

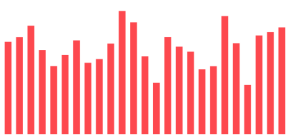


Book of Abstracts

Statistische Woche

08. – 11. September 2026 | Essen




Statistische Woche

UNIVERSITÄT
DUISBURG
ESSEN

Offen im Denken

BOOK OF ABSTRACTS

STATISTISCHE WOCHE

08. – 11. September 2026

Essen

Information und Technik
Nordrhein-Westfalen
Statistisches Landesamt



KonsortSWD
Konsortium für die
Sozial-, Verhaltens-, Bildungs- und
Wirtschaftswissenschaften

RatSWD
Rat für Sozial- und
Wirtschaftsdaten

RWE



Organisers / Organisierende:

University of Duisburg-Essen
Universität Duisburg-Essen



Offen im Denken

German Statistical Society – registered association (DStatG)
Deutsche Statistische Gesellschaft e.V. (DStatG)



Association of German Urban Statistics (VDSt)
Verband Deutscher Städtestatistik (VDSt)



German Association for Demography – registered association (DGD)
Deutsche Gesellschaft für Demographie e. V. (DGD)



With kind support of / Mit freundlicher Unterstützung von:



STATISTISCHE WOCHE

08. – 11. September 2026

Essen – Germany

www.statistische-woche.de

Contents

Programme – Overview	p. 7
Programme	p. 8
• Session Timetable	p. 8
• Tuesday	p. 11
• Wednesday	p. 17
• Thursday	p. 23
• Friday	p. 30
Abstracts	p. 32
Social Programme	p. 138
Participants	p. 139
Building Plans	p. 145

Inhaltsverzeichnis

Programm – Übersicht	S. 7
Programm	S. 8
• Sitzungszeitplan	S. 8
• Dienstag	S. 11
• Mittwoch	S. 17
• Donnerstag	S. 23
• Freitag	S. 30
Abstracts	S. 32
Rahmenprogramm	S. 138
Teilnehmende	S. 139
Gebäudepläne	S. 145

Partners / Partner:

- IT.NRW
- KonsortSWD / RatSWD
- RWE
- SAS
- Springer

Programme Committee / Programmkomitee:

**German Statistical Society – registered association /
Deutsche Statistische Gesellschaft e. V.:**

Yarema Okhrin (Chair / Vorsitz DStatG, Universität Trier)
 Sara Bleninger (Bayerisches Landesamt für Statistik)
 Hanna Brenzel (Statistisches Bundesamt)
 Christine Buchholz (Hochschule Bonn-Rhein-Sieg)
 Matei Demetrescu (Technische Universität Dortmund)
 Christoph Hanck (Universität Duisburg-Essen)
 Solveigh Jäger (Bundesverband der Deutschen Industrie)
 Carsten Jentsch (Technische Universität Dortmund)
 Göran Kauermann (LMU München)
 Hans Kiesel (OTH Regensburg)
 Philipp Otto (University of Glasgow)
 Ulrike Rockmann (Universität Oldenburg)
 Gaby Schneider (Goethe-Universität Frankfurt)
 Katharina Schüller (STAT-UP Statistical Consulting & Data Science GmbH)
 Markus Zwick (Statistisches Bundesamt)

Association of German Urban Statistics / Verband Deutscher Städtestatistik:

Uta Thien-Seitz (Chair / Vorsitz, Direktorium - Statistisches Amt München)
 Diana Andrä (Dortmunder Statistik)
 Jan Dohnke (Statistik und Stadtforschung Darmstadt)
 Andrea Schultz (Stadt Leipzig, Amt für Statistik und Wahlen)
 Ansgar Schmitz-Veltin (Landeshauptstadt Stuttgart, Statistisches Amt)

**German Society for Demography – registered association / Deutsche Gesellschaft für
Demographie e. V.:**

Bernhard Köppen (Chair / Vorsitz, Universität Koblenz)
 Patrizio Vanella (aQua-Institut)

Conference Organiser / Tagungsbeauftragter:

Arne Johannssen (Conference Commissioner / Tagungsbeauftragter DStatG, Hochschule Harz)

Local Organisers / Lokale Organisierende:

Christoph Hanck (Chair / Vorsitz, Universität Duisburg-Essen)
 Ilona Braun (Universität Duisburg-Essen)
 Yannick Hoga (Universität Duisburg-Essen)
 Jens Klenke (Universität Duisburg-Essen)
 Carolin Neuroth (Universität Duisburg-Essen)
 Nhi Ta (Universität Duisburg-Essen)

Programme Overview / Programmübersicht:**Tuesday, 08. September, 2026:**

08:00-09:00	Registration / Welcome Coffee
09:00-10:40	Sessions
10:40-11:10	Coffee Break
11:10-12:00	Opening
12:00-12:50	Plenary Session (Pierre Pinson): <i>AI-based and human aspects of forecasting</i>
12:50-14:10	Lunch Break
14:10-15:50	Sessions
15:50-16:20	Coffee Break
16:20-18:00	Sessions

Wednesday, 09. September, 2026:

09:00-10:40	Sessions
10:40-11:10	Coffee Break
11:10-12:00	Plenary Session (Monica Pratesi): <i>Beyond Accuracy: Rethinking Quality and the Future of Evidence in the Age of Citizen Data</i>
12:00-12:50	Heinz-Grohmann Lecture (Katharina Spieß): <i>Wo Bildungs- und Familienpolitik zusammentreffen: Was amtliche Daten und Paneldaten zeigen</i>
12:50-14:10	Lunch Break
14:10-15:50	Sessions
16:00-17:30	Sessions

Thursday, 10. September, 2026:

09:00-10:40	Sessions
10:40-11:10	Coffee Break
11:10-12:00	Plenary Session (Jonas Peters): <i>Learning under Change</i>
12:00-12:50	Gumbel Lecture (Julia Schaumburg): <i>Clustering extreme value indices in large panels</i>
12:50-14:10	Lunch Break
14:10-15:50	Sessions
15:50-16:20	Coffee Break
16:20-18:00	Sessions
18:00-20:00	Poster Session
18:45-20:00	General Assembly Meeting DStatG

Friday, 11. September, 2026:

09:00-10:40	Sessions
10:40-11:10	Coffee Break
11:10-12:50	Sessions

Session Timetable / Sitzungszeitplan:

	09:00-10:40	14:10-15:50	16:20-18:00	09:00-10:40	14:10-15:50	16:00-17:30	09:00-10:40	14:10-15:50	16:20-18:00	09:00-10:40	11:10-12:50
Section	Tue 08/09			Wed 09/09			Thu 10/09			Fri 11/09	
Main Topic: Sources, Quantification, Accumulation and Communication of Statistical Uncertainty				R11 T03 C54	R11 T03 C54						
Main Topic: Identification											
Main Topic: Statistics and Renewable Energy	R11 T03 C63	R11 T03 C63									
Young-Academics Mini-Symposium: Bayesian Econometrics		R11 T03 C54	R11 T03 C54								
Computational Statistics and Data Science							R11 T03 C82	R11 T03 C82	R11 T03 C82		
Empirical Economics and Applied Econometrics	R11 T03 C20	R11 T03 C20	R11 T03 C20	R11 T03 C20	R11 T03 C20						
Statistics in Finance				R11 T03 C75	R11 T03 C75						
Competence Development: Data Literacy and Statistics	R11 T03 C75	R11 T03 C75	R11 T03 C75								
Methodology of Statistical Surveys							R11 T03 C54	R11 T03 C54	R11 T03 C54	R11 T03 C54	R11 T03 C54
Regional Statistics	R11 T00 D01	R11 T00 D01	R11 T00 D01	R11 T00 D01							
Statistics in Environmental and Engineering Sciences			R11 T03 C63	R11 T03 C63							
Statistical Theory and Methods							R11 T03 C20	R11 T03 C20	R11 T03 C20	R11 T03 C20	R11 T03 C20
Nonparametric and Robust Statistics				R11 T03 C82	R11 T03 C82						
Economic, Social and Market Statistics						R11 T03 C63	R11 T03 C63	R11 T03 C63	R11 T03 C63		
Causal Machine Learning					R11 T03 C63						
Rent and Rental Indices	R11 T03 C54										
DGD							R11 T03 C75	R11 T03 C75			
VDSst				R11 T00 D01	R11 T00 D01	R11 T00 D01	R11 T00 D01	R11 T00 D01			

Plenary Sessions & Lectures

Tuesday, 08. September 2026

Pierre Pinson. „AI-based and human aspects of forecasting“.
Tuesday, 12:00-12:50 (page 32).

Wednesday, 09. September 2026

Monica Pratesi. „Beyond Accuracy: Rethinking Quality and the Future of Evidence in the Age of Citizen Data“.
Wednesday, 11:10-12:00 (page 32).

C. Katharina Spieß. „Wo Bildungs- und Familienpolitik zusammentreffen: Was amtliche Daten und Paneldaten zeigen“.
Wednesday, 12:00-12:50 (page 33).

Thursday, 10. September 2026

Jonas Peters. „Learning under Change“.
Thursday, 11:10-12:00 (page 33).

Chenhui Wang; Juan Juan Cai; Yicong Lin; **Julia Schaumburg.** „Clustering extreme value indices in large panels“.
Thursday, 12:00-12:50 (page 33).

Invited Speakers

Tuesday, 08. September 2026

Nils Weitzel. „Inferring the spatio-temporal structure of climate variability from observations, simulators, and reconstructions of past climate“.

Tuesday, 16:20-17:10 (page 57).

Wednesday, 09. September 2026

Chiara Amorino; Matheus Lima-Cornejo; **Cecilia Mancini.** „Drift burst test statistic: new local tools for asset price model selection“.

Wednesday, 09:00-09:50 (page 65).

Peter Boswijk; Jun Yu; Yang Zu. „Testing for Explosiveness in Financial Asset Prices Using High-Frequency Volatility“.

Wednesday, 14:10-15:00 (page 75).

Thursday, 10. September 2026

Stefano Maria Iacus. „Data Science for a Better World: Building a Safer Digital Environment“.

Thursday, 09:00-09:50 (page 89).

Andreas Anastasiou; **Tobias Kley;** Efstathios Paparoditis. „Spectral Mean Estimators of Time Series: Finite Sample Distributional Approximations“. (*page 108*).

Tuesday, 08. September 2026, 08:00 - 09:00

Welcome Coffee

Tuesday, 08. September 2026, 09:00 - 10:40

Empirical Economics and Applied Econometrics (Session 1)

Chair: Carsten Jentsch, Room: R11 T03 C20

Linus Nüsing; Jan Prüser. „High-Dimensional Nonparametric Proxy SVAR using Gaussian Processes“.

Tuesday, 09:00-09:25 (page 34).

Christian Wurtz; Carsten Jentsch. „Structural analysis and (bootstrap) inference in matrix-autoregressive models“.

Tuesday, 09:25-09:50 (page 35).

Statistics and Renewable Energy (Session 1)

Chair: Gaby Schneider, Room: R11 T03 C63

Antonia Milbert; **Carsten Neumann**; Rajasirpi Subramaniyan; **Elisabeth Stürmer**. „Von Kohleregionen zu regenerative Modellregionen – realistische Daten über die Entwicklung und räumliche Verteilung ihrer EE-Leistung“.

Tuesday, 09:00-09:25 (page 35).

Lisa Sieger; Dennis Schneider; Christoph Weber. „A Socio-Energetic Typology of Small and Medium-Sized Towns in Germany“.

Tuesday, 09:25-09:50 (page 36).

Regionalstatistik (Session 1): Registerzensus

Chair: Marcel Noack, Room: R11 T00 D01

René Söllner. „Zensus 2031: Ein Werkstattbericht“.

Tuesday, 09:00-09:25 (page 37).

Paul Fink; **Stella Klein**; Katrin Kräck. „Erprobung der Methodik eines Registerzensus - Eine Bayerische Perspektive“.

Tuesday, 09:25-09:50 (page 37).

Rent and Rental Indices

Chair: Ulrich Rendtel, Room: R11 T03 C54

Georg Wiegleb. „A Bayesian Distributional Approach to Qualified Local Rent Benchmarks“.

Tuesday, 09:00-09:25 (page 38).

Marco Schmandt. „„Krise? Welche Krise?“ Zur Messung der Bezahlbarkeit von Mietwohnungen in Deutschland“.

Tuesday, 09:25-09:50 (page 38).

Michael Windmann; Göran Kauermann; Oliver Trinkaus. „Wirksamkeit regulatorischer Mietpreisgrenzen – Eine empirische Annäherung“.

Tuesday, 09:50-10:15 (page 39).

Ingo Seifert; Martin Waschipky. „Stichprobenbasierte Fortschreibungen von Mietspiegeln: Methodische Abwägungen zu Stichprobenplanung und Fragebogenumfang am Beispiel der Stadt Leipzig“.

Tuesday, 10:15-10:40 (page 39).

Competence Development: Data Literacy and Statistics (Session 1)

Chair: Walter J. Radermacher, Room: R11 T03 C75

Walter J. Radermacher. „Ethik in der Statistikausbildung – Erfahrungsaustausch und Vernetzung“.

Tuesday, 09:00-09:25 (page 40).

Coffee Break: 10:40 - 11:10

Opening: 11:10 - 12:00

Room: R14 R00 A04

Plenary Session 1

Chair: Ansgar Steland, Room: R14 R00 A04

Pierre Pinson. „AI-based and human aspects of forecasting“.

Tuesday, 12:00-12:50 (page 32).

Tuesday, 08. September 2026, 14:10 - 15:50

Empirical Economics and Applied Econometrics (Session 2)

Chair: Carsten Jentsch, Room: R11 T03 C20

Manuel Frondel; Patrick Thiel; **Colin Vance.** „Bayesian Causal Estimation of German Fuel Price Policy: High Noon and the 2026 Tax Discount“.

Tuesday, 14:10-14:35 (page 40).

Nils Christian Hoenow; **Viola Helmers;** Eva Yang. „Electric bicycles and public transport tickets: Ownership and car use patterns“.

Tuesday, 14:35-15:00 (page 41).

Dörte Heger; Ingo W. K. Kolodziej; Luisa V. Licker; Arndt R. Reichert; **Ansgar Wübker**. „Effects of retirement age policies on informal caregiving in Europe“.

Tuesday, 15:00-15:25 (page 41).

Liam Vance; Kathrin Kaestner. „Warehouse Density and Local Air Quality: Causal Evidence from the Netherlands“.

Tuesday, 15:25-15:50 (page 42).

Statistics and Renewable Energy (Session 2)

Chair: Philipp Otto, Room: R11 T03 C63

Sultan Mahmud Chomon; Florian Ziel. „Networked Spatial Effects in European Electricity Price Forecasting“.

Tuesday, 14:10-14:35 (page 43).

Wenzhen Huang; Florian Ziel. „Interpretable Functional Modeling of Electricity Market Auction Curves“.

Tuesday, 14:35-15:00 (page 43).

Christine Müller. „Voltage state estimation in electrical power distribution grids and optimal sensor location“.

Tuesday, 15:00-15:25 (page 44).

Alexandra Fafali A. Kudjawu; Johannes Schwenzer; Simone Maxand. „Distributional Load Forecasting for Off-Grid Energy Systems“.

Tuesday, 15:25-15:50 (page 44).

Regionalstatistik (Session 2): Regionalplanung

Chair: Sara Bleninger, Room: R11 T00 D01

Jana Hoymann; Nadine Blätgen; Sylvie Dugay; Florian Bernardt; Philip Ulrich. „Regionaler Flächenbedarf für Siedlungs- und Verkehrsflächen“.

Tuesday, 14:10-14:35 (page 45).

Christoph Alfken; **Anna Grüne**. „Kleinräumige Analyse der Grünraumversorgung in Nordrhein-Westfalen“.

Tuesday, 14:35-15:00 (page 46).

Philip Ulrich; Florian Bernardt; Nadine Blätgen; Jana Hoymann; Sylvie Dugay. „Regionaler Flächenbedarf für den Ausbau erneuerbarer Energien“.

Tuesday, 15:00-15:25 (page 46).

Hanna Hoffmann; **Christoph Alfken**; **Katja von Eschwege**. „GeoStat: Ein offener, nutzerfreundlicher Zugang zu amtlichen georeferenzierten Statistikdaten – Neue Auswertungsmöglichkeiten am Beispiel der Daten des statistischen Unternehmensregisters“.

Tuesday, 15:25-15:50 (page 47).

Competence Development: Data Literacy and Statistics (Session 2)

Chair: Katharina Schüller, Room: R11 T03 C75

Andreas Kladroba; Markus Zwick. „Statistik im Spannungsfeld zwischen Theorie und Praxis: Zur Erinnerung an Peter von der Lippe“.

Tuesday, 14:10-14:35 (page 48).

Diana Rauwolf; Erhard Cramer; Udo Kamps; Maria Kateri. „Statistik und Data Literacy mit jamovi: Einsatzszenarien in der Hochschullehre“.

Tuesday, 14:35-15:00 (page 49).

Sigbert Klinke; Kleio Chrysopoulou Tseva. „Automating Statistical Exercise Feedback with LLMs: The sparrow Platform as Infrastructure for AI-Assisted Learning“.

Tuesday, 15:00-15:25 (page 49).

Kleio Chrysopoulou Tseva; Helena Julia Fikus; Sigbert Klinke. „Beyond Weekly Exercises: AI-Assisted Feedback for Independent Statistics Learning“.

Tuesday, 15:25-15:50 (page 50).

Young-Academics Mini-Symposium: Bayesian Econometrics (Session 1)

Chair: Jan Prüser, Room: R11 T03 C54

Ignacio Moreira Lara; Jan Prüser; Christoph Hanck. „A Structural Matrix Autoregressions Framework for International Spillovers“.

Tuesday, 14:10-14:35 (page 51).

Martin Fankhauser. „What can we learn from and about sufficient macro statistics under set-identification?“

Tuesday, 14:35-15:00 (page 51).

Annika Camehl; Maximilian Schröder. „Micro-based SVAR Identification“.

Tuesday, 15:00-15:25 (page 52).

Jan Prüser. „Assessing the Effects of Monetary Shocks on Macroeconomic Stars: A SMUC-IV Framework“.

Tuesday, 15:25-15:50 (page 52).

Coffee Break: 15:50 - 16:20

Tuesday, 08. September 2026, 16:20 - 18:00

Empirical Economics and Applied Econometrics (Session 3)

Chair: Carsten Jentsch, Room: R11 T03 C20

Qingyao Lai; Roxana Halbleib; Winfried Pohlmeier. „Predicting Intraday Price Direction with Conformal Methods“.

Tuesday, 16:20-16:45 (page 53).

Bayram Oruc; Sebastian R  th; Wouter Van der Veken; Ren Zhang. „When OPEC Speaks: The Dollar’s Evolving Reaction to Oil Supply News in the Age of Shale“.
Tuesday, 16:45-17:10 (page 53).

Florian Sch  tze. „Public Attention to Monetary Policy: Asymmetric Responses in Google Search Behavior Following Interest Rate Changes in the Euro Area“.
Tuesday, 17:10-17:35 (page 54).

Gabriel Arce-Alfaro; Dimitris Korobilis. „Unobserved Components with Stochastic Volatility in Trend (UC-SVT)“.
Tuesday, 17:35-18:00 (page 54).

Regionalstatistik (Session 3): Besch  ftigung und Einkommen

Chair: Andrea Buschner, Room: R11 T00 D01

Uwe Neumann. „Community-oriented urban policy and local earnings – a complicated relationship“.
Tuesday, 16:20-16:45 (page 55).

Anja Afentakis. „Aufbau eines regionalen Gesundheitspersonalmonitorings zur kleinr  umigen Planung und Steuerung personeller Ressourcen“.
Tuesday, 16:45-17:10 (page 55).

Sarah Kuhn; Holger Seibert; Antje Weyh. „Struktur und Entwicklung der Pendlerbewegungen in der Lausitz“.
Tuesday, 17:10-17:35 (page 56).

Sevin     zt  rk; Marit K  ffers; **Christoph Alfken; Johannes Rohde.** „Geheimhaltung georeferenzierter Daten: ein Methodenvergleich des Quadtree-Verfahrens und der Kerndichtesch  tzung“.
Tuesday, 17:35-18:00 (page 57).

Statistics in the Environmental Sciences, Natural Sciences and Technology (Session 1)

Chair: Gaby Schneider, Room: R11 T03 C63

Nils Weitzel. „Inferring the spatio-temporal structure of climate variability from observations, simulators, and reconstructions of past climate“.
Tuesday, 16:20-17:10 (page 57).

Florian Heinrichs. „Detecting Stationary States in Non-Stationary Time Series“.
Tuesday, 17:10-17:35 (page 58).

Sheila G  rz; Merle Mendel; Roland Fried. „Detecting Changes in the Mean of Spatial Data with Applications to Deforestation and Vorticity Extremes“.
Tuesday, 17:35-18:00 (page 59).

Competence Development: Data Literacy and Statistics (Session 3)

Chair: Gerald Seidel, Room: R11 T03 C75

Andreas Hildenbrand. „Anforderungen an die Statistikausbildung: Antwort auf Berger et al. (2026)“.

Tuesday, 16:20-16:45 (page 59).

Martin Missonig. „Daten, Data Literacy und Statistik I & II - Thesen zur grundständigen Statistikausbildung in der Wirtschaftswissenschaft“.

Tuesday, 16:45-17:10 (page 60).

Gerald Seidel. „Wie können wir angemessen über gesellschaftlich sensible Statistiken reden?“

Tuesday, 17:10-17:35 (page 60).

Katharina Schüller; Berthold Lausen; Eyke Hüllermeier; Jan Vahrenhold; Walter J. Radermacher; Katharina Weiß. „EDSEA: Eine europäische Allianz zur Definition von Data-Science-Kompetenzen“.

Tuesday, 17:35-18:00 (page 61).

Young-Academics Mini-Symposium: Bayesian Econometrics (Session 2)

Chair: Jan Prüser, Room: R11 T03 C54

Lukas Berend; Jan Prüser. „Sharpening Identification in Large Structural VARs Using Narrative Restrictions“.

Tuesday, 16:20-16:45 (page 62).

Catalina Martinez Hernandez. „Beware of large shocks! A non-parametric structural inflation model“.

Tuesday, 16:45-17:10 (page 62).

Matteo Luciani; **Michele Piffer;** Andrea Renzetti. „Theory-based priors for the output gap“.

Tuesday, 17:10-17:35 (page 63).

Other events on Tuesday

19:00 : Guided Tours at Zollverein Coal Mine and Coking Plant, see page 138

Wednesday, 09. September 2026, 09:00 - 10:40

Empirical Economics and Applied Econometrics (Session 4)

Chair: Carsten Jentsch, Room: R11 T03 C20

Johannes Ludsteck. „Do Low or High Wages Respond more Sensitive to Cyclical Fluctuations? – Empirical Evidence“.

Wednesday, 09:00-09:25 (page 63).

Simone Maxand; Franziska Dorn. „Heat or Eat: Distributional Subsidy Effects and the Joint Allocation of Essential Expenditures“.

Wednesday, 09:25-09:50 (page 64).

Roman Liesenfeld; Guilherme Valle Moura; **Julian Spies.** „Regularized Wishart Stochastic Volatility for High-Dimensional Covariance Forecasting“.

Wednesday, 09:50-10:15 (page 64).

Timo Dimitriadis; Christoph Hanck; Yannick Hoga; **Timo Puschmann.** „Monitoring Distributional Forecasts“.

Wednesday, 10:15-10:40 (page 65).

Statistics in Finance (Session 1)

Chair: Roxana Halbleib, Room: R11 T03 C75

Chiara Amorino; Matheus Lima-Cornejo; **Cecilia Mancini.** „Drift burst test statistic: new local tools for asset price model selection“.

Wednesday, 09:00-09:50 (page 65).

Lennart Dröge; Jos van Bommel. „Pricing Barrier Options in Discontinuous Markets: An Adaptive Step Monte Carlo Approach“.

Wednesday, 09:50-10:15 (page 66).

Philipp Haid; Julian Wustl; Yarema Okhrin. „Explainable Outlier Detection via Exponential Vine Copula Tilting“.

Wednesday, 10:15-10:40 (page 66).

Nonparametric and Robust Statistics (Session 1)

Chair: Melanie Birke, Room: ER11 T03 C82

Edwin Barthel; Carsten Jentsch. „Subgraph counts for hypergraphons with node covariates: asymptotic and bootstrap inference“.

