykcję. Nauką dostarczającą odpowiednich metod takiej analizy w ogólnym znaczeniu jest statystyka przestrzenna.

Dane przestrzenne i przestrzenno-czasowe mogą być danymi geograficznymi lub danymi zlokalizowanymi charakteryzującymi się wielością informacji. Wymiar przestrzenny danych zlokalizowanych wprowadza zależności dwuwymiarowe, bardziej złożone niż jednowymiarowe zależności czasowe. Cechą charakterystyczną zlokalizowanych obserwacji przestrzennych jest to, że są one współzależne i heterogeniczne. W takich przypadkach zastosowanie klasycznych miar statystycznych powoduje utratę istotnej części informacji, a klasyczne estymatory podstawowych parametrów struktury tracą niektóre własności. Stąd też statystyka przestrzenna proponuje odpowiednie estymatory charakteryzujące się dobrymi własnościami.

Zakres statystyki przestrzennej zmienia się wraz z możliwościami przetwarzania dużych baz danych, rozwojem specjalistycznego oprogramowania i zapotrzebowaniem praktyki gospodarczej na różnego rodzaju analizy stanowiące podstawę procesu decyzyjnego. Najczęściej związany jest z poszukiwaniem odpowiedzi na pytania dotyczące:

- charakteru analizowanego zjawiska, a więc określenia czy w przestrzeni ma ono charakter regularny, losowy czy klastrowy;
- wykazywania prawidłowości przestrzennej przez określone zdarzenia;
- skupienia lub rozproszenia podobnych wartości badanej zmiennej;
- stopnia skorelowania określonej zmiennej z innymi zmiennymi w danym obszarze;
- niezależności różnych wartości badanej zmiennej;
- przyciągania się lub odpychania się podobnych wartości;
- zmian w przestrzennym rozkładzie badanej zmiennej w czasie,
- znajomości funkcji kowariancji lub wariogramu, itd.

W ramach współczesnej statystyki przestrzennej wyróżnia się trzy podstawowe grupy metod:

- metody geostatystyczne (opisu statystycznych danych dotyczących powierzchni Ziemi);
- metody analizy danych punktowych bez atrybutów (ustalania sposobu przestrzennego rozmieszczenia zjawisk i zależności przestrzennych);
- metody analizy danych obszarowych i punktowych atrybutowych (umożliwiające analizę jednostek administracyjnych i obiektów przyrodniczych, które mają granice i są zorientowane według współrzędnych geograficznych).

W referacie omówione zostaną podstawowe zagadnienia i zastosowania nowoczesnych metod statystyki przestrzennej. W szczególności prezentowane będą metody opisu przestrzennej struktury zjawisk, estymacji parametrów i interpolacji danych przestrzennych oraz metody analizy danych punktowych.

Słowa kluczowe: dane przestrzenne, metody geostatystyczne, analiza danych punktowych i obszarowych

Spatial statistics – methods and applications

The main aim of the paper is to present the most important problems facing the solution before the spatial statistics and identify its role in the development of statistical research methodology.

The significant growth of interest in localized data (spatial and spatio-temporal) in economic and social analyses is observing only since the mid-'70s the twentieth century. It was an important incentive in the development of quantitative methods to enable the proper analysis of

spatial data. This is particularly true of methods that allow the description of spatial structures and spatial prediction. In a general sense spatial statistics is a branch of knowledge about these methods.

Spatial and spatio-temporal data can be geographical or localized data characterized by a multiplicity of information. The spatial dimension of data introduces the two-dimensional dependencies and interactions, more complex than the one-dimensional time-dependence. A characteristic feature of localized spatial observation is that they are interdependent and heterogeneous. In such cases, the use of classical statistical measures will lose a substantial part of the information, and the estimators of basic structural parameters are losing some properties. Hence, spatial statistics propose appropriate, good estimators.

Range of spatial statistics varies with the processing capabilities of large databases, specialized software development and business practice needs for different types of analysis underlying the decision making process. This range is frequently associated with seeking answers to questions about:

- analyzed the nature of the phenomenon, and thus determine whether it has a regular, random or clustered character in the space;
- demonstrating spatial regularity by specific events;
- concentration or dispersion of similar values of the test variable;
- the variable degree of correlation with other variables in the area;
- the independence of the different values of the test variable;
- attraction or repulsion of similar values;
- changes in the spatial distribution of the test variable at a time
- knowledge of the covariance function or variogram, etc.

Classification of issues facing the spatial statistics are caused by different types of geographic/localized data: surface data, point patterns data (with or without the quantitative and qualitative attributes), lattice, regional, and others data).

As a part of modern spatial statistics, there are three basic groups of methods:

- geostatistical methods (description of the statistical data on the surface);
- methods for analysis of point data without attributes (determining the spatial distribution of phenomena and spatial relationships);
- methods for analysis of area and point data with attributes (for an analysis of administrative units and natural objects that have boundaries and are oriented by latitude and longitude).

In this paper the basic issues and the application of modern methods of spatial statistics will be discussed. In particular, the methods describe the spatial structure of phenomena, estimation of parameters and interpolation of spatial data and point data analysis methods will be presented.

Key words: spatial data, geostatistical methods, point patterns and lattice data analysis

BIBLIOGRAPHY

1. Driesbeke J.J., M. Lejeune, G. Saporta (red), (2006), *Analyse statistique des données spatiales*, Ed. Technip, Paris
2. Arbia G., (1989), *Spatial Data Configuration in Statistical Analysis of Regional Economic and Related Problems*, Kluwer Academic Publishers, Dordrecht.
3. Suchecka J., Antczak E., (2010), *Elementy geostatystyki i metody analiz przestrzennych danych punktowych*, [w]: Bogdan Suchecki (red.), *Ekonometria Przestrzenna. Metody i modele analizy danych przestrzennych*, Wydawnictwo C. H. Beck, Warszawa, s. 70-99
4. Wingle W.L., E.P. Poeter, (1999), *Geostatistical Analysis and Training via the Internet*, APCOM'99: Computer Applications in the Minerals Industries 28th International Symposium, October 20–22, 1999, Colorado School of Mines, Golden, Colorado

Motoryn Ruslan (Ukraiński Państwowy Uniwersytet Finansów i Handlu Międzynarodowego w Kijowie, Politechnika Lubelska), **Motoryna Iryna** (Akademia Zarządzania Finansami w Kijowie), **Motoryna Tetiana** (Kijowski Narodowy Uniwersytet im. Tarasa Szewczenko)

Porównanie wskaźników wykorzystania oszczędności gospodarstw domowych Ukrainy i Polski

Każdy kraj ma właściwości stosunkowo obliczenia wskaźników wykorzystania oszczędności gospodarstw domowych. Dlatego istnieje problem zestawienia tych wskaźników. W artykule przedstawiono możliwości porównania wskaźników wykorzystania oszczędności gospodarstw domowych Ukrainy i Polski. Dlatego rozpatrzone czynniki, które wpływają na procesy formowania i wykorzystania oszczędności gospodarstw domowych oraz odpowiednie analityczne możliwości rachunków narodowych kształtują czynniki, które formują objętości i kierunki wykorzystania oszczędności gospodarstw domowych.

Słowa kluczowe: porównanie wskaźników, oszczędności gospodarstw domowych, Ukraina i Polska, czynniki oszczędności, rachunki narodowy

C Comparison indexes of the use of gross savings of households Ukraine and Poland

Every country has peculiarities regarding the calculation of indexes of households savings use. Therefore there is a problem of comparison of these indexes. Possibility of comparison indexes of the use of gross savings of households Ukraine and Poland are exposed in the articles. Factors, which influence on the processes of forming and use of gross savings of households and analytical possibilities of national accounts in relation to the estimation of factors use of gross savings of households are considered for this purpose.

Key words: comparison indexes, savings of households, Ukraine and Poland, factors use of savings, national accounts

Najman Krzysztof (Uniwersytet Gdański)

S Sieci samouczące się w analizie skupień

W grupowaniu danych społeczno-ekonomicznych badacz często staje się przed problemem braku wzorca grupowania. Gdy nie istnieje odpowiednio liczna grupa jednostek, dla których znana jest ich przynależność do skupień, badacz nie posiada wiedzy o tym czy jednostki posiadają strukturę grupową, jeżeli tak to ile jest skupień, jaki jest ich charakter i konfiguracja przestrzenna. Nie może skonstruować modelu, który na podstawie wartości cech opisujących badane jednostki i wiedzy o ich strukturze grupowej przyporządkowywałby poszczególne jednostki do istniejących skupień. W sytuacjach takich stosuje się w analizie skupień bezwzorcowe metody grupowania, do których zalicza się samouczące się sieci neuronowe.

Do samouczących się sieci neuronowych zalicza się sieci typu SOM (Self Organizing Map) [Kohonen 1995], sieci o zmiennej strukturze GSOM (Growing Self Organizing Map), sieci typu gaz neuronowy NG (Neural Gas) [Martinetz and Schulten 1991, 1994] oraz sieci typu gaz neuronowy o zmiennej strukturze GNG (Growing Neural Gas) [Fritzke 1994]. Wszystkie powyższe typy sieci mogą być stosowane w analizie skupień. Celem prezentowanych badań jest dokonanie analizy porównawczej powyższych czterech modeli samouczących się sieci neuronowych z punktu widzenia ich przydatności w analizie skupień. Zaprezentowana zostanie ich konstrukcja, wyjaśnione zostaną parametry sterujące algorytmem samouczenia się. Wskazane zostaną ich zalety, wady a także pewne szczególne własności. Rozważania zostaną oparte na ocenie teoretycznych podstaw konstrukcji wymienionych sieci a także na podstawie rezultatów badań empirycznych.

Słowa kluczowe: analiza skupień, nienadzorowane sieci neuronowe, samouczące się sieci neuronowe

BIBLIOGRAPHY

1. Fritzke B., *Growing cell structures - a self-organizing network for unsupervised and supervised learning*, Neural Networks, 1994, 7, 9, 1441-1460.
2. Kohonen T., *Self-Organizing Maps*, Springer-Verlag, Berlin Heidelberg, 1995, 1997, 2001.
3. Martinetz T. M., Schulten K. J., *A 'neural-gas' network learns topologies*, Artificial Neural Networks, ed. T. Kohonen, K. Makisara, O. Simula, J. Kangas, 1991, 397-402
4. Martinetz T. M., Schulten K. J., *Topology representing networks*, Neural Networks, 1994, 7, 507-22

Nowak Lucyna (Główny Urząd Statystyczny)

Rozwój badań statystycznych w obszarze demografii i migracji

W referacie podjęta zostanie próba oceny zmian w metodologii, zakresie przedmiotowym i podmiotowym badań, jakie miały miejsce w badaniach demograficznych oraz badaniach migracji w okresie 1990–2010 w kontekście zakresu pozyskiwanych danych statystycznych oraz ich użyteczności dla oceny sytuacji demograficznej oraz perspektyw rozwoju demograficznego.

Słowa kluczowe: bilans ludności, ruch naturalny, prognozy demograficzne, strumienie migracji, zasoby migracyjne, statystyka urodzeń, statystyka zgonów, statystyka małżeństw

Nowakowski Rafał (Urząd Statystyczny we Wrocławiu)

Wybory samorządowe na przykładzie stolic regionów dolnośląskiego, małopolskiego i wielkopolskiego - 20 lat doświadczeń.

W 2010 r. minęła 20 rocznica odbudowy samorządu gminnego w Polsce. Restytucja samorządności w Polsce z 1990 r. traktowana jest jako jedna z najważniejszych reform systemowych i administracyjnych III RP. Od tego czasu samorząd gminny potwierdził swoją istotną i doniosłą rolę w rozwoju kraju. W tym procesie szczególna rola przypadła największym ośrodkom miejskim. Poza dotychczasowymi kompetencjami i zadaniami skumulowały one w toku rozwoju nowe prerogatywy wynikające z ich podwójnego umocowania: jako miast na prawach powiatu i stolic nowych jednostek wojewódzkich.

Analiza funkcjonowania i rozwoju 16 stolic wojewódzkich pozwala wyróżnić miasta, które wiodą prym na tle innych ośrodków. Należą do nich Kraków, Poznań i Wrocław, które urosły do symboli odrodzonego samorządu gminnego (naturalnie pomija się tu Warszawę ze względu na jej funkcje stołeczne, przyjęte rozwiązania ustrojowe jak i na wielkość). Miasta te są nie tylko pozytywnym przykładem rozwoju lokalnego ale również centrami, które oddziałują na gminy z terenu własnych województw.

Wybór Krakowa, Poznania i Wrocławia nie jest przypadkowy. Przede wszystkim są to jednostki zaliczane do czołówki największych (terytorialnie i demograficznie) miast III RP. Należą również do najprężniej rozwijających się metropolii. Są także miastami, w których od 20 lat zaobserwować można interesujące zjawiska związane z wyborami samorządowymi.

Kraków jest miastem, który w okresie PRL-u powiększył się demograficznie i terytorialnie ponad 2,5 krotnie. Konsekwencją tego zjawiska jest występowanie istotnych różnic i podziałów socjopolitycznych, które przekładają się na kampanię wyborczą, wyniki wyborów samorządowych w stolicy Małopolski, a także stabilność struktur miejskich. Kraków jest więc mikroskalą podziałów politycznych, które odzwierciedlone są również na ogólnopolskiej scenie politycznej. Jest to niewątpliwie istotna cecha, która pozwala wyróżnić miasto w szeregu innych.

Poznań z kolei jest miastem, które od lat aktywnie wspiera i lobbuje na rzecz rozwoju samorządności. Podobnie jak Kraków miasto rozrosło się w ostatnim półwieczu dwukrotnie. Jednak w odróżnieniu od Krakowa nie spowodowało to jednak istotnych różnic społecznych i politycznych w funkcjonowaniu Poznania. Zarówno stolica Wielkopolski, jak i cały region, są bowiem jednostkami tworzonymi przez ludność o silnym poczuciu więzi z regionem i miastem, co przekłada się m.in. na specyfikę wyborów samorządowych.

Wrocław jest swoistym symbolem sukcesu odrodzonego samorządu. Pomimo braku tradycji, a zwłaszcza ciągłości samorządowej (w odróżnieniu Poznań posiadał ją od początku XIX w., a Kraków od lat 60. XIX w.), zdołał w przeciągu 20 lat stworzyć silne i stabilne struktury samorządowe. Zarówno miasto jak i region, złożone z konglomeratu narodowościowo-etnicznego, przejawia wyjątkową stabilność i prężność struktur samorządowych.

Powyższa charakterystyka wybranych miast przemawia niewątpliwie za ich wyborem do planowanej prezentacji. Również przebieg wyborów bezpośrednich na urząd prezydenta skłania do uwzględnienia tych miast w wystąpieniu. W Krakowie w latach 1990-2002 dochodziło do pięciokrotnej zmiany prezydenta podczas gdy we Wrocławiu i Poznaniu mamy do czynienia z jego stabilną i silną pozycją. Przy czym wyniki wyborów we Wrocławiu - w odróżnieniu od Krakowa i Poznania - kończą się już w i turze wyborów. W przypadku wyborów do rad miejskich wyjątkowym konsensusem i współpracą z prezydentem charakteryzuje się Rada Miejska Wrocławia. Najtrudniejsza sytuacja przedstawia się w Krakowie gdzie rady są silnie sfragmentaryzowane politycznie i ideowo, co wpływa na wyjątkowo trudną współpracę z prezydentem pozostającym bez zaplecza politycznego. z rozwiązaniem pośrednim mamy do czynienia w stolicy Wielkopolski.

Dla celów badawczych wybrano metodę komparatywną. Podyktowane to jest chęcią ukazania pomiędzy prezentowanymi miastami różnic i podobieństw, które są pochodną zarówno czynników społeczno-politycznych, ekonomicznych, kulturowych jak i historycznych.

Cezury czasowe wystąpienia obejmują lata 1990-2010 a więc okres od pierwszej do piątej kadencji rad gminnych, włącznie z wyborami samorządowymi z 2010 r.

Zakres tematyczny wystąpienia i referatu podzielony jest na dwie główne części:

- wybory do rad miejskich 1990-2010;
- wybory prezydentów 2002-2010;

W obu przypadkach analizie poddana zostanie:

- charakterystyka elekcji samorządowej;
- wyniki wyborów samorządowych;
- geografia wyborcza;
- relacje pomiędzy organami stanowiącymi a wykonawczymi (współpraca v. antagonizmy);
- specyfika i wyjątkowość omawianych miast w kontekście wyborów samorządowych;
- wybory samorządowe a inne doświadczenia (Łódź, Rzeszów, Szczecin).

Wystąpieniu towarzyszyć będzie prezentacja. Rozszerzona forma zamieszczona zostanie w przewidzianych periodykach pokonferencyjnych.

Słowa kluczowe: Elekcja, Samorząd, Rada Miasta, Prezydent, Frekwencja

The local elections on the example of regional capitals of Lower Silesia, Malopolska and Wielkopolska - 20 years of experience.

In 2010 passed the 20th anniversary of the reconstruction of the commune self-government in Poland. Restitution of self-government in Poland in 1990 is regarded as one of the most important administrative and systemic reform in the Republic of Poland. Since then commune self-government has confirmed its essential and significant role in the development of the country. The largest urban centers played the most important role in this process. Apart from their current competence and tasks, in the course of their development, they accumulated new prerogatives resulting from, they dual authorization: as cities on laws of the district and capital cities of new provincial units.

Analysis of the functioning and development of 16 provincial capital cities allows to single out cities which are in the lead relating to other centers. These include Cracow, Poznan, Wroclaw, which have become the symbols of revived commune self-government (of course Warsaw is being omitted because of its capital functions, accepted political solutions and its size). These cities are not only a positive example of the local development but also centers which influence communes from the area of their own provinces.

The choice of Cracow, Poznan and Wroclaw is not accidental. First of all they are the biggest (territorially and demographically) cities of the Republic of Poland. They also belong to

the most rapidly developing metropolis. For 20 years, the cities have been observing interesting phenomena associated with local elections.

Cracow is a city which in the period of the Polish People's Republic enlarged demographically and territorially more than 2.5 times. As a consequence of this phenomenon some major socio-political differences and divisions occur, which corresponds to elections campaign, the results of local elections in Cracow and the stability of urban structures. Therefore Cracow is a micro-scale of political divisions which are also reflected in the Polish political landscape. This is undoubtedly an important feature which lets the city stand out.

Poznan is a city which for years has actively supported and lobbied for the development of self-government. Similarly to Cracow the city expanded in the last half century twice. But unlike Cracow, it didn't result in important social and political differences in functioning of Poznan. Both Poznan and the entire region are units created by people strongly attached to the region and a city, which is then reflected, among other things, in the specificity of local elections.

Wroclaw is a specific symbol of the success of revived self-government. Despite the absence of tradition, especially the continuity of local government (as opposed to Poznan which had it from the beginning the 19th century and Cracow has had it since the 1960's. the 19th century), in the last 20 years Wroclaw has managed to create a strong and stable self-government structure. Both the city and the region, made up of national-ethnic conglomerate, is displaying the exceptional stability and the vitality of self-government structures.

Above characteristics of the selected cities undoubtedly justify their choice to the planned presentation. Also, the process of direct elections for city mayor makes us include those cities in the speech. In Cracow, in 1990-2002, the mayor was changed five times, while in Poznan and Wroclaw, the mayor's position is strong and stable. The elections in Wroclaw, in contrast to Cracow and Poznan, end already in the first round of elections. In the case of elections to city councils a unique consensus and cooperation with the city mayor is adopted by the Wroclaw City Council. The most difficult situation is in Cracow, where councils are highly fragmented politically and ideologically, which affects cooperation with the mayor who is without political base. An indirect solution was worked out in Poznan.

For purposes of research comparative method was chosen. This is dictated by the intention of showing differences and similarities between the cities presented, which are derived from both socio-political, economic, cultural as well as historical factors.

Time frames for the speech cover years 1990-2002, so the period from the first to the fifth term of communal councils, including the local elections in 2010.

The thematic scope of the paper and speech is divided into two main parts:

- Elections to city councils 1990-2010
- The election of mayors 2001-2010

In both cases, the following things will be analyzed:

- Characteristics of the local self-government elections,

- Results of local elections,
- Electoral geography,
- Relations between deciding and executive bodies (cooperation v. antagonism)
- The specificity and uniqueness of these cities in the context of local elections,
- Local elections and other experiences (Łódź, Szczecin, Rzeszów)

The speech will be accompanied by presentation. The extended form will be placed in conference magazines.

Okrasa Włodzimierz (Główny Urząd Statystyczny, Uniwersytet Kardynała Stefana Wyszyńskiego w Warszawie)

Statystyka i socjologia: współzależności rozwoju z perspektywy interdyscyplinarnej badań społecznych. Aspekty metodologiczne i instytucjonalne

Rozwój metod statystycznych i stały wzrost ich zastosowań w badaniach empirycznych przyczynił się do rozbudowy arsenału metodologicznego zarówno poszczególnych dyscyplin naukowych jak i do interdyscyplinarności badań podejmowanych na przecięciu różnych dziedzin. Z drugiej strony, problemy zarówno specyficzne dla poszczególnych dyscyplin, jak też trans-dyscyplinarne, doprowadziły do powstania szeregu nowych metod statystycznych - tak w sferze organizowania materiału obserwacyjnego (zbierania danych) jak i analizy danych - a także do integracji badań społecznych. Na szczególną uwagę wydają się zasługiwać owe interakcje rozwojowe w obszarze statystyki i socjologii, jako dziedziny wykorzystującej w stosunkowo najszerszym zakresie metody statystyczne spośród wszystkich nauk społecznych. Wyraża się to w twórczym rozwijaniu zastosowań praktycznie wszystkich z głównych działów statystyki, które tradycyjnie występowały też w innych dyscyplinach społecznych (tzn. pomijając nauki przyrodnicze), ale nie na tak szeroka skalę. Łącząc zainteresowanie np. próbkowaniem i wnioskowaniem z n. politycznymi i demografią; badaniem związków pomiędzy zmiennymi i technikami ich pomiaru z psychologią i pedagogiką; analizami przyczynowymi i ewaluacyjnymi z ekonomią (w tym, z ekonometrią), czy analizami przestrzennymi z regionalistyką i geografiami społeczną, itp.

Próba systematyzacji tych powiązań, przyczyniających się zarazem do wzrostu interdyscyplinarności (oraz integracji) badań społecznych stanowi przedmiot pierwszej części referatu. Jej uzupełnieniem jest krótkie omówienie historii instytucjonalnych związków socjologii i statystyki w Polsce, w tym przedsięwzięć podejmowanych pod auspicjami Głównego Urzędu Statystycznego (poczynając od powołania pod koniec lat 60. Zespołu ds. Badań Statystyczno-Socjologicznych).

Olbrot-Brzezińska Arleta (Urząd Statystyczny w Poznaniu)

Społeczne funkcje informacji statystycznej i odpowiedzialność statystyki publicznej

Artykuł 3 ustawy z 29 czerwca 1995 r. o statystyce publicznej stanowi, że statystyka publiczna zapewnia rzetelne, obiektywne i systematyczne informowanie społeczeństwa, organów państwa i administracji publicznej oraz podmiotów gospodarki narodowej o sytuacji ekonomicznej, demograficznej, społecznej oraz środowiska naturalnego. W ten sposób ustawodawca określił nie tylko misję organów statystycznych w Polsce, ale i podstawową społeczną funkcję - gwaranta bezpieczeństwa informacyjnego obywateli, organów państwa i podmiotów życia gospodarczego. Funkcję podstawową, ale nie jedyną.

Statystyka publiczna pełni także funkcje stabilizacyjne dla polskiej gospodarki. Dane i informacje statystyczne stają się bowiem punktem odniesienia i podstawą ocen dotyczących sytuacji i rozwoju demograficznego, społecznego i ekonomicznego kraju. Ważne zatem by informacje te pochodziły od organów statystyki publicznej, były wynikiem oficjalnych badań statystycznych, których jakość została zweryfikowana i poparta rzetelną metodologią i autorytetem organu państwa. Podmioty życia gospodarczego, organy administracji samorządowej i rządowej oraz społeczeństwo korzystające na co dzień z w pełni dostępnej, rzetelnej i obiektywnej informacji statystycznej są odporne na manipulację. Dane statystyczne powinny stać się zatem podstawą wszelkich decyzji społecznych i ekonomicznych obywateli, organów władzy, organizacji społecznych i gospodarczych, a także dla badań naukowych i rozwoju naukowo-technicznego.

Informacja statystyczna pomaga zrozumieć otaczający świat, tłumaczy zachodzące zjawiska społeczne i ekonomiczne. We współczesnym świecie jest podstawowym i najcenniejszym dobrem społecznym. Jednak aby móc z niego w pełni korzystać zaistnieć muszą dwa podstawowe warunki - ich dostarczyciel - statystyka publiczna - cieszyć się musi pełnym zaufaniem ze strony odbiorców danych, a także, a może przede wszystkim konsumenci informacji statystycznej muszą posiadać podstawowe umiejętności z zakresu wyszukiwania i odczytywania danych statystycznych w różnych formach prezentacji, a także ich interpretacji i dokonywania dodatkowych działań obliczeniowych na danych. Wobec braku przedmiotów statystycznych w podstawowej edukacji szkolnej funkcje i zadania edukacyjne musi wziąć na siebie nie tylko statystycy funkcjonujący w środowiskach naukowych, ale przede wszystkim, ze względu na szeroki zakres możliwości podejmowania działań - przedstawiciele statystyki publicznej. Skutkiem postępującej degradacji wiedzy w tym zakresie jest bowiem rosnąca marginalizacja statystyki publicznej i danych statystycznych w życiu społeczno-gospodarczym kraju.

Statystyka publiczna jest zatem odpowiedzialna nie tylko za dostępność, rzetelność i obiektywny charakter danych statystycznych, a więc wypełnienie funkcji gwaranta bezpieczeństwa informacyjnego obywateli i funkcję stabilizacyjną polskiej gospodarki, ale i za edukację statystyczną polskiego społeczeństwa tak by mogło w pełni korzystać z możliwości jakie otwiera przed nim pełna dostępność do informacji.

Słowa kluczowe: statystyka publiczna, społeczna funkcja, społeczna odpowiedzialność

Public functions of statistical information and responsibility of public statistics

Article 3 of the Act from 6-11-1995 about public statistics says, that public statistics provides reliable informing to the public and state authorities about economic, demographic, social and environmental situation. Basic social function of statistical authorities is to be guarantor of information security in Poland.

Statistics are reference points and base ratings for today's demographic, social and economical condition and future progress of them. It's important then for statistics, to come from statistical authorities. People can use data that are tamper-proof. Therefore, statistical data should become base for any social and economical decisions.

Statistical information helps to understand surroundings, explains social and economical phenomena. Two conditions must exist to use statistics properly: recipients of statistics have to trust it completely and they have to acquire special knowledge how to use it.

Public statistics is then responsible not only for accessibility, reliability and objective nature of statistical data, but also for statistical education of polish society, so that they can take full advantage of the opportunities opened for them by full access to information.

Key words: public statistics, public (social) function, social responsibility

Okupniak Magdalena (Uniwersytet Ekonomiczny w Poznaniu)

Zastosowanie analizy wewnątrz- i międzyblokowej w badaniach społecznych

Analiza blokowa to rodzaj wielowymiarowej analizy danych często stosowany w badaniach związanych z rolnictwem. Badania te zwykle przeprowadza się w celu oceny wpływu na wzrost i rozwój roślin różnych czynników takich jak nawozy mineralne, czy środki ochrony roślin. Głównym jej prekursorem był brytyjski statystyk Ronald Aylmer Fisher, który na początku XX wieku sformułował schemat układu bloków kompletnie zrandomizowanych. Jest to szczególny przypadek układu bloków, w którym badamy wpływ danego obiektu zastosowanego na losowo wybranych jednostkach eksperymentalnych podzielonych na równoliczne i możliwie jednorodne grupy.

Celem referatu jest zaprezentowanie wyników próby zastosowania analizy blokowej w badaniu dochodu ekwiwalentnego gospodarstwa domowego. Jako źródło informacji wykorzystano dane zgromadzone przez Radę Monitoringu Społecznego w ramach Diagnozy Społecznej 2011. Baza

zawierała informacje o 12387 gospodarstwach domowych. Każde z nich zostało opisane za pomocą 259 zmiennych. W przeprowadzonym badaniu, w którym dokonano analizy wewnątrz- i między - blokowej uwzględnione zostały zmienne takie jak: położenie gospodarstwa domowego, ilość osób wchodzących w skład gospodarstwa domowego, klasa miejscowości.

Słowa kluczowe: analiza wewnątrzblokowa, analiza międzyblokowa, dochód ekwiwalentny

Application the intra- and interblock analysis of block designs in social research

Analysis of the block designs is one of the multivariate data analysis methods which is often used in an agricultural studies. This type of research is usually carried out to assess the impact of various factors on the growth and development of plants, such as fertilizers, pesticides and others. The main precursor of block designs was a British statistician Ronald Aylmer Fisher who in an early twentieth century formulated a randomized block designs. It is a particular kind of the block designs, in which we examine the impact of an object used to randomly selected experimental units divided into equinumerous and possibly homogenous groups.

The aim of this study is to present results of application of the block analysis in an equivalent household income research. Data was obtained from the Council for Social Monitoring - The Social Diagnosis 2011. The database contained 12387 households which were described by 259 variables. The variables which were used in an intra- and interblock analysis were: location of household, number of people in the household and class of the place of residence.

Key words: interblock analysis, intrablock analysis, equivalised income

Ostasiewicz Walenty (Uniwersytet Ekonomiczny we Wrocławiu)

Rozwój myśli statystycznej w Polsce

Określenie początku powstania i rozwoju myśli statystycznej w Polsce nie jest łatwe. Chcąc mieć tę historię jak najdłuższą, niektórzy datują ją na okres życia kronikarzy. Jan Długosz (1415-1480) wymieniany jest wówczas jako pionier polskiego naukownictwa. Niedługo po nim żyjący inny duchowny, Marcin Kromer (1512-1589), także kronikarz, uważany jest za prekursora tej dziedziny.

Historię rozwoju myśli statystycznej i piśmiennictwa w języku polskim trzeba by rozpatrywać z perspektywy ogólnego rozwoju cywilizacyjnego Polski. Historię takiego rozwoju, Hugo Kołłątaj dzieli na cztery epoki. Pierwszą epokę obejmującą 350 lat stanowią lata od wprowadzenia chrześcijaństwa do założenia Akademii Krakowskiej w 1364 roku.

Koniec drugiej epoki wyznacza data sprowadzenia Jezuitów do Polski w 1564 roku. Skasowanie tego zakonu w całym kościele w 1773 roku zamyka trzeci etap obejmujący 213 lat. Czwarta epoka trwa do czasu „wymazania Polski z kart Europy”, czyli do roku 1795. Przedłużając tę klasyfikację wyodrębnimy więc okres niewoli, okres lat międzywojennych i ostatni etap lat powojennych.

W okresie niewoli obserwujemy ożywione zainteresowanie sprawami państwa. W tym okresie publikowane są pierwsze polskie prace statystyczne. W dość powszechnym mniemaniu początki statystyki w Polsce datowane są na lata Sejmu Wielkiego (1788-1792). Znalazło to swój wyraz w ogólnopolskiej konferencji naukowej zorganizowanej w 1993 „z okazji 75-lecia Głównego Urzędu Statystycznego i 200-lecia statystyki polskiej”. Jako „szteandarowe” nazwiska pierwszych statystyków wymienia się hrabię Moszczyńskiego i S. Stasica. Największą zasługą pierwszego z nich jest opracowanie metody statystycznej do określenia wymiaru podatkowego oraz sporządzenia odpowiednich tabel kalkulacyjnych. S. Staszic wymieniany jest powszechnie jako ten, dzięki któremu w języku polskim po raz pierwszy pojawiło się słowo „statystyka”, i który zapoczątkował publikacje na jej temat.

Słowa „statystyka” używał już H. Kołłątaj. Krytykując panującą w Polsce francuszczyznę pisał, że mowy polskiej używa się tylko z konieczności komunikowania się z warstwami niższymi społeczeństwa i w „listach statystycznych”.

Komentując w 1870 roku pracę S. Staszica, Z. Korzybski uznał ją raczej za memoriał polityczny niż pracę statystyczną.

Jeśli statystykę mielibyśmy rozumieć w sensie Conringa-Achenwalla, to raczej należałoby wskazać pracę J.K. Haura pt. *Oekonomika*. Praca ta wydana była w 1675 roku, a więc prawie równocześnie z pracami H. Conringa, który jako pierwszy z naukowców uczynił dyscyplinę akademicką. Prace Conringa były w Polsce dobrze znane. Opisując różne państwa, opisał on również państwo polskie, ale dość krytycznie. Pisał na przykład, że sejmy polskie są burzliwe i chaotyczne, że nie istnieją w Polsce przepisy podług, których miałyby zapadać decyzje. W odpowiedzi na fałszywy obraz Polski. Jan ze Wschowy napisał obszerny traktat o Polsce „De scopo reipublicae Poloniae”, który był wydany we Wrocławiu w 1665 roku, zaś w 1726 roku w Gdańsku wydane zostało tłumaczenie niemieckie tego dzieła.

Niezależnie od nurtu opisowego począwszy od drugiej połowy XVIII wieku rozwija się nurt matematyczny, zapoczątkowany przez A. Queteleta w Belgii i rozwijany przez Süßmilcha w Niemczech. Oba nurty były w Polsce doskonale znane. Świadczą o tym publikacje Marasiego (1866), Załęskiego (1868), Korzybskiego (1870) i inne. Poza nielicznymi przypadkami, w pracach tych, wyższość naukową przypisuje się statystyce teoretycznej podkreślając jej główne zadanie polegające na odkrywaniu praw „wielkiej liczby”. Przedmiot badań statystycznych upatrywano w każdej dziedzinie działalności ludzkiej. Już w owych czasach wyodrębniano statystykę lekarską, statystykę ludności, statystykę kredytu i pieniądza itp. Nurtu zapoczątkowanego przez arytmetyków politycznych w Anglii do statystyki nie zaliczano. Pojęcia „statystyka matematyczna” w języku polskim użył po raz pierwszy B. Danielewicz. W pracy „Z dziedziny statystyki matematycznej” (1884) pisze, że nauka przyszłości dla odróżnienia od tego, co dziś się przez statystykę rozumie, przybierze prawdopodobnie miano statystyki matematycznej lub analitycznej, zachodzić tu bowiem będzie nieco podobny stosunek jak np. pomiędzy fizyką doświadczalną

a matematyczną”. Najważniejszym zadaniem statystyki matematycznej uważał „wyznaczenie prawdopodobieństwa śmierci”, zakładając, że śmierć człowieka jest zjawiskiem przypadkowym. Stąd problematyka ta stanowi jedyny temat wymienionej wyżej pracy. W pracy zaś „zarys arytmetyki politycznej” napisanej wspólnie z S. Dicksteinem omawiana jest szczegółowo cała problematyka matematyki finansowej, rachunku prawdopodobieństwa, ubezpieczeń a także gier losowych. A. Danielewicz zamyka etap zniewolonej Polski, w okresie tym zrodziła się statystyka, powstało wiele pięknych prac. W początkowym okresie niepodległej Polski zaczyna się rozwijać statystyka publiczna. Powstaje GUS i formują się struktury instytucjonalne tej statystyki.

