

Poland, Bełchatów, 16-18.09.2026

MSA-SKAD 2026

Scientific Conference

**The International Scientific Conference
on Multivariate Statistical Analysis, the 44th ed.
and
The Scientific Conference
of the Classification and Data Analysis Section
of the Polish Statistical Association, the 35th ed.**

**The Conference is co-financed from the state budget
by the Ministry of Science and Higher Education
under the "Wektory Nauki" /**

**Projekt finansowany ze środków budżetu państwa,
przyznanych przez Ministra Nauki i Szkolnictwa Wyższego
w ramach Programu Wektory Nauki.**



**Ministerstwo Nauki
i Szkolnictwa Wyższego**



Co-Organizers / Współorganizatorzy

Department of Statistical Methods, University of Łódź /
Katedra Metod Statystycznych, Uniwersytet Łódzki



Classification and Data Analysis Section (SKAD)
of the Polish Statistical Association /
Sekcja Klasyfikacji i Analizy Danych (SKAD)
Polskiego Towarzystwa Statystycznego



Ministry of Science and Higher Education /
Ministerstwo Nauki i Szkolnictwa Wyższego



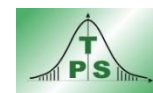
Polish Academy of Sciences /
Polska Akademia Nauk



Statistical Office in Łódź /
Urząd Statystyczny w Łodzi



Polish Statistical Association, Branch in Łódź /
Polskie Towarzystwo Statystyczne, Oddział w Łodzi



Honorary Patronage / Patronat honorowy

Honorary Patronage of the President of Statistics Poland /
Patronat Honorowy Prezesa Głównego Urzędu Statystycznego



Honorary Patronage of the Voivode of Łódź, Dorota Ryl /
Patronat Honorowy Wojewody Łódzkiego Doroty Ryl



Rector of University of Łódź /
Rektor Uniwersytetu Łódzkiego



Partners / Partnerzy

StatSoft Polska Sp. z o. o.



LabMasters sp. z o. o.



MSA-SKAD 2026

Scientific Committee / Komitet Naukowy

Chair / Przewodniczący

śp. prof. dr hab. Czesław Domański, University of Łódź

Vice-chair / Wiceprzewodniczący

prof. dr hab. Krzysztof Jajuga,
Wroclaw University of Economics and Business

Members / Członkowie

prof. Francesca Greselin, University of Milano Bicocca

prof. Marie Huskova, Charles University in Prague

prof. dr hab. Mirosław Krzyśko, Adam Mickiewicz University in Poznań

prof. dr hab. Paweł Lula, Krakow University of Economics

prof. dr hab. Włodzimierz Okrasa, The Cardinal Stefan Wyszyński University in Warsaw

prof. dr hab. Józef Pocięcha, Krakow University of Economics

prof. dr hab. Mirosław Szreder, University of Gdańsk

prof. dr hab. Marek Walesiak, Wroclaw University of Economics and Business

prof. dr hab. Janusz Wywiat, University of Economics in Katowice

dr hab. Grażyna Dehnel, prof. UEP, The Poznań University of Economics and Business

dr hab. Andrzej Dudek, prof. UWE, Wroclaw University of Economics and Business

dr hab. Jerzy Korzeniewski, prof. UŁ, University of Łódź

MSA-SKAD 2026

Organizing Committee / Komitet Organizacyjny

Chair / Przewodnicząca

**prof. dr hab. Alina Jędrzejczak,
Uniwersytet Łódzki**

Vice-chair / Wiceprzewodniczący

**prof. dr hab. Jacek Białek,
Uniwersytet Łódzki**

Secretaries / Sekretarze

dr Elżbieta Zalewska, Uniwersytet Łódzki

dr Artur Mikulec, Uniwersytet Łódzki

mgr Aleksandra Kupis-Fijałkowska, Uniwersytet Łódzki

Invited Lectures / Wykłady zaproszone:

Sesja plenarna I

Christian Hennig, Università di Bologna

What are outliers?

Outlier identification is an important task in data analysis, and there are many methods for doing so. The rationale behind these methods is however often not clear. What qualifies an observation to be an outlier, and how can we know whether a method is good at correctly identifying the outliers in the data? This presentation is about how to define outliers statistically in order to clarify the identification problem, and to provide a basis for constructing and assessing outlier identifiers. It turns out that such a definition cannot be unique and requires a number of decisions by the users that may depend on background knowledge and on what should be done with detected outliers. Without such input, the data on their own cannot decide what an outlier is. I will start from the simplest situation of one-dimensional data discussing the boxplot and other approaches, and particularly the role of statistical models for the non-outliers. If time allows, I will then go on to more complex problems such as outliers in cluster analysis and in mixed type data.

Sesja plenarna II

Stanisław Maciej Kot, Politechnika Gdańska

Empirical aspects of value judgements of income distributions: a theoretical background, estimation and applications

Assessing inequality, poverty, and related issues in income distribution is not a purely statistical exercise, since it requires invoking norms governing social preferences. For instance, social decision-makers' attitudes toward inequality are decisive for the shape of the disposable income distribution. Applied welfare economics has developed various quantitative measures of aversion to inequality that characterise the attitudes in question. The problem of great theoretical and practical importance is how to recover the levels of these measures when unobservable social preferences reveal themselves in the form of the actual distribution of disposable incomes? This paper offers the answer to this question. The paper proposes methods for eliciting aversion to income inequality, aversion to rank inequality, and for estimating the equivalence scale from empirical distributions of disposable incomes. The methods are applied to assess equivalence scales, inequality, social welfare and poverty in 27 countries of the European Union in 2021.

Sesja plenarna III

Marek Cierpiat-Wolan, Prezes Głównego Urzędu Statystycznego, Uniwersytet Rzeszowski

(Big) data integration methods: the potential of smart city data

Sesja plenarna IV

Eugeniusz Gatnar, Uniwersytet Ekonomiczny w Katowicach

Statystyka banków centralnych. Moje doświadczenia w NBP w latach 2007-2022.

Narodowy Bank Polski stanowi obok Głównego Urzędu Statystycznego drugi filar statystyki publicznej w Polsce. Departament Statystyki NBP gromadzi dane z sektora bankowego oraz prowadzi rachunki finansowe państwa. W prowadzeniu wielu badań statystycznych współpracuje z GUS.

W wystąpieniu zostaną przedstawione przykłady wykorzystania danych z zakresu statystyki pieniężnej, statystyki stóp procentowych oraz statystyki bilansu płatniczego w zarządzaniu sytuacjami kryzysowymi w latach 2007-2022, np. atakiem na złotówkę w 2009 roku, kredytami hipotecznymi w CHF, problemami z saldem błędów i opuszczeń w bilansie płatniczym, czy monetarnymi skutkami kryzysu pandemicznego.

W latach 2007-2022 pracowałem w NBP, a w szczególności w latach 2010-2016 nadzorowałem Departament Statystyki jako członek Zarządu NBP. Byłem także przedstawicielem NBP w Irving Fisher Committee for Central Bank Statistics oraz w Radzie Statystyki. W latach 2016-2022 byłem z kolei członkiem Rady Polityki Pieniężnej IV kadencji.

Słowa kluczowe: Bank centralny, statystyka monetarna, statystyka stóp procentowych, bilans płatniczy, inflacja

Central Bank Statistics. My experience at the National Bank of Poland (2007-2022).

The presentation will set out examples of how data from monetary statistics, interest rate statistics and balance of payments statistics were used in crisis management between 2007 and 2022, such as the attack on the zloty in 2009, mortgages denominated in CHF, problems with the errors and omissions balance in the balance of payments, and the monetary effects of the pandemic crisis.

Keywords: Central bank, monetary statistics, interest rate statistics, balance of payments, inflation

Sesja plenarna V

Michał Lewandowski, Szkoła Główna Handlowa w Warszawie

Jak zmierzyć subiektywną niepewność w złotówkach za pomocą różnicy między ceną kupna i sprzedaży?

Większość decyzji ekonomicznych dotyczy przyszłości, a więc podejmowana jest w warunkach niepewności. W klasycznej teorii ryzyka zakłada się, że znamy prawdopodobieństwa możliwych zdarzeń. Pozwala to mierzyć niechęć do ryzyka za pomocą premii za ryzyko, kwoty, jaką człowiek jest gotów zapłacić, aby zastąpić niepewny wynik jego wartością oczekiwaną. W praktyce jednak często nie znamy prawdopodobieństw lub nie potrafimy ich wiarygodnie określić. W takich sytuacjach mamy do czynienia z niepewnością, a nie jedynie ryzykiem.

W wykładzie przedstawię uogólnienie klasycznego pomiaru premii za ryzyko na przypadek ogólnej niepewności. Zamiast wartości oczekiwanej wykorzystamy koncepcję doskonałego zabezpieczenia (hedge), które całkowicie eliminuje niepewność. Następnie pokażę, jak na podstawie różnicy między ceną kupna a ceną (krótkiej-) sprzedaży tego samego niepewnego losu można zdefiniować i zmierzyć

premię za niepewność. Intuicyjnie jest to minimalna kwota potrzebna, aby obie strony tego samego zakładu, czyli kupno i wystawienie losu, stały się akceptowalne dla tej samej osoby.

Pokażę także, że proponowana metoda pozwala wyodrębnić dwie główne przyczyny różnicy między WTA i WTP: awersję do strat oraz niezdecydowanie. Na zakończenie zilustruję działanie tej metody za pomocą danych uzyskanych w eksperymencie.

How Can Subjective Uncertainty Be Measured in Dollars Using the WTA-WTP Gap?

Most economic decisions concern the future and are therefore made under uncertainty. In the classical theory of risk, it is assumed that the probabilities of possible outcomes are known. This allows risk aversion to be measured through the risk premium, the amount an individual is willing to pay to replace an uncertain prospect with its expected value. In practice, however, probabilities are often unknown or cannot be reliably assessed. In such situations, we face uncertainty rather than merely risk.

In this talk, I will present a generalization of the classical concept of the risk premium to the broader setting of uncertainty. Instead of the expected value, we use the notion of a perfect hedge, which completely eliminates uncertainty. I will then show how the difference between the buying price and the (short-)selling price of the same uncertain prospect can be used to define and measure an uncertainty premium. Intuitively, this premium is the minimum amount of money required to make both sides of the same bet, buying and issuing a prospect, acceptable to the same individual.

I will also show that the proposed method isolates the two main sources of the WTA–WTP disparity: loss aversion and preference imprecision. Finally, I will illustrate the proposed method using data from an experiment.

R Language Training/ Szkolenie R:

Piotr Ćwiakowski, LabMasters

Machine Learning w R

Modele predykcyjne oparte o techniki uczenia maszynowego są wciąż najpopularniejszym narzędziem zaawansowanej analityki biznesowej i coraz częściej również naukowej. Skuteczne i prawidłowe modelowanie z wykorzystaniem modeli uczenia maszynowego wymaga zastosowania nieco innego podejścia od klasycznej analizy statystyczno-ekonometrycznej. Na warsztacie omówiona zostanie kompletna procedura tworzenia modelu od podstaw do interpretacji wyników, na przykładzie lasu losowego. Program warsztatu:

- konstrukcja modelu lasów losowych,
- znaczenie walidacji krzyżowej i zbioru testowego w procesie modelowania,
- tuning hiperparametrów,
- wykorzystanie technik XAI w interpretacji wyników.

Papers / Referaty:

Tomasz Bartłomowicz, Uniwersytet Ekonomiczny we Wrocławiu; Andrzej Bąk, Uniwersytet Ekonomiczny we Wrocławiu

Pakiet conjoint w Pythonie – implementacja i zastosowanie w badaniu preferencji konsumentów

Metoda analizy conjoint znajduje szerokie zastosowanie w badaniach preferencji konsumentów, jednak jej implementacje są najczęściej dostępne w postaci narzędzi komercyjnych lub zamkniętych platform internetowych. W odpowiedzi na rosnące zapotrzebowanie na rozwiązania niekomercyjne oraz o otwartym kodzie źródłowym, w artykule przedstawiono pakiet conjoint opracowany w języku Python, stanowiący rozwinięcie wcześniej stworzonego przez autorów pakietu w środowisku R. Proponowane rozwiązanie rozszerza możliwości analizy conjoint na ekosystem Python.

Pakiet umożliwia kompleksową realizację analizy conjoint, obejmującą estymację użyteczności częściowych na poziomie indywidualnym i zagregowanym, ocenę ważności atrybutów, segmentację respondentów oraz symulację wyborów konsumenckich i estymację udziałów rynkowych. Ponadto, zawiera narzędzia wspomagające przygotowanie danych oraz projektowanie eksperymentów, w tym generowanie efektywnych planów czynnikowych oraz transformację danych preferencyjnych.

Zastosowanie pakietu zaprezentowano na przykładzie badania ankietowego typu rating-based, opartego na deklarowanych preferencjach konsumentów, którego celem była identyfikacja preferencji konsumentów względem ofert turystycznych. Respondenci dokonywali ocen przedstawionych profili produktów turystycznych, zróżnicowanych pod względem kluczowych atrybutów, takich jak kierunek podróży, standard zakwaterowania, forma wyżywienia, środek transportu, długość pobytu oraz cena. Uzyskane wyniki potwierdzają użyteczność zaproponowanego rozwiązania jako alternatywy o otwartym kodzie źródłowym dla komercyjnych narzędzi analizy conjoint.

Słowa kluczowe: conjoint analysis, preferencje wyrażone konsumentów, Python, R, oprogramowanie open source, rynek turystyczny

The conjoint Package in Python – Implementation and Application in Consumer Preference Research

Conjoint analysis is widely used in consumer preference research; however, its implementations are most often available as commercial software or closed online platforms. In response to the growing demand for non-commercial and open-source solutions, this paper presents the conjoint package developed in Python as an extension of the authors' previously developed package in the R environment. The proposed solution extends the capabilities of conjoint analysis to the Python ecosystem.

The package enables comprehensive conjoint analysis, including the estimation of part-worth utilities at both individual and aggregate levels, assessment of attribute importance, respondent segmentation, and simulation of consumer choices along with market share estimation. It also provides tools supporting data preparation and experimental design, including the generation of efficient factorial designs and transformation of preference data.

The functionality of the package is demonstrated through a rating-based survey study based on stated preferences, aimed at identifying consumer preferences with respect to tourism offers. Respondents

evaluated hypothetical product profiles characterized by key attributes such as travel destination, accommodation standard, meal plan, mode of transport, length of stay, and price. The results confirm the usefulness of the proposed solution as an open-source alternative to commercial tools for conjoint analysis.

Keywords: conjoint analysis, stated consumer preferences, Python, R, open-source software, tourism market

Jacek Białek, Uniwersytet Łódzki, Główny Urząd Statystyczny, Departament Cen i Usług; Dagmara Oprych-Franków, Główny Urząd Statystyczny, Departament Cen i Usług

Ocena wpływu poszczególnych produktów skanowanych na indeks cen typu GEKS

W ramach tzw. klasycznego (tradycyjnego) pomiaru inflacji, głównym źródłem danych o cenach dóbr i usług są ankieteryzy w terenie, natomiast informacje o poziomie i strukturze konsumpcji (wydatków konsumentów) pochodzą z Badania Budżetów Gospodarstw Domowych (BBGD). Metodologia szacowania wskaźnika cen towarów i usług konsumpcyjnych (CPI) czy zharmonizowanego wskaźnika cen konsumpcyjnych (HICP) jest w tym zakresie właściwie dopracowana, przy czym ma ona pewne ograniczenia wynikające z dostępności danych, częstotliwości ich kolekcji oraz aktualizacji koszyka inflacyjnego. Statystycy mają świadomość występowania pewnych obciążeń pomiaru inflacji, jednak ich całkowita eliminacja wydaje się niemożliwa z uwagi na dostępność „tradycyjnych” danych o cenach i strukturze konsumpcji.

Od blisko 30 lat, a ostatnio bardzo intensywnie, rośnie znaczenie alternatywnych źródeł danych o cenach produktów. Ponad 1/3 krajów UE, a także wiele państw spoza Europy (np. USA, Australia czy Japonia), do szacowania inflacji wykorzystuje dane skanowane, które pochodzą z elektronicznych terminali sieci handlowych. Tego rodzaju zbiory danych charakteryzują się nie tylko ogromnym wolumenem, ale też dostarczają informacji o cenach i ilościach kupowanych produktów już na najniższym poziomie agregacji danych (poziom kodu kreskowego). Z uwagi na dużą rotację produktów oraz sezonowość sprzedaży, „klasyczne” indeksy cenowe (indeksy bilateralne) nie nadają się do szacowania inflacji produktów skanowanych. Najczęściej rekomendowanymi indeksami cenowymi są w tym przypadku indeksy multilateralne, wśród których najpopularniejszymi metodami są formuły GEKS, CCDI, TPD czy indeks Geary-Khamisa. Indeksy multilateralne są jednak bardziej skomplikowane (działają na oknie czasowym), bardziej czasochłonne i trudniejsze w dekompozycji.

Niniejsza praca skupia się na multiplikatywnej dekompozycji indeksów multilateralnych typu GEKS. Taka dekompozycja pozwala ocenić wpływ poszczególnych produktów skanowanych na kształtowanie się indeksu cenowego. W pracy zaproponowano szereg multiplikatywnych dekompozycji wraz z ich znormalizowanymi wersjami, pozwalającymi na porównania wpływu produktów pomiędzy różnymi formułami indeksów typu GEKS. Wyniki teoretyczne zobrazowano z wykorzystaniem rzeczywistych zbiorów danych skanowanych wskazując tym samym na odmiennność indeksu GEKS-L w stosunku do pozostałych przedstawicieli klasy indeksów GEKS.

Słowa kluczowe: dane skanowane, wskaźnik cen towarów i usług konsumpcyjnych (CPI), indeksy multilateralne, dekompozycja indeksów cen

Assessing the Impact of Individual Scanner Data Products on the GEKS Price Index

Within the framework of the traditional approach to inflation measurement, the primary source of price data for goods and services is field price collection conducted by price collectors, while information on the level and structure of household consumption (consumer expenditures) is obtained from the Household Budget Survey (HBS). The methodology for compiling the Consumer Price Index (CPI) and the Harmonised Index of Consumer Prices (HICP) has been well established in this respect. Nevertheless, it is subject to certain limitations arising from data availability, the frequency of data collection, and the updating of the inflation basket. Statisticians are aware of the existence of various sources of measurement bias in inflation estimation; however, their complete elimination appears impossible given the limitations of traditional data on prices and consumption patterns.

For nearly three decades, and particularly in recent years, the importance of alternative sources of product price data has increased substantially. More than one-third of the European Union Member States, as well as many countries outside Europe (e.g., the United States, Australia, and Japan), use scanner data obtained from electronic point-of-sale (POS) systems of retail chains for inflation measurement. Such datasets are characterised not only by their massive volume but also by the availability of information on both prices and quantities purchased at the lowest level of data aggregation (i.e., barcode level). Owing to the high turnover of products and the seasonality of sales, conventional bilateral price indices are not suitable for measuring inflation based on scanner data. Instead, multilateral price indices are generally recommended, with the most widely used methods including the GEKS, CCDI, TPD, and Geary–Khamis indices. However, multilateral indices are more complex, operate over rolling time windows, require greater computational effort, and are more difficult to decompose.

This paper focuses on the multiplicative decomposition of GEKS-type multilateral indices. Such a decomposition makes it possible to assess the contribution of individual scanner data products to the overall price index. Several multiplicative decomposition approaches are proposed, together with their normalised versions, enabling comparisons of product contributions across different GEKS-type index formulae. The theoretical results are illustrated using real-world scanner datasets, highlighting the distinctive properties of the GEKS-L index compared with other representatives of the GEKS family of indices.

Keywords: scanner data, Consumer Price Index (CPI), multilateral price indices, price index decomposition

Beata Bieszk-Stolorz, Uniwersytet Szczeciński

Entropia populacji Keyfitza w ocenie czasu trwania w bezrobociu

Entropia populacji Keyfitza jest to kluczowa miara w demografii i analizie przeżycia, opisująca wrażliwość oczekiwanej długości życia na proporcjonalne zmiany śmiertelności. W przedstawionym badaniu została ona wykorzystana do oceny czasu trwania w bezrobociu osób zarejestrowanych w Powiatowym Urzędzie Pracy w Szczecinie w latach 2007-2024. W tym kontekście czasu trwania interpretacja entropii Keyfitza różni się od tej dla długości życia ludzkiego. Entropia informuje o dyspersji długości epizodów bezrobocia oraz o wrażliwości średniego czasu bezrobocia na zmiany hazardu wyjścia z bezrobocia. W przeprowadzonym badaniu analizowano modele parametryczne.

W większości przypadków lepiej dopasowane były modele zakładające lognormalny rozkład czasu trwania. Analizowano dwa zdarzenia kończące obserwację: podjęcie pracy i rezygnację z pośrednictwa urzędu. Oba zdarzenia miały wysokie wartości entropii populacji (powyżej 1), przy czym entropia wyznaczona dla podjęcia pracy była wyższa niż dla rezygnacji (za wyjątkiem 2020 roku). Świadczy to o mało przewidywalnym czasie wyjścia z kohorty, szczególnie dla podjęcia pracy. Część osób wychodzi z bezrobocia szybko, ale duża część pozostaje w nim bardzo długo. Niewielkie zmiany w dostępności pracy lub aktywizacji mogą nieść za sobą duży efekt. Standardowe, jednorodne programy aktywizacyjne są nieskuteczne, gdyż potrzebne są zindywidualizowane interwencje na rynku pracy.

Słowa kluczowe: analiza przeżycia, entropia populacji Keyfitza, bezrobocie rejestrowane

Keyfitz's Population Entropy in the Assessment of Unemployment Duration

Keyfitz's population entropy is a fundamental measure in demography and survival analysis, describing the sensitivity of life expectancy to proportional changes in mortality. In the present study, it was applied to evaluate the duration of unemployment among individuals registered at the Poviát Labour Office in Szczecin in the years 2007–2024. In the context of unemployment duration, the interpretation of Keyfitz's entropy differs from that used for human longevity. Here, entropy reflects the dispersion of unemployment spells and the sensitivity of the average unemployment duration to changes in the hazard of exiting unemployment. The parametric models were used in the analysis. In most cases, models assuming a lognormal distribution of unemployment duration provided a better fit. Two events terminating observation were considered: entering employment and resignation from the co-operation of the labour office. Both events exhibited high values of population entropy (above 1), with the entropy associated with entering employment being higher than that for resignation (except in 2020). This indicates a low predictability of exit times from the cohort, particularly for transitions into employment. Some individuals leave unemployment quickly, while a substantial proportion remain unemployed for a very long time. Even small changes in job availability or activation programmes may therefore produce large effects. Standard, uniform activation programmes are ineffective, as the labour market requires individualised interventions.

Keywords: survival analysis, Keyfitz's population entropy, registered unemployment

Marcin Błażejowski, WSB Merito University in Toruń

gretl4py: a package for calling gretl from Python

gretl4py jest pakietem zapewniającym interfejs pomiędzy libgretl – wieloplatformową, wysokowydajną, niezawodną i spójną biblioteką ekonometryczną – a językiem Python3.xx. Idea stworzenia takiego połączenia pomiędzy Python-em, powszechnie wykorzystywanym językiem programowania do obliczeń naukowych i analizy danych, a libgretl od pewnego czasu pojawiała się w społeczności użytkowników gretl.

gretl4py udostępnia wiązania (*bindings*) nie tylko dla natywnych estymatorów i testów statystycznych oferowanych przez libgretl, lecz także dla pakietów gretl, obejmujących rozszerzenia zarówno oficjalne, jak i te tworzone przez użytkowników. Udostępniając wszystkie te możliwości w środowisku Python,

gretl4py łączy funkcjonalność silnika ekonometrycznego libgretl z elastycznością i bogatym ekosystemem języka Python, tworząc praktyczne środowisko dla szerokiego zakresu zastosowań ekonometrycznych.

gretl4py: a package for calling gretl from Python

gretl4py is a software package that provides an interface between libgretl —a multi-platform, high-performance, reliable, and consistent econometric library—and Python3.xx. The idea of developing such a bridge between Python, a widely used programming language for scientific computing and data analysis, and libgretl had circulated within the gretl community for some time.

gretl4py provides Python bindings not only for native libgretl estimators and statistical tests, but also for gretl function packages, including both official and user-contributed extensions. By exposing these capabilities within the Python environment, gretl4py combines the functionality of the libgretl econometric engine with the flexibility and ecosystem of Python, providing a practical framework for a broad range of econometric applications.

Dawid Bonar, Uniwersytet Ekonomiczny w Krakowie; Justyna Borowiec-Kapek, Szkoła Doktorska Uniwersytet Ekonomiczny w Krakowie; Monika Papież, Uniwersytet Ekonomiczny w Krakowie; Michał Rubaszek, Szkoła Główna Handlowa w Warszawie

Wpływ regulacji rynku energii elektrycznej na transmisję zmienności cen energii elektrycznej z Niemiec do Polski

W artykule analizujemy źródła zmienności cen energii elektrycznej w Polsce, ze szczególnym uwzględnieniem powiązań między rynkiem polskim i niemieckim oraz wpływu mechanizmów integracji europejskiego rynku energii elektrycznej, obejmujących wprowadzenie mechanizmu Net Transfer Capacity (NTC), Flow-Based Market Coupling (FBMC) oraz przejście na 15-minutowe produkty handlowe w ramach Single Day-Ahead Coupling (SDAC). Na podstawie godzinowych danych dotyczących cen energii elektrycznej z lat 2020–2025 konstruujemy dzienną miarę zmienności cen i analizujemy jej dynamikę z wykorzystaniem zmiennych opisujących sytuację na rynku niemieckim, transgraniczne przepływy energii elektrycznej, udział odnawialnych źródeł energii, krańcowe koszty wytwarzania energii oraz zmienne regulacyjne odzwierciedlające strukturę instytucjonalną europejskiego rynku energii elektrycznej. Wyniki wskazują na istotną transmisję zmienności cen energii z rynku niemieckiego na rynek polski, która nasila się po wdrożeniu mechanizmu NTC. Otrzymane wyniki sugerują, że integracja rynku może wzmacniać transmisję zmienności pomiędzy sąsiednimi rynkami energii elektrycznej. Pokazujemy również, że wyższy udział energii wiatrowej wiąże się z niższą zmiennością cen energii elektrycznej, natomiast wzrost udziału energii słonecznej jest związany z wyższą zmiennością cen. Pokazujemy również, że wyższy udział energii wiatrowej wiąże się z niższą zmiennością cen energii elektrycznej, natomiast wzrost udziału energii słonecznej jest związany z wyższą zmiennością cen. Wyniki wskazują także, że zmienność cen rośnie zarówno w okresach niedoboru mocy, jak i nadpodaży energii, co podkreśla znaczenie elastyczności systemu elektroenergetycznego. Ponadto czynniki wpływające na zmienność różnią się pomiędzy godzinami szczytowymi i pozaszczytowymi.

Słowa kluczowe: zmienność cen energii elektrycznej; integracja rynku energii; mechanizmy market coupling; odnawialne źródła energii; transgraniczny handel energią elektryczną

Nierówności probabilistyczne w szacowaniu kwantylowych miar ryzyka

Celem badania jest ilościowa ocena ryzyka modelu, oparta na dolnych i górnych ograniczeniach kwantylowych miar ryzyka. Ograniczenia te uzyskano wykorzystując nowe propozycje uogólnień znanych w literaturze nierówności probabilistycznych. Teoretyczne wyniki zilustrowano na przykładach rozkładów, które są powszechnie stosowane w analizie finansowych szeregów czasowych. Następnie zaproponowane nierówności zastosowano do oszacowania kwantylowych miar ryzyka w modelach klasy GARCH. Uzyskane rezultaty wykorzystano do ilościowej oceny ryzyka rozważanych modeli. W tym celu obliczono bezwzględną oraz względną miarę ryzyka modelu (absolute measure of model risk, relative measure of model risk), a także tzw. warunkowy wkład w redukcję ryzyka modelu (conditional model risk contribution).

Słowa kluczowe: nierówność probabilistyczna, kwantylowa miara ryzyka, GARCH, ryzyko modelu

Probabilistic inequalities for the assessment of tail risk measures

The aim of the research is a quantitative assessment of model risk, based on lower and upper bounds on tail risk measures. These bounds are obtained using new proposals for generalisations of probabilistic inequalities known in the literature. The theoretical results are illustrated with examples of distributions commonly used in financial time series analysis. The proposed inequalities are then applied to estimate tail risk measures in GARCH models. The results are used for a quantitative assessment of the risk of the models considered. For this purpose, the absolute and relative measures of model risk, as well as the conditional model risk contribution, are estimated.

Keywords: probabilistic inequality, tail risk measure, GARCH, model risk

Model oparty na entropii w prognozowaniu emisji CO₂ na poziomie przedsiębiorstw

Prognozowanie emisji gazów cieplarnianych dla poszczególnych firm jest dziś kluczowym wyzwaniem, szczególnie w kontekście unijnych wymogów raportowania i planowania dekarbonizacji. W naszym badaniu proponujemy nowatorskie podejście, które łączy koncepcję entropii z efektywnością użycia materiałów i energii, by lepiej szacować emisje CO₂ w sektorze prywatnym. Stosujemy metodę hybrydową, łącząc rzeczywiste dane o emisjach firm z bazy CDP z danymi sektorowymi o przepływach materiałów i zużyciu energii z baz EU KLEMS i Eurostatu. Nasze podejście uwzględnia złożoność procesów w firmach w trzech wymiarach entropii: energetycznej, materiałowej i organizacyjnej. Wyniki pokazują, że nasz model dobrze przewiduje emisje i ujawnia kluczowe bariery efektywności. To podejście pozwala ominąć problem braku szczegółowych danych z firm, oferując precyzyjne narzędzie dla menedżerów i polityków klimatycznych.

Słowa kluczowe: entropia informacyjna; prognozowanie emisji CO₂; produktywność materiałowa; dekarbonizacja przedsiębiorstw; modelowanie hybrydowe.

An Entropy-Based Model for Forecasting Firm-Level CO₂ Emissions

Accurate forecasting of greenhouse gas emissions at the firm level is a critical challenge under the EU's Corporate Sustainability Reporting Directive (CSRD) and for corporate decarbonization planning. This study introduces a novel methodological framework that couples information entropy with material and energy productivity (MP) to estimate and forecast corporate CO₂ emissions. Using a hybrid approach, the model integrates real, verified firm-level emissions data (sourced from the CDP database) with sector-level material flow and energy consumption indicators (derived from EU KLEMS and Eurostat). Within the proposed equation, $CO_2 = f(S_{\text{energy}}, S_{\text{material}}, S_{\text{org}}, MP)$, the structural complexity of corporate processes is captured through three dimensions of entropy: energy, material, and organizational. The model calibration results demonstrate strong predictive capabilities and identify key structural barriers to emission efficiency. This framework effectively bypasses the common limitation of missing firm-level microdata on resource inputs, providing a robust, scalable tool to support corporate decision-making and targeted climate policy.