Wednesday, 09:00-09:25 (page 67).

Martin Wendler; Claudia Kirch. „Robust Change-Point Test for Panel Data Based on U-Statistics“.

Wednesday, 09:25-09:50 (page 67).

Eni Musta; Tsz Pang Yuen; Ingrid Van Keilegom. „Testing sufficient follow-up in cure models via monotone tail density constraints“.

Wednesday, 09:50-10:40 (page 68).

**Statistics in the Environmental Sciences, Natural Sciences and Technology
(Session 2)**

Chair: Ansgar Steland, Room: R11 T03 C63

Anja Schmiedt; Stefan Bedbur; Udo Kamps. „A tailored statistical model for non-metallic inclusion sizes in steels and its inference“.

Wednesday, 09:00-09:25 (page 68).

Jekaterina Zukovska. „Cognitive Flexibility Without Stress: Automated Home-Cage Profiling of Learning Dynamics in Mouse Behaviour“.

Wednesday, 09:25-09:50 (page 69).

**Sources, Quantification, Accumulation, and Communication of Statistical
Uncertainty (Session 1)**

Chair: Florian Dumpert, Room: R11 T03 C54

Alex Strebel; Franko Schmähling; Gerd Wübbeler; Max Siebenkotten; Bernd Kästner; Manuel Marschall. „Uncertainty evaluation for smart measurements applied to scanning hyperspectral imaging using ML“.

Wednesday, 09:00-09:25 (page 69).

Niklas Valentin Lehmann. „From Brier to 'h% correct': communicating forecast accuracy as hit rates“.

Wednesday, 09:25-09:50 (page 70).

Riko Kelter. „Optimal Bayesian sequential design of two-arm clinical phase II trials in a nutshell“.

Wednesday, 09:50-10:15 (page 71).

VDSt/Regionalstatistik: Bevölkerungsvorausberechnung, Prognosen

Chair: Diana Andrä, Room: R11 T00 D01

Valerie Leukert. „Bevölkerungsvorausberechnungen für Bayern bis 2044 - Methoden, Annahmen und kleinräumige Ergebnisse bis zur Gemeindeebene“.

Wednesday, 09:00-09:25 (page 71).

Moritz Zöphel. „Mikrosimulation als Instrument kleinräumiger Bevölkerungsvorausschätzung“.

Wednesday, 09:25-09:50 (page 72).

Elisabeth Zessin; Steffen Maretzke. „Analyse und Prognose regionaler Strukturen und Trends der Fertilität in Deutschland“.

Wednesday, 09:50-10:15 (page 72).

Coffee Break: 10:40 - 11:10

Plenary Session 2

Chair: Markus Zwick, Room: R14 R00 A04

Monica Pratesi. „Beyond Accuracy: Rethinking Quality and the Future of Evidence in the Age of Citizen Data“.

Wednesday, 11:10-12:00 (page 32).

Heinz-Grohmann-Lecture

Chair: Ulrich Rendtel, Room: R14 R00 A04

C. Katharina Spieß. „Wo Bildungs- und Familienpolitik zusammentreffen: Was amtliche Daten und Paneldaten zeigen“.

Wednesday, 12:00-12:50 (page 33).

Lunch Break: 12:50 - 14:10

Wednesday, 09. September 2026, 14:10 - 15:50

Causal Machine Learning

Chair: Jannis Kück, Room: R11 T03 C63

Martin Huber; Jannis Kueck; **Theresa M. A. Schmitz.** „Automatic Debiased Machine Learning and Sensitivity Analysis for Dynamic Treatment Effects“. (*page 73*).

Juejue Wang; **Julian Diefenbacher**; Jannis Kueck; Martin Spindler; Carlos Cinelli; Victor Chernozhukov. „Sensitivity Analysis for (Local) Average Treatment Effects in Instrumental Variable Models“. (*page 74*).

Alejandro Schuler; David Bruns-Smith; **Eva-Maria Oess**; Avi Feller. „Understanding the finite-sample relationship between Double Machine Learning and Targeted Maximum Likelihood Estimation“. (*page 74*).

Cathy Yi-Hsuan Chen; Victor Chernozhukov; **Jan Rabenseifner**; Martin Spindler. „Inference on Multiple Treatment Effects with an Application to Health Economics“. (*page 75*).

Empirical Economics and Applied Econometrics (Session 5)

Chair: Carsten Jentsch, Room: R11 T03 C20

Peter Boswijk; Jun Yu; Yang Zu. „Testing for Explosiveness in Financial Asset Prices Using High-Frequency Volatility“.

Wednesday, 14:10-15:00 (page 75).

Daniel Dzikowski; Kurt Lunsford; Carsten Jentsch. „Long-Run Prediction Uncertainty with Persistent Heteroskedasticity“.

Wednesday, 15:00-15:25 (page 76).

Statistics in Finance (Session 2)

Chair: Yarema Okhrin, Room: R11 T03 C75

Timo Dimitriadis; Marius Puke. „Statistical Inference for Score Decompositions“.

Wednesday, 14:10-14:35 (page 76).

Sven Lehmann; **Rainer Alexander Schüssler**. „Signal-Selected Subset Portfolios for Long-Run Growth“.

Wednesday, 14:35-15:00 (page 77).

Lukas Bauer; Ekaterina Kazak. „Conditional Method Confidence Set“.

Wednesday, 15:00-15:25 (page 78).

Theo Berger. „Conditional Volatility Modeling and Financial Risk Measurement via Neural Networks“.

Wednesday, 15:25-15:50 (page 78).

Nonparametric and Robust Statistics (Session 2)

Chair: Melanie Birke, Room: R11 T03 C82

Tim Greger. „Nonparametric quantile regression in the fully functional model“.

Wednesday, 14:10-14:35 (page 79).

Harry Haupt; **Joachim Schnurbus**; Ida Bauer. „Distribution-free prediction intervals from interval score optimised pairs of nonparametric regression quantiles“.

Wednesday, 14:35-15:00 (page 79).

Antonio Fioravanti; Carsten Jentsch. „Nonparametric covariance estimation with bias correction and bootstrapping for spatial lattice processes“.

Wednesday, 15:00-15:25 (page 80).

Melanie Birke; Natalie Neumeyer. „Testing independence of errors and covariates in functional linear models“.

Wednesday, 15:25-15:50 (page 80).

Sources, Quantification, Accumulation, and Communication of Statistical Uncertainty (Session 2)

Chair: Markus Zwick, Room: R11 T03 C54

Patrick Schenk; **Christoph Kern**; Trent D. Buskirk. „Examining the Role of Measurement and Representation in Uncertainty, Bias, and Unfairness Through Total Error Frameworks“.

Wednesday, 14:10-14:35 (page 81).

Florian Dumpert; Inken Veips; Maria Thurow; Markus Pauly. „Stabilität von Imputationsverfahren“.

Wednesday, 14:35-15:00 (page 82).

Pirmin Fessler. „Beyond Standard Errors: How Data Production Processes Shape Uncertainty and Inequality“.

Wednesday, 15:00-15:25 (page 82).

Anne Konrad. „Assessing Specification Uncertainty in Survey Weighting: A Multiverse Monte Carlo Approach“.

Wednesday, 15:25-15:50 (page 83).

VDSSt (Session 1): Wahlanalyse und statistische Helfer für die Wahlorganisation

Chair: Volker Holzendorf, Room: R11 T00 D01

Lutz Wallhorn. „Modelle, Milieus, Karten - Ein integrierter Ansatz für kleinräumige Wahlanalysen“.

Wednesday, 14:10-14:35 (page 83).

Boris Fischer. „Vergleich der Wahlstrukturanalysen der Bundestagswahlen 2025 und 2021 in München“.

Wednesday, 14:35-15:00 (page 84).

Johannes Kiess. „It's the social, economic and infrastructure, stupid! Erklärungsfaktoren auf Gemeindeebene und die Bundestagswahlergebnisse“.

Wednesday, 15:00-15:25 (page 84).

Leonie Schröder. „Vom Wahldatensatz zur Blaupause: Datengenerierung und -aufbereitung für Kommunalwahlen im Praxiseinsatz“.

Wednesday, 15:25-15:50 (page 85).

Wednesday, 09. September 2026, 16:00 - 17:30

Labour Markets and Social Security II

Chair: Bernd Hofmann, Room: R11 T03 C63

Ute Leber; **Barbara Schwengler.** „Forschungsbefunde zu Absprüngen von Bewerberinnen und Bewerbern auf Ausbildungsplätze“.

Wednesday, 16:00-16:25 (page 85).

Kerstin Bruckmeier; Böhme Fabian. „Monetäre Anreize im Sozialsystem und ihre Auswirkung auf die Beschäftigung von Grundsicherungsbeziehenden“.

Wednesday, 16:25-16:50 (page 86).

Ehsan Vallizadeh; Anton Klaus. „Arbeitsmarktintegration von Geflüchteten im Zeitverlauf: Erkenntnisse aus kohortenbasierten Verbleibsanalysen der Arbeitsmarktstatistik“.

Wednesday, 16:50-17:15 (page 87).

Janosch Korell; Torsten Harms. „Macht Pflege arm? Eine differenzierte Kausalanalyse mit Paneldaten des SOEP“.

Wednesday, 17:15-17:40 (page 88).

General Assembly Meeting VDSSt

Wednesday, 16:00-17:30, Room: R11 T00 D01

Other events on Wednesday

18:00 : Conference Dinner, see page 138

Thursday, 10. September 2026, 09:00 - 10:40

Computational Statistics and Data Science (Session 1)

Chair: Philipp Otto, Room: R11 T03 C82

Stefano Maria Iacus. „Data Science for a Better World: Building a Safer Digital Environment“.
Thursday, 09:00-09:50 (page 89).

Kai-Robin Lange; Tobias Schmidt; Matthias Reccius; Henrik Müller; Michael Roos; Carsten Jentsch. „Identifying economic narratives in large text corpora“.
Thursday, 09:50-10:15 (page 89).

Rainer Dyckerhoff; Erik Mendros; Stanislav Nagy. „Computing halfspace depth in $O(n^{d-1})$ via topological sweep“.
Thursday, 10:15-10:40 (page 90).

DGD (Session 1): Methodische Herausforderungen, Lösungsansätze und aktuelle Entwicklungen in der Regionaldemografie

Chair: Philipp Deschermeier, Room: R11 T03 C75

Patrizio Vanella; Julian Ernst. „Schätzung und Bias-Anpassung regionaler alters-, geschlechts- und nationalitätsspezifischer Mortalitätsraten bei unvollständigen und fehlerhaften Informationen - Anwendung für die deutschen Bundesländer“.
Thursday, 09:00-09:25 (page 90).

Falk Voit; Patrizio Vanella. „Sind Sie sicher? Eine Unsicherheitsanalyse am Beispiel regionaler Bevölkerungsprojektionen in Niedersachsen“.
Thursday, 09:25-09:50 (page 91).

Georg Wiegleb. „A Bayesian Approach to Municipal Cohort-Component Forecasting“.
Thursday, 09:50-10:15 (page 92).

Methodology of Statistical Surveys (Session 1)

Chair: Jannek Mühlhan, Room: R11 T03 C54

Gloria Deetjen; Armando Häring. „NeDaMo – Kalibrierung von Mobilfunkdaten mithilfe einer Sondererhebung“.
Thursday, 09:00-09:25 (page 92).

Anne Konrad. „Assessing the Quality of Survey Weights in Panel Studies: Limitations of Design Effects and Effective Sample Size“.
Thursday, 09:25-09:50 (page 93).

Björn Rohr; Christian Bruch; Christof Wolf. „Using Partial Least Squares Regression for Imputation of Multinomial Variables“.
Thursday, 09:50-10:15 (page 94).

Hoda Mohamed. „Microdata Linking in Official Statistics: Opportunities and Challenges“.
Thursday, 10:15-10:40 (page 94).

Statistical Theory and Methods (Session 1)

Chair: Jan-Lukas Wermuth, Room: R11 T03 C20

Paul Navas. „Characterizing M-Estimators using Canonical Divergences from Information Geometry“.

Thursday, 09:00-09:25 (page 95).

Marc-Oliver Pohle; **Tanja Zahn**; Sebastian Lerch. „Uncertainty Quantification in Forecast Comparisons“.

Thursday, 09:25-09:50 (page 96).

Matei Demetrescu; Christoph Hanck; Yannick Hoga. „A robust test for equal predictive accuracy“.

Thursday, 09:50-10:15 (page 96).

Timo Dimitriadis; **Jan-Lukas Wermuth**; Johanna Ziegel. „Sequential Comparison of Predictive Ability with Unbounded Score Differences“.

Thursday, 10:15-10:40 (page 97).

VDSt (Session 2): KI - Einsatzrahmenbedingungen und Nutzen

Chair: Andrea Schultz, Room: R11 T00 D01

Anika Groß. „Künstliche Intelligenz in der Kommunalverwaltung: Ein Überblick über den KI-Einsatz in den Kommunen“.

Thursday, 09:00-09:25 (page 97).

Jennifer Kreklow; Jan Lunge. „Agentic Coding am Beispiel GitHub Copilot. Ein Praxisbeispiel aus Wolfsburg“.

Thursday, 09:25-09:50 (page 98).

Annelie Heuser. „LLMs als Analysewerkzeug für offene Antworten bei Bürgerbefragungen“.

Thursday, 09:50-10:15 (page 98).

Labour Markets and Social Security I

Chair: Christian Zemann, Room: R11 T03 C63

Jonas Krinitz; **Christian Schneemann.** „Schuldenbremse und Investitionspaket - Wirkung der Grundgesetzänderung auf den Arbeitsmarkt“.

Thursday, 09:00-09:25 (page 99).

Davit Adunts; Ehsan Vallizadeh; Anton Klaus. „Demografie, Zuwanderung und Beschäftigungswachstum in Deutschland: Evidenz aus Arbeitsmarkt- und Decompositionsanalysen“.

Thursday, 09:25-09:50 (page 99).

Martina Rengers. „Wer arbeitet wie viel – und wie viel wäre gewünscht? Arbeitszeitrealitäten und -präferenzen im Geschlechtervergleich und nach Lebensform“.

Thursday, 09:50-10:15 (page 100).

Anja Sonnenburg; Ines Thobe. „Zukunft der Gesundheitsberufe: Erste Ergebnisse des Fachkräftemonitorings für das BMG bis 2040“.

Thursday, 10:15-10:40 (page 101).

Coffee Break: 10:40 - 11:10

Plenary Session 3

Chair: Yannick Hoga, Room: R14 R00 A04

Jonas Peters. „Learning under Change“.

Thursday, 11:10-12:00 (page 33).

Gumbel-Lecture

Chair: Yarema Okhrin, Room: R14 R00 A04

Chenhui Wang; Juan Juan Cai; Yicong Lin; **Julia Schaumburg.** „Clustering extreme value indices in large panels“.

Thursday, 12:00-12:50 (page 33).

Lunch Break: 12:50 - 14:10

Thursday, 10. September 2026, 14:10 - 15:50

Computational Statistics and Data Science (Session 2)

Chair: Christian Schulz, Room: R11 T03 C82

Monika Zimmermann; Claudia Collarin; Simon Wood; Florian Ziel. „Scalable Estimation of Large Generalized Additive Models: Extended Fellner-Schall via Hutch++ and Conjugate Gradients“.

Thursday, 14:10-14:35 (page 101).

Christian Schulz; Christoph Hanck; Simon Hirsch; Florian Ziel. „Online Regularized Generalized Linear Models for Vine Copulas“.

Thursday, 14:35-15:00 (page 102).

Btissame El Mahtout; Florian Ziel. „Fast Training of Mixture-of-Experts for Time Series Forecasting via Expert Loss Integration“.

Thursday, 15:00-15:25 (page 102).

DGD (Session 2): Methodische Herausforderungen, Lösungsansätze und aktuelle Entwicklungen in der Regionaldemografie

Chair: Patrizio Vanella, Room: R11 T03 C75

Michael Kalinowski. „QuBe - Projektion des zukünftigen Arbeitsangebots bis 2045“.

Thursday, 14:10-14:35 (page 103).

Timon Hellwagner; Patrizio Vanella. „Analyzing the Effects of Demographic Changes on Regional Human Capital Stocks in Germany“.

Thursday, 14:35-15:00 (page 103).

Steffen Maretzke; Jana Hoymann; Claus Schlömer. „Regionale Strukturen und Trends der Erwerbspersonenentwicklung in Deutschland 2022 bis 2045“.

Thursday, 15:00-15:25 (page 104).

Methodology of Statistical Surveys (Session 2)

Chair: Hanna Brenzel, Room: R11 T03 C54

Joseph Sakshaug; Jan Mackeben. „Cost-Error Trade-Offs in Introducing Web in an Ongoing Telephone Labor Market Survey“.

Thursday, 14:10-14:35 (page 105).

Michael Bergrab; Christian Assmann; Timo Gnambs. „Participation and Selectivity in a Randomised Mixed-Mode Panel Design: Evidence from NEPS Starting Cohort 4“.

Thursday, 14:35-15:00 (page 105).

Barbara Felderer; Anne Balz; **Julian Diefenbacher**; Jannis Kück; Aline Mohr; Tobias Rettig; Martin Spindler. „You’ve got Mail: Does sending Thank-you postcards increase response in a probability-based online panel?“

Thursday, 15:00-15:25 (page 106).

Daniel Meyer; **Elisabeth Stürmer**; Antonia Milbert. „Was Kommunalbefragungen zur Umfrageforschung beitragen: Methodische Befunde aus der BBSR-Kommunalbefragung „Kommunen im Wandel““.

Thursday, 15:25-15:50 (page 107).

Statistical Theory and Methods (Session 2)

Chair: Matei Demetrescu, Room: R11 T03 C20

Andreas Anastasiou; **Tobias Kley**; Efstathios Paparoditis. „Spectral Mean Estimators of Time Series: Finite Sample Distributional Approximations“. (*page 108*).

Miao Yu; Philipp Sibbertsen; Yeliz Özer. „Multivariate Local Whittle Estimation for Multivariate Functional Time Series“. (*page 108*).

Mustafa Kilinc; Michael Massmann. „Testing the order of fractional integration when smooth deterministic trends are possibly present“. (*page 109*).

VDSt (Session 3): Bauen, Wohnen, Verkehr - Angewandte Geodatenanalysen

Chair: Uwe Meer, Room: R11 T00 D01

Georg Wiegler. „Georeferenzierte Standortindikatoren: Kleinräumige Lageanalysen für kommunale Fragestellungen“.

Thursday, 14:10-14:35 (page 109).

Alexander Weigert; Kerstin Lange. „Bautätigkeitsstatistik im Wandel: Datenqualität, Digitalisierung und Auswertungen nach Gemeindegrößenklassen auf Basis des tatsächlichen Ereignisdatums“. *Thursday, 14:35-15:00 (page 110).*

Kerstin Lange; Alexander Weigert. „Bautätigkeitsstatistik im Wandel: Neues Auswertungskonzept mit methodischen Schätzungen am aktuellen Rand“. *Thursday, 15:00-15:25 (page 110).*

Stefan Kaup; Silas Eichfuss; **Lizanne Hauser.** „Schätzung kleinräumiger Armutsgefährdung in Deutschland mit GIS“. *Thursday, 15:25-15:50 (page 111).*

Economic, Social and Market Statistics

Chair: Solveigh Jäger, Room: R11 T03 C63

Verena Leyendecker. „Wie lassen sich forschende Unternehmen typisieren? Eine multivariate Analyse der FuE-Erhebung“. *Thursday, 14:10-14:35 (page 111).*

Jessica Ernst; Johannes Schmitt. „Forschung und Entwicklung in der Pandemie – Merkmale krisenresilienter Unternehmen“. *Thursday, 14:35-15:00 (page 112).*

Andreas Kladroba. „Fördert Demokratie die Innovationsfähigkeit eines Landes?“ *Thursday, 15:00-15:25 (page 112).*

Tobias Maier; Dorothea Krabbe; Peter Dreuw. „Erweiterungsmöglichkeiten der Unternehmensstrukturstatistik zur Analyse der Qualifikationsnachfrage von Unternehmen“. *Thursday, 15:25-15:50 (page 113).*

Coffee Break: 15:50 - 16:20

Thursday, 10. September 2025, 16:20 - 18:00

Computational Statistics and Data Science (Session 3)

Chair: Heiko Limberg, Room: R11 T03 C82

Heiko Limberg. „Von vorläufigen zu finalen Zahlen: Ein modellbasierter Ansatz zur Vervollständigung von Außenhandelsdaten“. *Thursday, 16:20-16:45 (page 114).*

Amelie Plöger; Colin Schmidt; Sarah Redlich. „Design und Evaluierung synthetisch generierter Mikrozensus-Campus Files für den Einsatz in der Hochschullehre“. *Thursday, 16:45-17:10 (page 114).*

Ulrich Reincke. „Vom Prompt zum Prozessfluss: Reproduzierbare Statistik und KI mit SAS Viya, Open Source und KI-Assistenz“. *Thursday, 17:10-17:35 (page 115).*

Methodology of Statistical Surveys (Session 3)

Chair: Steffen Moritz, Room: R11 T03 C54

Ariane Lestrade. „Textklassifikation mit Machine Learning im ESS: Einblicke aus dem AIML4OS-Projekt“.

Thursday, 16:20-16:45 (page 115).

Susanne Wegner. „Einsatz von LLMs zur Klassifikation der Wirtschaftszweige“.

Thursday, 16:45-17:10 (page 116).

Richard Bündgens; Holger Leerhoff; Birgit Pech; Schulz Julian; Christian Borgs; Mödinger Patrizia; Dohn Nicole. „Maschinelle Lernverfahren als Unterstützung für die Wirtschaftszweigschlüssel-Umstellung 2025 im Unternehmensregister URS“.

Thursday, 17:10-17:35 (page 117).

Murad Hajiyev; Christopher Nölting. „Optimierung der Datenverknüpfung mit überwachtem maschinellen Lernen: Random Forest“.

Thursday, 17:35-18:00 (page 117).

Statistical Theory and Methods (Session 3)

Chair: Dominik Wied, Room: R11 T03 C20

Sven Pappert. „Learning Copulas via Distributional Regressions“.

Thursday, 16:20-16:45 (page 118).

Harry Haupt. „Time series quantiles without moments“.

Thursday, 16:45-17:10 (page 119).

Philip Dörr; Holger Dette. „Model checks for copula regression“.

Thursday, 17:10-17:35 (page 119).

Holger Dette; Kathrin Möllenhoff; **Dominik Wied.** „Practically Significant Differences between Conditional Distribution Functions“.

Thursday, 17:35-18:00 (page 119).

National Accounts, Welfare Measurement

Chair: Albert Braakmann, Stefan Hauf, Room: R11 T03 C63

Stefan Hauf; Christoph-Martin Mai. „Werkstattbericht zum Einsatz von generativer KI in den Volkswirtschaftlichen Gesamtrechnungen“.

Thursday, 16:20-16:45 (page 120).

Ulf von Kalckreuth. „Navigating the value chain – from the top to the bottom and back again. A unified statistical framework for micro-level carbon accounting“.

Thursday, 16:45-17:10 (page 121).

Stefan Linz. „Konjunkturstatistik neu denken – Navigation durch die Vielfalt der Nutzerbedarfe“.

Thursday, 17:10-17:35 (page 121).