Owsiński Jan (Instytut Badań Systemowych PAN w Warszawie)

Optymalny podział rozkładu empirycznego (i kilka problemów z tym związanych)

Praca zajmuje się zagadnieniem podziału empirycznego rozkładu wielkości pewnego zjawiska - oznaczmy tę wielkość przez x_i , gdzie i jest indeksem jednostki, dla której obserwujemy daną wielkość (np. województwo, dla którego x_i jest poziomem PKB na głowę mieszkańca). Indeks i przybiera wartości od 1 do n . Załóżmy dalej, że wartości x_i zostały uporządkowane od najmniejszej do największej. Analizujemy rozkład w postaci skumulowanej, tj. W postaci wartości $z_i = \sum_{i'=1}^i x_{i'}$. Tak utworzony ciąg wartości z_i jest, oczywiście, wypukły. Naszym zadaniem jest taki podział tego rozkładu skumulowanego, na podzbiory, odpowiadające podzbiорom indeksów i , a więc podzbiорom zbioru $1, \dots, n$, by, z jednej strony, można było kształt rozkładu z_i przybliżyć z możliwie najmniejszym błędem przy pomocy odcinków linii prostej, wyznaczonych dla poszczególnych segmentów podziału, tworzących linię łamaną, a z drugiej - by tych segmentów było możliwie mało.

Zagadnienie to odpowiada kwestii bardzo często używanych kategoryzacji, stosowanych dla wielu podobnych rozkładów (np. „kraje rozwinięte”, „kraje rozwijające się”, „kraje ubogie” itp.). Przy stosowaniu takich kategoryzacji zazwyczaj albo w ogóle nie stosuje się metod statystycznych, a tylko pewne (na ogół słabo uzasadnione) przesłanki „merytoryczne”, albo też zastosowane metody statystyczne ograniczają się do ustalenia pewnych parametrów (kwantyli) rozkładu, bez uwzględniania jego faktycznego kształtu i możliwości istnienia obiektywnych przesłanek do wyznaczenia innej liczby przedziałów, być może optymalizującej wspomniane ogólne kryterium.

Zaproponowano ogólne podejście do optymalizacji podziału tak określonego rozkładu w duchu wspomnianego kryterium, sformułowano odpowiednią ogólną funkcję celu i jej konkretne realizacje dla prostego rozkładu empirycznego, wraz z odpowiednimi algorytmami. Zaproponowana metodyka pozwala dla dowolnych rozkładów otrzymać optymalne podziały na kategorie, uzasadnione w ramach przyjętego ogólnego kryterium.

Zarazem, na podstawie przykładów konkretnych rozkładów, otrzymanych - na przykład - w badaniach stopnia zamożności społeczeństw lub satysfakcji z życia w określonych krajach, zary-

sowano problemy, wynikające z faktu, że otrzymane w takich badaniach rozkłady mają często charakter, stawiający pod znakiem zapytania zarówno podstawy przyjętej metodyki, jak i w ogóle sens dokonywania podobnych ocen. W referacie przytoczono przykłady takich znanych rozkładów empirycznych i przeanalizowano zarówno ich możliwe pochodzenie, jak i konsekwencje dla ewentualnej kategoryzacji.

W podsumowaniu stwierdza się, że zaproponowana metodyka, wraz z odpowiadającą jej funkcją kryterium, stanowią nie tylko podstawę do ewentualnej kategoryzacji w odniesieniu do kumulatywnej funkcji rozkładu, ale i narzędzie do oceny racjonalności sposobu otrzymania rozkładów empirycznych przy pomocy, na przykład, ocen ekspertów.

Słowa kluczowe: rozkład empiryczny, kategoryzacja, funkcja kryterium

On the optimal division of an empirical distribution (and some related problems)

The paper deals with the question of dividing up an empirical distribution of a certain quantity - denote its respective values by x_i , where i is the index of a unit, for which we observe the quantity (like, e.g., a province, for which x_i is the value of GDP per capita). Index I takes values 1 through n . Assume that values x_i have been ordered increasingly. We now analyse the distribution in its cumulative form, i.e. the distribution of values z_i , where $z_i = \sum_{i'=1}^i x_{i'}$. The thus formed sequence of values z_i is, of course, convex. The task consists in dividing this cumulative distribution into subsets, corresponding to the subsets of values of index i , i.e. of the set $1, \dots, n$, in such a way that the shape of the distribution z_i is possibly well approximated by the segments of the straight line, determined for the particular segments of the division, forming a piecewise linear contour, but, on the other hand - the number of separate segments is possibly small.

This issue corresponds to the question of the very frequently used categorisations, applied for numerous similar distributions (like, e.g., „developed countries”, „developing countries”, „underdeveloped countries”, etc.). When using such categorisations, usually either formal methods are not applied at all, and only certain (mostly poorly justified) „substantive” prerequisites, or the methods applied are limited to establishment of some parameters of the distribution, without due attention to its actual shape and the possibility of existence of objective premises for determination of a different number of segments, including the potential optimisation of the criterion mentioned before.

A general approach is proposed for dealing with optimisation of division of this kind of distribution in the vein of the criterion mentioned. The respective general objective function is proposed, and its concrete realisations for a simple empirical distribution, along with corresponding algorithms. The methodology proposed allows for obtaining the optimum divisions into categories, justified in the framework of the criterion formulated, for arbitrary distributions.

At the same time, though, on the basis of concrete examples of distributions, obtained, for instance, in the studies of the wealth of societies or life satisfaction in definite countries, problems are outlined, resulting from the fact that the distributions, resulting from such studies, often display the features, which lead to both questioning the foundations of the methodology

proposed, and to questions concerning the very sense of such assessments. Examples of known distributions of this kind are quoted, and the way they arise, as well as the consequences for the potential categorisations are discussed.

In summary, it is stated that the methodology proposed, including the criterion function, constitutes not only a basis for the potential categorisation with respect to the cumulative distribution, but also a tool for evaluating the rationality of way, in which the distributions are obtained, through, for instance, expert assessments.

Key words: empirical distribution, categorisation, criterion function

Paradysz Jan (Centrum Statystyki Regionalnej, Uniwersytet Ekonomiczny w Poznaniu)

Statystyka regionalna: stan, problemy i kierunki rozwoju

Statystyka regionalna jest ilościowym opisem tożsamości administracyjnych i funkcjonalnych części danego kraju zwanych regionami. Istotą statystyki regionalnej jest to, że stanowi ona immanentną część statystyki ogólnokrajowej, jako hierarchicznego układu geograficzno-infrastrukturalnych powiązań gospodarczych, demograficznych i społecznych. Przedmiotem wystąpienia będzie ocena stanu pokrycia informacyjnego dla funkcjonowania Państwa i samorządu lokalnego w gospodarce rynkowej po 1990r w Polsce. System statystyki regionalnej powinien uwzględniać możliwości budowy układów informacyjnych wykraczających poza podziały administracyjne oparte na NUTS. Chodzi tutaj o zapewnienie informacji, na przykład, dla obszarów specjalnej troski (strefy ekonomiczne, grupy etniczne, walory środowiskowe, obszary zagrożone powodziami) lub zadaniowych (statystyka trans graniczna, euro regionalna, metropolitalna). Zwrócimy uwagę na rozwój metodologii badań statystycznych w przekrojach regionalnych na początku XXI wieku. Szczególnie uwypuklimy problemy, które powstały w czasie prac przygotowawczych do Powszechnego Spisu Rolnego 2011 oraz Narodowego Spisu Powszechnego 2012. W czasie tych prac poddano wnikliwej ocenie jakość informacji statystycznych pozyskiwanych w tradycyjnych spisach i badaniach reprezentacyjnych bez wykorzystywanie alternatywnych źródeł danych gromadzonych przez administrację państwową. Przedstawimy pierwsze wyniki bezpośredniego wykorzystania administracyjnych źródeł danych oraz omówimy wyniki wykonanej estymacji pośredniej w oparciu o BAEL i NSP. W ostatniej części wystąpienia wskażemy na dalsze kierunki rozwoju statystyki regionalnej w warunkach integracji źródeł i specjalnych badań reprezentacyjnych z wykorzystaniem efektu synergii. Pochylimy się także nad problemami związanymi z jakością danych regionalnych oraz kryteriami oceny wyników badań statystycznych.

Słowa kluczowe: statystyka regionalna, spis rolny, spis ludności, estymacja pośrednia, jakość danych statystycznych, administracyjne źródła danych

Regional statistics: condition, problems and development directions

Regional statistics is a quantitative description of the identity of the administrative and functional parts of the country called regions. The essence of regional statistics is that it is an inherent part of national statistics, as a hierarchical system of infrastructural linkages geo - economic, demographic and social facts. The subject of our paper will be an assessment of the information for the functioning of the State and local government in a market economy conditions after 1990 in Poland. The system of regional statistics should take into account the possibility of building information systems beyond administrative divisions based on the NUTS. It is here to provide information, for example, for special attention areas (economic zones, ethnic groups, environmental qualities, and flood risk areas) or task (cross border, euro regional and metropolitan statistics). We will pay attention to the development of statistical methodology in the regional cross sections at the beginning of the twenty-first century. Especially we accent problems that arose during the preparations for the 2011 Agricultural Census and the National Census 2012. During this work were carefully assessed the quality of statistical information obtained in traditional censuses and sample surveys without the use of alternative sources of data collected by the state administration. We will present the first results of the direct use of administrative data sources and discuss the results of indirect estimation made based on the LFS and the population census. In the last part we point the future directions of regional statistics in the conditions of the integration of source and special sample surveys with the use of synergy effects. We will identify also the problems relating to the quality of regional data and criteria for evaluation of survey results.

Key words: Regional statistics, agricultural census, census, indirect estimation, the quality of statistical data, administrative data sources

Parlińska Maria, Babiak Robert (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie)

Dokumentowanie i ryzyko asymetrii informacji

Największe korzyści do dyscypliny ekonomiki informacji wniosły badania noblisty Joseph E. Stiglitz'a, dotyczące technik dokumentowania (screening'u) używanych przez jednego agenta ekonomicznego celem uzyskania informacji od innego. Artykuł pokazuje jak w 2011 roku dwie instytucje, jedna z UE oraz druga instytucja rządowa, uzupełniają niejasności informacji przed podpisaniem Memorandum Porozumienia dotyczącego rozwoju budownictwa w Holandii. W artykule tym, autorzy analizują także przykład przemysłu transportowego w 2010 roku w Hiszpanii z interwencją rządową, bazującą na niepełnej informacji powodującej wysokie koszty z powodu takich braków informacji. Modele dokumentowania (screening'u) pozostają w kontraście z rządową polityką interwencjonizmu. Taka polityka ma limitowaną wartość, jeśli screening jest wykonany poprzedzając rządowy interwencjonizm. Faktyczny proces dokumentowania zależy

od natury scenariusza, ale jest także zależny od przyszłego związku między instytucjami.

Słowa kluczowe: asymetria informacji, ekonomika informacji, interwencjonizm, niepewność, ryzyko

Screening and risks of information asymmetry

The greatest benefit to the discipline of economics of information is possibly the Stiglitz research on screening techniques (ref:1) used by one economic agent to extract otherwise private information from another agent. The article cites how in 2011 two institutions, one being a European Union Institution and the other a Government Ministry, remedied information ambiguities before signing the Memorandum of Understanding for a major building development in the Netherlands. In this article the authors also analyzed the example of trucking industry in 2010 Spain with government interference based on imperfect information and the very high cost of such information imperfections.

The screening models are contrasted with the government interventionism policy. Such policy is of limited value if screening is carried out prior to any direct government intervention. The actual screening process depends on the nature of the scenario, but is also connected with the future relationship between the parties.

Key words: information asymmetry, economics of information, intervention, uncertainty, risk

BIBLIOGRAPHY

1. Joseph E., Stiglitz. (2001), "Information and the Change in the Paradigm in Economics. Aula Magna Nobel Prize Lecture, Columbia Business School, Columbia University.
2. Emilio Reineri. (2003), "The Challenges of Immigration and Integration in the European Union and Australia", Department of Sociology and Social Research University of Milan Bicocca National Europe Centre Paper No. 66 Paper presented to conference entitled 18-20 February 2003, University of Sydney.
3. Frédéric Marty, Arnaud Voisin (2008), "Partnership contracts, project finance and information asymmetries: from competition for the contract to competition within the contract ?", N° 2008-06 Février.
4. Tony Bunyan (2008), "The Shape of Things to come - EU Future", Group 2008 State-watch.
5. Geoff Coyle (2004), "The Analytic Hierarchy Process (AHP) Practical Strategy", Pearson Education Limited.
6. David Pearce Giles, Atkinson Susana Mourato (2006), "Cost-Benefit Analysis and the Environment Recent Developments", ISBN 92-64-01004-1 OECD.

7. Peter Guerra (2009), "How Economics and Information Security Affects Cyber Crime and What it Means in the Context of a Global Recession", White Paper.
8. Ross Anderson and Tyler Moore (2006), "The Economics of Information Security: A Survey and Open Questions", University of Cambridge Computer Laboratory.

Paszko Elżbieta, Jurek Anna Maria (Uniwersytet Łódzki)

Zastosowanie CQI - ciągłego doskonalenia jakości w procesie poprawy metod nauczania

Duże sukcesy stosowania zasady TQM - Total Quality Management w przemyśle spowodowały zainteresowanie wprowadzenia TQM w procesie nauczania. Pod koniec ubiegłego stulecia na podstawie TQM powstało nowe podejście - CQI (Continuous Quality Improvement), czyli Ciągłe Doskonalenie Jakości. W niniejszym artykule przedstawiona zostanie główna koncepcja TQM oraz zastosowanie CQI w szkolnictwie wyższym. Poprzez charakterystykę zasady, opis problemów towarzyszących wprowadzeniu tej koncepcji w dziedzinie nauczania akademickiego oraz krótką prezentację stosowania CQI na najlepszych uniwersytetach światowych zaprezentowane zostaną zarówno korzyści wynikające z wdrożenia tej zasady jak i towarzyszące temu trudności. Celem pracy jest zwrócenie uwagi na pozytywne skutki koncepcji CQI przy tworzeniu i wdrażaniu strategii rozwoju w Polskich szkołach wyższych.

Słowa kluczowe: TQM - System Sterowania Jakością, CQI - Ciągłe Doskonalenie Jakości, nauczanie, szkolnictwo wyższe

Application of TQM - total quality management in process of improvement teaching methods

Successes of application the TQM (Total Quality Management) approach in industry caused increasing interest of implementation TQM in process of teaching. At the end of 20th century CQI (Continuous Quality Improvement) appeared as a modification of TQM approach. In this paper TQM concept and CQI implementation are presented in higher education context. The characteristic of this approach, description of problems connected with application CQI in area of academic teaching methods and short presentation of working CQI in the best world universities will show not only all advantages of applying CQI but also attendant difficulties. The purpose of this paper is to focus on positive effect of CQI concept in creating and introducing development strategy in Polish universities.

Key words: TQM - Total Quality Management, CQI - Continuous Quality Improvement, teaching, higher education

Pawełek Barbara (Uniwersytet Ekonomiczny w Krakowie)

Badanie przydatności wyników testu koniunktury GUS do prognozowania krótkoterminowego zmian aktywności gospodarczej w Polsce na podstawie modelu wektorowej autoregresji

W referacie zostanie dokonany przegląd prezentowanych w polskiej i światowej literaturze przedmiotu propozycji metodycznych, dotyczących prognozowania krótkoterminowego zmian PKB na podstawie modelu wektorowej autoregresji z wykorzystaniem wskaźników budowanych w oparciu o test koniunktury.

Głównym celem wystąpienia jest przedstawienie wyników badania nad przydatnością wyników testu koniunktury GUS do prognozowania krótkoterminowego zmian aktywności gospodarczej w Polsce na podstawie modelu wektorowej autoregresji. W analizie zostaną uwzględnione wyniki testu koniunktury w przemyśle, budownictwie i handlu oraz testu koniunktury konsumenckiej.

Słowa kluczowe: koniunktura gospodarcza, prognozowanie krótkoterminowe, aktywność gospodarcza, model wektorowej autoregresji

Study of suitability of CSO business tendency survey to short-term forecasting of changes in economic activity in Poland based on vector autoregression model

In the paper will be reviewed methodological proposals presented in literature both Polish and world relating to short-term forecasting of changes in GDP based on vector autoregression model with using the indicators built based on a business tendency survey and consumer tendency survey.

The main objective is to present the results of a study on the usefulness of the CSO business tendency survey results to short-term predicting of changes in economic activity in Poland based on vector autoregressive model. In the analysis will be included results of business tendency survey in industry, construction and trade and consumer climate test.

Key words: business tendency, short-term forecasting, economic activity, vector autoregression model

Pełka Marcin (Uniwersytet Ekonomiczny we Wrocławiu)

Podejście wielomodelowe w klasyfikacji danych symbolicznych interwałowych

Obiekty symboliczne, w odróżnieniu od obiektów w ujęciu klasycznym, mogą być opisywane przez wiele różnych typów zmiennych. Oprócz zmiennych w ujęciu klasycznym (metrycznych lub niometrycznych) mogą być opisywane przez zmienne interwałowe, zmienne wielowariantowe i zmienne wielowariantowe z wagami, a także zmienne strukturalne [zob. np. Bock, Diday 2000, s. 2-3]. Pozwala to na dokładniejszy opis obiektów, ale utrudnia analizę skupień.

Podejście wielomodelowe było dotychczas z powodzeniem stosowane z zagadnieniami dyskryminacyjnymi i regresyjnymi [zob. np. Gatnar 2008]. Niemniej jednak idea podejścia wielomodelowego, tj. łączenia wyników wielu modeli, może być z powodzeniem zastosowana także w zagadnieniu klasyfikacji danych symbolicznych. Podejście wielomodelowe w klasyfikacji to nic innego jak łączenie (czyli agregacja) N klasyfikacji (modeli) bazowych $P_1, \dots, P_N$ w jedną klasyfikację złożoną P^* o k^* klasach [por. Fred, Jain 2005].

Celem artykułu jest przedstawienie wyników klasyfikacji danych symbolicznych interwałowych z wykorzystaniem podejścia wielomodelowego na przykładzie różnych zbiorów danych symbolicznych. W części empirycznej przedstawiono wyniki badań symulacyjnych (na podstawie wygenerowanych zbiorów danych o znanej strukturze klas) i porównano ich zgodność z klasyfikacją rzeczywistą.

Ensemble approach for clustering of interval-valued symbolic data

Symbolic objects, unlike classical objects, can be described by many different symbolic variable types. Besides well-known classical variables (metric or not) symbolic objects can be described by interval-valued variables, multivalued variables and multivalued variables with weights and also dependent variables [see for example Bock, Diday 2000, p. 2-3]. This allows to describe objects more accurately, but it makes cluster analysis more difficult.

Ensemble approach has been applied with a success to regression and discrimination tasks [see for example Gatnar 2008]. Nevertheless the idea of ensemble approach, that is combining (aggregating) the results of many base models, can be applied for cluster analysis of symbolic data. Ensemble clustering means combining (aggregating) N base clustering results (models) $P_1, \dots, P_N$ into one model P^* with K^* clusters [see: Fred, Jain 2005].

The aim of the article is to present the results of ensemble clustering of symbolic interval-valued data. The empirical part of the paper presents results simulation studies (based on generated data sets with known cluster structure) and agreement indices between two partitions were calculated.

Key words: ensemble clustering, interval-valued symbolic data

BIBLIOGRAPHY

1. Bock H.-H., Diday E. (red.) (2000), *Analysis of symbolic data. Explanatory methods for extracting statistical information from complex data*, Springer Verlag, Berlin-Heidelberg.
2. Fred A.L.N., Jain A.K. (2005), *Combining multiple clustering using evidence accumulation*, „IEEE Transaction on Pattern Analysis and Machine Intelligence”, Vol. 27, p. 835-850.
3. Gatnar E. (2008), *Podejście wielomodelowe w zagadnieniach dyskryminacji i regresji*, Wydawnictwo Na-ukowe PWN, Warszawa.
4. Hornik K. (2005), *A CLUE for CLUster Ensembles*, „Journal of Statistical Software, Vol. 14, Issue 12, s. 1-25.
5. Okun O., Valentini G. (2010) (red.), *Supervised and unsupervised ensemble methods and their applications*, Springer-Verlag, Berlin-Heidelberg.

Perek-Białas Jolanta (Szkoła Główna Handlowa w Warszawie, Uniwersytet Jagielloński w Krakowie)

Sytuacja starszych generacji w Europie Środkowo-Wschodniej na podstawie danych EU-SILC

W niniejszym artykule zaprezentujemy wyniki porównawczej analizy sytuacji społeczno-ekonomicznej starszych generacji dla wybranych krajów Europy Środkowo-Wschodniej: Polski, Litwy, Węgier, Czech, Słowacji, Łotwy, Estonii, które są w UE od 2004 oraz Bułgarii i Rumunii, które zostały członkami UE w 2007 z wykorzystaniem indywidualnych danych z badania EU-SILC (*European Union - Survey of Income and Living Condition*). Celem opracowania jest pokazanie, na ile badanie EU-SILC może być wykorzystane do oceny sytuacji społeczno-ekonomicznej starszych populacji w wymienionych krajach. Oprócz tego, zostanie podjęta dyskusja pokazująca zalety i wady badania w tego typu analizach, wraz ze wskazaniem braków informacyjnych w pomiarze tych zjawisk, które należy stosować względem starszych generacji, a które mogą być rekomendowane do zmiany/wprowadzenia w badaniach EU-SILC w przyszłości. Wskażemy też na aspekty badania, które są kluczowe do poprawnej analizy wyników z uwzględnieniem odpowiednio szacowanymi błędami standardowymi, jeśli mamy do czynienia z próbami złożonymi, co ma miejsce w badaniu EU-SILC.

Słowa kluczowe: EU-SILC, osoby starsze, sytuacja społeczno-ekonomiczna, kraje Europy Środkowo-Wschodniej

Situation of older generations based on EU-SILC in selected Central and Eastern European Countries

This paper offers results of comparative analysis of some countries of Central and Eastern Europe (CEE): Poland, Lithuania, Hungary, Czech Republic, Slovakia, Latvia, Estonia who joined the EU in 2004 and Bulgaria and Romania who joined the EU in 2007 by presenting the socio-economic situations of older generations using the microdata of EU-SILC (*European Union - Survey of Income and Living Condition*). The aim of this paper is to check if the EU-SILC survey could be used in analysis of socio-economic situations of older persons in these selected countries. We also present advantages and disadvantages of using this survey in analysis of this particular group with showing the missing gaps of information which should be applied for older generations and could be recommended to introduce in EU-SILC surveys in future. We will also show some methodological issues of the survey which are important to obtain a proper analysis with correct estimations of standard errors, if there is applied a complex sample design as such one could be found in EU-SILC.

Key words: EU-SILC, older generation, socio-economic situation, Central and Eastern Europe

Pietrzykowski Robert (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie)

Wykorzystanie regresji kwantylowej w analizach przestrzennych cen ziemi rolniczej

Podstawowym elementem analiz przestrzennych jest określenie interakcji pomiędzy badanymi jednostkami przestrzennymi: krajem, województwem, powiatem czy gminą. Zgodnie z prawem Toblera [Tobler 1970] jednostki przestrzenne sąsiadujące ze sobą powinny bardziej oddziaływać na siebie niż te, które znajdują się dalej. W praktyce do oszacowania funkcji ceny wykorzystuje się z reguły metodę regresji wielokrotnej z estymatorem MNK (metoda najmniejszych kwadratów), MNW (metoda największej wiarygodności), LMD (metoda najmniejszych odchyłeń bezwzględnych) [Greene 2000]. Nowym podejściem do analizy zależności jest regresja kwantylowa [Koenker i Bassetta 1978]. W modelach ekonomicznych często zdarza się, że reszty nie spełniają założenia o normalności rozkładu. Jeżeli chodzi o estymatory najmniejszych kwadratów to są one nadal nieobciążone i zgodne, ale przestają być efektywne. Powoduje to niepoprawne oszacowanie wariancji i obniżenie mocy używanych testów statystycznych. Metoda regresji kwantylowej jest pozbawiona tej wady i należy do odpornych metod estymacji. W analizach przestrzennych kluczowym elementem jest określenie zależności przestrzennych poprzez macierz wag. Na bazie ustalonych zależności można mówić o reżimach przestrzennych ze względu na badaną cechę.

Słowa kluczowe: analiza przestrzenna, regresja kwantylowa, NEG

Spatial analysis of agricultural land prices used to spatial quantile regression

In practice, to estimate regression function is used as a rule, the method of multiple linear regression with OLS estimator (least squares), MNW (the method of maximum likelihood), LMD (least absolute deviations) [Greene 2000]. New approach to the analysis of the relationship is a quantile regression [Koenker i Bassett 1978]. The aim of this study was to compare the spatial regression model estimated on the method of spatial OLS and quantile regression model using different weight matrices. The study is a continuation of the author analyses associated with the use of modified weight matrix [Pietrzykowski 2011]. The study was conducted on the basis of agricultural real estate market.

Key words: spatial analysis, quantile regression, NEG

BIBLIOGRAPHY

1. Anselin L. (2001): Spatial econometrics, Oxford, Basil Blackwell
2. Greene W. H. (2000), Econometric Analysis, 4th ed., Prentice Hall, Saddle River, N. J.
3. Koenker R., Bassett G. (1978): Regression Quantiles, „Econometrica”, No. 46
4. Kopczewska K. (2007): Ekonometria i statystyka przestrzenna, CEDEWU.PL, Warszawa
5. Pietrzykowski R. (2011): Koncepcja i zastosowanie modyfikacji macierzy wag w przestrzennych badaniach ekonomicznych (w druku)
6. Tobler W., (1970): A computer movie simulating urban growth in the Detroit region. Economic Geography, 46(2): 234-240
7. Suichecki B. (2010): Ekonometria przestrzenna. Metody i modele analizy danych przestrzennych, C.H. Beck, Warszawa

Piotrowska-Trybull Marzena, Sirko Stanisław (Akademia Obrony Narodowej)

Wpływ jednostki wojskowej na rozwój gminy

Rozwój społeczno-ekonomiczny poszczególnych gmin jest warunkowany czynnikami wewnętrznymi i zewnętrznymi. Na terenie niektórych gmin stacjonują jednostki wojskowe, które są elementem struktury organizacyjnej, różnych rodzajów Sił Zbrojnych RP. Jednostki te leżą na obszarach przyrodniczych, które obfitują w piękne krajobrazy oraz rzadką florę i faunę, inne rozmieszczone są na obszarach przemysłowych lub w dużych aglomeracjach. Ich lokalizację, na terenie poszczególnych gmin, determinuje szereg powiązań z otoczeniem lokalnym. Powiązania

mają charakter dwustronny. Jednostki czerpią z otoczenia zasoby (ludzkie, rzeczowe i informacyjne), w oparciu o które realizują przypisane im zadania. Natomiast do otoczenia dostarczają dobro publiczne, jakim jest bezpieczeństwo obywateli oraz generują efekty wynikające z użytkowania zasobów. Te zależności, w głównej mierze, skłoniły autorów do podjęcia badań zmierzających do określenia wpływu jednostki wojskowej na sytuację społeczno-ekonomiczną w gminie.

Słowa kluczowe: rozwój, społeczno-gospodarczy, gmina, jednostka wojskowa, wpływ, powiązania

Influence of military units on the community's development

Socio-economic development of particular communities are conditioned by internal and external factors. Military units are located on the territory of some of the Polish communities. They are the elements of organizational structure of the Army in Poland. Sometimes these units are situated on the territories which are rich in beautiful landscapes and rather rare flora and fauna, others are located in industrial districts or in big agglomerations. Their localization determines a lot of connections with local surroundings. These connections have a bilateral character. Military units use local resources from their surroundings (human, material and information) and they fulfill different tasks thanks to them. They also provide public goods to the local surroundings (e.g. safety of citizens) and they create effects for local society which depend on the ways of using resources. These are the main dependencies, that induced authors to conduct research aimed at identification of the influence of military units on the socio-economic development in the community.

Key words: development, socio-economic, community, military units, influence, connections

Pobłocka Agnieszka (Uniwersytet Gdański)

Polski rynek ubezpieczeń w latach 1991–2010

Ubezpieczenia pełnią istotną rolę w stabilizacji życia społeczeństw, gdyż stanowią zabezpieczenie przed skutkami niekorzystnych zdarzeń losowych. Polski rynek ubezpieczeń w gospodarce wolnorynkowej powstał w 1991 roku. W celu zbadania poziomu rozwoju ubezpieczeń w Polsce na tle innych państw UE zaprezentowana zostanie analiza podstawowych charakterystyk ubezpieczeniowych oraz wskaźników rozwoju ubezpieczeń w latach 1991-2010.

Słowa kluczowe: ubezpieczenia życiowe i majątkowe, gęstość ubezpieczeń, wskaźnik penetracji rynku

Market of insurance polish in 1991–2010

Insurances fulfill important role in stabilization of life of society, because they protect before results of fate damages. Market of insurance polish has emerged in free-market economy in 1991 year. Analysis of basic actuarial characteristic will be presented for researching level of development of insurance in Poland on background of other state UE in 1991-2010.

Key words: Insurance Life and Non-Life, destiny, penetration

BIBLIOGRAPHY

1. Gąsiorkiewicz L., Monkiewicz J. (red.), Ubezpieczenia w zarządzaniu ryzykiem przedsiębiorstwa. Tom. 2. Zastosowania, Wyd. POLTEX, Warszawa 2010
2. Handschke J., Monkiewicz J. (red.), Ubezpieczenia. Podręcznik akademicki, Wyd. POLTEX, Warszawa 2010
3. Hardyniak B., Monkiewicz J. (red.), Ubezpieczenia w zarządzaniu ryzykiem przedsiębiorstwa. Tom 1. Podstawy, Wyd. POLTEX, Warszawa 2010
4. Pajewska R., Wybrane cechy diagnostyczne - opis i charakterystyka [w:] Ubezpieczenia gospodarcze w Polsce i w Unii Europejskiej pod red. Michalski T., Wyd. Difin, Warszawa 2001
5. Ronka-Chmielowiec W. (red.), Zastosowanie metod ilościowych w analizie i ocenie ubezpieczeń dla działalności gospodarczej, Wyd. Uniwersytetu Ekonomicznego, Wrocław 2009
6. Sangowski T. (red.), Finansowe narzędzia zarządzania zakładem ubezpieczeń, Poltext, Warszawa 2005
7. Sangowski T., Szumlich T. (red.), Rynek ubezpieczeniowy po przystąpieniu Polski do Unii Europejskiej, Forum dyskusyjne ubezpieczeń i funduszy emerytalnych, Komisja Nadzoru Ubezpieczeń i Funduszy Emerytalnych, Zeszyt 2/2004, Wyd. Edytor, Warszawa 2004
8. Wierzbicka E., Ubezpieczenia non-life, Wyd. CeDeWu, Warszawa 2010

Pociecha Józef (Uniwersytet Ekonomiczny w Krakowie)

Okoliczności powstania polskiego towarzystwa statystycznego w Krakowie i jego pierwszy prezes

W drugiej połowie XIX wieku Kraków stał się duchową stolicą Polski, nieistniejącej od 1795 roku jako byt państwowy. Z tego też względu właśnie w Krakowie, w gronie statystyków i ekonomistów, powstał na początku 1912 roku projekt utworzenia odrębnego stowarzyszenia statystyków polskich. Głównym celem działalności tego stowarzyszenia miało być opracowywanie publikacji

statystycznych obejmujących swym zasięgiem ziemie polskie należące w tym czasie do trzech różnych państw zaborczych. Wobec braku państwowości polskiej zebrania danych statystycznych dla ziem polskich, doprowadzenia ich do porównywalności oraz ich opublikowania mogła podjąć się jedynie organizacja społeczna. Powstanie odpowiedniego stowarzyszenia mogło nastąpić jedynie w Galicji, cieszącej się wówczas autonomią i względnie liberalnymi prawami. W dniu 9 kwietnia 1912 r. C. K. Namiestnictwo w Galicji zatwierdziło projekt statutu PTS i datę tą przyjmuje się za początek istnienia Polskiego Towarzystwa Statystycznego.

Po formalnej rejestracji PTS przystąpiło do działalności. Na jego czele stanął 12 osobowy zarząd, a jego pierwszym prezesem (i jedynym w okresie działalności PTS w Krakowie) był prof. dr Juliusz Leo. Z szeroko zaprojektowanej działalności publikacyjnej PTS wydało tylko jedną pozycję, a mianowicie *Statystykę Polski*, wydaną w 1915 roku. W pracy tej starano się podać dane dla obszaru Polski w granicach przedrozbiorowych. Zakres tematyczny, chronologiczny i terytorialny warunkowało jednak istnienie wiarygodnych źródeł statystycznych. *Statystyka Polski* Wydana przez PTS była de facto pierwszym polskim rocznikiem statystycznym, uwzględniającym dane statystyczne z trzech zaborów. Ambitnie zaplanowaną działalność wydawniczą i naukową PTS przerwał wybuch pierwszej wojny światowej. Wobec mobilizacji wojsk, działań wojennych na terenie Galicji a w związku z tym ogromnych zniszczeń materialnych oraz gwałtownie pogarszającej się sytuacji materialnej ludności większość stowarzyszeń naukowych zaprzestała swojej działalności. Po wydaniu *Statystyki Polski*, PTS nie prowadził już prawdopodobnie żadnej działalności. Po 1918 roku, śmierci jego prezesa Profesora Juliusza Leo oraz zakończeniu wojny, krakowskie Polskie Towarzystwo Statystyczne nie wznowiło swojej działalności. Prawdopodobnie nie zostało jednak nigdy formalnie rozwiązane.

Ambitnie zarysowane plany działania powstałego w 1912 roku Polskiego Towarzystwa Statystycznego związane są nierozdzielnie z osobą jej prezesa, prof. Juliusza Leo. Ukończył on Wydział Prawa Uniwersytetu Jagiellońskiego. Studiował także ekonomię polityczną, nauki skarbowości i statystykę w Berlinie i Paryżu. Po habilitacji w roku 1888 został zatrudniony na Uniwersytecie Jagiellońskim w charakterze docenta a w 1892 został profesorem nauki skarbowości i prawa skarbowego UJ. W 1893 r. Juliusz Leo został wybranym radnym Miasta Krakowa, w 1901 r. został wiceprezydentem miasta, a w 1904 r. został wybranym prezydentem Miasta Krakowa. Swoją funkcję pełnił nieprzerwanie do swojej śmierci w dniu 21 lutego 1918 roku. Pełnił on urząd wójtka miasta przez 14 lat, najdłużej spośród sześciu prezydentów Krakowa doby autonomii Galicji. Jego dokonaniem było sfinalizowanie wykupienia w 1905 r. Wawelu od Austriaków, jemu Kraków zawdzięcza przeobrażenie w nowoczesny wielkomiejski ośrodek. Jest on nade wszystko twórcą Wielkiego Krakowa (obszar Krakowa zwiększył się z 7 km² do 47 km²). Rozszerzenie terytorium Krakowa przeprowadził Leo dzięki żelaznej konsekwencji i wytrwałości, realizując to przedsięwzięcie przez 13 lat.

Juliusz Leo był także wybitnym politykiem konserwatywnym tej doby. W 1901 roku został wybrany do Sejmu Krajowego we Lwowie. W 1911 roku odbyły się wybory do parlamentu austriackiego w Wiedniu w których J. Leo i zdobył mandat poselski. Już przy konstituowaniu Koła Polskiego przy parlamencie wiedeńskim J. Leo został wybranym jednym z czterech wiceprezesów koła a w 1912 r. został przewodniczącym Koła Polskiego w parlamencie austriackim.