Keywords: Information entropy; CO₂ forecasting; material productivity; corporate decarbonization; hybrid modeling

Katarzyna Cheba, Uniwersytet Zielonogórski, Park Technologii Kosmicznych w Zielonej Górze; Jarosław Januszewski, Park Technologii Kosmicznych w Zielonej Górze; Maciej Szymański, Park Technologii Kosmicznych w Zielonej Górze

Mapowanie usług ekosystemowych z wykorzystaniem technik satelitarnych

Artykuł podejmuje problematykę nowoczesnej oceny i monitorowania kapitału naturalnego przy użyciu zaawansowanych technik teledetekcyjnych oraz danych satelitarnych. Głównym celem artykułu jest zaprezentowanie metodologii operacyjnego mapowania tych usług z wykorzystaniem wieloźródłowych danych satelitarnych na przykładzie obszaru województwa lubuskiego. W artykule szczegółowo opisano zastosowanie nowoczesnych narzędzi obserwacji Ziemi, które pozwalają na obiektywną, ciągłą i powtarzalną analizę zmian środowiskowych. W pracy opisano praktyczne możliwości wykorzystania mapowania usług ekosystemowych w obszarach obejmujących: retencję wody, redukcję zanieczyszczeń powietrza oraz atrakcyjność krajobrazu. Pozyskane dane satelitarne i opracowane cyfrowe mapy usług ekosystemowych stanowią cenne narzędzie wspierające procesy decyzyjne. Pozyskana w ten sposób wiedza pozwala organom samorządowym, planistom przestrzennym oraz instytucjom ochrony przyrody na skuteczniejsze wdrażanie zasad zrównoważonego rozwoju, racjonalne gospodarowanie zasobami środowiska oraz efektywną adaptację do zmian klimatycznych.

Słowa kluczowe: mapowanie, usługi ekosystemowe, techniki satelitarne

Mapping ecosystem services using satellite techniques

The article addresses the modern assessment and monitoring of natural capital using advanced remote sensing techniques and satellite data. The main objective of the article is to present the methodology for operational mapping of these services using multi-source satellite data, illustrated with the example of the Lubusz Voivodeship. The article details the application of modern Earth observation tools that enable objective, continuous, and repeatable analysis of environmental changes. The paper

describes the practical possibilities of using ecosystem service mapping in areas including water retention, air pollution reduction, and landscape attractiveness. The acquired satellite data and developed digital maps of ecosystem services constitute a valuable tool supporting decision-making processes. The knowledge gained in this way enables local government authorities, spatial planners, and nature conservation institutions to more effectively implement the principles of sustainable development, rationally manage environmental resources, and efficiently adapt to climate change.

Keywords: mapping, ecosystem services, satellite techniques

Mariola Chrzanowska, Szkoła Główna Gospodarstwa Wiejskiego w Warszawie

Gdy etykieta nie ujawnia wyniku rzeczywistego: hybrydowe podejście PU-Heckmana do modelowania częściowo obserwowalnych danych

Metody uczenia maszynowego są powszechnie wykorzystywane w modelowaniu predykcyjnym. Większość standardowych podejść zakłada jednak, że obserwowana zmienna wynikowa wiernie odzwierciedla rzeczywisty stan analizowanego zjawiska. Założenie to może być naruszone w sytuacji częściowej obserwowalności etykiet. Brak etykiety pozytywnej nie musi bowiem oznaczać niewystąpienia zjawiska, lecz może wynikać z jego nieujawnienia, nieodnotowania lub działania innych mechanizmów selekcyjnych. W takich przypadkach obserwowana etykieta zależy zarówno od rzeczywistego wystąpienia zjawiska, jak i od procesu decydującego o tym, czy zostanie ono zaobserwowane. Pominięcie tego mechanizmu może prowadzić do obciążenia oszacowań oraz pogorszenia jakości predykcji.

W pracy zaproponowano hybrydowe podejście łączące metodologię Positive-Unlabeled Learning (PU Learning), rozwijaną między innymi przez Elkana i Noto [2], z komponentem selekcyjnym inspirowanym modelem Heckmana, odnoszącym się do klasycznego problemu błędu selekcji [1]. PU Learning wykorzystuje informacje zawarte w obserwacjach pozytywnych i nieoznaczonych, natomiast komponent selekcyjny modeluje proces ujawniania przypadków pozytywnych. Proponowana konstrukcja integruje statystyczne metody uczenia z możliwością jawnego modelowania mechanizmu selekcji [3]. Jej głównym celem jest rozdzielenie procesu generowania badanego zjawiska od procesu obserwacji jego etykiety, zamiast automatycznego klasyfikowania wszystkich obserwacji nieoznaczonych jako negatywne.

Szczególną uwagę poświęcono zagadnieniom identyfikacji modelu oraz interpretacji uzyskiwanych oszacowań. Komponent korekty selekcji jest wykorzystywany warunkowo, zależnie od przyjętych założeń dotyczących mechanizmu ujawniania oraz zakresu informacji dostępnych w danych. W sytuacji, gdy warunki identyfikacji nie pozwalają na jednoznaczne odtworzenie nieobserwowanego wyniku, oszacowane prawdopodobieństwa należy interpretować warunkowo, tj. w odniesieniu do przyjętej specyfikacji modelu i związanych z nią założeń.

Proponowane podejście zostanie poddane ocenie w badaniu symulacyjnym obejmującym scenariusze SCAR, SAR i MNAR. Wyniki zostaną porównane z metodą korekty g/c , modelami stosowanymi niezależnie oraz standardowym podejściem stackingowym (łączenia modeli). Analiza skoncentruje się na jakości predykcji, skali obciążenia oszacowań oraz stopniu, w jakim uzyskane prawdopodobieństwa mogą być interpretowane jako oszacowania nieobserwowanego wyniku rzeczywistego.

Literatura:

Heckman, J.: *Sample selection bias as a specification error*. *Econometrica* 47(1), s. 153–161 (1979).

Elkan, C., Noto, K.: *Learning classifiers from only positive and unlabeled data*. In: Proceedings of the 14th ACM SIGKDD, s. 213–220 (2008).

Hastie, T., Tibshirani, R., Friedman, J.: *The Elements of Statistical Learning*. Springer, New York (2009).

Słowa kluczowe: uczenie positive-unlabeled, błąd selekcji, niepełne etykiety, dane częściowo obserwowane, korekta Heckmana

When the Label Does Not Reveal the True Outcome: A Hybrid PU–Heckman Approach to Modeling Partially Observed Data

Machine learning methods are widely used for predictive modeling. Most standard approaches, however, assume that the observed outcome accurately reflects the true state of the phenomenon of interest. This assumption may be violated when labels are only partially observed. In such settings, the absence of a positive label does not necessarily indicate the absence of the underlying outcome. Instead, it may result from non-disclosure, underreporting, or other selection processes. Consequently, the observed label depends not only on whether the outcome occurs, but also on whether it becomes observable. Ignoring this mechanism may lead to biased estimates and degraded predictive performance.

This paper proposes a hybrid framework that combines Positive-Unlabeled (PU) Learning, following Elkan and Noto [2], with a selection component inspired by Heckman's classical sample selection model [1]. PU Learning exploits information contained in positive and unlabeled observations, while the selection component models the process through which positive cases become observed. The proposed framework integrates statistical learning methods with explicit modeling of the label-observation process [3]. Its primary objective is to disentangle the process generating the outcome of interest from the process governing its observation, rather than automatically treating all unlabeled observations as negative cases.

Particular attention is devoted to identification and interpretation. The selection-correction component is applied conditionally, depending on assumptions about the label-observation mechanism and the information available in the data. When identification conditions are insufficient to uniquely recover the unobserved outcome, the resulting probabilities should be interpreted conditionally, that is, in relation to the assumed model specification and its underlying assumptions.

The proposed approach will be evaluated through a simulation study covering SCAR, SAR, and MNAR settings. Its performance will be compared with the g/c correction method, stand-alone PU and selection models, and a standard stacking approach. The evaluation will focus on predictive performance, estimation bias, and the extent to which the resulting probabilities can be interpreted as estimates of the underlying true outcome.

References

Heckman, J.: *Sample selection bias as a specification error*. *Econometrica* 47(1), 153–161 (1979).

Elkan, C., Noto, K.: *Learning classifiers from only positive and unlabeled data*. In: Proceedings of the 14th ACM SIGKDD, pp. 213–220 (2008).

Hastie, T., Tibshirani, R., Friedman, J.: *The Elements of Statistical Learning*. Springer, New York (2009).

Keywords: positive-unlabeled learning, selection bias, incomplete labels, partially observed data, Heckman correction

Przestrzenne zróżnicowanie konwergencji wydajności pracy w sektorze rolnictwa, leśnictwa i rybołówstwa w regionach NUTS 3 Europy: podejście GWR i klasteryzacji przestrzennej

Celem artykułu jest zbadanie przestrzennego zróżnicowania procesu β -konwergencji wydajności pracy w sektorze A, obejmującym rolnictwo, leśnictwo i rybołówstwo, w europejskich regionach NUTS 3 w latach 2003–2023. Wydajność pracy będzie mierzona jako relacja wartości dodanej brutto wytworzonej w tym sektorze do liczby osób w nim pracujących.

Analiza zostanie przeprowadzona w kilku etapach. Najpierw oszacowany będzie globalny model β -konwergencji, w którym średnie roczne logarytmiczne tempo zmian wydajności pracy zostanie powiązane z jej początkowym poziomem. Następnie zastosowana zostanie geograficznie ważona regresja, pozwalająca uchwycić lokalne różnice w sile i kierunku zależności między wyjściowym poziomem wydajności a tempem jej wzrostu. Badanie uzupełni przestrzenna analiza skupień, w tym metoda spatial k-means, służąca wyodrębnieniu grup regionów o zbliżonych poziomach produktywności, dynamice wzrostu oraz lokalnych parametrach konwergencji.

Takie podejście pozwoli ustalić, czy proces konwergencji przebiega w Europie w sposób względnie jednorodny, czy też przyjmuje odmienne formy w poszczególnych częściach kontynentu, wskazując na obszary nadrabiania dystansu, stagnacji lub dywergencji. Wyniki badania mogą przyczynić się do lepszego zrozumienia terytorialnego zróżnicowania zmian produktywności w sektorze A, a także stanowić podstawę do formułowania bardziej dopasowanych przestrzennie rekomendacji dla polityki regionalnej, rolnej i polityki spójności.

Słowa kluczowe: wydajność pracy; β -konwergencja regionalna; GWR; klasteryzacja przestrzenna

Spatial Heterogeneity of Labour Productivity Convergence in Agriculture, Forestry and Fishing across European NUTS 3 Regions: A GWR and Spatial Clustering Approach

The aim of this article is to examine the spatial variation in the β -convergence of labour productivity in Sector A—which includes agriculture, forestry and fishing—across European NUTS 3 regions between 2003 and 2023. Labour productivity will be measured as the ratio of gross value added generated in the sector to the number of people employed in it.

The analysis will be conducted in several stages. First, a global β -convergence model will be estimated, linking the average annual logarithmic growth rate of labour productivity to its initial level. This will be followed by geographically weighted regression, which will make it possible to capture local differences in both the strength and direction of the relationship between initial productivity and subsequent growth. The analysis will be complemented by spatial clustering, including spatial k-means, to identify groups of regions with similar productivity levels, growth trajectories and local convergence parameters.

This approach will help determine whether convergence across Europe follows a broadly uniform pattern or takes different forms in different parts of the continent, revealing areas of catch-up, stagnation or divergence. The findings may contribute to a better understanding of the territorial

variation in productivity changes within Sector A and provide a basis for developing more place-sensitive recommendations for regional, agricultural and cohesion policy.

Keywords: labour productivity; regional β -convergence; GWR; spatial clustering

Svitlana Chugaievska, Andrzej Frycz Modrzewski Krakow University; Dariusz Fatula, Andrzej Frycz Modrzewski Krakow University; Kateryna Romanchuk, Poznań University of Economics and Business

Przedsiębiorczość ukraińska w Polsce po 2022 roku: czynniki rozwoju, wkład gospodarczy oraz perspektywy integracji

Masowa migracja obywateli Ukrainy do Polski po wybuchu pełnoskalowej wojny w Ukrainie w 2022 roku znacząco wpłynęła na krajobraz gospodarczy obu państw. Jednym z najbardziej widocznych zjawisk stał się dynamiczny rozwój przedsiębiorczości ukraińskiej w Polsce. Fenomen ten stanowi nie tylko mechanizm adaptacji ekonomicznej i samozatrudnienia migrantów, lecz także istotny czynnik wspierający rozwój polskiej gospodarki. Celem niniejszego badania jest identyfikacja kluczowych czynników sprzyjających rozwojowi przedsiębiorstw prowadzonych przez obywateli Ukrainy w Polsce po 2022 roku, ocena ich wkładu gospodarczego oraz analiza perspektyw długoterminowej integracji z polskim środowiskiem biznesowym.

Badanie koncentruje się na instrumentach wsparcia instytucjonalnego, uwarunkowaniach prawnych i regulacyjnych, dostępie do zasobów finansowych, kompetencjach przedsiębiorczych oraz roli sieci transnarodowych łączących Ukrainę i Polskę. W pracy zastosowano podejście mieszane (mixed-methods), obejmujące analizę oficjalnych danych statystycznych, raportów politycznych oraz najnowszej literatury naukowej, a także autorskie badanie ankietowe przeprowadzone wśród migrantów z Ukrainy mieszkających w Polsce. Wyniki ankiety stanowią podstawę empiryczną do opracowania modelu ekonometrycznego służącego identyfikacji determinant aktywności przedsiębiorczej wśród migrantów ukraińskich. Szczególną uwagę poświęcono cechom socjo-demograficznym, doświadczeniu zawodowemu, dostępowi do zasobów oraz czynnikom integracyjnym wpływającym na prawdopodobieństwo podjęcia działalności gospodarczej.

Uzyskane wyniki wskazują, że przedsiębiorstwa prowadzone przez obywateli Ukrainy przyczyniają się do tworzenia miejsc pracy, zwiększania wpływów podatkowych, dywersyfikacji rynku oraz wzmacniania więzi gospodarczych między Polską a Ukrainą. Jednocześnie przedsiębiorcy napotykają liczne wyzwania związane z nieznaną procedur administracyjnych, ograniczonym dostępem do finansowania, barierami językowymi oraz konkurencją rynkową. Analiza ekonometryczna pozwoliła zidentyfikować najważniejsze czynniki wpływające na zaangażowanie przedsiębiorcze i podkreśliła znaczenie wsparcia instytucjonalnego oraz kapitału społecznego w stymulowaniu aktywności gospodarczej.

W artykule wskazano, że przedsiębiorczość ukraińska stała się istotnym czynnikiem społeczno-ekonomicznej integracji migrantów oraz ważnym impulsem rozwoju dwustronnej współpracy gospodarczej. Długoterminowa trwałość tego zjawiska będzie zależeć od skuteczności polityk integracyjnych, kontynuacji wsparcia instytucjonalnego oraz szerszego otoczenia gospodarczego i politycznego. Wyniki badania przyczyniają się do pogłębienia wiedzy na temat przedsiębiorczości migrantów w Europie Środkowo-Wschodniej oraz dostarczają rekomendacji opartych na dowodach naukowych dla decydentów dążących do zwiększenia pozytywnego wpływu aktywności przedsiębiorczej w społecznościach migranckich.

Słowa kluczowe: przedsiębiorczość ukraińskich migrantów, integracja gospodarcza, działalność gospodarcza w Polsce

Ukrainian Entrepreneurship in Poland after 2022: Development Factors, Economic Contribution, and Integration Prospects

The large-scale migration of Ukrainian citizens to Poland following the outbreak of the full-scale war in Ukraine in 2022 has significantly influenced the economic landscape of both countries. Among the most notable developments has been the rapid growth of Ukrainian entrepreneurship in Poland. This phenomenon has emerged not only as a mechanism for economic adaptation and self-employment among migrants but also as an important contributor to the Polish economy. The purpose of this study is to examine the key factors that have facilitated the development of Ukrainian-owned businesses in Poland after 2022, assess their economic contribution, and analyze their prospects for long-term integration into the Polish business environment.

The research focuses on institutional support measures, legal and regulatory conditions, access to financial resources, entrepreneurial competencies, and the role of transnational networks connecting Ukraine and Poland. The study employs a mixed-method approach combining the analysis of official statistical data, policy reports, and recent academic literature with original survey data collected from Ukrainian migrants residing in Poland. The survey results serve as the empirical basis for the development of an econometric model aimed at identifying the determinants of entrepreneurial activity among Ukrainian migrants. Particular attention is paid to sociodemographic characteristics, professional experience, access to resources, and integration factors influencing the likelihood of business creation. The findings indicate that Ukrainian businesses contribute to job creation, tax revenues, market diversification, and the strengthening of economic ties between Poland and Ukraine. At the same time, entrepreneurs face challenges related to misunderstanding of administrative procedures, access to financing, language barriers, and market competition. The econometric analysis reveals the most significant factors affecting entrepreneurial engagement and highlights the importance of institutional support and social capital in facilitating business activity.

The paper argues that Ukrainian entrepreneurship has become an important factor in the socio-economic integration of migrants and a driver of bilateral economic cooperation. The long-term sustainability of this trend will depend on the effectiveness of integration policies, continued institutional support, and the broader economic and political environment. The results contribute to the understanding of migrant entrepreneurship in Central and Eastern Europe and provide evidence-based recommendations for policymakers seeking to enhance the positive impact of entrepreneurial activity among migrant communities.

Keywords: Ukrainian migrants' entrepreneurship, economic integration, Ukrainian migrants, business activity in Poland.

Czy reżimy wysokiej zmienności są zbieżne z aprecjacją franka szwajcarskiego? Analiza CHF/USD i CHF/AUD z wykorzystaniem modeli MS-GARCH

Celem badania jest ocena, czy reżimy wysokiej zmienności kursów franka szwajcarskiego są zbieżne z aprecjacją CHF oraz czy wyniki zależą od wyboru waluty odniesienia. Analiza obejmuje dzienne stopy zwrotu CHF/USD i CHF/AUD w latach 2000-2026. W badaniu zastosowano dwureżimowe modele MS-GARCH z warunkowym rozkładem t-Studenta, estymowane oddzielnie dla obu kursów. Reżimy wysokiej zmienności wyznaczono na podstawie wygładzonych prawdopodobieństw reżimowych, a następnie oceniono, czy zidentyfikowanym epizodom towarzyszyła aprecjacja franka oraz jakie było ich tło ekonomiczne. Wyniki wskazują na wyraźne różnice między analizowanymi kursami. Dla CHF/USD model identyfikuje pięć epizodów wysokiej zmienności, natomiast dla CHF/AUD 29 epizodów. Aprecjacja franka występuje w większości epizodów w obu parach, ale ich interpretacja nie jest jednolita. Najbardziej jednoznaczny sygnał zgodny z funkcją safe haven pojawia się w okresie pandemii COVID-19, natomiast epizody z lat 2011 i 2015 są silnie związane z decyzjami Szwajcarskiego Banku Narodowego dotyczącymi minimalnego kursu EUR/CHF. Porównanie CHF/USD i CHF/AUD pokazuje, że wybór waluty odniesienia wpływa nie tylko na ocenę skali aprecjacji franka, ale także na identyfikację reżimów wysokiej zmienności. Jest to istotne, ponieważ dolar amerykański często pełni rolę podstawowej waluty odniesienia, choć sam może przyciągać kapitał w okresach napięć rynkowych.

Słowa kluczowe: frank szwajcarski; waluty safe haven; MS-GARCH; reżimy wysokiej zmienności; szoki globalne

Do High-Volatility Regimes Coincide with Swiss Franc Appreciation? Markov-Switching Evidence from CHF/USD and CHF/AUD

The aim of the study is to assess whether high-volatility regimes in Swiss franc exchange rates are associated with CHF appreciation and whether the results depend on the choice of reference currency. The analysis covers daily CHF/USD and CHF/AUD returns for the period 2000-2026. The study applies two-regime MS-GARCH models with a conditional Student's t distribution, estimated separately for both exchange rates. High-volatility regimes are identified using smoothed regime probabilities. The identified episodes are then assessed in terms of Swiss franc appreciation and their economic background. The results show clear differences between the two exchange rates. For CHF/USD, the model identifies five high-volatility episodes, while for CHF/AUD it identifies 29 episodes. Swiss franc appreciation occurs in most episodes in both pairs, but their interpretation is not uniform. The clearest signal consistent with safe-haven behaviour appears during the COVID-19 pandemic, while the 2011 and 2015 episodes are strongly related to Swiss National Bank decisions concerning the EUR/CHF minimum exchange rate. The comparison between CHF/USD and CHF/AUD shows that the choice of reference currency affects not only the assessment of the scale of Swiss franc appreciation, but also the identification of high-volatility regimes. This is important because the U.S. dollar is often used as the main reference currency, although it may itself attract capital during periods of market stress.

Keywords: Swiss franc; safe-haven currencies; MS-GARCH; high-volatility regimes; global shocks

Czasowo-przestrzenne grafowe sieci neuronowe w modelowaniu ryzyka systemowego sektora ubezpieczeniowego z uwzględnieniem fizycznego ryzyka klimatycznego

Fizyczne ryzyko klimatyczne stało się jednym z głównych wyzwań dla stabilności finansowej, przy czym sektor ubezpieczeniowy pełni kluczową rolę w absorbowaniu strat związanych ze zmianami klimatu. Jednak dotychczasowe badania nad ryzykiem systemowym w sektorze ubezpieczeniowym koncentrują się przede wszystkim na danych na poziomie pojedynczych przedsiębiorstw, poświęcając stosunkowo niewiele uwagi współzależnościom między krajami oraz dynamice nieliniowej zależności. W niniejszym artykule analizowany jest wpływ fizycznego ryzyka klimatycznego na ryzyko systemowe sektora ubezpieczeniowego na poziomie krajowym z wykorzystaniem czasowo-przestrzennego modelu grafowych sieci neuronowych (Graph Neural Network, GNN).

Kraje są modelowane jako węzły dynamicznej sieci, natomiast krawędzie odzwierciedlają podobieństwo profili ryzyka klimatycznego poszczególnych krajów, wyznaczone na podstawie indeksów ND-GAIN oraz Global Climate Physical Risk Index (GCPRI). Zaproponowana architektura łączy Graph Attention Network (GAT), umożliwiającą endogeniczne uczenie się struktury powiązań i znaczenia poszczególnych krajów w sieci, z Gated Recurrent Unit (GRU), która pozwala uchwycić dynamikę w czasie. Ryzyko systemowe jest mierzone za pomocą podejścia DeltaCoVaR, co umożliwia ocenę wkładu sektora ubezpieczeniowego każdego kraju w ryzyko systemowe globalnego systemu ubezpieczeniowego.

Model pozwala jednocześnie uwzględnić nieliniowe zależności, strukturę sieci oraz ich ewolucję w czasie. Wyniki empiryczne wskazują, że fizyczne ryzyko klimatyczne istotnie wpływa na ryzyko systemowe i rozprzestrzenia się poprzez sieć. Zaproponowane podejście umożliwia identyfikację krajów o znaczeniu systemowym oraz potencjalnych kanałów transmisji ryzyka klimatycznego, przyczyniając się do wykorzystania informacji o fizycznym ryzyku klimatycznym w modelach ryzyka systemowego opartych na zaawansowanych metodach uczenia maszynowego.

Słowa kluczowe: ryzyko systemowe; fizyczne ryzyko klimatyczne; sektor ubezpieczeniowy; grafowe sieci neuronowe; Graph Attention Network; DeltaCoVaR.

A Spatio-Temporal Graph Neural Network Approach to Modeling Systemic Risk in the Insurance Sector under Physical Climate Risk

Physical climate risk has become a major challenge to financial stability, with the insurance sector acting as a key absorber of climate-related losses. However, existing studies on systemic risk in insurance focus primarily on firm-level data, with limited attention to cross-country interdependencies and nonlinear dynamics. This paper examines the impact of physical climate risk on systemic risk in the insurance sector at the national level using a spatio-temporal Graph Neural Network (GNN) framework.

Countries are modeled as nodes in a dynamic network, where edges reflect similarities in countries' climate risk profiles derived from ND-GAIN and the Global Climate Physical Risk Index. The proposed architecture combines a Graph Attention Network (GAT), enabling endogenous learning

of cross-country influence structures, with a Gated Recurrent Unit (GRU) to capture temporal dynamics. Systemic risk is quantified using the DeltaCoVaR framework, allowing us to assess each country's contribution to the systemic risk of the global insurance system.

The model jointly captures nonlinear dependencies, network structure, and their temporal evolution. Empirical results indicate that physical climate risk significantly affects systemic risk and propagates through the network. The approach identifies systemically important countries and reveals potential channels of climate risk transmission, contributing to the incorporation of physical climate risk information into systemic risk models based on advanced machine learning methods

Keywords: systemic risk; physical climate risk; insurance sector; Graph Neural Networks; Graph Attention Network; DeltaCoVaR.

Krzysztof Dmytrów, Uniwersytet Szczeciński

Wybór lokalizacji podczas kompletacji towarów za pomocą metod porządkowania liniowego – wpływ normalizacji zmiennych na długość trasy magazyniera

Jeżeli przedsiębiorstwo stosuje współdzielone składowanie towarów w magazynie, wówczas każdy indeks towarowy może być składowany w wielu lokalizacjach. Podczas procesu kompletacji istotne staje się zagadnienie wyboru lokalizacji, z których magazynier pobierze towary znajdujące się na zamówieniu. Podczas wyboru lokalizacji można realizować różne strategie wyboru wykorzystując w tym celu metody porządkowania liniowego, nadając zmiennym opisującym lokalizacje odpowiednie wagi. Do wyboru lokalizacji wykorzystano metodę TMAL (Taksonomiczna Miara Atrakcyjności Lokalizacji), opartą na syntetycznym mierniku rozwoju Z. Hellwiga z wykorzystaniem euklidesowych odległości obiektów (lokalizacji, w których składowane są kompletowane towary) od wzorca. Jednym z ważniejszych etapów każdej metody porządkowania liniowego jest normalizacja zmiennych. W analizowanym przykładzie zmienne opisujące lokalizacje są mierzone na skali ilorazowej. Dlatego porównano wyniki dla trzech formuł normalizacyjnych: standaryzacji, unitaryzacji zerowanej oraz normalizacji wektorowej (jednego z przekształceń ilorazowych). Metody te należą do trzech głównych grup metod normalizacyjnych: standaryzacji, unitaryzacji oraz przekształceń ilorazowych. Badanie przeprowadzono za pomocą metod symulacyjnych. Wygenerowano 1000 scenariuszy (zleceń kompletacyjnych). Za pomocą metody TMAL, z wykorzystaniem trzech analizowanych metod normalizacji dokonano wyboru lokalizacji do odwiedzenia przez magazyniera, a następnie wyznaczono długość pokonywanej przez magazyniera trasy. Wyniki uzyskane dla trzech omawianych metod normalizacji porównano za pomocą jednoczynnikowej analizy wariancji.

Słowa kluczowe: kompletacja, porządkowanie liniowe, metoda Hellwiga, normalizacja, metody symulacyjne

Locations selection during order picking using linear ordering methods – the impact of normalisation of variables on the picker's route length

If a company utilises shared storage of items in a warehouse, then each item code may be stored in multiple locations. During the order picking process, the issue of selecting the locations from which the picker will pick the items on the order becomes crucial. When selecting locations, various selection

strategies can be implemented using linear ordering methods, assigning appropriate weights to the variables describing the locations. The TMAL method (Polish abbreviation of the Taxonomic Measure of Location Attractiveness) was used for selection of locations, based on Z. Hellwig's synthetic measure of development using Euclidean distances of objects (locations where the items to be picked are stored) from a pattern. One of the most important stages of any linear ordering method is the normalisation of variables. In the analysed example, the variables describing the locations are measured on a ratio scale. Therefore, the results for three normalisation formulas were compared: z-score, min-max normalisation and vector normalisation (one of the quotient transformations). These methods belong to the three main groups of normalisation methods: standardisation, feature scaling and quotient transformations. The research was conducted using simulation methods. A total of 1,000 scenarios (orders) were generated. Using the TMAL method, and applying the three selected normalisation methods, the locations to be visited by the picker were selected, and the route length was determined. The results obtained for the three normalisation methods were compared using one-way analysis of variance.

Keywords: order picking, linear ordering, Hellwig's method, normalisation, simulation methods

Andrzej Dudek, Uniwersytet Ekonomiczny we Wrocławiu; Marcin Pełka Uniwersytet Ekonomiczny we Wrocławiu; Artur Skiba - kandydat do Szkoły Doktorskiej

Wyjaśnianie regionalnych różnic płacowych z wykorzystaniem uczenia maszynowego: podejście interpretacyjne oparte na SHAP

W prezentacji podjęty zostanie temat regionalnego zróżnicowania płac w Polsce. Badanie przeprowadzono z wykorzystaniem modeli uczenia maszynowego wzbogaconych o techniki wyjaśniania modeli (Explanatory Model Analysis). Jego celem było wyjaśnienie czynników, które w największym stopniu wpływają na różnice w poziomie wynagrodzeń między polskimi powiatami. Na podstawie danych na poziomie powiatów pochodzących z Banku Danych Lokalnych (BDL) dla lat 2010 i 2023 opracowano model służący do prognozowania poziomu wynagrodzeń w oparciu o zmienne ekonomiczne, demograficzne, infrastrukturalne i środowiskowe.

W celu interpretacji modelu zastosowano podejścia Variable Importance over Permutation (VIP) oraz SHapley Additive exPlanations (SHAP), które umożliwiają analizę zarówno globalnej ważności zmiennych, jak i lokalnego wkładu poszczególnych czynników. Wyniki wskazały, że udział ludności w wieku produkcyjnym, stopa bezrobocia oraz podatność społeczna pozostają kluczowymi determinantami różnic płacowych, choć ich względne znaczenie istotnie zmienia się w czasie.

Z kolei analiza SHAP pokazała, że konteksty regionalne dla powiatów jeleniogórskiego wrocławskiego, charakteryzują się odmienną dynamiką czynników, przy czym zmienne demograficzne i infrastrukturalne odgrywają różne role w badanych latach. Wyniki podkreśliły potencjał łączenia uczenia maszynowego z metodami wyjaśnialności w identyfikacji złożonych, nieliniowych determinant wynagrodzeń oraz w dostarczaniu bardziej przejrzystych podstaw analitycznych do zrozumienia zmieniających się nierówności regionalnych.