Christian Salwiczek. „Verwendung des statistischen Unternehmensregisters für daten-basierte Analysen. Restriktionen und Lösungsansätze für ausgewählte Fragestellungen“. *Thursday, 17:35-18:00 (page 122).*

Thursday, 04. September 2025, 18:00 - 20:00

Poster Session

Chair: Arne Johannssen, Room: R12 S00 H12

Valentin Schmidberger; Johannes Maucher; Julia Lapp; Albrecht Melan. „Code Generation for Open Data Statistics: Case Study on the Genesis Database“. *(page 123).*

Christine Buchholz. „Daten-Ethik als Element im digitalen Data-Literacy-Lehrprojekt“. *(page 123).*

Lisa Carius-Munz; Jannek Mühlhan. „Forschungsstrategie des Statistischen Bundesamtes: Kooperation mit der Wissenschaft“. *(page 124).*

Nadine Blätgen; Jana Hoymann; Sylvie Dugay; Fabian Dosch; Florian Bernardt; Philipp Ulrich; Martin Behnisch; Hannah Poglitsch; Jens-Martin Gutsche. „Kernergebnisse des Projektes "Freiraumverluste 2030"“. *(page 125).*

Judith Kaschowitz; Cornelia Müller; Dorothee Winkler. „Kleinräumige Entwicklungen in Städten analysieren: Datensatz „Innerstädtische Raumbewachung““. *(page 125).*

Tomás del Barrio Castro; Álvaro Escribano; **Yeliz Özer;** Philipp Sibbertsen. „Modeling Long Memory Cyclical Trends in the Cenozoic“. *(page 126).*

Marco Jeschke; Leia Betting; Roland Fried. „Probabilistic Modeling and Forecasting of Household Energy Consumption“. *(page 126).*

René Kremer. „Skalierbares Record Linkage für administrative Bevölkerungsdaten: Erfahrungen aus einem Methodentest“. *(page 127).*

Louise Drozdek; Franziska Dorn; Simone Maxand. „The Tail the Subsidy Doesn't Reach: Distributional Evidence on Extreme Energy Burdens in the US“. *(page 127).*

Maren Koehlmann; Markus Zwick. „Verbessern europäische Regelungen den Zugang zu Privately Held Data? Mobilfunk- und Plattformdaten im Kontext der revidierten Verordnung (EG) Nr. 223/2009“. *(page 128).*

Mareike Muth. „WISTA Wirtschaft und Statistik“. *(page 129).*

General Assembly Meeting DStatG

Thursday, 18:45-20:00, Room: R11 T00 D01

Friday, 11. September 2026, 09:00 - 10:40

Methodology of Statistical Surveys (Session 4)

Chair: Sara Schiesberg, Room: R11 T03 C54

Igor Franjić; Sara Schiesberg. „Record Linkage powered by LLMs: Mikrodatenverknüpfungen mit Transformer-basierten Sprachmodellen“.*Friday, 09:00-09:25 (page 129).***Yannik Buhl.** „Knochensägen, Katheter und Ibuprofen: Produktklassifikation mit LLMs im Krankenhausbereich“.*Friday, 09:25-09:50 (page 130).***Adrian Montag.** „Kontextbasierte Klassifikation in den freiwilligen Haushaltserhebungen: Ein agentenbasierter GenAI-Workflow“.*Friday, 09:50-10:15 (page 130).***Sara Schiesberg.** „Neuronale Nomenklatur: Retrieval-then-Reasoning für die Außenhandelsstatistik“.*Friday, 10:15-10:40 (page 131).*

Statistical Theory and Methods (Session 4)

Chair: Jennifer Simota, Room: R11 T03 C20

Daria Ovsyannikova. „Asymptotic Limits for Component-wise L2 Boosting: Theory, Derivation, and Simulation Evidence“.*Friday, 09:00-09:25 (page 131).***Johanna Einhorn; Timo Schmid.** „Distributional Random Forests in the context of hierarchical data“.*Friday, 09:25-09:50 (page 132).***Malte Jahn.** „A universal time series model (for continuous data)“.*Friday, 09:50-10:15 (page 132).***Jennifer Simota; Anne Leucht.** „High dimensional item selection via penalized multinomial logistic“.*Friday, 10:15-10:40 (page 133).*

Coffee Break: 10:40 - 11:10

Friday, 11. September 2025, 11:10 - 12:50

Methodology of Statistical Surveys (Session 5)

Chair: Hanna Brenzel, Room: R11 T03 C54

Johannes Baumstark; Michael Fürnrohr; Adrian Reichert; Judith Schwitulla. „Evaluierung der Haushaltegenerierung im Zensus 2022“.*Friday, 11:10-11:35 (page 133).*

Rainer Schnell; Rainer Lenz; Rolf Schmidt. „Eine methodische Kritik des Zensus 2022“.

Friday, 11:35-12:00 (page 134).

Rainer Schnell; **Severin V. Weiand**. „Die Illusion der Automatisierung des Zensus“.

Friday, 12:00-12:25 (page 134).

Statistical Theory and Methods (Session 5)

Chair: Mirko Alexander Jakubzik, Room: R11 T03 C20

Uwe Hassler; **Marc-Oliver Pohle**; Tanja Zahn. „Simultaneous Inference Bands for Autocorrelations“.

Friday, 11:10-11:35 (page 135).

Shaochen Wang; **Christian Weiß**. „Novel Stein-type Characterizations of Bivariate Count Distributions with Applications“.

Friday, 11:35-12:00 (page 135).

Angelika Silbernagel; Christian H. Weiß. „The Combined INAR Model for Count Random Fields“.

Friday, 12:00-12:25 (page 136).

Mirko Alexander Jakubzik. „Bootstrap and Minimum Distance Methods for Spatio-Temporal Counting Processes“.

Friday, 12:25-12:50 (page 136).

1 AI-based and human aspects of forecasting

Pierre Pinson

Imperial College London, Vereinigtes Königreich
p.pinson@imperial.ac.uk

SECTION: Plenary Session 1

Here is the abstract: “Weather, macroeconomic, sports and political forecasts are ubiquitous, while forecasting in business and industry drives tremendous operational and economic gains. Forecasting naturally combines quantitative approaches (from, e.g., statistics and machine learning) with human judgement and decision-making. Increasing availability of data and computing power has enhanced our forecasting capabilities exponentially. While many think that forecasting is an established science and practice, we currently see many elements of a revolution in that field. This (interactive) lecture will introduce the elements that are triggering and possibly driving that revolution, from new capabilities within AI, to tighter integration of forecasting and decision-making and to high-stakes applications. We will discuss forecasting from the scientific, business, and human points of view. We will ask ourselves whether we can forecast anything and everything. We will reflect on how data and information are becoming commodities.

2 Beyond Accuracy: Rethinking Quality and the Future of Evidence in the Age of Citizen Data

Monica Pratesi

University of Pisa, Italien
monica.pratesi@unipi.it

SECTION: Plenary Session 2

The data ecosystem is undergoing a profound transformation. Evidence is no longer produced exclusively by official statistical agencies and research institutions; citizens, communities, civil society organizations, digital platforms, and citizen science initiatives increasingly contribute to data production. This evolution challenges traditional notions of quality and calls for new approaches to assessing evidence.

Conventional quality frameworks emphasize relevance, accuracy, reliability, timeliness, coherence, comparability, and accessibility. While these dimensions remain essential, emerging frameworks such as the Copenhagen Framework on Citizen Data highlight additional aspects, including inclusion, transparency, trust, accountability, ethics, and citizen agency. Quality is therefore becoming both a technical and a social property of data.

Citizen-generated data illustrates this shift. Initiatives such as OpenStreetMap, participatory mapping, and citizen science platforms demonstrate how community engagement can produce valuable and trusted evidence, while also revealing populations and realities often overlooked by traditional data systems. However, challenges remain, including self-selection bias, uneven coverage, limited standardization, and ethical concerns.

Experiences such as Istat's homelessness census and UN Women's Gender-related Citizen Data Quality Assurance Framework show that high-quality evidence increasingly depends on combining methodological rigor with meaningful participation. The future of quality frameworks lies in

integrating statistical robustness with social legitimacy, ensuring that evidence is not only accurate, but also inclusive, trustworthy, and relevant to the communities it seeks to represent.

3 Learning under Change

Jonas Peters

ETH Zurich, Schweiz
jonas.peters@stat.math.ethz.ch

SECTION: Plenary Session 3

Many learning methods build on the paradigm that training and test data are independent draws from the same distribution. In this talk, we discuss when and how this assumption may be violated. We do this from a causal viewpoint and show how this view leads to stable learning methods. Specifically, we propose stable methods for prediction-intervention games and dimensionality reduction, and discuss connections to instrumental variables. The presented work is joint work with several co-authors who will be mentioned during the talk. The talk does not assume any prior knowledge about causality.

4 Wo Bildungs- und Familienpolitik zusammentreffen: Was amtliche Daten und Paneldaten zeigen

C. Katharina Spieß

Bundesinstitut für Bevölkerungsforschung (BiB), Deutschland
Direktorin@bib.bund.de

SECTION: Heinz-Grohmann-Lecture

tba

5 Clustering extreme value indices in large panels

Chenhui Wang, Juan Juan Cai, Yicong Lin, **Julia Schaumburg**

Vrije Universiteit Amsterdam, Niederlande
j.schaumburg@vu.nl

SECTION: Gumbel-Lecture

We study large panels of units grouped by shared extreme value indices (EVIs) and aim to identify these unknown groups. To achieve this, we order the Hill estimates of individual EVIs and segment them by minimizing the total squared distance between each estimate and its corresponding group mean. We then estimate the group-specific EVIs using a weighted group Hill estimator. We show that our method consistently recovers group memberships, and we establish the asymptotic normality of the group Hill estimator under possible cross-sectional tail dependence. The latter attains a faster convergence rate than the individual Hill estimator, leading to improved estimation accuracy. Simulation results reveal that our method achieves high empirical segmentation accuracy, and the resulting group EVI estimates have substantially reduced mean squared errors compared to individual estimates. We apply the proposed method to analyze a rainfall dataset covering the

winter observations of 4,735 weather stations across Europe. Splitting the data into two time periods (1950–1985 and 1986–2020), we find statistically significant evidence of an increase in the highest and a decrease in the lowest group EVI estimates, suggesting growing variability and intensification of extreme rainfall events across Europe.

6 High-Dimensional Nonparametric Proxy SVAR using Gaussian Processes

Linus Nüsing (1), Jan Prüser (2)

1: Universität Konstanz, Deutschland; 2: Technische Universität Dortmund, Deutschland
linus.nuesing@uni-konstanz.de

SECTION: Empirical Economics and Applied Econometrics (Session 1)

Structural Vector Autoregressions (SVARs) are a workhorse model for empirical macroeconomic analysis. Standard specifications, however, impose linearity in the conditional mean and constant transmission coefficients, ruling out state-dependent shock propagation and time variation in macroeconomic volatility. Moreover, the curse of dimensionality limits standard SVARs to small information sets, raising concerns about omitted variable bias.

We propose a Gaussian Process Factor Structural VAR (GP-FSVAR) that addresses these limitations within a single, unified framework. Building on the Factor-SVAR model in which the reduced-form innovations of a high-dimensional VAR are decomposed into a low-dimensional factor structure, we model the conditional mean and the factor loading matrix as a state-dependent function by placing Gaussian process priors on the basis representations. This yields a model that is nonparametric in both the conditional mean and the conditional variance. The state-dependent loading matrix simultaneously generates flexible impulse responses that vary continuously with economic conditions and implicit stochastic volatility, without requiring a pre-specified regime structure or a separate volatility process. Structural identification is achieved via an external proxy variable, which embeds naturally into the factor structure without additional computational burden. The model is order-invariant and scales to high-dimensional systems by keeping the number of structural parameters bounded independently of the number of variables.

We validate the approach in a Monte Carlo study using two complementary data-generating processes, demonstrating that the model captures genuine nonlinearities when present without introducing spurious ones when the true process is linear.

In an empirical application, we study the macroeconomic transmission of uncertainty shocks in the U.S. economy using systems of increasing dimension. We examine whether the effects of uncertainty shocks exhibit sign and size asymmetries, and whether their transmission varies systematically with the prevailing state of the economy.

7 Structural analysis and (bootstrap) inference in matrix-autoregressive models

Christian Wurtz, Carsten Jentsch

Technische Universität Dortmund, Deutschland
christian.wurtz@tu-dortmund.de

SECTION: Empirical Economics and Applied Econometrics (Session 1)

We consider a structural matrix-autoregressive (SMAR) model to conduct impulse response analysis for structural shocks to matrix-valued time series. The MAR model of order p offers a parsimonious and interpretable framework for these time series, thus addressing issues of high-dimensionality in corresponding vector-autoregressive (VAR) models. To interpret the dynamics, we resort to impulse response analysis as a popular tool from the SVAR context. Its conclusions rely on the valid identification of structural shocks that are mutually contemporaneously uncorrelated and interpretable. In contrast to the existing literature, the proposed SMAR model enables the identification of multiple structural shocks. To address the restrictive nature of the single-term MAR(p) model, we discuss the extension to a multi-term SMAR(p) model as a compromise between the single-term SMAR and the (unrestricted) SVAR model, trading off parsimony against flexibility. We discuss its identification, focusing in particular on issues that arise due to the typical Kronecker-product structure of the coefficient matrices in the MAR framework. Further, we discuss estimation and inference in the general multi-term SMAR(p) model. Specifically, we consider a suitable residual-based bootstrap method to compute confidence intervals for the impulse response curves and provide corresponding results. In this context, a key point concerns model misspecification and the use of MAR models to approximate more general SVAR data generating processes. Finally, we demonstrate the performance and practical use of our approach by Monte Carlo simulations and a real data application.

8 Von Kohleregionen zu regenerative Modellregionen – realistische Daten über die Entwicklung und räumliche Verteilung ihrer EE-Leistung

Antonia Milbert, Carsten Neumann, Rajasirpi Subramaniyan, Elisabeth Stürmer

Bundesinstitut für Bau-, Stadt- und Raumforschung (BBSR), Deutschland
Elisabeth.Stuermer@bbr.bund.de

SECTION: Statistics and Renewable Energy (Session 1)

Das Klimaschutzrechtliche Gebot der Energiewende und der beschlossene Kohleausstieg führen in den drei deutschen Braunkohleregionen zu erheblichen Strukturveränderungen. Das Lausitzer, Mitteldeutsche und Rheinische Revier nehmen den Ausstiegsbeschluss zum Anlass, für sich eine Zukunft als Modellregionen in der Energiewende zu entwickeln. Die Transformation der Energieversorgung von oligopolen Strukturen des fossil-nuklearen Energiesektors führt zu einem dezentralen und kleinteiligeren Anbieterfeld eines regenerativen Energiesystems. Dieses ist stärker von Privatpersonen, Kleinunternehmen, Bürgerinitiativen und Energiegenossenschaften geprägt. Zugleich entwickeln die drei großen Energieversorger Strategien zur Transformation ihres Geschäftsfeldes auf erneuerbare Energien in großem Maßstab. Flächen für Tagebaue werden beispielsweise in diesem

Zusammenhang vorrangig als Standorte für den Ausbau von Produktion, der Speicherung und Verteilung regenerativer Energie und den Aufbau alternativer Produktionen diskutiert.

Der Beitrag gibt einen Einblick in die aktuelle Entwicklung des Ausbaus erneuerbarer Energiekapazitäten in den drei Kohleregionen, auch im Vergleich zu Entwicklungen anderer Regionen. Er beschreibt die Kombination verschiedener Datenaufbereitungen des Marktstammdatenregisters. Neben der Verortung von Freiflächenanlagen nach Manske et al. (2022) werden raum-zeitliche Dynamiken der installierten EE-Leistung über kleinräumige, rasterbasierte Veränderungsanalysen von Kleinstanlagen wie beispielsweise „Balkonkraftwerken“ abgebildet. Diese werden über ein automatisiertes Verfahren der Geokodierung, unter approximierter Lokalitätsabschätzung und multi-kriterieller Informationsaggregation, aus dem Marktstammdatenregister konstruiert. Auf diese Weise wird ein neues Datenprodukt generiert, das die Entwicklungen der installierten Kapazitäten, privater und gewerblicher Eigentümerstrukturen berücksichtigt und somit ermöglicht die Flächeninanspruchnahme und die Niederspannungs-Netzbelastungen nachzuvollziehen.

9 A Socio-Energetic Typology of Small and Medium-Sized Towns in Germany

Lisa Sieger, Dennis Schneider, Christoph Weber

University of Duisburg-Essen, Germany
lisa.sieger@uni-due.de

SECTION: Statistics and Renewable Energy (Session 1)

The decarbonization of the heating sector represents one of the most pressing yet structurally complex challenges of the European energy transition. While national climate targets increasingly mandate municipal-level heat planning, the enormous heterogeneity of municipalities with respect to building stock conditions, renewable energy potentials, and socioeconomic contexts renders uniform policy approaches ineffective. This study develops a socio-energetic typology of 1,404 small and medium-sized towns in Germany, drawing on a multi-source dataset that integrates 67 indicators across four thematic dimensions: infrastructure, demographic and socioeconomic conditions, energy-demand side, and energy-supply side characteristics. Data are compiled from multiple official sources, including the German Federal Statistical Office and the Federal Institute for Research on Building, Urban Affairs and Spatial Development (BBSR), ensuring both comprehensiveness and reproducibility. Given the resulting high dimensionality, we employ a sequential two-step methodology: factor analysis reduces variable space to 13 latent factors capturing underlying structural variation across municipalities, which subsequently serve as inputs for a k-means cluster analysis identifying 13 internally homogeneous municipal archetypes. From these, four transition priority groups are derived, reflecting distinct combinations of urgency, fossil fuel dependency, socioeconomic capacity, and governance complexity. The analysis reveals that fossil fuel dependency, particularly reliance on natural gas, constitutes the dominant lock-in risk across German municipalities, that technical potential and socioeconomic capacity are frequently misaligned, and that settlement structure is a primary determinant of technology viability, especially for district heating network expansion. Beyond their diagnostic value, the identified archetypes provide a practical planning tool for intermunicipal cooperation, enabling municipalities with comparable structural profiles to benchmark their conditions, exchange transferable planning experience, and develop coordinated transition strategies that generate synergy effects. While the empirical application focuses on Germany, the analytical framework and underlying methodology are explicitly designed to be transferable to other European contexts where municipal heat planning is gaining policy salience.

10 Zensus 2031: Ein Werkstattbericht

René Söllner

Statistisches Bundesamt, Deutschland
rene.soellner@destatis.de

SECTION: Regionalstatistik (Session 1): Registerzensus

Der Konferenzbeitrag berichtet über den aktuellen Stand zur Vorbereitung des Zensus 2031. Die amtliche Statistik muss nach den Vorgaben der neuen europäischen Verordnung über europäische Bevölkerungs- und Wohnungsstatistiken (European Statistics on Population, ESOP) zur nächsten weltweiten Zensusrunde ab 2031 neue Datenanforderungen in den Bereichen Bevölkerung, Gebäude und Wohnen, Haushalte und Familien sowie Bildung und Arbeitsmarkt bedienen. Um die Belastung von Kommunen und Bevölkerung möglichst geringhalten, wird der Zensus 2031 erstmalig als registerbasiertes Verfahren durchgeführt. Das heißt im Vergleich zum Zensus 2022 werden noch stärker bereits vorhandene Statistik- und Verwaltungsdaten genutzt, die durch primärstatistische Elemente ergänzt werden. Die neuen Verfahren zur registerbasierten Ermittlung von Einwohnerzahlen wurden im Rahmen des Methodentests Bevölkerung umfangreich erprobt. Unter anderem wurde dort überprüft, ob Über- und Untererfassungen in den Melderegistern über den sogenannten Lebenszeichenansatz zuverlässig erfasst werden können.

Der Zensus 2031 ist ein gemeinsames Projekt des Statistischen Verbundes. Das Statistische Bundesamt berichtet über wichtige methodische Neuerungen, aktuelle Aufgabenschwerpunkte, Herausforderungen und das weitere Projektvorgehen.

11 Erprobung der Methodik eines Registerzensus - Eine Bayerische Perspektive

Paul Fink, **Stella Klein**, Katrin Kräck

Bayerisches Landesamt für Statistik, Deutschland
stellavklein@gmail.com

SECTION: Regionalstatistik (Session 1): Registerzensus

Mit dem Zensus werden Daten zur Bevölkerung und deren Zusammensetzung, Gebäuden und Wohnungen sowie Haushalten und Familien gewonnen. Um die Belastungen für Bürgerinnen und Bürger sowie für die Verwaltung zu reduzieren, ist eine Umstellung der Zensusmethodik von einem registergestützten Zensus, wie er in 2011 und 2022 durchgeführt wurde, hin zu einem registerbasierten Zensus notwendig. Zudem sollen Ergebnisse häufiger, aktueller und tiefer regional untergliedert zur Verfügung gestellt werden können. Seit der Entwicklung der Methodik des registergestützten Zensus in Deutschland, welche im Zensustest 2001 erprobt wurde, hat sich die Menge und Verknüpfbarkeit der in Deutschland vorliegenden Verwaltungsdaten so weit verändert, dass nun die registerbasierte Zensusmethodik eine Entlastung der Bürgerinnen und Bürger möglich macht.

Zur grundständigen Überprüfung der registerbasierten Methodik für Deutschland haben die statistischen Ämter des Bundes und der Länder auf Basis des Registerzensuserprobungsgesetzes einen Methodentest durchgeführt. Im Fokus stand hierbei insbesondere die Gewinnung von Informationen darüber, ob das Vorgehen dazu geeignet ist realitätsgerechte Bevölkerungszahlen zu ermitteln.

Der Vortrag beleuchtet diese Erprobungsphase aus Sicht des Bayerischen Landesamts für Statistik und stellt die dabei gewonnenen Erkenntnisse und Erfahrungen heraus. Zudem werden die identifizierten Herausforderungen dargestellt und erste Ansätze zur Optimierung aufgezeigt.

12 A Bayesian Distributional Approach to Qualified Local Rent Benchmarks

Georg Wiegleb

Landeshauptstadt Magdeburg, Deutschland
georg.wiegleb@stadt.magdeburg.de

SECTION: Rent and Rental Indices

Qualified regression-based rent benchmarks are subject to formal regulatory requirements while simultaneously aiming to meet scientific standards. Despite substantial methodological progress in recent years, their construction continues to involve important methodological challenges. In practice, these challenges are often addressed through pragmatic modeling strategies, including multi-stage procedures or separately specified model components. While such approaches may be well suited to applied settings, they can make it difficult to formulate a single coherent statistical model and to propagate uncertainty consistently across model components. A further methodological question concerns the construction and interpretation of rent ranges derived from regression-based benchmarks.

This talk discusses a Bayesian distributional approach to local rent benchmarks as a possible framework for addressing these methodological challenges in an integrated way. The approach builds on joint estimation, allowing central model components to be inferred within a structured model and uncertainty to be propagated consistently across them. Semi-parametric smoothers are used to capture complex covariate effects and spatial structures, while high-dimensional categorical covariates can be incorporated as regularized structured blocks. By jointly modeling multiple parameters of the conditional rent distribution, the approach may provide a basis for constructing rent ranges directly from the estimated distribution, with an interpretation that follows naturally from the model itself.

13 „Krise? Welche Krise?“ Zur Messung der Bezahlbarkeit von Mietwohnungen in Deutschland

Marco Schmandt

TU Berlin, Deutschland
m.schmandt@tu-berlin.de

SECTION: Rent and Rental Indices

Trotz der verbreiteten Wahrnehmung einer Wohnungskrise in Deutschland sind die gemessenen Mietkostenbelastungsquoten in Haushaltsbefragungen in den vergangenen 20 Jahren konstant geblieben. Vor diesem Hintergrund wird ein alternatives Maß der Bezahlbarkeit von Mietwohnungen entwickelt, das den Marktzugang zu öffentlich inserierten Mietwohnungen misst. Das Maß ist definiert als der Anteil der angebotenen Wohnungen, der entlang der lokalen Verteilung der Nettoeinkommen als bezahlbar gelten kann.

Eine notwendige Datengrundlage, um dieses Maß zu definieren, sind lokale Nettoeinkommensverteilungen von Mieterhaushalten. Diese werden auf Basis von Daten aus der Einkommensteuerstatistik und Haushaltsbefragungen als „kontrafaktische“ lokale Einkommensverteilungen konstruiert. Basierend auf diesen Verteilungen wird das verfügbare Wohnbudget durch eine Kombination aus Residualeinkommens- und Einkommensanteilsansatz definiert.

Das Bezahlbarkeitsmaß wird auf Ebene der Landkreise und kreisfreien Städte für den Zeitraum 2013 bis 2022 berechnet. Die Bezahlbarkeit, definiert über den Marktzugang, ist für Haushalte mit niedrigen Einkommen zurückgegangen, allerdings nur in Regionen mit überdurchschnittlich steigenden Mieten. Einpersonenhaushalte mit niedrigen Einkommen mit und ohne Kinder haben nahezu keinen Marktzugang.