Rok 1912 jest szczytem osiągnięć politycznych i administracyjnych Juliusza Lea. W tym też kontekście należy popatrzeć na jego wybór na prezesa nowo powstałego Polskiego Towarzystwa

Statystycznego w Krakowie. Profesor Juliusz Leo miał zrozumienie dla ważności statystyki, w czasie swoich studiów w Berlinie uczęszczał na wykłady ze statystyki. Czynnych badań statystycznych jednak nie uprawiał. Jego wybór na prezesa PTS zapewniał prestiż i szerokie możliwości działania nowopowstałemu towarzystwu, jednoczącemu statystyków polskich ze wszystkich zaborów, gdyż jego prezes był prezydentem Krakowa oraz znanym i wpływowym politykiem, posłem we Lwowie i Wiedniu, przywódcą polskich członków parlamentu austriackiego.

Podogrodzka Małgorzata (Szkola Główna Handlowa w Warszawie)

Niedostosowanie strukturalne rynku matrymonialnego determinantą spadku procesu zawierania małżeństw w Polsce w latach 1990–2010

Obserwowany od początku lat 90. spadek natężenia zawieranych małżeństwa wiąże się głównie ze zmianą postaw potencjalnych małżonków oraz wzrostem ich wzajemnych oczekiwań. Celem artykułu jest ukazanie, w jakim stopniu niedopasowanie strukturalne tego rynku tj. nieodpowiedni wiek partnerów, wykształcenie, miejsc zamieszkania czy zmiany w proporcji płci może wpływać na spadek natężenia małżeńskości. W analizie wykorzystana zostanie procedura standaryzacji bezpośredniej oraz pośredniej.

Słowa kluczowe: małżeństwo

Non-structural market marriage as a factor of fall in the process of the marriage in Poland in 1990–2010

Observed since the beginning of the 1990s. decrease in the intensity of the marriage is related primarily to the change in attitudes of potential spouses and increases their mutual expected. The purpose of article is to demonstrate the extent to which the structural mismatch of this market. inadequate age partners, training, residences or changes in aspect of gender may influence the decrease in intensity of marry. In the analysis procedure used is the standardisation of the direct and indirect.

Key words: marriage

Ptak-Chmielewska Aneta (Szkoła Główna Handlowa w Warszawie)

Uwarunkowania rynku przedsiębiorstw w Polsce. Statystyka przeżywalności firm.

Zainteresowanie tematem demografii przedsiębiorstw w Polsce w ostatnich latach znacząco wzrosło. Wzrost popularności demografii przedsiębiorstw (business demography) związany był głównie z projektem Eurostatu i zapisem podstawowych celów strategii lizbońskiej. w ostatnich latach w Polsce obserwuje się spowolnienie dynamiki rozwoju populacji przedsiębiorstw (dynamika rozumiana jako różnica między współczynnikiem „urodzeń” i współczynnikiem „zgonów” przedsiębiorstw) z nieznacznym wzrostem po 2009 roku. w badaniu podjęta zostanie próba oceny wpływu otoczenia rynkowego i zmian zachodzących na rynku na dynamikę populacji przedsiębiorstw. Do oceny dynamiki populacji przedsiębiorstw wykorzystane zostaną współczynniki demograficzne oraz analizy na poziomie makro czyli na poziomie populacji przedsiębiorstw i rynku.

W referacie zostaną przytoczone i omówione podstawowe źródła informacji do analizy przeżycia przedsiębiorstw w Polsce, oraz mocne i słabe strony źródeł takich jak REGON, KRS (Krajowy Rejestr Sądowy), VAT, BJS (Baza Jednostek Statystycznych), i badań ankietowych GUS. Ostatnie zmiany ustawodawcze wskazują na możliwość poprawy aktualności danych w rejestrze REGON jako najbardziej pełnym źródle danych o liczbie i strukturze podmiotów gospodarczych w Polsce.

W referacie przedstawiona zostanie analiza współzależności podstawowych współczynników demograficznych odniesionych do przedsiębiorstw i zmiennych makroekonomicznych. Współczynniki demograficzne jako podstawowe wskaźniki dynamiki rozwoju populacji przedsiębiorstw służą jako miara wyprzedzająca zmiany w cyklu koniunkturalnym dzięki swojej zależności od zmian w czynnikach makroekonomicznych (PKB, Inflacja, Bezrobocie). Analiza przekrojowa uwzględnia kraje europejskie dla których dostępne były dane publikowane przez Eurostat. Dla Polski uwzględniono również zmiany w czasie analizowanych zmiennych. Okres analizy obejmuje lata 1997-2010.

Słowa kluczowe: demografia przedsiębiorstw, wskaźniki makroekonomiczne, statystyka przeżywalności firm

C

onditions of enterprises' market in Poland. Enterprises' survival statistics.

The popularity of business demography in Poland increased in recent years. The growing popularity of business demography was caused by Eurostat project and basic Lisbon strategy goals. Recently in Poland decreasing dynamics of enterprises' population is observed (dynamics defined as difference between „birth” rate and „death” rate of enterprises) with a small increase after 2009. In this paper we try to assess the influence of economic conditions and changes

on the market on enterprises' population dynamics. For assessment of enterprises' population dynamics we use demographic indicators and macroeconomic analysis on the population and economy level.

Basic sources of data for enterprises' survival analyses are discussed in the paper. Also discussion about strengths and weaknesses of sources like REGON, KRS, VAT, BJS and CSO surveys is included. Recent changes show possibility of quality improvements in REGON as basic and full source of data for number and structure of active enterprises in Poland.

Analysis of correlation between demographic indicators for enterprises and macroeconomic indicators was presented in the paper. Demographic indicators as basic indicators of enterprises' population dynamics may be used as a measure forwarded changes in the economic cycle due to correlation with basic economic indicators (GDP, Inflation, Unemployment). Cross-sectional analysis cover European countries for which databases published by Eurostat were available. For Poland also time series were included for analyzed indicators. Analysis covers period 1997-2010.

Key words: business demography, macroeconomic indicators, enterprises' survival statistics

Rogalińska Dominika (Główny Urząd Statystyczny)

Wyzwania statystyki regionalnej na tle debaty o polityce spójności

Wraz z przyjęciem Traktatu Lizbońskiego zmienił się zakres polityki spójności - obok kwestii społecznych i gospodarczych pojawił się wymiar przestrzenny. W ramach toczącej się na ten temat dyskusji konieczne jest silniejsze akcentowanie roli statystyki publicznej jako źródła danych niezbędnych do formułowania celów, monitoringu i ewaluacji polityki.

W wystąpieniu przedstawione zostaną kluczowe kierunki rozwoju polskiej statystyki regionalnej - od przemian struktury organizacyjnej po podejmowane działania. Organizacja statystyki w Polsce stale dopasowuje się do zmieniającego się zapotrzebowania odbiorców danych. Podstawą efektywnej obsługi informacyjnej użytkowników są zintegrowane działania wielu jednostek służb statystyki publicznej. Efektem zmian jest utworzenie Wojewódzkich Ośrodków Badań Regionalnych, których zadaniem jest nie tylko inicjowanie badań regionalnych, ale, przede wszystkim, szybkie reagowanie na zgłaszane potrzeby informacyjne, głównie ze strony jednostek samorządu terytorialnego. Nie można także zapominać o stałym doskonaleniu publicznie dostępnych baz danych, takich jak Bank Danych Lokalnych (BDL). Ostatnio wprowadzone zmiany polegały na udostępnieniu wybranych cech dla poziomu miejscowości statystycznych oraz poszerzeniu zakresu danych krótkookresowych. Znacząco poprawiono także funkcjonalność interfejsu. Należy również wspomnieć o podejmowanych pracach metodologicznych z zakresu statystyki regionalnej. W kontekście monitorowania realizacji założeń strategicznych w ramach prowadzonej polityki rozwoju istotną inicjatywą jest opracowanie nowej typologii i gmin wiej-

skich, która znajduje się już w końcowej fazie. Inne prace dotyczą kwestii i opracowywania wskaźników.

Ważnym elementem upowszechniania informacji i danych statystycznych jest działalność wydawnicza. Wśród wielu publikacji regionalnych warto wymienić Obszary wiejskie w Polsce. Ponadto, drugi już rok z rzędu zostanie wydane Statystyczne Vademecum Samorządowca (SVS), które spotkało się z bardzo dobrymi recenzjami grupy docelowej. Podsumowaniem prezentacji będzie refleksja nad wyzwaniami determinowanymi nie tylko przemianami natury społeczno-gospodarczej czy politycznej, ale również wynikającymi z nowych możliwości technicznych. Propozycje rozwiązań w wielu obszarach tematycznych mogą zostać w istotny sposób wzbogacane dzięki współdziałaniu środowiska naukowego z praktyką statystyczną - dotyczy to zwłaszcza problemów natury meto- dologicznej, bardziej precyzyjnego określenia potrzeb informacyjnych poszczególnych grup odbiorców, a także, co jest chyba najtrudniejsze, szeroko rozumianych prac analitycznych.

Roszak Magdalena, Kołodziejczak Barbara (Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu)

E-ewaluacja wiedzy statystycznej studentów medycyny

Ewaluacja wiedzy statystycznej studentów medycyny prowadzona jest na większości polskich uczelni tradycyjnie, w formie pisemnej na papierze lub ustnej. Obie formy egzaminowania wymagają dużego nakładu pracy od strony merytorycznej (ułożenie kilku różnych zestawów pytań egzaminacyjnych) oraz organizacyjnej (zapewnienie warunków do samodzielnej pracy każdemu studentowi) podczas egzaminu. Po egzaminie sporą ilość czasu wykładowca musi poświęcić na sprawdzenie prac studenckich i przesłanie ocen do wiadomości zdających. W przypadku kilku terminów tego samego egzaminu i egzaminów poprawkowych cała procedura się powtarza. Podobnie jest w kolejnym semestrze i w nowym roku akademickim. Korzystając z mechanizmów dostępnych w portalach e-learningowych można tradycyjny egzamin ze statystyki zamienić na elektroniczny egzamin testowy, który na początku wymagają większego nakładu pracy jednak będzie służył na lata. Właściwą metodykę planowania egzaminu testowego gwarantuje stosowanie standardu Question and Test Interoperability (QTI). Pytania egzaminacyjne mogą być dzielone na grupy o różnym stopniu trudności i z różną punktacją. Standard gwarantuje indywidualizację egzaminu przez losowanie pytań dla każdego studenta oraz wspólne dla egzaminowanych kontrolowanie czasu testu. Nakład pracy niezbędny dla przygotowania bazy pytań testowych z dziedziny rozkłada się na wielokrotnie prowadzoną ewaluację procesu dydaktycznego w ciągu kolejnych lat. Zakładamy, że merytoryczna żywotność bazy pytań z przedmiotu to 5 lat. Całość pracuje z dowolnym portalem e-learnigowym spełniającym międzynarodowe standardy edukacji zdalnej w zakresie egzaminowania. Pytania z multimediami (grafiką, dźwiękiem, filmem czy wzorami np. chemicznymi) są przenaszalne na zewnątrz, w przypadku zmiany portalu. Standard QTI 1.2 zawiera wzorce 4 typów pytań, natomiast standard QTI 2.1 aż 19

wzorców budowy treści edukacyjnych. Bazy pytań mogą być stosowane podczas zaliczeń i egzaminu z przedmiotu oraz posłużyć do powtórek czy self-testów wykonywanych samodzielnie przez studentów w dowolnym czasie i miejscu. Referat będzie prezentował przykładowy egzamin ze statystyki zbudowany w standardzie QTI dla studentów medycyny wraz z propozycją powtórek wykorzystanych jako self-testy. Zostały one opracowane z myślą o wdrożeniu elektronicznego egzaminowania na zajęciach z przedmiotu biostatystyka, oferowanego studentom medycyny. Przedstawione zostaną także opinie studentów z wstępnych badań dotyczących testowego egzaminowania z przedmiotu. Badania te, stanowią przykład wdrożeniem nowatorskich metod ewaluacji w proces kształcenia akademickiego oraz prezentują przesłanki do dalszych działań w kierunku wspomagania nauczania statystyki na uczelni medycznej elektronicznym procesem monitorowania postępów wiedzy studenta.

Słowa kluczowe: biostatystyka, e-learning, egzaminowanie, kształcenie zdalne, kształcenie medyczne

Evaluation of the statistical knowledge of medical students

Evaluation of statistical knowledge of medical students is conducted in the majority of Polish universities traditionally written on paper or oral. Both forms of the examinations are labor intensive and the merits (arrangement of several different sets of exam questions) and organization (providing conditions to work independently to each student) during the exam. After the exam fair amount of time the teacher must devote to check student work and send it to the ratings of candidates. For several periods of the same exam and retake examinations the whole procedure is repeated. Similarly, in the next semester of the new academic year. By using the mechanisms available in e-learning portals can be a traditional exam in statistics by replacing the electronic test exam, which at the beginning requires more work but will serve for years.

The proper method of scheduling an exam ensures the use of standard Question and Test Interoperability (QTI). The test questions can be divided into groups of varying difficulty and with different score. Standard guarantees individualization of the test by drawing questions for each student and the common control test time for test takers. The amount of work necessary to prepare the database of the test questions in the field of distributed on several occasions is conducting an evaluation of the teaching process in the following years. We assume that the viability of the database of questions on the merits of the subject is 5 years. The whole works with any e-learning portal meets international standards of education in the remote examination. Questions from the media (graphics, sound, film, or patterns such as chemicals) are transferable to the outside, if you change the portal. QTI 1.2 standard contains the references of four types of questions, and QTI 2.1 to 19 patterns of the construction of educational content. Base questions may be used during the credits and examination of the subject and used for repetition or self-tests performed independently by students at any time and place.

The paper will present a sample exam with statistics built in QTI standard for medical students, along with the proposal repetition that are used as self-tests. They were designed to implement an electronic examination of the subject in the class of biostatistics, offered to medical students. Presented are the opinions of students from pre-test studies examining the subject. These

studies are an example of implementation of innovative methods of evaluation in the academic process and present evidence for further action to support the teaching of statistics in medical school electronic process for monitoring the progress of student's knowledge.

Key words: biostatistics, e-learning, examination, medical education, virtual learning

BIBLIOGRAPHY

1. Helic D. (12.03.2010), „Template-based Approach to Development of Interactive Tests with IMS Question & Test Interoperability”, <http://coronet.iicm.tugraz.at/denis/pubs/edmedia06a.pdf>.
2. Ren-Kurc A., Roszak M. (2011), „Ewaluacja procesu dydaktycznego. Organizacja egzaminów testowych i ankietowania”, Technologie Informacyjne w Warsztacie Nauczyciela, „Nowe Wyzwania Edukacyjne Implikowane Rozwojem Technologii Informacyjnej”, Wydawnictwo Naukowe Uniwersytetu Pedagogicznego w Krakowie, 253-260
3. Wills G., Davis H., Gilbert L., Hare J., Howard Y., Jeyes S., Millard D., Sherratt R. (10th-11th July 2007), „Delivery of QTIv2 Question Types”, International CAA Conference, Loughborough UK.

Roszka Wojciech (Uniwersytet Ekonomiczny w Poznaniu)

Szacowanie ubóstwa w ujęciu lokalnym na podstawie zintegrowanych źródeł danych

Zwiększające się zapotrzebowanie na rzetelną i aktualną informację na możliwie niskim poziomie agregacji jest rosnącym wyzwaniem dla polskiej statystyki publicznej. Postulaty udostępniania wiarygodnych szacunków na niskim poziomie agregacji przestrzennej wymusza, w dotychczasowym podejściu do badań statystycznych, zwiększanie liczebności próby, co podnosi koszty i czas przeprowadzenia badania. Te trudności, a także zwiększone obciążenie respondentów mogą uniemożliwiać spełnienie postulatów odbiorców komunikatów statystycznych.

Zastosowanie metodologii statystycznej integracji danych umożliwia jednocześnie wykorzystanie informacji pochodzących z różnych źródeł. Za jej pomocą tworzone są zintegrowane zbiory danych o zwiększonej liczebności, co może stwarzać przesłanki do uogólniania wyników na niższym poziomie agregacji przestrzennej. Dodatkowo możliwa jest łączna obserwacja cech obserwowanych w różnych badaniach. Generuje to efekt synergii zwiększający zasób wiedzy pochodzący z badań społeczno-ekonomicznych.

Celem artykułu jest oszacowanie odsetka gospodarstw domowych zagrożonych ubóstwem w ujęciu województw. Do tej pory takie szacunki dostarczane były jedynie w ujęciu makroregionów

(grup województw - poziom NUTS 1). Cel zostanie osiągnięty poprzez integrację zbiorów danych Badania Budżetów Gospodarstw Domowych i oraz Badania Dochodów i Jakości Życia (EU-SILC) z 2005 roku. Jakość szacunków oceniona zostanie w badaniu symulacyjnym.

Słowa kluczowe: statystyczna integracja danych, wielokrotna imputacja, metoda reprezentacyjna

Estimation of poverty on local level based on integrated data sources

Increasing demand for reliable and current information at the lowest possible level of aggregation is a growing challenge for the Polish public statistics. Postulates to provide reliable estimates at a low level of spatial aggregation forces, in the current approach to the socio-economical statistics, increasing the sample size, which escalates the cost and time of performing the survey. These difficulties, as well as an increased burden on respondents, may prevent the fulfillment of demands of customers of statistical information.

Application of statistical data integration methodology allows simultaneous use of information from different sources. That allows to create integrated data sets with increased number of units, which may predispose the conditions to generalize the results at a lower level of spatial aggregation. In addition, it is possible to observe jointly the characteristics from different studies. This creates a synergy effect which increases knowledge resources derived from socio-economic research.

The aim of this article is to estimate the percentage of households at a risk of poverty on a voivodship level. Until now, such estimates were provided only on the NUTS 1 level. This will be achieved through the integration of the following data sets: Household Budget Survey and Survey on Income and Living Conditions (EU-SILC) from 2005. The quality of the estimates will be assessed in a simulation study.

Key words: statistical data integration, multiple imputation, sample survey

Rozkrut Dominik (Urząd Statystyczny w Szczecinie)

Differentiation of innovation strategies

Appropriate indicators that can capture different aspects of innovation are crucial from the point of view of policy-making and policy evaluation. Classical indicators, based on results of innovation surveys, constructed using a single variables such as the "innovation rate", are of limited information capacity. Since innovation is a multidimensional process, application of exploratory data analysis may give additional insight into its nature. Especially, these methods may be applied to reveal so called "innovation modes", i.e. common innovation-related patterns

across firms. When revealed, these empirical modes are useful concepts in analysis of actual business models both in micro and macro scale.

The problem may be described by means of example of the results of the innovation studies for Europe. These based on the innovation survey found that percentage of firms in Portugal that introduced a technological innovation between 1998 and 2000, was similar to that of Finland. Such a result may be considered doubtful as Finland performs very well on such indicators related to R&D, and a range of other innovation measures. In reality, a significantly larger percentage of Finnish firms innovate through in-house creative capabilities, as Portuguese firms more often turn out to be technology adopters or at least modifiers. The problem is that the simple indicators are too aggregated, averaging out different methods of innovation with different innovation outputs and outcomes. These simple indicators that combines all information, regardless of the way firms innovate turn out to be oversimplified.

The goal of the study is to effectively exploit the potential of innovation studies to produce disaggregated indicators that identify how firms innovate. As the extent to which different practices are adopted by firms is varying, the study tries to shed light on this by applying multidimensional statistical analysis. Application of multidimensional exploratory techniques like factor analysis used here raise the potential of revealing modes of innovation. The research reveals differences in innovation practices observed on the regional level when compared with national results. The spectrum of observed innovation practices seems to be wider in case of national dataset, at least for innovation active firms. The interpretation of underlying modes of innovation activity increases our understanding of what innovation strategies are prevalent.

Rozmus Dorota (Uniwersytet Ekonomiczny w Katowicach)

Podejście zagregowane w taksonomii

Stosując metody taksonomiczne w jakimkolwiek zagadnieniu klasyfikacji ważną kwestią jest zapewnienie wysokiej poprawności wyników grupowania. Od niej bowiem zależeć będzie skuteczność wszelkich decyzji podjętych na ich podstawie. Stąd też w literaturze wciąż proponowane są nowe rozwiązania, które mają przynieść poprawę dokładności grupowania w stosunku do tradycyjnych metod (np. k-średnich, metod hierarchicznych). Przykładem mogą tu być metody polegające na zastosowaniu podejścia zagregowanego (łączenia wyników uzyskanych w wyniku wielokrotnego grupowania).

Stabilność algorytmu taksonomicznego w odniesieniu do niewielkich zmian w zbiorze danych (np. wybór podzbioru obiektów, nieduże zmiany w wartościach cech), czy też parametrów (np. losowa inicjalizacja) jest pożądaną cechą algorytmu. Badania empiryczne pokazują, że podejście zagregowane przynosi bardziej stabilne rozwiązania niż pojedyncze algorytmy taksonomiczne.

Głównym celem tego referatu jest zaprezentowanie głównych zagadnień związanych z podej-

ściem zagregowanym w taksonomii. Szczegółowo zaprezentowanych zostanie kilka wybranych metod wraz z empirycznymi wynikami ich zastosowań.

Słowa kluczowe: taksonomia, podejście zagregowane

Ensemble approach in taxonomy

High accuracy of the results is very important task in any grouping problem (clustering). It determines effectiveness of the decisions based on them. Therefore in the literature there are proposed methods and solutions that main aim is to give more accurate results than traditional clustering algorithms (e.g. k-means or hierarchical methods). Examples of such solutions can be cluster ensembles.

A desirable quality of any clustering algorithm is also stability of the method with respect to small perturbations of data (e.g. data subsampling, small variations in the feature values) or the parameters of the algorithm (e.g. random initialization). Empirical results shown that cluster ensembles are more stable than traditional clustering algorithms.

The main aim of this research is to present main ideas and concepts connected with cluster ensemble approach. Few particular algorithms will be presented together with some empirical results of application of these methods.

Key words: taxonomy, cluster ensemble

BIBLIOGRAPHY

1. DUDOIT S., FRIDLYAND J.: Bagging to Improve the Accuracy of a Clustering Procedure, „Bioinformatics”, 2003, Vol. 19, No. 9, 1090-1099.
2. FRED A., JAIN A.K.: Data clustering using evidence accumulation, „Proceedings of the 16th International Conference on Pattern Recognition”, 2002, 276-280.
3. LEISCH F.: Bagged Clustering, „Adaptive Information Systems and Modeling in Economics and Management Science”, 1999, Working Paper 51.

**Rozpędowska-Matraszek Danuta (Państwowa Wyższa Szkoła Zawodowa w Skier-
niewicach)**

Ocena skuteczności funkcjonowania opieki zdrowotnej w Polsce - analiza według grup powiatów podobnych

Zgodnie z istniejącym prawem, ludność naszego kraju ma zapewniony powszechny i równy dostęp do opieki zdrowotnej. Wymiar przestrzenny jest zazwyczaj drugoplanowy. Funkcjonowanie systemu opieki zdrowotnej w powiatach pozwala określić rolę, jaką pełni system ochrony

zdrowia w regionie, w celu właściwego zapewnienia usług w zakresie opieki medycznej. Jest to również punkt wyjścia dla analizy prezentowanej w opracowaniu. Podmiotami tej analizy są powiaty w Polsce. Poziom ich rozwoju gospodarczego jest jak wiadomo zróżnicowany. Wyodrębnienie grup powiatów o podobnych cechach wchodzących w skład funkcjonującego systemu ochrony zdrowia, pozwala na ocenę skuteczności funkcjonowania systemu. Zasadniczym celem opracowania jest wyodrębnienie grup powiatów podobnych ze względu na poziom rozwoju opieki zdrowotnej oraz statystyczno-ekonometryczna analiza czynników, które w istotny sposób wpływają na dostęp do opieki zdrowotnej w tych grupach. Jako narzędzia analizy zastosowano metody taksonomiczne i statystyczno-ekonometryczne. W szczególności są to aglomeracyjne metody grupowania i modele ekonometryczne szacowane na podstawie danych panelowych. Badanie empiryczne, jako kontynuacja badań, przeprowadzono na podstawie danych rocznych GUS z lat 1999-2010 dla grup powiatów i listy zmiennych, które mogą charakteryzować powiaty pod względem funkcjonowania systemu opieki zdrowotnej.

Słowa kluczowe: analiza regionalna, opieka zdrowotna, metody taksonomiczne, modele panelowe

The assessment of health care system operation efficiency based on the example of the Poland - an analysis using clustering of poviats

Pursuant to applicable law, our country's population is offered universal and equal access to health care. The functioning of the health care system in poviats allows to determine the role played by the health care system in the region in order to properly provide health care services. It is also the starting point for an analysis presented in the study. Subjects of that analysis are poviats of the Poland. The poviats are known to differ in their economic development levels. Isolating groups of poviats in the region with similar characteristics, being part of the functioning health care system, makes it possible to assess the functioning efficiency of the system. The main objective of the study is to isolate groups of poviats with similar health care development levels and perform an econometric analysis of factors that significantly affect access to health care in those groups. As its tools, the analysis uses taxonomic and econometric methods. Those methods include, in particular, agglomerative clustering methods and econometric models estimated on the basis of panel data. Examination of the empirical, as a continuation of the research based on Central Statistical Office data from 1999 to 2010 for groups describe the poviats they are characterised considering of functioning of the health care system.

Key words: regional analysis, health care, taxonomic analysis, panel models

BIBLIOGRAPHY

1. Dańska-Borsiak B., (2011), Dynamiczne modele panelowe w badaniach ekonomicznych, Wydawnictwo Uniwersytetu Łódzkiego, Łódź.
2. Jajuga K., (1998), Ekonometria. Metody i analiza problemów ekonomicznych, Wyd. AE Wrocław.
3. Kwiatkowski E. (red.), (2008), Zróżnicowanie rozwoju polskich regionów. Elementy teorii i próba diagnozy, Wydawnictwo Uniwersytetu Łódzkiego, Łódź.

4. Rozpędowska-Matraszek D., (2012), The Assessment of Health Care System Operation Efficiency Based on the Example of Poviats of the Łódź and West Pomeranian Voivodships [in:] Folia Oeconomica Stentinensia (EC:) W. Tarczyński, Publisher Versita, University of Szczecin, Szczecin.

Ryś-Jurek Roma (Uniwersytet Przyrodniczy w Poznaniu)

Zróźnicowanie regionalne produkcji rolniczej w krajach UE-27

Wspólna Polityka Rolna dotychczas nie zmniejszyła istniejących różnic pomiędzy poszczególnymi krajami członkowskimi UE-27. Dlatego pilną potrzebą jest ustalenie obszarów rolniczych o podobnej strukturze i skali produkcji rolniczej, w tym miejsca polskiego rolnictwa w UE-27. Podmiotem badań będzie przeciętne gospodarstwo rolne. Celem opracowania jest próba delimitacji regionów w UE-27 o podobnej strukturze i skali produkcji pochodzącej z gospodarstw rolnych. W badaniu położono nacisk na uwypuklenie podobieństwa struktur produkcji typowych dla obszarów klimatycznych UE, a także wskazanie miejsca polskiego rolnictwa wśród poszczególnych krajów UE-27 w 2008 roku.

Do przeprowadzenia badań zastosowane będą dane z sieci danych rachunkowości rolnej gospodarstw rolniczych FADN (ang. Farm Accountancy Data Network). Zawiera ona średnie ważone przeliczone na gospodarstwo rolne z 27 krajów członkowskich UE w 2008 roku. Przedmiotem badań będzie delimitacja regionów w UE-27 o podobnej strukturze i skali produkcji pochodzącej z gospodarstw rolnych wyznaczona na podstawie struktury produkcji, wielkości ekonomicznej wyrażonej w ESU, powierzchni użytków rolnych (UR), produkcji i dochodów z gospodarstwa rolnego. Badanie to wykonane będzie za pomocą hierarchicznej klasyfikacji aglomeracyjnej metodą Warda. Strukturę produkcji reprezentują 22 zmienne, a pozostałe cechy to pojedyncze zmienne. Metoda Warda charakteryzuje się dużą efektywnością i wyróżnia się tworzeniem skupień o małej liczebności [Walesiak 2004, s. 322-324]. Prowadzi do wyznaczenia skupień o zbliżonych liczebnościach, charakteryzujących się minimalną wariancją, dlatego też często jest wykorzystywana do klasyfikacji jednostek przestrzennych [Filipiak 2006, s. 57].

Optymalna liczba klas krajów o podobnej strukturze i skali produkcji wyznaczona będzie na podstawie analizy wielkości przyrostów wariancji wewnątrzklasowej w procesie łączenia klas między poszczególnymi poziomami łączy [Wysocki 2010, s. 128]. Dodatkowo podczas badań zastosuje się analizę opisową, porównawczą oraz podstawowe metody statystyki opisowej (w tym: średnią harmoniczną i standaryzacją cech prostych).

Słowa kluczowe: produkcja rolnicza, metoda Warda, wariancja

Regional differences in agricultural production in the EU-27 countries

The Common Agricultural Policy has not yet diminished the existing differences between the country member of the EU-27. Therefore, the urgent need is to determine the agricultural areas with similar structure and scale of the agricultural production, and also the place of Polish agriculture in the EU-27. The subject of the research will be the average farm. The aim of the article is an attempt of regions' delimitation in the EU-27 with similar structure and scale of the production from farms. In this study, the emphasis is placed on highlighting the similarity of production structures typical of the climate's areas of the EU, and also on indicating the place of Polish agriculture among the countries from the EU-27 in the year 2008 roku.

For this research, the data from FADN database (ang. Farm Accountancy Data Network) will be used. This database contains weighted averages per individual farm in every of 27 UE member countries in 2008. The aim of the research will be the delimitation of regions in the UE-27 distinguished by similar structure and scale of production achieved by agricultural farms, implemented on the basis of the structure of production, economic size expressed in terms of the ESU, production and income of agricultural farm. Research will be conducted using the agglomeration hierarchic classification based on the Ward's method. The structure of production is represented by 22 variables; the remaining characters are represented by single variables. The distinctive feature of the Ward's method is its high efficiency, observed in the process of connecting classes between different levels of connections [Wysocki 2010, s. 128]. This allows to achieve the small clusters [Walesiak 2004, s. 322-324] of similar number of elements and minimal variance. That's why the Ward's method is often used when classifying spatial items [Filipiak 2006, s. 57].

The optimal number of classes of countries of similar structure and scale of production will be estimated on the basis of the analysis of the amount of the intra-class variance's increment. While conducting research, the descriptive analysis will be used, as well as comparative analysis and the basic methods of the descriptive statistics (including harmonic average and standardization of variables).

Key words: agricultural output, Ward's method, variance

BIBLIOGRAPHY

1. FADN (2011): http://ec.europa.eu/agriculture/rca/database/database_en.cfm,do step: listopad2011.
2. Filipiak K., Metody statystyczne stosowane do oceny regionalnego zróżnicowania rolnictwa, [w:] Regionalne zróżnicowanie produkcji rolniczej w Polsce, Raporty PIB nr 3, red. A. Harasim, IUNG - PIB, Puławy 2006: 53-60.
3. Walesiak M., Metody klasyfikacji, [w:] Metody statystycznej analizy wielowymiarowej w badaniach marketingowych, red. E. Gatnar i M. Walesiak, Wydawnictwo Akademii Ekonomicznej im. O. Langego we Wrocławiu, Wrocław 2004: 316-350.
4. Wysocki F., Metody taksonomiczne w rozpoznawaniu typów ekonomicznych rolnictwa i obszarów wiejskich, Wydawnictwo Uniwersytetu Przyrodniczego w Poznaniu, Poznań 2010.

Siwiński Włodzimierz (Uniwersytet Warszawski, Akademia L. Koźmińskiego)

Statystyczny obraz kryzysu gospodarczego 2008–2010

Ostatni największy powojenny kryzys gospodarczy był powszechnym zaskoczeniem. Dotyczy to może nie tyle samego wystąpienia kryzysu, które przecież są stałym zjawiskiem cyklicznie pojawiającym się w gospodarce kapitalistycznej, co raczej niespotykanej po wojnie skali i gwałtowności załamania oraz faktu wystąpienia w krajach najwyżej rozwiniętych. Referat dokumentuje zawodność prognoz opracowywanych przez najważniejsze instytucje międzynarodowe monitorujące sytuację makroekonomiczną gospodarki światowej i głównych grup krajów. Przykładem są prognozy Międzynarodowego Funduszu Walutowego. Prognozy opracowywane już w czasie trwania kryzysu finansowego prawie do połowy IV kwartału 2008 r. wskazywały, że w 2009 r. nastąpi jedynie spowolnienie wzrostu gospodarczego w krajach rozwiniętych. Pierwsze prognozy wskazujące na absolutny spadek PKB w tej grupie krajów pojawiły się dopiero pod koniec IV kwartału 2008 r i właściwie do końca pierwszego kwartału 2009 r daleko odbiegały od rzeczywistej skali załamania jaka w tym roku wystąpiła.

Problem zawodności w prognozowaniu obecnego kryzysu w jakimś stopniu jest następstwem pogłębiającej się luki informacyjnej dotyczącej rynków finansowych ograniczającej monitorowanie ewentualnego ryzyka wystąpienia nierównowag w gospodarce światowej i w poszczególnych krajach. Tak więc, statystyczny obraz narastających problemów gospodarczych jest wysoce nieadekwatny, co utrudnia nie tylko monitorowanie gospodarki ale przede wszystkim podejmowanie właściwych działań ewentualnie korygujących narastanie nierównowag grożących kryzysem.

Luka informacyjna dotyczy przede wszystkim statystyki finansowej i jest następstwem dynamicznego rozwoju instytucji oraz mechanizmów funkcjonowania rynków finansowych. Jednym z głównych problemów, który wywołał głębokie załamanie produkcji był krach finansowy objawiający się bardzo głębokim spadkiem międzynarodowej płynności finansowej. Ten spadek spowodował drastyczne ograniczenie finansowania nawet bieżącej działalności gospodarczej firm i instytucji finansowych wywołując falę bankructw wskutek rzeczywistej niewypłacalności lub nawet tylko czasowej utraty płynności finansowej. Kryzys nie był oczywiście spowodowany brakiem adekwatnej statystyki finansowej ale ujawnił te braki, które nie pozwoliły właściwie ocenić narastania ryzyka systemowego (ryzyka krachu finansowego) i ewentualnie podjąć właściwych działań zaradczych.

Referat koncentruje się na źródłach i ocenie nieadekwatności statystycznego obrazu ryzyka systemowego rynków finansowych. Trudności w ocenie systemowego ryzyka wynikają z ograniczonych możliwości pomiaru poszczególnych składników płynności finansowej, która ma dzisiaj globalny charakter. Trudności te są następstwem zmian instytucjonalnych oraz mechanizmów funkcjonowania systemu finansowego. Chodzi np. o trudności w ocenie rzeczywistych rozmiarów zagregowanej dźwigni finansowej, trudności w ocenie ryzyka wynikającego z często krańcowego niedostosowania struktury czasowej pasywów i aktywów instytucji finansowych zwłaszcza w seg-

mentach rynku finansowego, w których instytucje nie podlegają żadnym lub bardzo niewielkim regulacjom (segment tzw. para bankowości lub bankowości alternatywnej). Ocena ryzyka z tym związana jest szczególnie utrudniona wobec lawinowego rozwoju nowych instrumentów finansowych (tzw. inżynieria finansowa), które transferują ryzyko pomiędzy różnymi podmiotami, co niezwykle utrudnia jego rzeczywistą ocenę. W referacie analizowane jest rola rozwijającej się bardzo szybko tzw. sekurytyzacji i związanego z tym "przepakowywania" ryzyka oraz rola niezwykle dynamicznego rozwoju zupełnie nieregulowanego rynku nowych instrumentów finansowych, zwłaszcza swapów na zwłokę w spłacie kredytów (CDS *credit default swaps*), które stały się źródłem ryzyka systemowego zagrażając stabilności całego systemu finansowego.