Słowa kluczowe: uczenie maszynowe, metody analizy wyjaśnialności modelu, analiza różnic płacowych

Explaining regional wage disparities with machine learning: A SHAP-based interpretation approach

The aim of the study is to provide an explanation for the factors that most influence the differences in wage levels between Polish powiats (equivalent to counties). This study investigates regional wage disparities in Poland by applying machine learning models enhanced by Explanatory Model Analysis techniques. Using powiat-level data from the Local Data Bank (Pol. Bank Danych Lokalnych – BDL) for 2010 and 2023, a machine learning framework was developed to predict wage levels based on economic, demographic, infrastructural and environmental variables. To interpret the model, we employed the Variable Importance over Permutation (VIP) and SHapley Additive exPlanations (SHAP) approaches, which provide insights into both the global feature importance and the local contributions of individual variables. The results indicate that the share of the productive population, unemployment rates and social vulnerability remain key determinants of wage differences, although their relative influence shifts significantly over time. The SHAP analysis demonstrates how regional contexts such as the Jelenia Góra and Wrocław powiats exhibit distinct factor dynamics, with demographic and infrastructural variables playing varying roles across the studied years. The findings highlight the potential of combining machine learning with explainability methods to uncover complex, nonlinear determinants of wages, offering a more transparent analytical basis for understanding evolving regional disparities.

Keywords: machine learning, explanatory model analysis, wage disparities

Marta Dziechciarz, Uniwersytet Ekonomiczny we Wrocławiu; Marcin Pełka, Uniwersytet Ekonomiczny we Wrocławiu

Różne podejścia skalowania wielowymiarowego danych symbolicznych w wizualizacji wyników porządkowania liniowego

W artykule zaprezentowano rozwiązania metodyczne pozwalające na przeprowadzenie porządkowania liniowego wraz z wizualizacją wyników dla danych symbolicznych. Podstawą do zastosowania porządkowania dla tego typu danych jest odpowiednia miara odległości (m.in. de Carvalho, Wassersteina), które pozwalają na graficzną reprezentację wyników oraz wyznaczenie izokwant rozwoju. W części empirycznej na podstawie danych regionalnych OECD przygotowano dane opisujące kraje Unii Europejskiej z wykorzystaniem danych symbolicznych, które poddano dwustopniowej procedurze porządkowania liniowego.

Słowa kluczowe: porządkowanie liniowe, dane symboliczne, skalowanie wielowymiarowe

Different multidimensional scaling methods for symbolic data for visualization of linear ordering results

The article presents methodological steps enabling linear ordering together with the visualization of results for symbolic data. The application of linear ordering methods to this type of data is based on appropriate distance measures (for example de Carvalho and Wasserstein), which allow for the graphical representation of results and the determination of development isoquants. In the empirical section, based on OECD regional data, symbolic data describing European Union countries were prepared and subjected to a two-stage linear ordering procedure.

Keywords: linear ordering, symbolic data, multidimensional scaling

Wielowymiarowa analiza porównawcza stabilności rankingów efektywności środowiskowej krajów UE w obliczu globalnych szoków: perspektywa PBA vs. CBA

Celem referatu jest ocena stabilności rankingów względnej efektywności środowiskowej państw Unii Europejskiej oraz ustalenie, w jakim stopniu obraz tej stabilności zależy od zastosowania dwóch alternatywnych sposobów rachunkowości emisji: opartego na produkcji (PBA – Production-Based Accounting) oraz na konsumpcji (CBA – Consumption-Based Accounting). Szczególna uwaga zostanie poświęcona zmianom względnych pozycji państw w okresach globalnych zaburzeń gospodarczych, zdrowotnych i energetyczno-geopolitycznych oraz w następujących po nich fazach dostosowawczych.

W badaniu zostanie wykorzystana nieparametryczna metoda Data Envelopment Analysis (DEA), w wariancie nieradialnej, nieorientowanej kierunkowej funkcji odległości (NDDF), przy zmiennych efektach skali i słabej dyspozycyjności emisji CO₂. Model obejmie trzy nakłady: nakłady brutto na środki trwałe, zatrudnienie i zużycie energii pierwotnej; PKB będzie stanowił efekt pożądany, a emisje CO₂ – efekt niepożądany. Dla każdego roku zostaną wyznaczone odrębne rankingi PBA i CBA, a ich zgodność i stabilność ocenione z wykorzystaniem współczynnika Kendalla τ_b oraz miar zmian rang.

Analiza będzie bazować na danych panelowych dla państw UE-27 z lat 2005–2023, pochodzących z baz Eurostat oraz Global Carbon Project. Przyjęty okres obejmuje globalny kryzys finansowy, pandemię COVID-19 oraz szok energetyczno-geopolityczny związany z rosyjską inwazją na Ukrainę, co pozwoli porównać stabilność rankingów w fazach wejścia w szok i późniejszego dostosowania.

Badanie pozwoli zidentyfikować państwa o najbardziej i najmniej stabilnych pozycjach rankingowych, wskazać przypadki największej asymetrii reakcji rankingów PBA i CBA oraz ocenić, czy większe przetasowania pojawiają się bezpośrednio w momencie wystąpienia szoku, czy raczej w następujących po nim okresach dostosowawczych. Uzupełniającą zostanie przeanalizowany związek asymetrii zmian rankingów z wybranymi charakterystykami emisyjnymi i strukturalnymi gospodarek.

Słowa kluczowe: Data Envelopment Analysis (DEA), efektywność środowiskowa, PBA vs. CBA, stabilność rankingów, globalne szoki.

Multivariate Comparative Analysis of the Stability of Environmental Efficiency Rankings in EU Countries under Global Shocks: A PBA vs. CBA Perspective

The aim of the paper is to assess the stability of relative environmental-efficiency rankings for European Union countries and to determine the extent to which this stability depends on two alternative approaches to emission accounting: production-based accounting (PBA) and consumption-based accounting (CBA). Particular attention will be paid to changes in countries' relative positions during global economic, health, and energy-geopolitical disturbances and during the subsequent adjustment phases.

The study will employ the non-parametric Data Envelopment Analysis (DEA) method using a non-radial, non-oriented directional distance function (NDDF), with variable returns to scale and weak disposability of CO₂ emissions. The model will include three inputs—gross fixed capital formation, employment, and primary energy consumption—GDP as the desirable output, and CO₂ emissions

as the undesirable output. Separate PBA and CBA rankings will be constructed for each year, and their agreement and stability will be assessed using Kendall's τ_b and measures of rank changes.

The analysis will be based on panel data for EU-27 countries covering 2005–2023, drawn from Eurostat and the Global Carbon Project. The period includes the global financial crisis, the COVID-19 pandemic, and the energy-geopolitical shock associated with Russia's invasion of Ukraine, allowing ranking stability to be compared between shock-entry and subsequent adjustment phases.

The study will identify countries with the most and least stable ranking positions, highlight cases of the strongest asymmetry between PBA and CBA ranking responses, and examine whether larger reshufflings occur immediately at shock entry or during subsequent adjustment periods. In addition, the relationship between ranking asymmetry and selected emission-related and structural characteristics of national economies will be explored.

Keywords: Data Envelopment Analysis (DEA), environmental efficiency, PBA vs. CBA, ranking stability, global shocks

Andrzej Geise, Uniwersytet Mikołaja Kopernika w Toruniu

Ukryta struktura zależności między rynkiem ropy a wzrostem gospodarczym: analiza gospodarek UE przy użyciu dekompozycji falkowej i analizy skupień

Od czasu globalnego kryzysu finansowego lat 2008–2009 problem niepewności stał się jednym z kluczowych zagadnień w badaniach nad funkcjonowaniem gospodarki, rynków finansowych, czy rynków surowców energetycznych. W ostatnich latach znaczenie to dodatkowo wzrosło w związku z nakładaniem się szoków o charakterze ekonomicznym, pandemicznym i geopolitycznym, wynikających z pandemii COVID-19, rosyjskiej agresji na Ukrainę oraz eskalacji napięć na Bliskim Wschodzie. Wydarzenia te pokazały, że relacja między cenami ropy naftowej a aktywnością gospodarczą nie może być analizowana w oderwaniu od szerszego środowiska ryzyka i niepewności. Szczególnie istotne staje się zatem pytanie, czy różne źródła niepewności w podobny sposób oddziałują na współzależność między rynkiem ropy a wzrostem gospodarczym, czy też ich wpływ zależy od horyzontu czasowego i częstotliwości. Z tego względu analiza tej relacji wymaga podejścia umożliwiającego jednoczesne uchwycenie zmian w czasie oraz na różnych skalach częstotliwości wahań (Chen, 2023).

Celem artykułu jest identyfikacja i stworzenie profili grup gospodarek Unii Europejskiej według podobieństwa reakcji wzrostu gospodarczego na szoki cen ropy naftowej w warunkach podwyższonego ryzyka geopolitycznego oraz niepewności cen ropy. W badaniu empirycznym założono, że wpływ szoków naftowych na wzrost gospodarczy nie ma charakteru jednorodnego, lecz zależy od horyzontu czasowego, fazy cyklu koniunkturalnego oraz źródła niepewności. Dotychczasowe badania wskazują, że relacja między cenami ropy a aktywnością gospodarczą różni się w zależności od częstotliwości i może być szczególnie istotna w dłuższym horyzoncie (Redin et al., 2018), a narzędzia falkowe pozwalają uchwycić zmienność tej zależności zarówno w czasie, jak i przy różnych częstotliwościach (Aloui et al., 2018).

W badaniu zostanie zastosowana koherencja falkowa między wzrostem gospodarczym a cenami ropy dla krajów UE, a następnie cząstkowa koherencja falkowa, pozwalająca ocenić, czy obserwowana

zależność utrzymuje się po uwzględnieniu ryzyka geopolitycznego oraz niepewności cen na rynku ropy. Z uzyskanych map koherencji zostaną wyodrębnione syntetyczne wskaźniki opisujące siłę relacji w krótkim, średnim i długim horyzoncie, jej trwałość w okresach kryzysowych oraz skalę osłabienia po kontroli czynników niepewności. Wskaźniki te posłużą do grupowania gospodarek UE zgodnie z podejściem, w którym grupowanie opiera się na podobieństwie dynamiki, a nie wyłącznie na klasycznych cechach strukturalnych badanych gospodarek (Jokubaitis & Celov, 2023).

Pomimo rosnącej liczby badań dotyczących wpływu cen ropy naftowej, ryzyka geopolitycznego i niepewności na aktywność gospodarczą, wciąż brakuje analiz, które pozwalałyby grupować gospodarki UE według charakteru ich podatności na szoki zewnętrzne. Dotychczasowa literatura koncentruje się głównie na identyfikacji samej relacji między cenami ropy a wzrostem gospodarczym lub na ocenie wpływu pojedynczych źródeł niepewności; rzadziej natomiast pokazuje, czy kraje UE tworzą odmienne grupy pod względem bezpośredniej wrażliwości na szoki naftowe oraz pośredniej ekspozycji na kanały geopolityczne i niepewności. Wypełnieniem tej luki jest zaproponowanie nowej typologii gospodarek UE opartej na czasowo-częstotliwościowych profilach zależności między wzrostem gospodarczym, cenami ropy i czynnikami niepewności.

Bibliografia

- Aloui, C., Hkiri, B., Hammoudeh, S., & Shahbaz, M. (2018). A Multiple and Partial Wavelet Analysis of the Oil Price, Inflation, Exchange Rate, and Economic Growth Nexus in Saudi Arabia. *Emerging Markets Finance and Trade*, 54(4), 935–956. <https://doi.org/10.1080/1540496X.2017.1423469>
- Chen, X. (2023). Are the shocks of EPU, VIX, and GPR indexes on the oil-stock nexus alike? A time-frequency analysis. *Applied Economics*, 55(48), 5637–5652. <https://doi.org/10.1080/00036846.2022.2140115>
- Jokubaitis, S., & Celov, D. (2023). Business Cycle Synchronization in the EU: A Regional-Sectoral Look through Soft-Clustering

Słowa kluczowe: koherencja falkowa, grupowanie gospodarek, dynamika relacji, ropa naftowa, wzrost gospodarczy

Hidden Patterns in the Oil-Economic Growth Nexus: A EU Economy's Look through Wavelet Decomposition and Clustering

Since the global financial crisis of 2008–2009, uncertainty has emerged as a central issue in research on the functioning of the economy, financial markets, and energy commodity markets. In recent years, its importance has further intensified due to overlapping economic, pandemic, and geopolitical shocks, including Russia's aggression against Ukraine, escalating tensions in the Middle East, and increases in oil and energy prices. These developments have demonstrated that the relationship between oil prices and economic activity cannot be analyzed in isolation from the broader environment of risk and uncertainty. It thus becomes particularly relevant to ask whether different sources of uncertainty affect the interdependence between the oil market and economic growth in a similar manner, or whether their impact depends on the time horizon and frequency. Consequently, the analysis of this relationship requires an approach that simultaneously captures variation over time and across different frequency scales (Chen, 2023).

The objective of this research is to identify and construct profiles of European Union economies according to the similarity of their economic growth responses to oil price shocks under conditions of elevated geopolitical risk and oil price uncertainty. The empirical analysis assumes that the impact

of oil shocks on economic growth is heterogeneous and depends on the time horizon, the phase of the business cycle, and the source of uncertainty. Existing studies indicate that the relationship between oil prices and economic activity varies across frequencies and may be particularly significant in the long run (Redin et al., 2018), while wavelet-based methods capture the variability of this relationship both over time and across frequencies (Aloui et al., 2018).

The study employs wavelet coherence between economic growth and oil prices for EU countries, followed by partial wavelet coherence to assess whether the observed relationship persists after accounting for geopolitical risk and oil price uncertainty. From the resulting coherence maps, synthetic indicators will be derived to describe the strength of the relationship in the short-, medium-, and long-term, its persistence during crisis periods, and the extent of its attenuation after controlling for uncertainty factors. These indicators will then be used to cluster EU economies, in line with an approach that groups based on similarities in dynamic behavior rather than solely on conventional structural characteristics of economies (Jokubaitis & Celov, 2023).

Despite the growing body of literature on the effects of oil prices, geopolitical risk, and uncertainty on economic activity, there remains a lack of analyses that classify EU economies according to the nature of their vulnerability to external shocks. The existing literature predominantly focuses on identifying the relationship between oil prices and economic growth or on assessing the impact of individual sources of uncertainty; less often does it examine whether EU countries form distinct groups in terms of their direct sensitivity to oil price shocks and indirect exposure to geopolitical and uncertainty channels. This study addresses this gap by proposing a new typology of EU economies based on time–frequency profiles of the relationships between economic growth, oil prices, and uncertainty factors.

Literature:

Aloui, C., Hkiri, B., Hammoudeh, S., & Shahbaz, M. (2018). A Multiple and Partial Wavelet Analysis of the Oil Price, Inflation, Exchange Rate, and Economic Growth Nexus in Saudi Arabia. *Emerging Markets Finance and Trade*, 54(4), 935–956. <https://doi.org/10.1080/1540496X.2017.1423469>

Chen, X. (2023). Are the shocks of EPU, VIX, and GPR indexes on the oil-stock nexus alike? A time-frequency analysis. *Applied Economics*, 55(48), 5637–5652. <https://doi.org/10.1080/00036846.2022.2140115>

Jokubaitis, S., & Celov, D. (2023). Business Cycle Synchronization in the EU: A Regional-Sectoral Look through Soft-Clus

Keywords: wavelet coherence, clustering, relation dynamics, crude oil, economic growth

Andrzej Geise, Uniwersytet Mikołaja Kopernika w Toruniu; Marta Dober, Uniwersytet Mikołaja Kopernika w Toruniu

Ryzyko geopolityczne i niepewność cen ropy naftowej w prognozowaniu inflacji w Polsce. Analiza porównawcza modeli ekonometrycznych i modeli uczenia maszynowego

The Role of Geopolitical Risk and Oil Price Uncertainty in Forecasting Inflation in Poland. The Comparative Analysis of Econometric and Machine Learning Models

Accurate forecasts of inflation are crucial for monetary policy, macroeconomic stability, and the formation of expectations by households, firms, and financial markets. In recent years, inflation

forecasting has become particularly challenging due to the overlapping effects of COVID-19, the war in Ukraine, the war in Iran, energy price shocks, and heightened geopolitical uncertainty (Alomani et al., 2025; Al-Thaqeb & Algharabali, 2019; Caldara et al., 2026; Joseph et al., 2024). These developments have weakened macroeconomic stability and heightened the need for forecasting approaches that capture both domestic inflation dynamics and external sources of uncertainty.

The aim of this study is to compare the accuracy of selected econometric models and machine learning methods in forecasting inflation in Poland. This research focuses on several questions. First, are there significant differences in inflation forecast accuracy between econometric models and machine learning models when the limited number of monthly observations constrains the analysis? Second, does the inclusion of variables reflecting geopolitical risk, economic policy uncertainty, and oil price uncertainty improve forecast accuracy? Third, do the predictive contributions of these variables differ between econometric and machine learning modeling frameworks?

The study uses both basic and advanced econometric models (such as ARIMA, SARIMA-GARCH, and ARDL-GARCH specifications) and machine learning models (such as XGBoost and Deep RNN). Forecasting performance is evaluated in an out-of-sample framework using standard forecast accuracy measures. The empirical design also compares model variants with and without exogenous predictors in order to assess the informational value of geopolitical risk and oil price uncertainty.

The contribution of the study is as follows. First, it provides evidence on the relative forecasting performance of time-series econometric models and machine learning methods for inflation in Poland. Second, it examines whether uncertainty-related variables improve inflation forecasts in a small open economy exposed to external geopolitical and energy market shocks.

Literature

- Alomani, G., Kayid, M., & Abd El-Aal, M. F. (2025). Global inflation forecasting and Uncertainty Assessment: Comparing ARIMA with advanced machine learning. *Journal of Radiation Research and Applied Sciences*, 18(2), 101402. <https://doi.org/10.1016/J.JRRAS.2025.101402>
- Al-Thaqeb, S. A., & Algharabali, B. G. (2019). Economic policy uncertainty: a literature review. *J. Econ. Asymmetries*, 20, e00133. <https://doi.org/10.1016/j.jeca.2019.e00133>
- Caldara, D., Conlisk, S., Iacoviello, M., & Penn, M. (2026). Do geopolitical risks raise or lower inflation? *Journal of International Economics*, 159, 104188. <https://doi.org/10.1016/J.JINTECO.2025.104188>
- Joseph, A., Potjagailo, G., Chakraborty, C., & Kapetanios, G. (2024). Forecasting UK inflation bottom up. *International Journal of Forecasting*, 40(4), 1521–1538. <https://doi.org/10.1016/J.IJFORECAST.2024.01.001>

Keywords: forecasting, inflation, time series, machine learning, risk, uncertainty

Romana Głowicka-Wołoszyn, Uniwersytet Przyrodniczy w Poznaniu; Andrzej Wołoszyn, Uniwersytet Przyrodniczy w Poznaniu

Sytuacja finansowa młodych gospodarstw domowych w Polsce w dobie kryzysów

Młode gospodarstwa domowe to fundament przyszłego rozwoju społecznego i gospodarczego kraju. Powszechny dostęp do edukacji, rozwój nowych technologii sprzyja lepszemu startowi młodych osób na rynku pracy i uzyskiwaniu wyższych wynagrodzeń mimo braku zawodowego doświadczenia. Często jednak brak własnego mieszkania i koszty związane z wynajmem lub spłatą kredytu mieszkaniowego

stają się znaczącym obciążeniem budżetów młodych gospodarstw domowych. Jednocześnie ograniczone oszczędności lub ich brak, przy gwałtownych zmianach sytuacji ekonomicznej spowodowanej nagłym wzrostem inflacji czy cenami energii i paliw, nie zapewniają odpowiedniego poziomu bezpieczeństwa finansowego. Kluczową kwestią jest więc diagnozowanie sytuacji finansowej młodych gospodarstw domowych obejmujące ich sytuację dochodową, wydatki, oszczędności, zadłużenie i gospodarowanie budżetem zarówno w wymiarze obiektywnym, jak i subiektywnym. Pozwala bowiem ocenić możliwości rozwojowe młodych gospodarstw domowych w Polsce, ich stabilność finansową i odporność na kryzysy ekonomiczne.

Celem badań była wielowymiarowa ocena sytuacji finansowej młodych (do 35 lat) gospodarstw domowych w Polsce w porównaniu do pozostałych klas wiekowych. Badania przeprowadzono w latach 2018-2024 w celu uwzględnienia kryzysów ekonomicznych wywołanych pandemią COVID-19, wybuchem wojny w Ukrainie, nadmiernym wzrostem inflacji czy cen energii i paliw.

Badania przeprowadzono na podstawie danych jednostkowych, nieidentyfikowalnych pochodzących z *Badania Budżetów Gospodarstw Domowych* przeprowadzonych przez Główny Urząd Statystyczny w Warszawie w latach 2018-2024. Gospodarstwa domowe podzielono na klasy wieku na podstawie wieku osoby referencyjnej tj. głowy gospodarstwa domowego. Dla każdej z wyodrębnionych klas wieku wyznaczono w badanych latach wartości ustalonych częściowych wskaźników ich sytuacji finansowej oraz wartości wskaźnika syntetycznego. Do konstrukcji syntetycznego wskaźnika sytuacji finansowej gospodarstw domowych zastosowano wielowymiarową metodę taksonomii relatywnej. Metoda ta polega na konstrukcji wskaźnika syntetycznego w oparciu o zrelatywizowane względem pozostałych grup wiekowych wartości wskaźników częściowych. Uzyskane wartości wskaźnika syntetycznego prezentowały syntetyczną ocenę sytuacji finansowej poszczególnych grup wiekowych gospodarstw w relacji do oceny tego zjawiska we wszystkich pozostałych grupach łącznie. Tę relatywną ocenę można też określić jako ocenę dysproporcji w sytuacji finansowej danej grupy gospodarstw domowych w odniesieniu do pozostałych. Uzyskane wartości wskaźnika syntetycznego w każdym roku pozwoliły na analizę procesu niwelowania lub narastania różnic pomiędzy grupami wiekowymi w zakresie sytuacji finansów gospodarstw domowych, a także pozwoliły na ocenę odporności sytuacji finansowej gospodarstw domowych na pojawiające się kryzysy ekonomiczne.

Na podstawie przeprowadzonych badań oceniono, że sytuacja finansowa młodych gospodarstw domowych (do 35 lat) w latach 2018-2024 była relatywnie nieco lepsza niż wszystkich pozostałych klas wiekowych łącznie. Przyczyniła się do tego wyższa relatywna ocena dochodów, wydatków, stopy oszczędności i subiektywna ocena gospodarowania budżetem. Relatywnie gorsze wartości odnotowano jedynie w przypadku wskaźnika zadłużenia. Wśród młodych gospodarstw domowych odnotowano drugi najwyższy wskaźnik zadłużenia, ale jednocześnie w latach dobrej sytuacji dochodowej – najwyższy odsetek oszczędzających regularnie. Młode gospodarstwa domowe cechowała najmniejsza stabilność dochodu rozporządzalnego i najbardziej dynamiczne jego zmiany w badanym okresie: największy spadek po wybuchu pandemii i największy wzrost po 2023 roku. Zmiany te przyczyniły się do spadku relatywnego wskaźnika sytuacji finansowej młodych gospodarstw domowych w latach 2019-2022.

Słowa kluczowe: młode gospodarstwa domowe, sytuacja finansowa, relatywny wskaźnik syntetyczny, taksonomia relatywna

Financial Situation of Young Households in Poland in Times of Crisis

Young households constitute an important foundation for a country's future social and economic development. Widespread access to education and the development of new technologies facilitate a better start for young people in the labor market and enable them to earn higher wages despite their limited professional experience. However, the lack of owner-occupied housing and the costs associated with renting or repaying a mortgage often place a substantial burden on young household budgets. At the same time, limited savings or the absence of savings, combined with rapid changes in economic conditions caused by sudden increases in inflation or energy and fuel prices, may undermine an adequate level of financial security. It is therefore essential to assess the financial situation of young households by considering their income, expenditure, savings, indebtedness, and household budget management from both objective and subjective perspectives. Such an assessment makes it possible to evaluate the development potential of young households in Poland, as well as their financial stability and resilience to economic crises.

The aim of the study was to provide a multidimensional assessment of the financial situation of young households in Poland, defined as households with a reference person aged up to 35, in comparison with other age groups. The analysis covered the period 2018–2024 to account for the economic crises associated with the COVID-19 pandemic, the outbreak of the war in Ukraine, the sharp rise in inflation, and increases in energy and fuel prices.

The study was based on non-identifiable microdata from the Household Budget Survey conducted by Statistics Poland in Warsaw between 2018 and 2024. Households were divided into age groups according to the age of the reference person, i.e. the head of the household. For each age group and each year under study, values of selected partial indicators of the household financial situation and a synthetic indicator were calculated. The synthetic indicator of household financial situation was constructed using the multidimensional method of relative taxonomy. This method involves constructing a synthetic indicator based on partial indicators expressed relative to the corresponding values observed in the other age groups. The resulting values of the synthetic indicator provided an overall assessment of the financial situation of each household age group relative to all other groups considered jointly. This relative assessment can also be interpreted as a measure of disparities in the financial situation of a given household group compared with the remaining groups. The annual values of the synthetic indicator made it possible to analyze whether differences in the financial situation of households across age groups were narrowing or widening over time and to assess the resilience of household finances to emerging economic crises.

The results indicate that the financial situation of young households aged up to 35 was, in relative terms, slightly better in 2018–2024 than that of all other age groups considered jointly. This was primarily attributable to relatively more favorable assessments of income, expenditure, the savings rate, and the subjective evaluation of household budget management. Relatively less favorable results were recorded only for the indebtedness indicator. Young households had the second-highest level of indebtedness, but at the same time, in years characterized by favorable income conditions, they also recorded the highest proportion of households saving regularly. Young households were characterized by the lowest stability of disposable income and the most dynamic changes in this variable over the study period: they experienced the largest decline following the outbreak of the pandemic and the strongest increase after 2023. These changes contributed to a decline in the relative financial situation indicator for young households between 2019 and 2022.

Keywords: young households, financial situation, relative synthetic indicator, relative taxonomy

Wioletta Grzenda, Szkoła Główna Handlowa w Warszawie; Olga Momot, Szkoła Główna Handlowa w Warszawie

Ocena stabilności rankingów ważności zmiennych w modelach drzewiastych na przykładzie inkluzji finansowej w późniejszym okresie życia

Cyfryzacja usług bankowych może tworzyć nowe bariery dla inkluzji finansowej w późniejszym okresie życia, szczególnie poprzez mechanizmy wykluczenia cyfrowego. Wykorzystując dane z badania Survey of Health, Ageing and Retirement in Europe (SHARE), analizujemy, w jaki sposób wskaźniki wykluczenia cyfrowego, w szczególności korzystanie z Internetu oraz umiejętności komputerowe, są powiązane z posiadaniem rachunku bankowego wśród osób w wieku 65 lat i więcej. Ponadto w analizie inkluzji finansowej jako czynniki uzupełniające uwzględniamy cechy demograficzne, społeczno-ekonomiczne, zdrowotne oraz subiektywny dobrostan.

W celu uchwycenia złożonych i potencjalnie nieliniowych zależności zastosowano trzy modele uczenia maszynowego oparte na drzewach: drzewa decyzyjne, lasy losowe oraz XGBoost. Ważność zmiennych dla poszczególnych modeli oceniono przy użyciu metod Mean Decrease in Impurity (MDI) oraz SHapley Additive exPlanations (SHAP). Chociaż modele te osiągnęły zbliżoną trafność predykcyjną, jedynie korzystanie z Internetu w niedawnym okresie oraz całkowity dochód konsekwentnie należą do najważniejszych zmiennych we wszystkich metodach, podczas gdy pozycje pozostałych czynników w rankingach znacznie się różnią.

W odpowiedzi na tę niestabilność proponujemy nowe podejście do oceny stabilności rankingów ważności zmiennych w różnych modelach. Wnioski merytoryczne formułujemy na podstawie wyników modelu charakteryzującego się największą stabilnością w identyfikacji kluczowych czynników.

Słowa kluczowe: Wykluczenie cyfrowe, inkluzja finansowa, SHARE, ważność zmiennych, drzewa decyzyjne, XGBoost, lasy losowe, MDI, SHAP

Assessing the Stability of Feature Importance Rankings in Tree-Based Models: Evidence from Financial Inclusion in Later Life

The digitalisation of banking services may create new barriers to financial inclusion in later life, particularly through mechanisms of digital exclusion. Using data from the Survey of Health, Ageing and Retirement in Europe (SHARE), this study examines how indicators of digital exclusion, specifically Internet use and computer skills, are associated with bank account ownership among individuals aged 65 and older. We further consider demographic, socioeconomic, health, and subjective well-being characteristics as complementary covariates in analysing financial inclusion.

To capture complex and potentially non-linear relationships, we employ three tree-based machine learning models: decision trees, random forests, and extreme gradient boosting. Feature importance is assessed using Mean Decrease in Impurity (MDI) and SHapley Additive exPlanations (SHAP), producing rankings of variable importance across models. While the models achieve similar predictive performance, only recent Internet use and total income consistently emerge among the most

important variables across methods, whereas the ranking positions of other variables vary substantially.

To address this instability, we introduce a novel approach to assessing the robustness of feature importance rankings across models. Substantive conclusions are then drawn based on the model exhibiting the highest explanatory stability.

Keywords: Digital exclusion, financial inclusion, SHARE, feature importance, decision trees, XGBoost, random forests, MDI, SHAP

Klaudia Jarno, Yeldbird Sp. z o.o.; Piotr Niedziela, Yeldbird Sp. z o.o.; Aleksandra Wilczewska, Yeldbird Sp. z o.o.; Łukasz Sidor, Yeldbird Sp. z o.o.; Witold Abramowicz, Yeldbird Sp. z o.o., Uniwersytet Ekonomiczny w Poznaniu

Ocena wpływu przekazywania cen minimalnych na różnych etapach procesu aukcjonowania w reklamie programatycznej z wykorzystaniem eksperymentu A/B i porównania modeli zagnieżdżonych

Wydawcy cyfrowi stosują ceny minimalne (ang. floor price) w aukcjach programatycznych w systemie Google Ad Manager (dalej: GAM). Ceny te mogą być dodatkowo przekazywane uczestnikom aukcji, odbywających się na wcześniejszym etapie zbierania popytu, w tzw. header biddingu, np. poprzez odpowiedni moduł narzędzia o nazwie Prebid. Celem badania była weryfikacja hipotezy, że przekazanie ceny minimalnej do Prebida zwiększa przychód programatyczny (wyrażony wskaźnikiem rCPM), oraz hipotez wtórnych dotyczących liczby ofert (ang. bids) i współczynnika wygranych (ang. win rate). W eksperymencie A/B wybrano sześć jednostek reklamowych i podzielono ruch na dziesięć równych prób. W pięciu próbach cenę minimalną ustawiano wyłącznie w GAM, a w pozostałych pięciu identyczne poziomy cen przekazywano dodatkowo do Prebida. Dla każdej jednostki zastosowano pięć poziomów cen: od zera do wartości dostosowanej do danej powierzchni reklamowej.

