

German Polish Seminar on Data Analysis and Applications (GPSDAA) Stralsund, Germany – v1.0

Convener: Gero Szepannek and André Grüning, University of Applied Sciences

2026-09-11

Outer Schedule

The GPSDAA takes place in the afternoon of Friday 2026-09-11 after the ECDA closes:

- 1245 ECDA closes / lunch in mensa
- ~13:45 shuttle bus to the venue of the GPSDAA seminar at the Störtebeker brewery
- 1400 GPSDAA begins
- ~1700 tour of the brewery, followed by dinner in the brewery's restaurant.
- each presentation takes 20min+5min for discussion

Inner Schedule

- 14:00 Welcome by Gero Szepannek and André Grüning
- 14:10 Talk 1
- 14:30 Talk 2
- 14:50 Talk 3
- 15:10 Talk 4
- 15:30 Talk 5 (online)
- 15:50 Talk 6
- 16:10 Talk 7
- 16:30 Talk 8
- 16:50 Talk 9
- 17:10 Closing remarks, followed by tour of brewery

Abstracts

Talk 1: 7th Goal of Sustainable Development in Poland, Germany, and the rest of Europe

- Małgorzata Markowska, Wrocław University of Economics and Business, Uniwersytet Ekonomiczny we Wrocławiu, malgorzata.tarczynska-luniewska@usz.edu.pl
- Andrzej Sokołowski, Andrzej Frycz Modrzewski Krakow University, Uniwersytet Andrzeja Frycza Modrzewskiego w Krakowie, asokolowski@uafm.edu.pl
- Jacek Wychowanek, Higher School of Management and Entrepreneurship in Wałbrzych, Wałbrzyska Wyższa Szkoła Zarządzania i Przedsiębiorczości

The Sustainable Development Goals (SDGs) were adopted by the United Nations in 2015 as a universal call to action to end poverty, protect the planet, and ensure that by 2030 all people enjoy peace and prosperity. Goal No 7 “Affordable and Clean Energy” declares to ensure access to affordable, reliable, sustainable and modern energy for all. The aim of the paper is to analyze Goal 7 performance in Germany and Poland, comparing to the rest of Europe, in 2005-2024.

Eurostat provides seven datasets related to SDG 7. We took the following original variables (one from each set) calculated per capita: primary energy consumption in tonnes of oil equivalent, final energy consumption in tonnes of oil equivalent, final energy consumption in households, energy productivity purchasing power standard per kilogram of oil equivalent, share of renewable energy in gross final energy consumption, and energy oil import dependency in percentage. Finally, SDG 7 fulfillment has been, in this paper, measured by the proposed Composite SD7 Index. Primary variables were positively skewed, except export dependence with negative asymmetry. The latter variable has been transferred to positive asymmetry. Then the log of variables were transformed into (0;1) interval, taking reference points as minimum and maximum+SD/10. Reference points were calculated for the whole analyzed period globally. Composite SDG7 Index has been calculated as geometric average of normalized variables, weighted by population size for the rest of Europe. It is important, that we compared individual countries not with the European average, but for the average for the rest of Europe – in our case without Germany and Poland. The rest of Europe Composite SDG7 Index was better than for Germany until 2017, then the German index was going down from 0.54 to 0.47 in 2024, while for the rest of Europe index was still going up. Poland was generally much farther, but growing from 0.34 to 0.44. Trend function were estimated and forecast calculated for 2025-2027 period.