Rozwiązanie problemu luki informacyjnej w zakresie światowych finansów jest obecnie wielkim wyzwaniem dla międzynarodowych instytucji regulujących międzynarodowe finanse, dla rządów krajowych i banków centralnych. Od tego bowiem będzie zależało ewentualne wzmocnienie regulacji rynków finansowych mające zapobiegać powstawaniu tak gwałtownych i głębokich kryzysów finansowych i gospodarczych, jak to obecnie miejsce. Działania te zostały zainicjowane o czym świadczy raport wraz z rekomendacjami podjęcia odpowiednich działań przygotowany przez Radę Stabilności Finansowej (FSB, *Financial Stability Board*), powołaną przez tzw. grupę 20.

Statistical Image of Economic Crisis 2008–2010

The last most severe crisis was a great surprise almost for everybody. It is true may be not for the crisis as such but for its severity and violence but also for the fact that it was developed in most advanced countries. It is evident that all the forecast made by the most eminent international organizations and academic institutions completely failed to anticipate the production collapse in 2009. It is documented in this paper by the forecasts made by IMF: in its reports and updates' In 2008 when the financial crisis was already unfolded the prediction published by IMF anticipated only the slowdown of economic growth but not the collapse. Only the last updates released in the end of forth quarter of 2008 anticipated the decline of GDP in advanced countries in 2009 but in much smaller scale than later occurred.

The failure to forecast the current crisis reflects much deeper problem of growing data gaps related especially to financial sector that limit the possibility to identify the risk of growing mismatches in the financial system yhat may lead to financial crunch and severely downgrade the whole financial system (systemic risk) on global as well as on national levels. Lack of timely and adequate information hinders the ability of policymakers and financial institutions to develop the effective responses to market disturbances.

Data gaps in financial statistics are consequence of the development of markets and institutions. One of the key factor contributing to the production collapse was liquidity crunch. This crunch strongly limited available financing causing for many institutions and companies great difficulties to finance their current activities jeopardizing their surveillance and leading in many cases to bankruptcy due to insolvency. or even only to lack of liquidity. Of course, financial crisis was not a result of lack of proper economic or financial statistics but it highlighted the obvious truth that the lack of proper information and analysis on important financial vulnerabilities

substantially hinder the possibility for assessing the growing systemic risk and eventually to develop effective responses.

The paper discuss the reasons and implications of the financial data gaps. The difficulties in assessing the systemic risk are partly due to the lack of proper statistical data of different parts of financial liquidity (for instance, the willingness of market participant to supply funding or trade in securities markets). These difficulties are also due to institutional developments and changes in operation of financial markets. For instance, there is a lack of proper statistical data of aggregate leverage and maturity mismatches in the financial system especially in the unregulated or weakly regulated segments of financial markets (ex. in the shadow banking sector). The proper capture of development of the risk in financial sector is also difficult due to very dynamic development of new financial instruments: for instance different types of asset-backed commercial papers transferring the risk into different parts of financial markets or especially credit default swaps (CDS) traded in completely unregulated market which strongly contributed to building-up of systemic risk.

Solving the problems related to this information gap is a great challenge for international community including the international organizations (especially IMF, BIS,), central banks and national governments. The improvement in providing of timely and accurate statistical information and analysis is a basic precondition for surveillance and policy responses at both international and national levels. In 2009 the Group of Twenty Finance Ministers and Central Bank Governors call on Financial Stability Board and International Monetary Bank to explore the information gap and provide the appropriate proposal to deal with these problems. The staff from both organizations prepared report which include a list of several recommendations providing the proposals for actions indispensable to fill the information gaps.

Sobczak Elżbieta (Uniwersytet Ekonomiczny we Wrocławiu)

Pracujący wg sektorów intensywności działalności B+R w krajach UE - analiza przestrzenno-strukturalna

Przemiany strukturalne mają charakter wielopłaszczyznowy i mogą być analizowane z wykorzystaniem różnorodnych metod statystycznych. Znajdują one również odzwierciedlenie w przekształceniach sektorowej struktury pracujących na poziomie krajowym. Sektorowa struktura pracujących powszechnie traktowana jest jako jeden z podstawowych wskaźników poziomu rozwoju społeczno-gospodarczego oraz dojrzałości systemu rynkowego.

Celem opracowania jest analiza zróżnicowania i przemian struktury pracujących w państwach członkowskich Unii Europejskiej w latach 2004-2009 w sektorach ekonomicznych wyodrębnionych wg intensywności działalności B+R oraz identyfikacja jednorodnych grup państw. Układ referatu obejmuje:

1. pojęcie sektorów wysokiej techniki,

2. podstawy informacyjne i metodologiczne badań,
3. wyniki klasyfikacji państw członkowskich Unii Europejskiej ze względu na sektorową strukturę pracujących .

Przeprowadzona analiza statystyczna prowadzi do klasyfikacji państw członkowskich UE ze względu na strukturę pracujących wg intensywności działalności B+R z wykorzystaniem metod analizy skupień (Cluster Analysis) w latach 2004, 2007 i 2009.

W celu klasyfikacji państw UE przeprowadzono poniższy schemat postępowania:

- określenie zróżnicowania między badanymi krajami za pomocą kwadratu odległości euklidesowej,
- klasyfikacja hierarchiczna krajów na homogeniczne klasy z wykorzystaniem metody Warda,
- wstępna wielowariantowa propozycja dotycząca liczby klas na podstawie analizy wstępnych wyników klasyfikacji przedstawionych na dendrogramie oraz wykresie odległości wiązania względem etapów wiązania,
- wielowariantowa klasyfikacja krajów metodą k-średnich,
- wybór klasyfikacji optymalnej z wykorzystaniem wskaźnika jakości klasyfikacji Calińskiego-Harabasa
- przedstawienie składu otrzymanych klas państw i ich charakterystyka,
- ocena jednorodności otrzymanych klas regionów z wykorzystaniem odległości międzyklasowej oraz odległości poszczególnych regionów od środków ciężkości klas,

Otrzymane wyniki badań pozwalają na określenie specyfiki klas państw wyodrębnionych ze względu na strukturę pracujących, jak również ocenę zmian strukturalnych jakie miały miejsce w okresie 2004-2009.

Słowa kluczowe: struktura pracujących, sektory wg intensywności działalności B+R, metody analizy skupień

Workforce by R&D activity intensity in EU countries - space-structural analysis

Structural changes are of multilevel nature and may be analyzed by means of different statistical methods. They are also reflected in transformations referring to sector structure of workforce at national level. Sector structure of workforce is commonly referred to as one of basic indicators for social and economic development level, as well as the maturity of market system.

The objective of the hereby study is to analyze sector structure of workforce diversification and transformation in EU countries in the years 2004-2009 in economic sectors by R&D activity intensity, as well as the identification of homogenous groups of countries. The study includes:

1. concept of high-tech sectors ,
2. information and methodology background for research,
3. classification results of EU countries with regard to sector structure of workforce.

The conducted statistical analysis leads up to classification EU member countries according to workforce sectors structure by R&D activity intensity applying methods of cluster analysis in the years 2004, 2007 and 2009.

In order to classify EU countries the following procedure was performed:

- specifying diversification between studied countries by means of Squared Euclidean Distance,
- hierarchical classification of countries into homogenous groups by means of Ward method,
- initial multivariate position referring to the number of groups of countries,
- multivariate classification of countries by means of k-means method,
- selection of optimal classification applying Caliński-Harabasz classification quality indicator,
- presentation of the obtained groups of countries composition and their characteristics by applying basic descriptive parameters,
- assessment of the obtained groups of countries homogeneity using intergroup distance and the distance of particular countries from groups' gravity centres.

The obtained research results facilitate defining the specific nature of groups of countries distinguished with regard to sector structure of workforce, as well as the assessment of changes observed in 2004-2009.

Key words: workforce structure, sectors by R&D activity intensity, cluster analysis methods

BIBLIOGRAPHY

1. Anderberg M.R., Cluster Analysis for Application, Academic Press, New York, San Francisco, London, 1973.
2. Brdulak H., Gołębiowski T., Rola innowacyjności w budowaniu przewagi konkurencyjnej [w:] H. Brdulak, T. Gołębiowski (red.), Wspólna Europa innowacyjność w działalności przedsiębiorstw, Difin, Warszawa 2003.
3. Caliński R.B., Harabasz J., a Dendrite Method for Cluster Analysis, "Communications in Statistics" 1974, nr 3, s. 1-27.

4. Dunn E.S., a statistical and analytical technique for regional analysis, Papers of the Regional Science Association, 1960, nr 6, s. 97-112.
5. Hartigan J.A., Clustering Algorithms, John Wiley & Sons, New York 1975.
6. Hatzichronoglou T., Revision of the High-Technology Sektor and Produkt Classification, OECD, Paris 1996.
7. Nauka i technika w 2007r., Informacje i opracowania statystyczne, GUS, Warszawa 2009.
8. Perloff H.S., Dunn E.S., Lampard E.E., Mutha R.F., Regions, resources and economic growth, John Hopkins Press, Baltimore 1960.
9. Pocięcha J., Podolec B., Sokołowski A., Zając K., Metody taksonomiczne w badaniach społeczno-ekonomicznych, PWE, Warszawa 1988.
10. Sneath P.H., Sokal R.R., Numerical Taxonomy, Freeman, San Francisco 1973.
11. Stern S., Porter M., Furman J.L., The Determinants of National Innovative Capacity, National Bureau of Economic Research, Working Paper No 7876, September 2000.
12. Weresa M. A., Zdolność innowacyjna polskiej gospodarki; pozycja w świecie i regionie konkurencyjnej [w:] H. Brdulak, T. Gołębiowski (red.), Wspólna Europa innowacyjność w działalności przedsiębiorstw, Difin, Warszawa 2003.
13. Wojnicka E. (red.) Perspektywy rozwoju małych i średnich przedsiębiorstw wysokich technologii w Polsce do 2020 roku, Polska Agencja Rozwoju Przedsiębiorczości, Warszawa 2006.

Sokołowski Andrzej (Uniwersytet Ekonomiczny w Krakowie), Basiura Beata (Akademia Górniczo-Hutnicza w Krakowie)

Taksonomiczna metoda Warda

Aglomeracyjna metoda Warda to jedna z najpopularniejszych metod taksonomicznych. Nasz referat ma charakter głównie historyczny. Pokazujemy pierwsze publikacje w których zaproponowano tę metodę i jej podstawowy, oryginalny opis. W obliczu pojawiających bardzo wielu propozycji nowych metod taksonomicznych, w latach 80-tych ubiegłego wieku kilku autorów podjęło próby ich oceny poprzez badania symulacyjne. W tych porównaniach metoda Warda wypadła na ogół bardzo dobrze. Do jej popularności przyczyniła się też pewna właściwość powodująca, że otrzymywane grupy miały dość zrównoważone liczebności i praktycznie nie otrzymywano podziałów zawierających wiele grup jednoelementowych.

W pracy przedstawiamy trzy kierunki w których następowały udoskonalania i ewolucja metody Warda:

- odejście od oryginalnego postulatu minimalizacji wariancji wewnątrzgrupowej na rzecz minimalizacji sumy odległości wewnątrzgrupowych, co pozwala na wykorzystywanie różnych miar odległości (a nie tylko kwadratu odległości euklidesowej)
- próby obiektywizacji decyzji wyboru liczby podgrup, czyli cięcia dendrogramu
- badania stochastycznych własności metody dla zadanych modeli generujących mieszanki.

Podsumowaniem będzie wskazanie zalet metody Warda, szczególnie w porównaniu z innymi metodami aglomeracyjnymi analizy skupień, ale również niedoskonałości tej metody.

Stanimir Agnieszka (Uniwersytet Ekonomiczny we Wrocławiu)

Zróznicowane techniki prezentacji zmiennych niemetrycznych

Zmienna nominalna jest to zmienna zmierzona na najniższej skali pomiaru. Jest ona opisana dwiema lub kilkoma rozłącznymi kategoriami. Jeśli w badaniu wystąpi obiekt zmierzony na tego typu skali to można go przyporządkować tylko do jednej z kategorii danej zmiennej. Kategorie są uznawane jako symbole a nie liczby. Dla zmiennych zmierzonych na skali nominalnej można wykonać niewiele działań matematycznych. Najczęściej dokonuje się zliczenia wystąpień kategorii i na tej podstawie wyznaczenie częstości lub proporcji. Miarą położenia jest modalna. Zmienne porządkowe to druga grupa zmiennych, które zalicza się do niemetrycznych. Skala porządkowa powstaje przez wzbogacenie skali nominalnej o relację porządku (dla kategorii). Między kategoriami można ustalić relację mniejszości lub większości jednak nie jest możliwe określenie wielkości różnic między kategoriami. Miarą położenia jest mediana.

Celem prezentacji jest wskazanie graficznych rozwiązań poszukiwania interakcji między zmiennymi. Z tego względu należy wskazać, że między wyróżnionymi skalami występuje kumulatywność działań. Działania dopuszczalne na skali nominalnej można wykonać na skali porządkowej, a pomiar wykonany na skali porządkowej można przekształcić na pomiar wykonany na skali nominalnej. Zainteresowanie analityka w badaniach statystycznych, jak również z innych dziedzin, często koncentruje się wokół rozpoznawania zależności między wieloma zmiennymi, a nie tylko wokół jednowymiarowej analizy zmiennych. Powszechnie znane są współczynniki mierzące zależności między dwiema zmiennymi nominalnymi oparte na statystyce χ^2 , np. ϕ -Yule'a, V Cramera, Q Kendalla, lub zmiennymi porządkowymi: korelacji rang Spearmana, korelacji rang Kendalla, Konkordancji Kendalla. Po zastosowaniu tych miar zależności pozostaje jednak pytanie czy występują jakieś relacje między kategoriami zmiennych? Wnioski o dogłębnej analizie powiązań między kategoriami dwóch lub wielu zmiennych nominalnych można wyciągnąć na podstawie: analizy korespondencji (zarówno w podejściu klasycznym jak i dla wielu zmiennych), wykresów mozaikowych (konstruowanych na podstawie tablicy kontyngencji lub wielowymiarowej tablicy kontyngencji) oraz wykresów czteropolowych (konstruowane dla tablic kontyngencji 2×2).

W celu zobrazowania proponowanych graficznych metod prezentacji zmiennych niemetrycznych wykorzystano wyniki egzaminu gimnazjalnego przeprowadzonego w 2010r. Dane, które wykorzystano dotyczą wyników z dwóch części i 6 obszarów egzaminu, płci, oraz podziału terytorialnego obszaru Dolnego Śląska i Opolszczyzny.

Słowa kluczowe: zmienne niemetryczne, wykresy mozaikowe, analiza korespondencji

Different techniques of graphical presentation of categorical data

Nominal scale is the weakest of measurement scales. Variables measured on this scale are described by two or more mutually exclusive categories. If the object of study is measured on this scale, that can be assigned to only one of the categories of the variable. Categories of nominal variables shall be recognized as symbols rather than numbers. For variables measured on a nominal scale permit only the most rudimentary of mathematical operations. The most popular is counting of numbers of categories. Ordinal variables are the second group of variables that are classified as categorical. Ordinal scale is created by a nominal scale with category ordering. This scales require the relationship of minority or majority between categories, however, with undetermined size of the differences between categories. The positional measure is the median.

The aim the presentation is to show graphical solutions, studying the interaction between variables. For this reason, it should be noted that mathematical operations used on lower scales may be used on higher scales - it is cumulating of mathematical operations. Thus, the measurement performed on an ordinal scale can be converted to a measurement made on a nominal scale. The interest in statistical research analyst, as well as in other fields, often revolves around the recognition of relationships between many variables, not just around the univariate analysis of variables. Commonly known are the coefficients measuring the relationship between two nominal variables based on a χ^2 statistics, such as Yule's ϕ , Cramer's V, Kendall Q, or ordinal variables: Spearman rank correlation coefficient, Kendall's rank correlation coefficient, Kendall's coefficient of concordance. After applying these measures, however, remains one question. Whether there are any relationships between categories of variables? Results of in-depth analysis of the relationship between two or more categories of nominal variables can be drawn from: analysis of correspondence (both in the classical approach, as for many variables), mosaic plots (constructed on the basis of a multidimensional contingency table or contingency table) and a four-fold diagrams (constructed for contingency tables 2×2).

The Gymnasium Examination results are taken to illustrate the proposed graphical methods of presentation of CATEGORICAL variables. Data refer to the results of the two parts and 6 areas of examination, gender, and the territorial division of the Lower Silesia Region and Opolskie Voivodship.

Key words: categorical data, mosaics displays, correspondence analysis

Stańczyk Elżbieta (Urząd Statystyczny we Wrocławiu)

Kwalifikacje zawodowe osób bezrobotnych w świetle danych GUS. Międzywojewódzka analiza porównawcza

Celem referatu jest porównawcza analiza w układzie przestrzennym stanu i struktury osób bezrobotnych według posiadanych kwalifikacji, tj. według wykształcenia i doświadczenia zawodowego - według wybranych grup zawodu ostatniego miejsca pracy.

Badaniem objęto bezrobotnych, którzy posiadają kwalifikacje do wykonywania jakiegokolwiek zawodu poświadczonych dyplomem, świadectwem, zaświadczeniem instytucji szkoleniowej lub innym dokumentem uprawniającym do wykonywania zawodu.

Szczególną uwagę zwrócono w referacie na klasyfikację zawodów i specjalności na potrzeby rynku pracy (KZiS) oraz zakresu jej stosowania w badaniach prowadzonych przez GUS. KZiS jest pięciopoziomowym, hierarchicznie usystematyzowanym zbiorem zawodów i specjalności występujących na rynku pracy. Struktura klasyfikacji oparta jest na systemie pojęć, z których najważniejsze to: zawód, specjalność, umiejętności oraz kwalifikacje zawodowe.

W klasyfikacji, uwzględniono cztery szerokie poziomy kwalifikacji, które zdefiniowano w odniesieniu do poziomów wykształcenia określonych w Międzynarodowej Klasyfikacji Standardów Edukacyjnych (ISCED 97).

W opracowaniu wykorzystano dane Ministerstwa Pracy i Polityki Społecznej dotyczące bezrobotnych zarejestrowanych w urzędach pracy oraz dane uzyskane z reprezentacyjnego Badania Aktywności Ekonomicznej Ludności (BAEL).

Słowa kluczowe: Zarejestrowany bezrobotny, oferty pracy, deficytowe i nadwyżkowe zawody, specjalność, Klasyfikacja Zawodów i Specjalności, Krajowe Ramy Kwalifikacji

Occupational qualifications of the unemployed on the basis of research by CSO. Interoivodship comparative analysis

The aim of this paper is comparative analysis in spatial arrangements of the state and structure of the unemployed by education, length of service and professional experience - by selected occupational groups of their last job.

The research has been conducted among unemployed with occupational qualifications - persons who do have qualifications for performing any occupation confirmed by a diploma, certificate, certificate issued by a training institution or any other document entitling them to perform an occupation.

Special attention was paid to Polish Classification of Occupations and Specializations for the needs of Labour Market (Klasyfikacja Zawodów i Specjalności - KZiS, published in 2010) and

the scope of its application in research done by CSO.

KZiS organises occupations in a hierarchical framework. They are based on two main concepts: the concept of kind of work performed, the concept of skill, and skill specialization. Four skill levels are defined at the most aggregate level, the major groups. These four skill levels are defined in terms of the educational categories and levels of the International Standard Classification of Education (ISCED 97).

The paper presents data on the unemployed from the Ministry of Labour and Social Policy. There are unemployed persons registered in the labour offices. Data from the sample research Labour Force Survey (LFS) were also included in the publication.

Key words: Unemployed persons, job offers, shortage and surplus occupations, specialization, Polish Classification of Occupations and Specializations, National Qualifications Framework

BIBLIOGRAPHY

1. Klasyfikacja zawodów i specjalności dla potrzeb rynku pracy, opracowanie: Strojna, E. , Piotrowicz, J. Żywiec-Dąbrowska, E. , Ministerstwo Pracy i Polityki Społecznej, Departament Rynku Pracy, Warszawa, 2010,
2. Rozporządzenie Ministra Pracy i Polityki Społecznej z dnia 27 kwietnia 2010 r. W sprawie klasyfikacji zawodów i specjalności na potrzeby rynku pracy oraz zakresu jej stosowania, Dz. U. Z dnia 17 maja 2010 r., Nr 82 poz. 537,
3. Suknarowska-Drzewiecka, E., Wojciechowski, P., Podnoszenie i uzupełnianie kwalifikacji zawodowych pracowników, wyd. C. H. Beck, Warszawa, 2010,
4. Zapotrzebowanie na kwalifikacje i szkolenia pracowników przedsiębiorstw z Dolnego Śląska, red. nauk. Kupczyk, T., wyd. Wyższa Szkoła Handlowa, Wrocław, 2011.

Stonawski Marcin (Uniwersytet Ekonomiczny w Krakowie, IIASA)

Kapitał ludzki w warunkach starzenia się ludności w Polsce

Kapitał ludzki należy do najważniejszych zasobów warunkujących procesy, które kształtują poziom i jakość życia ludności. Współcześnie symptomem znaczenia kapitału ludzkiego jest powstawanie gospodarek „opartych na wiedzy”. Odpowiednio wykorzystany, wysokiej jakości kapitał ludzki jest równocześnie wartościowym czynnikiem wytwórczym i podstawowym źródłem przewagi konkurencyjnej na świecie.

Poziom kapitału ludzkiego kształtuje się pod wpływem różnorodnych czynników. Wśród nich szczególne miejsce zajmują zjawiska demograficzne, takie jak rozrodczość, umieralność i ruchy

wędrówkowe. Kształtują one bowiem liczebność oraz strukturę populacji zaangażowanej w różnej mierze w procesach ekonomicznych. Proces starzenia się populacji bezpośrednio dotyczy rynku pracy. Starzeją się bowiem zasoby pracy aktualnie obecne na tym rynku oraz zmniejsza się liczebność potencjalnych pracowników.

Celem pracy jest oszacowanie poziomu kapitału ludzkiego oraz ocena wpływu przemian demograficznych na kształowanie się tego zasobu w Polsce w latach 2010-2050.

Oszacowanie kapitału ludzkiego wykonano przy pomocy autorskiej metody opartej na szacunkach akumulacji i deprecjacji, uwzględnia zarówno ilościowy, jak i jakościowy aspekt tych zasobów. Istotą tej metody stanowi pomiar zasobu nagromadzonego w jednostce uwzględniającej szerokie spektrum czynników wpływających na indywidualny poziom kapitału ludzkiego kształtowany w ciągu całego życia jednostki. Do rozważanych czynników należą między innymi: formalna edukacja, doświadczenie, starzenie się wiedzy i umiejętności, proces zapominania oraz stan zdrowia. Metoda ta umożliwia również pomiar kapitału na poziomie zbiorowości. W proponowanym ujęciu kapitał ludzki w populacji nie jest traktowany jedynie jako suma zasobów zgromadzonych w poszczególnych jednostkach. Dodatkowo rozpatrywane są czynniki mające wpływ na kształtowanie się kapitału ludzkiego na poziomie społeczeństwa, na przykład efekt synergii podczas współdziałania jednostek.

Na podstawie wyżej wymienionej metody oszacowano indywidualne profile kapitału ludzkiego oraz poziom kapitału ludzkiego na poziomie zagregowanym. Na potrzeby pracy stworzono demograficzne projekcje struktury ludności według poziomu jej wykształcenia w Polsce w latach 2002-2052. Zastosowano tutaj metodę projekcji wielostanowych. Opracowano również wielowariantowe projekcje kapitału ludzkiego dla populacji Polski na lata 2010-2050.

Na podstawie przeprowadzonego postępowania badawczego sformułowano następujące wnioski ogólne:

- Postępujący proces starzenia się populacji powoduje zmniejszenie się kapitału ludzkiego w ujęciu ilościowym. Dzieje się tak dlatego, że rozmiary populacji w wieku produkcyjnym kurczą się przy równoczesnym starzeniu się tej populacji.
- Jakościowe zmiany w akumulacji kapitału ludzkiego przejawiające się w dążeniu do podwyższania poziomu wykształcenia oraz do polepszania stanu zdrowia zwiększają kapitał ludzki przypadający na jednego pracownika.
- Wykonane dla Polski projekcje demograficzne wskazują, że inwestycje w kapitał ludzki łagodzą negatywne skutki oddziaływania starzenia się ludności na zasoby pracy. Procesy te nie są jednak w stanie ani zahamować ani też odwrócić tendencji wynikających ze zmian demograficznych.
- w Polsce kapitał ludzki w przeliczeniu na pracownika będzie rosł do 2050 roku, lecz tempo jego przyrostu znacząco się obniży. Kapitał ludzki na poziomie populacji będzie wzrastać tylko do około 2025-2030 roku. W dłuższym horyzoncie przewiduje się odwrócenie trendu wzrostowego kapitału ludzkiego.

Słowa kluczowe: kapitał ludzki, starzenie populacji, zjawiska demograficzne

Human capital and population ageing in Poland

Human capital is one of the most important resources that affect level and quality of life in populations. Its importance is emphasized in the formation of knowledge-based economies. Properly used high quality human capital can be both valuable production factor and a source of competitive advantage in the modern world.

Human capital formation is influenced by many factors. Among them demographic processes, such as fertility, mortality and migration, play a significant role. They influence size and structure of population actively involved in economic processes. Population ageing directly affects the labour market. It causes ageing and depopulation of the labour force.

The aim of the article is to estimate human capital in Poland and evaluate the consequences of population ageing on human capital in the period 2010-2050.

The estimation of human capital is done using the author's method based on accumulation and depreciation of this capital. It takes into account both its quantitative and qualitative dimensions. We include the following factors that influence human capital on individual level: formal education, experience, obsolescence of knowledge and skills, process of forgetting and health. The method can be also used for estimation of this resource on the aggregated level. In this approach we do not assume that total capital is just a sum of human capital accumulated by individuals. Additionally we add another effect - human capital accumulation as a result of synergy from collaboration of individuals.

We use this method for calculation of human capital in Poland on the individual and aggregated level in Poland. We also prepare multi-state population projection by age, sex and education and multi-scenario projection of human capital for the period 2010-2050.

We formulate the following general findings:

- Acceleration of population ageing causes a decrease of human capital in quantitative terms because of shrinkage and ageing working population.
- Investments in education and level of health - changes in quality of human capital - increase the human capital on the individual level.
- Population projections for Poland show that investments in human capital can slow down negative consequences of population ageing but they cannot stop or reverse these tendencies, caused by demographic change.
- According to projections, human capital per working person in Poland rises in the period 2010-2050, but its pace decreases. On the aggregated level human capital increases only until the years 2025-30. In the long term human capital diminishes.

Key words: human capital, population ageing, demographic processes

Strahl Danuta, Markowska Małgorzata (Uniwersytet Ekonomiczny we Wrocławiu)

Możliwości oceny innowacyjności regionów szczebla NUTS 2 w świetle zasobów informacyjnych Eurostatu

Referat przedstawia dorobek metodologiczny Eurostatu dotyczący pomiaru innowacyjności regionów państw Unii Europejskiej na szczeblu NUTS 2. Omówiona zostanie ewolucja podejść do pomiaru innowacyjności a także zmiany w zasobach informacyjnych ilustrujących dynamikę i poziom charakterystyk innowacyjności.

Dorobek Eurostatu zostanie oceniony na tle podejść stosowanych do oceny innowacyjności w świetle badań prowadzonych w literaturze przedmiotu.

Oceniona zostanie możliwość prowadzenia statystycznych analiz porównawczych w ujęciu dynamicznym w świetle zasobów informacyjnych dostępnych w bazach danych Eurostatu.

Słowa kluczowe: innowacyjność regionów, bazy danych, Eurostat, pomiar innowacyjności, regiony szczebla NUTS 2

The possibility of NUTS 2 level regions innovation assessment in the perspective of Eurostat information resources

The paper presents Eurostat methodological output referring to innovation measurement of NUTS 2 level regions representing EU member states. It also discusses the evolution of approaches towards innovation measurement and changes in information resources illustrating the dynamics and level of innovation characteristics.

Eurostat output will be evaluated at the background of approaches applied for the assessment of innovation in the perspective of research conducted in professional literature.

The assessment will refer to the possibility of performing statistical comparative analyses in dynamic perspective based on information resources available in Eurostat data bases.

Key words: innovation in regions, data bases, Eurostat, innovation measurement, NUTS 2 level regions

Styrc Marta (Szkola Główna Handlowa w Warszawie)

Czynniki wpływające na stabilność pierwszych małżeństw w Polsce

Stabilność małżeństw jest ważnym obszarem badań dotyczących rodziny ze względu na jej znaczenie dla płodności, dobrostanu dzieci oraz partnerów. Wzrost częstotliwości rozwodów obserwowany w Polsce od drugiej połowy lat 1990. prowadzi do pytania o czynniki związane ze zwiększonym ryzykiem rozwodów. Czynniki te mogą zostać podzielone na te, które dotyczą cech indywidualnych partnerów, charakterystyk związku oraz charakterystyk kontekstualnych. W opracowaniu opis czynników różnicujących skłonność do rozwodu, obecnych w rozważaniach teoretycznych i badaniach empirycznych demografów, poprzedza analizy empiryczne. Dane z badania ankietowego „Biografie rodzinne, zawodowe i edukacyjne” z 2006 r. zostały wykorzystane do oszacowania modelu hazardu rozpadu pierwszych małżeństw w Polsce w okresie od połowy lat 1980. do 2006 r. Zastosowano model wykładniczy przedziałami stały z proporcjonalnym wpływem zmiennych.

Zgodnie z oczekiwaniami małżeństwa z dziećmi przedmałżeńskimi, małżeństwa zawierane przez pary oczekujące dziecka, małżeństwa kobiet wychowanych w większych miejscowościach oraz kobiet pracujących są bardziej narażone na rozpad. Stwierdzono także, że ryzyko rozpadu małżeństwa wzrosło po 1999 r. Analizy wskazują także na prawdopodobną zmianę zależności pomiędzy stabilnością małżeństw a wykształceniem kobiet – we wcześniejszych badaniach małżeństwa kobiet lepiej wykształconych miały mniejsze ryzyko rozpadu, natomiast w niniejszym badaniu są to małżeństwa bardziej narażone na rozpad w porównaniu do średniego i podstawowego poziomu edukacji. Małżeństwa poprzedzone kohabitacją nie były mniej stabilne niż małżeństwa bezpośrednie. Słabszy niż oczekiwany okazał się związek ryzyka rozpadu małżeństwa z posiadaniem dzieci – jedynie małżeństwa z bardzo małymi dziećmi (0-2 lata) są bardziej stabilne, obecność starszych dzieci jak również liczba dzieci nie różnicuje ryzyka rozpadu w porównaniu do braku dzieci w związku.

W części podsumowującej sformułowano postulaty dotyczące źródeł danych do przyszłych badań nad stabilnością związków.

Determinants of the marital stability of first marriages in Poland

Marital stability is an important topic in studies on family because of its meaning for fertility and for the well-being of children and partners. The rise of the divorce rate observed in Poland since the second half of the 1990s poses a question about factors correlated with the risk of marital disruption. In the relevant body of research, one can distinguish between factors related to partners' characteristics, features of the relationship, and the context. The paper starts with theoretical considerations on the correlates of divorce, backed with empirical findings from studies for other countries. Next, the event history regression of first marriages disruption is

estimated. The model is specified as a piecewise constant exponential model with proportional relative risks. The data used come from the Education, Family and Employment Survey from 2006.

Most of the results of this study are consistent with findings for other countries: marriages with premarital children or entered into while expecting a child, as well as marriages of women brought up in bigger cities and those of employed women were less stable. The change of the educational gradient of divorce is an important finding – in the studies pertaining to the period before the 1990s the women with higher education showed a higher risk of divorce. In the current study, which refers to the period between the mid-1980s and 2006, marriages of women with highest education levels have the lowest risk of disruption. Surprisingly, marriages preceded by cohabitation do not have an elevated disruption risk compared to direct marriages. The impact of the presence of children on the disruption risk is lower than expected – only marriages with very small children (0-2 years old) are more stable, while parity and presence of older children do not make a difference, when compared with couples without children. In conclusion, some suggestions have been formulated regarding the data sources for the future research on the stability of unions.

Key words: divorce, marital disruption, stability of marriage

Styrc Marta (Szkoła Główna Handlowa w Warszawie)

Gradient edukacyjny rozwodów w Polsce

Zgodnie z hipotezą Goode'go (1962) wraz z upowszechnianiem się rozwodów maleje różnica pomiędzy intensywnością rozwiązywania małżeństw kobiet z wyższym i niższym wykształceniem. W sytuacji, gdy rozwody są relatywnie rzadkie w danym społeczeństwie, jedynie kobiety mające lepszy dostęp do zasobów (ekonomicznych, społecznych, interpersonalnych) są w stanie pokonywać trudności związane z rozwodem. Edukacja jest dobrym przybliżeniem dostępu do zasobów (Perelli-Harris i in. 2010) i stąd też w sytuacji małej dostępności rozwodów kobiety lepiej wykształcone mają większe ryzyko rozwiązania małżeństwa w porównaniu do kobiet gorzej wykształconych. Gdy rozwód staje się bardziej powszechny, a jego koszty ekonomiczne i społeczne maleją, różnice według poziomu wykształcenia zmniejszają się lub wręcz może dojść do zmiany gradientu edukacyjnego rozwodów z pozytywnego na negatywny, tj. im niższy poziom wykształcenia, tym wyższa intensywność rozwodów.

W Polsce przemiana systemowa z gospodarki centralnie sterowanej do rynkowej rozpoczęta w 1989 r. wpłynęła także na instytucjonalne uwarunkowania rozwodów, co mogło wpłynąć na zmianę gradientu edukacyjnego. Sugerują to wyniki wcześniejszych badań: Harkonen i Dronkers (2006) na danych z FFS (Family and Fertility Survey) odnoszących się do okresu sprzed 1991 r. stwierdzili pozytywny gradient rozwodów w Polsce. Natomiast Styrc (2010) na danych z EFES (Employment, Family and Education Survey) obejmujących okres 1990-2006 stwierdziła, że

ryzyko rozpadu małżeństwa (rozpad definiowany przez respondenta) jest najniższe wśród kobiet najlepiej wykształconych. Zestawienie tych dwóch wyników badań potencjalnie wskazuje na zmianę gradientu edukacyjnego w Polsce w latach 1990., jednak ze względu na nieporównywalność danych i metody potwierdzenie tej hipotezy wymaga przeprowadzenia odrębnych analiz. Umożliwiają to dane z badania "Generacje, rodziny i płeć kulturowa GGS-PL" (Generations and Gender Survey), które pozwalają na obserwację historii małżeńskich od lat 1980. do lat 2000., tym samym pokrywając okres potencjalnej zmiany gradientu edukacyjnego. W celu przetestowania hipotezy dotyczącej zmiany gradientu edukacyjnego zostaną oszacowane modele analizy historii zdarzeń. Wpływ wykształcenia zostanie oszacowany odrębnie dla wyróżnionych okresów czasu kalendarzowego, co pozwoli na obserwację zmiany zależności w czasie.