Najpierw przeprowadzono analizę wizualną zależności między ceną minimalną a rCPM, liczbą ofert i współczynnikiem wygranych w obu warunkach eksperymentalnych. Następnie dla każdej jednostki dopasowano dwa modele regresji wielomianowej: bazowy, opisujący zależność wyniku od ceny minimalnej, oraz rozszerzony o zmienną zero-jedynkową oznaczającą jednoczesne przekazywanie cen minimalnych w Prebidzie i GAM. Modele porównano testem F dla modeli zagnieżdżonych. Analiza wizualna wskazała brak systematycznego wpływu przekazywania cen do Prebida na relację między ceną minimalną a rCPM, natomiast wykazała różnice w relacjach między ceną minimalną a liczbą ofert oraz między ceną minimalną a współczynnikiem wygranych, przy czym wzorec tych różnic zmieniał się między jednostkami. Testy F potwierdziły brak istotnych różnic w dopasowaniu modeli dla rCPM we wszystkich sześciu jednostkach. Dla liczby ofert różnice były istotne statystycznie w dwóch jednostkach. Podobne wyniki uzyskano dla współczynnika wygranych.

Wnioski wskazują, że przekazywanie cen minimalnych do Prebida nie zwiększa skuteczności monetyzacji mierzonej rCPM, lecz może zmieniać zachowanie licytantów w aukcji. Badanie demonstruje zastosowanie porównania modeli zagnieżdżonych w eksperymencie ekonomicznym z wieloma poziomami ceny minimalnej.

Słowa kluczowe: testy statystyczne, eksperyment A/B, aukcje programatyczne, ceny minimalne

Assessment of the impact of passing floor prices at different stages of the programmatic advertising auction process using an A/B experiment and nested model comparison

Digital publishers use floor prices in programmatic auctions in Google Ad Manager (hereafter: GAM). These prices may additionally be passed to auctions at an earlier stage of demand collection, known as header bidding, for example through an appropriate module of a tool called Prebid. The aim of this study was to verify the hypothesis that passing the floor price to Prebid increases programmatic revenue, expressed as rCPM, as well as secondary hypotheses concerning the number of bids and the win rate. In the A/B experiment, six ad units were selected and traffic was divided into ten equal samples. In five samples, the floor price was set only in GAM, while in the remaining five samples, identical price levels were additionally passed to Prebid. For each unit, five price levels were applied, ranging from zero to a value adjusted to the given ad inventory.

First, a visual analysis was conducted of the relationship between the floor price and rCPM, the number of bids, and the win rate under both experimental conditions. Next, two polynomial regression models were fitted for each unit: a baseline model describing the relationship between the outcome and the floor price, and an extended model including a binary variable indicating the simultaneous passing of floor prices in both Prebid and GAM. The models were compared using an F-test for nested models. The visual analysis indicated no systematic impact of passing prices to Prebid on the relationship between the floor price and rCPM. However, it revealed differences in the relationships between the floor price and the number of bids, as well as between the floor price and the win rate, with the pattern of these differences varying across units. The F-tests confirmed the absence of significant differences in model fit for rCPM across all six units. For the number of bids, the differences were statistically significant in two units. Similar results were obtained for the win rate.

The conclusions indicate that passing floor prices to Prebid does not improve monetization effectiveness as measured by rCPM, but it may change bidder behavior in the auction. The study demonstrates the use of nested model comparison in an economic experiment with multiple floor price levels.

Keywords: statistical tests, A/B experiment, programmatic auctions, floor prices

Przemysław Jaśko, Uniwersytet Ekonomiczny w Krakowie; Jerzy Rydlewski, Akademia Górniczo-Hutnicza w Krakowie

Metoda Parameter Cascade w estymacji parametrów strukturalnych układów zwyczajnych równań różniczkowych - własności, symulacje oraz zastosowanie empiryczne w modelu Goodwina

W pracy zastosowano metodę Parameter Cascade (kaskady parametrów) w estymacji parametrów modelu Goodwina, zadanego jako układ zwyczajnych równań różniczkowych z audytywnym błędem pomiarowym, który to model opisuje dynamikę współzależności udziału wynagrodzeń w PKB oraz stopy zatrudnienia.

W analizie podkreśla się zarówno aparat metodologiczny jak i ograniczenia natury praktycznej. W szczególności porównuje się działanie metody estymacji Parameter Cascade z wynikami estymacji metodą nieliniowych najmniejszych kwadratów (NLS). Wyniki symulacyjne wskazują, że w skończonej próbie, obydwa estymatory są obciążone, ponadto metoda nieliniowych najmniejszych kwadratów

może dawać oszacowania, które znacznie różnią się od prawdziwych wartości parametrów i generują odstające trajektorie.

Zastosowanie empiryczne bazuje na danych dotyczących udziału wynagrodzeń w PKB oraz stopy zatrudnienia. Model zdaje się nie w pełni odwzorowywać zaobserwowaną dynamikę, co może być związane z ograniczoną dostępnością danych w ujęciu czasowym, bądź też może wskazywać na konieczność rozszerzenia bądź respecyfikacji układu zwyczajnych równań różniczkowych dla modelu Goodwina.

Słowa kluczowe: model Goodwina, zwyczajne równania różniczkowe w ekonomii, analiza danych funkcjonalnych, metody kolokacji, metoda Parameter Cascade

Parameter Cascade Estimation Method for ODE System Structural Parameters: Properties, Simulations, and Empirical Application to the Goodwin Model

The parameter cascade method is employed to estimate the parameters of the Goodwin model, formulated as a system of ordinary differential equations with measurement error, capturing the dynamic interaction between the labour income share and the employment rate.

The analysis highlights both methodological advances and practical limitations. In particular, the performance of the parameter cascade approach is benchmarked against the nonlinear least squares (NLS) estimator. Simulation results indicate that, in finite samples, both estimators exhibit bias; moreover, the NLS method may yield parameter estimates that deviate substantially from the true values and generate implausible trajectories.

An empirical application based on data for labour income share and employment rate is also conducted. However, the model does not fully reproduce realistic dynamics, which may reflect constraints related to data availability or point to the necessity of respecifying or extending the underlying ODE system of the Goodwin model.

Keywords: Goodwin model, ODE systems in economics, functional data analysis, collocation methods, parameter cascade method

Stanisław Jaworski, Szkoła Główna Gospodarstwa Wiejskiego w Warszawie; Wojciech Zieliński, Szkoła Główna Gospodarstwa Wiejskiego w Warszawie

Odporność na kłamstwo jednostronnego przedziału ufności dla odsetka pytań drażliwych

W badaniach ekonomiczno-społecznych badacz jest zainteresowany oszacowaniem odsetka pozytywnych odpowiedzi na pytanie drażliwe. Pytanie drażliwe jest takim pytaniem, na które respondent może nie chcieć udzielić rzetelnej odpowiedzi. W celu uzyskania wiarygodnych informacji, w badaniach ankietowych zadawane jest dodatkowo pytanie neutralne w taki sposób by zagwarantować osobie ankietowanej jak najlepszą ochronę prywatności, czyli by jak najbardziej utrudnić ankietownikowi odgadnięcie rzeczywistej odpowiedzi na pytanie drażliwe. Niestety mimo takich zabezpieczeń często osoba ankietowana kłamie. W pracy zaproponowany zostanie sposób modelowania zjawiska oszukiwania oraz zostanie zbadana odporność przedziału ufności dla odsetka

Słowa kluczowe: pytania drażliwe, przedział ufności, odporność

Robustness against lying of one-sided confidence interval for a fraction of sensitive question

Keywords: sensitive questions, confidence interval, robustness

Bartłomiej Jefmański, Uniwersytet Ekonomiczny we Wrocławiu

Miary oceny jakości procedur porządkowania liniowego obiektów opisanych intuicjonistycznymi zbiorami rozmytymi

Celem referatu jest przedstawienie modyfikacji miar oceny jakości procedur porządkowania liniowego dla obiektów opisanych intuicjonistycznymi zbiorami rozmytymi w zastosowaniu do danych pochodzących z badań ankietowych. Dane ankietowe reprezentowane są w postaci intuicjonistycznych zbiorów rozmytych, gdzie stopień przynależności, nieprzynależności i niepewności odpowiadają udziałom odpowiedzi pozytywnych, negatywnych i neutralnych z porządkowych skal pomiaru typu Likerta.

Użyteczność zaproponowanych miar zilustrowano na podstawie badania symulacyjnego obejmującego 10000 replikacji z wykorzystaniem rozkładu Dirichleta. Zdefiniowano 9 scenariuszy generowania danych odpowiadających wybranym wzorcom odpowiedzi ankietowych. Scenariusze różniły się między sobą poziomem stopni niepewności oraz relacją między stopniami przynależności i nieprzynależności. Obejmowały przypadki z wyraźną przewagą ocen pozytywnych i niskim stopniem niepewności, neutralne z symetrycznym rozkładem ocen oraz niepewne charakteryzujące się wysokim stopniem niepewności respondentów. Jako przykład procedury porządkowania liniowego zastosowano Intuicjonistyczną Rozmytą Miarę Syntetyczną.

Słowa kluczowe: Porządkowanie liniowe, intuicjonistyczne zbiory rozmyte, miary oceny jakości, Intuicjonistyczna Rozmyta Miara Syntetyczna

Quality Assessment Measures for Linear Ordering Procedures of Objects Described by Intuitionistic Fuzzy Sets

The aim of the paper is to present modifications of quality assessment measures for linear ordering procedures applied to objects described by intuitionistic fuzzy sets in the context of survey data. Survey data are represented in the form of intuitionistic fuzzy sets, where the degrees of membership, non-membership and hesitancy correspond to the proportions of positive, negative and neutral responses obtained from ordinal measurement scales of the Likert type.

The usefulness of the proposed measures is illustrated on the basis of a simulation study comprising 10000 replications using the Dirichlet distribution. Nine data generation scenarios were defined, corresponding to selected survey response patterns. The scenarios differed in terms of the level of hesitancy and the relationship between the degrees of membership and non-membership. They covered cases with a clear predominance of positive assessments and low hesitancy, neutral cases with a symmetric distribution of responses, and uncertain cases characterized by a high degree

of respondent hesitancy. The Intuitionistic Fuzzy Synthetic Measure was used as an example of a linear ordering procedure.

Keywords: Linear ordering, intuitionistic fuzzy sets, quality assessment measures, Intuitionistic Fuzzy Synthetic Measure

Małgorzata Just, Uniwersytet Przyrodniczy w Poznaniu, Katedra Finansów i Rachunkowości; Krzysztof Echaust, Uniwersytet Ekonomiczny w Poznaniu, Katedra Badań Operacyjnych i Ekonomii Matematycznej

Reakcje światowych rynków akcji na konflikty geopolityczne: analiza wpływu wojny rosyjsko-ukraińskiej oraz konfliktu Stanów Zjednoczonych i Izraela z Iranem

Celem badania była analiza reakcji światowych rynków akcji na wstrząsy geopolityczne: wojnę rosyjsko-ukraińską oraz konflikt Stanów Zjednoczonych i Izraela z Iranem. Wykorzystując metodę analizy zdarzeń oraz modele regresji przekrojowej, zbadano reakcje 77 międzynarodowych rynków akcji, analizując skumulowane nadzwyczajne stopy zwrotu oraz czynniki determinujące ich zróżnicowanie. Analizę przeprowadzono z uwzględnieniem poziomu rozwoju gospodarczego, pozycji w handlu energią oraz przynależności regionalnej.

Wyniki wskazują, że oba konflikty wywołały istotne statystycznie i negatywne reakcje światowych rynków akcji, jednak ich skala i dynamika różniły się między krajami. W przypadku wojny rosyjsko-ukraińskiej średnia nadzwyczajna stopa zwrotu w dniu zdarzenia wyniosła -2,68%, a średnia skumulowana nadzwyczajna stopa zwrotu w analizowanym oknie zdarzenia (24.02–08.03.2022) wyniosła około -5%. Najsilniejsze spadki odnotowano w Europie oraz w krajach będących importerami netto energii, podczas gdy eksporterzy energii wykazywali większą odporność, a w niektórych przypadkach dodatkowo nadzwyczajne stopy zwrotu. Konflikt Stanów Zjednoczonych i Izraela z Iranem wywołał podobną, choć bardziej stopniową reakcję, prowadząc do średniej skumulowanej nadzwyczajnej stopy zwrotu na poziomie około -5,7% (28.02–20.03.2026). Największe negatywne efekty wystąpiły w regionie Azji i Pacyfiku oraz Europy, szczególnie w krajach zależnych od importu energii.

Analiza regresji potwierdziła, że pozycja kraju w handlu energią, położenie geograficzne oraz poziom rozwoju gospodarczego istotnie wpływają na skalę reakcji rynków. Porównanie obu konfliktów wskazuje na podobny kierunek oddziaływania, lecz różny przebieg czasowy i intensywność efektów. W przypadku wojny rosyjsko-ukraińskiej większe znaczenie miały czynniki regionalne, natomiast w konflikcie Stanów Zjednoczonych i Izraela z Iranem większą rolę mogło odgrywać ryzyko zakłóceń na światowym rynku energii.

Słowa kluczowe: analiza zdarzeń, nadzwyczajna stopa zwrotu, rynki akcji, wojna rosyjsko-ukraińska, konflikt na Bliskim Wschodzie

Global Stock Market Reactions to Geopolitical Conflicts: An Analysis of the Impact of the Russia–Ukraine War and the United States–Israel–Iran Conflict

The aim of this study was to analyze the reactions of global stock markets to geopolitical shocks: the Russia–Ukraine war and the United States–Israel–Iran conflict in 2022. Using an event study methodology and cross-sectional regression models, the study examined the responses of 77 international stock markets by analyzing cumulative abnormal returns and the factors explaining their variation. The analysis accounted for differences in the level of economic development, countries' positions in international energy trade, and regional location.

The results indicate that both conflicts triggered statistically significant negative reactions in global stock markets; however, the magnitude and dynamics of these reactions varied across countries. In the case of the Russia–Ukraine war, the average abnormal return on the event day was -2.68% , while the average cumulative abnormal return over the event window (24 February–8 March 2022) was approximately -5% . The largest declines were observed in Europe and in net energy-importing countries, whereas energy exporters showed greater resilience and, in some cases, positive abnormal returns. The United States–Israel–Iran conflict triggered a similar, although more gradual, reaction, resulting in an average cumulative abnormal return of approximately -5.7% (28 February–20 March 2022). The strongest negative effects occurred in the Asia-Pacific region and Europe, particularly in countries dependent on energy imports.

The regression analysis confirmed that a country's position in international energy trade, geographical location, and level of economic development significantly influenced the magnitude of market reactions. A comparison of the two conflicts indicates a similar direction of impact but differences in the timing and intensity of the effects. In the case of the Russia–Ukraine war, regional factors were more important, whereas in the United States–Israel–Iran conflict, the risk of disruptions in global energy markets may have played a greater role.

Keywords: event study, abnormal return, stock markets, Russia–Ukraine war, Middle East conflict

Oleksij Kelebaj, Uniwersytet Ekonomiczny w Krakowie; Anna Połec, Uniwersytet Jagielloński w Krakowie, Szkoła Doktorska Nauk Społecznych; Daniela Pietraszek, Uniwersytet Jagielloński w Krakowie

Zróżnicowanie sytuacji na rynku pracy a przepływy migracyjne w Polsce i Niemczech – analiza przestrzenna i taksonomiczna

Celem wystąpienia jest analiza zależności pomiędzy zróżnicowaniem sytuacji na rynku pracy a przepływami migracyjnymi w Polsce i Niemczech z wykorzystaniem metod analizy przestrzennej oraz taksonomicznej. W badaniu uwzględniono regionalne zróżnicowanie wybranych wskaźników rynku pracy, takich jak stopa bezrobocia, poziom zatrudnienia, PKB czy kierunki migracji wewnętrznych i zagranicznych.

Zastosowanie metod taksonomicznych umożliwi wyodrębnienie grup regionów o podobnej charakterystyce społeczno-ekonomicznej, natomiast narzędzia analizy przestrzennej pozwolą na identyfikację wzorców przestrzennych oraz zależności pomiędzy sąsiadującymi jednostkami terytorialnymi. Porównanie Polski i Niemiec pozwoli wskazać podobieństwa i różnice w mechanizmach

kształtujących mobilność ludności w państwach o odmiennym poziomie rozwoju gospodarczego i strukturze rynku pracy.

Wyniki badania mogą przyczynić się do lepszego zrozumienia przestrzennych uwarunkowań migracji oraz dostarczyć istotnych wniosków dla polityki regionalnej i rynku pracy. Przeprowadzona analiza pozwoli również ocenić, w jakim stopniu zróżnicowanie sytuacji ekonomicznej regionów wpływa na decyzje migracyjne mieszkańców oraz wskazać obszary wymagające szczególnego wsparcia w zakresie rozwoju społeczno-gospodarczego.

Słowa kluczowe: migracje, rynek pracy, bezrobocie, wzrost gospodarczy

Spatial and Taxonomic Analysis of Labour Market Disparities and Migration Flows in Poland and Germany

The aim of this presentation is to examine the relationship between labour market disparities and migration flows in Poland and Germany using spatial and taxonomic analysis methods. The study considers regional differences in selected labour market indicators, including the unemployment rate, employment level, gross domestic product (GDP), and the directions of internal and international migration.

The application of taxonomic methods will enable the identification of groups of regions with similar socio-economic characteristics, while spatial analysis tools will facilitate the detection of spatial patterns and relationships between neighbouring territorial units. A comparative analysis of Poland and Germany will make it possible to identify similarities and differences in the mechanisms shaping population mobility in countries with different levels of economic development and labour market structures.

The findings are expected to contribute to a better understanding of the spatial determinants of migration and provide valuable insights for regional and labour market policy. Furthermore, the analysis will assess the extent to which regional economic disparities influence migration decisions and identify areas requiring targeted support for socio-economic development.

Keywords: Migration, Labour Market, Unemployment, Economic Growth

Błażej Kochański, Politechnika Gdańska

Wszystkie twarze AUC: wspólne podstawy miar stosowanych w statystyce, uczeniu maszynowym, medycynie i bankowości

Pole pod krzywą ROC (AUC, Area Under the Curve) jest jedną z najczęściej wykorzystywanych miar jakości modeli klasyfikacyjnych. Celem referatu jest przedstawienie AUC jako uniwersalnej miary asocjacji między zmienną ilościową lub porządkową a zmienną dychotomiczną, której znaczenie wykracza daleko poza klasyczną analizę krzywych ROC.

W referacie zostaną przedstawione interpretacje AUC oparte na krzywej ROC, na porównywaniu parami (probabilistyczna) oraz na rangach. Omówione zostaną związki tej miary z innymi pojęciami stosowanymi w różnych dyscyplinach naukowych, w szczególności z testami nieparametrycznymi

Manna-Whitneya i Brunnera-Munzela, współczynnikiem Giniego i Accuracy Ratio wykorzystywanymi w analizie ryzyka kredytowego, miarą D Somersa, współczynnikiem korelacji rangowo-dwuseryjnej Glassa, deltą Cliffa oraz innymi miarami opartymi na porównaniach rang i prawdopodobieństwie przewagi.

Celem wystąpienia jest pokazanie, że pomimo różnorodności terminologii stosowanej w medycynie, ekonomii, statystyce i uczeniu maszynowym wiele powszechnie wykorzystywanych miar opiera się na tej samej matematycznej idei. Przedstawienie tych zależności pozwala lepiej zrozumieć interpretację AUC oraz jego rolę jako uniwersalnego narzędzia oceny jakości modeli predykcyjnych i siły zależności między zmiennymi w różnych obszarach zastosowań.

Słowa kluczowe: AUC, krzywa ROC, miary asocjacji, modele klasyfikacyjne

Many Faces of AUC: Common Foundations for Measures Used in Statistics, Machine Learning, Medicine, and Banking

The area under the ROC curve (AUC) is one of the most commonly used measures of classification model quality. The aim of this paper is to present AUC as a universal measure of association between a quantitative or ordinal variable and a dichotomous variable, whose significance extends far beyond classical ROC curve analysis.

This paper will present interpretations of AUC based on the ROC curve, pairwise comparison (probabilistic), and rank-based analysis. We will discuss the relationship of this measure to other concepts used in various scientific disciplines, in particular the nonparametric tests Mann-Whitney and Brunner-Munzel, the Gini coefficient and Accuracy Ratio used in credit risk analysis, Somers' D, Glass's rank-biserial correlation coefficient, Cliff's delta, and other measures based on rank comparisons and probability of superiority.

The aim of this presentation is to demonstrate that despite the diversity of terminology used in medicine, economics, statistics, and machine learning, many commonly used measures are based on the same mathematical idea. Presenting these relationships allows for a better understanding of the interpretation of AUC and its role as a universal tool for assessing the quality of predictive models and the strength of relationships between variables in various application areas.

Keywords: AUC, ROC curve, measures of association, classification models

Grzegorz Kończak, Uniwersytet Ekonomiczny w Katowicach

Dyskretność rozkładu statystyki testu Davida-Hellwiga i propozycja jej uciąglenia

Test Davida-Hellwiga jest nieparametryczną metodą weryfikacji hipotezy zgodności rozkładu empirycznego z zadaniem rozkładem teoretycznym. Statystyka testowa ma rozkład dyskretny, co prowadzi do konserwatywnych wniosków oraz obniżonej mocy testu, zwłaszcza przy małych liczebnościach prób. Dyskretność rozkładu wynika z kombinatorycznego charakteru statystyki, opartej na zliczaniu liczby pustych cel.

W artykule przedstawiono dyskretny charakter statystyki testowej Davida-Hellwiga oraz wynikające z niego implikacje dla testowania hipotez. Analizowany był wpływ rozkładu prawdopodobieństwa statystyki na rozmiar testu i jego moc. Na podstawie analiz teoretycznych i badań symulacyjnych określono skalę konserwatyzmu związanego z dyskretnością statystyki, dla różnych rozkładów teoretycznych i liczebności prób. W celu przezwyciężenia tych ograniczeń zaproponowano metodę korekty ciągłości ECA (empty cells area), stanowiącą ciągły odpowiednik dyskretniej statystyki testu Davida-Hellwiga. Przedstawiono teoretyczne własności proponowanej statystyki oraz wykazano jej skuteczność za pomocą symulacji Monte Carlo dla różnych rozkładów teoretycznych.

Przeprowadzone analizy symulacyjne potwierdzają skuteczność proponowanej modyfikacji statystyki testowej. Wyniki wskazują, że skorygowana postać statystyki istotnie poprawia rozmiar i moc testu, zwłaszcza dla prób o niewielkich liczebnościach, przy jednoczesnym zachowaniu nieparametrycznego charakteru oraz właściwości odpornościowych oryginalnego testu.

Słowa kluczowe: testy nieparametryczne, rozkład dyskretny, poprawka na ciągłość, test Davida-Hellwiga, moc testu

Discreteness of the David-Hellwig Test Statistic and a Proposal for Its Continuity Correction

The David-Hellwig test is a nonparametric method for assessing the goodness-of-fit of an empirical distribution to a specified theoretical distribution. Its test statistic has a discrete distribution, which leads to conservative inference and reduced size, especially for small samples. This discreteness arises from the combinatorial nature of the statistic, based on counting the number of empty cells.

This paper examines the discrete nature of the David-Hellwig test statistic and its implications for hypothesis testing. We analyze how the probability distribution of the statistic affects the size and power of the test. Using theoretical arguments and simulation studies, we quantify the degree of conservativeness induced by discreteness for various theoretical distributions and sample sizes. To mitigate these limitations, we propose a continuity-corrected statistic, ECA (empty cells area), which provides a continuous analogue of the original discrete David-Hellwig test statistic. We present the theoretical properties of the proposed statistic and demonstrate its performance by Monte Carlo simulations for different theoretical distributions.

The simulation results confirm the effectiveness of the proposed modification of the test statistic. The findings indicate that the corrected statistic improves the test size and power, particularly for small samples, while preserving the nonparametric nature and robustness properties of the original David-Hellwig test.

Keywords: non-parametric tests, discrete distributions, continuity correction, David-Hellwig test, test power

Jerzy Korzeniewski, Uniwersytet Łódzki

Ocena pakietów języka R do badania wydźwięku tekstu angielskojęzycznego

W środowisku języka R jest dostępnych kilka pakietów służących do ustalania wydźwięku tekstu angielskojęzycznego (meanr, syuzhet, sentiment, SentimentAnalysis, Rsentiment). W referacie

przedstawiona jest ocena tych pakietów pod kątem jakości ustalania wydźwięku opinii o obiekcie. Ocena ma charakter empiryczny i została przeprowadzona na standardowym zbiorze booking.com opinii o obiektach hotelowych. Ponadto, została przedstawiona własna propozycja modyfikacji zbioru opinii oraz algorytmu ustalania wydźwięku opinii o obiekcie. Zasadniczą ideą algorytmu jest badanie bigramów słów występujących w sąsiedztwie obiektu po uprzedniej adaptacji tekstu polegającej na wyszukiwaniu słów istotnych. Zaproponowanego algorytm nie wymaga stosowania stoplisty ani lematyzacji.

Słowa kluczowe: text mining, sentyment obiektu, pakiety języka R

An evaluation of R packages for analyzing the sentiment of English-language text

Several packages are available in the R environment for determining the sentiment of English-language text (meanr, syuzhet, sentiment, SentimentAnalysis, Rsentiment). This paper presents an evaluation of these packages in terms of the quality of their sentiment analysis of opinions about an object i.e. entity sentiment. The evaluation is empirical in nature and was conducted on a standard dataset of booking.com reviews of hotel businesses. In addition, we present our own proposal for modifying the review dataset and the algorithm used to determine the entity sentiment. The fundamental idea behind the algorithm is to analyze bigrams of words appearing in the vicinity of an object after preprocessing the text to identify relevant words. The proposed algorithm does not require the use of a stoplist or lemmatization.

Keywords: text mining, entity sentiment, R packages

Jadwiga Kostrzewska, Uniwersytet Ekonomiczny w Krakowie; Maciej Kostrzewski, Uniwersytet Ekonomiczny w Krakowie

Ujemne ceny na wybranych europejskich rynkach energii elektrycznej

Proces liberalizacji rynków energii elektrycznej doprowadził do uwolnienia cen oraz wzmocnienia roli mechanizmów rynkowych w bilansowaniu podaży i popytu. Rezultatem transformacji energetycznej jest rosnący udział odnawialnych źródeł energii w miksie energetycznym, do czego przyczyniło się również wprowadzenie mechanizmu merit order. Ze względu na niskie koszty krańcowe produkcji energii z wiatru i słońca wzrost energii wyprodukowanej ze źródeł odnawialnych prowadzi, w określonych warunkach, do spadku cen hurtowych, a niekiedy także do pojawienia się ujemnych cen. Zjawisko to jest obecnie przedmiotem rosnącej uwagi regulatorów i uczestników rynku, zwłaszcza w kontekście rosnącej liczby przypadków ujemnych cen obserwowanych na europejskich rynkach energii elektrycznej. W badaniu porównujemy wybrane rynki energii elektrycznej pod względem częstotliwości, skali oraz dynamiki występowania ujemnych cen oraz składu miks energetyczny. Przedmiotem zainteresowania jest również zjawisko ultra niskich cen, rozumianych jako ceny dodatnie, lecz zbliżone do zera, które mogą stanowić istotny sygnał narastającej nadpodaży energii elektrycznej. W pracy omawiamy konsekwencje pojawiania się ujemnych cen dla odbiorców końcowych oraz możliwe znaczenie taryf dynamicznych, magazynowania energii i elastyczności popytu.

Słowa kluczowe: ujemne ceny energii elektrycznej, ultra niskie ceny, odnawialne źródła energii, transformacja energetyczna, europejskie rynki energii

Negative prices in selected European electricity markets

The liberalisation of electricity markets has led to the deregulation of electricity prices and strengthened the role of market mechanisms in balancing supply and demand. The ongoing energy transition has resulted in a growing share of renewable energy sources in the energy mix, a development also supported by the operation of the merit-order mechanism. Due to the low marginal costs of electricity generation from wind and solar energy, an increase in renewable electricity generation may, under certain conditions, lead to a decline in wholesale electricity prices and, in some cases, to the occurrence of negative prices. That phenomenon has attracted increasing attention from regulators and market participants, particularly in the light of the growing number of negative price occurrences observed in European electricity markets. The study compares selected electricity markets in terms of the frequency, magnitude, and dynamics of negative price occurrences, as well as the composition of their energy mixes. The analysis also considers the phenomenon of ultra-low prices, defined as positive prices close to zero, which may constitute an important indicator of a growing oversupply of electricity. In addition, the paper discusses the implications of negative electricity prices for final consumers and examines the potential role of dynamic tariffs, energy storage, and demand-side flexibility.

Keywords: negative electricity prices, ultra-low prices, renewable energy sources, energy transition, European electricity markets

Dominik Krężolek, Uniwersytet Ekonomiczny w Katowicach

Agregacja estymatorów indeksu ogona rozkładu w pomiarze ryzyka ekstremalnego: zastosowanie na rynku metali

Estymacja ryzyka ekstremalnego na rynkach surowcowych jest istotnie wrażliwa na dwie arbitralne decyzje analityka: wybór estymatora indeksu ogona rozkładu (Hilla, Dekkersa-Einmahl-de Haana, jądrowego, o zredukowanym obciążeniu) oraz wybór liczby statystyk pozycyjnych k użytych w estymacji. Na rynku metali rozbieżność oszacowań ξ między estymatorami sięga 30-50%, co przekłada się na różnice w prognozach Expected Shortfall rzędu 20-40%. W niniejszym artykule proponujemy nowatorski schemat agregacji, który eliminuje konieczność arbitralnego wyboru, łącząc wiele estymatorów na siatce wartości k w jeden estymator zagregowany.

Schemat agregacji opiera się na ważonej kombinacji czterech estymatorów: Hilla, Dekkersa-Einmahl-de Haana, jądrowego Csörgő-Deheuvels-Masona z optymalnym pasmem Groenebooma (i współautorów) oraz estymatora o zredukowanym obciążeniu Caeiro-Gomesa, obliczanych na siatce wartości $k \in \{k_{min}, \dots, k_{max}\}$. Wagi agregacji wyznaczone są przez minimalizację pozapróbkowej funkcji straty Fisslera-Ziegla (FZ loss) dla par (VaR, ES). Jest to kluczowa innowacja: optymalizacja odbywa się ze względu na jakość prognoz miar ryzyka (downstream task), a nie ze względu na błąd estymacji samego ξ , który nie jest obserwowalny. Rozważamy trzy warianty: wagi stałe (wyznaczone na próbie treningowej), wagi adaptacyjne (rolling-window) oraz wagi z wykładniczym wygaszaniem.