14 Wirksamkeit regulatorischer Mietpreisgrenzen - Eine empirische Annäherung

Michael Windmann (1,2), Göran Kauermann (1), Oliver Trinkaus (2)

1: LMU München, Deutschland; 2: EMA-Institut für empirische Marktanalysen
michael.windmann@lmu.de

SECTION: Rent and Rental Indices

Die Einhaltung gesetzlicher Regulierungen von Wohnungsmieten, wie z.B. die Mietpreisbremse, steht aktuell im gesellschaftlichen, politischen und auch wissenschaftlichen Fokus. Aus empirischer Sicht zeigen wir Methoden und deren Grenzen, um „überhöhte“ Mieten mittels statistischer Modellierung zu identifizieren. Wir nutzen dazu exemplarisch Daten des Münchner Mietspiegels. Eine „überhöhte“ Miete verstößt dabei rein juristisch betrachtet nicht direkt gegen die Mietpreisbremse, was wir kurz skizzieren. Eine statistische Betrachtung erlaubt aber das Ausmaß von Mieten zu schätzen, die oberhalb gesetzlicher Schwellenwerte liegen. Ebenso kann untersucht werden, welche Wohnungstypen hier stärker betroffen sind und ob es regionale Auffälligkeiten gibt. Weiter zeigen wir den Einfluss der als „überhöht“ identifizierten Wohnungen im Mietspiegeldatensatz, wenn diese gelöscht oder auf den zulässigen Wert gekappt werden. Der Vortrag beleuchtet damit ein komplexes und juristisch kniffliges Thema aus empirischer Sicht.

15 Stichprobenbasierte Fortschreibungen von Mietspiegeln: Methodische Abwägungen zu Stichprobenplanung und Fragebogenumfang am Beispiel der Stadt Leipzig

Ingo Seifert, Martin Waschipky

Stadt Leipzig, Amt für Statistik und Wahlen, Deutschland
ingo.seifert@leipzig.de

SECTION: Rent and Rental Indices

Sollen Mietspiegel mit Hilfe von Stichproben fortgeschrieben werden, stehen potenziell folgeschwere, weil zum Zeitpunkt der Erhebung prinzipiell unumkehrbare, Entscheidungen in der methodischen Konzeption an: Wie groß soll die zu erhebende Stichprobe sein? Soll der „alte“ Fragebogen aus der ursprünglichen Erstellung des Mietspiegels im Ganzen wieder verwendet werden oder auf möglichst wenige relevante Fragen reduziert werden? Anhand der Fortschreibung des Leipziger Mietspiegels werden in unserem Vortrag potentielle Antworten auf diese Fragen aus praktikabler

wie statistischer Sicht dargestellt. Dabei gehen wir auch darauf ein, wie relevante Fragen zur Kosten der Unterkunft in der Konzeption der Mietspiegelfortschreibung adäquat berücksichtigt werden können. Anhand der für Leipzig erhobenen Daten illustrieren wir die Auswirkungen der methodischen Entscheidungen.

16 Ethik in der Statistikausbildung – Erfahrungsaustausch und Vernetzung

Walter J. Radermacher

LMU, Deutschland
radermacher.wa@lmu.de

SECTION: Competence Development: Data Literacy and Statistics (Session 1)

Die Arbeit von Forschenden und Fachleuten, die sich mit Daten befassen oder Modelle für maschinelles Lernen entwickeln, hat immer größere Auswirkungen auf den Einzelnen und die Gesellschaft. Dennoch werden die ethischen Implikationen dieser Arbeit oft erst nachträglich berücksichtigt. Der Grund: Die meisten Universitätsstudiengänge in datenbezogenen Disziplinen wie Datenwissenschaft, Statistik und Teilen der Informatik konzentrieren sich hauptsächlich auf die Vermittlung technischer Fähigkeiten. Infolgedessen sind die Studierenden auf die ethischen Fragen, mit denen sie konfrontiert werden, nur unzureichend vorbereitet – oder sind sich dieser Fragen gar nicht bewusst.

Diese Sitzung zielt darauf ab, diese Mängel anzugehen und die Vernetzung einer Gemeinschaft voranzutreiben, die daran interessiert ist, Datenethik in die Lehre zu integrieren. Es werden Erfahrungen und konkrete Projekte einer Roundtable-Arbeitsgruppe vorgestellt, in der Hochschulen aus Deutschland und Österreich seit Anfang 2026 zusammenarbeiten.

17 Bayesian Causal Estimation of German Fuel Price Policy: High Noon and the 2026 Tax Discount

Manuel Frondel (1,2), **Patrick Thiel** (1), **Colin Vance** (1,3)

1: RWI - Leibniz Institute for Economic Research, Germany; 2: Ruhr University Bochum; 3: Constructor University Bremen
vance@rwi-essen.de

SECTION: Empirical Economics and Applied Econometrics (Session 2)

In Spring 2026, the German government introduced two measures to increase fuel price transparency and reduce consumer prices. The High Noon regulation, effective April 1, restricts stations to a single price increase per day at 12:00. One month later, a two-month excise tax discount reduced the per-litre duty on petrol and diesel by €0.14 (€0.167 inclusive of VAT). This paper estimates the causal impact of both measures on retail fuel prices.

We employ a Bayesian Structural Time Series (BSTS) model within the potential outcomes framework for causal inference. The model constructs a probabilistic counterfactual – what prices would have prevailed absent each intervention – by fitting a state-space model to the pre-intervention period and projecting it forward. The specification includes a semi-local linear trend, weekly seasonality, and a regression component in which Rotterdam (or Brent) spot prices serve as an exogenous covariate. As Germany is a price-taker on wholesale energy markets, these

prices are unaffected by domestic policy and provide a real-time signal of counterfactual input costs, separating global market movements from policy effects.

A key advantage of the Bayesian framework is that posterior estimates from the comparable 2022 fuel tax discount can be encoded as informative priors, disciplining inference while the post-intervention series remains short. The full posterior over the time path of the effect, rather than a single average treatment effect, tracks the dynamic evolution of pass-through. Data cover daily median retail prices across German stations from 2023 through May 20, 2026.

Preliminary results indicate the High Noon regulation had no discernible effect on price levels, despite a transitory upward movement reported in the media. The tax discount appears largely passed through to consumers, consistent with the 2022 episode, though the pass-through rate shows early signs of dissipating – a pattern similarly documented for the earlier discount.

18 Electric bicycles and public transport tickets: Ownership and car use patterns

Nils Christian Hoenow, **Viola Helmers**, Eva Yang

RWI - Leibniz Institute for Economic Research, Deutschland
vhelmers@rwi-essen.de

SECTION: Empirical Economics and Applied Econometrics (Session 2)

Riding electric bicycles and using public transport are popular alternatives to private car use. Utilizing data from a 2022 survey of 6,285 participants in Germany, we examine who typically owns e-bikes or public transport tickets. We firstly employ regression analyses to identify correlations between individual characteristics and e-bike as well as ticket ownership, respectively. We find that e-bike owners tend to be older, earn higher incomes, and often reside in more rural areas. For public transport ticket ownership, these associations are largely reversed. Through latent class analyses, we identify distinct groups of e-bike and ticket owners. In a second step, we investigate associations between e-bike or ticket ownership and several car use measures through propensity score matching and regression analyses. We find that, compared to non-owners, owners of either alternative exhibit lower car use, with the difference being larger for public transport ticket owners.

19 Effects of retirement age policies on informal caregiving in Europe

Dörte Heger (1,2), Ingo W. K. Kolodziej (2,3), Luisa V. Licker (4,5), Arndt R. Reichert (4), **Ansgar Wübker** (2,5)

1: Hochschule Bochum, Germany; 2: RWI – Leibniz Institute for Economic Research, Germany; 3: Hochschule Fresenius – University of Applied Sciences, Germany; 4: Center for Health Economics Research Hannover (CHERH), Leibniz University Hannover, Germany; 5: Harz University of Applied Science, Germany
awuebker@hs-harz.de

SECTION: Empirical Economics and Applied Econometrics (Session 2)

Population ageing is placing pressure on pension systems, leading governments to implement reforms aimed at extending working lives and securing fiscal sustainability, most notably through increases in statutory retirement ages. While these policies are designed to stabilize public finances and labour supply, their broader social implications remain underexplored. This study investigates

how retirement age policies affect the provision of informal care using data from the Survey of Health, Ageing and Retirement in Europe (2004–2020) and exploiting country-specific early and full retirement age thresholds in twelve countries using a regression discontinuity design. We analyse a sample of approximately 21,600 individuals aged 50–70. Informal care is measured as providing care at least weekly (less intensive), and intensive informal care as providing daily care. Our results show that eligibility to early retirement significantly increases the likelihood of providing informal care, while no significant effects are observed at the full retirement age. These results indicate that early retirement eligibility frees up time resources that are reallocated to family caregiving. From a policy perspective, early retirement age reforms therefore have important spill-over effects on the availability of informal care and the functioning of long-term care systems. Extending working lives may unintentionally weaken family-based care capacities, potentially increasing demand for formal care services and public care expenditures. Sustainable ageing policy thus requires integrated reform strategies that jointly consider pension policy, labour market participation, and long-term care provision.

20 Warehouse Density and Local Air Quality: Causal Evidence from the Netherlands

Liam Vance (1), Kathrin Kaestner (2)

1: Independent Researcher, Germany; 2: RWI – Leibniz-Institut Für Wirtschaftsforschung
liam.jvance@web.de

SECTION: Empirical Economics and Applied Econometrics (Session 2)

PM2.5 is a leading source of morbidity and mortality in industrialized countries, accounting for around 239,000 premature deaths in the EU in 2022. While warehousing operations have been identified as a potential PM2.5 source, existing evidence is concentrated in the US context and largely correlational. We address this gap using a 20-year panel from industrial areas in the Netherlands to estimate the causal effect of hinterland warehouse density on local PM2.5 concentrations.

We employ a two-way fixed effects (TWFE) estimator alongside a doubly robust machine learning (DML) estimator as a mutual robustness check; their convergence strengthens the causal interpretation of the results. The Netherlands is a particularly relevant case study given its intense logistical expansion since the 1980s, distinctive topography, and high population density. Areas designated as industrial in 1989 form the identification backbone of the analysis, ensuring overlap between treated and control units while avoiding post-treatment bias in sample selection. The study is conducted at a fine spatial resolution of approximately 0.76 km², combining satellite-derived PM2.5 measurements with a dataset of nearly 27,000 classified warehouse buildings.

Descriptively, PM2.5 concentrations fell by around 28% over the 2000–2019 study period, notwithstanding steady growth in warehouse density. The econometric results indicate that transport-oriented warehouse density increases local annual PM2.5 by approximately 0.022 µg/m³ per additional warehouse per km². A cluster of ten such warehouses thus contributes roughly 4.4% of the WHO annual guideline value of 5 µg/m³. While moderate in absolute terms, this magnitude is policy-relevant in densely populated or heavily industrialized areas. Warehouses focused on storage rather than transshipment, by contrast, show smaller and statistically insignificant effects.

These findings suggest that warehouse siting decisions should differentiate between facility types according to their freight-activity profiles, and explicitly account for population exposure in surrounding areas.

21 Networked Spatial Effects in European Electricity Price Forecasting

Sultan Mahmud Chomon, Florian Ziel

University of Duisburg-Essen, Germany
sultan.chomon@uni-due.de

SECTION: Statistics and Renewable Energy (Session 2)

As European bidding zones are highly interconnected by physical transmission lines, spatial influences propagate across neighboring nodes through a network. It is reflected in the day-ahead electricity prices across European bidding zones, as the auction algorithm also uses information beyond each bidding zone's geographic boundary. To capture how this interconnection affects the neighboring bidding zone's electricity prices, we have used a metric graph to map the spatial coverage of information using a well-defined neighborhood measure. We propose the Networked Spatio-Temporal Model (NSTM), which maps irregular spatial nodes into an ordered network, enabling the systematic incorporation of neighborhood information. We implement the NSTM across 39 bidding zones covering the majority of European electricity markets in a high-resolution, streaming-forecasting setup. The model uses autoregressive, cross-hour, and seasonal effects, along with fuel and emission prices and day-ahead forecasts of fundamentals, as interconnected information to predict the day-ahead prices for each bidding zone. A Europe-wide study presented in this paper shows that the NSTM consistently outperforms traditional island-based pure local models. This paper provides a framework that demonstrates the critical role the networked structure plays in propagating information across interconnected markets and its vast implications on day-ahead electricity price forecasting.

22 Interpretable Functional Modeling of Electricity Market Auction Curves

Wenzhen Huang, Florian Ziel

Universität Duisburg-Essen, Germany
wenzhen.huang@uni-due.de

SECTION: Statistics and Renewable Energy (Session 2)

This study addresses the challenge of forecasting supply auction curves in day-ahead electricity markets. The increasing penetration of renewable energy sources introduces profound structural complexities to the grid. Notably, this influx fundamentally alters the shape and slope of the auction curve, driving frequent negative prices and shifting merit order dynamics. Such high volatility motivates the need to model the full dynamic shape of the supply curve rather than focusing solely on point price forecasts.

We propose an interpretable functional modeling framework for representing and forecasting auction curves across the full price spectrum. The approach constructs the cumulative supply curve from the bottom up by estimating the marginal bid density using a flexible basis function expansion. To guarantee economic reality, we strictly enforce curve monotonicity through a non-linear link function. This structure dynamically maps market covariates to the curve's shape, with all parameters optimized end-to-end via gradient descent. Furthermore, physics-informed constraints ensure that variables only influence the curve within their valid economic domains.

The forecasting study results show that the proposed framework can effectively capture key features of observed auction curves, outperforming benchmarks in out-of-sample accuracy while guaranteeing physical monotonicity. In particular, the model facilitates a proportional decomposition of the curve dynamics into economically meaningful drivers, offering precise insights into the role of renewable generation and available thermal capacity in shaping bidding behavior at every price level.

23 Voltage state estimation in electrical power distribution grids and optimal sensor location

Christine Müller

TU Dortmund, Deutschland
cmueller@statistik.tu-dortmund.de

SECTION: Statistics and Renewable Energy (Session 2)

In dynamic electrical power distribution grids, the voltage states are strongly time dependent, but unknown and of high importance for grid stability. The unknown voltage states can be estimated with power measurements using the Kalman filter. However, it is usually not possible to obtain these power measurements by sensors at all nodes of the grid. To get information at the other nodes, so-called pseudo-measurements using historical data or weather data must be used. These measurements have a much less precision than the measurements at the nodes with sensors. We consider here the problem of improving the voltage state estimation by determining the nodes of the grid, where the measurements should be taken with the high precision of sensors, and the nodes, where lower precision given by pseudo-measurements is sufficient. Since the covariance matrices of state estimation and state prediction are time dependent and are additionally of complicated form, we propose a simplified design criteria based on a simplified matrix. Although this matrix is still time dependent, it is much easier to treat and methods developed for other networks can be applied.

24 Distributional Load Forecasting for Off-Grid Energy Systems

Alexandra Fafali A. Kudjawu (1), Johannes Schwenzer (1), Simone Maxand (1,2)

1: Europa-Universität Viadrina, Germany; 2: Berlin School of Economics, Germany
kudjawu@europa-uni.de

SECTION: Statistics and Renewable Energy (Session 2)

Despite the expansion of off-grid energy systems as a solution to close the universal electrification gap and advance Sustainable Development Goal 7, many households still lack access to electricity, underscoring the continued importance of research in this area. While prior studies have focused on point-load forecasting for off-grid energy systems, distributional load forecasting remains underexplored. Using real-world mini-grid data, with special focus on an African perspective, a variety of probabilistic forecasting methods (classical and novel) for short-term electricity demand forecasting were employed. Among these methods are flexible semi-parametric regression methods, quantile regressions, and deep neural network architectures, which have been shown to capture the unique characteristics of mini-grids such as irregular, low and highly variable demand spikes. Finally, based on forecasting metrics such as CRPS, coverage, sharpness, and pinball loss, the best-performing method was identified, and differences among the methods were highlighted.

25 Regionaler Flächenbedarf für Siedlungs- und Verkehrsflächen

Jana Hoymann (1), Nadine Blätgen (1), Sylvie Dugay (1), Florian Bernardt (2), Philip Ulrich (2)

1: Bundesinstitut für Bau-, Stadt- und Raumforschung, Deutschland; 2: Gesellschaft für wirtschaftliche Strukturforschung
jana.hoymann@bbr.bund.de

SECTION: Regionalstatistik (Session 2): Regionalplanung

In Deutschland steigt die Konkurrenz um die knappe Ressource Fläche. Neben der Expansion von Siedlungs- und Verkehrsflächen ergeben sich neue Raumansprüche aus der Energiewirtschaft, veränderter Rahmenbedingungen in der Agrarproduktion, dem Naturschutz und Anpassungsmaßnahmen an den Klimawandel. Das flächenpolitische Ziel, die Flächeninanspruchnahme durch Siedlung und Verkehr auf 30 Hektar pro Tag bis 2020 zu reduzieren wurde klar verfehlt. Das fortgeschriebene Ziel verlangt eine Reduktion auf unter 30 Hektar am Tag bis 2030.

Vor dem Hintergrund dieses Ziels, stellt sich die Frage, wie sich die zukünftige Flächeninanspruchnahme durch Siedlung und Verkehr differenziert abbilden lässt. In diesem Beitrag wird ein Modellansatz vorgestellt, der die kleinräumige Quantifizierung und Verortung der Flächenneuinanspruchnahme durch Siedlung und Verkehr in Deutschland ermöglicht.

Das Modell QMORE_PR bildet die Entwicklung der Siedlungs- und Verkehrsfläche im demografischen und ökonomischen Kontext auf Ebene der Landkreise und kreisfreien Städte in Deutschland ab. Die Grundlagen des Modells liegen in der Identifizierung überwiegend nachfragebezogener Einflussfaktoren (z.B. Wirtschaftswachstum, sektoraler Strukturwandel, Demografie) und deren konsequenter Verknüpfung mit zu erwartenden gesamtwirtschaftlichen und regionalen Entwicklungen.

Im Land-Use-Scanner des BBSR wird auf Basis dieser Projektionen eine kleinräumige Verortung der in QMORE_PR ermittelten Flächenbedarfe vorgenommen. Der Land Use Scanner ist ein räumlich explizites Simulationsmodell, dass auf Basis von räumlichen Standortfaktoren Flächenbedarfe auf geeigneten Flächen verteilt.

Die Kombination dieser beiden Modelle ermöglicht eine kleinräumige Projektion der Flächendynamik ausgehend von verschiedenen Szenarien der Siedlungs- und Verkehrsflächenentwicklung. Die Ergebnisse liefern damit eine empirische Grundlage zur Bewertung einer nachhaltigen Flächenpolitik auf regionaler und kommunaler Ebene.

Die Flächeninanspruchnahme für Siedlungs- und Verkehrsflächen beträgt zunächst bis 2030 43,5 ha pro Tag. In den darauffolgenden Fünfjahreszeiträumen geht dieser Veränderungswert sukzessive zurück, bis etwa 21,5 ha pro Tag von 2045 bis 2050. Das flächenpolitische Ziel, die Flächeninanspruchnahme durch Siedlung und Verkehr auf unter 30 Hektar pro Tag zu reduzieren wird damit erst nach 2040 erreicht.

26 Kleinräumige Analyse der Grünraumversorgung in Nordrhein-Westfalen

Christoph Alfken, Anna Grüne

IT.NRW - Landesamt für Statistik Nordrhein-Westfalen
christoph.alfken@it.nrw.de

SECTION: Regionalstatistik (Session 2): Regionalplanung

Die Versorgung der Bevölkerung mit öffentlich zugänglichen Grünräumen stellt vor dem Hintergrund von Trends wie dem Klimawandel und einer zunehmend alternden Bevölkerung eine zentrale Herausforderung der Daseinsvorsorge dar. Viele Städte und Gemeinden in Nordrhein-Westfalen sind hochverdichtet, gleichzeitig gibt es auch ländlich geprägte Kommunen. Der Beitrag untersucht die Versorgung mit Grünräumen für die Naherholung in den Kommunen in Nordrhein-Westfalens. Methodisch wird eine auf der Two-Step Floating Catchment Area-Methode (2SFCA) basierende Auswertung eingesetzt, um den Versorgungsgrad und potenziellen Nutzungsdruck auf Grünräume in den Kommunen Nordrhein-Westfalens abzubilden. Im Fokus steht dabei die differenzierte Betrachtung der Erreichbarkeit und Verfügbarkeit von Grünräumen für unterschiedliche Bevölkerungsgruppen.

Als Datengrundlage dienen ausgewiesene Flächen aus den amtlichen Daten zur Flächenstatistik sowie Bevölkerungsdaten des Zensus 2022 auf Gitterzellenebene. Auf dieser Basis werden räumliche Disparitäten der Grünversorgung auch unterhalb der Gemeindeebene sichtbar.

Der Beitrag geht insbesondere der Frage nach, wie die Versorgung unterschiedlicher Bevölkerungsgruppen – differenziert nach z. B. der Altersstruktur – ist und inwiefern die Ergebnisse kleinräumiger Analysen neue Erkenntnisse für die amtliche Statistik liefern können. Darüber hinaus werden besonders belastete Räume identifiziert und Beispiele für eine vergleichsweise günstige Versorgungssituation herausgearbeitet, wobei zugleich der Mehrwert kleinräumiger amtlicher Daten verdeutlicht und die Anwendung neuer Methoden in der amtlichen Statistik erprobt wird.

27 Regionaler Flächenbedarf für den Ausbau erneuerbarer Energien

Philip Ulrich (1), Florian Bernardt (1), Nadine Blätgen (2), Jana Hoymann (2), Sylvie Dugay (2)

1: GWS mbH, Deutschland; 2: Bundesinstitut für Bau-, Stadt- und Raumforschung
ulrich@gws-os.com

SECTION: Regionalstatistik (Session 2): Regionalplanung

Die fortschreitende Energiewende in Deutschland bringt einen signifikanten Bedarf an zusätzlicher Fläche mit sich, insbesondere für den Ausbau von Photovoltaik (PV) und Windkraft. Vor dem Hintergrund des Ziels, den Flächenverbrauch bis 2050 deutlich zu reduzieren, stellt sich die Frage, wie sich die zukünftige Flächeninanspruchnahme durch erneuerbare Energien räumlich differenziert abbilden lässt. In diesem Beitrag wird ein Modellansatz vorgestellt, der die kleinräumige Quantifizierung und Verortung der Flächenneuinanspruchnahme im Kontext des Ausbaus erneuerbarer Energien in Deutschland ermöglicht.

Das Modell QMORE_PR der GWS kombiniert regionale Ausbaupfade für PV- und Windenergieanlagen mit spezifischen Kennwerten zur Flächeninanspruchnahme je installierter Megawatt-Leistung. Dabei wird zwischen der gesamten Windparkfläche und der tatsächlich versiegelten Fläche unterschieden, um differenzierte Aussagen über die reale Flächenbeanspruchung treffen zu können.

Zusätzlich wird die Wirkung von Repowering-Maßnahmen berücksichtigt: Durch den Ersatz älterer Anlagen mit moderner, leistungsstärkerer Technologie kann der Flächenbedarf je erzeugter Energieeinheit reduziert werden, während sich gleichzeitig die räumliche Verteilung der Anlagen verändert.

Im Land-Use-Scanner des BBSR wird auf Basis dieser Projektionen eine kleinräumige Verortung der in QMORE_PR ermittelten Flächenbedarfe vorgenommen. Der Land Use Scanner ist ein räumlich explizites Simulationsmodell, dass auf Basis von räumlichen Standortfaktoren Flächenbedarfe auf geeigneten Flächen verteilt. Die Kombinationen dieser beiden Modelle ermöglicht eine kleinräumige Projektion der Flächendynamik ausgehend von verschiedenen Szenarien des EE-Ausbaus. Die Ergebnisse liefern damit eine empirische Grundlage zur Bewertung der Zielkonflikte zwischen Energiewende, Flächenschutz und Raumordnung und bieten Ansatzpunkte für eine nachhaltige Flächenpolitik auf regionaler und kommunaler Ebene.

Die Flächeninanspruchnahme für PV-Freiflächenanlagen beträgt laut den Projektionen zunächst bis 2030 17 ha pro Tag. In den darauffolgenden Fünfjahreszeiträumen geht dieser Veränderungswert sukzessive zurück. Für die Windparkfläche steigt der Flächenveränderungswert noch bis Mitte der 30er Jahre an. Nach 2045 besteht kein Bedarf mehr an zusätzlichen Windparkflächen. Für beide Technologien werden genügend Flächen durch Stilllegungen frei, um neue Anlagen zu installieren.