Słowa kluczowe: rozwód, wykształcenie, małżeństwo

Educational gradient of divorce in Poland

According to Goode (1962) the rise of divorce is associated with a change in educational gradient of divorce. If divorce is relatively rare in a society only women having access to broadly defined resources (economic, social, interpersonal) are able to cope with difficulties associated with marital disruption. The access to resources may be measured through education achieved (Perelli-Harris et al. 2010) and thus while the access to divorce is limited women with higher education have greater risk of marital disruption compared to less educated women. The spread of divorce together with its decreasing economic and social cost reduce the differentials in the risk of marital dissolution. In particular, it is possible that the relation between education and marital stability reverses, i.e. the risk of disruption becomes higher for lower educated women.

In Poland the system transformation from centrally planned to market economy initiated in 1989 has modified also the institutional settings for marital dissolution what may have influenced the educational gradient of divorce. Previous results indicate on a possible reversal of educational gradient: Harkonen and Dronkers (2006) using the FFS data (Family and Fertility Survey) reported positive correlation between education and the risk of divorce before 1991. Styrc (2010) using the EFES data (Employment, Family and Education Survey) that covered period 1990-2006 stated that the risk of marital dissolution is the lowest for women with higher education. In order to test the Goode's hypothesis in Poland an event history study covering period since 1980s has been conducted. The data come from the Polish 2011 Generation and Gender Survey.

Key words: divorce, education, marriage

BIBLIOGRAPHY

1. Goode W. J., 1962, Marital satisfaction and instability: a cross-cultural class analysis of divorce rates, [w:] Bendix R., Lipset S. M. (red.), *Class, status, and power. social stratification in comparative perspective*, The Free Press, New York: 377-387.
2. Härkönen J., Dronkers J., 2006, Stability and change in the educational gradient of divorce. A Comparison of seventies countries, „*European Sociological Review*”, 22(5): 501-517.

3. Perelli-Harris B., Sigle-Rushton W., Kreyenfeld M., Lappergard T., Keizer R., Berghammer C., 2010, The educational gradient of childbearing within cohabitation in Europe, „Population and Development Review”, 36(4): 775-801.
4. Styrz M. (2010), Czynniki wpływające na stabilność pierwszych małżeństw w Polsce, „Studia Demograficzne”, 157-158, w druku.

Styrz Marta, Matysiak Anna (Szkoła Główna Handlowa w Warszawie)

Status społeczno-ekonomiczny a stabilność małżeństw

Ekonomiczna teoria rodziny, sformułowana przez Garry'ego Beckera (Becker i in. 1977), przewiduje, że zatrudnienie mężczyzn ma pozytywny wpływ na stabilność małżeństw, podczas gdy zatrudnienie kobiet destabilizuje małżeństwo. Wnioski te wynikają z założenia o maksymalizacji użyteczności gospodarstwa domowego poprzez specjalizację kobiet w produkcji domowej a mężczyzn w dostarczaniu dochodów. Coraz częściej są one jednak poddawane w wątpliwość. W literaturze przedmiotu pojawiają się argumenty mówiące, że w nowoczesnych społeczeństwach zmieniają się role kulturowe kobiet i mężczyzn (Sigle-Rushton 2010), a decyzje o zawarciu związku i trwaniu w nim w dużej mierze zależą od satysfakcji czerpanej ze związku (Ross i Sawhill 1975). Co więcej, praca zawodowa kobiet prowadzi do poprawy sytuacji materialnej pary, co może zmniejszać napięcia na tle finansowym (Oppenheimer 1997, Cherlin 2000).

Celem niniejszego opracowania jest wkład do dyskusji na temat związku między zatrudnieniem kobiet i mężczyzn a stabilnością małżeństw poprzez analizę tej zależności w Polsce. Do oszacowania modeli historii zdarzeń wykorzystano dane z badania GGS-PL (Generation and Gender Survey).

Wyniki wstępnych analiz są w dużej mierze zgodne z teorią Beckera. Małżeństwa mężczyzn pracujących charakteryzują się wyższą stabilnością niż mężczyźni niepracujących, natomiast wśród kobiet obserwuje się zależność odwrotną. Jednak o ile wśród mężczyzn zależność ta utrzymywała się zarówno w okresie gospodarki planowanej, jak i po 1989 roku, o tyle w przypadku kobiet ujawniła się ona dopiero w latach 1990-tych. Sugeruje to, że specjalizacja ról kobiet i mężczyzn jest nadal silnie zakorzeniona w społeczeństwie polskim a przemiany społeczno-ekonomiczne, jakie zaszły w Polsce w okresie transformacji ustrojowej, uwypukliły tylko jej oddziaływanie na stabilność małżeństw.

Słowa kluczowe: rozwód, rozpad małżeństwa, stabilność małżeństwa, status społeczno-ekonomiczny, aktywność zawodowa

Socio-economic status and marital stability

The economic theory of family formulated by Garry Becker (Becker et al. 1977) predicts that men's employment stabilises marriage, contrary to the effect of women's employment which is expected to destabilise marriage. This prediction is based on the assumption of sex role specialisation within a couple which presupposes that couple's utility is maximised if a wife specialises in home production and a husband in paid work. Nevertheless, it has been increasingly questioned recently, as women have been more and more present in the labour market, there has been an ongoing change in gender roles and household organisation has been shifting from production to consumption (Sigle-Rushton 2010). It has been argued that in modern societies decisions to remain married depend more on satisfaction from the quality of the union and that similarity of economic activities and interests may improve understanding between spouses (Ross & Sawhill 1975). Moreover, the additional income provided by a woman leads to higher living standards and thus should reduce marital strains (Oppenheimer 1997, Cherlin 2000).

The aim of the study is to contribute to the discussion on the association between men's and women's economic activity and marital stability through investigating the case of Poland. To this end, we estimate an event-history model for marital disruption using the Polish GGS (2011).

The preliminary results are consistent with Becker's theory. Employment of men is associated with higher marital stability and employment of women's with lower marital stability. Interestingly, for men the relation held both in the centrally planned and market economy, but for women it became apparent only in the 1990s. The results suggest that gender specialisation roles are still valid in the Polish society and their effects for marital stability have been strengthened following the transition from the centrally planned to the market economy.

Key words: divorce, marital disruption, stability of marriage, socio-economic status, labour force participation

BIBLIOGRAPHY

1. Becker, G. S., Landes, E. M., Michael, R. T. (1977). An economic analysis of marital instability. *Journal of Political Economy*, 85, 1141-1188.
2. Cherlin, A.J. (2000). Toward a new home socioeconomics of union formation. In L. J. Waite, C. Bachrach, M. Hindin, E. Thomson, & A. Thornton (Eds.), *The ties that bind - Perspectives on marriage and cohabitation* (pp. 126-144). Hawthorne, NY: Aldine De Gruyter.
3. Oppenheimer, Valerie Kincade, (1997), Women's employment and the gain to marriage: The specialization and trading model. *Annual Review of Sociology*, 23: 431-453.
4. Ross, H.L. and I.V. Sawhill. (1975). *Time of Transition. The Growth of Families Headed by Women*. Washington, DC: The Urban Institute.
5. Sigle-Rushton, W. (2010). Men's unpaid work and divorce: reassessing specialization and trade in British Families. *Feminist Economics*, 16, 1-26.

Suchecka Jadwiga, Modranka Emilia (Uniwersytet Łódzki)

Statystyka przestrzenna - metody i zastosowania

Zasadniczym celem referatu jest przedstawienie najważniejszych problemów stojących do rozwiązania przed statystyką przestrzenną i wskazanie jej roli w rozwoju metodologii badań statystycznych.

Znaczący wzrost zainteresowania danymi zlokalizowanymi (przestrzennymi i przestrzenno-czasowymi) w badaniach statystycznych obserwujemy dopiero od połowy lat 70. XX wieku. Był to istotny impuls w rozwoju metod ilościowych umożliwiających ich prawidłową analizę. Szczególnie dotyczy to metod, które umożliwiają opis struktur przestrzennych oraz przestrzenną predykcję. Nauką dostarczającą odpowiednich metod takiej analizy w ogólnym znaczeniu jest statystyka przestrzenna.

Dane przestrzenne i przestrzenno-czasowe mogą być danymi geograficznymi lub danymi zlokalizowanymi charakteryzującymi się wielością informacji. Wymiar przestrzenny danych zlokalizowanych wprowadza zależności dwuwymiarowe, bardziej złożone niż jednowymiarowe zależności czasowe. Cechą charakterystyczną zlokalizowanych obserwacji przestrzennych jest to, że są one współzależne i heterogeniczne. W takich przypadkach zastosowanie klasycznych miar statystycznych powoduje utratę istotnej części informacji, a klasyczne estymatory podstawowych parametrów struktury tracą niektóre własności. Stąd też statystyka przestrzenna proponuje odpowiednie estymatory charakteryzujące się dobrymi własnościami.

Zakres statystyki przestrzennej zmienia się wraz z możliwościami przetwarzania dużych baz danych, rozwojem specjalistycznego oprogramowania i zapotrzebowaniem praktyki gospodarczej na różnego rodzaju analizy stanowiące podstawę procesu decyzyjnego. Najczęściej związany jest z poszukiwaniem odpowiedzi na pytania dotyczące:

- charakteru analizowanego zjawiska, a więc określenia czy w przestrzeni ma ono charakter regularny, losowy czy klastrowy;
- wykazywania prawidłowości przestrzennej przez określone zdarzenia;
- skupienia lub rozproszenia podobnych wartości badanej zmiennej;
- stopnia skorelowania określonej zmiennej z innymi zmiennymi w danym obszarze;
- niezależności różnych wartości badanej zmiennej;
- przyciągania się lub odpychania się podobnych wartości;
- zmian w przestrzennym rozkładzie badanej zmiennej w czasie,
- znajomości funkcji kowariancji lub wariogramu, itd.

Klasyfikacja problemów stojących przed statystyką przestrzenną wynika z występowania różnych typów danych geograficznych i zlokalizowanych (powierzchniowe z atrybutami ilościowymi i jakościowymi, punktowe bez atrybutów, obszarowe, regionalne i inne).

W ramach współczesnej statystyki przestrzennej wyróżnia się trzy podstawowe grupy metod:

- metody geostatystyczne (opisu statystycznych danych dotyczących powierzchni Ziemi);
- metody analizy danych punktowych bez atrybutów (ustalania sposobu przestrzennego rozmieszczenia zjawisk i zależności przestrzennych);
- metody analizy danych obszarowych i punktowych atrybutowych (umożliwiające analizę jednostek administracyjnych i obiektów przyrodniczych, które mają granice i są zorientowane według współrzędnych geograficznych).

W referacie omówione zostaną podstawowe zagadnienia i zastosowania nowoczesnych metod statystyki przestrzennej. W szczególności prezentowane będą metody opisu przestrzennej struktury zjawisk, estymacji parametrów i interpolacji danych przestrzennych oraz metody analizy danych punktowych.

Słowa kluczowe: dane przestrzenne, metody geostatystyczne, analiza danych punktowych i obszarowych

Spatial statistics - methods and applications

The main aim of the paper is to present the most important problems facing the solution before the spatial statistics and identify its role in the development of statistical research methodology.

The significant growth of interest in localized data (spatial and spatio-temporal) in economic and social analyses is observing only since the mid-'70s the twentieth century. It was an important incentive in the development of quantitative methods to enable the proper analysis of spatial data. This is particularly true of methods that allow the description of spatial structures and spatial prediction. In a general sense spatial statistics is a branch of knowledge about these methods.

Spatial and spatio-temporal data can be geographical or localized data characterized by a multiplicity of information. The spatial dimension of data introduces the two-dimensional dependencies and interactions, more complex than the one-dimensional time-dependence. A characteristic feature of localized spatial observation is that they are interdependent and heterogeneous. In such cases, the use of classical statistical measures will lose a substantial part of the information, and the estimators of basic structural parameters are losing some properties. Hence, spatial statistics propose appropriate, good estimators.

Range of spatial statistics varies with the processing capabilities of large databases, specialized software development and business practice needs for different types of analysis underlying the decision making process. This range is frequently associated with seeking answers to questions about:

- analyzed the nature of the phenomenon, and thus determine whether it has a regular, random or clustered character in the space;
- demonstrating spatial regularity by specific events;
- concentration or dispersion of similar values of the test variable;
- the variable degree of correlation with other variables in the area;
- the independence of the different values of the test variable;
- attraction or repulsion of similar values;
- changes in the spatial distribution of the test variable at a time
- knowledge of the covariance function or variogram, etc.

Classification of issues facing the spatial statistics are caused by different types of geographic/localized data: surface data, point patterns data (with or without the quantitative and qualitative attributes), lattice, regional, and others data).

As a part of modern spatial statistics, there are three basic groups of methods:

- geostatistical methods (description of the statistical data on the surface);
- methods for analysis of point data without attributes (determining the spatial distribution of phenomena and spatial relationships);
- methods for analysis of area and point data with attributes (for an analysis of administrative units and natural objects that have boundaries and are oriented by latitude and longitude).

In this paper the basic issues and the application of modern methods of spatial statistics will be discussed. In particular, the methods describe the spatial structure of phenomena, estimation of parameters and interpolation of spatial data and point data analysis methods will be presented.

Key words: spatial data, geostatistical methods, point patterns and lattice data analysis

BIBLIOGRAPHY

1. Droesbeke J.J., M. Lejeune, G. Saporta (red), (2006), *Analyse statistique des données spatiales*, Ed. Technip, Paris
2. Arbia G., (1989), *Spatial Data Configuration in Statistical Analysis of Regional Economic and Related Problems*, Kluwer Academic Publishers, Dordrecht.
3. Suhecka J., Antczak E., (2010), *Elementy geostatystyki i metody analiz przestrzennych danych punktowych*, [w]: Bogdan Suhecki (red.), *Ekonometria Przestrzenna. Metody i modele analizy danych przestrzennych*, Wydawnictwo C. H. Beck, Warszawa, s. 70-99
4. Wingle W.L., E.P. Poeter, (1999), *Geostatistical Analysis and Training via the Internet*, APCOM'99: Computer Applications in the Minerals Industries 28th International Symposium, October 20-22, 1999, Colorado School of Mines, Golden, Colorado.

Szkutnik Zbigniew (Akademia Górniczo-Hutnicza w Krakowie)

Algorytm EM i jego modyfikacje

Algorytm EM jest często stosowaną metodą iteracyjną wyznaczania estymatorów największej wiarygodności. Jego źródła można wprawdzie szukać w różnych starszych i nieformalnych pracach aplikacyjnych, ale prawdziwym początkiem jego popularności stała się praca [1]. To w niej zaproponowano nazwę algorytmu (od ang. Expectation-Maximization, co odpowiada dwóm krokom powtarzanym w kolejnych iteracjach), podjęto próbę formalnej analizy jego własności i wskazano wiele możliwości zastosowań. Okres kolejnych 20 lat intensywnych badań algorytmu został podsumowany w poświęconej mu książce [3]. W wykładzie zostaną przedstawione główne idee algorytmu, jego podstawowe własności, w szczególności delikatna kwestia zbieżności ([2,6]), oraz różne zaproponowane w literaturze modyfikacje głównej idei (np. [4,5]). Wskazane będą także przykłady zastosowań, zarówno te z tzw. brakującymi danymi, do których algorytm EM został stworzony (np. tzw. problemy odwrotne, [5]), jak i te, w których przed zastosowaniem algorytmu EM trzeba problem odpowiednio przeformułować. algorytm EM, estymatory największej wiarygodności, dane brakujące

Słowa kluczowe: algorytm EM, estymatory największej wiarygodności, dane brakujące

EM algorithm and its modifications

EM algorithm is a popular iterative method of finding maximum likelihood estimators. Although its roots may be traced back to some older and rather informal applied papers, the great career of the EM algorithm was really launched by the seminal paper [1], in which the term EM was coined (from Expectation-Maximization, corresponding to two basic steps repeated iteratively), the first (although partly flawed) analysis of the convergence was given and several possible applications were pointed to. The next 20 years of intensive research on the EM properties and applications were summarized in the book [3]. In this lecture, basic ideas related to EM will be presented and the somewhat delicate question of its convergence will be discussed (see, e.g., [2,6]). Some of the proposed modifications of the algorithm (e.g., [4,5]) and some exemplary applications will also be pre-sented. The examples will include both problems with missing data (e.g. inverse problems, [5]), for which EM was mainly created, and some other problems which have to be suitably reformulated to become amenable to the EM technology.

Key words: EM algorithm, maximum likelihood estimators, missing data

BIBLIOGRAPHY

1. Dempster A.P., Laird N.M., Rubin D.B., Maximum likelihood from incomplete data via

- the EM algorithm (with discussion), "Journal of the Royal Statistical Society B", 1977, 39, 1-38.
2. Latham G.A., Andersen R.S., Assessing quantification for the EMS algorithm, "Linear Algebra and its Applications", 1994, 210, 89-122.
 3. McLachlan G.J., Krishnan T., The EM Algorithm and Extensions, Wiley, New York, 1997.
 4. Silverman B.W., Jones M.C., Wilson J.D., Nychka D.W., a smoothed EM approach to indirect estimation problems with particular reference to stereologz and emission tomography (with discussion), "Journal of the Royal Statistical Society B", 1990, 52, 271-324.
 5. Szkutnik Z., Doubly smoothed EM algorithm, "Journal of the American Statistical Association", 2003, 98, 178-190.
 6. Wu C.F.J., On the convergence properties of the EM algorithm, "Annals of Statistics", 1983, 11, 95-103.

Szmuksta-Zawadzka Maria, Zawadzki Jan (Zachodniopomorski Uniwersytet Technologiczny w Szczecinie)

O metodach prognozowania brakujących danych w szeregach czasowych z wahaniami okresowymi (sezonowymi)

Referat zostanie poświęcony omówieniu metod prognozowania w warunkach braku pełnej informacji w szeregach czasowych dla okresów jednostkowych o różnej długości - poczynając od miesięcy a na danych krótszych niż dzień kończąc. Rozważania odnosić się będą do dwóch rodzajów luk w danych: systematycznych i niesystematycznych. Z lukami systematycznymi mamy do czynienia wtedy, gdy nie będą dostępne informacje liczbowe przynajmniej o jednym podokresie w całym przedziale czasowym próby. Rozpatrywane będą metody prognozowania zarówno dla danych oryginalnych (z sezonowością) jak i danych z których wyeliminowano wahania sezonowe. Egzemplifikacją rozważań o charakterze teoretycznym będą przykłady prognozowania brakujących danych w ekonomicznych szeregach czasowych o różnej długości.

Słowa kluczowe: szeregi czasowe, wahania sezonowe, brakujące dane, prognozowanie

On methods of forecasting the missing data in time series of periodical (seasonal) fluctuations

The subject of the paper is to discuss prediction methods in case of missing information in the time series of different length - from months to the ones shorter than a day. Considerations shall cover two kinds of gaps in systematic and non-systematic data. The systematic gaps occur where

the numerical information is not available at least on one sub-period in the entire time interval. The prediction methods for the original data (seasonal data) and the data without seasonal fluctuations would be considered. Examples of the prediction of missing data in economic time series of different length would be given to illustrate theoretical considerations.

Key words: time series, seasonal fluctuations, missing data, prediction

Sztanderska Urszula (Uniwersytet Warszawski)

Polska statystyka publiczna jako inspiracja i bariera rozwoju badań rynku pracy

Zmiana systemu ekonomicznego postawiła przed statystyką publiczną zupełnie nowe zadania wynikające m.in. z potrzeby obserwacji samoczynnych procesów gospodarczych i diagnozowania na ich podstawie kondycji ekonomicznej kraju. Badania rynku pracy stanowiły w tym zakresie potężne wyzwanie, któremu – po części – udało się sprostać. Uruchamiane od lat 90. nowe badania i stare badania, ale znacznie przeorientowane, pozwalały (i nadal pozwalają) na diagnozę funkcjonowania rynku pracy, zaspakajającą nie tylko ogólny głód wiedzy ale stwarzającą podstawy działań, zwłaszcza dla polityki ekonomicznej (monetarnej, fiskalnej, polityki rynku pracy jak i iniektorach innych powiązanych polityk, zwłaszcza edukacyjnej).

Interesariusze statystyki publicznej to nie tylko służby publiczne kreujące politykę ekonomiczną i społeczną w skali kraju, regionalnej i lokalnej, ale i w podmioty życia gospodarczego – firmy, instytucje i obywatele. Wiedza o kondycji, kierunkach rozwoju rynku pracy stanowi bowiem nieodłączny składnik decyzji tych podmiotów; utrudniony dostęp do niej generuje błędne decyzje (a więc i np. nietrafione inwestycje w edukację, migracje, lokalizacje działalności gospodarczej, poszukiwanie pracy lub poszukiwanie pracowników etc. albo przynajmniej podnosi ich koszty transakcyjne). Z tego punktu widzenia można stworzyć szeroką listę potrzeb informacyjnych, których zaspokojenie poza statystyką publiczną jest niemożliwe lub utrudnione i zbyt kosztowne. Należą do nich m.in. badania losów absolwentów na rynku pracy jako źródło wiedzy o przydatności i wykorzystaniu nowo tworzonego kapitału ludzkiego, poszerzone analizy bezrobocia i, zwłaszcza w kontekście bezrobocia równowagi i inne. W tej roli – dostarczania zestawu informacji dla decyzji ekonomicznych – statystyka publiczna jest więc klasycznym dobrem publicznym, a wydaje się często niedocenianym.

Referat dostarcza przeglądu podstawowych potrzeb analitycznych dla diagnozowania rynku pracy z wykorzystaniem statystyki publicznej i wskazuje bariery ich zaspokojenia oraz potencjalne możliwości relatywnie niskonakładowego ich pokonania, zwłaszcza w powiązaniu z istniejącymi (ewentualnie wymagającymi nieznacznych modyfikacji) publicznymi zasobami informacyjnymi takimi jak np. informacje Zakładu Ubezpieczeń Społecznych. Wskazuje też na słabości niektórych zasobów takich np. jak należące do Publicznych Służb Zatrudnienia, czy Systemu Informacji Oświatowej, których modyfikacja i otwarcie na potrzeby statystyki publicznej mogłoby znacząco podnieść ich wartość analityczną. Wreszcie dostrzega obszary, o których dane

publiczne i w ślad za tym statystyczne pozostają fragmentaryczne lub niedostosowane do potrzeb badawczych i analitycznych, co obniża trafność decyzji i efektywność funkcjonowania rynku pracy.

Z drugiej strony referat prezentuje używanie utworzonych lub zmodyfikowanych badań statystyki publicznej i analityczny potencjał, jaki zawierają, a który nie zawsze jest dostatecznie wykorzystywany przez użytkowników. Takimi podstawowymi zasobami informacji są kwartalne Badania Aktywności Ekonomicznej Ludności, okresowo poszerzane o modułowe badania niektórych aspektów funkcjonowania rynku pracy jak np. na temat powiązania pracy zawodowej z wypełnianiem obowiązków rodzinnych, czy wcześniej pasywnych i aktywnych polityk rynku pracy, przeprowadzane co 2 lata Badania Wynagrodzeń, znacznie zmodyfikowane, ciągłe Badania Budżetów Gospodarstw Domowych i inne. Ich zasadniczą cechą jest możliwość bardziej gruntownych wieloprzekrojowych analiz statystycznych, czego nie sposób uzyskać w badaniach agregujących informacje jednostkowe. W badaniach tych szczególnie uwzględnia się stronę popytową rynku pracy (osoby), skąpiej zaś reprezentowane są podmioty ze strony popytowej – przedsiębiorstwa i instytucje. Wprawdzie istnieją badania i tej strony jak np. Badania Popytu na Pracę, jednak nie pozwalają one na równie precyzyjną analizę popytowej sytuacji na rynku pracy, co poważnie ogranicza możliwości analiz użytkowych i badań naukowych.

Szukielójc-Bieńkuńska Anna (Główny Urząd Statystyczny)

Jakość życia w badaniach GUS

Od statystyki publicznej oczekuje się aby nadążała za zmieniającymi się uwarunkowaniami i pozwalała na bieżąco monitorować najważniejsze zjawiska i procesy gospodarce i życiu społecznym. Dotyczy to zarówno kontekstu krajowego jak i międzynarodowego. Wymusza to stałą modernizację i doskonalenie systemów informacji statystycznych. W ostatnich latach priorytetowym zadaniem dla statystyki publicznej jest udoskonalenie metod pomiaru i analizy różnych aspektów rozwoju społecznego, w tym jakości życia.

Jakość życia jest pojęciem złożonym i podobnie jak wiele innych pojęć stosowanych w naukach społecznych nie mającym jednej powszechnie akceptowanej definicji. W referacie przedstawiona zostanie koncepcja badań i analiz jakości życia przyjęta przez GUS. Koncepcja ta nawiązuje do doświadczeń i tradycji polskich badań społecznych jak również do zaleceń międzynarodowych, w szczególności zaleceń opracowanych w ramach Europejskiego Systemu Statystycznego.

Na potrzeby prowadzonych przez GUS prac badawczych i analitycznych przyjęto, że jakość życia obejmuje zarówno obiektywną ocenę materialnych i niematerialnych wymiarów życia i sytuacji społecznej jednostek oraz grup społecznych jak również subiektywne poczucie satysfakcji życiowej w kontekście własnych potrzeb i możliwości.

Podstawowym warunkiem dokonania wielowymiarowych ocen jakości życia jest konieczność in-

tegracji wiedzy dotyczącej różnych wymiarów życia. Cel ten można osiągnąć w dwojaki sposób: poprzez łączenie informacji pochodzących z wielu źródeł danych bądź poprzez wdrażanie wieloaspektowych ankietowych badań jakości życia ludności. Biorąc pod uwagę zarówno ograniczenia jak i pozytywne strony każdego z tych rozwiązań GUS stara się stosować je równolegle.

W referacie najwięcej uwagi poświęcone będzie nowemu, zrealizowanemu przez GUS w 2011 roku, wieloaspektowemu badaniu jakości życia i spójności społecznej. Omówione zostaną zakres tematyczny oraz wybrane wyniki badania. W stosunku do realizowanych wcześniej przez GUS badań, w badaniu jakości życia i spójności społecznej rozszerzono zakres przedmiotowy o nowe bądź niedostatecznie wcześniej uwzględniane aspekty jakości życia, a także wprowadzono na szerszą skalę zmienne subiektywne - zarówno w zakresie oceny sytuacji materialnej jak również poziomu zadowolenia z różnych sfer życia, percepcji niektórych zjawisk społecznych. W konsekwencji, w badaniu uwzględniony został zestaw podstawowych pytań dotyczących najważniejszych aspektów jakości życia badanej jednostki.

Słowa kluczowe: jakość życia, wskaźniki obiektywne, subiektywne oceny, poziom zadowolenia, integracja

Quality of life in the surveys of the Central Statistical Office of Poland

Public statistics is expected to keep pace with changing conditions and to allow for a current monitoring of the most important phenomena and processes occurring in economy and social life. This concerns both domestic and international contexts and requires a permanent modernization and refinement of statistical information systems. In recent years, the priority task for public statistics has been the refinement of measurement methods and of analysis regarding various aspects of social development, including quality of life.

Quality of life is a complex notion and, similarly to a number of other notions used in social science, it does not have a one, commonly accepted definition. This paper presents the concept of quality of life surveys and analyses adopted by the Central Statistical Office of Poland (CSO). This concept draws on experiences and traditions of Polish social surveys as well as on international recommendations, particularly these developed within the European Statistical System.

For the needs of research and analytical works carried out by the CSO, it was assumed that 'quality of life' covers both objective assessment of material and non-material dimensions of individuals' and social groups' life and social situation, as well as subjective feeling of satisfaction with life in the context of own needs and capabilities.

The necessity to integrate knowledge on various dimensions of life is the basic condition for making multidimensional assessments of quality of life. This goal can be attained in two ways: by combining information from a number of sources or by implementing multidimensional questionnaire surveys of quality of life. Considering both limitations and positive sides to each of these solutions, the CSO endeavors to apply both of them at the same time.

In the paper, most of the attention is paid to the new multidimensional survey of quality of

life and social cohesion, carried out by the CSO in 2011. Thematic scope and some results of the survey are discussed. Its thematic scope is wider than in the surveys carried out by the CSO earlier, so that it encompasses new aspects of quality of life and these which were not sufficiently included earlier. Also, subjective variables are introduced on a larger scale - concerning the assessment of material situation, the level of satisfaction with various spheres of life and the perception of some social phenomena. Consequently, a set of basic questions on the most important aspects of interviewee's life is taken into account in the survey.

Key words: quality of life, objective indicators, subjective assessment, integration

Szukielój-Bieńkuńska Anna (Główny Urząd Statystyczny)

Pomiar ubóstwa i wykluczenia społecznego w polskiej statystyce publicznej

Ubóstwo jako kwestię społeczną i ekonomiczną można rozpatrywać w kilku wymiarach. Istnieje ona w wymiarze społeczności lokalnej, regionalnym, państwowym, ponadpaństwowym oraz globalnym. Na każdym z tych poziomów niezbędnym warunkiem podjęcia skutecznych działań mających na celu przeciwdziałanie ubóstwu oraz walkę z jego negatywnymi skutkami jest rzetelna diagnoza tego zjawiska, w tym ocena jego rozmiarów oraz identyfikacja zagrożonych grup. Podstawowym źródłem wiedzy z tego zakresu są badania i analizy statystyczne.

Referat zawierać będzie syntetyczny przegląd stosowanych dotychczas przez polską statystykę publiczną metod pomiaru ubóstwa i wykluczenia społecznego, zarówno w kontekście potrzeb krajowych, jak i wymagań Unii Europejskiej w tym zakresie. Omówiona zostanie także propozycja kompleksowego podejścia do pomiaru oraz analizy ubóstwa i wykluczenia społecznego w oparciu o wyniki nowego wieloaspektowego badania spójności społecznej, zrealizowanego przez GUS w 2011 r. Analizowane będą trzy formy ubóstwa: ubóstwo monetarne (dochodowe), ubóstwo oceniane z punktu widzenia złych warunków życia oraz ubóstwo subiektywne. Pokazane zostanie czy i w jakich grupach osób mamy do czynienia w Polsce z nakładaniem się różnych form ubóstwa, a także czy i w jakim stopniu ubóstwo związane jest z izolacją społeczną, co w konsekwencji prowadzi do wykluczenia społecznego. W referacie przedstawione zostaną także wyniki badania dotyczące społecznej percepcji ubóstwa i społecznego wykluczenia.

Słowa kluczowe: ubóstwo, wielowymiarowe ubóstwo, izolacja społeczna, społeczne wykluczenie

The measurement of poverty and social exclusion in Polish public statistics

Poverty, as a social and economic problem, can be considered at several dimensions. It exists at the level of local community, as well as at regional, national, supranational and global levels. At each of these, a reliable diagnosis of this phenomenon, including an assessment of its scale

and the identification of groups at risk of it, is an indispensable condition to the effectiveness of actions aimed at counteracting poverty and combating its negative consequences. Statistical surveys and analyses are the basic source of knowledge in this field.

The paper contains a synthetic review of methods of poverty and social exclusion measurement hitherto applied by Polish public statistics, both in the context of domestic needs and European Union's requirements in this regard. It also discusses a proposal for a comprehensive approach to the measurement of poverty and social exclusion, based on the results of a new, multidimensional, social cohesion survey carried out by the CSO in 2011. Three forms of poverty are analyzed: monetary (income) poverty, poverty assessed from the standpoint of bad living conditions, as well as subjective poverty. The paper shows whether the overlapping of various forms of poverty occurs in Poland and if so, in which groups of individuals (social groups); also, it shows whether and to what extent is poverty linked with social isolation, which leads to social exclusion. The paper presents also the results of the survey as far as social perception of poverty and social exclusion is concerned.

Key words: poverty, multidimensional poverty, social isolation, social exclusion

Szwarc Krzysztof (Uniwersytet Ekonomiczny w Poznaniu)

Poziom życia ludności a zróżnicowanie sytuacji demograficznej w powiatach województwa wielkopolskiego

Struktura demograficzna danego rejonu ma znaczący wpływ na procesy ludnościowe zachodzące w danym społeczeństwie, a te z kolei warunkują funkcjonowanie innych dziedzin życia. Mają one wpływ między innymi na sytuację na rynku pracy, edukację itp. To właśnie pod kątem struktury demograficznej powinno przygotowywać się różnego rodzaju reformy. Z sytuacją demograficzną nierozzerwalnie związany jest poziom życia ludności. Osoby z rodzin wielodzietnych w mniejszym zakresie korzystają z usług kulturalnych, sportowych, rekreacyjnych niż osoby nie posiadające dzieci. Wiele wskaźników dobrobytu jest odnoszonych do liczby mieszkańców danego obszaru. Wystarczy wspomnieć PKB na mieszkańca, liczbę lekarzy na 100 osób czy liczbę mieszkań oddanych do użytku na 1000 mieszkańców. Inne zmienne, np. związane z zanieczyszczeniem środowiska czy dostępem do usług medycznych wpływają na umieralność.

Celem referatu jest zbadanie rozwoju poziomu życia ludności w powiatach województwa wielkopolskiego, a także zróżnicowania sytuacji demograficznej w badanych jednostkach administracyjnych. Postawiono następujące hipotezy badawcze:

- istnieje przestrzenne zróżnicowanie poziomu życia ludności w powiatach województwa wielkopolskiego;
- procesy demograficzne zachodzące w województwie wielkopolskim charakteryzują się znaczącym zróżnicowaniem;

- niektóre aspekty sytuacji demograficznej badanych obszarów są ściśle związane z poziomem życia ludności.

Do oceny zróżnicowania badanych zjawisk zostanie wykorzystany wzorcowy miernik Hellwiga. Poziom życia ludności zostanie oceniony na podstawie szeregu zmiennych charakteryzujących następujące dziedziny:

- ochrona zdrowia i opieka socjalna;
- rynek pracy, warunki i bezpieczeństwo pracy;
- wynagrodzenia i dochody ludności;
- warunki mieszkaniowe;
- oświata i edukacja;
- komunikacja i łączność;
- degradacja i ochrona środowiska naturalnego.

Z kolei sytuacja demograficzna zostanie oceniona na podstawie danych charakteryzujących ruch naturalny ludności:

- współczynnik obciążeń demograficznych,
- wskaźnik struktury starości populacji,
- współczynnik feminizacji,
- współczynnik zgonów niemowląt,
- współczynnik zawierania małżeństw,
- współczynnik dynamiki demograficznej,
- współczynnik dzietności ogólnej.

Do oceny siły związku badanych zmiennych zostanie wykorzystany współczynnik korelacji Pearsona.

Dane wykorzystane w opracowaniu pochodzą z Głównego Urzędu Statystycznego, a zakres czasowy badania obejmuje rok 2010.

Słowa kluczowe: poziom życia, demografia, rozwój

The level of life and the diversity of the demographic situation in the districts of the Wielkopolska voivodship

The demographic structure of the region has a significant impact on population processes occurring in a given society, and these in turn determine other areas of life functioning. They affect,

inter alia, on the situation on the labor market, education etc. The demographic situation is inextricably linked to the level of life of the population. People with large families to a lesser extent use the services of cultural, sports, recreation than people without children. Many indicators of prosperity is connected to the number of inhabitants of the area. Just mention the GDP per capita, the number of doctors per 100 people, or the number of completed dwellings per 1000 inhabitants. Other variables, such as those associated with environmental pollution or access to medical services affect mortality.