Na gruncie teoretycznym pokazujemy zgodność estymatora zagregowanego pod warunkami regularnej zmienności i β -mixing zależności szeregu czasowego. Dowodzimy ponadto nierówności typu oracle: w najgorszym przypadku estymator zagregowany nie pogarsza błędu średniokwadratowego względem najlepszego estymatora składowego.

Własności proponowanej procedury weryfikujemy w rozbudowanym studium symulacyjnym Monte Carlo z procesami generującymi dane klasy AR(1)-GARCH(1,1) z innowacjami o rozkładzie t-Studenta (znany ξ) oraz z modelem przełączania reżimów Markowa (time-varying ξ). Estymator zagregowany dominuje każdy estymator indywidualny względem funkcji straty FZ i częstotliwości wchodzenia do zbioru Model Confidence Set Hansena, Lundego i Nasona (2011). Przewaga agregacji wzmacnia się w scenariuszu ze zmiennym w czasie indeksem ogona.

W badaniu empirycznym stosujemy procedurę do dziennych stop zwrotu sześciu metali notowanych na LME i COMEX (złoto, srebro, platyna, pallad, miedź, nikiel) w okresie 2000-2025, wykorzystując schemat rolling-window z oknem 1000 obserwacji. Estymator zagregowany generuje prognozy VaR i ES, które:

- 1) przechodzą backtesty VaR (Kupiec, Christoffersen) i ES (Bayer-Dimitriadis ESR) na konwencjonalnych poziomach istotności częściej niż jakikolwiek estymator indywidualny,
- 2) wchodzą do zbioru Model Confidence Set w ponad 90% okien rolling-window oraz
- 3) redukują wrażliwość oszacowań ES na wybór k o czynnik 3-5 w porównaniu z samym estymatorem Hilla.

Analiza ujawnia różnice między metalami szlachetnymi i przemysłowymi: dla tych drugich korzyści z agregacji są większe, co wynika z silniejszej heterogeniczności struktury ogona.

Wyniki mają bezpośrednie implikacje praktyczne. Agregacja estymatorów indeksu ogona stanowi prostą w implementacji, a jednocześnie skuteczną metodę redukcji ryzyka modelowego (model risk) w estymacji miar ryzyka ekstremalnego na rynkach surowcowych.

Słowa kluczowe: indeks ogona, agregacja estymatorów, uśrednianie modeli, Value-at-Risk, Expected Shortfall, funkcja straty Fisslera-Ziegla, Model Confidence Set, rynek metali, ryzyko modelowe

Aggregating Tail Index Estimators for Robust Extreme Risk Measurement: Evidence from Metal Commodity Markets

Extreme risk estimation in commodity markets is highly sensitive to two arbitrary decisions made by the analyst: the choice of the tail index estimator (Hill, Dekkers-Einmahl-de Haan, kernel, or bias-reduced) and the choice of the number of upper order statistics k used in the estimation. In metal markets, the discrepancy in ξ estimates across estimators reaches 30-50%, translating into differences in Expected Shortfall forecasts on the order of 20-40%. In this paper, we propose a novel aggregation scheme that eliminates the need for arbitrary selection by combining multiple estimators over a grid of k values into a single consensus estimator.

The aggregation scheme is based on a weighted combination of four estimators - Hill, Dekkers-Einmahl-de Haan, the kernel estimator of Csörgő-Deheuvels-Mason with the optimal bandwidth of Groeneboom et al., and the bias-reduced estimator of Caeiro-Gomes, computed over a grid of values $k \in \{k_{min}, \dots, k_{max}\}$. The aggregation weights are determined by minimizing the out-of-sample Fissler-Ziegel loss function (FZ loss) for (VaR, ES) pairs. This constitutes the key innovation: the optimization targets the quality of risk measure forecasts (downstream task) rather than the estimation error of ξ itself, which is unobservable. We consider three variants: fixed weights

(determined on a training sample), adaptive weights (rolling-window), and exponentially forgetting weights.

On theoretical grounds, we establish the consistency of the aggregated estimator under regular variation conditions and β -mixing dependence of the time series. We further prove an oracle-type inequality: in the worst case, the aggregated estimator does not deteriorate the mean squared error relative to the best individual component estimator.

The properties of the proposed procedure are verified in an extensive Monte Carlo simulation study using data-generating processes of the AR(1)-GARCH(1,1) class with Student-t innovations (known ξ) and a Markov regime-switching model (time-varying ξ). The aggregated estimator dominates every individual estimator in terms of the FZ loss function and the frequency of inclusion in the Model Confidence Set of Hansen, Lunde, and Nason (2011). The advantage of aggregation is amplified in the scenario with a time-varying tail index.

In the empirical application, we apply the procedure to daily returns of six metals listed on the LME and COMEX (gold, silver, platinum, palladium, copper, and nickel) over the period 2000-2025, using a rolling-window scheme with a window of 1,000 observations. The aggregated estimator produces VaR and ES forecasts that:

- 1) pass VaR backtests (Kupiec, Christoffersen) and ES backtests (Bayer-Dimitriadis ESR) at conventional significance levels more frequently than any individual estimator,
- 2) enter the Model Confidence Set in over 90% of rolling windows, and
- 3) reduce the sensitivity of ES estimates to the choice of k by a factor of 3-5 compared to the Hill estimator alone.

The analysis reveals differences between precious and industrial metals: for the latter, the benefits of aggregation are greater, reflecting stronger heterogeneity in the tail structure.

The results carry direct practical implications: aggregation of tail index estimators constitutes a straightforward-to-implement yet effective method for reducing model risk in extreme risk measure estimation in commodity markets.

Keywords: tail index, estimator aggregation, model averaging, Value-at-Risk, Expected Shortfall, Fissler-Ziegel loss function, Model Confidence Set, metal markets, model risk

Małgorzata Krzciuk, Uniwersytet Ekonomiczny w Katowicach

Wnioskowanie w oparciu o próby nielosowe - podejście quasi-randomizacyjne i oparte na modelu

W opracowaniu rozważany jest problem wnioskowania w oparciu o próby nielosowe. Omówione zostaną dwa główne podejścia w ramach tego zagadnienia tj. podejście quasi randomizacyjne i nadpopulacyjne. W przypadku każdego z nich zaprezentowane będą wybrane metody, takie jak: ważenie ze względu na skłonność do udzielenia odpowiedzi, próbkowanie odwrotne, metody odporne, wygładzanie jądrowe oraz metody wykorzystujące techniki uczenia maszynowego. Omówione zostaną także zalety oraz ograniczenia przedstawionych metod, wraz z obszarami ich zastosowań. Rozważania uzupełnione będą również badaniami symulacyjnymi własności omawianych estymatorów wybranych charakterystyk.

Słowa kluczowe: próby nielosowe, wnioskowanie statystyczne, podejście quasi randomizacyjne, podejście nadpopulacyjne

Inference based on non-probability samples - quasi-randomization approach and superpopulation modeling

This paper examines the issue of inference using non-probability samples. Two main approaches will be discussed: quasi-randomization and superpopulation modeling. For each approach, selected methods will be presented, including propensity score weighting, inverse sampling, double robust methods, kernel smoothing, and methods using certain machine learning techniques. The advantages and limitations of the presented methods, along with their application areas, will also be discussed. The discussion will also be complemented by simulation studies that examine the properties of the estimators of selected characteristics that are being considered.

Keywords: non-probability sampling, statistical inference, quasi randomisation approach, superpopulation modeling

Mariusz Kubus, Politechnika Opolska

Dobór zmiennych jako problem optymalizacji ciągłej – alternatywa dla metod sekwencyjnego przeszukiwania

Selekcja zmiennych przed estymacją modelu (filters) jest popularnym i chętnie wybieranym podejściem w modelowaniu predykcyjnym, zwłaszcza w problemach wysokowymiarowych. Kryteria oceniające indywidualny wpływ predyktorów na zmienną objaśnianą są efektywne obliczeniowo, ale nie uwzględniają interakcji między zmiennymi oraz oceniają podobnie zmienne redundantne umieszczając je obok siebie w rankingu. Z tego względu zaproponowano kryteria oceniające podzbiory zmiennych (np. integralną pojemność informacji Hellwiga), które maksymalizują związek predyktorów ze zmienną objaśnianą i jednocześnie minimalizują redundancję. Maksymalizacja takiego kryterium jest zadaniem optymalizacji kombinatorycznej, a do jego rozwiązania zwykle stosuje się heurystyczne metody przeszukiwania. Alternatywę stanowi metoda HSIC Lasso, formułująca problem selekcji zmiennych jako zadanie optymalizacji ciągłej z regularyzacją. Metoda ta wykorzystuje miarę HSIC (Hilbert-Schmidt Independence Criterion), pozwalającą wykrywać zarówno liniowe, jak i nieliniowe zależności między zmiennymi z wykorzystaniem metod jądrowych.

Celem referatu jest ocena, czy podejście oparte na optymalizacji ciągłej może zapewnić przewagę względem klasycznych metod sekwencyjnego przeszukiwania stosowanych w filtrach wielowymiarowych. Analiza zostanie przeprowadzona na danych symulacyjnych oraz zbiorze danych rzeczywistych.

Słowa kluczowe: selekcja zmiennych, filtry wielowymiarowe, HSIC, optymalizacja ciągła, regularyzacja

Variable selection as a continuous optimisation problem – an alternative to sequential search methods

Filtering of the variables before to model estimation is a popular and widely adopted approach in predictive modelling, particularly in high-dimensional problems. Criteria evaluating the individual relevance of predictors are computationally efficient, but they do not account for interactions between variables and tend to rank redundant variables similarly, placing them close together in the ranking. For this reason, criteria evaluating subsets of variables (e.g. Hellwig's integral information capacity) have been proposed, which maximise the relationship between predictors and the response variable whilst minimising redundancy. Maximising such a criterion is a combinatorial optimisation problem, and heuristic search methods are typically used to solve it. An alternative is the HSIC Lasso method, which formulates the variable selection problem as a continuous optimisation task with regularisation. This method utilises the HSIC (Hilbert-Schmidt Independence Criterion) measure, which allows for the detection of both linear and non-linear relationships between variables using kernel methods.

The aim of this paper is to assess whether an approach based on continuous optimisation can offer an advantage over traditional sequential search methods used in multivariate filters. The analysis will be carried out using simulation data and a real-world dataset.

Keywords: variable selection, multivariate filters, HSIC, continuous optimisation, regularisation

Paweł Kufel, Uniwersytet WSB Merito w Toruniu

Prezentacja pakietu realizującego bayesowskie uśrednianie modeli wektorowej autoregresji BMAVAR.gfn w gretlu

Referat przedstawia działanie pakietu w programie gretl do bayesowskiego uśredniania modeli wektorowej autoregresji (VAR). Koncepcja uśredniania modeli pozwala uwzględnić wiele specyfikacji i uśredniać wyniki na ich podstawie. Procedura generuje miary takie jak prawdopodobieństwo modelu (posterior model probability) oraz prawdopodobieństwo włączenia zmiennej do modelu (posterior inclusion probability). Otrzymywane są także prognozy i odpowiedzi funkcji impulsowych, które opierają się na połączeniu wyników z wielu modeli.

Omówione zostaną założenia implementacji pakietu w gretl oraz jego interfejs graficzny. Na przykładzie empirycznym zaprezentowane zostaną okna wynikowe, prognozy oraz funkcje odpowiedzi impulsowych.

Słowa kluczowe: uśrednianie modeli, modele wektorowej autoregresji, pakiet gretl

Presentation of the BMAVAR.gfn Package for Bayesian Model Averaging of Vector Autoregression Models in Gretl

This paper describes the functionality of the package in Gretl for Bayesian Model Averaging (BMA) of Vector Autoregression Models (VAR). The concept of model averaging allows consideration of multiple specifications and averaging the results based on them. The procedure generates measures

such as the posterior model probability and the posterior inclusion probability. Forecasts and impulse response functions are also obtained, which are derived from a combination of results from accepted models.

The assumptions underlying the package's implementation in gretl and its graphical user interface will be outlined. Using an empirical example, the results windows, forecasts, and impulse response functions will be presented.

Keywords: model averaging, vector autoregression models, gretl package

Joanna Landmesser-Rusek, Szkoła Główna Gospodarstwa Wiejskiego w Warszawie

Konstrukcja i formalna weryfikacja grafów przyczynowych na rynku walutowym z wykorzystaniem ekosystemu PyWhy

Klasyczne podejścia do analizy powiązań na rynku walutowym opierają się głównie na modelach korelacji lub współzależności liniowych, co często prowadzi do identyfikacji pozornych relacji i uniemożliwia precyzyjne wskazanie kierunku przepływu impulsów informacyjnych. Niniejsze badanie podejmuje próbę przezwyciężenia tych ograniczeń poprzez zastosowanie paradygmatu wnioskowania przyczynowego (Causal AI) do budowy strukturalnych grafów walutowych. Analizie poddano dzienne szeregi logarymicznych stóp zwrotu dla koszyka głównych walut światowych. Proces badawczy zrealizowano w oparciu o otwarte biblioteki ekosystemu PyWhy (DoWhy oraz EconML), łączące teorię grafów Judea Pearla z zaawansowanymi metodami uczenia maszynowego. W pierwszym etapie, z uwagi na nieliniowy charakter i grube ogony rozkładów stóp zwrotu, zastosowano algorytmy Causal Discovery z uwzględnieniem opóźnień czasowych (Time-Lagged Graphs) w celu eliminacji pętli zwrotnych i konstrukcji acyklicznych grafów skierowanych (DAG). Badanie przeprowadzono w oknie pełzającym, co pozwoliło na uchwycenie dynamiki zmian struktury rynkowej w czasie. W drugim etapie, wykorzystując technikę Double Machine Learning (DML) zaimplementowaną w EconML, wyizolowano czyste efekty przyczynowe (CATE) między poszczególnymi walutami, kontrolując wpływ globalnych czynników zakłócających (tzw. confounders), takich jak indeks zmienności VIX czy rentowności obligacji skarbowych. Kluczowym elementem pracy jest formalna weryfikacja stabilności uzyskanej struktury grafu za pomocą modułu Refutation w bibliotece DoWhy, gdzie odporność modeli przetestowano poprzez testy placebo oraz symulacje niezaobserwowanych zmiennych zakłócających. Wyniki badań pozwalają na identyfikację stabilnych w czasie, strukturalnych powiązań przyczynowo-skutkowych na rynku FX, co stanowi istotny krok naprzód w obszarze zarządzania ryzykiem walutowym oraz budowy odpornych strategii inwestycyjnych.

Słowa kluczowe: Causal AI, PyWhy, grafy walutowe, wnioskowanie przyczynowe, Double Machine Learning

Construction and Formal Refutation of Causal Graphs in the Foreign Exchange Market Using the PyWhy Ecosystem

Classical approaches to analyzing linkages in the foreign exchange market primarily rely on correlation models or linear dependencies, which often lead to the identification of spurious relationships and preclude precise directionality of information impulse flows. This study attempts to overcome

these limitations by applying the causal inference paradigm (Causal AI) to construct structural currency graphs. Daily log-return series for a basket of major global currencies were analyzed. The research workflow was implemented using open-source libraries from the PyWhy ecosystem (DoWhy and EconML), which integrate Judea Pearl's graph theory with advanced machine learning methods. In the first stage, addressing the non-linear characteristics and heavy-tailed distributions of log-returns, causal discovery algorithms incorporating time-lagged graphs were employed to eliminate feedback loops and construct directed acyclic graphs (DAGs). The analysis was conducted using a rolling window approach, thereby capturing the temporal dynamics of the market structure. In the second stage, leveraging the Double Machine Learning (DML) technique implemented in EconML, conditional average treatment effects (CATE) between individual currencies were isolated while controlling for global confounding factors, such as the VIX volatility index and government bond yields. A pivotal element of this work is the formal validation of the resulting graph structure stability using the Refutation module in the DoWhy library, where model robustness was evaluated through placebo tests and simulations of unobserved confounders. The empirical findings enable the identification of time-stable, structural cause-and-effect relationships in the FX market, representing a significant advancement in the domains of currency risk management and the design of resilient investment strategies.

Keywords: Causal AI, PyWhy, currency graphs, causal inference, Double Machine Learning

Samuele Lo Piano, Gdansk University of Technology, Universitat Pompeu Fabra; Marta Kuc-Czarnecka, Gdansk University of Technology; Roger Pielke Jr., American Enterprise Institute, Washington, University of Colorado Boulder; Andrea Saltelli, Universitat Pompeu Fabra, University of Bergen

Paradigm Before Parameter: A Nested Sensitivity Design for Quantifying Structural Uncertainty in World Population Projections

Demographic projections routinely condition high-stakes policy decisions, yet the uncertainty communicated around them is almost always parametric — fertility, mortality, and migration assumptions — while the choice of forecasting paradigm itself is treat.

Keywords: sensitivity analysis; structural uncertainty; demographic forecasting; variance decomposition; model validation

Paweł Lula, Uniwersytet Ekonomiczny w Krakowie

Wspierana za pomocą ontologii identyfikacja odwołań do konceptów składających się na system klasyfikacji wiedzy

Celem pracy jest opracowanie metody identyfikacji fragmentów tekstów naukowych zawierających odwołania do konceptów składających się na przyjęty system klasyfikacji wiedzy. Podstawą proponowanego rozwiązania jest system klasyfikacji wiedzy. Z uwagi na przyjęte założenie dotyczące potrzeby uwzględnienia wszystkich obszarów wiedzy wykorzystano system stosowany w systemie OpenAlex uwzględniający istnienie ponad 4500 tematów (topics) pogrupowanych w 252 podobszary (subfields), które tworzą 26 obszarów (fields) istniejących w ramach 4 dziedzin nauki (domains).

Przypisanie fragmentów tekstu do poszczególnych konceptów odbywa się poprzez porównanie charakterystyk poszczególnych elementów przyjętego systemu klasyfikacji z kolejnymi fragmentami analizowanego tekstu.

W części empirycznej pracy zaprezentowane działanie systemu w odniesieniu do opisów projektów finansowanych przez Komisję Europejską.

Słowa kluczowe: przetwarzanie języka naturalnego, duże modele językowe, klasyfikacja wiedzy, klasyfikacja dokumentów

Ontology-Based Identification of References to Concepts Composing a Knowledge Classification System

The aim of this work is to develop a method for identifying fragments in scientific texts that contain references to concepts forming part of established knowledge classification system. The proposed solution is based on a knowledge classification system. Given the assumption that all areas of knowledge should be taken into account, the system used OpenAlex classification system, which includes over 4,500 topics grouped into 252 subfields, which in turn form 26 fields within 4 scientific domains. Text fragments are assigned to specific concepts by comparing the semantic characteristics of nodes from knowledge classification system with successive fragments of the analyzed text.

The empirical part of this study presents the application of the proposed solution for analysis of descriptions of research projects financed by the European Commission.

Keywords: natural language processing, large language models, knowledge classification system, document classification

Aleksandra Łuczak, Wydział Ekonomiczny, Uniwersytet Przyrodniczy w Poznaniu; Sławomir Kalinowski, Instytut Rozwoju Wsi i Rolnictwa PAN

Wielowymiarowy barometr ubóstwa jako narzędzie identyfikacji wczesnych sygnałów biedy

Tradycyjne metody pomiaru ubóstwa, oparte głównie na danych dochodowych lub wybranych wskaźnikach deprivacji materialnej, często prowadzą do niejednoznacznych ocen jego skali i charakteru. Brak spójnego, wielowymiarowego podejścia utrudnia zarówno rzetelną diagnozę sytuacji społecznej, jak i projektowanie skutecznych interwencji publicznych. W odpowiedzi na te ograniczenia opracowano wielowymiarowy barometr ubóstwa – narzędzie integrujące różnorodne obszary życia społecznego oraz umożliwiające systematyczny monitoring zagrożenia ubóstwem na poziomie gmin. Metodologia badania wykorzystuje zmodyfikowaną metodę TOPSIS, pozwalającą na konstrukcję syntetycznego wskaźnika opartego na danych dotyczących m.in. zdrowia, edukacji i warunków mieszkaniowych. Analizy przeprowadzono na podstawie danych z Banku Danych Lokalnych GUS dla gmin województwa mazowieckiego w latach 2012–2023. Zakres badania rozszerzono o ocenę tempa poprawy lub pogorszenia sytuacji w poszczególnych gminach, co umożliwia wnioskowanie o skuteczności podejmowanych interwencji publicznych. Wyniki ujawniły wyraźny gradient przestrzenny: gminy położone bardziej peryferyjnie częściej charakteryzują się wyższymi wartościami barometru, co wskazuje na większe ryzyko ubóstwa. Opracowany barometr stanowi narzędzie

wspierające formułowanie polityk publicznych dostosowanych do lokalnych uwarunkowań oraz umożliwia precyzyjniejsze ukierunkowanie działań przeciwdziałających ubóstwu i wykluczeniu społecznemu.

Słowa kluczowe: wielowymiarowy pomiar, metoda TOPSIS, ubóstwo, syntetyczny wskaźnik

Wojciech Łukaszonek, Uniwersytet Kaliski; Marcin Szymkowiak, Uniwersytet Ekonomiczny w Poznaniu; Waldemar Wołyński, Uniwersytet Adama Mickiewicza w Poznaniu

Grupowa ważność zmiennych w RF jako narzędzie interpretacji skupień MFPCA na przykładzie powiatowych rynków pracy

Analiza regionalnych rynków pracy wymaga uwzględnienia jednocześnie wielu wskaźników oraz ich dynamiki w czasie. Podejścia oparte wyłącznie na danych przekrojowych pomijają ciągłość zachodzących procesów, co ogranicza możliwości analizy rozwoju sytuacji na rynku pracy. W referacie zaproponowano podejście łączące wielowymiarową funkcjonalną analizę głównych składowych (MFPCA) z lasem losowym (Random Forest, RF) jako narzędziem niezależnej walidacji i interpretacji wyników.

Badaniem objęto wszystkie powiaty w Polsce ($n = 380$) w latach 2004-2025 ($T = 22$), opisane 12 zmiennymi (m.in. stopa bezrobocia rejestrowanego, udział bezrobotnych długotrwale pozostających bez pracy, przeciętne wynagrodzenie, liczba podmiotów gospodarczych na 1000 mieszkańców). Zmienne poddano unitaryzacji zerowanej z uwzględnieniem ich charakteru (stymulanta/destymulanta). Trajektorie czasowe wygładzono z wykorzystaniem dwóch baz funkcyjnych: B-sklejanej (B-spline) oraz Fouriera. Funkcjonalna reprezentacja zbioru danych była podstawą analizy głównych składowych, a otrzymane ładunki stanowiły podstawę hierarchicznej analizy skupień metodą Warda ($k = 5$).

Etykiety wyłonionych grup wykorzystano jako zmienną celu w modelu lasu losowego, którego wejściem była zwektoryzowana macierz obserwacji pierwotnych w układzie powiat $\times$ zmienna_rok (380×264), niezależna od wyboru bazy funkcyjnej. Zastosowano grupową ważność permutacyjną zmiennych (Gregorutti i in., 2015), pozwalającą ocenić łączny wkład całej trajektorii danej zmiennej w różnicowanie skupień, z pominięciem zniekształceń wynikających z autokorelacji czasowej.

Modele osiągnęły wysoką dokładność OOB (87,9% dla bazy B-sklejanej, 87,1% dla bazy Fouriera), co potwierdziło merytoryczną spójność podziału wyznaczonego przez MFPCA. Analiza rozkładu różnic klasyfikacji ($\Delta = \text{skupienie RF} - \text{skupienie MFPCA}$) pozwoliła zidentyfikować powiaty graniczne, wymagające pogłębionej analizy. Ranking ważności zmiennych, wyznaczony osobno dla obu modeli ($Y_{\text{B-spline}}$ i Y_{Fourier}), umożliwił określenie charakterystyk rynku pracy najsilniej różnicujące wyznaczone grupy.

Uzyskane wyniki wskazują, że połączenie MFPCA z lasem losowym umożliwia jednocześnie metodologicznie zaawansowane grupowanie uwzględniające dynamikę czasową oraz merytorycznie zrozumiałą interpretację w kategoriach oryginalnych wskaźników rynku pracy.

Słowa kluczowe: dane funkcyjne, analiza głównych składowych, las losowy, grupowa ważność zmiennych, rynek pracy

Group Variable Importance in Random Forests as a Tool for Interpreting MFPCA-Based Classification The Case of Polish County Labour Markets

The analysis of regional labour markets requires the simultaneous consideration of multiple indicators and their dynamics over time. Approaches based solely on cross-sectional data neglect the continuity of underlying processes, thereby limiting the ability to analyse labour market development. This paper proposes a framework that combines Multivariate Functional Principal Component Analysis (MFPCA) with Random Forests (RF) as a tool for independent validation and interpretation of the classification results.

The study covered all Polish counties ($n = 380$) over the period 2004–2025 ($T = 22$), described by 12 variables, including the registered unemployment rate, the share of long-term unemployed individuals, average wages, and the number of business entities per 1000 inhabitants. All variables were normalised using zero unitarisation, taking into account their nature (stimulants or destimulants). Trajectories were smoothed using two basis systems: B-spline and Fourier. The resulting functional representation served as the input for principal component analysis, and the obtained component scores were used in hierarchical clustering based on Ward's method ($k = 5$).

The cluster labels were then used as the response variable in a Random Forest model, while the predictors consisted of the vectorised matrix of the original observations arranged in a county $\times$ variable-year format (380×264), independent of the choice of functional basis. Group permutation variable importance (Gregorutti et al., 2015) was applied to assess the joint contribution of the entire temporal trajectory of each variable to cluster discrimination, while reducing distortions caused by temporal autocorrelation.

The models achieved high out-of-bag (OOB) classification accuracy (87.9% for the B-spline basis and 87.1% for the Fourier basis), confirming the substantive consistency of the clustering obtained by MFPCA. The analysis of the distribution of classification differences ($\Delta = \text{RF cluster} - \text{MFPCA cluster}$) enabled the identification of borderline counties requiring further investigation. Variable importance rankings, computed separately for the two models ($Y_{\text{B-spline}}$ and Y_{Fourier}), made it possible to identify the labour market characteristics that contributed the most to distinguishing the identified clusters.

The results demonstrate that combining MFPCA with Random Forests enables both, methodologically advanced clustering that captures temporal dynamics and substantively meaningful interpretation in terms of the original labour market indicators.

Keywords: functional data, principal component analysis, Random Forest, group variable importance, labour market

Justyna Majewska, Uniwersytet Ekonomiczny w Katowicach

Przestrzenne zróżnicowanie dostępności usług a wyludnienie starzejących się obszarów metropolitalnych na przykładzie Górnośląsko-Zagłębiowskiej Metropolii

Starzenie się ludności i spadek jej liczby często występują jednocześnie w obszarach metropolitalnych, jednak sama struktura wieku nie przesądza o kierunku zmian: część regionów o wysokim udziale osób starszych utrzymuje ludność lub ją zwiększa (Eurostat, 2024). O tym, czy gmina traci mieszkańców, współdecydują więc warunki lokalne, a wśród nich dostępność usług publicznych. Celem opracowania jest ustalenie, czy i w jakim stopniu dostępność usług różnicuje natężenie odpływu migracyjnego

między gminami Górnośląsko-Zagłębiowskiej Metropolii – obszaru poprzemysłowego o wyraźnym zróżnicowaniu demograficznym.

Analizę oparto na danych zastanych dla 41 gmin metropolii: statystyce ludności z Banku Danych Lokalnych (GUS) oraz lokalizacjach placówek opieki, ochrony zdrowia i transportu, na podstawie których wyznaczono przestrzenną dostępność usług. Związek dostępności z natężeniem odpływu oszacowano przestrzennym modelem bayesowskim, który pozwala oddzielić wpływ dostępności od struktury wieku i innych cech gmin.

Analiza pozwoli wskazać, w których typach gmin niska dostępność usług współwystępuje z podwyższonym odpływem mieszkańców, oraz wyodrębnić obszary metropolii najbardziej narażone na dalszy ubytek ludności.

Słowa kluczowe: dostępność usług, wyludnienie, starzenie się ludności, odporność demograficzna

Spatial variation in service accessibility and the depopulation of ageing metropolitan areas: the case of the Upper Silesian Metropolis (GZM)

Population ageing and population decline often occur together in metropolitan areas, yet age structure alone does not determine which way a place goes: some regions with a high share of older residents hold their population or even grow (Eurostat, 2024). Whether a municipality loses residents therefore depends in part on local conditions, among them the accessibility of public services. This paper sets out to establish whether, and to what extent, service accessibility accounts for differences in out-migration across the municipalities of the Upper Silesian Metropolis (GZM) – a post-industrial area marked by pronounced demographic contrasts.

The analysis draws on existing data for the 41 municipalities of the metropolis: population statistics from the Local Data Bank of Statistics Poland (GUS), together with the locations of care, health and transport facilities, from which spatial service accessibility was derived. The relationship between accessibility and out-migration was estimated using a spatial Bayesian model, which separates the effect of accessibility from age structure and other municipal characteristics.

The analysis will identify the types of municipality in which low service accessibility coincides with elevated out-migration, and locate the areas of the metropolis most exposed to further population loss.

Keywords: public service accessibility, population ageing, metropolitan areas, spatial statistics

Aleksandra Matuszewska-Janica, Szkoła Główna Gospodarstwa Wiejskiego w Warszawie

Zróżnicowanie komponentów luki płacowej w krajach UE: trzyskładnikowa dekompozycja Oaxaca-Blindera z wykorzystaniem danych Eurostat SES 2022

Luka płacowa między kobietami i mężczyznami (gender pay gap) jest jednym ze wskaźników wykorzystywanych przez Eurostat do monitorowania realizacji Celów Zrównoważonego Rozwoju (wskaźnik SDG 5.20). Wskaźnik ten odzwierciedla różnicę w przeciętnych wynagrodzeniach kobiet i mężczyzn, jednak nie pozwala na pełną ocenę rzeczywistej skali nierówności płacowych. Część

obserwowanej luki wynika bowiem z odmiennych charakterystyk obu grup, takich jak poziom wykształcenia, doświadczenie zawodowe, wykonywany zawód czy sektor zatrudnienia. Dlatego istotnym zagadnieniem jest identyfikacja zakresu, w jakim różnice w wynagrodzeniach mogą zostać wyjaśnione zróżnicowaniem cech pracowników, oraz określenie części luki niewyjaśnionej, która może odzwierciedlać wpływ innych czynników.