28 GeoStat: Ein offener, nutzerfreundlicher Zugang zu amtlichen georeferenzierten Statistikdaten – Neue Auswertungsmöglichkeiten am Beispiel der Daten des statistischen Unternehmensregisters

Hanna Hoffmann (1), Christoph Alfken (1), Katja von Eschwege (2)

1: Statistisches Landesamt Nordrhein-Westfalen, IT.NRW; 2: Statistisches Bundesamt, Deutschland
katja.eschwege@destatis.de

SECTION: Regionalstatistik (Session 2): Regionalplanung

Wie können amtliche georeferenzierte Statistikdaten noch besser zugänglich, nutzbar und vernetzt werden? Der GeoStat-Prozess – inspiriert vom bewährten RegioStat-Prozess im Statistischen Verbund – schafft hierfür eine innovative Lösung: Er bündelt bisher verstreute, georeferenzierte Daten der Statistischen Ämter des Bundes und der Länder in der Regionaldatenbank Deutschland. Damit wird ein zentraler, nutzerfreundlicher Zugang geschaffen, der Daten in verschiedenen Formaten (CSV, Excel, GeoJSON) und über standardisierte APIs bereitstellt – und das in kleinräumiger Form (INSPIRE-Gitter) und unterschiedlicher fachlicher Tiefe (z.B. WZ-Abschnitte).

Zudem bietet das Datenangebot neue Möglichkeiten für Analysen und Visualisierungen. Die Daten aus dem statistischen Unternehmensregister zeigen beispielhaft den konkreten Mehrwert für Nutzende. Kleinräumige Auswertungen zu Niederlassungen nach WZ-Klassifizierung machen regionale Wirtschaftsstrukturen sichtbar. So lassen sich z. B. Standortentscheidungen fundieren oder kommunale Förderprogramme zielgenau ausrichten.

Der Vortrag stellt vor, wie der GeoStat-Prozess die Datenbeschaffung vereinfacht und gleichzeitig neue Analysemöglichkeiten für Verwaltung, Forschung und Wirtschaft eröffnet. Es werden erste kartografische Darstellungen kleinräumiger Daten des statistischen Unternehmensregisters für ausgewählte Wirtschaftsbereiche gezeigt.

29 Statistik im Spannungsfeld zwischen Theorie und Praxis: Zur Erinnerung an Peter von der Lippe

Andreas Kladroba (1), Markus Zwick (2)

1: Stifterverband für die Deutsche Wissenschaft; 2: Statistisches Bundesamt
andreas.kladroba@stifterverband.de

SECTION: Competence Development: Data Literacy and Statistics (Session 2)

Der Zufall wollte es, dass sich in dem Jahr, in dem die Statistische Woche in Essen stattfindet, der Todestag von Peter Michael von der Lippe zum zehnten Mal jährt. Hier in Essen hat er über 30 Jahre geforscht und gelehrt.

Vielen ist Peter von der Lippe als streitbarer Fachmann für Wirtschafts- und Sozialstatistik ein Begriff. Sein „rotes Buch“ war über viele Jahre das deutschsprachige Standardwerk zur Wirtschaftsstatistik.

Wir als seine akademischen ‚Kinder‘ kennen Peter von der Lippe aber auch als unermüdlichen Kämpfer für „sein“ Fach, die Statistik. Daher möchten wir die Gelegenheit nutzen und in diesem Vortrag an ihn erinnern.

Wir werden dabei drei wesentliche Aspekte ansprechen, die uns nicht zuletzt auch im Data Literacy Ausschuss der DStatG nach wie vor beschäftigen. Mit dem Blick zurück und einigen der damaligen Thesen, werden wir versuchen den gegenwärtigen Stand der Diskussion zu bewerten:

1. Vom „Niedergang“ zur Data Science: Die Entwicklung der Statistik

Vor allem im Rahmen der sogenannten Bologna-Reform sind viele Statistikveranstaltungen stark gekürzt und ausgedünnt worden. Peter von der Lippe sah darin einen „unaufhaltsamen Niedergang“ des Faches Statistik. Heute stehen wir im Zeitalter der „Data Science“. Ist von der Lippe damit widerlegt oder sehen wir nur die logische Konsequenz des „Niedergangs“?

2. Wer „macht“ eigentlich Statistik?

Während Peter von der Lippe noch die Spanne der Statistik“macher“ zwischen „Nur-“ und „Auch-Statistikern“ sah (mit beiden hatte er so seine Probleme), sagt heute die Data Literacy Charta: Jeder (sollte es machen).

3. Digitalisierung in der Statistikausbildung

Peter von der Lippe zeigte sich kritisch gegenüber dem aufkommenden „E-Learning“ (wie es damals noch hieß) in der Statistik. Heute sind wir weiter, aber gilt das nur technisch oder auch didaktisch?

Der Vortrag ist ein Beitrag zur von Berger et al (2026) angestoßenen Diskussion im AStA-WiSoStA und es ist vorgesehen die verschriftlichte Form hier ebenfalls zu publizieren.

30 Statistik und Data Literacy mit jamovi: Einsatzszenarien in der Hochschullehre

Diana Rauwolf, Erhard Cramer, Udo Kamps, Maria Kateri

Institut für Statistik und Wirtschaftsmathematik, RWTH Aachen
rauwolf@isw.rwth-aachen.de

SECTION: Competence Development: Data Literacy and Statistics (Session 2)

Datenbasierte Entscheidungen haben in Wissenschaft, Wirtschaft und Gesellschaft eine zunehmend größere Bedeutung und erfordern ein fundiertes Verständnis statistischer Denkweisen und Methoden. In der Statistikausbildung an Hochschulen stellt die Arbeit mit möglichst realitätsnahen Fragestellungen und Beispielen einen zentralen Aspekt dar, um Studierende bei der Entwicklung der erforderlichen Kompetenzen zu einer erfolgreichen Anwendung statistischen Wissens gezielt zu unterstützen.

Zur Realisierung dieses Ausbildungsziels wurde am Institut für Statistik und Wirtschaftsmathematik (ISW) der RWTH Aachen das Projekt jamovi.RWTH auf Basis der Open-Source Software jamovi umgesetzt. Ziel der Lernumgebung ist es, Studierenden statistische Analysen in einer direkten und niederschweligen Weise zugänglich zu machen, Lerninhalte zu vertiefen und so den Erwerb praxisrelevanter Kompetenzen im Bereich statistischer Methodik und Data Literacy sicherzustellen. Die vorhandenen Angebote sind auf die Inhalte der jeweiligen Lehrveranstaltungen abgestimmt – aber nicht darauf beschränkt. Die Verknüpfung zu den Lehrinhalten erfolgt durch speziell entwickelte jamovi-Module. jamovi.RWTH bietet in Studiengängen mit Statistikanteilen ein breites Spektrum an konkreten Anwendungsmöglichkeiten und ermöglicht es Studierenden von Beginn an, praxisbezogene Erfahrungen mit Datensätzen zu gewinnen, ohne eine Programmiersprache zu nutzen. jamovi.RWTH kann auch gezielt dazu eingesetzt werden, das Erlernen der Software R durch die Programmierung einführender statistischer Verfahren zu unterstützen.

Im Vortrag werden Einsatzmöglichkeiten von jamovi.RWTH in der Lehre, insbesondere für Studierende der Wirtschaftswissenschaften, des Wirtschaftsingenieurwesens und der Informatik vorgestellt. Das Lehrkonzept kann problemlos modifiziert und auf andere Statistikausbildungen übertragen werden, um die Lehre in den Bereichen Statistik und Data Literacy praxisnah und kompetenzorientiert zu gestalten.

Das Projekt wurde aus Mitteln der RWTH zur Verbesserung der Infrastruktur im Bereich Lehre gefördert.

31 Automating Statistical Exercise Feedback with LLMs: The sparrow Platform as Infrastructure for AI-Assisted Learning

Sigbert Klinke, Kleio Chrysopoulou Tseva

Humboldt-Universität zu Berlin, Deutschland
sigbert@wiwi.hu-berlin.de

SECTION: Competence Development: Data Literacy and Statistics (Session 2)

Building on TutorIAIn as a system for AI-assisted feedback in statistics education, we present 'sparrow', a lightweight, web-based platform for automated correction and feedback generation in statistical education, built on a robust and scalable preprocessing pipeline. Developed at Humboldt-Universität zu Berlin, 'sparrow' hosts nearly 1,000 exercises from Statistics I and II

- covering descriptive statistics, probability, estimation, and hypothesis testing - delivered via a simple interface. Unlike full adaptive learning systems, 'sparrow' focuses on efficient, LLM-assisted correction: students submit solutions, and the system uses large language models (LLMs) to evaluate correctness and generate explanations, while preserving pedagogical integrity through ground-truth validation.

The technical backbone of 'sparrow' lies in its preprocessing pipeline, which prepares exercises for dynamic delivery and intelligent feedback. Input files - originally in LaTeX or R Markdown - are first converted to HTML using pandoc and exams::exams2html, respectively. The pipeline then applies LLMs to: (1) extract the number of subtasks per exercise, (2) assign hierarchical metadata (topic → subtopic → concept → variant), and (3) generate multilingual versions (17 languages, including German, English, Greek, and Chinese). Exercises with a high number of subtasks are automatically excluded from the system to ensure manageable and high-quality learning experiences. This structured metadata enables targeted exercise selection and supports future personalization. The system's architecture is minimal and secure: a Flask server mediates LLM access via a reverse proxy, restricting usage to the university network. All student interactions are stateless, with no persistent data storage. The LLMs are used solely for feedback generation and solution checking—not for content creation or adaptive pathing.

'sparrow' demonstrates that a focused, well-structured preprocessing pipeline—combined with controlled LLM integration—can efficiently support scalable, accurate, and privacy-preserving automated correction. It offers a practical, maintainable alternative to complex AI-driven platforms, ideal for statisticians seeking to automate feedback without sacrificing control or correctness.

32 Beyond Weekly Exercises: AI-Assisted Feedback for Independent Statistics Learning

Kleio Chrysopoulou Tseva, Helena Julia Fikus, Sigbert Klinke

Humboldt-Universität zu Berlin, Germany
kleio.chrysopoulou.tseva@hu-berlin.de

SECTION: Competence Development: Data Literacy and Statistics (Session 2)

Preparing for statistics exams requires more than attending lectures and weekly exercise classes. Students need opportunities to test their understanding independently, receive feedback on their reasoning, and identify gaps before the exam. This creates a need that is difficult to meet at scale through traditional teaching formats alone.

Over the Summer Semester 2026, we integrated TutorIAIn into our statistics course as a supplementary tool for exam preparation. TutorIAIn is a web-based platform developed at Humboldt University of Berlin by Prof. Dr. Joachim Gassen that provides AI-generated feedback on open-ended questions and quizzes. Crucially, its responses are based on instructor-developed solutions and feedback guidelines, making the feedback substantially more targeted than what students would obtain from general-purpose AI tools.

During the semester, students work through guided exercises in weekly sessions. TutorIAIn is designed to complement this structure by giving students a way to practice independently in the lead-up to the exam: they can submit answers to statistics problems and receive immediate, individualized feedback on their reasoning, at any time, without instructor involvement or waiting for an e-mail response.

At the end of the semester, we will collect student feedback to assess how the tool was received and whether students found it useful for exam preparation. We present the course design, the

rationale for integrating TutorIAIn, and our expectations regarding its impact on self-directed learning in statistics.

This contribution invites a discussion on how AI-assisted feedback tools can be embedded meaningfully into existing course structures. Not to replace teaching, but to extend its reach beyond the classroom.

33 A Structural Matrix Autoregressions Framework for International Spillovers

Ignacio Moreira Lara (1,2), Jan Prüser (2), Christoph Hanck (1)

1: Universität Duisburg Essen, Germany; 2: Technische Universität Dortmund
ignacio.moreira-lara@tu-dortmund.de

SECTION: Young-Academics Mini-Symposium: Bayesian Econometrics (Session 1)

Structural identification in Vector Autoregression models faces a fundamental challenge: the number of structural restrictions required to identify shocks proliferates rapidly as the number of variables and their interactions increase. This becomes particularly acute when extending to multi-country settings where country-specific shocks must be separately identified. Existing approaches in the literature either require an impractical number of restrictions or impose restrictive assumptions that limit their applicability to large-scale multi-country problems. We propose Bayesian Structural Matrix Autoregressions (BSMAR), a novel framework that solves these limitations by exploiting a bilinear matrix structure to jointly identify shocks originating from different countries while reducing the number of required structural restrictions. The approach incorporates economically meaningful restrictions and employs efficient sampler methods to address the high dimensional nature of the data. We demonstrate the method's capabilities through an empirical application in which we simultaneously identify a large number of country-specific supply and demand shocks using data from OECD and key partner countries. The application recovers sensible impulse responses and reveals strong levels of connectedness across countries. We derive connectivity measures to quantify the magnitude of international spillovers and highlight the model's ability to capture shock transmission in a parsimonious yet interpretable way. We also provide possible extensions that relax restrictive assumptions in the baseline model.

34 What can we learn from and about sufficient macro statistics under set-identification?

Martin Fankhauser

Bocconi, Italy
martin.fankhauser@unibocconi.it

SECTION: Young-Academics Mini-Symposium: Bayesian Econometrics (Session 1)

Modern tools for macroeconomic policy evaluation and causal inference often rely on sufficient statistics that are nonlinear, non-bijective functions of impulse responses or other parameters; also historical decompositions and unit-effect normalizations can be expressed in the same way. We first show how to adjust posterior inference on such statistics via a simple reweighting step so the induced prior is corrected to match a desired single prior. Eliciting the latter, however, requires precise probability statements for every event involving the statistic, typically stronger assumptions

than intended. Moreover, updating a single precise prior via Bayes' rule does not guarantee yielding proper posterior probabilities. We therefore propose a framework that constructs classes of priors on the statistic of interest in a novel and easily operationalizable way. We show in two application how these novel identification strategies can substantially sharpen inference.

35 Micro-based SVAR Identification

Annika Camehl (1), Maximilian Schröder (2)

1: Erasmus University Rotterdam; 2: European Central Bank
camehl@ese.eur.nl

SECTION: Young-Academics Mini-Symposium: Bayesian Econometrics (Session 1)

This paper introduces a composite likelihood approach to structural identification in SVAR models that integrates micro- and macro-level estimates. We assume that a common structural elasticity exists and estimate it jointly using macro and micro likelihood components. Micro treatment effects enter as probabilistic restrictions rather than fixed calibrations, preserving uncertainty and allowing for potential tension between data sources. Joint estimation improves recovery of structural parameters and reduces multimodality. Applications to tax and hurricane shocks show that integrating micro information sharpens and stabilizes aggregate impulse responses.

36 Assessing the Effects of Monetary Shocks on Macroeconomic Stars: A SMUC-IV Framework

Jan Prüser

EBS University, Deutschland
J.Prueser@gmx.net

SECTION: Young-Academics Mini-Symposium: Bayesian Econometrics (Session 1)

This paper proposes a structural multivariate unobserved components model with external instrument (SMUC-IV) to investigate the effects of monetary policy shocks on key U.S. macroeconomic stars namely, the potential GDP growth, trend inflation, and the neutral interest rate. A key feature of our approach is the use of an external instrument to identify monetary policy shocks within the multivariate unobserved components modelling framework. We develop an MCMC estimation method to facilitate posterior inference within our proposed SMUC-IV framework. In addition, we propose a marginal likelihood estimator to enable model comparison across alternative specifications. Our empirical analysis shows that contractionary monetary policy shocks have negative effects on the macroeconomic stars, highlighting the non-zero long-run effects of transitory monetary policy shocks.

37 Predicting Intraday Price Direction with Conformal Methods

Qingyao Lai (1), Roxana Halbleib (1), Winfried Pohlmeier (2)

1: Faculty of Economics and Behavioural Sciences, University of Freiburg; 2: Department of Economics, Zeppelin University
qingyao.lai@vwl.uni-freiburg.de

SECTION: Empirical Economics and Applied Econometrics (Session 3)

This paper studies the intraday predictability of return-sign and the associated uncertainty quantification using conformal inference with high-frequency trade and quote data. We focus on upward, flat, and downward price movements at three sampling frequencies during the opening, mid-day, and closing periods. As base classifiers, we implement multinomial logistic regression and XGBoost, and evaluate their point predictions using standard classification metrics.

Beyond point prediction, we aim to provide confidence measures through conformal inference, specifically Mondrian conformal prediction. Since conformal methods rely on exchangeability, their application to high-frequency financial data is challenging due to temporal dependence and pronounced intraday heterogeneity. We therefore examine not only the performance of the two base classifiers, but also how empirical coverage varies with the calibration-window length and the choice of conformity score.

Consistent with well-documented findings in the literature, our results show that directional signals are strongest during the opening hour, with downward movements exhibiting relatively stable predictability across sampling frequencies; flat price changes dominate at midday, while directional predictability becomes more fragile toward the end of the trading day. After applying Mondrian conformal prediction, coverage differs across classes, frequencies, and intraday periods. Target coverage is achieved in most cases, except for the flat class at 60s frequency during the opening hour. Among the models, XGBoost delivers the most favorable validity-efficiency trade-off, producing smaller prediction sets while maintaining high conditional coverage.

38 When OPEC Speaks: The Dollar's Evolving Reaction to Oil Supply News in the Age of Shale

Bayram Oruc (1), Sebastian R  th (2), Wouter Van der Veken (3), Ren Zhang (4)

1: Statistisches Bundesamt, Deutschland; 2: University of Erfurt; 3: National Bank of Belgium; 4: Texas State University
bayram.oruc@destatis.de

SECTION: Empirical Economics and Applied Econometrics (Session 3)

Over the past two decades, the dollar-oil correlation moved from negative to positive as the U.S. became a net petroleum exporter. We study whether the dollar's response to OPEC announcements changed with this transition. In a daily event study, the reaction to lower-than-expected oil supply news turns from depreciation to appreciation. Structural time-varying-parameter VARs indicate that the reversal is gradual and parallels a terms-of-trade response that stops deteriorating. Interacted panel local projections tie this time variation to U.S. petroleum trade and response differentials between commodity importer-exporter currencies. The evidence suggests a shale-era change in the transmission of oil supply news.

39 Public Attention to Monetary Policy: Asymmetric Responses in Google Search Behavior Following Interest Rate Changes in the Euro Area

Florian Schütze

Helmut Schmidt Universität, Deutschland
florian.schuetze.hsu@gmail.com

SECTION: Empirical Economics and Applied Econometrics (Session 3)

Monetary policy decisions play an important role in shaping the economic landscape and fuel wide public debates. This paper explores the public perception of such decisions with the help of Google Trends to explicate the link between policy fluctuations and internet search intensities related to key economic variables. The analysis focuses on search activity surrounding monetary policy decision dates and compares the responses to interest rate increases and decreases. The results indicate a statistically significant asymmetric reaction in public search actions. In the majority of the euro area's founding countries, interest rate reductions have been observed to result in a more pronounced increase in search intensity compared to the effects of interest rate increases. This finding contradicts the research hypothesis that search intensity reacts equally to both types of policy changes. The findings persist when the observation window is extended beyond the policy decision date and when the sample is restricted to the period following the global financial crisis. These findings contribute to the expanding body of literature that utilizes digital trace data to measure public attention and economic expectations. The authors posit that expansionary monetary policy decisions elicit a more robust information-seeking behavior among the public, underscoring the pivotal role of monetary policy in shaping public awareness of economic conditions, possibly amplified by media.

40 Unobserved Components with Stochastic Volatility in Trend (UC-SVT)

Gabriel Arce-Alfaro (1), Dimitris Korobilis (2)

1: Central Bank of Ireland, Ireland, Republik; 2: University of Glasgow, Scotland
gabriel.arcealfaro@centralbank.ie

SECTION: Empirical Economics and Applied Econometrics (Session 3)

We propose a new unobserved components model for inflation in which changes in inflation volatility can affect the evolution of trend inflation directly. The model is designed to study whether periods of heightened inflation uncertainty are associated with changes in the permanent component of inflation, rather than only with larger short-run fluctuations. We estimate the model using Markov chain Monte

Carlo methods and also allow the effect of volatility on trend inflation to vary over time. Using U.S. inflation data from 1959 onward, we find that higher volatility is associated with downward movements in trend inflation for CPI inflation, while the evidence is weaker for PCE inflation. In an out-of-sample forecasting exercise, the model performs well for headline inflation at medium and long horizons, especially during more volatile periods. These forecasting gains should be interpreted

as evidence in favour of the overall specification, rather than as evidence on a single mechanism in isolation.

41 Community-oriented urban policy and local earnings – a complicated relationship

Uwe Neumann

RWI-Leibniz-Institut für Wirtschaftsforschung, Deutschland
uwe.neumann@rwi-essen.de

SECTION: Regionalstatistik (Session 2): Beschäftigung und Einkommen

The literature on regional agglomeration suggests that local economic revitalisation is likely to involve a rise in local wages. In the context of urban regeneration, community-oriented policy envisages to improve prosperity among the residential population of deprived neighbourhoods. Yet, due to an ever-increasing preference of households to reside at central locations this policy may spur gentrification if outsiders are attracted to new jobs and upgraded housing environments. Using Germany as a case study, the analysis explores whether local economies have received a boost that may have affected household sorting and local household income during the past two decades. The study reveals no considerable shift in sorting that would indicate gentrification. With a view to income over the past decade in programme areas local households with a low income have performed slightly better than their counterparts elsewhere. Moderate funding of urban regeneration in combination with support to local communities is not capable of providing a remarkable boost, but it may bring about improvements for the residential population without accelerating gentrification.

42 Aufbau eines regionalen Gesundheitspersonalmonitorings zur kleinräumigen Planung und Steuerung personeller Ressourcen

Anja Afentakis

Statistisches Bundesamt, Deutschland
anja.afentakis@destatis.de

SECTION: Regionalstatistik (Session 2): Beschäftigung und Einkommen

Vor dem Hintergrund des demografischen Wandels erhöht sich in den kommenden Jahren der Bedarf an qualifiziertem medizinischen Fachpersonal. Gelingt es in Deutschland eine ausreichende Zahl an medizinischen Fachkräften zu beschäftigen, um den steigenden Bedarf zu decken?

Am 1. Oktober 2021 ist das Gesetz über die Statistiken zu Gesundheitsausgaben und ihrer Finanzierung, zu Krankheitskosten sowie zum Personal im Gesundheitswesen (Gesundheitsausgaben- und -personalstatistikgesetz – GAPStatG) in Kraft getreten. Im §5 wird das Statistische Bundesamt beauftragt, ein regionales Gesundheitspersonalmonitoring für die Bereiche Öffentlicher Gesundheitsdienst, ambulante und (teil-)stationäre Pflegeeinrichtungen sowie Krankenhäuser und Vorsorge- oder Rehabilitationseinrichtungen aufzubauen und fortzuführen. Das regionale Gesundheitspersonalmonitoring schafft eine Grundlage um zu beurteilen, ob die Bevölkerung in Deutschland flächendeckend durch qualifiziertes Gesundheitspersonal adäquat versorgt wird. Durch das Monitoring können Regionen identifiziert werden, in denen ein Mangel oder Überschuss an qualifiziertem Gesundheitspersonal besteht.

Das Statistische Bundesamt nutzt für das regionale Gesundheitspersonalmonitoring vorrangig bereits erhobene Daten, die im Statistischen Verbund durch Vollerhebungen in regionaler Tiefe erfasst wurden. Hierzu zählen die Pflegestatistik sowie die Grund- und Diagnosedaten der Krankenhäuser und Vorsorge- oder Rehabilitationseinrichtungen des Bundes und der Länder. Ausschließlich das Personal im Öffentlichen Gesundheitsdienst wird vom Statistischen Bundesamt für das regionale Gesundheitspersonalmonitoring durch eine jährlich Vollerhebung erfasst. Neben den genannten Hauptdatenquellen werden auch weitere Daten wie beispielsweise die Bevölkerungsfortschreibung der Statistischen Ämter des Bundes und der Länder genutzt, um Indikatoren zur Versorgung der Bevölkerung im regionalen Vergleich zu bilden.

Die Ergebnisse des regionalen Gesundheitspersonalmonitorings werden auf der Ebene der 96 Raumordnungsregionen in Form von Tabellen und Karten dargestellt. Sie sind im Informationssystem der Gesundheitsberichterstattung des Bundes abrufbar. Zu jedem Bereich steht eine Basistabelle mit erhobenen Merkmalen untergliedert nach Geschlecht sowohl für die Anzahl des Personals als auch für Vollzeitäquivalente zur Verfügung. Zudem werden Indikatoren für den Regionalvergleich angeboten.