The aim of the paper is to explore the development of the level of life in the counties of Wielkopolska voivodship, as well as diversify the demographic situation in the studied administrative units. It has been following research hypotheses:

- there is a spatial differentiation of population living in districts of the wielkopolska;
- demographic processes occurring in the wielkopolska province are characterized by considerable diversity;
- some aspects of the demographic study areas are closely related to the level of living of the population.

To assess the diversity of the phenomena studied will use the standard gauge by Hellwig. The level of life will be evaluated on the basis of a number of variables that characterize the following areas:

- protect the health and welfare;
- labor market conditions and occupational safety;
- wages and incomes of the population;
- housing conditions;
- education;
- transport and communications;
- degradation and environmental protection.

The demographic situation will be assessed on the basis of data characterizing the natural movement of population:

- demographic dependency ratio,
- percentage of population in old age,
- rate of feminization,
- infant mortality rate,
- coefficient of marriages,

- coefficient of demographic dynamics,
- the total fertility rate.

To assess the strength of the studied variables will be used Pearson's correlation coefficient. The data used in the development originate from the Central Statistical Office, and the time range covers the year 2010 survey.

Key words: level of life, demography, development

Szymkowiak Marcin (Uniwersytet Ekonomiczny w Poznaniu)

Konstrukcja estymatorów kalibracyjnych wartości globalnej dla różnych funkcji odległości

W badaniach statystycznych prowadzonych przez urzędy statystyczne braki odpowiedzi stanowią jeden z istotnych problemów, który wpływa na jakość zebranych danych. Jedną z metod umożliwiających redukcję obciążenia i zwiększenie precyzji szacunku na skutek występowania braków informacji jest kalibracja, której podstawy teoretyczne zostały zaproponowane przez Devilla i Särndala (1992). W klasycznym ujęciu wyznaczanie wag kalibracyjnych oparte jest na odpowiednio dobranej funkcji odległości, która minimalizuje odległość między wyjściowymi wagami wynikającymi ze schematu losowania próby a tzw. wagami kalibracyjnymi. W praktycznych zastosowaniach wykorzystywana jest przy tym funkcja odległości oparta na statystyce chi-kwadrat. W artykule przedstawione zostaną inne funkcje odległości, które można wykorzystać na etapie konstrukcji wag kalibracyjnych oraz metoda ich wyznaczania. W części empirycznej z wykorzystaniem rzeczywistych danych ukazane zostaną wyniki badań symulacyjnych, których celem było porównanie przedstawionych funkcji odległości służących do wyznaczania wag kalibracyjnych.

Słowa kluczowe: kalibracja, wagi kalibracyjne, estymatory kalibracyjne, braki odpowiedzi, funkcja odległości

Construction of calibration estimators of total for different distance measures

Missing data are one of the major types of non-random errors in statistical surveys. They produce significantly biased results and can considerably affect the survey quality. As a rule, this problem is evident in all kinds of surveys conducted by statistical offices of many countries where lack of response to certain survey questions is quite normal, although definitely undesirable from the point of view of estimation. In view of the above, recent years have seen a growing interest in various methods, which are designed to offset the negative effect of missing data. One of these methods is calibration, which is successfully used by statistical offices of many countries to handle missing data and was proposed by Deville and Särndal (1992).

In its classical form, calibration is a method, in which calibrated weights are computed by minimizing a distance measure between the initial sampling weights and new weights, which need to satisfy certain calibration constraints. In practice, one of the most common distance measures is based on chi-square statistics. The main goal of this paper is to present the construction of calibration estimators for total for different types of distance measures. Its empirical part, based on a simulation study using real data, provides a comparison of different types of distance measures that can be used in the construction of calibration weights.

Key words: calibration, calibration weights, calibration estimators, nonresponse, distance measures

BIBLIOGRAPHY

1. Deville J-C., Särndal C-E. (1992), "Calibration Estimators in Survey Sampling", Journal of the American Statistical Association, Vol. 87, 376–382.
2. Särndal C-E., Lundström S. (2005), "Estimation in Surveys with Nonresponse", John Wiley & Sons, Ltd.
3. Särndal C-E. (2007), "The Calibration Approach in Survey Theory and Practice", Survey Methodology, Vol. 33, No. 2, 99–119.

Śliwicki Dominik (Urząd Statystyczny w Bydgoszczy)

Czynniki determinujące poziom wynagrodzenia. Analiza ekonometryczna

Wynagrodzenie jest zapłatą za wykonaną przez pracownika na rzecz pracodawcy pracę z tytułu zawartej umowy. Wyróżnia się cztery podstawowe funkcje wynagrodzeń: społeczną, kosztową, motywacyjną, dochodową. Funkcja społeczna wiąże się ze znaczeniem płac dla określania pozycji społecznej, prestiżu, poczucia ważności i użyteczności. Funkcja kosztowa dotyczy w szczególności pracodawców, ponieważ dla nich płaca jest kosztem prowadzenia działalności gospodarczej. Funkcja motywacyjna wynagrodzeń przejawia się w skłonności ludzi do podejmowania pracy. Funkcja dochodowa dotyczy przede wszystkim pracowników, ponieważ płaca jest dla nich źródłem dochodu, który służy do zaspakajania potrzeb materialnych i niematerialnych pracowników oraz ich rodzin.

Badania rynku pracy realizowane przez statystykę publiczną, instytucje badawcze, ośrodki naukowe dostarczają informacji o poziomie i strukturze wynagrodzeń, ich rozmieszczeniu przestrzennym i rozwoju w czasie. Niejednokrotnie badania te zmierzają w kierunku identyfikacji czynników wpływających na wielkość wynagrodzeń. Jednym z najbardziej użytecznych narzędzi w badaniu takiego typu jest model ekonometryczny, który identyfikuje istotne czynniki wpływające na poziom wynagrodzeń oraz określa kierunek i wielkość tego wpływu.

W referacie zostanie przedstawiony ekonometryczny model czynników wpływających na wielkość wynagrodzeń, przy czym wielkość wynagrodzeń zostanie przedstawiona w postaci zmiennej skategoryzowanej. Kategoriami będą przedziały wyznaczone przez klasyczne i pozycyjne miary średnie oraz miary rozproszenia. Dane wykorzystane do oszacowania modelu będą pochodzić z badania struktury wynagrodzeń według zawodów.

Słowa kluczowe: uporządkowany logitowy model ekonometryczny, wynagrodzenie brutto

Factors determining the level of wages. The econometric analysis

Salary is wages made by an employee for an employer to work under the contract. There are four basic functions of wages: a social, cost, motivational, profitable. Social function associated with the importance of wages in determining social status, prestige, a sense of validity and usefulness. The cost function applies in particular to employers because their wages is the cost of doing business. Incentive pay function is manifested in the tendency of people to work. Profitable function concerns the income of workers, because wages are a source of income, which is used to cater for the needs of tangible and intangible workers and their families.

Labor market researches carried out by public statistics, research institutions, scientific centers provide information about the level and structure of remuneration, their spatial distribution and development over time. Often these studies are aimed towards identifying the factors affecting the amount of wages. One of the most useful tools in the study of this type is an econometric model that identifies significant factors affecting the level of wages and determines the direction of the magnitude of this effect.

The paper will be presented an econometric model of the factors affecting the amount of wages, the amount of wages will be presented as a categorical variable. Categories will be determined by the classic and position measures of the average and dispersion. The data used to estimate the model will come from the structure of earnings survey by occupation.

Key words: ordered logit econometric model, the gross wage

BIBLIOGRAPHY

1. Borkowska S., Wynagrodzenie godziwe, Instytut Pracy i Spraw Socjalnych, Warszawa 1999.
2. Gruszczyński M., Mikroekonometria. Modele i metody analizy danych indywidualnych, pod red., Wolters Kluwer Polska, Warszawa, 2010.
3. Kufel T., Ekonometria-Rozwiązywanie problemów z wykorzystywaniem programu Gretl, PWN, Warszawa, 2011.
4. Maddala G.S., Ekonometria, PWN, Warszawa, 2008.

Traat Imbi (University of Tartu)

Domain estimators calibrated on reference survey

It happens frequently nowadays that several surveys are carried out on the same population. Typically, some variables are common in two or more surveys. It is natural to require that estimates based on the common variables are consistent/coherent with each other in the different surveys. More specifically, in our case, population totals, or totals of larger domains, for certain variables are estimated in one survey, called the reference survey (RFS). Estimates based on the same variables but at a more detailed domain level are obtained from a second (later) survey, called the present survey (PRS). Consistency is required for the PRS estimates; domain totals for common study variables have to sum up to the corresponding estimated totals in the RFS. In a special case, the latter totals need no estimation; they are known exactly from registers. The RFS, or the registers, cannot by themselves provide the desired domain estimates due to the lacking domain identifiers.

We assume that the finite population U is divided into disjoint and exhaustive domains U_d , $d = 1, 2, \dots, D$. The probability sample from U is s and its part falling into U_d is s_d . The aim is to estimate the vector of domain totals $\mathbf{Y}_d = \sum_{U_d} \mathbf{y}_k$, for $d \in \{1, 2, \dots, D\}$, consistently with $\mathbf{Y}_0$, the population total of $\mathbf{y}_k$, known exactly or estimated from another source. This means that we construct the weights w_{ck} such that the estimators

$$\hat{\mathbf{Y}}_d = \sum_{s_d} w_{ck} \mathbf{y}_k \text{ satisfy } \sum_{d=1}^D \hat{\mathbf{Y}}_d = \mathbf{Y}_0.$$

We use the calibration framework in a more general setting. Correspondingly, the expression for weights is

$$w_{ck} = w_k [1 + (\mathbf{Z} - \hat{\mathbf{Z}})' \mathbf{M}^{-1} \mathbf{z}_k] \text{ with } \mathbf{M} = \sum_s w_k \mathbf{z}_k \mathbf{z}_k', \quad (1)$$

where $\mathbf{Z} = \sum_U \mathbf{z}_k$, $\hat{\mathbf{Z}} = \sum_s w_k \mathbf{z}_k$ and $\mathbf{z}_k$ is a vector-variable. Our aim is achieved by choosing $\mathbf{z}_k$ and w_k in a suitable way. Three sets of weights are considered and compared in the presentation. An ordinary auxiliary variable calibration weight is a special case of (1), but it does not give the desired consistency with $\mathbf{Y}_0$.

Key words: Common variables, Coherence, Multi-survey, Weighting

BIBLIOGRAPHY

1. Särndal, C.-E., Traat, I. (2011). Domain estimators calibrated on information from another survey. *Acta et Commentationes Universitatis Tartuensis de Mathematica*, vol. 15, pp. 43-60.

2. Traat, I. (2010). Comparison of Calibration-based consistent domain estimators. In Carlson, Nyquist and Villani (eds), *Official Statistics - Methodology and Applications in Honour of Daniel Thorburn*, pp. 137-147. Available at officialstatistics.wordpress.com.

Trzpiot Grażyna, Orwat-Acedańska Agnieszka (Uniwersytet Ekonomiczny w Katowicach)

Klasyfikacja funduszy inwestycyjnych ze względu na styl zarządzania - podejście regresji kwantylowej

Modele Analizy Stylu umożliwiają badanie wpływu pewnych czynników (udziałów stylu) reprezentujących inwestycje w określone klasy aktywów na osiągnięte przez fundusz stopy zwrotu. Klasyczny model analizy Stylu Sharpe'a daje niepełny obraz zależności, koncentrując uwagę jedynie na centralnej części rozkładu zmiennej objaśnianej - jest to model warunkowej wartości oczekiwanej z ograniczeniami na parametry, które szacowane są metodą najmniejszych kwadratów. Fakt ten może mieć poważne konsekwencje dla poprawnego wnioskowania na temat wpływu czynników modelu na zmiany zmiennej objaśnianej, zwłaszcza w sytuacjach braku normalności rozkładu składnika losowego stóp zwrotu, asymetrii rozkładu, istnienia grubych ogonów, obserwacji nietypowych lub niepewności, co do kształtu i rodzaju rozkładu.

Uogólniamy klasyczną analizę stylu Sharpe'a na kwantylową analizę stylu, która pozwala na obserwację zależności pomiędzy stopami zwrotu funduszy a czynnikami ryzyka w odniesieniu do całego rozkładu prawdopodobieństwa.

Celem pracy jest badanie wpływu pewnych czynników na cały warunkowy rozkład stop zwrotu jednostek uczestnictwa dwóch klas funduszy inwestycyjnych oraz ocena przydatności podejścia kwantylowego w analizie stylu tych funduszy. Realizacji celu służy analiza porównawcza struktury udziałów stylu dla różnych rzędów kwantyla oraz analiza porównawcza wyników klasyfikacji funduszy względem udziałów stylu oszacowanych klasycznie i kwantylowo. Implementacja kwantylowej analizy stylu w odniesieniu do polskich funduszy inwestycyjnych zrównoważonych oraz funduszy inwestycyjnych akcji ma na celu badanie, czy kwantylowa analiza stylu w równym stopniu jest przydatna do analizy stylu zarządzania tych dwóch klas funduszy.

Słowa kluczowe: kwantylowa Analiza Stylu, regresja kwantylowa, fundusze inwestycyjne zrównoważone, fundusze inwestycyjne akcji

The classification of mutual funds based on the management style - quantile regression approach

Style Analysis allows to assess impact of same risk factors (style weights) representing investment in various asset classes on funds' rate of returns. Classical Sharpe Style Analysis model offers only partial view on the dependencies since it focuses only on the central part of dependent variable endogenous distribution - it is conditional expected value model with restrictions and

parameters are estimated by least squares method. When error term is non-normal, asymmetric, fat-tailed or in presence of outliers it may have serious consequences for correct inference on the factors impact on endogenous variable. We implement extension the Sharpe style analysis to Quantile Style Analysis which allows to investigate dependencies between fund returns and the risk factors for the whole quantile distribution of the returns. This approach is robust to deviations from the classical assumptions necessary for the Sharpe style analysis (a distribution of error terms is left unspecified, which is the main virtue of the method as far as robustness to outliers is concerned).

The main aim of the paper is to assess usefulness of the quantile approach for the style analysis of balanced and equity mutual funds. We do this by classifying the funds according to estimated style weights for different quantile orders in quantile style analysis models. Employing hierarchical classification methods we check whether the procedure generates similar clusters for different quantile orders of a given fund. For the balanced funds the clustering procedure generates totally different classes for different quantiles. The differences are significantly smaller among equity funds. Since clusters can refer to different levels of risk and the differences depend on quantile orders it is incorrect to restrict one's attention to the classical style analysis for the analysed mutual funds.

Key words: Quantile Style Analysis, quantile regression, mutual balanced funds, equity mutual funds

BIBLIOGRAPHY

1. Hao L., Najman D.Q : Quantile Regression, „Quantitative Application in the Social Sciences”, (2007).
2. Kim T., White H., Stone D.: Asymptotic and Bayesian Confidence Intervals for Sharpe-Style Weights. *Journal of Financial Econometric*. 3(3). 315-343 (2005)
3. Trzpiot G., Orwat - Acedańska A.. The classification of Polish mutual balanced funds on the management style - quantile regression approach *Theory and Applications of Quantitative Methods, Econometrics* 31, *Prace Naukowe UE* nr 194, str. 9-23, (2011)

Ulman Paweł (Uniwersytet Ekonomiczny w Krakowie)

Wpływ aktywności zawodowej osób niepełnosprawnych na sytuację ekonomiczną ich gospodarstw domowych

Aktywność ekonomiczna osób niepełnosprawnych w Polsce jest jedną z najniższych w Europie, co potwierdzają wyniki wielu badań statystycznych pokazując, że sytuacja osób niepełnosprawnych na rynku pracy jest trudniejsza niż sytuacja osób sprawnych. Relatywnie osoby niepełnosprawne są częściej bezrobotne i zdecydowanie częściej bierne zawodowo niż osoby sprawne nawet wtedy, gdy weźmiemy pod uwagę osoby w wieku produkcyjnym.

Niska aktywność zawodowa osób niepełnosprawnych może wpływać negatywnie na sytuację ekonomiczną ich gospodarstw domowych. Z badań nad dochodami i wydatkami gospodarstw z osobami niepełnosprawnymi wynika, że sytuacja dochodowa tych gospodarstw jest istotnie gorsza niż gospodarstw osób sprawnych. Różny jest również poziom i struktura wydatków tych gospodarstw. W szczególnej sytuacji są gospodarstwa, w których niepełnosprawną osobą jest dziecko. Wtedy często bierność zawodowa dotyczy jednego z rodziców, który jest zmuszony do rezygnacji z pracy w celu opieki nad dzieckiem. Niewątpliwie jest to ukryty koszt niepełnosprawności.

W pracy podjęto próbę określenia jak sytuacja zawodowa niepełnosprawnych członków gospodarstwa wpływa na jego dochody oraz wydatki. W tym celu wykorzystane zostały charakterystyki rozkładu dochodów i wydatków - przy uwzględnieniu skali ekwiwalentności - wyznaczone dla różnych grup gospodarstw domowych ze względu na skład osobowy gospodarstwa i sytuację jego członków na rynku pracy. Wykorzystano również modele regresyjne dla dochodów na osobę lub na jednostkę ekwiwalentną, w których wprowadzono zmienne opisujące aktywność ekonomiczną osób niepełnosprawnych. W ostatnim kroku podjęto próbę estymacji parametrów modelu logitowego opisującego prawdopodobieństwo zdarzenia, że osoba niepełnosprawna jest bezrobotna przy wykorzystaniu jako jednej ze zmiennych objaśniających sytuacji materialnej jej gospodarstwa domowego. Otwartym pozostaje problem na ile nisko oceniana sytuacja materialna gospodarstwa jest wynikiem niepodjęcia pracy przez osobę niepełnosprawną, a na ile impulsem do poszukiwania tejże pracy.

Wszystkie analizy zostały przeprowadzone na podstawie indywidualnych danych z badania budżetów gospodarstw domowych zrealizowanego w 2010 r. W ramach tego badania GUS zbiera informacje nie tylko na temat dochodów i wydatków gospodarstw domowych, ale również na temat cech społeczno-demograficznych członków tych gospodarstw oraz ich sytuacji na rynku pracy. Między innymi uzyskujemy informacje o niepełnosprawności w gospodarstwach domowych, co umożliwia dogłębną analizę związków tego zjawiska z sytuacją ekonomiczną osób jak i gospodarstw domowych, w których żyją.

Wawrowski Łukasz (Urząd Statystyczny w Poznaniu)

Analiza ubóstwa w przekroju powiatów w województwie wielkopolskim z wykorzystaniem metod statystyki małych obszarów

Ubóstwo jest zjawiskiem bardzo złożonym. Problem badany jest w Polsce już od ponad stu lat, ale gwałtowne przyspieszenie prac nastąpiło w momencie transformacji gospodarczej. Ograniczanie tego niekorzystnego zjawiska jest obecnie najważniejszym celem Banku Światowego. Jednak, aby tę misję realizować potrzebne są metody identyfikacji ubogich. W praktyce stosowanych jest wiele wskaźników mierzących m.in. zasięg, głębokość czy dotkliwość ubóstwa. Charakterystyki te są jednak dostarczane na bardzo ogólnym poziomie. Otrzymanie bardziej szczegółowych informacji jest możliwe dzięki zastosowaniu statystyki małych obszarów (SMO). Jest to zbiór metod pozwalających na estymację parametrów przy małej liczebności próby z wy-

korzystaniem wszelkich dostępnych źródeł informacji. Jej dodatkowym atutem jest możliwość zapewnienia pokrycia informacyjnego na niskich poziomach agregacji przestrzennej.

Głównym celem artykułu jest próba oszacowania wskaźników opisujących ubóstwo na poziomie powiatów w województwie wielkopolskim przy użyciu wybranych metod SMO. Obliczenie tych miar odbędzie się na podstawie danych udostępnianych przez GUS charakteryzujących poziom życia gospodarstw domowych. Taka estymacja pozwoli uzyskać kompleksowej informacji na poziomie lokalnym dotyczącej sfery ubóstwa.

Słowa kluczowe: estymacja, mała próba, statystyka małych obszarów, ubóstwo

Poverty analysis across districts of Wielkopolska using small area estimation methods

Poverty is a very complex phenomenon. It has been researched in Poland for over one hundred years, but the start of the economic transformation has been followed by a sudden acceleration in research. A reduction of this negative phenomenon is the main goal of the World Bank now. To achieve this objective, however, it is necessary to develop methods of identifying people affected by poverty. In practice, many indicators are applied, which measure among others the scope, depth or intensity of poverty. Unfortunately, these indicators are produced at a very general level. Getting more detailed information is possible thanks to the application of small area estimation (SAE). It is a set of methods which enable the estimation of parameters for a small sample size by making use of all available information sources. Its additional advantage is information coverage at a low level area.

The main goal of this paper is to present an attempt to estimate poverty indicators at district-level in Wielkopolska using selected methods of SAE. These measures were calculated on the basis of statistical data about the standard of living of households in Poland provided by the Central Statistical Office. Such estimation can provide comprehensive information about poverty at a local level.

Key words: estimation, small sample size, small area estimation, poverty

BIBLIOGRAPHY

1. Rao J.N.K., 2003, Small Area Estimation, Wiley, New York.
2. Panek T., 2011, Ubóstwo, wykluczenie społeczne i nierówności, Szkoła Główna Handlowa w Warszawie, Warszawa.

Wawrzyniak Katarzyna (Zachodniopomorski Uniwersytet Technologiczny w Szczecinie)

O cena stopnia realizacji celów głównych strategii rozwoju kraju według województw z wykorzystaniem analizy korespondencji

W pracach (Wawrzyniak 2011a, 2011b) wykazano, że analiza korespondencji dla małych prób może stanowić użyteczne narzędzie diagnostyczne zarówno w makro-, jak i w mikroskali. W obu przypadkach wykorzystano procedurę podwajania obserwacji i uzyskano informacje o mocnych i słabych stronach badanych obiektów (województw oraz spółek giełdowych) ze względu na wyróżnione zmienne diagnostyczne. W obecnym artykule w analizie korespondencji (Gatnar, Walesiak 2004; Stanimir 2005) wykorzystana zostanie złożona macierz znaczników utworzona na podstawie porównania wartości wskaźników osiągniętych w poszczególnych województwach z ich wartością zakładaną (normą) w Strategii na konkretny moment czasu. Zaproponowana metoda umożliwi realizację głównego celu artykułu, czyli ocenę stopnia realizacji celów głównych Strategii Rozwoju Kraju (SRK) w województwach, a tym samym wskazanie tych województw, w których cele osiągnięto oraz tych, w których cele nie zostały zrealizowane. Dzięki zastosowanej metodzie będzie możliwa również wizualizacja wyników w przestrzeni dwuwymiarowej. Badanie zostanie przeprowadzone na podstawie danych statystycznych Głównego Urzędu Statystycznego dotyczących wskaźników monitorujących realizację celów SRK według województw [5].

Słowa kluczowe: analiza korespondencji, złożona macierz znaczników, diagnozowanie ilościowe, Strategia Rozwoju Kraju

The evaluation of main goals of the national development strategy according to voivodships with using the correspondence analysis

In the articles (Wawrzyniak 2011a, 2011b) one showed that the correspondence analysis for small samples could be a useful diagnostic tool in macro- and microeconomics. In both cases the procedure with doubling was used and the weaknesses and strengths of the analyzed objects (voivodships and companies listed on Warsaw Stock Exchange) were found. In the current paper the data in correspondence analysis are organized in the multiple indicator matrix (Gatnar, Walesiak 2004; Stanimir 2005). The multiple indicator matrix is created by comparing the real values of indices in the voivodships with the norms from the Strategy. The aim of this paper is to evaluate the degree of the realization of the goals of the Strategy in voivodships and to identify those voivodships in which the Strategy have been implemented and those ones in which have been not. Application of the correspondence analysis enables also the two-dimensional visualization of the results. The data used in research came from the Central Statistical Office [5].

Key words: correspondence analysis, multiple indicator matrix, quantitative diagnosing, the National Development Strategy

BIBLIOGRAPHY

1. Gatnar E. (red.), Walesiak M. (red.) (2004), „Metody statystycznej analizy wielowymiarowej w badaniach marketingowych”, Wydawnictwo Akademii Ekonomicznej we Wrocławiu, Wrocław.
2. Stanimir A. (2005), „Analiza korespondencji jako narzędzie do badania zjawisk ekonomicznych”, Wydawnictwo Akademia Ekonomiczna we Wrocławiu, Wrocław.
3. Wawrzyniak K. (2011a), „Analiza korespondencji jako narzędzie diagnostyczne w makroskali”, w: J. Dziechciarz (red.), *Ekonometria. Zastosowanie metod ilościowych*, nr 30, *Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu* nr 163, s. 19-27, Wydawnictwo Uniwersytetu Ekonomicznego we Wrocławiu, Wrocław.
4. Wawrzyniak K. (2011b), „Diagnoza sytuacji finansowo-ekonomicznej spółek giełdowych z wykorzystaniem klasycznej analizy korespondencji”, w: K. Brzozowska (red.), *Oeconomia, Folia Pomeranae Universitatis Technologiae Stetinensis* nr 285 (62), s. 105-116, Wydawnictwo Zachodniopomorskiego Uniwersytetu Technologicznego w Szczecinie, Szczecin.
5. http://www.stat.gov.pl/gus/wskazniki_monitorujace_PLK_HTML.htm

Wiatrak Andrzej Piotr (Uniwersytet Warszawski)

Aktualne problemy statystyki rolniczej

Akcesja Polski do Unii Europejskiej spowodowała zmianę w prowadzonych badaniach statystycznych, w tym dotyczących rolnictwa. Obecnie trwa proces dostosowań w tym zakresie, co jest przedmiotem referatu. Ukazano następujące zagadnienia: istotę i specyfikę statystyki rolniczej, gospodarstwa rolne, ich typologia i struktura, produkcja rolnicza i bilanse produktów rolnych, rachunki ekonomiczno-rolnicze, zatrudnienie, praca i dochody w rolnictwie oraz ceny produktów rolniczych w powiązaniu ze specyfiką polskiego rolnictwa, a zwłaszcza występującego rozdrobnienia, bezrobocia utajonego i braku związków większości gospodarstw z rynkiem. Przedstawiając te zagadnienia wskazano na potrzebę ściślejszego powiązania informacji gromadzonych przez GUS oraz Ministerstwo Rolnictwa i Rozwoju Wsi wraz z jego agendami. Ponadto wskazano, że procesy doskonalenia bazy statystyczno-rolniczej powinny być oparte na zapewnieniu reprezentatywności badań w skali kraju i regionu oraz opracowaniu operatorów losowania. Takie podejście będzie sprzyjało zapewnieniu jakości wyników wykorzystywanych dla potrzeb statystyki i polityki rolnej.

Wierzbicka Anna (Instytut Statystyki i Demografii)

Analiza taksonomiczna stanu zdrowia publicznego Polski na tle wybranych krajów europejskich

Działanie Unii Europejskiej, jako organizacji międzynarodowej, nakierowane jest na poprawę zdrowia publicznego, zapobieganie chorobom i dolegliwościom ludzkim oraz usuwanie źródeł zagrożeń dla zdrowia fizycznego i psychicznego. W dążeniu do prowadzenia skutecznej polityki zdrowotnej, wspierającej działania państw członkowskich, na szczeblu unijnym gromadzone są dane oraz materiały fundamentalne dla oceny stanu zdrowia publicznego poszczególnych krajów. Dzięki temu możliwe jest dokonywanie różnorodnych analiz, istotnych z punktu widzenia oceny przemian w systemach leczniczych oraz wskazania stopnia podobieństwa poszczególnych członków Unii Europejskiej.

Celem artykułu jest analiza porównawcza stanu zdrowia publicznego Polski oraz wybranych państw europejskich. Badanie zostanie przeprowadzone w oparciu o wskaźniki medyczne, ekonomiczne i społeczne zawarte w bazie danych EUROSTATU. Zastosowanie metod taksonomicznych pozwoli na odpowiednie uporządkowanie oraz wskazanie krajów, które charakteryzują się najlepszym poziomem stanu zdrowia publicznego. Dokonanie bardziej szczegółowej oceny możliwe będzie dzięki ogólnemu omówieniu systemów opieki zdrowotnej poszczególnych państw. Oprócz Polski analizą objęta zostanie Bułgaria, Estonia, Litwa, Rumunia oraz Czechy i Węgry. Istotnym elementem będzie okres badawczy obejmujący lata 2004 – 2009, tj. od początku akcesji większości tych krajów do Unii Europejskiej. EUROSTAT nie dysponuje odpowiednimi danymi dotyczącymi pozostałych państw byłego Bloku Wschodniego, dlatego zostały one pominięte. Referat odniesie się także do kilku wybranych członków Unii Europejskiej, zaliczanych do grupy krajów rozwiniętych. Jednakże, ze względu na skomplikowaną przeszłość historyczną i związane z nią trudności społeczno – gospodarcze, analiza krajów postsocjalistycznych, w tym Polski, wydaje się ciekawsza z porównawczego punktu widzenia, dlatego też stanowić ona będzie podstawę referatu.

Słowa klucze: stan zdrowia publicznego, taksonomiczna miara rozwoju, analiza porównawcza

Taxonomic analysis of the Polish public health in comparison with selected European countries

Activity of the European Union as an international organization is directed towards improving public health, preventing human illness and diseases and obviating sources of danger to physical and mental health. In order to conduct effective health policies, supporting actions of the member states, the data and materials which are fundamental to public health assessments of individual countries are being collected at the Union level. This makes it possible to conduct a variety of analysis relevant to the valuation of changes in medical systems, and indicate the degree of similarity of individual members of the European Union.

The purpose of this article is to analyze the Polish public health in comparison with selected European countries. The study will be conducted on the basis of medical, economic and social indicators from the EUROSTAT database. The use of taxonomic methods will allow for proper arrangement and indication of the countries that are characterized by the highest level of public health. A more detailed assessment will be possible thanks to a general presentation of health care systems of each country. Apart from Poland the study will include Bulgaria, Estonia, Lithuania, Romania, Hungary, and Czech Republic. An important element will be the period of study covering the years 2004 - 2009, that is since the beginning of the accession of most of these countries into the European Union. EUROSTAT does not have adequate data on other countries of the former Eastern Bloc, which is why they were omitted. The paper will also refer to a few selected members of the European Union, belonging to the group of developed countries. However, because of the complicated historical past and connected with it socio - economic difficulties, the analysis of post-socialist countries, including Poland, seems more interesting from a comparative point of view, therefore, it will be the basis of this article.

Key words: public health, taxonomic measure of development, comparative analysis

BIBLIOGRAPHY

1. Berbeka J., *Poziom życia ludności a wzrost gospodarczy w krajach Unii Europejskiej*, Kraków: Akademia Ekonomiczna, 2006.
2. Grabiński T., Wydymus S., Zeliaś A., *Metody taksonomii numerycznej w modelowaniu zjawisk społeczno-gospodarczych*, PWN, Warszawa 1989.
3. Hellwig Z. *Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom rozwoju oraz zasoby i strukturę wykwalifikowanych kadr*. Przegląd Statystyczny 15.4.1968.
4. Zeliaś A., Malina A., *O budowie taksonomicznej miary jakości życia. Syntetyczna miara rozwoju jest narzędziem statystycznej analizy porównawczej*. Taksonomia z. 4, 1997.

Wilk Justyna (Uniwersytet Ekonomiczny we Wrocławiu)

Podejście symboliczne w analizach regionalnych

Analizy regionalne prowadzi się najczęściej z wykorzystaniem wskaźników bazujących na danych udostępnianych przez statystykę publiczną. Problemem wielu analiz tego typu jest precyzja opisu badanych zjawisk. Dane statystyczne dla jednostek terytorialnych stanowią wartości agregatowe i nie uwzględniają struktury terytorialnej zjawiska. Może to zniekształcać obraz rzeczywistej sytuacji regionu. Na przykład stopa bezrobocia w 2009 roku w województwie dolnośląskim wynosiła 12,5%, przy czym w $\frac{2}{3}$ powiatów tego województwa odnotowano znacznie wyższe wartości.

Propozycją rozwiązania tego problemu jest zastosowanie podejścia symbolicznego, w którym jest możliwość:

- opisanie obiektów (np. gmin) za pomocą zmiennych o realizacjach w postaci przedziałów wartości, zbiorów kategorii i struktur udziałowych (np. struktura procentowa ludności gminy według wieku demograficznego),
- reprezentacji obiektów niższego rzędu (np. powiatów) za pomocą obiektów wyższego rzędu, np. województw (np. przedział liczbowy przeciętnego miesięcznego wynagrodzenia brutto w powiatach województwa).

Taki sposób interpretacji zjawisk wymaga wprowadzenia specyficznych rozwiązań metodologicznych, wypracowanych przede wszystkim na gruncie analizy danych symbolicznych (zob. Diday i Noirhomme-Fraiture [2008], Gatnar i Walesiak [2011]), a także adaptacji innych metod statystycznych (zob. np. Hair i in. [2006]).

Celem referatu jest propozycja zastosowania podejścia symbolicznego w regionalnych analizach porównawczych jako sposobu na uwzględnienie terytorialnej struktury zjawisk społecznych, gospodarczych i środowiskowych. Zarówno w literaturze polskiej, jak i zagranicznej brakuje opracowania na ten temat. W referacie zaprezentowany zostanie:

- sposób reprezentacji zjawisk i jednostek terytorialnych z wykorzystaniem danych symbolicznych,
- metody jakie można zastosować w analizach regionalnych na podstawie danych symbolicznych,
- oprogramowanie w analizie danych symbolicznych,
- przykład zastosowania podejścia symbolicznego w badaniach regionalnych.

Słowa kluczowe: dane symboliczne, badania regionalne, analizy statystyczne

Symbolic approach in regional analyses

Regional analyses are most often conducted with using indicators being based on data disclosed by public statistics. A description precision of examined phenomena is a problem of many analyses. Statistical data for regional units constitute aggregate values and they are not taking a territorial structure of the phenomenon into account. It can distort an image of the real situation of a region. For example unemployment rate was 12,5% in dolnośląskie voivodeship in 2009 but there were denoted much higher values for two-thirds of its poviats.

A suggestion to solve this problem is applying symbolic approach that makes possible to:

- characterize objects (e.g. gminas) with using interval-valued variables, multivalued variables and modal variables (e.g. gmina's population structure according to demographic age),

- describe lower-level objects (e.g. gminas) by higher-level objects, e.g. voivodeships, poviats (e.g. an interval of average monthly gross wages and salaries of poviats located in a voivodeship).

This way of phenomena description requires implementing peculiar methodological solutions, in particular symbolic data analysis (see Diday and Noirhomme-Fraiture [2008], Gatnar and Walesiak [2011]) as well as adaptation of other statistical methods (see e.g. Hair et al. [2006]).

The aim of the paper is to make a proposal of using symbolic approach in regional analyses as the way to take the territorial structure of economic, social and environmental phenomena into account. In Polish as well as foreign publications drawing up to this subject is missing. The paper discusses:

- a way of phenomena and regional units description with application of symbolic data,
- methods appropriated in regional analyses being based on symbolic data,
- symbolic data analysis software,
- an example of using symbolic approach in regional investigations.

Key words: symbolic data, regional research, statistical analyses

BIBLIOGRAPHY

1. Diday E., Noirhomme-Fraiture M. (Ed.), Symbolic data analysis and the Sodas software, John Wiley & Sons, Chichester 2008.
2. Gatnar E., Walesiak M. (red.), Analiza danych jakościowych i symbolicznych z wykorzystaniem programu R [Qualitative and symbolic data analysis with using R software], PWN, Warszawa 2011.
3. Hair J.F., Black W.C., Babin B.J., Anderson R.E., Tatham R.L. (2006), Multivariate Data Analysis, Pearson Prentice Hall, New Jersey.