Celem prezentowanego badania jest ocena stopnia wyjaśnienia luki płacowej między kobietami i mężczyznami w krajach Unii Europejskiej oraz identyfikacja podobieństw i różnic w strukturze komponentów dekompozycji pomiędzy państwami. W analizie zastosowano trzyskładnikową dekompozycję Oaxaca–Blindera. Badanie przeprowadzono z wykorzystaniem danych jednostkowych pochodzących z Badania Struktury Wynagrodzeń (Structure of Earnings Survey – SES) Eurostatu z 2022 roku. W celu identyfikacji grup państw o podobnej strukturze komponentów dekompozycji zostanie wykorzystana metoda k-średnich, co umożliwi identyfikację profili państw UE pod względem potencjalnych mechanizmów kształtujących lukę płacową.

Słowa kluczowe: luka płacowa, dekompozycja Oaxaca–Blindera, analiza skupień

Variations in the components of the gender wage gap across EU countries: a three-component Oaxaca-Blinder decomposition based on Eurostat SES 2022 data

The gender wage gap is one of the indicators used by Eurostat to monitor progress towards the Sustainable Development Goals (SDG indicator 5.20). This indicator reflects the difference in average pay between women and men; however, it does not provide a full assessment of the scale of pay inequality. This is because the observed differences in the two groups' characteristics, such as level of education, work experience, occupation, or sector of employment, are a contributing factor. It is therefore important to identify the extent to which differences in pay can be explained by variations in workers' characteristics, and to determine the unexplained portion of the gap, which may reflect the influence of other factors.

The aim of this study is to assess the extent to which the gender wage gap in European Union countries can be explained, and to identify similarities and differences in the structure of the decomposition components across countries. The analysis employed the Oaxaca–Blinder decomposition, which is a three-fold decomposition. The present study was conducted using individual-level data from Eurostat's 2022 Structure of Earnings Survey (SES). The k-means method was employed to identify groups of countries with analogous structures of decomposition components, thereby enabling the identification of profiles of EU countries in terms of the mechanisms shaping the pay gap.

Keywords: gender wage gap, Oaxaca–Blinder decomposition, cluster analysis

*Izabela Miechowicz, Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu;
Kacper Grudniewski, Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu (student);
Natalia Pieprz, Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu (studentka);
Jakub Wesołowski, Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu (student);
Joanna Pawlak, Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu; Katarzyna
Dąbrowska, Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu (studentka);*

Od objawu do prawdopodobieństwa: matematyczne wspomaganie diagnozy choroby afektywnej

Celem pracy było opracowanie i walidacja modelu predykcyjnego opartego na danych z kwestionariusza OPCRIT (Operational Criteria Checklist for Psychotic Disorders), umożliwiającego różnicowanie zaburzeń afektywnych jednobiegunowych i dwubiegunowych. Dodatkowym założeniem było ograniczenie częstości błędnych rozpoznań w przypadkach o niejednoznacznym obrazie klinicznym. Model miał na celu estymację prawdopodobieństwa występowania choroby afektywnej dwubiegunowej na podstawie wybranych objawów nieswoistych dla epizodów manii i hipomanii.

W badaniu uczestniczyło 585 osób z rozpoznaniem choroby afektywnej dwubiegunowej oraz 225 osób z chorobą afektywną jednobiegunową. U wszystkich badanych zastosowano kwestionariusz OPCRIT. Następnie całą grupę podzielono losowo na zbiór uczący i testowy w stosunku 80:20, korzystając z generatora liczb losowych.

Na podstawie danych ze zbioru uczącego skonstruowano model statystyczny (regresję logistyczną) oparty na ośmiu pytaniach zawartych w kwestionariuszu OPCRIT. Kolejnym krokiem była ocena jakości tego modelu przy użyciu metody U-smile, w której porównano go z modelem zerowym, zakładającym jednakowe prawdopodobieństwo wystąpienia choroby dwubiegunowej dla każdego pacjenta, odpowiadające częstości tej choroby w całej kohorcie.

Następnie przeprowadzono walidację działania modelu na zbiorze uczącym, obliczając jego czułość, swoistość oraz wartości predykcyjne dodatnią i ujemną wraz z 95% przedziałami ufności. Na końcu ponownie zastosowano metodę U-smile, aby sprawdzić, czy opracowany model klasyfikacyjny przewyższa model referencyjny w zakresie jakości klasyfikacji pacjentów.

Model regresji logistycznej uzyskał na zbiorze treningowym czułość równą 93,33% oraz swoistość na poziomie 40,78%. Wyniki na zbiorze testowym potwierdziły jego dobrą skuteczność w wykrywaniu przypadków choroby afektywnej dwubiegunowej.

Z kolei analiza przeprowadzona metodą U-smile wskazała na wyraźnie lepszą jakość klasyfikacji niż w modelu odniesienia. Poprawa ta była widoczna zarówno w danych uczących, jak i testowych, i dotyczyła obu badanych jednostek chorobowych.

Opracowany model może stanowić wartościowe narzędzie wspierające proces diagnostyczny zaburzeń afektywnych, zwłaszcza w sytuacjach charakteryzujących się niejednoznacznym obrazem klinicznym. Jego wdrożenie jest stosunkowo proste oraz nie wymaga znacznych nakładów czasowych, a jednocześnie umożliwia estymację indywidualnego prawdopodobieństwa rozpoznania choroby afektywnej dwubiegunowej u pacjenta.

Słowa kluczowe: choroba afektywna jednobiegunowa, choroba afektywna dwubiegunowa, OPCRIT, model regresji logistycznej, metoda U-smile

From Symptoms to Probability: A Mathematical Approach to Supporting the Diagnosis of Affective Disorders

The aim of this study was to develop and validate a predictive model based on data obtained from the Operational Criteria Checklist for Psychotic Disorders (OPCRIT) to differentiate between unipolar depressive disorder and bipolar disorder. An additional objective was to reduce the rate of diagnostic misclassification in patients presenting with an ambiguous clinical picture. The model was designed to estimate the probability of bipolar disorder based on selected clinical features that are not specific to manic or hypomanic episodes.

The study included 585 patients diagnosed with bipolar disorder and 225 patients with unipolar depressive disorder. All participants were assessed using the OPCRIT questionnaire. The entire cohort was then randomly divided into training and test datasets in an 80:20 ratio using a random number generator.

A statistical model based on logistic regression was developed using data from the training dataset. The model incorporated eight variables derived from the OPCRIT questionnaire. Its performance was initially evaluated using the U-smile method by comparing it with a null model, which assumed an identical probability of bipolar disorder for every patient corresponding to the overall prevalence of bipolar disorder in the study cohort.

The model was subsequently validated by calculating its sensitivity, specificity, positive predictive value (PPV), and negative predictive value (NPV), together with their 95% confidence intervals. Finally, the U-smile method was applied again to determine whether the proposed classification model outperformed the reference model in terms of overall classification quality.

In the training dataset, the logistic regression model achieved a sensitivity of 93.33% and a specificity of 40.78%. Performance in the test dataset confirmed its good ability to identify patients with bipolar disorder.

The U-smile analysis demonstrated substantially better classification performance than the reference model. This improvement was observed in both the training and test datasets and was evident for both diagnostic categories.

The proposed model may serve as a valuable decision-support tool in the diagnostic assessment of affective disorders, particularly in patients with an ambiguous clinical presentation. Its implementation is straightforward, requires minimal time and resources, and enables estimation of an individual patient's probability of having bipolar disorder.

Keywords: unipolar depressive disorder; bipolar disorder; OPCRIT; logistic regression model; U-smile method

Małgorzata Misztal, Uniwersytet Łódzki; Dorota Pekasiewicz, Uniwersytet Łódzki

Zróżnicowanie regionalne dochodów gospodarstw domowych emerytów w Polsce – wielowymiarowa analiza statystyczna

W ostatnich latach obserwowany jest postępujący proces starzenia się ludności w Polsce. Według informacji Głównego Urzędu Statystycznego na koniec 2024 roku liczba osób w wieku 60 i więcej lat

wynosiła około 10 mln, co oznacza wzrost o 0,9% w porównaniu do roku 2023. Wskaźnik ludności w wieku poprodukcyjnym na 100 osób w wieku produkcyjnym był równy 40,8, podczas gdy w roku 2023 wynosił 40,0.

Wydłużająca się średnia długość życia przekłada się na wzrost liczby osób, które przechodzą na emeryturę, a zatem zwiększa się liczba gospodarstw domowych, w których głównym źródłem utrzymania są świadczenia emerytalne. Uważa się, że sytuacja dochodowa gospodarstw domowych emerytów jest zwykle bardziej stabilna niż gospodarstw utrzymujących się z pracy ale poziom dochodów jest zazwyczaj niższy niż w gospodarstwach osób aktywnych zawodowo.

W pracy dokonano analizy sytuacji ekonomicznej gospodarstw domowych emerytów w województwach Polski w kontekście ich miesięcznych dochodów ekwiwalentnych. Wykorzystano wybrane parametry rozkładu dochodów, w tym przeciętny dochód ekwiwalentny oraz wskaźniki nierówności, ubóstwa i zamożności. Celem badania była ocena regionalnego zróżnicowania dochodów gospodarstw domowych emerytów w Polsce w układzie wojewódzkim. Wykorzystując metody wielowymiarowej analizy danych (analizę głównych składowych i analizę skupień) wskazano różnice w sytuacji dochodowej gospodarstw domowych emerytów w poszczególnych województwach i zidentyfikowano grupy szczególnie narażone na ubóstwo.

Słowa kluczowe: rozkłady dochodów, nierówności dochodowe, ubóstwo, emeryci, analiza wielowymiarowa

Regional differentiation of pensioners' household incomes in Poland – a multivariate statistical analysis

In recent years, the population of Poland has been ageing progressively. According to statistics from the Central Statistical Office, at the end of 2024 the number of people aged 60 and over reached approximately 10 million, representing an increase of 0.9 per cent compared with 2023. The ratio of the post-working-age population to every 100 people of working age reached 40.8, whereas in 2023 it stood at 40.0.

Longer life expectancy is leading to an increase in the number of people retiring, and consequently there is a rise in the number of households where pension benefits are the main source of income. It is generally considered that the income situation of pensioner households is usually more stable than that of households relying on employment, but income levels are typically lower than in households of working people.

The study analyzed the economic situation of pensioners' households in Poland's voivodeships in terms of their monthly equivalent incomes. Selected parameters of the income distribution were used, including average equivalent income and inequality, poverty and richness indices. The objective of the study was to assess the regional differentiation of pensioners' household incomes in Poland at the voivodeship level. Using methods of multivariate data analysis (including principal component analysis and cluster analysis), differences in the income situation of pensioners' households across individual voivodeships were highlighted, and groups particularly vulnerable to poverty were identified.

Keywords: income distributions, income inequality, poverty, pensioners, multivariate analysis

Warunkowa heteroskedastyczność w tygodniowej liczbie zgonów w Polsce. Zastosowanie modeli ARMA–GARCH w prognozowaniu umieralności

W badaniu analizowano dynamikę tygodniowej liczby zgonów w Polsce w latach 2011–2025 z wykorzystaniem modeli warunkowej heteroskedastyczności. Modelowanie prowadzono dla logarytmicznie transformowanych danych o liczbie zgonów. Równanie średniej obejmowało deterministyczny trend, efekty sezonowe modelowane za pomocą zmiennych zero-jedynkowych, zmienną opisującą okres pandemii COVID-19 oraz składnik ARMA odpowiedzialny za krótkookresową zależność czasową. Warunkową wariancję modelowano za pomocą procesów GARCH.

W badaniu porównano różne warianty specyfikacji modeli, obejmujące alternatywne struktury równania średniej i wariancji oraz różne sposoby uwzględniania efektu pandemii. Ocenę modeli przeprowadzono z wykorzystaniem kryteriów informacyjnych oraz miar jakości prognoz i dopasowania. Uzyskane wyniki wskazują na występowanie istotnej warunkowej heteroskedastyczności oraz klastrowania zmienności w tygodniowych szeregach liczby zgonów nawet po uwzględnieniu trendu, sezonowości, efektów pandemicznych i zależności autoregresyjnych. Prognozy uzyskane dla danych logarytmicznych poddano odwrotnej transformacji logarytmicznej do pierwotnej skali liczby zgonów. Uwzględnienie modelowania zmiennej wariancji w czasie poprawiło dokładność prognoz oraz jakość przedziałów predykcji.

Badanie pokazuje, że modele zmienności stosowane w ekonometrii finansowej mogą stanowić wartościowe narzędzie analizy wysokoczęstotliwościowych procesów demograficznych i stanowić użyteczne rozszerzenie klasycznych modeli prognozowania umieralności stosowanych w demografii.

Słowa kluczowe: modele GARCH, warunkowa heteroskedastyczność, szeregi czasowe, prognozowanie demograficzne, umieralność

Conditional heteroscedasticity in weekly death rates in Poland. The use of ARMA–GARCH models in mortality forecasting

The study analysed trends in weekly death rates in Poland between 2011 and 2025 using conditional heteroscedasticity models. Modelling was carried out on logarithmically transformed data on the number of deaths. The mean equation included a deterministic trend, seasonal effects modelled using binary variables, a variable describing the COVID-19 pandemic period, and an ARMA component responsible for short-term time dependence. Conditional variance was modelled using GARCH processes.

The study compared various model specification variants, including alternative structures of the mean and variance equations and different ways of accounting for the pandemic effect. The models were evaluated using information criteria and measures of forecast quality and fit. The results indicate the presence of significant conditional heteroscedasticity and clustering of variability in the weekly death count series, even after accounting for trend, seasonality, pandemic effects and autoregressive relationships. The forecasts obtained for the logarithmic data were subjected to inverse logarithmic transformation back to the original scale of death counts. The inclusion of time-varying variance modelling improved the accuracy of the forecasts and the quality of the prediction intervals.

The study shows that volatility models used in financial econometrics can serve as a valuable tool for analysing high-frequency demographic processes and constitute a useful extension of the classical mortality forecasting models used in demography.

Keywords: GARCH models, conditional heteroscedasticity, time series, demographic forecasting, mortality

Kamila Migdał Najman, Uniwersytet Gdański; Krzysztof Najman, Uniwersytet Gdański

Pokolenie Z na pomorskim rynku pracy

Postępujący coraz szybciej rozwój technologii informatycznych, pogłębiająca się cyfryzacja oraz rosnąca liczba zastosowań sztucznej inteligencji stały się kluczowymi czynnikami transformacji społeczno-gospodarczej, prowadząc do istotnych zmian na rynku pracy. Procesy te sprzyjają zarówno automatyzacji zadań rutynowych i ograniczaniu zatrudnienia w niektórych sektorach, jak i powstawaniu nowych, wysoko wyspecjalizowanych zawodów, których zakres i znaczenie ulegają dynamicznym przeobrażeniom. Wzrasta znaczenie kompetencji cyfrowych, analitycznych i interpersonalnych oraz zdolności adaptacyjnych, co wymusza ciągłe podnoszenie kwalifikacji i elastyczność zawodową.

W konsekwencji zmienia się także profil kompetencji oczekiwanych w szczególności od najmłodszego pokolenia, które zaczyna wchodzić na rynek pracy – pokolenia Z. Młodzi pracownicy wyróżniają się potrzebą samorealizacji, dążeniem do równowagi między życiem zawodowym a prywatnym oraz pragnieniem uczestnictwa w procesie decyzyjnym, a jednocześnie oczekują środowiska sprzyjającego współpracy, zaufaniu i kreatywności.

Celem prezentowanych badań jest analiza porównawcza oczekiwań pracowników należących do pokolenia Z, z oczekiwaniami rynku pracy w woj. pomorskim w 2025 roku. Oczekiwania pracodawców sformułowano w oparciu o treść ofert pracy publikowanych na portalach internetowych scrapowanych od lipca 2024 do lipca 2025. Dane te skonfrontowano z wynikami Barometru Zawodów oraz Badaniem Losów Absolwentów w woj. pomorskim a także badaniami realizowanymi wśród studentów Uniwersytetu Gdańskiego w ramach badań nad koncepcją organizacji turkusowych w 2025 roku. W badaniach zastosowano metody EAD (Eksploracyjnej Analizy Danych), ML (machine learning).

Badanie pozwoliło na identyfikację luk kompetencyjnych pokolenia Z, a także problemy niedopasowania oczekiwań rynku pracy z jednej strony do oczekiwań młodych pracowników z drugiej.

Słowa kluczowe: rynek pracy, luki kompetencyjne, eksploracyjna analiza danych

Generation Z in the Pomeranian Labour Market

The accelerating development of information technologies, the ongoing digitalization of economic and social processes, and the expanding applications of artificial intelligence have become key drivers of socio-economic transformation, resulting in profound changes in labour markets. These processes facilitate both the automation of routine tasks and the reduction of employment in certain sectors, while simultaneously creating new, highly specialized occupations whose scope and significance continue to evolve dynamically. As a consequence, digital, analytical, and interpersonal competencies,

as well as adaptability, are gaining increasing importance, necessitating continuous skills development and professional flexibility.

Accordingly, the profile of competencies expected from the youngest generation entering the labour market - Generation Z - is also undergoing significant transformation. Young employees are characterized by a strong need for self-actualization, a desire to maintain work - life balance, and an expectation to participate in organizational decision making processes. At the same time, they seek working environments that foster collaboration, trust, autonomy, and creativity.

The aim of the present study is to conduct a comparative analysis of the expectations of Generation Z employees and the requirements of the labour market in the Pomeranian Voivodeship in 2025. Employers' expectations were identified on the basis of job advertisements published on online recruitment portals and collected through web scraping between July 2024 and July 2025. These data were subsequently compared with the results of the Occupational Barometer (Barometr Zawodów), the Graduate Tracking Survey (Badanie Losów Absolwentów) conducted in the Pomeranian Voivodeship, and research carried out among students of the University of Gdańsk within studies devoted to the concept of teal organizations in 2025. The study employed Exploratory Data Analysis (EDA) and Machine Learning (ML) methods.

The research enabled the identification of competency gaps among Generation Z, as well as areas of mismatch between labour market expectations and the needs and aspirations of young employees. The findings indicate the existence of significant discrepancies between the competencies sought by employers and the values, working preferences, and career expectations expressed by representatives of the youngest workforce generation.

Keywords: labour market, competency gaps, exploratory data analysis

Alicja Olejnik, Uniwersytet Łódzki; Agata Żółtaszek, Uniwersytet Łódzki

Przestrzenne zróżnicowanie wykorzystania sztucznej inteligencji w europejskich przedsiębiorstwach: podział Wschód–Zachód

W artykule analizowano przestrzenne zróżnicowanie wdrażania sztucznej inteligencji przez przedsiębiorstwa w regionach europejskich. Szczególną uwagę poświęcono cyfrowym efektom przestrzennego rozlewania się oraz utrzymującym się dysproporcjom między Europą Środkowo-Wschodnią (CEE) a pozostałą częścią kontynentu. Zastosowanie uciętego przestrzennego modelu Durbina (partial Spatial Durbin Model) umożliwia włączenie endogenicznego przestrzennego procesu dyfuzji wykorzystania AI, ale również oddzielenie wpływu własnych cech technologicznych i ekonomicznych regionu od oddziaływań przenoszonych z regionów sąsiednich. Podejście to pozwala wykazać, że współzależność przestrzenna nie jest jedynie statystyczną właściwością danych, lecz stanowi integralny element wdrażania sztucznej inteligencji w Europie.

Wyniki wskazują, że wykorzystaniu AI przez przedsiębiorstwa sprzyja przede wszystkim regionalna cyfrowa zdolność absorpcyjna, odzwierciedlona w powszechnym dostępie do Internetu, rozwiniętych kompetencjach w zakresie analityki danych, zaangażowaniu przedsiębiorstw w handel elektroniczny oraz w strukturze sektora przedsiębiorstw sprzyjającej modernizacji technologicznej. Jednocześnie

cyfrowe efekty przestrzennego rozlewania się mają charakter heterogeniczny. Silna endogeniczna zależność przestrzenna wdrażania AI wskazuje na dyfuzję technologii zachodzącą za pośrednictwem międzyregionalnych sieci produkcyjnych, wymiany wiedzy, mobilności pracowników, imitacji technologicznej oraz wspólnego otoczenia instytucjonalnego. Efekty przestrzennie opóźnionych zmiennych reprezentujących wybrane zdolności cyfrowe sugerują jednak, że sąsiadujące regiony mogą również konkurować o wykwalifikowanych pracowników, przedsiębiorstwa innowacyjne i inwestycje cyfrowe. Bliskość regionów zaawansowanych cyfrowo nie prowadzi zatem automatycznie do uzyskiwania porównywalnych korzyści.

Istotna niekorzystna pozycja regionów Europy Środkowo-Wschodniej utrzymuje się także po uwzględnieniu regionalnych zasobów innowacyjnych, gotowości cyfrowej, struktury przedsiębiorstw oraz interakcji przestrzennych. Wynik ten wskazuje, że europejski podział w zakresie wykorzystania sztucznej inteligencji wynika z szerszych asymetrii strukturalnych i instytucjonalnych, a nie wyłącznie ze zróżnicowania infrastruktury cyfrowej. Rozdzielenie bezpośrednich efektów regionalnych, międzyregionalnych efektów rozlewania oraz przestrzennych mechanizmów sprzężeń zwrotnych pozwala wykazać, dlaczego polityka skoncentrowana wyłącznie na wzmacnianiu potencjału poszczególnych regionów może okazać się niewystarczająca bez skuteczniejszych mechanizmów dyfuzji międzyregionalnej i spójności terytorialnej.

Słowa kluczowe: wdrażanie sztucznej inteligencji; Europa Środkowo-Wschodnia; regionalna luka cyfrowa; przestrzenne efekty rozlewania; regiony NUTS 2; przestrzenny model Durбина

The Geography of Enterprise AI Adoption in Europe: An East–West Divide

This study investigates the regional geography of enterprise artificial intelligence adoption in Europe, with particular attention to digital spillovers and the persistent divide between Central and Eastern Europe (CEE) and the rest of the continent. A partial Spatial Durbin Model is employed to distinguish the effects of a region's own technological and economic characteristics from those transmitted through neighbouring regions, while simultaneously accounting for the endogenous spatial diffusion of AI adoption. This approach reveals that spatial interdependence is not merely a statistical feature of the data but an integral component of the European AI landscape.

The results indicate that enterprise AI adoption is supported primarily by regional digital absorptive capacity, reflected in widespread internet access, data-analytics capabilities, e-commerce engagement and an enterprise structure conducive to technological upgrading. At the same time, digital spillovers are heterogeneous. The strong endogenous spatial dependence of AI adoption points to diffusion through interregional production networks, knowledge exchange, labour mobility, technological imitation and shared institutional environments. However, the spatially lagged effects of selected digital capabilities suggest that neighbouring regions may also compete for skilled labour, innovative firms and digital investment. Consequently, proximity to digitally advanced regions does not automatically translate into equivalent benefits for their neighbours.

A substantial CEE disadvantage persists after controlling for regional innovation resources, digital readiness, enterprise structure and spatial interactions. This finding suggests that the European AI divide reflects broader structural and institutional asymmetries rather than differences in digital

infrastructure alone. By separating direct regional effects, cross-regional spillovers and spatial feedback mechanisms, the model demonstrates why policies focused exclusively on improving individual regions may be insufficient without stronger mechanisms for interregional diffusion and cohesion.

Keywords: artificial intelligence adoption; Central and Eastern Europe; regional digital divide; spatial spillovers; NUTS 2 regions; Spatial Durbin Model

Natalia Pawelec, Uniwersytet Łódzki

Problem wyboru indeksu cenowego przy szacowaniu inflacji na bazie danych skrapowanych

W pomiarze inflacji coraz większą rolę odgrywają alternatywne źródła danych o cenach, rozwijane równolegle do tradycyjnych metod zbierania notowań. Jednym z takich źródeł są dane scrapowane, czyli ceny pozyskiwane automatycznie ze stron internetowych sprzedawców. Wystąpienie koncentruje się na tym typie danych, ponieważ umożliwia on obserwację dużej liczby cen z wysoką częstotliwością, ale zwykle nie zawiera informacji o liczbie sprzedanych produktów ani wartości sprzedaży. Celem referatu jest omówienie problemu wyboru indeksu cenowego dla danych scrapowanych, ze szczególnym uwzględnieniem nieważonych indeksów bilateralnych, łańcuchowych i multilateralnych. Przedstawione zostaną także wybrane problemy przygotowania danych, w tym filtrowanie ekstremalnych zmian cen czy brakujące obserwacje.

Słowa kluczowe: dane scrapowane, pomiar inflacji, CPI, indeksy cenowe

The problem of choosing a price index formula for estimating inflation based on web-scraped data

In inflation measurement, alternative sources of price data are playing an increasingly important role, developing alongside traditional methods of price collection. One such source is web-scraped data, i.e. prices collected automatically from retailers' websites. The presentation focuses on this type of data, as it allows a large number of prices to be observed at high frequency, but usually does not include information on the number of products sold or sales values. The aim of the paper is to discuss the problem of choosing a price index formula for web-scraped data, with particular emphasis on unweighted bilateral, chained and multilateral indices. Selected issues related to data preparation will also be presented, including the filtering of extreme price changes and missing observations.

Keywords: web-scraped data, inflation measurement, CPI, price indices

Barbara Pawełek, Uniwersytet Ekonomiczny w Krakowie; Maria Sadko, Uniwersytet Ekonomiczny w Krakowie

Zastosowanie metod nadzorowanego uczenia statystycznego w analizie wybranych wydatków gospodarstw domowych w Polsce: problem danych z dużą liczbą zer

Wydatki gospodarstw domowych są ważnym czynnikiem stymulującym wzrost gospodarczy, dlatego ich rzetelne badanie jest dla ekonomistów kwestią kluczową. Celem głównym pracy jest ocena przydatności metod nadzorowanego uczenia statystycznego do analizy wydatków gospodarstw

domowych w Polsce na wyposażenie mieszkania i prowadzenie gospodarstwa domowego. Celem dodatkowym jest identyfikacja determinant wspomnianych wydatków. W badaniu wykorzystano dane pochodzące z Badania budżetów gospodarstw domowych z 2022 roku. Zbiór danych będący podstawą przeprowadzonych analiz zawiera wiele zer, co uwzględniono w zastosowanej procedurze badawczej. Zbudowano także klasyczny model selekcji próby, którego wyniki były punktem odniesienia w ocenie przydatności metod uczenia statystycznego do analizy tytułowego zagadnienia. W celu poprawy interpretowalności wyników otrzymanych za pomocą algorytmów uczenia maszynowego zastosowano techniki wyjaśnialnej sztucznej inteligencji (eXplainable AI). Przeprowadzone wstępne analizy potwierdziły przydatność podejść bazujących na metodach nadzorowanego uczenia statystycznego do analizy danych z dużą liczbą zer. Natomiast techniki XAI umożliwiły sformułowanie użytecznych wniosków ekonomicznych. W referacie zostaną przedstawione główne wyniki przeprowadzonych badań oraz omówione zalety i ograniczenia zastosowanych podejść badawczych.

Słowa kluczowe: nadzorowane uczenie statystyczne, wydatki gospodarstw domowych, dane z dużą liczbą zer, wyjaśnialna sztuczna inteligencja

Michał Bernard Pietrzak, Politechnika Gdańska; Jarosław Krajewski, Politechnika Gdańska

Modelowanie poziomu produkcji energii słonecznej oraz energii wiatrowej w Polsce

Artykuł podejmuje problematykę modelowania poziomu produkcji energii słonecznej oraz wiatrowej w Polsce. W ostatnich latach obserwuje się systematyczny i dynamiczny wzrost produkcji energii słonecznej oraz energii wiatrowej, co przekłada się na wzrost udziału tego rodzaju energii odnawialnej w krajowym miksie energetycznym. Cel artykułu stanowi modelowanie poziomu produkcji energii słonecznej i energii wiatrowej oraz identyfikacja okresów niedoborów i nadprodukcji energii elektrycznej. W ramach realizacji celu przedstawiono specyfikację trendowo-sezonowego modelu ekonometrycznego dla każdego ze źródeł energii, a następnie na podstawie oszacowanych modeli wyznaczono prognozy produkcji energii. Wyznaczenie wartości teoretycznych modelu oraz wartości prognoz pozwoliło na identyfikację miesięcy charakteryzujących się niedoborem produkcji energii elektrycznej albo jej nadprodukcją. Niewątpliwie występowanie okresów nadprodukcji oraz niedoborów produkcji stanowi jeden z głównych mankamentów produkcji energii słonecznej oraz wiatrowej, gdzie poziom produkcji mocno uzależniony jest od krajowych warunków atmosferycznych. Na koniec przedstawiono wnioski oraz rekomendacje dotyczące konsumpcji energii odnawialnej w Polsce oraz krajowego bezpieczeństwa energetycznego. Wyniki badań stanowią wkład w dyskusję nad kierunkami dotyczące zarządzania miksem energetycznym oraz prowadzenia zrównoważonej polityki energetycznej w Polsce.

Słowa kluczowe: energia odnawialna, energia słoneczna, energia wiatrowa, krajowy miks energetyczny, bezpieczeństwo energetyczne

Modeling the level of solar energy production and wind energy production in Poland

The article addresses the issue of modeling the level of solar energy production and wind energy production in Poland. In recent years, a systematic and dynamic increase in the production of solar energy and wind energy has been observed, which translates into an increase in the share of this type

of renewable energy in the national energy mix. The aim of the article is to model the level of solar energy production and wind energy production and to identify periods of shortages and overproduction of electrical energy. As part of the implementation of the aim, the specification of a trend-seasonal econometric model for each of the energy sources was presented, and then, on the basis of the estimated models, forecasts of energy production were determined. The determination of the theoretical values of the model and the forecast values allowed the identification of months characterized by a shortage of electrical energy production or its overproduction. Undoubtedly, the occurrence of periods of overproduction and shortages of production constitutes one of the main drawbacks of solar energy and wind energy production, where the level of production is strongly dependent on national atmospheric conditions. At the end, conclusions and recommendations regarding the consumption of renewable energy in Poland and national energy security were presented. The results of the research constitute a contribution to the discussion on directions concerning the management of the energy mix and the conduct of a sustainable energy policy in Poland.

Keywords: renewable energy, solar energy, wind energy, national energy mix, energy security

Dominika Polko-Zajac, Uniwersytet Ekonomiczny w Katowicach

Kierunkowe hipotezy alternatywne w analizie korelacji między trzema zbiorami zmiennych

W artykule zaproponowano testy statystyczne służące do badania istotności zależności korelacyjnych między trzema zbiorami zmiennych. Proponowane podejście opiera się na metodzie nieparametrycznej kombinacji zależnych testów permutacyjnych (NPC), co umożliwia elastyczne i odporne wnioskowanie bez konieczności przyjmowania założeń dotyczących rozkładów badanych zmiennych. Własności proponowanych procedur testowania zostały zbadane za pomocą symulacji Monte Carlo dla różnych scenariuszy generowania danych. Uzyskane wyniki wskazują, że podejście oparte na NPC stanowi wiarygodne narzędzie weryfikacji złożonych hipotez, zwłaszcza w przypadku małych prób lub w sytuacji naruszenia klasycznych założeń. W artykule zawarto również przykład empiryczny prezentujący identyfikację istotnych zależności między zbiorami danych.