43 Struktur und Entwicklung der Pendlerbewegungen in der Lausitz

Sarah Kuhn, Holger Seibert, Antje Weyh

Institut für Arbeitsmarkt- und Berufsforschung, Deutschland
sarah.kuhn@iab.de

SECTION: Regionalstatistik (Session 2): Beschäftigung und Einkommen

Der demografische Wandel und damit verbundene Fachkräfteengpässe stellen für die Lausitz eine große Herausforderung dar. Verschiedene Initiativen sehen in der Rückholung von Auspendler*innen eine Möglichkeit diese Engpässe zu reduzieren.

Im Jahr 2024 pendelten etwas mehr als 102.000 Beschäftigte mit Wohnort in der Lausitz zu einer Arbeitsstätte außerhalb der Lausitz. Wir analysieren mittels Daten der Statistik der Bundesagentur für Arbeit die potenzielle Fachkräfteressource der Auspendler*innen genauer und zeigen, dass von dem insgesamt recht großen Potenzial schlussendlich nur ein geringer Teil als realistisches Arbeitskräftepotenzial für die Lausitz angesehen werden kann. Besonders im Fokus steht das Entgelt der Pendler*innen im Vergleich zu denjenigen, die nicht pendeln, denn aus ökonomischer Sicht bestimmt sich der Gesamtnutzen des Pendelns zu einem nicht unerheblichen Teil durch den monetären Faktor Entgelt. Im Mittel verdienen Beschäftigte mit Wohn- und Arbeitsort in der Lausitz 3.222 Euro, Auspendler*innen hingegen 4.011 Euro (Stand 2024). Pendeln lohnt sich finanziell noch einmal mehr für höher qualifizierte Beschäftigte. Bei Beschäftigten, die eine Helfer- bzw. Anlern Tätigkeit ausüben, macht es kaum einen Unterschied, ob sie in der Lausitz arbeiten oder auspendeln. Auspendelnde Spezialist*innen verdienen im Schnitt ca. 1.000 Euro mehr pro Monat, bei Expert*innen sind es gut 900 Euro.

Als realistisches Arbeitskräftepotenzial für die Lausitz können vor allem Personen angesehen werden, die weit pendeln und deren Entgeltdifferenz zu einem vergleichbaren Job in der Lausitz relativ gering sind, wie dies z.B. bei den Erziehungs- oder Logistikberufen der Fall ist. Aber auch für dieses Potenzial, welches schlussendlich nur rund ein Zehntel der Gesamtauspendler*innen umfasst, muss eine passende attraktive Arbeit innerhalb der Lausitz zur Verfügung stehen.

Bei der Verwendung der Beschäftigungsstatistik für Pendleranalysen sind einige Besonderheiten zu beachten. So ist beispielsweise nur erfasst, ob sich Wohn- und Arbeitsort unterscheiden. Es kann also keine Aussage darüber getroffen werden, ob und wie oft der Arbeitsort tatsächlich aufgesucht wird.

44 Geheimhaltung georeferenzierter Daten: ein Methodenvergleich des Quadtree-Verfahrens und der Kerndichteschätzung

Sevinç Öztürk, Marit Köffers, **Christoph Alfken**, Johannes Rohde

IT.NRW - Statistisches Landesamt, Deutschland
johannes.rohde@it.nrw.de

SECTION: Regionalstatistik (Session 2): Beschäftigung und Einkommen

Aus Wissenschaft, Politik und Gesellschaft wird zunehmend ein Bedarf an kleinräumigen Analysen auf Basis von georeferenzierten Daten an die Statistischen Ämter herangetragen. Insbesondere Auswertungen auf Ebene von Gitterzellen ermöglichen detaillierte Aussagen zur räumlichen Verteilung von gesellschaftlich oder politisch relevanten Sachverhalten und können dazu beitragen, die Lebenswirklichkeit der Bürgerinnen und Bürger adäquater abzubilden.

Der Auswahl eines geeigneten Geheimhaltungsverfahrens für Auswertungen auf Gitterzellenebene kommt dabei eine hohe Bedeutung zu: Nach Anwendung des Verfahrens sollen die Auswertungen weiterhin relevante Aussagen über räumliche Strukturen ermöglichen, ohne dass hohe Informationsverluste entstehen. Gleichzeitig dürfen kleinräumige Analysen keine Rückschlüsse auf Einzelangaben einzelner Merkmalsträger erlauben.

Im Rahmen der Statistischen Woche 2025 wurde die Erprobung von zwei Ansätzen für die Veröffentlichung von Ergebnissen aus der Lohn- und Einkommensteuerstatistik des Berichtsjahres 2021 auf Gitterzellenebene vorgestellt: das Statistische Bundesamt stellte die Anwendung der Quadtree-Methode vor, bei der Gitterzellen mit kritischen Informationen mit benachbarten Gitterzellen aggregiert werden bis keinerlei Aufdeckungsrisiken mehr vorliegen. IT.NRW präsentierte die Nutzung der Kerndichteschätzung, wobei die Geheimhaltung durch die Verwischung von Informationen auf benachbarte Gitterzellen sichergestellt wird. Beide Anwendungen erzielten vielversprechende Ergebnisse.

Diese Arbeit baut auf den genannten Ergebnissen auf: Die Quadtree-Methode und die Kerndichteschätzung werden hinsichtlich der Nutzung für Auswertungen mit georeferenzierten Daten auf Gitterzellenebene systematisch gegenübergestellt. Hierfür werden verschiedene Kennzahlen betrachtet, die den Geheimhaltungsschutz sowie den auftretenden Informationsverlust beider Methoden miteinander vergleichen. Die Ergebnisse tragen zur weiteren Orientierung für Anwenderinnen und Anwender bei der Auswahl eines Geheimhaltungsverfahrens für die künftige Erstellung kleinräumiger Analysen bei.

45 Inferring the spatio-temporal structure of climate variability from observations, simulators, and reconstructions of past climate

Nils Weitzel (1,2)

1: Research Center Trustworthy Data Science and Security of the University Alliance Ruhr, Germany; 2: TU Dortmund University, Germany
weitzel@statistik.tu-dortmund.de

SECTION: Statistics in Environmental and Engineering Sciences (Session 1)

The climate system is a high-dimensional, chaotic, and interacting system exhibiting variations on timescales from minutes to millions of years, driven by both changes in boundary conditions and internal dynamics. Projections of future climate, therefore, rely on process-based simulators

accurately reproducing these variations across temporal and spatial scales. However, assessing the consistency of simulated and observed climate variability patterns is challenging due to heterogeneous data sources, increasing data sparsity further back in time, and multiple sources of uncertainty.

Here, we first introduce two complementary frameworks for evaluating simulators against temperature reconstructions from natural archives such as sediment cores, which are the primary evidence for climate changes prior to the last two centuries. The first framework uses spectral analysis, while the second applies time series decompositions and the integrated quadratic distance for comparing probability distributions. For global mean temperature, we find consistency between simulators and reconstructions on timescales from one year to hundreds of thousands of years. In contrast, discrepancies in regional temperature variability suggest that the spatial covariance structure of state-of-the-art climate simulators is misspecified on multidecadal and longer timescales.

To further investigate the mechanisms responsible for these mismatches, we propose a Bayesian hierarchical framework for reconstructing spatio-temporal fields from sparse and noisy temperature proxies. This framework combines a data stage modeling local proxy-climate relationships with a process stages controlling the spatio-temporal climate evolution. We illustrate our framework with simulation studies and a network of temperature reconstructions from pollen assemblages. Future work aims at using statistical postprocessing methods to incorporate the insights from climate field reconstructions into refined projections of regional temperatures for the next centuries.

46 Detecting Stationary States in Non-Stationary Time Series

Florian Heinrichs

FH Aachen, Deutschland
f.heinrichs@fh-aachen.de

SECTION: Statistics in Environmental and Engineering Sciences (Session 1)

Detecting intervals where a drifting signal is effectively constant is essential in fields like process control, healthcare monitoring, and AI (e.g., early stopping in deep learning). Existing steady-state detectors assume i.i.d. errors, an unrealistic assumption once the data exhibits serial dependence or a time-varying distribution.

An additive model with a smooth, possibly non-constant mean and locally stationary errors is considered. The statistical task is to test whether a segment of prescribed length exists in which the mean is "almost" constant, depending on a user-defined tolerance.

The contributions are threefold. (1) A scan statistic for the detection of constant regions is introduced. (2) The asymptotic distribution of the scan statistic is established, which yields critical values and gives an asymptotically exact level- α -test. (3) Finite sample properties are studied through simulations and a real data application.

Unlike change-point methods, which flag any departure from stationarity, the proposed framework explicitly targets stable segments within non-stationary time series, offering a practical solution in many applications.

47 Detecting Changes in the Mean of Spatial Data with Applications to Deforestation and Vorticity Extremes

Sheila Görz, Merle Mendel, Roland Fried

Technische Universität Dortmund, Deutschland
sheila.goerz@tu-dortmund.de

SECTION: Statistics in Environmental and Engineering Sciences (Session 1)

Reliably identifying sudden changes in the structure of data is crucial for any application. Our focus is on detecting changes in the mean function of spatial data. In our work, we extend the approach from Schmidt (2024) of partitioning a time series into blocks to spatial data. The proposed test statistic replaces the originally used Gini's mean difference to compare the arithmetic block means by the ordinary variance. We examine the asymptotic behavior of such test statistics under the hypothesis of a constant mean and prove the asymptotic consistency of the resulting test for a large set of non-constant mean functions.

The methodology is illustrated using two application domains. First, we analyze satellite images of the Amazonian rainforest to identify regions exhibiting signals of deforestation. Second, we investigate changes in parameter estimates of the distribution of a fractional compound Poisson process fitted to the recurrence of mid-latitude vorticity extremes on the northern hemisphere. One interest lies in discovering whether parameter estimates based on more recent data differ from those based on older data. Another interest is to detect homogeneous parameter regions through spatial monitoring using a moving window approach.

48 Anforderungen an die Statistikausbildung: Antwort auf Berger et al. (2026)

Andreas Hildenbrand

IHK Rhein-Neckar, Deutschland
andreas.hildenbrand@rhein-neckar.ihk24.de

SECTION: Competence Development: Data Literacy and Statistics (Session 3)

Berger et al. (2026) formulieren vor dem Hintergrund von Digitalisierung und Künstlicher Intelligenz (KI) neun Thesen zur Weiterentwicklung der Statistikausbildung an Hochschulen. Sie betonen insbesondere die Bedeutung von Data und Statistical Literacy, den kritischen Umgang mit Datenherkunft und -qualität, die Reflexion methodischer Grenzen sowie eine transparente und adressatengerechte Kommunikation statistischer Ergebnisse.

Ich untersuche, inwieweit diese Anforderungen bereits in etablierten, auf die allgemeine Öffentlichkeit ausgerichteten Lehrbüchern der Statistik realisiert sind. Dazu werden zentrale Aspekte der Thesen inhaltsanalytisch mit zwei weit verbreiteten populärwissenschaftlichen Büchern verglichen: Huff (2023), erstmals 1954 erschienen, sowie Krämer (2015), erstmals 1991 erschienen. Beide Werke vermitteln statistische Grundprobleme anhand anschaulicher Beispiele und haben eine große Reichweite.

Es kann gezeigt werden, dass wesentliche Forderungen in den Thesen bereits weitgehend erfüllt sind. Insbesondere wird in beiden Werken ein ausgeprägtes Bewusstsein für Datenherkunft und -qualität sowie für die Grenzen statistischer Methoden vermittelt. Auch die Nutzung realitätsnaher Beispiele trägt zur Förderung kritischer Datenkompetenz bei. Defizite zeigen sich hingegen in der

expliziten Kommunikation statistischer Unsicherheit sowie im Fehlen von aktivierenden Lehr- und Lernmethoden.

Dafür gibt es bei Huff (2023) abschließende Fragen, die er Leserinnen und Lesern mitgibt, um statistische Darstellungen kritisch zu hinterfragen und gegenüber irreführender oder gar manipulativer Statistik sensibler zu werden. Dieser fragenbasierte Ansatz kann als eine Form eines „statistischen Coaching-Ansatzes“ interpretiert werden. Dieser erscheint besonders geeignet, um Reflexionskompetenz zu fördern. Als solcher könnte er weiterverfolgt werden.

Literatur

Berger, Ursula; Ertz, Florian; Hotz, Thomas; Huber, Sarah; Ickstadt, Katja; Lübke, Karsten; Münnich, Ralf; Schüller, Katharina; Skill, Thomas; Weihs, Claus; Weinert, Henrike (2026): Daten, Künstliche Intelligenz und Evidenz – neue Anforderungen an die Statistikausbildung an Hochschulen. AStA Wirtschafts- und Sozialstatistisches Archiv, <https://doi.org/10.1007/s11943-025-00363-7>.
Huff, Farrell (2023): How to Lie with Statistics. Dublin: Penguin Books.
Krämer, Walter (2015): So lügt man mit Statistik. Frankfurt: Campus-Verlag.

49 Daten, Data Literacy und Statistik I & II - Thesen zur grundständigen Statistikausbildung in der Wirtschaftswissenschaft

Martin Missong

Universität Bremen, Deutschland
missong@uni-bremen.de

SECTION: Competence Development: Data Literacy and Statistics (Session 3)

Die Ansprüche an eine kompetenzorientierte universitäre Ausbildung im Bereich der Quantitativen Methoden steigen kontinuierlich. Beredtes Beispiel hierfür ist das wachsende Spannbreite der zu vermittelnden „Literacies“. Gleichzeitig wird gefordert, über Gemeinsamkeiten und Unterschiede von Statistischer Methodenlehre auf der einen sowie Maschinellern und Künstlicher Intelligenz auf der anderen Seite zu informieren, eingebettet in weiter greifende Initiativen zur Erhöhung der Sichtbarkeit und des Ansehens in der Statistik in der Gesellschaft. Im Vortrag wird gezeigt, dass aus beiden Anforderungslinien grundlegende Konsequenzen für die Pflichtkurse in der Studieneingangsphase wirtschaftswissenschaftlicher Studiengänge („Statistik I & II“) resultieren. Diese beziehen sich auf Lehrinhalte, Lehrformen und Prüfungsaspekte und werden wo möglich anhand von Beispielen illustriert und zur Diskussion gestellt.

50 Wie können wir angemessen über gesellschaftlich sensible Statistiken reden?

Gerald Seidel

c/o Bundesagentur für Arbeit, Deutschland
gerald-seidel@web.de

SECTION: Competence Development: Data Literacy and Statistics (Session 3)

"Die Statistik sagt ..." ist eine häufige Formulierung in der veröffentlichten Kommunikation. Nicht selten kommen dabei unterschiedliche Autoren auf Basis derselben statistischen Quelle zu deutlich unterschiedlichen Ergebnissen - bei zum Teil starker persönlicher Betroffenheit von Angehörigen einzelner Personengruppen, über die ggf. eine gravierend wertende "statistische"

Aussage getroffen wird. Gleichzeitig ist ein Verzicht auf eine personengruppenspezifische Analyse i.d.R. nicht angezeigt, da eine entsprechende Unterscheidung grundlegend für die zu Grunde liegende Frage- bzw. Problemstellung ist oder zumindest sein könnte.

Vor dem Hintergrund einzelner Beispiele soll erörtert werden, wie die Ergebnisse personengruppenspezifischer statistischer Analysen gleichzeitig inhaltlich sachgerecht und ethisch vertretbar kommuniziert werden können. Insbesondere wird dabei auf die Interdependenz von Datenethik und Statistical Literacy eingegangen. Entsprechende didaktische Ansätze werden skizziert.

51 EDSEA: Eine europäische Allianz zur Definition von Data-Science-Kompetenzen

Katharina Schüller (1), Berthold Lausen (2), Eyke Hüllermeier (3), Jan Vahrenhold (4), Walter J. Radermacher (5), Katharina Weiß (6)

1: STAT-UP Statistical Consulting & Data Science GmbH; 2: University of Essex; 3: European Association for Data Science (EuADS); LMU München; 4: Informatics Europe (IE); Universität Münster; 5: Federation of European National Statistical Societies (FENStatS); LMU München; 6: Universität Bielefeld
katharina.schueller@stat-up.com

SECTION: Competence Development: Data Literacy and Statistics (Session 3)

In einer zunehmend daten- und KI-geprägten Gesellschaft fehlt bislang eine europaweit anerkannte, disziplinübergreifende Definition beruflicher Data-Science-Kompetenzen. Bestehende Kompetenzrahmen variieren erheblich – je nachdem, ob sie maschinelles Lernen, statistische Modellierung oder domänenspezifische Anwendungen in den Vordergrund stellen. Diese Heterogenität erschwert sowohl die Professionalisierung des Berufsfeldes als auch die Entwicklung vergleichbarer Hochschulcurricula und Akkreditierungsstandards.

Die European Data Science Education Alliance (EDSEA) setzt hier an. Im Rahmen eines Memorandum of Understanding schließen sich die European Association for Data Science (EuADS), die Federation of European National Statistical Societies (FENStatS) und Informatics Europe (IE) zusammen, um ein gemeinsames, grundlegendes Kompetenzprofil für Data-Science-Fachkräfte zu entwickeln. Das angestrebte Rahmenwerk strukturiert sich entlang entlang der Domänen der beitragenden Disziplinen sowie übergreifender Kompetenzbereiche wie bspw. Ethik und Kommunikation. Dabei knüpft EDSEA an bestehenden Initiativen an – darunter das Framework der Alliance for Data Science Professionals sowie Datenethik-Standards des International Statistical Institute (ISI) und des Institute of Electrical and Electronic Engineers (IEEE) – mit dem Ziel, Transparenz, Reproduzierbarkeit und gesellschaftliche Verantwortung als Grundprinzipien datenbasierter Arbeit zu verankern.

Mittelfristig verfolgt die Allianz das Ziel, das erarbeitete Kompetenzprofil als Grundlage für die Akkreditierung von Bachelor- und Masterstudiengängen nutzbar zu machen. Der Beitrag stellt die Ziele, Struktur und den aktuellen Entwicklungsstand von EDSEA vor und diskutiert, welche Rolle die Statistik-Community in diesem europäischen Definitionsprozess einnehmen kann. Dabei wird insbesondere die Frage adressiert, wie spezifisch statistische Kompetenzen – etwa im Bereich Inferenz, Modellkritik und Unsicherheitskommunikation – im Kern eines solchen Rahmens verankert werden sollten, um einer rein technikgetriebenen Verengung des Kompetenzprofils entgegenzuwirken.

52 Sharpening Identification in Large Structural VARs Using Narrative Restrictions

Lukas Berend (1), Jan Prüser (2)

1: University of Hagen, Germany; 2: TU Dortmund, Germany
lukas.berend@fernuni-hagen.de

SECTION: Young-Academics Mini-Symposium: Bayesian Econometrics (Session 2)

We propose a high-dimensional structural vector autoregression framework that features a factor structure in the error terms and accommodates a large number of linear inequality restrictions on impact impulse responses and structural shocks. In particular, we demonstrate that narrative restrictions can be imposed through constraints on the structural shocks, which can be used to sharpen inference and enhance the structural interpretation of the identified shocks. To estimate the model, we develop a highly efficient sampling algorithm that scales well with both the model dimension and the number of inequality restrictions on impact responses and structural shocks. We apply our approach to a large-scale structural VAR to identify ten structural shocks and analyze the dynamic responses of thirty-nine macroeconomic variables, thereby illustrating the model's flexibility and practical feasibility in complex empirical applications. Our empirical findings show that incorporating narrative restrictions narrows the identification set and enhances the interpretability of the estimated effects of structural shocks.

53 Beware of large shocks! A non-parametric structural inflation model

Catalina Martinez Hernandez

European Central Bank, Germany
Catalina.Martinez_Hernandez@ecb.europa.eu

SECTION: Young-Academics Mini-Symposium: Bayesian Econometrics (Session 2)

We propose a novel empirical structural inflation model that captures non-linear shock transmission using a Bayesian machine learning framework that combines VARs with non-linear structural factor models. Unlike traditional linear models, our approach allows for non-linear effects at all impulse response horizons. Identification is achieved via sign, zero, magnitude and historical restrictions within the factor model. Applying our method to euro area energy shocks, we find that inflation reacts disproportionately to large shocks, while small shocks trigger no significant response. These non-linearities are present along the pricing chain, more pronounced upstream for commodity and producer prices and gradually attenuating downstream for consumer prices. For policy makers, the finding that large shocks transmit differently implies that they may require a differentiated response.

54 Theory-based priors for the output gap

Matteo Luciani (3), **Michele Piffer** (1,2), Andrea Renzetti (1)

1: Bank of England; 2: Kings College London; 3: Federal Reserve Board
m.b.piffer@gmail.com

SECTION: Young-Academics Mini-Symposium: Bayesian Econometrics (Session 2)

Estimating potential output and the output gap requires identifying restrictions on the dynamics of trend and cycle. We develop a general econometric framework that specifies theory-based restrictions directly over the joint distribution of the latent trajectory, rather than through a particular state evolution equation. The approach nests standard unobserved-components models, while allowing external quantitative or qualitative information to be incorporated in a transparent and precision-weighted manner. Simulation evidence suggests that our prior can recover potentially asymmetric cyclical evolutions without having to estimate potentially challenging non-linear models. An empirical application to U.S. GDP shows that the Covid pandemic was not associated with a drop in the potential level of output, but the Great Financial Crisis was.

55 Do Low or High Wages Respond more Sensitive to Cyclical Fluctuations? – Empirical Evidence

Johannes Ludsteck

IAB, Deutschland
johannes.ludsteck@online.de

SECTION: Empirical Economics and Applied Econometrics (Session 4)

The sensitivity of wages responses to cyclical fluctuations may depend on the position of workers in the wage distribution. This is interesting from a practical point of view by telling which workers bear greater wage adjustment costs. It allows us, however, also to test the relevance of leading wage setting theories.

According to information asymmetry models, incentives and bonus payments are more important for high wage

workers since it is more difficult and costly to monitor the tasks performed by them. This in turn induces greater wage responses for high wages. The predictions union-firm bargaining theories go in the same direction.

Contrary predictions are given by hiring standards adjustment hypothesis. It claims that firms adjust hiring standards during booms in order to lessen the necessity for even larger wage increases. This triggers promotion chains which in turn increase labor demand at the bottom of the wage distribution and greater wage changes at the bottom.

In order to test these theories empirically, we develop an estimation approach which allows to control for composition bias by inclusion of fixed person effects and person-specific coefficients. Our estimator solves a fundamental problem of quantile regression: inclusion of fixed effects in quantile regression changes the interpretation of all other coefficients in a way that renders them worthless for our application.

The test is conducted using a large representative panel data set of West-German prime age full-time workers. We find confirmative evidence of the Reder hypothesis.

56 Heat or Eat: Distributional Subsidy Effects and the Joint Allocation of Essential Expenditures

Simone Maxand (1), Franziska Dorn (2)

1: European University Viadrina, Deutschland; 2: University of Duisburg-Essen
maxand@europa-uni.de

SECTION: Empirical Economics and Applied Econometrics (Session 4)

This paper develops a copula-based bivariate quantile regression framework to evaluate how multivariate policies reshape the joint distribution of essential expenditures. The method identifies conditional directional bivariate quantiles from the full joint distribution and embeds an event-study design within a multivariate setting. Unlike conventional mean-based or univariate quantile approaches, it captures non-linear trade-offs and policy-induced changes in the dependence structure between expenditure categories, enabling a multidimensional assessment of heterogeneous treatment effects across the entire distribution. Using data from the US Panel Study of Income Dynamics (2003–2019), we examine how food and energy subsidies affect the joint allocation of household energy and food spending and how they reshape the heat-or-eat trade-off. We find that subsidies not only shift marginal expenditure shares but also alter the interdependence between essential goods, with pronounced effects among vulnerable households and differences across state policy regimes. These results underscore the importance of multidimensional distributional methods for evaluating policies aimed at ensuring minimum living standards.