Witkowski Janusz, Walczak Tadeusz, Berger Jan (Główny Urząd Statystyczny)

S Statystyka publiczna - rozwój historyczny i aktualne wyzwania

W referacie postanowiliśmy przypomnieć korzenie polskiej statystyki publicznej oraz zwrócić uwagę na jej niełatwe drogi rozwoju, którego celem nadrzędnym była zawsze służba dla kraju oraz jego obywateli. Z tego względu wyodrębniono kilka najważniejszych okresów historii statystyki publicznej, akcentując te wydarzenia, które wywarły największy wpływ na realizację jej zadań oraz na jej przyszły rozwój.

Punktem wyjścia rozważań jest charakterystyka warunków, w jakich powstawały pierwsze założenia organizacji i zadań służb statystycznych po utworzeniu niepodległego Państwa Polskiego po 123 latach zaborów. Podkreślono rolę polskich statystyków, w tym działaczy Polskiego Towarzystwa Statystycznego, tworzących pierwsze opracowania i publikacje statystyczne charakteryzujące sytuację społeczno-gospodarczą na ziemiach polskich pod panowaniem mocarstw zaborczych. Opracowania te stały się jakby załącznikiem przyszłych założeń organizacji służb statystyki publicznej po odzyskaniu przez Polskę niepodległości w 1918 r.

Ważny etap rozwoju statystyki obejmuje okres międzywojenny, w którym ukształtowały się jej organizacyjne podstawy. Ogromne znaczenie miało tu utworzenie Głównego Urzędu statystycznego podległego bezpośrednio Prezydium Rady Ministrów. Przyjęto wówczas tzw. scentralizowany model organizacji badań statystycznych realizowanych w oparciu o program badań zatwierdzany przez Radę Ministrów. GUS był nie tylko głównym realizatorem tych badań, ale także otrzymał uprawnienia do koordynacji badań statystycznych prowadzonych w całym kraju. Zainicjowano wówczas wiele badań statystycznych, podjęto szereg prac metodologicznych i w zakresie standardów klasyfikacyjnych. Dzięki temu stworzono podstawę dla lepszej spójności informacji wynikowych.

O uzyskaniu wysokiego profesjonalnego poziomu pracy służb statystycznych w okresie międzywojennym świadczy przygotowanie, przeprowadzenie i opracowanie wyników dwóch kolejnych powszechnych spisów ludności w latach 1921 i 1931. Spotkało się to z uznaniem środowiska międzynarodowego czego wyrazem było przyznanie Polsce przez Międzynarodowy Instytut Statystyczny funkcji organizatora 18. Kongresu tej organizacji. Kongres pod patronatem prezydenta Rzeczypospolitej Ignacego Mościckiego odbył się w dniach 21–24 sierpnia 1929 r.

W referacie omówiono także krótko tragiczne losy polskiej statystyki w latach okupacji hitlerowskiej 1939 - 1945, a następnie rozwój statystyki w okresie po II wojnie światowej, podzielony na trzy okresy historyczne: lata 1945 do 1989, 1989 do 2004 oraz po 2004 roku.

Mimo niezwykle trudnych warunków lokalowych i kadrowych tuż po II wojnie światowej, służby statystyczne przystąpiły do organizacji pierwszych badań szczególnie użytecznych do oceny stanu ludności i majątku narodowego pozostałego po zniszczeniach wojennych. Najpilniejszą sprawą było oszacowanie liczby ludności zamieszkałej w kraju i jej terytorialnego rozmieszczenia. Pierwszy, tzw. sumaryczny spis według stanu z dnia 14 lutego 1946 r. wykazał, że ogólna liczba ludności Polski w jej nowych granicach wyniosła 23930 tys. osób. Kolejne, pełne spisy ludności przeprowadzono z częstotliwością generalnie co dziesięć lat (1950, 1960, 1970, 1978, 1988, 2002 oraz ostatni - w 2011 roku).

W okresie powojennym przystąpiono także do opracowania systemu badań statystycznych odpowiadających najpilniejszym potrzebom organów zarządzania gospodarką w nowych warunkach ustrojowych, co oznaczało m. in. konieczność dostosowania tematyki badań do wymagań wieloletnich, rocznych i kwartalnych planów gospodarczych, zgodnie z przepisami dekretu z dnia 31 lipca 1946.

Od roku 1956 rozwinęła się działalność publikacyjna, wznowiono także wydawanie roczników statystycznych, których publikowanie zaprzestano w 1951 r. Znacznie ożywiono współpracę międzynarodową. Podjęto wówczas dwu i wielostronne prace metodologiczne dotyczące porównań ważniejszych kategorii społeczno-ekonomicznych, w tym w dziedzinie rachunków narodo-

wych oraz spożycia i cen. Ważną rolę w tych pracach odgrywał Zakład Badań Statystyczno-Ekonomicznych GUS. Wyrazem uznania międzynarodowej społeczności statystyków dla aktywności GUS była decyzja Międzynarodowego Instytutu Statystycznego przyznająca Polsce zadanie zorganizowania kolejnego - 40 Kongresu tej organizacji w Warszawie w dniach 1-9 września 1975 r.

W drugiej połowie lat 1980. przystąpiono do wprowadzania bardziej wszech-stronnych zmian w statystyce publicznej. Zmiany te objęły niemal wszystkie aspekty funkcjonowania systemu informacji statystycznej. Bardziej precyzyjnie określono rolę statystyki w kraju zmierzającym do wprowadzenia gospodarki rynkowej oraz przestrzegania zasad społeczeństwa demokratycznego. Szczególny akcent położono na zapewnienie warunków uzyskiwania obiektywnych i rzetelnych informacji bazujących na społecznym zaufaniu do statystyki i partnerskich stosunkach z respondentami.

Prace nad wprowadzeniem zasadniczych zmian w statystyce nasiliły się zwłaszcza po 1990 roku w ramach przygotowania Polski do wstąpienia do UE. Oprócz głębokich zmian metodologicznych i prawnych dostosowano zakres badań i opracowań do obowiązujących programów badań i terminów dostarczania informacji w Unii Europejskiej. Prawnym fundamentem tych zmian było przyjęcie ustawy z dnia 29 czerwca 1995 r. o statystyce publicznej.

Od 1 maja 2004 r., po wstąpieniu Polski do Unii Europejskiej polska statystyka publiczna funkcjonuje jako pełnoprawny członek Europejskiego Systemu Statystycznego (ESS). Z tego faktu wynika cały szereg nowych obowiązków, ale także uprawnień i możliwości rozwojowych. W obecnych warunkach do zadań tych należą w szczególności:

- ciągłe doskonalenie zakresu i metodologii badań w celu zaspokojenia zmieniających się potrzeb informacyjnych użytkowników,
- dalsza modernizacja procesu badań statystycznych z uwzględnieniem różnicowania źródeł danych, w tym szerszego wykorzystania źródeł administracyjnych oraz zastosowania osiągnięć technologii informatycznych,
- wzmocnienie pozycji polskiej statystyki na arenie międzynarodowej,
- zagwarantowanie odpowiedniej jakości statystyki zgodnie z wymogami Europejskiego Kodeksu Praktyk Statystycznych,
- doskonalenie pomiaru postępu społecznego, dobrobytu oraz zrównoważonego rozwoju, wychodząc poza tradycyjny pomiar wzrostu gospodarczego mierzonego PKB.

Słowa kluczowe: statystyka oficjalna, historia statystyki, metodologia badań, Europejski System Statystyczny, jakość statystyki

The Polish official statistics its historical development and current challenges

The paper presents a historical overview of the Polish Official Statistics and focuses on challenges encountered in its development, which has always been primarily treated as service for

the country and its citizens. Most significant periods in the history of Official Statistics are identified, each highlighting events with the biggest influence on the performance of its duties and its future development.

The paper begins with a description of conditions that shaped the underlying assumptions about the organization and tasks of statistical services in the newly formed independent state of Poland after 123 years of partitions. It stresses the role of Polish statisticians, including members of the Polish Statistical Association, who pre-pared the first reports and statistical materials about the socio-political situation in the ethnically Polish regions under foreign rule. This early work constituted the basis of the future organizing principles of statistical services after Poland regained independence in 1918.

Another important period comprises the years before World War II, when the organizational structure of Official Statistics was established. This refers, in particular, to the creation of the Central Statistical Office (CSO), reporting directly to the Prime Minister. It was then that the so-called 'centralized model of organizing statistical surveys' was adopted, which was to be conducted according to a survey programs approved by the Council of Ministers. CSO was not only charged with the task of conducting surveys but also coordinating surveys across the whole country. Many survey programs were started at that time, accompanied by methodological studies and development of classification standards. All this helped to lay the foundations for the production of more coherent statistical output.

The high professional standard of statistical services in the period of 1918-1939 was demonstrated during the preparation, administration and analysis of two censuses in 1921 and 1931. This achievement received international recognition, which was marked by International Statistical Institute's decision to grant Poland the task of organizing its 18th Congress. The Congress, enjoying the patronage of the President of Poland, Ignacy Mościcki, was held between 21 and 24 of August 1929.

The paper continues with a brief look at the tragic circumstances in the Nazi-occupied Poland and progress made in the field of statistics in the years following the war, focusing on three periods: 1945-1989, 1989-2004 and after 2004.

Despite extremely difficult operating conditions and limited staff after World War II, statistical services started to conduct their first surveys in order to assess the state of population and national wealth and resources in the war-ravished country. The most urgent need was to estimate the population size and its territorial distribution. The first, so-called 'summary' census at 14 February 1946 put the country's population at 23,930,000 people. Consecutive censuses were conducted more or less every decade (1950, 1960, 1970, 1978, 1988, 2002 and the last one in 2011).

During the post-war period work was started to develop a system of statistical surveys that would meet the needs of authorities responsible for managing the country's economy; this meant adapting the survey program in accordance with long-term, annual and quarterly economic plans.

In 1956 publishing activity was restarted, including the publication of statistical yearbooks, which was discontinued in 1951. Several bilateral and multilateral comparative methodological

studies were initiated, for example in the field of national accounts, consumption and prices. In recognition of the work conducted by CSO, International Statistical Institute selected Poland as the organizer of its 40th congress, which was held in Warsaw, between 1–9 September 1975.

In the late 1980s Official Statistics underwent sweeping changes. The changes were intended to redefine the role of statistics in a country, which was about to switch to free market economy and become a democratic society. The matter of particular significance was to create conditions for obtaining objective and reliable information and to build up trust towards statistics and develop partner relationships with respondents.

Following 1990, the reforming efforts within the field of Official Statistics were intensified in preparation for Poland's accession to the EU. In addition to far-reaching methodological and legal changes, survey programs and information provision deadlines were adapted to those existing in the EU.

Since the moment of joining the EU on 1 May, 2004, Polish Official Statistics has been a full member of the European Statistical System (ESS). This entails a number of duties as well as privileges and development prospects. Under new circumstances, these tasks include:

- ongoing development in terms of survey scope and methodology to meet the changing needs of information users;
- further modernization of the statistical survey process and making use of various data sources, particularly more reliance on administrative data and advances in information technology;
- strengthening the position of Polish statistics internationally;
- ensuring the required quality of statistical information in accordance with the European Statistics Code of Practice
- improving ways of measuring social progress, welfare and sustainable development, which go beyond the traditional measure of economic progress by means of GDP.

Key words: Official Statistics, history of statistics, survey methodology, European Statistical System, statistical quality

Wróblewska Wiktoria (Szkół Główna Handlowa w Warszawie)

Zmiany maksymalnego trwania życia i długowieczność - wyzwania dla statystyki

Od blisko dwustu lat notowany jest w krajach rozwiniętych systematyczny wzrost trwania życia. Od 1950 roku liczba osób dożywających setnych urodzin wzrasta dwukrotnie w ciągu kolejnych dziesięciu lat. W literaturze przedmiotu toczy się dyskusja nad perspektywami dalszego wzrostu

trwania życia. Środowisko badaczy dzieli się na dwie grupy. Jedną, której głównym przedstawicielem jest James Vaupel, uważa, że nie ma naturalnego limitu dla trwania życia i poziom e_0 będzie w dalszym ciągu wzrastać, o czym świadczy elastyczność granicy trwania życia widoczna w przekraczaniu dotychczas wyznaczanych limitów (Oeppen i Vaupel 2002, Vaupel i in. 2003, Vaupel i in. 1998, Thatcher i in. 1998). Druga, z S.Jay Olshanskim, wskazuje na ograniczenia utrzymania dotychczasowego tempa wzrostu trwania życia w przyszłości, na co wskazuje m.in. występowanie zjawiska entropii tablic trwania życia oraz rektangularyzacja krzywej przeżycia, co świadczy o zbliżaniu się do biologicznej granicy trwania życia (Olshansky i in. 2001, Olshansky i in. 1990, Fries 1980, 2003, Hayflick 1981).

Celem pracy jest analiza długowieczności i zmian, które zaszły od połowy XIX wieku w maksymalnym oczekiwanym trwaniu życia na świecie. W szczególności omówione zostaną granice życia osiągnięte przez najdłużej żyjące osoby oraz dyskusja, która toczy się w literaturze przedmiotu nad granicami trwania życia. Na przykładzie "niebieskich stref" zaprezentowany zostanie stan wiedzy w zakresie behawioralnych przyczyn długowieczności. Przedstawione zostaną także wyzwania wynikające dla statystyki w związku z analizowanymi zmianami w trwaniu życia i perspektywami długowieczności.

Słowa kluczowe: trwanie życia, długowieczność, determinanty, perspektywy

Changes in life expectancy and longevity - challenges for statistics

Human life expectancy has increased at a steady pace in the industrialized countries for almost two hundred years. Since 1950 the number of people celebrating their 100th birthdays has at least doubled each decade. There is a discussion going on in the literature on the further increase of life expectancy. There are two different approaches to the future evolution of human longevity. The first one, represented by James Vaupel, assumes that there is no biological limit to life expectancy and level of e_0 will continue to grow. Evidence of this is the flexibility of the estimates of maximum life expectancy at birth made by experts in the past (Oeppen i Vaupel 2002, Vaupel et al. 2003, Vaupel et al. 1998, Thatcher . 1998). The second approach represented by S.Jay Olshansky, shows the limitations on the current growth rate of life expectancy in the future, including entropy of the life table and rectangularisation of the survival curve (Olshansky et al. 2001, Olshansky et al. 1990, Fries 1980, 2003, Hayflick 1981).

The main goal of this study is to analyze the human longevity and its changes since the mid-nineteenth century in the world. The biological limit of life achieved by the longest living people in the world and the literature discussion on the maximum life expectancy (e_0) will be presented. We will describe the state of knowledge about the behavioral determinants of longevity on the example of the blue zones. The challenges for statistics in relation to changes in life expectancy and longevity prospects will also be presented.

Key words: life expectancy, longevity, determinants, future

BIBLIOGRAPHY

1. Hayflick L. 1981, Prospects for Human Life Extension by Genetic Manipulation, [in:]

- „Aging: a Challenge to Science and Society”, Vol. 1, D. Danon, N.W. Shock, M. Marois (eds). New York: Oxford University Press.
2. Fries J.F., 1980, Aging, natural death, and the compression of morbidity, „The New England Journal of Medicine” 303:130-135.
 3. Fries F.J., 2003, Measuring and Monitoring Success in Compressing Morbidity, „Annals of Internal Medicine” 139: 455-459.
 4. Oeppen J., Vaupel J.W., 2002, Broken Limits to Life Expectancy, „Science” 296:1029-1031.
 5. Olshansky S.J., Carnes B.A., Casseli C., 1990, In Search of Methuselah: Estimating the Upper Limits to Human Longevity, „Science” 250:634-640.
 6. Olshansky S.J., Carnes B.A., Désesquelles A., 2001, Prospects for Human Longevity, „Science” 291:1491-1492.
 7. Thatcher A.R., Kannisto V., Vaupel J.W., 1998, The Force of Mortality at Ages 80 to 120, Monographs on Population Aging Series V5. Odense, Denmark: Odense University Press.
 8. Vaupel J., Carey J., Christensen K., Johnson T., Yashin A., Holm V., Iachine J., Kannisto V., Khazaeli A., Liedo P., Longo V., Zeng Y., Manton K., Curtsinger J., 1998, Biodemographic Trajectories of Longevity, „Science” 280:855-860.
 9. Vaupel J.W., Carey J.R., Christensen K., 2003, Aging. It's Never Too Late, „Science” 301:1679-1681.

Wyszkowska Dorota (Uniwersytet w Białymstoku, Urząd Statystyczny w Białymstoku)

Statystyka regionalna jako instrument wsparcia rozwoju polskich województw

Od momentu akcesji naszego kraju do Unii Europejskiej nowego znaczenia nabrały regiony oraz rozwój regionalny. W krajach Wspólnoty regionom przypisuje się bowiem istotną rolę w ogólnej polityce rozwojowej. W historii istniała nawet koncepcja utworzenia Europy Regionów, w której narodowe państwa europejskie byłyby zastąpione regionami, odpowiednikami dzisiejszych polskich województw. Tego typu samorządowe jednostki powinny być, zdaniem twórców tej koncepcji, podmiotem w stosunkach z UE. Umacnianiu roli regionów w Unii służy postępująca decentralizacja w państwach członkowskich oraz powszechnie obowiązująca zasada subsydiarności.

W wyniku reaktywowania w Polsce w roku 1999 samorządowych województw pojawiła się w naszym kraju jednostka, która swoim potencjałem i zakresem przydzielonych zadań odpowiada europejskim regionom. Ustawodawca nałożył na nią obowiązek realizacji zadań związanych z prowadzeniem polityki rozwoju regionalnego, zmierzającej do poprawy jej konkurencyjności.

Pojęcie rozwoju regionalnego w literaturze przedmiotu definiowane jest w sposób różnorodny. Za A. Klasikiem można je określić jako trwały wzrost trzech elementów: potencjału gospodarczego regionów, ich siły konkurencyjnej oraz poziomu i jakości życia mieszkańców.

Władze samorządowych województw działając w warunkach ograniczonych środków finansowych, dążąc do zapewnienia rozwoju regionu, stale stają przed dylematami: jaki kierunek zaangażowania środków wybrać, jakie inicjatywy wesprzeć, czy wydatkowanie środków na dany cel przynosi oczekiwane rezultaty, jakie działania wesprzeć ze środków bezzwrotnej pomocy unijnej?

W odpowiedzi na tak postawione pytania powinna być przydatna statystyka regionalna, którą można zdefiniować jako część statystyki publicznej, dla której przedmiotem badań jest region (w Polsce - województwo) lub mniejsze jednostki samorządowe. Wykorzystywanie różnorodnych metod statystycznych w prowadzonych badaniach powinno służyć ocenie przestrzennych aspektów zmian gospodarczych, społecznych i ekologicznych. Zbierane i przetwarzane informacje dotyczące zjawisk występujących w regionach powinny zapewnić rzetelną podstawę podejmowania decyzji rozwojowych, jak też oceny wcześniej realizowanych działań.

Zakres rozwoju polskiej statystyki regionalnej oraz niska świadomość możliwości wykorzystania jej zasobów powoduje jednak, że jest ona dość rzadko wykorzystywana w procesach podejmowania decyzji rozwojowych. Należy także zauważyć, iż w niektórych zakresach użytkownicy informacji statystycznych dostrzegają braki danych lub ich znaczne opóźnienia w dostępności, co dodatkowo potęguje niski stopień wykorzystania statystyki regionalnej w procesach decyzyjnych.

Stąd celem niniejszego artykułu będzie wskazanie możliwości wykorzystania danych statystycznych (statystyki regionalnej) w realizacji polityki rozwoju regionalnego i lokalnego oraz zasygnalizowanie niedostatków statystyki regionalnej w tym zakresie.

Słowa kluczowe: rozwój regionalny, miary rozwoju, statystyka regionalna

Regional statistics as a tool of supporting the development of Polish voivodships

Since the accession of Poland to the European Union, regions and regional development have taken on a new meaning. Regions play an important role in the general development policy of the member states. Historically speaking, there was once an idea of creating a Europe of the Regions, in which states could have been replaced by regions - equivalents of Polish voivodships existing at present. Self-government entities of this type should have been, according to the creators of the idea, a party in the relations with the EU. Advancing decentralization and a binding principle of subsidiarity make the role of the regions in the EU stronger.

As a result of coming back to self-government voivodships in 1999, an entity equal to European

regions in terms of its potential and the scope of functions has emerged. Legislators obliged the entity to perform the tasks connected with carrying out the regional development policy leading to the improvement of its competitiveness.

The term of 'regional development' has been defined in a various ways in the literature on the subject. Following A. Klasik's definition, it is 'a constant development of three elements: economic potential of regions, their competitive power as well as the level and quality of life of their inhabitants.

The authorities of self-government voivodships act under financial pressure while carrying out the task of regional development. They constantly need to solve the dilemmas: where to invest funds, which initiatives to support, whether the funding of a certain aim is bringing the expected results, which projects to support from the non-refundable Union financial aid.

To answer questions posed in this way, regional statistics comes in handy. This statistics is defined here as part of public statistics dealing with regions (voivodships in Poland) or smaller self-government entities. The use of various statistical methods in research should help to evaluate spatial aspects of economic, social and ecological changes. Collected and processed data on phenomena in regions should form a solid base for taking development decisions as well as for evaluating activities previously carried out.

The scope of development of Polish regional statistics together with a low awareness of the use of its resources still result in its infrequent use in the processes of development decision making. It should be stated that users notice the data unavailability or a significant delay in the publishing of statistical data concerning certain fields, which makes for the even less frequent use of regional statistics in decision making processes.

The aim of the article is to present the possibilities of use of statistical data (regional statistics) in carrying out the regional and local development policies as well as pointing out the weaknesses of regional statistics within the same scope.

Key words: regional development, measures of development, regional statistics

BIBLIOGRAPHY

1. Chądzyński J., Nowakowska A., Przygodzki Z, Region i jego rozwój w warunkach globalizacji, CEDEWU, Warszawa 2007
2. Ekonomiczne i organizacyjne instrumenty wspierania rozwoju lokalnego i regionalnego Rozwój, innowacyjność, infrastruktura, Dylewski Marek (red.), Zeszyty Naukowe nr 501 Ekonomiczne Problemy Usług nr 22, Wydawnictwo Naukowe Uniwersytetu Szczecińskiego, Szczecin 2008
3. Filipiak B., Kogut M., Szewczuk A., Ziolo M, Rozwój lokalny i regionalny. Uwarunkowania, finanse, procedury, Fundacja na rzecz Uniwersytetu Szczecińskiego, Szczecin 2005
4. Kokocińska K., Polityka regionalna w Polsce i w Unii Europejskiej, Wydawnictwo Naukowe Uniwersytetu im. Adama Mickiewicza, Poznań 2010

5. Ministerstwo Rozwoju Regionalnego, Krajowa Strategia Rozwoju Regionalnego 2010 - 2020. Regiony - Miasta - Obszary wiejskie, Warszawa 2010
6. Ministerstwo Rozwoju Regionalnego, Raport o rozwoju i polityce regionalnej, Warszawa 2007
7. Wlazlak K., Rozwój regionalny jako zadanie administracji publicznej, Oficyna Wydawnicza Wolters Kluwers, Warszawa 2010
8. Oręziak L., Finansowanie rozwoju regionalnego w Polsce, Wydawnictwo Wyższej Szkoły Handlowej im. R. Łazarskiego, Warszawa 2008
9. Strahl D., Metody oceny rozwoju regionalnego, Wydawnictwo Akademii Ekonomicznej im. Oskara Langego we Wrocławiu, Wrocław 2006

Wywił Janusz (Uniwersytet Ekonomiczny w Katowicach)

Estymacja średniej na podstawie próby prostej porządkowanej za pomocą zmiennej dodatkowej

Problem estymacji wartości średnie (globalnej) w populacji skończonej i ustalonej jest rozważany. Estymacja jest prowadzona na podstawie próby dobieranej zwrotnie za pomocą schematu losowania zależnego od statystyki pozycyjnej zmiennej dodatkowej, która jest silnie skorelowana ze zmienną badaną. Wyznaczono parametry zaproponowanego estymatora oraz wyznaczono jego rozkład graniczny. Do dokładność estymatora porównano z precyzją znanych estymatorów, takich jak regresyjny, na podstawie odpowiednio zaprojektowanego komputerowego badania symulacyjnego. Wynikiem tej analizy są wnioski dotyczące praktycznego stosowania proponowanej strategii estymacji.

Estimation of population mean on the basis of a simple sample ordered by auxiliary variable

Estimation of the population average in a finite population by means of sampling strategies dependent on an order statistic of an auxiliary variable is considered. The inclusion probabilities are dependent on the probability function of the order statistic of the auxiliary variable highly correlated with a variable we are interested in. The simple estimator of the population mean (total) of the variable under study has been proposed. Its basic parameters were derived. The limit distribution of the estimator has also been considered. Finally, on the basis of simulation analysis, the accuracy of the estimator has been compared with precision of the well known estimators like the simple regression one. This leads to the conclusion that the proposed statistic can be a useful estimator in a lot of practical statistical research.

Zdaniewicz Witold (Instytut Statystyki Kościoła Katolickiego w Warszawie)

Wkład Instytutu Statystyki Kościoła Katolickiego w polską statystykę publiczną

Kościół katolicki obejmuje większość mieszkańców Polski. Jednak już sama przynależność do Kościoła wiąże się z wieloma statystycznymi (pomiarowymi) problemami. Katolicy są włączeni w społeczny organizm Kościoła i zajmują w nim odpowiednie pozycje społeczne. W ramach tych pozycji odgrywają indywidualne role społeczne, gdzie szczególnie znaczenie ma ich aktywność religijna. W społecznym organizmie Kościoła szczególną rolę pełni parafia, która w pewnym sensie odpowiada społeczności lokalnej. Stanowi ona również bardzo dogodny kanał obserwacji statystycznej polskiego społeczeństwa. Wykorzystuje go od 40 lat Instytut Statystyki Kościoła Katolickiego. W wyniku jego pracy parafia stała się jednostką badawczą dla opracowań statystycznych i socjologicznych. Instytut podejmuje we współpracy z wieloma ośrodkami naukowymi wiele badań m.in. badania dominantes i communicantes (od 1980 r.) , badania trzeciego sektora oraz kultury.

Słowa kluczowe: religia, statystyka, Kościół katolicki

The contribution of the Institute of Catholic Church Statistics to the Polish public statistics

The Catholic Church covers most of the Polish population. However, the membership of the Church involves a number of statistical (measurement) problems. Catholics are involved in the social organism of the Church and take its proper position in society. In these positions play individual roles in society, with particular importance of their religious activity. In the social body, the Church's special role plays the parish, which in a sense corresponds to the local community. It is also very convenient channel of statistical observation of the Polish society. The Institute of the Catholic Church Statistics uses his Chanel as a research unit for statistical and sociology studies. The Institute in cooperation with many research centers conducts many research such as dominantes and communicantes (since 1980), research on the culture and on the third sector.

Key words: religion, statistics, Catholic Church

Zhang Li-Chun (Statistics Norway)

Micro calibration for data integration

The upcoming register-based Census 2011 in several European countries presents an intriguing setting of data integration. On the one hand, a register-based population data file is made available for cross tabulation of a large set of social-demographic variables on detailed aggregation levels; on the other hand, many similarly measured statistical variables are collected in the ongoing national survey programs which, typically, are regarded as somewhat closer to the target concept of interest. In micro calibration, one adjusts the census variables on the micro data level, by a minimum amount in some sense, so that the distribution of each adjusted census variable agrees with what can be inferred from the corresponding ‘target’ variable in a relevant survey, possibly jointly with a chosen set of covariates, in which sense the adjusted census variable achieves statistical validity (Zhang, 2011). Moreover, by means of micro calibration, the statistical variables that initially exist only in separate surveys become ‘integrated’ and jointly available in a single data set.

Micro calibration represents thus an alternative to statistical matching, which constructs the joint information of multiple variables, the parts of which are only observed in separate (usually, two) datasets that have no, or virtually no, overlapping units. Given that there are actually no observations of the target joint distribution, statistical matching proceeds under the conditional independence assumption (CIA). That is, conditional on the covariates that are observable in both data sources, the statistical variables from separate sources are independent of each other. It follows that the actual integrated dataset is ‘centred’ about the expected joint data set under the CIA, whereas in micro calibration it would be close to the initial register-based census data set. In other words, the information in the different survey data is combined with the information in the initial census data in micro calibration, whereas in statistical matching only with the hypothetical CIA but no empirical information.

In this talk I will explain how micro calibration can be achieved for categorical as well as continuous data. Also to be discussed are some uses of the micro calibrated data, such as for the purpose of small area estimation.

BIBLIOGRAPHY

1. Zhang, L.-C. (2011). Topics of statistical theory for register-based statistics and data integration. *Statistica Neerlandica*, Early view available online.

Zieliński Wojciech (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie)

Dobór liczności próby w kontrolach prowadzonych przez NIK

W statystycznej kontroli jakości badane obiekty poddane są ocenie alternatywnej i badaczka interesuje odsetek obiektów ocenionych negatywnie. Jednym z takich zastosowań są kontrole jakości prowadzone przez Najwyższą Izbę Kontroli, których celem jest m. in. wykrywanie nieprawidłowości w pracy urzędów skarbowych. Karliński (2003) podaje matematyczne szczegóły prowadzenia kontroli dla cech ocenianych alternatywnie, a w szczególności wymagania odnośnie do liczebności próbki, by uzyskać żadaną dokładność oceny oraz ryzyko błędu. Podany sposób wyznaczania minimalnej liczności próby oparty jest na przybliżeniach asymptotycznych. W pracy pokazany jest sposób wyznaczania minimalnej liczności próby bazujący na dokładnych rozkładach skończenie próbkowych.

Słowa kluczowe: statystyczna kontrola jakości, ocena alternatywna, plan badania

Statistical properties of a control design of controls provided by Supreme Chamber of Control

In statistical quality control objects are alternatively rated. It is of interest to estimate a fraction of negatively rated objects. One of such applications is a quality control provided by Supreme Chamber of Control to find out a percentage of abnormalities in the work among others of tax offices. Mathematical details of experimental designs for alternatively rated phenomena are given in Karliński (2003). There are given requirements for sample sizes, numbers of negative rates, accuracy of estimation and error risks. In the paper, some statistical properties of given plans are investigated.

Key words: statistical quality control, alternative rating, experimental design

BIBLIOGRAPHY

1. Karliński W. 2003: <http://nik.gov.pl/docs/kp/kontrolapanstwowa-200305.pdf>

Źródło Tomasz (Uniwersytet Ekonomiczny w Katowicach)

O predykcji wartości globalnych dla domen skokrelowanych przestrzennie

Rozważany jest problem predykcji wartości globalnych w domenach w badaniach wielookresowych. Zaproponowany model nadpopulacji jest przypadkiem szczególnym Ogólnego Mieszanego Modelu Liniowego, gdzie zakładana jest korelacja przestrzenna i korelacja w czasie składników losowych. Ponadto zakłada się, że populacja oraz przynależność elementów populacji do domen mogą się zmieniać w czasie.

W statystyce małych obszarów niezależność pomiędzy domenami jest standardowym założeniem, co jest szczególnie istotne w związku z estymacją błędu średniokwadratowego. Ponadto, wiadomo, że przy tym założeniu i występowaniu korelacji przestrzennej pomiędzy domenami „zyski z uwzględnienia struktury przestrzennej w statystyce małych obszarów nie okazują się duże” jak zauważają w przypadku badań prowadzonych w jednym okresie np. Chandra, Salvati, Chambers (2007) s. 343. W artykule będzie rozważany problem korelacji przestrzennej w grupach domenach w przypadku badań wielookresowych.

Słowa kluczowe: statystyka małych obszarów, badania wielookresowe, korelacja czasowo-przestrzenna

On prediction of totals for spatially correlated domains

The problem of prediction of domain totals based on longitudinal data is considered. The proposed superpopulation model is special case of the General Linear Mixed Model where spatial and temporal correlation of random components is assumed. Moreover, it is assumed that the population can change in time and domains' affiliation of population elements can change in time.

In the small area estimation the standard assumption is the independence between domains, which is very important because of the problem of the MSE estimation. Moreover, under this assumption and spatial correlation within domains it is known, that "the gains from inclusion of spatial structure in small area estimation do not appear to be large" as indicated for surveys in one period for example by Chandra, Salvati and Chambers (2007) p. 343. In the paper the problem of spatial correlation but within groups of domains will be studied for longitudinal data.

Key words: small area estimation, longitudinal surveys, spatial and temporal correlation

BIBLIOGRAPHY

1. Chandra H., Salvati N., Chambers R., Small area estimation for spatially correlated populations - a comparison of direct and indirect model-based method, „Statistics in Transition - new series”, 2007, 8, 2, 331-350.