Słowa kluczowe: testy permutacyjne, funkcje łączące, złożone hipotezy alternatywne, symulacja Monte Carlo

Directional alternative hypotheses in the analysis of correlations among three sets of variables

This paper proposes statistical tests for assessing the significance of correlation relationships among three sets of variables. The proposed approach is based on the Nonparametric Combination (NPC) of dependent permutation tests, which enables flexible and robust inference without the need to assume distributional forms of the analyzed variables. The properties of the proposed testing procedures are examined through a Monte Carlo simulation study under various data-generating scenarios. The obtained results indicate that the NPC-based approach provides a reliable tool for testing complex hypotheses, particularly in the case of small sample sizes or when classical assumptions are violated. The paper also includes an empirical example illustrating the identification of significant correlations among the sets of data.

Keywords: permutation tests, combining functions, complex alternative hypotheses, Monte Carlo simulation

Aneta Ptak-Chmielewska, Szkoła Główna Handlowa w Warszawie; Małgorzata Grzywińska-Rąpca, Uniwersytet Warmińsko-Mazurski w Olsztynie

Dobrobyt finansowy, oszczędności i zadłużenie europejskich gospodarstw domowych

Ocena sytuacji finansowej gospodarstw domowych determinowana jest zarówno przez czynniki obiektywne jak i przez czynniki subiektywne. Celem artykułu była identyfikacja determinantów odróżniających gospodarstwa domowe doświadczające problemów finansowych od tych które takich trudności nie odczuwają. W pierwszym etapie dane z Eurostatu o oszczędnościach, zadłużeniu gospodarstw oraz dochodach i konsumpcji pozwoliły podzielić grupy państw UE (metodą k-średnich) na państwa posiadające problemy finansowe (wysoki udział with „great difficulty”, „with difficulty”, „with some difficulty”) i te, które nie mają problemów finansowych (wysoki udział „fairly easily”, „easily”) według subiektywnej oceny. W dalszym etapie analizy dochodów, stóp zadłużenia i oszczędności gospodarstw uzupełniły obraz sytuacji subiektywnej o sytuację obiektywną. Do analizy wykorzystano między innymi modele panelowe.

Słowa kluczowe: gospodarstwa domowe, finanse, czynniki subiektywne

Financial Well-Being, Savings, and Debts of European Households

The assessment of the financial situation of households is determined by both objective and subjective factors. The aim of the article was to identify the determinants that distinguish households experiencing financial difficulties from those that do not face such problems. In the first stage, Eurostat data on household savings, debt, income, and consumption were used to classify groups of EU countries (using the k-means method) into countries experiencing financial difficulties (a high share of responses “with great difficulty,” “with difficulty,” and “with some difficulty”) and those without financial difficulties (a high share of responses “fairly easily” and “easily”) based on subjective assessment. In a subsequent stage, analyses of household incomes, debt ratios, and savings complemented the subjective assessment with an objective perspective. The analysis employed, among other methods, panel data models.

Keywords: households, finances, subjective factors

Elżbieta Roszko-Wójtowicz, Uniwersytet Łódzki; Barbara Dańska-Borsiak, Uniwersytet Łódzki

Prognozy ruchu turystycznego w głównych miastach Polski z uwzględnieniem dynamiki i sezonowości

Turystyka miejska stanowi istotny i dynamicznie zmieniający się segment rynku turystycznego, a jej rozwój charakteryzuje się znacznym zróżnicowaniem przestrzennym i sezonowym. Celem badania jest identyfikacja wzorców sezonowości ruchu turystycznego w głównych miastach Polski oraz ocena jego przewidywanego kształtowania się w krótkim okresie. Analizę przeprowadzono dla 20 miast, obejmujących wszystkie stolice województw oraz wybrane dodatkowe ośrodki miejskie, na podstawie

miesięcznych danych dotyczących liczby turystów w latach 2017–2026. Tak szczegółowy poziom przestrzenny i czasowy umożliwia porównanie miast o odmiennej skali i funkcjach turystycznych oraz identyfikację różnic w charakterze i natężeniu sezonowości.

W badaniu wykorzystano analizę trendu oraz addytywne i multiplikatywne modele sezonowości, uwzględniając również zakłócenia związane z pandemią COVID-19. Dobór postaci modelu dla poszczególnych miast oparto na ocenie jakości prognoz ex post. Wyniki wskazują na wyraźne zróżnicowanie wzorców sezonowych pomiędzy badanymi ośrodkami. Najsilniejsza koncentracja ruchu turystycznego w miesiącach letnich występuje przede wszystkim w miastach, w których istotną rolę odgrywają walory wypoczynkowe i przyrodnicze, podczas gdy w dużych ośrodkach o bardziej zróżnicowanych funkcjach sezonowość jest na ogół słabsza. Na podstawie zidentyfikowanych trendów i wzorców sezonowych opracowano prognozy ruchu turystycznego do końca 2026 r. Wyniki dostarczają wiedzy o zróżnicowaniu czasowym turystyki miejskiej w Polsce i mogą stanowić podstawę zarówno dalszych badań porównawczych, jak i działań związanych z zarządzaniem rozwojem turystyki na poziomie miejskim.

Słowa kluczowe: turystyka miejska, ruch turystyczny, sezonowość, prognozowanie, miasta, Polska

Forecasts of Tourist Flows in Major Polish Cities: Accounting for Dynamics and Seasonality

Urban tourism represents an important and dynamically evolving segment of the tourism market, characterised by considerable spatial and seasonal variation. The aim of the study is to identify patterns of seasonality in tourist flows across major Polish cities and to assess their expected short-term development. The analysis covers 20 cities, including all regional capitals and selected additional urban destinations, using monthly data on tourist arrivals for the period 2017–2026. This detailed spatial and temporal perspective makes it possible to compare cities differing in size and tourism functions and to identify differences in the nature and intensity of seasonality.

The study employs trend analysis and additive and multiplicative seasonal models, while also accounting for the disruption caused by the COVID-19 pandemic. The model specification for individual cities is selected on the basis of ex post forecast accuracy. The results reveal substantial differences in seasonal patterns across the analysed urban destinations. The strongest concentration of tourist flows during the summer months is observed primarily in cities where recreational and natural attractions play an important role, whereas seasonality tends to be less pronounced in major urban centres with more diversified functions. Based on the identified trends and seasonal patterns, forecasts of tourist flows through the end of 2026 are developed. The findings provide new insights into the temporal differentiation of urban tourism in Poland and may support both further comparative research and evidence-based tourism management at the city level.

Keywords: urban tourism, tourist flows, seasonality, forecasting, cities, Poland

Zastosowanie analizy historii zdarzeń do badań dojrzałości technologii odnawialnych źródeł energii

Proces transformacji energetycznej gospodarki wymaga koordynacji różnych działań na poziomie krajowym i międzynarodowym, spośród których kluczowe jest wdrożenie odnawialnych źródeł energii. Kraje bazując na dostępnych zasobach naturalnych i uwarunkowaniach geograficzno-ekonomicznych mogą rozwijać technologie OZE wykorzystujące energię wiatru, słońca oraz wody. Osiągnięcie odpowiedniej efektywności wdrażanego OZE jest możliwe po uzyskaniu dojrzałości technologicznej. Kraje mogą osiągać tę dojrzałość w różnych momentach, ale są i takie, które nie uzyskują jej w ogóle. Można więc doszukać się tutaj analogii do występowania obserwacji cenzurowanych i kompletnych. Stwarza to możliwość zastosowania analizy historii zdarzeń do modelowania czasu potrzebnego do uzyskania dojrzałości kraju w zakresie pełnego wdrożenia określonego rodzaju OZE. Celem artykułu jest modelowanie hazardu osiągnięcia różnych wymiarów dojrzałości krajów w zakresie technologii energetyki fotowoltaicznej i wiatrowej. Ponadto w referacie będą zidentyfikowane czynniki, które wpływają na tempo osiągania tej dojrzałości w poszczególnych krajach. Otrzymane wyniki mogą być pomocne w ocenie transformacji energetycznej na świecie oraz mogą stanowić wsparcie dla procesu decyzji w obszarze planowania energetycznego.

Słowa kluczowe: model proporcjonalnego hazardu Coxa, OZE, fotowoltaika, energetyka wiatrowa

Application of event history analysis to renewable energy technology maturity testing

The process of energy transformation requires the coordination of various activities at the national and international levels, of which the implementation of renewable energy sources is crucial. Based on available natural resources and geographical and economic conditions, countries can develop renewable energy technologies utilizing wind, solar, and hydropower. Achieving adequate efficiency in the implementation of renewable energy is possible after achieving technological maturity. Countries may achieve this maturity at different times, but some do not achieve it at all. Therefore, an analogy can be drawn here to the occurrence of censored and complete observations. This creates the possibility of using event history analysis to model the time required for a country to achieve maturity in terms of full implementation of a given type of renewable energy. The aim of this paper is to model the hazard of achieving various dimensions of country maturity in terms of photovoltaic and wind energy technologies. Furthermore, the paper will identify factors that influence the pace of achieving this maturity in individual countries. The obtained results can be helpful in assessing the global energy transition and can support decision-making in the area of energy planning.

Keywords: Cox proportional hazards model, renewable energy, photovoltaics, wind energy

Grzegorz Sitek, Uniwersytet Ekonomiczny w Katowicach; Janusz L. Wywił, Uniwersytet Ekonomiczny w Katowicach

Testy na normalność rozkładu prawdopodobieństwa wykorzystujące funkcje momentów z próby

Problem testowania hipotezy o normalności rozkładu zmiennej badanej należy do ciągle aktualnych. O przydatności praktycznej testów na normalność rozkładu prawdopodobieństwa zależy od własności

ich statystyk testowych lecz przede wszystkim od ich mocy. W pracy zajmujemy się klasą testów, których statystyki testowe są funkcjami momentów centralnych z próby. W szczególności są to statystyki testowe które są funkcjami ocen z próby wskaźników skośności lub kurtozy rozkładu zmiennej jednowymiarowej lub wielowymiarowej. Analizowane będą nie tylko klasyczne wskaźniki asymetrii ale również ich modyfikacje. Głównym celem referatu jest symulacyjna ocena i porównanie mocy tych testów. W tym wybrano odpowiednie rozkłady alternatywne różniące się od rozkładu normalnego skośnością lub kurtozą. Przeprowadzone analizy pozwoliły na porównanie mocy testów opartych na klasycznych współczynnikach skośności lub kurtozy z mocami testów wykorzystujących zmodyfikowane miary skośności lub kurtozy.

Słowa kluczowe: Testy normalności, momenty centralne, moc testu.

Normality tests of probability distributions using sample moment functions

The problem of testing the normality hypothesis for a given variable remains relevant. The practical utility of normality tests depends on the properties of their test statistics, but primarily on their power. In this article, we address a class of tests whose test statistics are functions of the sample central moments. More precisely, these are test statistics that are functions of sample estimates of skewness or kurtosis indices for univariate or multivariate distributions. We will analyze not only the classic skewness or kurtosis indices but also their modifications. The main goal of this article is to simulate and compare the power of these tests. In this article, we select appropriate alternative distributions that differ from the normal distribution in terms of skewness or kurtosis. The analysis allows for a comparison of the power of tests using classic skewness or kurtosis indices with the power of tests using modified measures of skewness or kurtosis.

Keywords: Normality tests, central moments, power of the test.

Małgorzata Snarska, Uniwersytet Ekonomiczny w Krakowie; Jarosław Duda, Uniwersytet Jagielloński

Probabilistyczne prognozowanie dynamiki krzywej dochodowości z wykorzystaniem hierarchicznej rekonstrukcji korelacji

Celem badania jest opracowanie nieparametrycznej metody gęstościowego prognozowania dynamiki krzywej dochodowości. Zamiast założeń o rozkładzie normalnym, proponujemy Hierarchiczną Rekonstrukcję Korelacji (HCR) – estymację łącznej gęstości zmian czynników Diebold-Li jako wielomianu w bazie ortonormalnej na $[0,1]^3$. Badamy, czy HCR poprawia trafność prognoz w porównaniu z benchmarkiem gaussowskim, szczególnie w okresach kryzysowych.

Dane obejmują dzienne notowania rentowności obligacji skarbowych USA dla dziesięciu terminów zapadalności w okresie 1993–marzec 2026 ($n = 8\,127$ obserwacji), pochodące z bazy Fed/Bloomberg. Czynniki poziomu, nachylenia i krzywizny wyznaczamy modelem Diebold-Li, przy czym parametr zaniku λ estymujemy metodą wiarygodności profilowej, co – w przeciwieństwie do standardowej kalibracji $\lambda = 0,0609$ – daje optymalne $\hat{\lambda} = 0,491$ i przesuną szczyt ładunku krzywizny z 30 do 17 miesięcy. Łączna gęstość dziennych zmian trzech czynników modelowana jest przez HCR: wielomian w bazie iloczynu wielomianów Legendre'a stopnia do dziewiątego, co daje potencjalnie $10^3 = 1\,000$ współczynników wyznaczanych analitycznie jako średnie próbkowe. Sparse HCR redukuje przestrzeń modelu do mniej

niż 100 istotnych współczynników przy użyciu testu t i korekcji Bonferroniego. Adaptive HCR śledzi nietrwałość strukturalną poprzez wykładniczo ważoną estymację z parametrem zaniku pamięci; analizę toczącego się okna stosujemy do detekcji przełomów strukturalnych. Wartość ekonomiczną weryfikujemy przez: (i) model probit recesji z HCR jako regresorem, (ii) Value-at-Risk ważony gęstością HCR porównany z symulacją historyczną, (iii) prognozy przedziałowe rentowności 10-letnich obligacji metodą Monte Carlo.

Estymowane $\hat{\lambda} = 0,491$ istotnie różni się od wartości kanonicznej; prognozy gęstościowe poza próbą wyrażnie koncentrują się wokół wartości estymowanej i monotonicznie pogarszają się wraz ze zbliżaniem do $\lambda = 0,0609$. Sparse HCR redukuje bazę o ponad 90% bez istotnej utraty trafności prognoz. Analiza rolująca dokumentuje nietrwałość strukturalną w okolicach globalnego kryzysu finansowego, pandemii COVID-19 i cyklu zacieśniania 2022–2023. HCR jako regresor w modelu probit podnosi AUC wskaźnika recesji z 0,594 do 0,653. VaR ważony HCR szybciej dostosowuje się do zmian reżimu niż symulacja historyczna. Przedziały prognoz Monte Carlo dla rentowności 10-letniej są o 17% szersze niż gaussowskie, lepiej odwzorowując empiryczne grube ogony rozkładu.

HCR stanowi elastyczną, nieparametryczną alternatywę dla rozkładu gaussowskiego w modelowaniu krzywej dochodowości. Profilowo-wiarygodnościowa estymacja λ poprawia dopasowanie do danych i powinna zastąpić kanoniczną kalibrację. Sparse HCR czyni metodę praktyczną obliczeniowo, a Adaptive HCR pozwala modelować epizody nietrwałości. Wyniki wskazują, że informacja zawarta w łącznej gęstości czynników ma wymierną wartość dla prognozowania recesji i zarządzania ryzykiem rynkowym, co wzbogaca literaturę z zakresu probabilistycznego prognozowania krzywej dochodowości i statystycznych metod oceny ryzyka.

Słowa kluczowe: krzywa dochodowości, prognozowanie probabilistyczne, estymacja gęstości, model Diebold-Li, Value-at-Risk

Probabilistic Forecasting of Yield Curve Dynamics Using Hierarchical Correlation Reconstruction

This study develops a non-parametric density forecasting framework for yield curve dynamics. Rather than assuming Gaussian factor innovations, we propose Hierarchical Correlation Reconstruction (HCR), which estimates the joint density of Diebold-Li factor changes as a polynomial in an orthonormal basis on $[0,1]^3$. We ask whether HCR improves probabilistic forecast accuracy relative to Gaussian benchmarks, particularly during crisis and regime-change episodes.

The dataset comprises daily US Treasury yields at ten maturities spanning 1993 to March 2026 ($n = 8,127$ observations), sourced from the Federal Reserve and Bloomberg. Level, slope, and curvature factors are extracted via the Diebold-Li model; crucially, the decay parameter λ is estimated by profile likelihood rather than fixed at the canonical $\lambda = 0.0609$, yielding $\hat{\lambda} = 0.491$ and shifting the curvature-loading peak from 30 to 17 months. The joint density of daily changes in all three factors is modelled by HCR: a polynomial expansion in the tensor product of Legendre polynomials up to degree nine, producing up to $10^3 = 1,000$ coefficients estimated analytically as sample averages. Sparse HCR retains fewer than 100 significant coefficients via t -tests with Bonferroni correction. Adaptive HCR tracks structural non-stationarity through exponentially weighted estimation with a memory decay parameter; rolling-window analysis identifies structural breaks. Economic value is assessed via three applications: (i) a probit recession model augmented with the HCR density; (ii) HCR-weighted

Value-at-Risk compared with historical simulation; and (iii) Monte Carlo forecast intervals for the 10-year Treasury yield drawn from the estimated HCR density.

The estimated $\lambda = 0.491$ departs significantly from the canonical value; out-of-sample density forecasts peak sharply at the estimated λ and deteriorate monotonically toward $\lambda = 0.0609$. Sparse HCR reduces the basis by over 90% with negligible loss of predictive accuracy. Rolling-window analysis documents structural non-stationarity around the Global Financial Crisis, the COVID-19 pandemic, and the 2022–2023 tightening episode. Adding the HCR density as a regressor raises the recession probit AUC from 0.594 to 0.653. HCR-weighted VaR adapts more rapidly to regime changes than historical simulation. Monte Carlo forecast intervals for the 10-year yield are 17% wider than Gaussian, better reproducing the empirically fat-tailed distribution of daily yield changes.

HCR offers a flexible, non-parametric alternative to Gaussian assumptions in yield curve modelling. Profile-likelihood estimation of λ improves data fit and should replace canonical calibration in practice. Sparse HCR keeps the method computationally tractable, while Adaptive HCR accommodates regime changes. The findings demonstrate that the joint factor density contains measurable economic value for recession forecasting and market risk management, contributing to the literature on probabilistic yield curve forecasting and statistical risk assessment methods.

Keywords: yield curve forecasting, probabilistic forecasting, density estimation, Diebold-Li model, Value-at-Risk.

Piotr Szczepocki, Uniwersytet Łódzki; Ewa Feder-Sempach, Uniwersytet Łódzki; Stanislav Uryasev, Stony Brook University

Miary Drawdown Beta w ocenie ryzyka systematycznego aktywów na rynku akcji

Celem referatu jest prezentacja współczynników ryzyka systematycznego opartych na koncepcji obsunięcia kapitału (drawdown), określanych zbiorczo mianem Drawdown Beta. Miary te stanowią alternatywę dla klasycznej bety wywodzącej się z modelu CAPM i koncentrują się na zachowaniu aktywów w okresach największych spadków rynku, które są szczególnie istotne z punktu widzenia inwestorów.

W referacie omówione zostaną dwa podejścia rozwijane w literaturze przedmiotu: Conditional Drawdown-at-Risk Beta (CDaR Beta) oraz Expected Regret of Drawdown Beta (ERoD Beta). W przeciwieństwie do tradycyjnej bety, bazującej na współzmienności stóp zwrotu, proponowane wskaźniki uwzględniają zależności występujące podczas skrajnie niekorzystnych faz rynku.

Część empiryczna opiera się na dziennych danych dla spółek wchodzących w skład indeksu S&P 500 w latach 1995–2024. Przeprowadzona analiza obejmuje porównanie różnych specyfikacji współczynnika beta oraz ocenę ich przydatności w konstrukcji portfeli inwestycyjnych. Uzyskane wyniki pozwalają ocenić, czy miary Drawdown Beta dostarczają dodatkowych informacji o ryzyku aktywów względem klasycznej bety CAPM, zwłaszcza w okresach silnych spadków rynkowych.

Słowa kluczowe: Drawdown Beta, ryzyko systematyczne, obsunięcie kapitału, S&P 500, portfele inwestycyjne

Drawdown Beta Measures in the Assessment of Systematic Risk in the Stock Market

The aim of this paper is to present systematic risk measures based on the concept of drawdowns, collectively referred to as Drawdown Beta. These measures provide an alternative to the traditional beta coefficient derived from the Capital Asset Pricing Model (CAPM) and focus on asset behavior during periods of severe market declines, which are of particular importance to investors.

The paper discusses two approaches proposed in the literature: Conditional Drawdown-at-Risk Beta (CDaR Beta) and Expected Regret of Drawdown Beta (ERoD Beta). Unlike the traditional beta, which is based on the covariance of returns, these measures capture the relationships between assets and the market during extremely adverse market conditions.

The empirical analysis is based on daily data for companies included in the S&P 500 index over the period 1995–2024. The study compares different beta specifications and evaluates their usefulness in portfolio construction. The results provide evidence on whether Drawdown Beta measures offer additional information about asset risk beyond that captured by the traditional CAPM beta, particularly during periods of substantial market downturns.

Keywords: Drawdown Beta, systematic risk, drawdown, S&P 500, investment portfolios

Marcin Szymkowiak, Statistics Poland, Poznań University of Economics and Business; Kamil Wilak, Statistics Poland, Poznań University of Economics and Business

Mapowanie ubóstwa na poziomie podregionów w Polsce: podejście z wykorzystaniem wielowymiarowego modelu Faya-Herriota

Europejskie Badanie Dochodów i Warunków Życia (EU-SILC) stanowi podstawowe źródło danych w Polsce wykorzystywane przez Główny Urząd Statystyczny do publikowania wskaźników ubóstwa, zarówno w skali całego kraju, jak i na poziomie regionalnym. Sytuacja ta dotyczy również wielu innych krajów, w których rośnie zapotrzebowanie na wiarygodne informacje na temat ubóstwa w ujęciu przestrzennym. Prowadzenie właściwej polityki społecznej, zgodnej z wytycznymi polityki spójności, wymaga pomiaru ubóstwa i raportowania go na niższych poziomach agregacji przestrzennej. Tak uzyskane mapy ubóstwa wspierają istotne decyzje polityczne, w tym alokację środków rozwojowych przez rządy, ministerstwa odpowiedzialne za infrastrukturę i rozwój czy organizacje międzynarodowe, takie jak Bank Światowy.

Estymacja bezpośrednia wskaźnika zagrożenia ubóstwem (ang. At-Risk-of-Poverty Rate - AROP) na podstawie danych EU-SILC staje się jednak niewiarygodna, gdy liczebność próby w odpowiednich domenach spada poniżej pewnego progu, co ma miejsce zazwyczaj przy niskich poziomach agregacji przestrzennej. Wymaga to zastosowania pośrednich, opartych na modelach metod estymacji, które „pożyczają moc” z informacji dodatkowych, w tym danych skorelowanych z samym wskaźnikiem AROP. W Polsce liczebność próby EU-SILC jest wystarczająca do publikowania wiarygodnych oszacowań wskaźnika zagrożenia ubóstwem jedynie na poziomie krajowym oraz na poziomie NUTS 2. Oszacowania na niższych poziomach, takich jak NUTS 3 (podregiony), nie mogą być uzyskane z akceptowalną precyzją przy zastosowaniu wyłącznie estymacji bezpośredniej.

Wobec rosnącego zapotrzebowania na oszacowania wskaźnika zagrożenia ubóstwem na niższych poziomach agregacji przestrzennej, statystyka małych obszarów stanowi odpowiednie uzasadnienie metodyczne, łączące dane EU-SILC z informacjami pochodzącymi ze spisu powszechnego oraz rejestrów administracyjnych. Wśród rozmaitych podejść szczególnie dobrze sprawdza się wielowymiarowy model Faya-Herriota: zamiast modelować wskaźnik zagrożenia ubóstwem osobno model ten szacuje go w połączeniu ze skorelowanymi zmiennymi pomocniczymi, „pożyczając” siłę z odpowiednich zależności w celu poprawy precyzji uzyskiwanych oszacowań wskaźnika AROP w porównaniu z klasycznym jednowymiarowym modelem Faya-Herriota.

Głównym celem prezentacji jest omówienie zastosowania wielowymiarowego modelu Faya-Herriota w badaniu prowadzonym wspólnie przez Urząd Statystyczny w Poznaniu, Główny Urząd Statystyczny i Bank Światowy, na potrzeby opracowania wskaźnika zagrożenia ubóstwem na poziomie NUTS 3, z wykorzystaniem danych pochodzących z różnych źródeł statystycznych, tj. poziomie agregacji przestrzennej, który dotychczas nie był oficjalnie publikowany w Polsce.

Słowa kluczowe: mapowanie ubóstwa, wskaźnik zagrożenia ubóstwem, statystyka małych obszarów, wielowymiarowy model Faya-Herriota

Poverty Mapping at the Subregional Level in Poland: A Multivariate Fay-Herriot Approach

The European Survey on Income and Living Conditions (EU-SILC) remains the primary source used by Statistics Poland to report poverty indicators, both nationally and at the regional level. This is also the case in many other countries facing an increasing demand for reliable, spatially detailed poverty maps. Following a sound social strategy, aligned with cohesion policy guidelines, requires measuring poverty and reporting it at lower levels of spatial aggregation. Poverty maps at this finer resolution support key policy decisions, including the allocation of development funds by national governments, ministries responsible for infrastructure and development, or international bodies such as the World Bank. Such decisions call for the most reliable and precise estimates available, delivered at the lowest feasible level of spatial aggregation.

Direct estimation of the At-Risk-of-Poverty Rate (AROP) from EU-SILC data, however, becomes unreliable once sample sizes fall below a certain threshold in relevant domains, which is typically the case at low levels of spatial aggregation. This calls for indirect, model-based estimation methods that draw strength from auxiliary information, including data correlated with AROP itself. In Poland, EU-SILC sample sizes are sufficient to publish reliable AROP estimates only at the national level and at NUTS 2. Estimates at lower levels, such as NUTS 3 (subregions), cannot be obtained with acceptable precision using direct estimation alone.

Given the growing demand for AROP estimates at lower levels of spatial aggregation, small area estimation offers an appropriate framework, combining EU-SILC data with information from the census and administrative registers. Among the available approaches, the multivariate Fay-Herriot model is particularly well suited to this task: rather than modelling AROP in isolation, it estimates it jointly with a correlated auxiliary variable, borrowing strength from that correlation to improve the precision of the resulting AROP estimates relative to a univariate approach.

The main aim of the presentation is to discuss the application of the multivariate Fay-Herriot model in a study conducted jointly by Statistics Poland and the World Bank, aimed at producing AROP maps

at the NUTS 3 level using data from multiple statistical sources - a level of spatial aggregation not yet officially published in Poland.

Keywords: poverty mapping, at-risk-of poverty rate, small area estimation, Multivariate Fay-Herriot model

Sławomir Śmiech, Uniwersytet Ekonomiczny w Krakowie; Lilia Karpińska, Uniwersytet Ekonomiczny w Krakowie; Kornelia Kłopecka-Miętka, Uniwersytet Ekonomiczny w Krakowie

Lokalne profile ubóstwa energetycznego na obszarach peryferyjnych i kurczących się demograficznie: analiza klas ukrytych w trzech gminach Dolnego Śląska

Obszary peryferyjne i kurczące się demograficznie są szczególnie narażone na współwystępowanie niskich dochodów, starzenia się ludności, pogarszającego się stanu zasobu mieszkaniowego oraz ograniczonego dostępu do usług publicznych. Celem badania jest identyfikacja lokalnych profili ubóstwa energetycznego oraz ocena ich związku ze zdrowiem i korzystaniem z pomocy społecznej. Analizę przeprowadzono na podstawie danych ankietowych obejmujących 3355 mieszkańców trzech gmin Dolnego Śląska, reprezentujących odmienne konteksty lokalne: postindustrialny, uzdrowiskowo-turystyczny i wiejski.

Zastosowanie analizy klas ukrytych pozwoliło wyodrębnić trzy profile: relatywnie niskiej deprywacji z umiarkowaną presją kosztową, deprywacji mieszkaniowej i ciepłej oraz wyrzeczeń ekonomicznych związanych z ograniczaniem wydatków na żywność i leczenie. Dalsze analizy, uwzględniające niepewność przypisania do klas, wykazały, że osoby należące do klasy deprywacji mieszkaniowej i ciepłej najczęściej deklarowały pogorszenie stanu zdrowia w okresie zimowym oraz korzystanie z pomocy społecznej. Podwyższone prawdopodobieństwo korzystania z pomocy społecznej odnotowano również w klasie wyrzeczeń ekonomicznych.

Wyniki nie wskazują na prostą hierarchię, w której gmina postindustrialna jest zawsze najbardziej zagrożona. Pokazują natomiast, że nawet w obrębie szerszego obszaru peryferyjnego i depopulacyjnego ubóstwo energetyczne przyjmuje odmienne lokalne formy. Polityka publiczna powinna zatem łączyć wsparcie dochodowe z modernizacją mieszkań, ochroną zdrowia i lokalnie dopasowaną pomocą społeczną.

Słowa kluczowe: ubóstwo energetyczne, obszary peryferyjne, depopulacja, analiza klas ukrytych, zdrowie, pomoc społeczna, Dolny Śląsk

Local Profiles of Energy Poverty in Peripheral and Demographically Shrinking Areas: A Latent Class Analysis of Three Municipalities in Lower Silesia

Peripheral and demographically shrinking areas are particularly exposed to the combined effects of low incomes, population ageing, deteriorating housing conditions, and limited access to public services. This study identifies local profiles of energy poverty and examines their associations with winter health deterioration and the use of social assistance. The analysis is based on survey data from 3,355 residents of three municipalities in Lower Silesia representing contrasting local contexts: post-industrial, spa-tourism, and rural.

Latent class analysis identified three distinct profiles: relatively low deprivation with moderate cost pressure, housing and thermal deprivation, and economic sacrifice involving reductions in food and medical spending. Subsequent analyses accounting for classification uncertainty showed that respondents in the housing and thermal deprivation class had the highest probability of reporting poorer health during winter and using social assistance. The economic-sacrifice class was also associated with a higher probability of social-assistance use.

The results do not reveal a simple hierarchy in which the post-industrial municipality is uniformly the most disadvantaged. Instead, they show that energy poverty takes different forms even within a broader peripheral and shrinking area. Effective policy should therefore combine income support with housing renovation, health protection, and locally tailored social assistance.