Talk 2: Title TBA

- Andreas Karasenko, Chair of Marketing & Innovation, University of Bayreuth, andreas.karasenko@uni-bayreuth.de
- Daniel Baier, Chair of Marketing & Innovation, University of Bayreuth, daniel.baier@uni-bayreuth.de

Recent developments in machine learning (ML), especially transformer-based discriminative and generative deep learning, transform the marketing landscape. So, e.g., marketers predict with high accuracy sentiment scores from online customer review (OCR) comments in natural language and gain valuable insights whether, when, and how apps, products, or services should be improved. However, oftentimes, OCR comments contain additional interesting information that goes beyond sentiment indications. In this work, we propose a new approach to predict – based on the well-known Technology Acceptance Model (TAM) – extended TAM construct scores from OCRs and compare the accuracy of this prediction with various ML models for this purpose. The comparison is based on a dataset with $n = 5,356$ OCR comments for the Ikea app, labeled by three human experts ($n = 3$), and 18 ML models. Following this we conduct a case study on the Ikea dataset and show how to use these TAM construct scores in conjunction with topic modeling to identify various usability issues of the Ikea app. Additionally, we propose an approach that leverages TAM constructs to identify OCRs with complex and rich content that would not be identifiable with sentiment alone.

Talk 3: The degree of implementation of the location selection strategies during order picking in a warehouse, with shared and class-based storage – a comparison of results obtained assuming that the variables are measured on ratio and ordinal scales

- Krzysztof Dmytrów. Uniwersytet Szczeciński, krzysztof.dmytrow@usz.edu.pl
- Beata Bieszk-Stolorz, Uniwersytet Szczeciński

If a company uses shared storage of items in a warehouse, then each item index can 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, by assigning appropriate weights to the variables describing the locations. In the analysed example, the variables describing the locations are measured on a ratio scale. There may be outliers amongst the variable values, which could distort the ranking of locations to be visited. Therefore, the analysis compared the degree of implementation of the location selection strategies using the classical TOPSIS method, which utilises the distances of objects from the pattern and antipattern determined using the Euclidean metric and the TOPSIS method in which the variables were rescaled

to a weaker ordinal scale through ranking, and the distances from the pattern and antipattern were determined using the generalised distance measure GDM2 (Generalised Distance Measure; the ‘2’ denotes the version of the measure for an ordinal scale). The study was conducted using simulation methods. A total of 1,000 scenarios (orders) were generated, and the ABC class-based storage was assumed. For both approaches, locations to be visited by the picker were selected, and the degree to which the location selection strategies were implemented was then examined.

Talk 4: Automatic analysis of expert opinions

- Paweł Lula, Uniwersytet Ekonomiczny w Krakowie / Krakow
University of Economics, lulap@uek.krakow.pl

Presentation of the solution for automatic analysis of expert opinions is the main purpose of the presentation. The proposed method consists of two main stages. During the first stage, the main topics addressed by the expert are identified. During the second stage, the sentiment of each identified topic is determined. In the empirical section, the operation of the algorithm for analyzing expert opinions written in Polish and English will be demonstrated.

Talk 5: From Self-Checkout to Just-Walk-Out: Understanding Customer Preferences for Cashierless Technologies (online)

- Friederike Paetz, Carsten D. Schultz, Hochschule Anhalt,
friederike.paetz@hs-anhalt.de

The growing popularity of online grocery shopping has increased competitive pressure on traditional retailers, making innovative in-store technologies an important means of improving the shopping experience. Cashierless checkout systems promise greater convenience and efficiency, yet customer acceptance varies across technologies. Previous research has often focused on individual systems or treated different technologies as equivalent, limiting insights into system-specific preferences.

This study compares customer preferences for self-checkout, smart trolley, just-walk-out, and facial recognition with traditional cashier-based checkout. It examines the effects of convenience, customer experience, perceived effort, performance, data privacy, and socio-demographic characteristics. Data from a discrete choice experiment with 236 German consumers are analysed using a Mixed Logit model to capture heterogeneous preferences.

The results reveal distinct preference patterns. Just-walk-out and facial recognition are perceived similarly, whereas smart trolleys occupy an intermediate position. Customer experience and income increase preferences for AI-based systems, while privacy concerns, age, and female gender are associated with lower

acceptance. In contrast, larger households show lower preferences for traditional checkout, potentially reflecting greater sensitivity to waiting times.

The findings demonstrate that customer preferences differ substantially across cashierless technologies. Retailers should therefore consider system-specific customer characteristics when selecting and implementing checkout solutions. Combining different technologies may help accommodate heterogeneous customer preferences and improve the overall shopping experience.