Kontakt

PRZEWODNICZĄCA KOMITETU ORGANIZACYJNEGO

dr hab. Elżbieta Gołata prof. nadzw. UEP
e-mail: elzbieta.golata@ue.poznan.pl

SEKRETARIAT KOMITETU ORGANIZACYJNEGO

Urząd Statystyczny w Poznaniu
ul. J. H. Dąbrowskiego 79, 60-959 Poznań
e-mail: kongres2012@stat.gov.pl
+48 61 27 98 325 - (pon-pt: 8-15)
+48 61 27 98 343 - (pon: 9-12, wt: 7-10, śr: 7-10)
+48 61 27 98 266
fax +48 61 27 98 101

Lista uczestników

Nazwisko i Imię	E-mail	Instytucja
Adamczewski Wojciech	<i>w.adamczewski@stat.gov.pl</i>	Zakład Wydawnictw Statystycznych
Andrzejczak Emilia	<i>e.andrzejczak@stat.gov.pl</i>	Główny Urząd Statystyczny
Andrzejczak Karol	<i>karol.andrzejczak@put.poznan.pl</i>	Politechnika Poznańska
Augustyniak Zbigniew	<i>z.augustyniak@stat.gov.pl</i>	Urząd Statystyczny w Poznaniu
Auksztol Jerzy	<i>SekretariatUSGDK@stat.gov.pl</i>	Urząd Statystyczny w Gdańsku
Bal-Domańska Beata	<i>beata.bal – domanska@ue.wroc.pl</i>	Uniwersytet Ekonomiczny we Wrocławiu
Barczak Andrzej Stanisław	<i>abar@ae.katowice.pl</i>	Uniwersytet Ekonomiczny w Katowicach
Bartniczak Bartosz	<i>bbartniczak@gmail.com</i>	Urząd Statystyczny we Wrocławiu
Bartosiewicz Stanisława	<i>stanisawa.bartosiewicz@gmail.com</i>	Wyższa Szkoła Bankowa we Wrocławiu
Basiura Beata	<i>bbasiura@zarz.agh.edu.pl</i>	Akademia Górniczo-Hutnicza w Krakowie
Batóg Barbara	<i>batog@wneiz.pl</i>	Uniwersytet Szczeciński
Batóg Jacek	<i>batog@wneiz.pl</i>	Uniwersytet Szczeciński
Bąk Iwona	<i>iwona.bak@zut.edu.pl</i>	Zachodniopomorski Uniwersytet Technologiczny w Szczecinie
Bednarski Tadeusz	<i>t.bednarski@prawo.uni.wroc.pl</i>	Uniwersytet Wrocławski
Beręsewicz Maciej	<i>maciej.beresewicz@ue.poznan.pl</i>	Uniwersytet Ekonomiczny w Poznaniu
Berger Jan	<i>J.Berger@stat.gov.pl</i>	Główny Urząd Statystyczny
Białas Tomasz	<i>T.Bialas@stat.gov.pl</i>	Główny Urząd Statystyczny
Białek Jacek	<i>jbialek@uni.lodz.pl</i>	Uniwersytet Łódzki
Bielak Renata	<i>r.bielak@stat.gov.pl</i>	Główny Urząd Statystyczny
Bielecka Anna	<i>bielenda@alk.edu.pl</i>	Akademia Leona Koźmińskiego
Bieniek Monika	<i>M.Bieniek@stat.gov.pl</i>	Główny Urząd Statystyczny
Bieszk-Stolorz Beata	<i>stolorz@interia.pl</i>	Uniwersytet Szczeciński
Billari Francesco	<i>francesco.billari@unibocconi.it</i>	Bocconi University, Milan
Błackowska Anna	<i>anna.blackowska@wsb.wroclaw.pl</i>	Wyższa Szkoła Bankowa we Wrocławiu

Błażejowski Marcin	<i>marcin.blazejowski@wsb.torun.pl</i>
Bojarski Jacek	<i>j.bojarski@wmie.uz.zgora.pl</i>
Borawska Monika	<i>M.Borawska@stat.gov.pl</i>
Borucka Jadwiga	<i>jadwiga.borucka@gmail.com</i>
Borys Tadeusz Józef	<i>tadeusz.borys@ue.wroc.pl</i>
Brzezińska Justyna	<i>justyna.brzezinska@ue.katowice.pl</i>
Budzyński Ireneusz	<i>i.budzynski@stat.gov.pl</i>
Bulska Ewa Bogusława	<i>ebulska@wp.pl</i>
Caliński Tadeusz	<i>calinski@up.poznan.pl</i>
Ceranka Bronisław	<i>bronicer@up.poznan.pl</i>
Chambers Raymond	<i>ray@uow.edu.au</i>
Chlebicki Paweł	<i>pchlebicki@esripolska.com.pl</i>
Chłoń-Domińczak Agnieszka	<i>Agnieszka.Chlon@gmail.com</i>
Chromińska Maria	<i>rektor@wshir.pl</i>
Cierpiał-Wolan Marek	<i>SekretariatUSrze@stat.gov.pl</i>
Cmela Piotr Ryszard	<i>sekretariatUSLDZ@stat.gov.pl</i>
Czyżewski Andrzej	<i>a.czyzewski@ue.poznan.pl</i>
Decker Reinhold	<i>rdecker@wiwi.uni – bielefeld.de</i>
Dehnel Grażyna	<i>g.dehnel@ue.poznan.pl</i>
Deville Jean-Claude	<i>deville@ensae.fr</i>
Diering Magdalena	<i>magdalena.diering@put.poznan.pl</i>
Dmochowska Halina	<i>h.dmochowska@stat.gov.pl</i>
Dobrowolska Anna	<i>sekretariat – pk@stat.gov.pl</i>
Domański Czesław	<i>czedoman@uni.lodz.pl</i>
Doniec Dorota	<i>D.Doniec@stat.gov.pl</i>
Dudek Hanna	<i>hanna_dudek@sggw.pl</i>
Dyzmann-Sroka Agnieszka	<i>agnieszka.dyzmann – sroka@wco.pl</i>
Dziechciarz Józef	<i>joze.f.dziechciarz@ue.wroc.pl</i>

Wyższa Szkoła Bankowa w Toruniu
Uniwersytet Zielonogórski
Urząd Statystyczny w Olsztynie
Szkoła Główna Handlowa w Warszawie
Uniwersytet Ekonomiczny we Wrocławiu
Uniwersytet Ekonomiczny w Katowicach
Główny Urząd Statystyczny
Oddział Warszawski PTS
Uniwersytet Przyrodniczy w Poznaniu
Uniwersytet Przyrodniczy w Poznaniu
University of Wollongong
ESRI
Szkoła Główna Handlowa w Warszawie
Wyższa Szkoła Handlu i Rachunkowości
Urząd Statystyczny w Rzeszowie
Urząd Statystyczny w Łodzi
Uniwersytet Ekonomiczny w Poznaniu
Bielefeld University
Uniwersytet Ekonomiczny w Poznaniu
ENSAE
Politechnika Poznańska
Główny Urząd Statystyczny
Główny Urząd Statystyczny
Uniwersytet Łódzki
Urząd Statystyczny w Katowicach
Szkoła Główna Gospodarstwa Wiejskiego w Warszawie
Wielkopolskie Centrum Onkologii
Uniwersytet Ekonomiczny we Wrocławiu

Fattorini Lorenzo	<i>lorenzo.fattorini@unisi.it</i>
Filas-Przybył Sylwia	<i>s.filas@stat.gov.pl</i>
Franceschi Sara	<i>franceschi2@unisi.it</i>
Frątczak Ewa Zofia	<i>ewaf@sgh.waw.pl</i>
Furmańczyk Konrad	<i>konfur@wp.pl</i>
Gabińska Magdalena	<i>m.gabinska@stat.gov.pl</i>
Gamrot Wojciech	<i>wojciech.gamrot@ue.katowice.pl</i>
Ganczarek-Gamrot Alicja	<i>alicja.ganczarek – gamrot@ue.katowice.pl</i>
Gatnar Eugeniusz	<i>eugeniusz.gatnar@nbp.pl</i>
Genge Ewa	<i>ewa.witek@ue.katowice.pl</i>
Gerasymenko Sergiy	<i>serguey@wp.pl</i>
Ghosh Malay	<i>ghoshm@stat.ufl.edu</i>
Głowiński Cezary	<i>cezary.glowinski@sas.com</i>
Gołata Elżbieta	<i>elzbieta.golata@ue.poznan.pl</i>
Gorzelańczyk Krystyna	<i>kgorz@wp.pl</i>
Goujon Anne	<i>Anne.Goujon@oeaw.ac.at</i>
Górecki Tomasz	<i>tomasz.gorecki@amu.edu.pl</i>
Groenen Patrick	<i>groenen@few.eur.nl</i>
Gruchociak Hanna	<i>hanna.gruchociak@ue.poznan.pl</i>
Grzelak Aleksander	<i>agrzelak@interia.pl</i>
Grzenda Wioletta	<i>wgrzend@sgh.waw.pl</i>
Grześkowiak Alicja	<i>alicja.grzeskowiak@ue.wroc.pl</i>
Hanusik Krystyna	<i>khanusik@uni.opole.pl</i>
Hetmańska Aurelia	<i>SekretariatUsKce@stat.gov.pl</i>
Hozér Józef	<i>mhk@wneiz.pl</i>
Hozér-Koćmiel Marta	<i>mhk@wneiz.pl</i>
Idzik Marcin	<i>midzik@pentor.pl</i>
Jajuga Krzysztof	<i>krzysztof.jajuga@ue.wroc.pl</i>

University of Siena
Urząd Statystyczny w Poznaniu
University of Tuscia
Szkoła Główna Handlowa w Warszawie
Szkoła Główna Gospodarstwa Wiejskiego w Warszawie
Urząd Statystyczny w Białymstoku
Uniwersytet Ekonomiczny w Katowicach
Uniwersytet Ekonomiczny w Katowicach
Narodowy Bank Polski
Uniwersytet Ekonomiczny w Katowicach
Wyższa Szkoła Bankowa w Poznaniu
University of Florida
SAS Institute Sp. z o.o.
Uniwersytet Ekonomiczny w Poznaniu
Kolegium Pracowników Służb Społecznych
Vienna Institute of Demography
Uniwersytet im. Adama Mickiewicza w Poznaniu
Siena University
Uniwersytet Ekonomiczny w Poznaniu
Uniwersytet Ekonomiczny w Poznaniu
Szkoła Główna Handlowa w Warszawie
Uniwersytet Ekonomiczny we Wrocławiu
Uniwersytet Opolski
Urząd Statystyczny w Katowicach
Uniwersytet Szczeciński
Uniwersytet Szczeciński
TNS Pentor, SGGW w Warszawie
Uniwersytet Ekonomiczny we Wrocławiu

Jakóbiak Krzysztof	<i>krakow@stat.gov.pl</i>	Urząd Statystyczny w Krakowie
Janczur-Knapke Magdalena	<i>M.Janczur@stat.gov.pl</i>	Główny Urząd Statystyczny
Januszewska Jovita	<i>m.smialek@stat.gov.pl</i>	Centrum Informatyki Statystycznej
Jarosz Lidia	<i>l.jarosz@stat.gov.pl</i>	Urząd Statystyczny w Kielcach
Jaworski Stanisław	<i>stanislaw_jaworski@sggw.pl</i>	Szkoła Główna Gospodarstwa Wiejskiego w Warszawie
Jefmański Bartłomiej	<i>bartlomiej.jefmanski@wp.pl</i>	Uniwersytet Ekonomiczny we Wrocławiu
Jeznach Maria	<i>m.jeznach@stat.gov.pl</i>	Główny Urząd Statystyczny
Jędrzejczak Alina	<i>jedrzej@uni.lodz.pl</i>	Uniwersytet Łódzki
Józefowski Tomasz	<i>t.jozefowski@stat.gov.pl</i>	Urząd Statystyczny w Poznaniu
Jóźwiak Janina	<i>ninaj@sgh.waw.pl</i>	Szkoła Główna Handlowa w Warszawie
Jurečková Jana	<i>jurecko@karlin.mff.cuni.cz</i>	Charles University in Prague
Jurek Anna	<i>annamariajurek@gmail.com</i>	Uniwersytet Łódzki
Jurkiewicz Tomasz	<i>jurkiewicz@wzr.ug.edu.pl</i>	Uniwersytet Gdański
Kalinowski Marcin	<i>mkalinowski@pluton.wsb.gda.pl</i>	Wyższa Szkoła Bankowa w Gdańsku
Kamińska-Gawryluk Ewa	<i>E.Kaminska@stat.gov.pl</i>	Urząd Statystyczny w Białymstoku
Katulska Krystyna	<i>krakat@amu.edu.pl</i>	Uniwersytet im. Adama Mickiewicza
Kaźmierczak Maciej	<i>m.kazmierczak@stat.gov.pl</i>	Urząd Statystyczny w Poznaniu
Keilman Nico	<i>nico.keilman@econ.uio.no</i>	University of Oslo
Kisielińska Joanna	<i>joanna_kisielinska@sggw.pl</i>	Szkoła Główna Gospodarstwa Wiejskiego w Warszawie
Klimanek Tomasz	<i>t.klimanek@ue.poznan.pl</i>	Uniwersytet Ekonomiczny w Poznaniu
Kobus Paweł	<i>pawel_kobus@sggw.pl</i>	Szkoła Główna Gospodarstwa Wiejskiego w Warszawie
Kołodziejczak Barbara	<i>bkolodziejczak@ump.edu.pl</i>	Uniwersytet Medyczny im. K. Marcinkowskiego w Poznaniu
Komar-Morawska Agnieszka	<i>A.Komar@stat.gov.pl</i>	Główny Urząd Statystyczny
Kończak Grzegorz	<i>grzegorz.konczak@ue.katowice.pl</i>	Uniwersytet Ekonomiczny w Katowicach
Kopczewska Katarzyna	<i>kkopczewska@wne.uw.edu.pl</i>	Uniwersytet Warszawski
Korczyński Adam	<i>korczynskiadam@gmail.com</i>	Szkoła Główna Handlowa w Warszawie
Kordos Jan	<i>jan1kor2@aster.pl</i>	Wyższa Szkoła Menedżerska
Kornecki Janusz	<i>kornecki@uni.lodz.pl</i>	Urząd Statystyczny w Łodzi

Koronacki Jacek	<i>korona@ipipan.waw.pl</i>
Kosiorowski Daniel	<i>daniel.kosiorowski@uek.krakow.pl</i>
Koska Anna	<i>a.koska@stat.gov.pl</i>
Kot Stanisław Maciej	<i>skot@zie.pg.gda.pl</i>
Kowaleski Jerzy Tadeusz	<i>kowaleski@lodz.msk.pl</i>
Kowalewski Grzegorz	<i>grzegorz.kowalewski@ue.wroc.pl</i>
Kowalewski Jacek	<i>j.kowalewski@stat.gov.pl</i>
Kowalka Ewa	<i>e.kowalka@stat.gov.pl</i>
Kowalska Małgorzata	<i>M.Kowalska@stat.gov.pl</i>
Kowalski Krzysztof	<i>sekretariatUSWAW@stat.gov.pl</i>
Kowerski Mieczysław	<i>mkowerski@wszia.edu.pl</i>
Kozłowska Zofia	<i>sekretariatUSWAW@stat.gov.pl</i>
Kraj Mariusz	<i>M.Kraj@stat.gov.pl</i>
Krasucki Piotr Paweł	<i>piotr_krasucki@yahoo.com</i>
Kraśniewska Wacława	<i>w.krasniewska@stat.gov.pl</i>
Królikowska Barbara	
Kruszka Kazimierz	<i>k.kruszka@stat.gov.pl</i>
Krzanowski Wojciech	<i>W.J.Krzanowski@exeter.ac.uk</i>
Krzykowski Grzegorz	<i>grzegorzkrzykowski@gmail.com</i>
Krzyśko Mirosław	<i>mkrzyisko@amu.edu.pl</i>
Kufel Tadeusz	<i>tadeusz.kufel@umk.pl</i>
Kujawińska Agnieszka	<i>agnieszka.kujawinska@put.poznan.pl</i>
Kukło Cezary	<i>cz.kuklo@life.pl</i>
Kuna-Broniowska Izabela	<i>izabela.kuna@up.lublin.pl</i>
Kupis-Fijałkowska Aleksandra	<i>a.fijałkowska@uni.lodz.pl</i>
Kupiszewski Marek	<i>m.kupisz@twarda.pan.pl</i>
Kurkiewicz Jolanta	<i>kurliej@uek.krakow.pl</i>
Kurkowski Krzysztof	<i>M.Kozłowska@stat.gov.pl</i>

Instytut Podstaw Informatyki PAN
Uniwersytet Ekonomiczny w Krakowie
Urząd Statystyczny w Opolu
Politechnika Gdańska
Uniwersytet Łódzki
Uniwersytet Ekonomiczny we Wrocławiu
Urząd Statystyczny w Poznaniu
Urząd Statystyczny w Poznaniu
Główny Urząd Statystyczny
Urząd Statystyczny w Warszawie
Wyższa Szkoła Zarządzania i Administracji w Zamościu
Urząd Statystyczny w Warszawie
Centrum Badań i Edukacji Statystycznej
Instytut Spraw Publicznych
Główny Urząd Statystyczny
Polskie Towarzystwo Informatyczne
Urząd Statystyczny w Poznaniu
University of Exeter
Wyższa Szkoła Bankowa w Gdańsku
Uniwersytet im. Adama Mickiewicza w Poznaniu
Uniwersytet Mikołaja Kopernika w Toruniu
Politechnika Poznańska
Uniwersytet w Białymstoku
Uniwersytet Przyrodniczy w Lublinie
Uniwersytet Łódzki
Środkowoeuropejskie Forum Badań Migracyjnych i Ludnościowych
Uniwersytet Ekonomiczny w Krakowie
Główny Urząd Statystyczny

Kuropka Ireneusz	<i>ireneusz,kuropka@ue.wroc.pl</i>
Kuźmicka Janina	<i>J.Kuzmicka@stat.gov.pl</i>
Kwiatkowska-Ciotucha Dorota	<i>dorota.kwiatkowska@ue.wroc.pl</i>
Laczka Éva	<i>eva.laczka@ksh.hu</i>
Laskowska Joanna	<i>joanna_ewa_stec@yahoo.com</i>
Ledwina Teresa	<i>ledwina@imipan.pan.wroc.pl</i>
Lemmi Achille	<i>lemmi@unisi.it</i>
Lewiński vel Iwański Stanisław	<i>stanislaw.lewinski@pwr.wroc.pl</i>
Liberda Barbara	<i>barbara.liberda@uw.edu.pl</i>
Łagodziński Władysław Wiesław	<i>w.lagodzinski@stat.gov.pl</i>
Łangowska-Szczeńiak Urszula	<i>uls@uni.opole.pl</i>
Łazowska Bożena	<i>b.lazowska@stat.gov.pl</i>
Łączyński Artur	<i>A.Laczynski@stat.gov.pl</i>
Łomiak Anna	<i>A.Lomiak@stat.gov.pl</i>
Łuczak Aleksandra	<i>luczak@up.poznan.pl</i>
Łyson Piotr	<i>p.lyson@stat.gov.pl</i>
Majdzińska Anna	<i>a_majdzinska@uni.lodz.pl</i>
Malaga Krzysztof	<i>k.malaga@ue.poznan.pl</i>
Malesa Ewa	<i>e.malesa@stat.gov.pl</i>
Malecka Marta	<i>marta.malecka@uni.lodz.pl</i>
Markowicz Iwona	<i>imarkowicz@wneiz.pl</i>
Markowska Małgorzata	<i>malgorzata.markowska@ue.wroc.pl</i>
Markowski Krzysztof	<i>k.markowski@stat.gov.pl</i>
Matuszewska-Janica Aleksandra	<i>aleksandra.matuszewska@sggw.pl</i>
Mejza Stanisław	<i>smejza@up.poznan.pl</i>
Meller Ewa	<i>ewa.meller@yahoo.com</i>
Mielniczuk Jan	<i>miel@ipipan.waw.pl</i>
Migdał-Najman Kamila	<i>kmn@wzr.pl</i>

Uniwersytet Ekonomiczny we Wrocławiu
Urząd Statystyczny w Opolu
Uniwersytet Ekonomiczny we Wrocławiu
Hungarian Central Statistical Office
Oddział Warszawski PTS
Instytut Matematyczny PAN
Siena University
Politechnika Wrocławska
Uniwersytet Warszawski
Polskie Towarzystwo Statystyczne
Uniwersytet Opolski
Centralna Biblioteka Statystyczna
Główny Urząd Statystyczny
Główny Urząd Statystyczny
Uniwersytet Przyrodniczy w Poznaniu
Główny Urząd Statystyczny
Uniwersytet Łódzki
Uniwersytet Ekonomiczny w Poznaniu
Główny Urząd Statystyczny
Uniwersytet Łódzki
Uniwersytet Szczeciński
Uniwersytet Ekonomiczny we Wrocławiu
Urząd Statystyczny w Lublinie
Szkoła Główna Gospodarstwa Wiejskiego w Warszawie
Uniwersytet Przyrodniczy w Poznaniu
Urząd Statystyczny w Poznaniu
Instytut Podstaw Informatyki PAN
Uniwersytet Gdański

Mikulec Artur	<i>amikulec@uni.lodz.pl</i>
Młodak Andrzej	<i>a.mlodak@stat.gov.pl</i>
Moczko Jerzy Andrzej	<i>jmoczko@ump.edu.pl</i>
Modranka Emilia	<i>emodranka@uni.lodz.pl</i>
Mojsiewicz Magdalena	<i>m.mojiewicz@stat.gov.pl</i>
Monse Patricia	<i>patriciamonse@yahoo.fr</i>
Morze Marek	<i>M.Morze@stat.gov.pl</i>
Motoryn Ruslan	<i>motoryn@i.ua</i>
Motoryna Tetiana	<i>motoryn@i.ua</i>
Motyl Krystyna	<i>k.motyl@stat.gov.pl</i>
Mroczkowski Marek	<i>m.mroczkowski@stat.gov.pl</i>
Nadege Lohatame Kabongo	<i>sala.mondo@yahoo.fr</i>
Najman Krzysztof	<i>krzysztof.najman@wzr.pl</i>
Natkowska Monika	<i>m.natkowska@stat.gov.pl</i>
Nowak Lucyna	<i>l.nowak@stat.gov.pl</i>
Nowakowski Rafał	<i>r.nowakowski@stat.gov.pl</i>
Ntumba Lukusa Joseline	<i>papyjos@yahoo.fr</i>
Obrębalski Marek	<i>marek.obrebalski@ue.wroc.pl</i>
Okrasa Włodzimierz	<i>wlodek.okrasa@wanadoo.fr</i>
Okupniak Magdalena	<i>magdalena.okupniak@ue.poznan.pl</i>
Olbrót-Brzezińska Arleta	<i>olbrotbrzezinska@poczta.onet.pl</i>
Orwat Agnieszka	<i>agnieszka.orwat@ue.katowice.pl</i>
Ostasiewicz Walenty	<i>walenty.ostasiewicz@ue.wroc.pl</i>
Owsiński Jan	<i>owsinski@ibspan.waw.pl</i>
Panek Tomasz	<i>tompa10@interia.pl</i>
Paradysz Jan	<i>jan.paradysz@ue.poznan.pl</i>
Parlińska Maria	<i>maria-parlinska@sggw.pl</i>

Uniwersytet Łódzki
 Urząd Statystyczny w Poznaniu
 Uniwersytet Medyczny im. K. Marcinkowskiego w Poznaniu
 Uniwersytet Łódzki
 Urząd Statystyczny w Szczecinie
 BEL CAMPUS
 Urząd Statystyczny w Olsztynie
 Ukraiński Państwowy Uniwersytet
 Finansów i Handlu Międzynarodowego w Kijowie
 Kijowski Narodowy Uniwersytet Tarasa Shewchenko
 Urząd Statystyczny w Zielonej Górze
 Główny Urząd Statystyczny
 DEPOT BOISSON
 Uniwersytet Gdański
 Urząd Statystyczny w Poznaniu
 Główny Urząd Statystyczny
 Urząd Statystyczny we Wrocławiu
 BEL CAMPUS
 Uniwersytet Ekonomiczny we Wrocławiu
 Główny Urząd Statystyczny, UKSW w Warszawie
 Uniwersytet Ekonomiczny w Poznaniu
 Urząd Statystyczny w Poznaniu
 Uniwersytet Ekonomiczny w Katowicach
 Uniwersytet Ekonomiczny we Wrocławiu
 Instytut Badań Systemowych PAN
 Szkoła Główna Handlowa w Warszawie
 Uniwersytet Ekonomiczny w Poznaniu
 Szkoła Główna Gospodarstwa Wiejskiego w Warszawie

Parys Dariusz	<i>parysd@wp.pl</i>	Uniwersytet Łódzki
Paszko Elżbieta	<i>paszko.ela@gmail.com</i>	Uniwersytet Łódzki
Pawełek Barbara	<i>barbara.pawelek@uek.krakow.pl</i>	Uniwersytet Ekonomiczny w Krakowie
Pelka Marcin	<i>marcin.pelka@ue.wroc.pl</i>	Uniwersytet Ekonomiczny we Wrocławiu
Perek-Białas Jolanta	<i>jolanta.perek – bialas@uj.edu.pl</i>	Szkoła Główna Handlowa w Warszawie
Pieniążek Marek	<i>m.pieniazek@stat.gov.pl</i>	Główny Urząd Statystyczny
Pietrzykowski Robert	<i>robert_pietrzykowski@sggw.pl</i>	Szkoła Główna Gospodarstwa Wiejskiego w Warszawie
Plenikoska Teresa	<i>plenikowska@wzr.pl</i>	Uniwersytet Gdański
Pobłocka Agnieszka	<i>ekoap@univ.gda.pl</i>	Uniwersytet Gdański
Pociecha Józef	<i>pociecha@uek.krakow.pl</i>	Uniwersytet Ekonomiczny w Krakowie
Podogrodzka Małgorzata	<i>mpodog@sgh.waw.pl</i>	Szkoła Główna Handlowa w Warszawie
Podolec Barbara	<i>barbara.podolec@uek.krakow.pl</i>	Uniwersytet Ekonomiczny w Krakowie
Polińska Anna	<i>a.polinska@stat.gov.pl</i>	Urząd Statystyczny w Poznaniu
Potkański Tomasz	<i>tpotkanski@zmp.poznan.pl</i>	Związek Miast Polskich
Potrykowska Alina	<i>a.potrykowska@stat.gov.pl</i>	Główny Urząd Statystyczny, Rządowa Rada Ludnościowa
Ptak-Chmielewska Aneta	<i>aptak@sgh.waw.pl</i>	Szkoła Główna Handlowa w Warszawie
Pytalska Aleksandra	<i>a.pytalska@stat.gov.pl</i>	Główny Urząd Statystyczny
Roeske-Słomka Iwona	<i>ksid@ue.poznan.pl</i>	Uniwersytet Ekonomiczny w Poznaniu
Rogalińska Dominika	<i>d.rogalinska@stat.gov.pl</i>	Główny Urząd Statystyczny
Romaniuk Joanna	<i>asia.romaniuk@gmail.com</i>	Szkoła Główna Handlowa w Warszawie
Rossa Agnieszka	<i>agrossa@uni.lodz.pl</i>	Uniwersytet Łódzki
Roszak Magdalena	<i>mmer@ump.edu.pl</i>	Uniwersytet Medyczny im. K. Marcinkowskiego w Poznaniu
Roszka Wojciech	<i>wojciech.roszka@ue.poznan.pl</i>	Uniwersytet Ekonomiczny w Poznaniu
Rozkrut Dominik	<i>D.Rozkrut@stat.gov.pl</i>	Urząd Statystyczny w Szczecinie
Rozmus Dorota	<i>drozmus@ue.katowice.pl</i>	Uniwersytet Ekonomiczny w Katowicach
Rozpędowska-Matraszek Danuta	<i>daroma@uni.lodz.pl</i>	Państwowa Wyższa Szkoła Zawodowa w Skierniewicach
Ryś-Jurek Roma	<i>rys – jurek@up.poznan.pl</i>	Uniwersytet Przyrodniczy w Poznaniu
Siatkowski Idzi	<i>idzi@up.poznan.pl</i>	Uniwersytet Przyrodniczy w Poznaniu

Simisi Alex	<i>egec.rdc@gmail.com</i>	Enterprise de Genie Civil et Expertise
Sirko Stanisław	<i>s.sirko@aon.edu.pl</i>	Akademia Obrony Narodowej
Siwiński Włodzimierz	<i>siwinski@alk.edu.pl</i>	Uniwersytet Warszawski, Akademia Leona Koźmińskiego
Słaby Teresa	<i>tsl@pnet.pl</i>	Szkoła Główna Handlowa w Warszawie
Sobczak Elżbieta	<i>elzbieta.sobczak@ue.wroc.pl</i>	Uniwersytet Ekonomiczny we Wrocławiu
Sojkin Bogdan	<i>bogdan.sojkin@ue.poznan.pl</i>	Uniwersytet Ekonomiczny w Poznaniu
Sokołowski Andrzej	<i>sokolows@uek.krakow.pl</i>	Uniwersytet Ekonomiczny w Krakowie
Stachowiak Dorota	<i>d.stachowiak@stat.gov.pl</i>	Urząd Statystyczny w Poznaniu
Stanimir Agnieszka	<i>agnieszka.stanimir@ue.wroc.pl</i>	Uniwersytet Ekonomiczny we Wrocławiu
Stańczyk Elżbieta	<i>E.Stanczyk@stat.gov.pl</i>	Urząd Statystyczny we Wrocławiu
Stepień Lechosław	<i>czeslaw.stepien@gmail.com</i>	Uniwersytet Łódzki
Stepień Radosław	<i>R.Stepien@stat.gov.pl</i>	Główny Urząd Statystyczny
Stonawski Marcin	<i>stonawsm@uek.krakow.pl</i>	Uniwersytet Ekonomiczny w Krakowie, IIASA
Strahl Danuta	<i>danuta.strahl@ue.wroc.pl</i>	Uniwersytet Ekonomiczny we Wrocławiu
Strzelecki Zbigniew	<i>zstrzelecki@mbpr.pl</i>	Główny Urząd Statystyczny, Rządowa Rada Ludnościowa
Styrc Marta	<i>mstyrc@o2.pl</i>	Szkoła Główna Handlowa w Warszawie
Suchecką Jadwiga	<i>suchecką@uni.lodz.pl</i>	Uniwersytet Łódzki
Sucheckie Bogdan	<i>sucheckie@uni.lodz.pl</i>	Uniwersytet Łódzki
Szabelska Alicja	<i>aszab@up.poznan.pl</i>	Uniwersytet Przyrodniczy w Poznaniu
Szałtys Dorota	<i>S.Szaltys@stat.gov.pl</i>	Główny Urząd Statystyczny
Szkie-Czech Ewa	<i>sekretariatU Sopl@stat.gov.pl</i>	Urząd Statystyczny w Opolu
Szkutnik Zbigniew	<i>szkutnik@agh.edu.pl</i>	Akademia Górniczo-Hutnicza w Krakowie
Szreder Mirosław	<i>mszreder@wzr.ug.edu.pl</i>	Uniwersytet Gdański
Sztanderska Urszula	<i>sztanderska@wne.uw.edu.pl</i>	Uniwersytet Warszawski
Szubert Tomasz	<i>tomasz.szubert@ue.poznan.pl</i>	Uniwersytet Ekonomiczny w Poznaniu
Szukielój-Bieńkuńska Anna	<i>A.Bienkunska@stat.gov.pl</i>	Główny Urząd Statystyczny
Szulczyńska Elżbieta	<i>Elzbieta.Szulczynska@mg.gov.pl</i>	Ministerstwo Gospodarki
Szwarc Krzysztof	<i>k.szwarc@ue.poznan.pl</i>	Uniwersytet Ekonomiczny w Poznaniu

Szymkowiak Marcin	<i>m.szymkowiak@ue.poznan.pl</i>
Śliwicki Dominik	<i>D.Sliwicki@stat.gov.pl</i>
Śmilowska Teresa	<i>t.smilowska@stat.gov.pl</i>
Śmilowski Eugeniusz	<i>eugensmilowski@gmail.com</i>
Środa-Murawska Stefania	<i>stef fi@umk.pl</i>
Świeczewska Iwona	<i>iswiecz@uni.lodz.pl</i>
Tkaczyk Wanda	<i>W.Tkaczyk@stat.gov.pl</i>
Tomaszewicz Łucja	<i>ltomasz@uni.lodz.pl</i>
Traat Imbi	<i>imbi.traat@ut.ee</i>
Trzpiot Grażyna	<i>trzpiot@ae.katowice.pl</i>
Ulman Paweł	<i>ulmanp@uek.krakow.pl</i>
Walczak Tadeusz	<i>T.Walczak@stat.gov.pl</i>
Walesiak Marek	<i>marek.walesiak@ae.jgora.pl</i>
Walska Ewa	<i>e.walska@stat.gov.pl</i>
Wałęga Agnieszka	<i>agnieszka.walega@uek.krakow.pl</i>
Wardzińska-Sharif Amelia	<i>A.Wardzinska@stat.gov.pl</i>
Wasserstein Ronald	<i>ron@amstat.org</i>
Waszak Łukasz	<i>lwaszak@amu.edu.pl</i>
Wawrowski Łukasz	<i>lukasz8989@gmail.com</i>
Wawrzyniak Katarzyna	<i>katarzyna.wawrzyniak@zut.edu.pl</i>
Wesołowska-Janczarek Mirosława	<i>mirosława.wesolowska@up.lublin.pl</i>
Wiatrak Andrzej Piotr	<i>apw@mail.wz.uw.edu.pl</i>
Wierzbicka Anna	<i>demo@uni.lodz.pl</i>
Wilk Justyna	<i>justyna.wilk@ue.wroc.pl</i>
Wilkin Jerzy	<i>wilkin@wne.uw.edu.pl</i>
Witkowski Janusz	<i>j.witkowski@stat.gov.pl</i>
Wojciechowska Urszula	<i>wojciechowskau@coi.waw.pl</i>
Wojtkowiak-Jakacka Małgorzata	<i>SekretariatUSWRO@stat.gov.pl</i>

Uniwersytet Ekonomiczny w Poznaniu
Urząd Statystyczny w Bydgoszczy
Urząd Statystyczny w Łodzi
Rada Statystyki
Uniwersytet Mikołaja Kopernika w Toruniu
Uniwersytet Łódzki
Główny Urząd Statystyczny
Uniwersytet Łódzki
University of Tartu
Uniwersytet Ekonomiczny w Katowicach
Uniwersytet Ekonomiczny w Krakowie
Główny Urząd Statystyczny
Uniwersytet Ekonomiczny we Wrocławiu
Główny Urząd Statystyczny
Uniwersytet Ekonomiczny w Krakowie
Główny Urząd Statystyczny
American Statistical Association
Uniwersytet im. Adama Mickiewicza w Poznaniu
Urząd Statystyczny w Poznaniu
Zachodniopomorski Uniwersytet Technologiczny w Szczecinie
Uniwersytet Przyrodniczy w Lublinie
Uniwersytet Warszawski
Instytut Statystyki i Demografii
Uniwersytet Ekonomiczny we Wrocławiu
Uniwersytet Warszawski
Główny Urząd Statystyczny
Centrum Onkologii Instytut
Urząd Statystyczny we Wrocławiu

Wojtyś Małgorzata
 Woźniak Halina
 Woźniak Witold
 Wróblewska Wiktoria
 Wyłupek Grzegorz
 Wysocka Anna
 Wysocki Feliks
 Wyszowska Dorota
 Wywiał Janusz
 Załuska Urszula
 Zawadzki Jan
 Zdaniewicz Witold
 Zgierska Agnieszka
 Zhang Li-Chun
 Zieliński Wojciech
 Zmysłony Roman
 Zwara Wioletta
 Zyprych-Walczak Joanna
 Żądło Tomasz

M.Wojtys@mini.pw.edu.pl
SekretariatUSWRO@stat.gov.pl
W.Wozniak@stat.gov.pl
wwrobl@sgh.waw.pl
wylupek@impan.pan.wroc.pl
A.Wysocka@stat.gov.pl
wysocki@up.poznan.pl
d.wyszowska@stat.gov.pl
janusz.wywial@ue.katowice.pl
urszula.zaluska@ue.wroc.pl
jan.zawadzki@zut.edu.pl
zdaniewicz@ecclesia.org.pl
a.zgierska@stat.gov.pl
lcz@ssb.no
wojtek.zielinski@statystyka.info
r.zmyslony@wmie.uz.zgora.pl
sekretariatUSBDG@stat.gov.pl
zjoanna@up.poznan.pl
tomasz.zadlo@ue.katowice.pl

Politechnika Warszawska
 Urząd Statystyczny we Wrocławiu
 Główny Urząd Statystyczny
 Szkoła Główna Handlowa w Warszawie
 Instytut Matematyczny PAN
 Główny Urząd Statystyczny
 Uniwersytet Przyrodniczy w Poznaniu
 Uniwersytet w Białymstoku, Urząd Statystyczny w Białymstoku
 Uniwersytet Ekonomiczny w Katowicach
 Uniwersytet Ekonomiczny we Wrocławiu
 Zachodniopomorski Uniwersytet Technologiczny w Szczecinie
 Instytut Statystyki Kościoła Katolickiego
 Główny Urząd Statystyczny
 Statistics Norway
 Szkoła Główna Gospodarstwa Wiejskiego w Warszawie
 Uniwersytet Zielonogórski
 Urząd Statystyczny w Bydgoszczy
 Uniwersytet Przyrodniczy w Poznaniu
 Uniwersytet Ekonomiczny w Katowicach



z okazji jubileuszu 100-lecia
 Polskiego Towarzystwa
 Statystycznego

NBP

Narodowy Bank Polski

Projekt edukacyjny pt. „Wydawnictwa okolicznościowe oraz wystawy związane z jubileuszem 100-lecia Polskiego Towarzystwa Statystycznego przygotowane na Kongres Statystyki Polskiej” został dofinansowany ze środków **Narodowego Banku Polskiego**