57 Regularized Wishart Stochastic Volatility for High-Dimensional Covariance Forecasting

Roman Liesenfeld (1), Guilherme Valle Moura (2), Julian Spies (1)

1: Universität zu Köln, Deutschland; 2: Universidade Federal de Santa Catarina, Brazil
julian.spies@wiso.uni-koeln.de

SECTION: Empirical Economics and Applied Econometrics (Session 4)

We extend Uhlig's Wishart stochastic volatility (WSV) model by introducing a regularized state transition for the precision matrix that shrinks the covariance forecasts toward a prior reference matrix. This regularization ensures stationarity of the return process and stabilizes the eigenvalues of the covariance forecasts, while preserving closed-form expressions for filtering, prediction, and likelihood evaluation. We provide conditions for stationarity and for the existence of second- and fourth-order moments, show that the model admits a multivariate GARCH representation, and derive bounds that illustrate how regularization prevents degeneracy and excessive dispersion in the eigenvalues of the covariance matrix forecasts. To account for regime shifts in the correlation structure, we further embed a time-varying directional forgetting scheme into the regularized model, allowing the forgetting rate to differ over time and across directions in the return space. In a high-dimensional application with up to 1,000 assets, the regularized models deliver significantly more accurate covariance forecasts than the baseline WSV and perform well compared to several benchmark models.

58 Monitoring Distributional Forecasts

Timo Dimitriadis (1), Christoph Hanck (2), Yannick Hoga (2), **Timo Puschmann** (2)

1: Goethe Universität Frankfurt, Deutschland; 2: Universität Duisburg-Essen, Deutschland
timo.puschmann@vwl.uni-due.de

SECTION: Empirical Economics and Applied Econometrics (Session 4)

Distributional forecasts play an increasingly important role across diverse fields, including meteorology, economics, and finance. In this paper, we propose “online” tests to monitor the adequacy of density forecasts. In contrast to one-shot tests, this allows forecast failure to be detected quickly, which is important as a timely review of the forecast methodology may be restored. Our monitoring procedure not only allows to evaluate univariate density forecasts, but also multivariate ones. Over some prespecified horizon, our tests hold size exactly even in finite samples. Hence, under the null of correctly specified distributional forecasts, our monitoring scheme only rejects with a given probability of, say, 5%. Monte Carlo simulations demonstrate that the power of our tests to detect misspecified forecasts is higher compared to the only other extant proposal. Finally, we apply our test to financial returns and macroeconomic series and find that our test can quickly detect forecast failure and provide insights into its causes.

59 Drift burst test statistic: new local tools for asset price model selection

Chiara Amorino (1), Matheus Lima-Cornejo (2), **Cecilia Mancini** (2)

1: UPF Barcellona, Spain; 2: University of Verona, Italy
cecilia.mancini@univr.it

SECTION: Statistics in Finance (Session 1)

We consider the problem of model selection for the continuous-time evolution of a financial asset price within the broad class of Itô semimartingales. These models rely on four basic components and are widely used in financial econometrics because they capture the expected trend, driven by the firm’s fundamentals, perturbed by standard Brownian fluctuations and possibly by abrupt large price changes (big jumps) as well as frequent small discontinuities (small jumps). Within this framework, asset risk is measured and derivative pricing, hedging strategies, and forecasts are developed. Correct identification of the components needed for an accurate description of each asset is therefore crucial.

Itô semimartingales are compatible with discrete observations, provided they are not excessively high-frequency, but this entails a loss of information; for example, jump times cannot be precisely identified without continuous observation.

Christensen, Oomen, and Renò (2022) introduced and studied the drift burst statistic to explain sudden market crashes over short time intervals by testing whether the triggering factor is an explosion of the trend component. The statistic is defined as the normalized ratio of a local estimator of the drift coefficient to a local estimator of the volatility coefficient.

Mancini (2023) further investigated the behavior of this statistic and showed that fundamentally different asymptotic behaviors arise depending on which components are present in the model. This leads to a battery of tests for the presence of specific components and features of the model, as well as local estimators of the coefficients. Each test requires careful implementation to be

effective. We begin by presenting a new local test for the presence of a jump at a given time, including its theoretical foundation, simulation performance, and application to financial data.

60 Pricing Barrier Options in Discontinuous Markets: An Adaptive Step Monte Carlo Approach

Lennart Dröge (1), Jos van Bommel (2)

1: Universität Augsburg, Deutschland; 2: Université du Luxembourg, Luxembourg
lennart.droege@uni-a.de

SECTION: Statistics in Finance (Session 1)

Continuous pricing models neglect significant overnight jumps. To price this "gap risk" efficiently, we introduce two novel simulation techniques, the Adaptive Step by Knock-out Probability and Adaptive Step by Overshooting Return, which dynamically adjust time steps based on the asset's proximity to the barrier. We validate these algorithms by pricing down-and-out calls under a hybrid diffusion-jump process. Results confirm our methods achieve high-frequency precision with drastically reduced computational load compared to fixed-step simulations. Empirically, we demonstrate that neglecting gap risk causes systematic undervaluation, with our model correctly capturing premiums near the barrier that can exceed theoretical benchmarks by a factor of two.

61 Explainable Outlier Detection via Exponential Vine Copula Tilting

Philipp Haid, Julian Wustl, Yarema Okhrin

University of Augsburg, Deutschland
yarema.okhrin@wiwi.uni-augsburg.de

SECTION: Statistics in Finance (Session 1)

Pointwise continuous multivariate outlier detection requires calibrated false-alarm control, sensitivity to relevant local departures, and interpretable evidence for a flag. We cast this problem as a likelihood-ratio comparison between a clean multivariate reference law and a class of interpretable local alternatives whose departures are intrinsically explainable relative to the reference. A vine copula reference model provides marginal and pair-copula likelihood coordinates, on which bounded positive exponential tilts define proper alternative laws. Positivity implements a "focus-on-bad-news" geometry: excess fitted surprise can support a local departure explanation, whereas unusually low surprise is not pursued as standalone evidence against the reference.

Profiling over the tilt class selects the marginal, dependence, or mixed alternative with strongest likelihood-ratio evidence. The resulting tilt identifies the fitted-vine-relative departure, and, for attribution, an intrinsic path decomposition allocates the same evidence across the coordinates. Split-conformal rank calibration converts the score into finite-sample valid compatibility p-values under exchangeability. Synthetic experiments illustrate calibrated detection, localization of departures in marginal behavior, dependence structure, or both, and component-level explanations.

62 Subgraph counts for hypergraphons with node covariates: asymptotic and bootstrap inference

Edwin Barthel, Carsten Jentsch

TU Dortmund Fakultät Statistik, Deutschland
barthel@statistik.tu-dortmund.de

SECTION: Nonparametric and Robust Statistics (Session 1)

As a natural extension of the graphon random graph model, we consider the hypergraphon model, which is an inhomogeneous model for exchangeable uniform hypergraphs. These are also called m -graphs and allow for edges that connect multiple (exactly m many) nodes. For $m=2$, the classical graphon model is recovered. In this setup, we equip the nodes with covariates and study weighted subhypergraph counts, where the weights are determined by the node covariates. Specifically, each found copy of a fixed subhypergraph is weighted by a function of its node covariates. While the node covariates and the (node-wise) latent positions of the hypergraphon are assumed to be jointly i.i.d., the dependence of the covariates and the latent positions may be arbitrary. For instance, our setup allows to count only those copies of a subhypergraph, whose node covariates satisfy certain properties. In this general setup, using the theory of generalized U-statistics, we give a full characterization of the joint limiting distributions of covariate-weighted subhypergraph counts. While the resulting limiting distribution is often normal, in general, it is a sum of random variables in different orders of the Wiener chaos. In particular, this extends known results on the asymptotic distribution of ordinary subgraph counts in the classical graphon random graph model (with $m=2$). Furthermore, assuming estimable hypergraphons, we show first-order consistency of a class of network bootstrap approaches for vectors of such counts in the case of a limiting normal distributions. In the special case of no node covariates, we extend these consistency results also to the non-normal limiting distribution case. To establish the latter, we propose a novel joint limiting theorem for generalized U-statistics on a triangular array type sample with sample-size dependent kernels, which may be of independent interest. The performance of these bootstrap methods is explored in simulations.

63 Robust Change-Point Test for Panel Data Based on U-Statistics

Martin Wendler (1), Claudia Kirch (2)

1: Otto-von-Guericke Universität, Deutschland; 2: Heinrich-Heine-Universität Düsseldorf
martin.wendler@ovgu.de

SECTION: Nonparametric and Robust Statistics (Session 1)

Many methods to detect change-points, i.e. abrupt shifts in time series are based on the CUSUM-statistic, which uses partial sums and thus is not robust against outliers. In this talk, we will discuss why 2-sample U-statistics are suitable for change-point detection and often more robust, and an overview over results on this topic will be given. Furthermore, the methods will be extended to functional data and to panel data. It is often challenging to obtain critical values, so we introduce a dependent wild bootstrap for U-statistics. With simulation results, we will demonstrate the better behavior of the robust procedures for more heavy-tailed random variables.

64 Testing sufficient follow-up in cure models via monotone tail density constraints

Eni Musta (1), Tsz Pang Yuen (2), Ingrid Van Keilegom (2)

1: University of Amsterdam, Netherlands; 2: KU Leuven, Belgium
e.musta@uva.nl

SECTION: Nonparametric and Robust Statistics (Session 1)

Cure models are used in survival analysis when a fraction of subjects will never experience the event of interest. A typical example arises in oncology: some patients, after successful treatment, may be cured of their cancer and thus never relapse, while others remain at risk. Because the study duration is limited some of the event times are right-censored and they may correspond either to cured patients or to still-susceptible patients who would relapse if followed longer. Hence, estimating cure rates reliably requires sufficiently long follow-up, but determining how long is "long enough" remains challenging. In practice, follow-up can be considered sufficient if the probability of an event occurring after the study ends is negligibly small. We develop a nonparametric test for this condition by exploiting monotonicity of the density function of the survival times in the tail region and then extend the approach to account for categorical covariates. A naive intersection-union test, requiring rejection of insufficient follow-up for every covariate value, is overly conservative in practice and lacks power. To improve upon this, we propose a new test procedure that relies on the test decision for one properly chosen covariate value. We show that both methods yield tests of asymptotically level α and investigate their finite sample performance through simulations.

65 A tailored statistical model for non-metallic inclusion sizes in steels and its inference

Anja Schmiedt (1), Stefan Bedbur (2), Udo Kamps (2)

1: OTH Regensburg, Deutschland; 2: RWTH Aachen, Deutschland
anja.schmiedt@oth-regensburg.de

SECTION: Statistics in Environmental and Engineering Sciences (Session 2)

In a dataset of non-metallic inclusion sizes in samples of engineering steel, standard order statistics are not suitable for modelling ascendingly ordered measurements within individual samples. This talk presents a flexible model of ordered random variables, which allows for changes in the distributions described by the model parameters. The joint maximum likelihood estimation of these parameters and the shape parameter of an underlying left-truncated Weibull distribution is considered. Furthermore, a model test is presented for the null hypothesis that common order statistics are an adequate model.

To reduce the number of model parameters, e.g., in situations involving small samples, a link-function approach is presented to minimize the number of parameters involved in the model. We discuss how a link function parameter can also serve as a material indicator.

Additionally, we propose an asymptotic test to check for the presence of a linear link function. Other topics include hypothesis tests concerning the parameters of the link function and the construction of simultaneous confidence regions for these parameters, as well as confidence bands for the entire graph of the link function. Finally, we discuss future research topics that may involve multi-sampling settings for detecting material differences.

The findings are illustrated by applying them to a real metallurgical dataset.

66 Cognitive Flexibility Without Stress: Automated Home-Cage Profiling of Learning Dynamics in Mouse Behaviour

Jekaterina Zukovska

Humboldt Universität zu Berlin, Deutschland
zukovskj@hu-berlin.de

SECTION: Statistics in Environmental and Engineering Sciences (Session 2)

Cognitive flexibility is a core executive function frequently disrupted in neurological and psychiatric disorders. However, few behavioural paradigms allow continuous, stress-free assessment of discrimination and reversal learning (DL/RL) in the home cage. Using the automated CognitionWall DL/RL task, we examined strain-dependent learning dynamics in C57BL/6JRj and DBA/2JRj mice without food deprivation or experimenter interaction.

Learning performance was quantified using four independent Kaplan-Meier measures: entries, sessions, hours, and days. Each of them capturing a distinct behavioural dimension. Across all metrics, reversal learning was consistently slower and more variable than discrimination learning, reflecting the increased cognitive load associated with inhibiting a previously learned rule, resolving conflict, and updating behavioural strategies. Despite this shared pattern, the measures differed in sensitivity: entries captured behavioural microstructure, days reflected overall duration, while sessions and hours most clearly differentiated strain-specific learning dynamics.

DBA/2JRj mice showed faster DL acquisition and more efficient RL adaptation, indicating enhanced cognitive flexibility. In contrast, C57BL/6JRj mice required more attempts and more real time to adjust to rule changes. Among all criteria, sessions and hours emerged as the most biologically informative combination, jointly describing how many learning attempts are required and how much real time animals invest in cognitive adaptation.

Both strains maintained stable circadian activity patterns, underscoring the welfare-friendly nature of the continuous home-cage setup. Together, these findings demonstrate that the CognitionWall DL/RL task, combined with multi-metric survival analysis, provides a sensitive and ethologically relevant platform for dissecting genetic influences on executive function and for preclinical modelling of cognitive flexibility deficits.

67 Uncertainty evaluation for smart measurements applied to scanning hyperspectral imaging using ML

Alex Strebel, Franko Schmähling, Gerd Wübbeler, Max Siebenkotten, Bernd Kästner, Manuel Marschall

Physikalisch Technische Bundesanstalt Berlin, Deutschland
alex.strebel@ptb.de

SECTION: Sources, Quantification, Accumulation, and Communication of Statistical Uncertainty (Session 1)

Modern material characterization at the nanoscale and high-resolution defect identification are key challenges in material science, advanced manufacturing, and quality control. Scanning hyperspectral imaging (sHSI) has emerged as a powerful tool for this purpose by providing spatially resolved

spectral information about a sample under test. While high-resolution spectroscopic metrology is required in, e.g. the semiconductor industry, sHSI measurements are often associated with long acquisition times and high measurement costs. Simply increasing the scanning speed decreases the signal-to-noise ratio. To address this limitation, compressed sensing (CS) and low-rank reconstruction methods have recently been proposed to accelerate data acquisition and improve spatial resolution in sHSI while significantly reducing the amount of observed data. These approaches typically yield point-estimates only and do not provide uncertainty evaluation, which is essential in metrology. In this work, we propose a Bayesian formulation of low-rank matrix reconstruction for sHSI that enables uncertainty evaluation while simultaneously improving reconstruction quality by introducing informative prior knowledge. Prior information is incorporated through a machine learning model, allowing the learned structure to be interpreted probabilistically rather than as a purely black-box predictor. A Gibbs-type sampling scheme is developed to perform posterior inference, where a Laplace approximation around a linearization point is introduced, enabling efficient and easy-to-implement sampling in the resulting high-dimensional model.

Experimental results on synthetic sHSI data demonstrate that the proposed Gibbs-type sampler provides accurate posterior mean reconstructions together with uncertainty evaluation for each spatial position and interferogram coordinate. Because of its computational efficiency, the approach is promising for metrological sHSI applications in which only limited measurements are available or samples evolve or degrade over time.

68 From Brier to 'h% correct': communicating forecast accuracy as hit rates

Niklas Valentin Lehmann

TU Bergakademie Freiberg, Deutschland
niklasl.2306@gmail.com

SECTION: Sources, Quantification, Accumulation, and Communication of Statistical Uncertainty (Session 1)

Communicating the accuracy of probabilistic forecasts and estimates to non-quantitative audiences remains a persistent obstacle in applied work. Strictly proper scoring rules are the principled tools for measuring error, but their numerical values rarely carry intuitive meaning: neither lay decision-makers nor domain experts typically possess a frame of reference for what constitutes a "low" or "high" score. We propose a general reparametrization. For any strictly proper scoring rule with finite expected loss, the score of a probabilistic forecast admits a certainty-equivalent hit rate h : the accuracy that a binary classifier would have to achieve to deliver the same expected utility to the population of decision-makers implicitly weighted by the scoring rule. The construction rests on the Schervish-type representation of proper scoring rules as mixtures of cost-loss decision problems. Instead of reporting the expected loss, we suggest reporting the certainty-equivalent hit rate — "this forecast is $h\%$ accurate". Most importantly, this framing retains propriety — h is maximized by truthful reporting — and preserves a defensible decision-theoretic meaning. For symmetric rules such as the Brier score, h is simply 1 minus the (normalized) error; for asymmetric rules, h also depends on the specification of a type-I/type-II error split. Non-finite rules such as the logarithmic score admit no such finite equivalent. The approach is legitimate but not harmless: because the certainty-equivalent hit rate depends on the chosen scoring rule, h can be inflated or deflated by the analyst's rule choice. We discuss the resulting scope for manipulation and argue that a transparent disclosure of the underlying scoring rule is necessary. Whether certainty-equivalent accuracy scores improve uncertainty communication in practice is an open empirical question.

69 Optimal Bayesian sequential design of two-arm clinical phase II trials in a nutshell

Riko Kelter

Institute of Medical Statistics and Computational Biology, University of Cologne, Germany
rkelter@uni-koeln.de

SECTION: Sources, Quantification, Accumulation, and Communication of Statistical Uncertainty (Session 1)

Two-arm phase II clinical trials with binary endpoints often benefit from an interim analysis that allows early stopping for futility, but Bayesian calibration of such designs is usually based on computationally intensive Monte Carlo simulation. In this talk, a simulation-free methodology is developed for Bayesian optimal two-stage designs in two-arm phase II trials with binary endpoints using Bayes factors as the primary measure of statistical evidence. Building on recent matrix-search methods for fixed-sample two-arm Bayes factor designs and earlier correction formulas for one-arm two-stage designs, the proposed approach derives exact prior-predictive expressions for the operating characteristics of a two-stage two-arm design with a single futility interim analysis. In particular, Bayesian power and Bayesian type-I error are obtained by correcting the corresponding fixed-sample quantities for trajectories that would have been removed by early stopping, yielding a fully numerical calibration procedure that avoids Monte Carlo error entirely. The resulting algorithm searches over admissible interim and final sample sizes to identify designs that satisfy prespecified constraints on Bayesian power, Bayesian type-I error, and the probability of compelling evidence in favour of the null hypothesis, while minimizing the expected sample size under the null hypothesis. The methodology is illustrated in realistic phase II settings, and examples show that the proposed design framework can accommodate interim futility monitoring while preserving interpretable operating characteristics. Aspects like how prior specification and evidence thresholds affect feasibility, efficiency and runtime, and how statistical uncertainty is mitigated by the novel methodology are highlighted. Overall, the approach extends simulation-free Bayes factor design methodology to the practically important setting of two-arm two-stage phase II trials and provides a transparent basis for Bayesian design calibration and sensitivity analysis, providing a powerful tool for the medical statistician's toolbox.

70 Bevölkerungsvorausberechnungen für Bayern bis 2044 - Methoden, Annahmen und kleinräumige Ergebnisse bis zur Gemeindeebene

Valerie Leukert

Bayerisches Landesamt für Statistik, Deutschland
DyanneValerie.Leukert@statistik.bayern.de

SECTION: Regionalstatistik/VDSt: Bevölkerungsvorausberechnung, Prognosen

Im Jahr 2044 werden nach der aktuellen Regionalisierten Bevölkerungsvorausberechnung des Bayerischen Landesamts für Statistik etwa 13,53 Millionen Menschen in Bayern leben – gegenüber dem Jahr 2024 ein Plus von circa 278 000 Personen (+2,1 %). Damit fällt das Bevölkerungswachstum gemäß der nun vorliegenden Ergebnisse etwas schwächer aus, als in der zuvor veröffentlichten

Berechnung (Bayern 2043 gegenüber 2023: +4,3 %; 2044 gegenüber 2024: +2,1 %). Die Unterschiede sind im Wesentlichen auf die Annahme eines gegenüber dem Ausland leicht reduzierten Wanderungssaldos und niedrigerer Geburtenzahlen zurückzuführen.

Der Vortrag gibt einen kurzen Überblick über die Methodik der Vorausberechnungen des Landesamts und die auf Basis der aktuellen demographischen Entwicklungen in Bayern getroffenen Annahmen. Daneben werden die zentralen Ergebnisse der jährlich aktualisierten Regionalisierten Bevölkerungsvorausberechnung bis 2044 für Bayern und die bayerischen Landkreise sowie kreisfreien Städte und des Demographie-Spiegels bis 2038 bzw. 2044 vorgestellt, mit welchem nun erstmals auch für die Gemeinden Bayerns Vorausberechnungsergebnisse unter Berücksichtigung der neuen Datenbasis des Zensus 2022 zur Verfügung stehen.

71 Mikrosimulation als Instrument kleinräumiger Bevölkerungsvorausschätzung

Moritz Zöphel

Humboldt-Universität zu Berlin, Deutschland
moritz.zoephel@student.hu-berlin.de

SECTION: Regionalstatistik/VDSt: Bevölkerungsvorausberechnung, Prognosen

Kommunale Planungsprozesse können von kleinräumigen Bevölkerungsprognosen profitieren, die bestehende Verfahren um eine individuenzentrierte Perspektive ergänzen. Der Vortrag stellt einen Mikrosimulationsansatz vor, der Fertilität, Mortalität, Zu- und Fortzüge sowie Binnenmigration auf Ebene einzelner Ortsteile bzw. Verkehrszellen modelliert. Anders als bei makroanalytischen Verfahren wird dabei jede Person mit ihren Merkmalen (Alter, Geschlecht, Wohnort und Migrationshintergrund) als eigenständige Simulationseinheit abgebildet. Demografische Ereignisse werden als stochastische Übergänge modelliert, deren Eintrittsraten von diesen Individualmerkmalen abhängen, wodurch sich heterogene Verhaltensweisen und Merkmalswechselwirkungen erfassen lassen. Am Beispiel der Stadt Leipzig wird gezeigt, wie sich durch die Kombination migrationspezifischer Zuzugsverteilungen mit alters und herkunftsspezifischer Binnenmobilität Segregationsszenarien entwickeln lassen, die räumliche Konzentrations- und Durchmischungsdynamiken abbilden. Abschließend diskutiert der Vortrag methodische Herausforderungen, insbesondere Skalierbarkeit bei feiner räumlicher Gliederung, Ratenschätzung und Herausforderungen bei der technischen Umsetzung von Mikrosimulationen.

72 Analyse und Prognose regionaler Strukturen und Trends der Fertilität in Deutschland

Elisabeth Zessin, Steffen Maretzke

Bundesinstitut für Bau-, Stadt- und Raumforschung im BBR (BBSR)
steffen.maretzke@bbr.bund.de

SECTION: Regionalstatistik/VDSt: Bevölkerungsvorausberechnung, Prognosen

Regionale Unterschiede im Fertilitätsniveau sind aktuell ein wesentlicher Treiber regionaler Unterschiede der Bevölkerungsentwicklung in Deutschland. Der Beitrag informiert über die identifizierten regionalen Muster und Trends der Fertilitätsentwicklung in der Vergangenheit und zeigt, wie diese Informationen für die Entwicklung regional differenzierter Fertilitätsannahmen für die nächste

BBSR-Bevölkerungsprognose (2025 bis 2050) genutzt werden. Ein inhaltlicher Schwerpunkt dieses Beitrages liegt darin, die differenzierte methodische Vorgehensweise zur Entwicklung der Annahmen zum Wandel der regionalen Fertilitätsstrukturen und -muster bis 2050 vorzustellen. Weil diese Prognose auf der Ebene von Stadt- und Landkreisen gerechnet wird, erfolgten auch die hier präsentierten Fertilitätsanalysen auf dieser regionalen Ebene.