Talk 6: The Many Faces of AUC: Common Foundations of Measures Used in Statistics, Machine Learning, Medicine, and Banking

- Błażej Kochański, Politechnika Gdańska,
blazej.kochanski@pg.edu.pl

The area under the ROC curve (AUC, Area Under the Curve) is one of the most widely used measures for evaluating the performance of classification models. The aim of this presentation is to introduce AUC as a general measure of association between a continuous or ordinal variable and a binary variable, whose significance extends far beyond the classical ROC curve framework.

The presentation will discuss three complementary interpretations of AUC: the ROC curve interpretation, the pairwise comparison (probabilistic) interpretation, and the rank-based interpretation. It will also examine the relationships between AUC and several well-established concepts from different scientific disciplines, including the Mann–Whitney and Brunner–Munzel nonparametric tests, 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 the probability of superiority.

The presentation aims to demonstrate that, despite the diversity of terminology used in medicine, economics, statistics, and machine learning, many widely applied measures are founded on the same underlying mathematical concept. Highlighting these relationships provides a deeper understanding of the interpretation of AUC and its role as a general measure of predictive performance and the strength of association between variables across a wide range of applications.

Talk 7: Linear Programmimg: Bridge between Classification and Simple Testing

- Ulrich Müller-Funk , University of Münster (Germany),
funk@wi.uni-muenster.de

We primarily deal with binary classification with a focus on unbalanced data. Here, the "positives" form the interesting but rare cases. That includes problems like anomaly detection. In that situation, concepts like Bayes classifiers optimizing accuracy, ROC curves providing a threshold etc. are no longer valuable.

Instead, a Neyman-Pearson or a Minimax formulation seems to be adequate. Such an approach is no novelty and has been discussed in a few papers.

There, the underlying idea is that the solution can be read directly from a Neyman-Pearson lemma. If F denotes the law underlying the predictor variables X and Y the binomial label, then the conditional distributions of X given $Y = 0$ resp. $Y = 1$ are the two simple alternatives to be tested. Unfortunately, F happens to be a convex combination of both these laws. Either hypothesis now implies that all three distributions must be equal. Nevertheless, one feels that properly formulated optimization tasks – the one in terms of odds ratios, the other one based on likelihood ratios – should be of the same nature. Neyman-Pearson lemmas are framed in terms of their applications in hypothesis testing, yet—quite independently of this context—they characterize solutions to linear programming problems in function spaces. In this context, the handling of class probabilities must also be elaborated upon. Furthermore, aspects such as unbiasedness, risk space, reversed classifiers, and performance measurement are to be addressed.

Talk 8: `xegaMigration` – Island Models for the R-Package `xega`.

- Andreas Geyer-Schulz, Karlsruher Institut für Technologie (KIT), andreas.geyer-schulz@kit.edu

In this contribution a set of highly configurable communication protocols is presented which govern gene migration between islands of extended evolutionary or genetic algorithms in the R-package `xega`. With the current version, a seamless migration of island models from Linux notebooks without openMPI to high-performance computing environments with thousand of CPUs is possible. Innovations in the communication protocols become necessary in order to use all available CPU processing power on a multi-processor system.

First, a survey of the currently available communication protocols in `xegaMigration` is given. Next, an innovative variant of the communication protocols which provides completely wait-free asynchronous communication between islands with run-time prediction based self-adapting process termination is presented. This variant aims for the best exploitation of CPU processing power possible (doing more function evaluations given a fixed time slot). However, especially for heterogenous island models, this may not be the best choice, because for the speed of convergence of the algorithm the time of information exchange also matters.

See <https://CRAN.R-project.org/package=xega>

**Talk 9: The ESG-Adjusted Fundamental Strength Index as
a Tool for Assessing the Condition of Listed Companies:
Integrating Financial Ratios and Sustainability Data**

- Małgorzata Tarczyńska-Łuniewska, Uniwersytet Szczeciński,
malgorzata.tarczynska-luniewska@usz.edu.pl

TBA