Keywords: energy poverty, peripheral areas, demographic shrinkage, latent class analysis, health, social assistance, Lower Silesia

Natalia Torlińska-Walkowiak, Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu; Katarzyna Anna Majewska, Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu; Anna Sowińska, Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu; Andrzej Kędzia, Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu; Justyna Opydo-Szymaczek, Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu

Precyzja poprzez dopasowanie: wpływ niskorosłości na odontogenezę i osteogenezę

Ocena wieku zębowego i kostnego w odniesieniu do wieku chronologicznego stanowi istotny element diagnostyki rozwoju dzieci. U pacjentów z izolowanym niedoborem hormonu wzrostu (GHD) analiza tych zależności pozwala na lepsze zrozumienie charakteru opóźnienia dojrzewania oraz monitorowanie skuteczności leczenia.

Celem badania była analiza korelacji wieku zębowego, kostnego i chronologicznego u dzieci z GHD z wykorzystaniem dopasowania propensity score.

Badaniem objęto 48 dzieci (30 chłopców, 18 dziewcząt) w wieku $10,57 \pm 2,24$ roku, spełniających kryteria włączenia spośród 93 zakwalifikowanych pacjentów. Przeprowadzono ocenę kliniczną i biochemiczną. Wiek zębowy określano na podstawie badania klinicznego, natomiast wiek kostny oceniano na podstawie radiogramów ręki i nadgarstka metodą Greulich'a i Pyle'a. Metodą propensity score dopasowano do grupy badanej grupę kontrolną i przeprowadzono szereg porównań i zbadano zależności.

Wykazano istotne statystycznie korelacje pomiędzy wiekiem zębowym, kostnym i chronologicznym ($p < 0,05$). Opóźnienie wieku zębowego było istotnie większe u dzieci z GHD w porównaniu do populacji zdrowej ($-0,37$ vs $0,28$; $p = 0,0008$). Opóźnienie wieku kostnego było bardziej nasilone niż opóźnienie wieku zębowego i wynosiło średnio $1,73$ roku ($p < 0,0001$), przy czym było większe przed rozpoczęciem leczenia niż w jego trakcie.

U dzieci z GHD występują istotne rozbieżności pomiędzy wiekiem kostnym, zębowym i chronologicznym, które zmniejszają się w trakcie terapii hormonem wzrostu. Uwzględnienie tych różnic jest szczególnie ważne w planowaniu leczenia stomatologicznego oraz interdyscyplinarnej opiece nad pacjentem.

Słowa kluczowe: Izolowany niedobór hormonu wzrostu, hormon wzrostu, wiek, zębowy, wiek kostny, propensity score

Precision Through Matching: The Impact of Short Stature on Odontogenesis and Osteogenesis

Assessment of dental and skeletal age in relation to chronological age is an important component of evaluating growth and development in children. In patients with isolated growth hormone deficiency (GHD), analysis of these relationships provides a better understanding of the nature of delayed maturation and facilitates monitoring of treatment efficacy.

The aim of the study was to analyze the correlations between dental, skeletal, and chronological age in children with GHD using propensity score matching.

The study included 48 children (30 boys and 18 girls), aged 10.57 ± 2.24 years, who met the inclusion criteria among 93 eligible patients. Clinical and biochemical assessments were performed. Dental age was determined based on clinical examination, while skeletal age was assessed using hand and wrist radiographs according to the Greulich and Pyle method. A control group was matched to the study group using propensity score matching, followed by a series of comparative analyses and assessment of the relationships between the variables.

Statistically significant correlations were found between dental, skeletal, and chronological age ($p < 0.05$). Dental age delay was significantly greater in children with GHD compared with the healthy population (-0.37 vs. 0.28 ; $p = 0.0008$). Skeletal age delay was more pronounced than dental age delay and averaged 1.73 years ($p < 0.0001$). The skeletal age delay was greater before the initiation of treatment than during treatment.

Children with GHD show significant discrepancies between skeletal, dental, and chronological age, which decrease during growth hormone therapy. Consideration of these differences is particularly important when planning dental treatment and providing interdisciplinary care for patients.

Keywords: Isolated growth hormone deficiency, growth hormone, age, dental age, skeletal age, propensity score

Joanna Trębska, Uniwersytet Łódzki; Witold Śmigielski, Uniwersytet Łódzki

Wpływ aktywności fizycznej osób w późniejszej dorosłości na ich aktywność zawodową

Referat stanowi wstępny etap badań nad znaczeniem aktywności fizycznej podejmowanej przez osoby w późniejszej dorosłości (tutaj w wieku 55-74 lata). Analiza empiryczna opiera się na wynikach Wieloośrodkowego Ogólnopolskiego Badania Stanu Zdrowia Ludności przeprowadzonego na przełomie 2025 i 2026 roku na reprezentatywnej próbie ponad 9 tys. dorosłych Polaków, w tym 3,5 tys. w wieku 55-74 lata. Wstępne wyniki wskazują na istotny wpływ aktywności fizycznej, szczególnie o dużej intensywności, na aktywność zawodową osób tuż przed osiągnięciem wieku emerytalnego, a także w pierwszej fazie okresu emerytalnego.

Słowa kluczowe: aktywność fizyczna, aktywność zawodowa, ekonomia seniorów, regresja logistyczna

The Effect of Physical Activity in Later Life on Employment Activity

This paper presents the preliminary stage of a research project examining the role of physical activity in the lives of individuals in later adulthood, defined here as adults aged 55–74 years. The empirical analysis is based on data from the Multi-Center Nationwide Study of the Health Status of the Polish Population, conducted between 2025 and 2026 on a representative sample of 9,425 adult Poles, including 3,566 respondents aged 55–74 years.

The preliminary findings indicate a significant relationship between physical activity and labor market participation among older adults. In particular, high-intensity physical activity appears to have a positive impact on the professional activity of individuals approaching retirement age, as well as those in the early stages of retirement. These results suggest that maintaining an active lifestyle may contribute to extending labor market engagement and promoting active aging.

Keywords: physical activity, economics of elderly, labor market participation, active aging, logistic regression

Joanna Trzęsiok, Uniwersytet Ekonomiczny w Katowicach

Kiedy model odpływa? O monitorowaniu dryfu modeli uczenia maszynowego

Dynamiczny rozwój metod uczenia maszynowego oraz sztucznej inteligencji sprawił, że modele predykcyjne są coraz powszechniej wykorzystywane do wspomagania decyzji w wielu obszarach życia społecznego i gospodarczego. Wraz z rosnącą liczbą zastosowań wzrasta jednak znaczenie zagadnienia zaufania do modeli, rozumianego nie tylko jako ich wysoka skuteczność predykcyjna, lecz także jako zapewnienie odpowiedniej jakości danych, wyjaśnialności podejmowanych decyzji, transparentności działania, odporności na zmiany oraz zgodności z obowiązującymi regulacjami prawnymi, w tym z wymogami AI Act. W konsekwencji ocena modeli uczenia maszynowego nie może ograniczać się wyłącznie do analizy klasycznych mierników jakości uzyskiwanych na etapie ich budowy. Coraz większego znaczenia nabiera monitorowanie modeli po ich wdrożeniu, pozwalające na wykrywanie zmian mogących prowadzić do pogorszenia jakości ich działania.

Jednym z najważniejszych wyzwań związanych z eksploatacją modeli wykorzystujących metody uczenia maszynowego jest zjawisko dryfu (drift), rozumiane jako zmiana właściwości danych lub zależności pomiędzy zmiennymi. Zmiany te mogą prowadzić do stopniowej utraty trafności predykcji, mimo że model w momencie wdrożenia charakteryzował się wysoką jakością.

Celem referatu jest przedstawienie problematyki monitorowania dryfu modeli uczenia maszynowego w środowiskach produkcyjnych. Omówione zostaną podstawowe rodzaje dryfu oraz najczęściej stosowane podejścia do jego wykrywania i monitorowania, ze szczególnym uwzględnieniem mierników oraz metod oceny zmian zachodzących w modelach i danych. W dalszej części zaprezentowany zostanie przegląd wybranych narzędzi wspierających monitorowanie modeli po wdrożeniu, ze wskazaniem ich podstawowych funkcjonalności oraz możliwości praktycznego wykorzystania.

Słowa kluczowe: dryf, monitorowanie dryfu, mierniki jakości, machine learning

When the Model Starts to Drift: On Monitoring Machine Learning Model Drift

The dynamic development of machine learning and artificial intelligence has led to the widespread adoption of predictive models to support decision-making across numerous areas of social and economic life. As the number of applications continues to grow, increasing attention has been paid to the issue of trustworthiness of machine learning models. This concept extends beyond high predictive performance to include data quality, explainability of model decisions, transparency, robustness to changing conditions, and compliance with legal regulations, including the requirements of the AI Act. Consequently, the evaluation of machine learning models can no longer be limited to the analysis of traditional performance metrics obtained during model development. Increasing importance is therefore attached to post-deployment model monitoring, which enables the detection of changes that may lead to a deterioration in model performance.

One of the major challenges associated with the deployment of machine learning models is model drift, understood as changes in data characteristics or in the relationships between variables. Such changes may gradually reduce predictive performance, even when the model achieved high accuracy at the time of deployment.

The aim of the presentation is to discuss the problem of monitoring machine learning model drift in production environments. The presentation will introduce the main types of drift and the most commonly used approaches to their detection and monitoring, with particular emphasis on metrics and methods for assessing changes in both data and model behaviour. It will also provide an overview of selected tools supporting post-deployment model monitoring, highlighting their key functionalities and potential practical applications.

Keywords: drift, drift monitoring, quality metrics, machine learning

Marek Walesiak, Uniwersytet Ekonomiczny we Wrocławiu; Grażyna Dehnel, Uniwersytet Ekonomiczny w Poznaniu

Analiza wrażliwości metody hybrydowej w porządkowaniu liniowym

Celem artykułu jest analiza wrażliwości metody hybrydowej w porządkowaniu liniowym. W metodzie hybrydowej porządkowanie liniowe poprzedzone jest zastosowaniem skalowania wielowymiarowego w przestrzeni dwuwymiarowej. Analizę wrażliwości scharakteryzowano z punktu widzenia wyboru procedury skalowania wielowymiarowego obejmującej wybór metody normalizacji, miary odległości i modelu skalowania. W skalowaniu wielowymiarowym następuje utrata informacji związana z przejściem konfiguracji obiektów z pierwotnej przestrzeni wielowymiarowej w przestrzeń dwuwymiarową. Dla tych procedur konieczna jest zatem analiza wyników skalowania wielowymiarowego w przestrzeni dwuwymiarowej pokazująca rozkład błędów związanych z rozmieszczeniem poszczególnych obiektów. Najlepsza jest sytuacja, w której rozkład błędów jest jednostajny (każdy z obiektów rozmieszczony jest w przestrzeni dwuwymiarowej z tym samym błędem). Im bardziej w skalowaniu wielowymiarowym wyniki rozkładu błędów odbiegają od rozkładu jednostajnego tym gorsze otrzymuje się wyniki porządkowania liniowego. Może się to przejawiać istotnymi zmianami w rankingu obiektów, koncentracji lub rozproszeniu wartości miary agregatowej na podstawie, której przeprowadza się porządkowanie liniowe. Pomocnymi narzędziami analizy są tutaj współczynnik rho Spearmana oraz wykresy pudełkowe.

Słowa kluczowe: analiza wrażliwości, miary agregatowe, skalowanie wielowymiarowe, program R, podejście hybrydowe

Sensitivity Analysis of the Hybrid Method in Linear Ordering

The aim of the article is to analyse the sensitivity of the hybrid method in linear ordering. In the hybrid method, linear ordering is preceded by the use of multidimensional scaling in two-dimensional space. The sensitivity analysis is characterized from the point of view of the choice of multidimensional scaling procedure including the selection of the normalization method, distance measure, and scaling model. In multidimensional scaling, there is a loss of information associated with the transition of object configurations from the original multidimensional space to two-dimensional space. For these procedures, it is therefore necessary to analyse the results of multidimensional scaling in two-dimensional space showing the distribution of errors related to the placement of individual objects. The best situation is when the distribution of errors is uniform (each of the objects is distributed in two-dimensional space with the same error). The more the results of the error distribution differ from the uniform distribution in multidimensional scaling, the worse the results of linear ordering are obtained. This can manifest itself in significant changes in the ranking of objects, concentration, or dispersion of the values of the aggregate measure on the basis of which the linear ordering is performed. Helpful analysis tools here are Spearman's rho coefficient and boxplots.

Keywords: sensitivity analysis, aggregate measures, multidimensional scaling, R program, hybrid approach

Janusz Wątroba, StatSoft Polska

Prezentacja nowego rozwiązania wspomagającego nauczanie statystyki w programie Statistica

Zestaw dydaktyczny został zaprojektowany w celu wsparcia nauczania metod statystycznych na różnych kierunkach studiów. Rozwiązanie jest adresowane do wykładowców prowadzących zajęcia dydaktyczne i szkolenia warsztatowe, a także o studentach oraz uczestnikach kursów, którzy chcą uczyć się statystyki w praktyczny, uporządkowany sposób. Zestaw pozwala płynnie realizować zajęcia krok po kroku, zgodnie z logicznie zaplanowaną ścieżką pracy, dzięki czemu uczestnicy przechodzą przez cały proces analizy statystycznej w sposób spójny i zrozumiały — od przygotowania danych, przez dobór odpowiednich metod, aż po interpretację wyników. Dzięki uporządkowanej strukturze i jasno poprowadzonej sekwencji działań. Zestaw dydaktyczny ułatwia organizację pracy na zajęciach, ogranicza liczbę błędów oraz zmniejsza czas poświęcany na kwestie techniczne, pozwalając skupić się na rozumieniu metod i wyciąganiu wniosków. W efekcie wspiera nauczanie w formule praktycznej, zwiększa samodzielność uczestników i poprawia efektywność pracy w trakcie warsztatów. Moduł będzie dostępny w wersji polskiej i angielskiej co umożliwi poszerzenie zakresu jego zastosowania.

W trakcie prezentacji zostanie zaprezentowany moduł wspierający zastosowanie metod wnioskowania statystycznego.

Słowa kluczowe: wspomaganie nauczania statystyki, oprogramowanie statystyczne

The Didactic Bundle is designed to support the teaching of statistical methods in various fields of study. The solution is addressed to lecturers conducting didactic classes and workshop training, as well as to students and course participants who want to learn statistics in a practical, structured way. The kit allows you to smoothly carry out the classes step by step, according to a logically planned work path, thanks to which participants go through the entire statistical analysis process in a coherent and understandable way – from data preparation, through the selection of appropriate methods, to the interpretation of the results. Thanks to its structured structure and clear sequence of actions. The teaching kit makes it easier to organize your work, reduce the number of mistakes and reduce the time spent on technical issues, allowing you to focus on understanding methods and drawing conclusions. As a result, it supports teaching in a practical formula, increases the independence of participants and improves the effectiveness of work during workshops. The module will be available in Polish and English, which will allow for expanding the scope of its application.

During the presentation, a module supporting the use of statistical inference methods will be presented.

Keywords: supporting teaching statistics, statistical software

Janusz Wątroba, StatSoft Polska

Hierarchiczne modele mieszane – idea i zastosowania

Jednym z kluczowych założeń modeli liniowych jest niezależność obserwacji (pomiarów). Założenie to nie jest spełnione w przypadku pomiarów dokonywanych wielokrotnie na tych samych jednostkach (pomiarów powtarzanych) oraz w sytuacji gdy jednostki tworzą strukturę hierarchiczną (zagnieżdżoną). W obydwu przypadkach należy zastosować tzw. hierarchiczne modele mieszane, które w odróżnieniu od klasycznych modeli liniowych (np. model regresji liniowej czy modele analizy wariancji) dostarczają dokładniejsze oceny badanych efektów, zwiększają moc statystyczną oraz nie powodują wzrostu błędów I rodzaju.

Modele te znajdują zastosowanie w różnych obszarach badań empirycznych dlatego też w literaturze można spotkać się z różnymi nazwami, np. modele wielopoziomowe, modele mieszane, liniowe modele mieszane, modele z efektami losowymi, modele komponentów wariancyjnych czy modele split-plot.

Typowe przykłady zastosowań można znaleźć w badaniach pedagogicznych (gdzie badane jednostki, np. uczniowie są zgrupowane w klasach a klasy w szkołach) czy też w badaniach medycznych (gdzie pacjenci mogą być leczeni na różnych oddziałach a oddziały są zlokalizowane w różnych szpitalach).

Słowa kluczowe: hierarchiczne modele mieszane, efekty stałe i losowe

Hierarchical mixed models – idea and applications

One of the key assumptions of linear models is the independence of observations (measurements). This assumption is not fulfilled in the case of measurements made repeatedly on the same units (repeated measurements) and in a situation where the units form a hierarchical (nested) structure. In both cases, the so-called hierarchical mixed models should be used, which, unlike classic linear models (e.g. linear regression model or variance analysis models), provide more accurate assessments of the investigated effects, increase statistical power and do not cause an increase in type I errors.

These models are used in various areas of empirical research, which is why various names can be found in the literature, e.g. multilevel models, mixed models, linear mixed models, random effect models, variance component models or split-plot models.

Typical examples of applications can be found in pedagogical research (when the individuals studied, e.g. pupils are grouped in classes and classes in schools) or in medical research (where patients can be treated in different departments and the wards are located in different hospitals).

Keywords: hierarchical mixed models, fixed and random effects

Janusz Wywił, Uniwersytet Ekonomiczny w Katowicach

O estymacji przeciętnych w populacji i jej domenach

Zakłada się, że populacja jest podzielona na rozłączne i niepuste podzbiory zwane domenami. Zakłada się, że w próbie są obserwowane wartości zmiennych pomocniczej i badane, a w populacji tylko wartości zmiennej pomocniczej. W celu estymacji przeciętnych w populacji i domenach sugeruje się wykorzystanie modelu nadpopulacji. Przeciętne są estymowane na podstawie metody największej wiarygodności. W szczególności przy założeniu modelu regresyjnego nadpopulacji wyprowadzono estymatory typu ilorazowego i regresyjnego. Rozważania są zilustrowane przykładem empirycznym.

Słowa kluczowe: estymacja średnich, metoda największej wiarygodności, estymatory ilorazowe i regresyjne

On the estimation of averages in population and its domains

The population is assumed to be divided into disjoint and non-empty subsets called domains. It is assumed that the sample contains observed values of the auxiliary variable and variable under study, while the population contains only the values of the auxiliary variable. To estimate population and domain averages, it is suggested to use a superpopulation model. Averages are estimated using the maximum likelihood method. In particular, assuming the regression model of superpopulation, ratio and regression estimators are derived. The discussion is illustrated with an empirical example.

Keywords: estimation, maximum likelihood method, ratio and regression estimators

Artur Zaborski, Uniwersytet Ekonomiczny we Wrocławiu; Marcin Pełka, Uniwersytet Ekonomiczny we Wrocławiu;

Metody skalowania wielowymiarowego danych symbolicznych

Skalowanie wielowymiarowe danych symbolicznych ma na celu zaprezentowanie związków, relacji między obiektami symbolicznymi z przestrzeni m -wymiarowej w przestrzeni r -wymiarowej ($r < m$). Obiekty symboliczne w przestrzeni r wymiarowej, gdzie zwykle $r=2$, są zwykle przedstawione jako prostokąty (rectangle model) (zob. Gorenen i in. 2005; Groenen i in. 2006; Brito i Dias 2022) lub okręgi (circle model) (Brito i Dias 2022). Większość metod skalowania danych symbolicznych wymaga jako danych wejściowych albo odległości w formie przedziałów liczbowych. Macierz ta może być wynikiem

ocen, zastosowania wielu odległości między obiektami i następnie ich agregacji lub wynikać z tablicy danych symbolicznych w której obiekty opisywane są przez zmienne symboliczne interwałowe.

W artykule dokonano przeglądu i porównania różnych podejść w skalowaniu wielowymiarowym danych symbolicznych, wskazano zalety i wady poszczególnych metod. W części empirycznej zastosowano metody Interscal, SymScal, I-Scal oraz histogramowe skalowanie wielowymiarowe.

Literatura:

Bock, H.-H., Diday, E. (Eds.) (2000). Analysis of symbolic data. Explanatory methods for extracting statistical information from complex data. Berlin-Heidelberg: Springer Verlag.

Brito, P., Dias, S. (Eds.) (2022). Analysis of distributional data. CRC Press, Boca Raton.

Groenen, P. J. F., Winsberg, S., Rodríguez, O., Diday, E. (2005). Multidimensional scaling of interval dissimilarities. Econometric Report, 2005–15, Rotterdam: Erasmus University.

Groenen, P. J. F., Winsberg, S., Rodríguez, O., Diday, E. (2006). I-Scal: multidimensional scaling of interval dissimilarities. Computational Statistics and Data Analysis, 51, 360–378.

Słowa kluczowe: skalowanie wielowymiarowe, dane symboliczne interwałowe, dane symboliczne histogramowe

Multidimensional scaling methods for symbolic data

Symbolic multidimensional scaling aims to represent relationships between objects from m -dimensional space in r -dimensional space ($r < m$). In the r -dimensional space, where $r=2$, symbolic objects are represented by rectangles (using rectangle model - see Groenen et. al. 2005, 2006; Brito and Dias 2022) or circles (using circle model - see Brito and Dias 2022). All multidimensional scaling methods for symbolic data require interval-valued distance matrix as the input. This matrix can be obtained by aggregating different judgements (opinions) or by calculating interval-valued distance (where symbolic objects are described by interval-valued symbolic variables). The paper presents and compares different multidimensional scaling methods for symbolic data. It presents their strengths and weaknesses. The empirical part of the paper presents results of the I-Scal, Interscal, Symscal and also histogram multidimensional scaling.

Bibliography:

Bock, H.-H., Diday, E. (Eds.) (2000). Analysis of symbolic data. Explanatory methods for extracting statistical information from complex data. Berlin-Heidelberg: Springer Verlag.

Brito, P., Dias, S. (Eds.) (2022). Analysis of distributional data. CRC Press, Boca Raton.

Groenen, P. J. F., Winsberg, S., Rodríguez, O., Diday, E. (2005). Multidimensional scaling of interval dissimilarities. Econometric Report, 2005–15, Rotterdam: Erasmus University.

Groenen, P. J. F., Winsberg, S., Rodríguez, O., Diday, E. (2006). I-Scal: multidimensional scaling of interval dissimilarities. Computational Statistics and Data Analysis, 51, 360–378.

Keywords: multidimensional scaling, interval-valued symbolic data, symbolic histogram data

Dobór współczynnika korelacji do selekcji zmiennych diagnostycznych w procedurach taksonometrycznych

Celem referatu jest zwrócenie uwagi na dyskusyjność apriorycznego stosowania współczynnika korelacji liniowej Pearsona do wyboru zmiennych diagnostycznych w ramach metod taksonometrycznych. Paradoksalnie bowiem, pomimo fundamentalnego znaczenia etapu doboru zmiennych diagnostycznych dla wyników metod taksonometrycznych, w literaturze niemal nie podejmuje dyskusji nad zasadnością wyboru współczynnika korelacji wykorzystywanego w tym celu. Dobór zmiennych diagnostycznych do modelu taksonometrycznego musi odbywać się poprzez badanie związków korelacyjnych. Ażeby jednak wartości współczynnika korelacji liniowej Pearsona charakteryzowały się wiarygodnością, analizowane cechy powinny odznaczać się symetrią rozkładu oraz brakiem występowania obserwacji odstających. Jednak badanie ich symetrii oraz kurtozy można zastąpić badaniem normalności. Oczywiście badanie normalności cech nie ma na celu wnioskowania statystycznego, gdyż metody taksonometryczne są częścią tzw. deterministycznego (opisowego) podejścia w statystyce, a więc niezakładającego losowego wektora zmiennych. W referacie przeprowadzono badania symulacyjne, określające ile oraz jakie zmienne zostają wybrane do budowania modelu taksonometrycznego w zależności od tego, jaki współczynnik korelacji zostaje użyty, a także jaka metoda doboru zmiennych do modelu zostanie wykorzystana. Łącznie przeprowadzono więc dziewięć wariantów, powstałych z łączenia wykorzystywania współczynnika korelacji liniowej Pearsona, korelacji rang Spearmana oraz korelacji rang tau-b Kendalla z metodą parametryczną Hellwiga, jego pozycyjną odmianą oraz metodą odwróconej macierzy Maliny i Zeliasia. Warianty te zostały obliczone w odmiennych scenariuszach, podobnych w zakresie liczby cech oraz obserwacji, lecz różniących się w zakresie występowania wielowymiarowego oraz jednowymiarowych rozkładów normalnych. Ważnym zastrzeżeniem jest, iż generowanie liczb pseudolosowych powoduje, że tak wytworzone cechy charakteryzują się niskim stopniem skorelowania, dlatego w ramach analizowanych metod doboru zmiennych diagnostycznych do metod taksonometrycznych obniżono wartości krytyczne. Niniejszy referat składa się z dwóch części – teoretycznej i praktycznej. W części teoretycznej przedstawia się ważność problematyki wyboru odpowiedniego współczynnika korelacji przy budowaniu agregacyjnych mierników oraz przedstawia się analizowane zagadnienie w świetle literatury przedmiotu. W części praktycznej zaś za pomocą metod symulacyjnych wykazuje się, jak stosowanie różnych współczynników korelacji wraz z odmiennymi metodami doboru cech wpływa na ostateczny zbiór wybranych zmiennych diagnostycznych. Tym zaś czynnikiem, jaki różnicuje poszczególne scenariusze jest charakteryzowanie się – lub też nie – przez cechy wielowymiarowym lub jednowymiarowym rozkładem normalnym. Uzyskane wyniki pozwalają nie tylko udowodnić różnorodność wyników w zakresie doboru zmiennych do modelu taksonometrycznego przy stosowaniu różnych współczynników korelacji, lecz przede wszystkim – zależność tych różnic od charakteryzowania – bądź nie – rozkładem normalnym przez badane cechy. Umożliwiają także wysnucie konkluzji mających praktyczne zastosowanie w procedurze doboru zmiennych do modelu w taksonometrii.

Słowa kluczowe: taksonometria, współczynnik korelacji liniowej Pearsona, metody symulacyjne.

Selecting an Appropriate Correlation Coefficient for the Selection of Diagnostic Variables in Taxonometric Procedures

The aim of this presentation is to draw attention to the questionable practice of the a priori use of Pearson's product-moment correlation coefficient for the selection of diagnostic variables in taxonometric methods. Paradoxically, despite the fundamental importance of diagnostic variable selection for the results of taxonometric analyses, the literature contains very little discussion regarding the appropriateness of the correlation coefficient employed for this purpose. The selection of diagnostic variables for a taxonometric model must be based on the analysis of correlations among variables. However, for Pearson's product-moment correlation coefficient to provide reliable results, the analysed variables should exhibit approximately symmetric distributions and be free from outliers. The assessment of skewness and kurtosis may, however, be replaced by testing for normality. It should be emphasized that normality testing is not intended for statistical inference in this context, since taxonomic methods belong to the deterministic (descriptive) approach to statistics and therefore do not assume a random vector of variables. The study employs simulation methods to determine the number and type of diagnostic variables selected for the construction of a taxonometric model depending on both the correlation coefficient used and the variable selection method applied. In total, nine variants were analysed by combining three correlation measures - Pearson's product - moment correlation coefficient, Spearman's rank correlation coefficient, and Kendall's tau-b rank correlation coefficient - with three variable selection methods: Hellwig's parametric method, its positional counterpart, and the inverse matrix method proposed by Malina and Zeliaś. The variants were evaluated under different simulation scenarios that were comparable with respect to the number of variables and observations but differed in terms of the presence or absence of univariate and multivariate normal distributions. It should be noted that pseudorandom number generation produces variables with relatively low levels of correlation. Therefore, the critical threshold values adopted in the analysed variable selection methods were reduced to ensure the applicability of the procedures. The presentation consists of two parts. The theoretical part discusses the importance of selecting an appropriate correlation coefficient in the construction of composite taxonomic measures and reviews the existing literature on this issue. The empirical part employs simulation methods to demonstrate how different correlation coefficients, combined with different variable selection procedures, influence the final set of selected diagnostic variables. The distinguishing factor across the simulation scenarios is whether the analysed variables follow univariate and/or multivariate normal distributions. The obtained results demonstrate not only that different correlation coefficients may lead to different sets of selected diagnostic variables but, more importantly, that these differences depend on whether the analysed variables satisfy the assumption of normality. The findings also provide practical recommendations for the procedure of diagnostic variable selection in taxonometric modelling.

Keywords: taxonometric analysis, Pearson's product-moment correlation coefficient, simulation methods

Metoda różnicy w różnicach z uwzględnieniem intensywności oddziaływania czynnika w analizie rynku pracy w Polsce

Metoda różnicy w różnicach jest jedną z podstawowych metod analizy wyników quasi-eksperymentów związanych z badaniem przyczynowości zjawisk o charakterze ekonomicznym. Klasyczny model metody różnicy w różnicach uwzględnia próby pochodzące z dwóch populacji: grupy badanej i grupy kontrolnej oraz dwa okresy czasowe: przed i po wystąpieniu oddziaływania czynnika. Model ten pozwala na wyznaczenie wartości estymatora, który stanowi ocenę przeciętnego efektu wpływu ustalonego czynnika na zmienną wynikową. Współcześnie rozwijane są alternatywne ujęcia metody różnicy w różnicach. Jednym z nich jest uwzględnienie zróżnicowanego poziomu intensywności oddziaływania czynnika potencjalnie wpływającego na zmienną wynikową. Celem referatu jest zaprezentowanie możliwości aplikacyjnych metody różnicy w różnicach z uwzględnieniem intensywności oddziaływania czynnika w analizie rynku pracy w Polsce. W badaniu zostaną wykorzystane dane dotyczące wybranych aspektów rynku pracy w powiatach Polski.

Słowa kluczowe: quasi-eksperyment, metoda różnicy w różnicach, rynek pracy

The difference-in-differences method, considering the intensity of the factor's impact, in the analysis of the labour market in Poland

The difference in differences method is one of the fundamental techniques used to analyse the outcomes of quasi experiments related to the study of causal relationships in economic phenomena. The classical model considers samples drawn from two populations—the treated (or exposed) group and the control group—and two time periods: before and after the occurrence of the factor's impact. This model allows estimating a parameter whose value represents the average effect of a given factor on the outcome variable.

Contemporary approaches to the difference in differences method have been extended in various directions. One such extension involves accounting for a differentiated level of the factor's intensity of impact on the outcome variable under analysis. The aim of this paper is to present the potential applications of the difference in differences method, taking into account the intensity of a factor's impact, in the analysis of the labour market in Poland. The study will use data on selected aspects of the labour market at the Polish county level.

Keywords: difference in differences method, labour market, quasi experiments.

Książka streszczeń on-line