73 Automatic Debiased Machine Learning and Sensitivity Analysis for Dynamic Treatment Effects

Martin Huber (1), Jannis Kueck (2), **Theresa M. A. Schmitz** (3)

1: Department of Economics, University of Fribourg, Switzerland; 2: Düsseldorf Institute for Competition Economics, Heinrich Heine University Düsseldorf, Germany; 3: Chair of Statistics and Econometrics, Heinrich Heine University Düsseldorf, Germany
theresa.schmitz@hhu.de

SECTION: Causal Machine Learning

We develop an Automatic Debiased Machine Learning (AutoDML) estimator for causal effects in dynamic treatment models, and propose a sequential sensitivity analysis to assess robustness to violations of the identifying assumptions. Our estimation approach relies on efficient influence functions and Riesz representers for dynamic average treatment effects. However, rather than explicitly estimating treatment propensity scores or conditional densities, we directly learn the relevant debiasing weights through a dynamic Riesz representation. Our estimation approach yields Neyman-orthogonal, doubly robust estimators that are root-n consistent under specific regularity conditions. We extend the nonparametric sensitivity analysis framework for static treatments to dynamic settings. This yields interpretable bounds and enables us to sequentially decompose the omitted variable bias into time-dependent scaling factors, outcome confounding, and selection confounding strength. We assess the finite sample properties of our proposed estimator and the validity of our bias bounds through a simulation study. In overlap as well as weak overlap settings, we compare the performance of our estimation approach to dynamic double machine learning and document differences in sensitivity to tuning parameters, bias, and stability. Overall, our approach effectively links identification, estimation, and sensitivity analysis within a unified sequential framework, yielding robust and computationally efficient results for dynamic treatment analysis.

74 Sensitivity Analysis for (Local) Average Treatment Effects in Instrumental Variable Models

Juejue Wang (1), **Julian Diefenbacher** (2,3), Jannis Kueck (4), Martin Spindler (2), Carlos Cinelli (1), Victor Chernozhukov (5)

1: Department of Statistics, University of Washington, USA; 2: Faculty of Business Administration, Institute of Statistics, University of Hamburg, Germany; 3: GESIS - Leibniz Institute for the Social Sciences, Germany; 4: Düsseldorf Institute for Competition Economics, Heinrich Heine University Düsseldorf, Germany; 5: Department of Economics, MIT, USA
julian.diefenbacher@gesis.org

SECTION: Causal Machine Learning

We develop a nonparametric framework for sensitivity analysis of instrumental variable estimates that handles both violations of the exogeneity assumption—via unobserved confounders of the instrument—and violations of the exclusion restriction assumption—via unobserved side-effects of the instrument. Specifically, we generalize the standard local average treatment effect (LATE) identification result, showing that, in the presence of post-instrument mediators, the IV estimand still identifies a weighted average of controlled LATEs—a result that is useful for sensitivity analysis and is also of independent interest. Based on these characterizations, we show how to obtain closed-form bounds on the target parameter, both by bounding the reduced form and first stage separately and by a pseudo-outcome construction that recasts the estimand as the root of a single linear moment. These bounds are expressed in terms of familiar and interpretable partial R^2 values, and we also provide novel characterizations of these parameters in terms of also familiar covariate imbalance measures. We provide flexible estimation procedures based on debiased machine learning, along with Anderson–Rubin–type confidence sets that remain valid under weak identification. Simulation evidence illustrates the finite-sample performance of our methods. We illustrate our method by estimating the effect of 401(k) participation on net financial assets, using 401(k) eligibility as an instrument.

75 Understanding the finite-sample relationship between Double Machine Learning and Targeted Maximum Likelihood Estimation

Alejandro Schuler (1), David Bruns-Smith (2), **Eva-Maria Oess** (3), Avi Feller (1)

1: UC Berkeley; 2: Stanford University; 3: University of Cologne
oess@wiso.uni-koeln.de

SECTION: Causal Machine Learning

Double machine learning (DML; a generalization of augmented inverse propensity score weighting) and Targeted Maximum Likelihood Estimation (TMLE) are two influence function-based estimators popular in causal inference. While asymptotically equivalent, the two can differ substantially in finite samples, prompting extensive debate over their relative merits. In this paper, we study DML and a natural TMLE implementation (TMLE with a “linear update”) in a large class of causal estimands; in this case, the two estimators differ by a single scaling factor applied to the re-weighted residuals. Our main contribution is to analyze the properties of this factor and demonstrate that it directly controls a bias-variance trade-off. When overlap is good or when the Riesz representer is estimated via Riesz regression, the scaling factor is close to one and the two estimators are nearly

identical. When overlap is poor and the Riesz representer is estimated via traditional propensity score methods, however, the scaling factor can be far from one and the estimators diverge sharply; in this regime, TMLE de-biases the naive plug-in estimate much more aggressively than DML, at the cost of greater variance. We then propose a family of alternative estimators that trace out different points on this bias-variance trade-off, nesting DML and TMLE as special cases. Finally, we separate the choice of bias-variance trade-off from the substitution-estimator property often cited as a key benefit of TMLE. We show that DML is always a substitution estimator under a linear submodel and propose a natural modification that preserves this property under nonlinear submodels as well.

76 Inference on Multiple Treatment Effects with an Application to Health Economics

Cathy Yi-Hsuan Chen (1), Victor Chernozhukov (2), **Jan Rabenseifner** (3), Martin Spindler (3)

1: University of Glasgow, UK; 2: MIT, US; 3: University of Hamburg, Germany
jan.rabenseifner@uni-hamburg.de

SECTION: Causal Machine Learning

With the rise of digitization more complex data sets are available for empirical research. In causal inference the focus has been traditionally on estimation of the causal effect of a selected treatment variable or very few. But in many relevant applications, there are very many treatment variables and their interactions can be very complex. An example are the causes for high blood pressure. In this paper we develop methods for estimation and inference of the marginal treatment effect of variables in such complex settings with modern machine learning methods. Based on the concept of Shapley values, we provide debiased estimators for the marginal contribution and derive their theoretical properties which allows for valid inference on the marginal effects. Because of the complex nature of the underlying estimation problem, also a simulation based debiased estimator is proposed and analyzed. In a simulation study the small sample properties for the estimators are compared. Despite the challenging setting, our estimator overall gives good results. In an application we estimate the marginal effects of different lifestyle risks for high blood pressure using data from The National Health and Nutrition Examination Survey (NHANES).

77 Testing for Explosiveness in Financial Asset Prices Using High-Frequency Volatility

Peter Boswijk (1), Jun Yu (2), Yang Zu (2)

1: University of Amsterdam, Niederlande; 2: University of Macau, Macau
h.p.boswijk@uva.nl

SECTION: Empirical Economics and Applied Econometrics (Session 5)

Based on a continuous-time stochastic volatility model with linear drift and jumps, we develop a test for explosive behavior in financial asset prices at a low frequency when prices are sampled at a higher frequency. The test exploits the volatility information in the high-frequency data, while maintaining robustness to jumps using truncation. The method involves first removing jumps and then devolatilizing low-frequency logarithmic asset price increments using truncated realized volatility measures. Subsequently, a supremum-type recursive Dickey-Fuller test is performed on the

reconstructed sample. The proposed test has a nuisance-parameter-free asymptotic null distribution and is easy to implement. We study the size and power properties of the test in Monte Carlo simulations. A real-time date-stamping strategy based on the reconstructed sample is proposed for the origination date of the explosive regime. Conditions under which the real-time date-stamping strategy is consistent are established. The test and the date-stamping strategy are used to study explosive behavior in recent major cryptocurrency markets and in the S&P 500 index during the 1990s dotcom bubble period.

78 Long-Run Prediction Uncertainty with Persistent Heteroskedasticity

Daniel Dzikowski (1), Kurt Lunsford (2), Carsten Jentsch (1)

1: TU Dortmund, Deutschland; 2: Federal Reserve Bank of Cleveland, USA
daniel.dzikowski@tu-dortmund.de

SECTION: Empirical Economics and Applied Econometrics (Session 5)

Precise statements about the long-run uncertainty of economic and financial variables are essential contributions for policy, planning, and portfolio decisions made today. The US Social Security Administration, for example, makes projections for several macroeconomic variables out to 75 years, producing intermediate, high cost, and low cost scenarios. Müller and Watson (2016) offer a framework for measuring the long-term uncertainty of variables caused by persistent patterns in their means. However, their framework assumes non-persistent structures in the second moment of the variables. This assumption is questionable because heteroskedasticity is often observed to be very persistent in economic and financial variables. We extend their framework by incorporating persistent structures in the second moments and show that these structures must be taken into account to make precise statements for the long-run uncertainty of economic and financial variables.

79 Statistical Inference for Score Decompositions

Timo Dimitriadis (1), Marius Puke (2)

1: Goethe Universität Frankfurt, Deutschland; 2: Universität Hohenheim, Deutschland
Dimitriadis@econ.uni-frankfurt.de

SECTION: Statistics in Finance (Session 2)

We introduce inference methods for score decompositions, which partition scoring functions for predictive assessment into three interpretable components: miscalibration, discrimination, and uncertainty. Our estimation and inference relies on a linear recalibration of the forecasts, which is applicable to general multi-step ahead point forecasts such as means and quantiles due to its validity for both smooth and non-smooth scoring functions. This approach ensures desirable finite-sample properties, enables asymptotic inference, and establishes a direct connection to the classical Mincer-Zarnowitz regression. The resulting inference framework facilitates tests for equal forecast calibration or discrimination, which yield three key advantages. They enhance the information content of predictive ability tests by decomposing scores, deliver higher statistical power in certain scenarios, and formally connect scoring-function-based evaluation to traditional calibration tests, such as financial backtests. Applications demonstrate the method's utility. We find that for survey

inflation forecasts, discrimination abilities can differ significantly even when overall predictive ability does not. In an application to financial risk models, our tests provide deeper insights into the calibration and information content of volatility and Value-at-Risk forecasts. By disentangling forecast accuracy from backtest performance, the method exposes critical shortcomings in current banking regulation.

80 Signal-Selected Subset Portfolios for Long-Run Growth

Sven Lehmann (1), Rainer Alexander Schüssler (2)

1: University of Hagen, Germany; 2: Aalborg University Business School, Denmark
raineralexanderschuessler@gmail.com

SECTION: Statistics in Finance (Session 2)

Long-horizon portfolio choice is often formulated as a problem of estimating optimal continuous weights. In high-dimensional asset universes, however, this architecture is fragile: small errors in expected returns or covariances can translate into large weight distortions and poor compounded performance. This paper proposes an alternative structure-first approach for long-run growth investing: Signal-Selected Subset Portfolios (S³P).

The framework combines cardinality-constrained subset selection with equal weighting within the selected subset. Conditional information is used only to decide which assets enter the portfolio, while capital is allocated equally across the chosen assets. This separates two margins that are usually conflated: support selection and within-support weighting. The central argument is that, under realistic estimation noise, the robust first-order use of predictive information is to select the right assets, not to fine-tune continuous weights. Equal weighting eliminates sensitivity to magnitude errors, while cardinality constraints convert noisy forecasts into a discrete, stable portfolio decision.

The paper develops a value decomposition that separates the economic contribution of support selection from the incremental value of continuous weighting on a fixed support. It also introduces a controlled-information experiment in which the return environment is held fixed while the strength and persistence of predictive signals are varied. The results show that even weak but directionally informative signals can generate long-run value when embedded in the S³P architecture, whereas sparse continuous-weight portfolios are more vulnerable to estimation noise.

Empirically, the framework is illustrated in high-dimensional U.S. equity settings with monthly rebalancing and realistic subset sizes. The applications use predictive inputs of varying complexity, ranging from simple return-based signals to richer information sets and machine-learning forecasts. Across these realistic settings, the results show that conditional information creates robust economic value when used for subset selection, but not when translated into continuous weight magnitudes.

81 Conditional Method Confidence Set

Lukas Bauer (1), Ekaterina Kazak (2)

1: University of Freiburg, Deutschland; 2: University of Birmingham
lukas.bauer@vwl.uni-freiburg.de

SECTION: Statistics in Finance (Session 2)

This paper proposes a Conditional Method Confidence Set (CMCS) which allows to select the best subset of forecasting methods with equal predictive ability conditional on a specific economic regime. The test resembles the Model Confidence Set and is adapted for conditional forecast evaluation. We show the asymptotic validity of the proposed test and illustrate its properties in a simulation study. The proposed testing procedure is particularly suitable for stress-testing of financial risk models required by the regulators. We showcase the empirical relevance of the CMCS using the stress-testing scenario of Expected Shortfall. The empirical evidence suggests that the proposed CMCS procedure can be used as a robust tool for forecast evaluation of market risk models for different economic regimes.

82 Conditional Volatility Modeling and Financial Risk Measurement via Neural Networks

Theo Berger

Hochschule Hannover, Deutschland
theo.berger@hs-hannover.de

SECTION: Statistics in Finance (Session 2)

This paper presents a comprehensive residual-analysis of recurrent neural networks (RNNs) in the context of conditional volatility and financial risk modeling. Using both simulated and empirical economic time series, we evaluate the ability of RNNs to capture stylized facts of financial volatility, such as volatility clustering and serial dependence, under varying sample sizes. We conduct a rigorous assessment of the statistical properties of model residuals to compare the performance of RNNs against traditional econometric benchmarks (e.g., GARCH models).

Our results show that RNNs effectively capture key features of financial time series, particularly in modeling volatility clustering and dynamic dependence structures, even in small-sample settings. However, we identify a persistent limitation: RNNs struggle to adequately represent asymmetric volatility effects. A central finding is a trade-off between high in-sample predictive accuracy and the ability to generate residuals that are statistically well-behaved and free from autocorrelation and heteroscedasticity. This suggests that while RNNs excel in fitting complex patterns, their out-of-sample risk forecasting performance may be compromised if residual diagnostics are neglected. The study emphasizes the importance of residual analysis as a critical step in validating machine learning models in financial econometrics, and provides practical guidance for the application of RNNs in volatility and risk modeling.

83 Nonparametric quantile regression in the fully functional model

Tim Greger

Universität Bayreuth, Deutschland
tim.greger@uni-bayreuth.de

SECTION: Nonparametric and Robust Statistics (Session 2)

Compared to mean regression, quantile regression has several advantages. Concerning quantile regression, we look at a double functional model, i.e. a conditional model where both variables are functional. In functional settings the definition of a quantile itself is not straightforward. This problem arises because ordering is not possible for multivariate or even functional elements, which yields that the Lebesgue measure does not exist for infinite dimensions. In this talk, we discuss the projection of functional random variables onto one of their directions, when the direction is either fixed or random. Asymptotic properties are given for a quantile estimator, which is based on a smoothed kernel estimator of the distribution function. The performance of the estimator in finite samples is in addition checked in a simulation study.

84 Distribution-free prediction intervals from interval score optimised pairs of nonparametric regression quantiles

Harry Haupt, Joachim Schnurbus, Ida Bauer

Universität Passau, Deutschland
joachim.schnurbus@uni-passau.de

SECTION: Nonparametric and Robust Statistics (Session 2)

We propose a distribution-free method for constructing prediction intervals for trending seasonal time series over time horizons of up to one seasonal cycle. In contrast to related approaches, the proposed method does not require iteration between analysing seasonal and trend components. Prediction intervals are determined by a pair of nonparametric regression quantiles using a local linear approach. Bandwidth selection balances the trade-off between coverage and length by minimising the average training-sample interval score for a reordered set of quantiles. The theoretical part explains the extension of nonparametric trend regression to seasonal processes and derives a Bahadur representation under assumptions for locally stationary processes. Besides Monte Carlo simulations, the proposed method is applied to short-term predictions of monthly and quarterly industrial production for OECD countries.

85 Nonparametric covariance estimation with bias correction and bootstrapping for spatial lattice processes

Antonio Fioravanti, Carsten Jentsch

TU Dortmund, Germany
antonio.fioravanti@tu-dortmund.de

SECTION: Nonparametric and Robust Statistics (Session 2)

The accurate estimation of covariances from spatial data is crucial in many applications. While parametric approaches have been widely used for modeling spatial covariances, the parametric assumption may not be correctly specified, which may lead to wrong conclusions.

On the contrary, non-parametric approaches are often neglected in practice, because their finite sample estimation accuracy suffers from the large number of covariance parameters to be estimated and they are often prone to severe bias issues.

In this paper, we develop a novel framework for bias correction for non-parametric sample covariance estimators of stationary lattice processes. We derive the exact finite sample biases of different versions of sample covariance estimators for spatial data. Our results show that the estimators' expectations can be expressed as linear combinations of the population covariances that only depend on the spatial lag of the sample covariance and on the sample size of the observed lattice data. Exploiting this characterization enables the construction of (jointly) bias-corrected covariance estimators that are nearly unbiased, depending on the strength of spatial dependence in the lattice process. Additionally, we derive exact formulas for the (joint) mean-squared error (MSE) and provide asymptotic normality results. In particular, we show that the bias-corrected estimators are asymptotically equivalent to their uncorrected versions. As an application, we use the bias-corrected covariance estimates for bootstrap inference and propose a novel spatial bootstrap procedure. Simulation studies for (Gaussian) lattice processes demonstrate that our proposed bias correction can achieve substantial bias reduction, while often reducing also the MSE, particularly when the underlying process exhibits strong spatial persistence. In particular in these cases, the proposed bootstrap procedure benefits from the bias-correction and regularization of the covariance estimation in finite samples.

86 Testing independence of errors and covariates in functional linear models

Melanie Birke (1), Natalie Neumeyer (2)

1: Universität Bayreuth, Deutschland; 2: Universität Hamburg, Deutschland
melanie.birke@uni-bayreuth.de

SECTION: Nonparametric and Robust Statistics (Session 2)

In functional linear models the assumption of independence of errors and covariates is often essential for data analysis procedures. Neglecting certain dependencies will result in inconsistent parameter estimates and, as a consequence, misleading conclusions in applications. In the case of functional covariates and scalar responses we suggest two different hypothesis tests based on empirical residual characteristic functionals. In a first one the squared distance between the characteristic functionals in general and under independence is integrated with respect to some measure on the Borel-sigma field on a Hilbert space. The second approach uses a truncated basis

representation of the Hilbert space valued observations. Independence can then be expressed via the characteristic functions of residuals and projections onto the basis functions which results in a sum of double integrals of the squared difference of the characteristic functions over the real numbers. For both approaches we show consistency of the resulting tests as well as the asymptotic distributions under the null hypothesis. As well known, bootstrap procedures work better in such settings. Therefore we base our tests on a residual bootstrap. In a simulation study we investigate the finite sample performance of both tests and discuss the limitations and advantages.

87 Examining the Role of Measurement and Representation in Uncertainty, Bias, and Unfairness Through Total Error Frameworks

Patrick Schenk (1,2), Christoph Kern (1,2), Trent D. Buskirk (3)

1: Department of Statistics, LMU München; 2: MCML, Munich Center for Machine Learning; 3: School of Data Science, Old Dominion University
p.o.s.on.stats@gmail.com

SECTION: Sources, Quantification, Accumulation, and Communication of Statistical Uncertainty (Session 2)

The quality of the underlying data limits how good any descriptive analysis, statistical inference, or prediction model can be. At the same time, data may often be treated as given – and their uncertainties and biases as inevitable.

Here, we explore two total error frameworks originating from the social sciences, Total Survey Error and its generalization, Total Data Quality, which map the entire data pipeline – from a researcher’s thoughts about their analysis goals to the pre-processed data that are eventually fed into analysis. The frameworks group sources of uncertainty and bias into errors of measurement (i.e., actual values deviating from true values) and errors of representation (i.e., the statistical distribution in the data deviating from that in the target population). At each step in the data pipeline, the frameworks provide a scaffolding to consider data quality – relative to either the previous step or a general reference point.

Our overarching main objective here is “group fairness”: that is, for all groups in a population, analyses and models should work well enough and about equally well. Differences in data quality between groups are one of the main underlying obstacles to group fairness. We utilize the two frameworks to illustrate group differences in errors of measurement and errors of representation at each step of the data pipeline, drawing on examples from empirical research spanning traditional and more modern data types.

By locating the sources of unfairness, uncertainty, and bias within the data pipeline, the frameworks show how data can be better designed and how to assess “gathered data” (i.e., data created by someone else). Thus, they can help data scientists avoid preparing unfair data and communicate data-based sources of uncertainty and biases to both external data users and colleagues within the same statistical organization.

88 Stabilität von Imputationsverfahren

Florian Dumpert (1), Inken Veips (2), Maria Thurow (2,3), Markus Pauly (2,3)

1: Statistisches Bundesamt; 2: TU Dortmund; 3: Research Center Trustworthy Data Science and Security
florian.dumpert@destatis.de

SECTION: Sources, Quantification, Accumulation, and Communication of Statistical Uncertainty (Session 2)

Bei der Datenaufbereitung ist der Umgang mit fehlenden Werten ein wichtiger Aspekt, der spätere Forschungsergebnisse und amtliche Zahlen beeinflusst. Insbesondere in der amtlichen Statistik sollten standardmäßig die Eigenschaften von Imputationsmethoden auf Datensätze untersucht werden, bevor erstere produktiv eingesetzt werden. Der Vortrag stellt zu diesem Zweck eine prototypische Simulationsstudie zu Imputationsmethoden vor. Anhand der deutschen Verdienststrukturerhebung aus den Jahren 2010 bis 2018 zeigen wir, wie sich Imputationsmethoden anhand ihrer (multivariaten) Imputationsgenauigkeit bewerten lassen. Neben üblichen Kennzahlen wird zusätzlich auch die Datenplausibilität betrachtet, indem vordefinierte Plausibilitätsregeln in die Studie einbezogen werden. Auch wird prototypisch untersucht, inwiefern ausgewählte Subgruppen gleichermaßen gut imputiert werden (Stichwort deskriptive Fairness als Qualitätskriterium).

89 Beyond Standard Errors: How Data Production Processes Shape Uncertainty and Inequality

Pirmin Fessler

Oesterreichische Nationalbank, Österreich
pirmin.fessler@oenb.at

SECTION: Sources, Quantification, Accumulation, and Communication of Statistical Uncertainty (Session 2)

Official statistics typically quantify uncertainty in terms of sampling variability, confidence intervals, or standard errors. In the production of microdata-based indicators, however, uncertainty is not only quantified and reported, but also modified through a sequence of data processing steps and practical decisions.

This presentation examines how procedures such as editing, weighting, trimming, and imputation shape both the level and the distribution of measured wealth and inequality. Using the Household Finance and Consumption Survey (HFCS) as a central example, it highlights the trade-off between variance and bias, particularly for adjustments affecting the upper tail of the distribution. Measures such as weight trimming or constrained imputations tend to dampen extreme values, often reducing measured inequality while increasing statistical precision.

At the same time, many influential decisions in statistical data production are not fully captured within formal uncertainty frameworks. Choices related to data cleaning, model specification, or the treatment of extreme values rely on practical judgement and institutional conventions, yet can have substantial effects on resulting indicators.

These limitations become more pronounced when survey data are integrated into frameworks such as Distributional Wealth Accounts (DWA), where uncertainty is further transformed and becomes difficult to quantify.

The presentation concludes that communicating uncertainty requires going beyond standard error-based measures. Greater emphasis should be placed on transparently documenting data production processes and communicating their implications for the interpretation of inequality statistics.

90 Assessing Specification Uncertainty in Survey Weighting: A Multiverse Monte Carlo Approach

Anne Konrad

Leibniz-Institut für Bildungsverläufe, Deutschland
anne.konrad@lifbi.de

SECTION: Sources, Quantification, Accumulation, and Communication of Statistical Uncertainty (Session 2)

Assessing uncertainty is a central task in survey statistics and is typically framed in terms of sampling variability. However, an additional and often underestimated source of uncertainty arises from modeling and specification choices made during data processing and analysis. In particular, the construction of survey weights involves different decisions, such as the choice of auxiliary variables, model specification for nonresponse adjustment, and calibration constraints, that may substantially affect resulting estimates.

This paper proposes a multiverse Monte Carlo framework to systematically quantify this form of specification uncertainty in the context of survey weighting. Instead of evaluating methods under a small number of fixed data-generating processes (DGPs), we consider a structured space of plausible DGPs and modeling choices. For each Monte Carlo iteration, multiple weighting approaches are applied across alternative specifications, resulting in a distribution of estimates that reflect both sampling variability and variability induced by modeling decisions.

Our approach allows us to decompose total uncertainty into within-specification (classical Monte Carlo variance) and between-specification components. This provides new insights into the robustness of weighting methods and highlights scenarios in which conclusions about method performance depend strongly on researcher degrees of freedom.

The proposed framework contributes to a more comprehensive understanding of uncertainty in survey statistics by extending the focus beyond sampling error. Moreover, it shows how sensitivity analyses over plausible specifications can be integrated into routine Monte Carlo simulations.

91 Modelle, Milieus, Karten - Ein integrierter Ansatz für kleinräumige Wahlanalysen

Lutz Wallhorn

Red Data Analytics GmbH, Deutschland
lutz@red-data.eu

SECTION: VDSSt (Session 1): Wahlanalyse und statistische Helfer für die Wahlorganisation

Der Vortrag stellt einen Ansatz vor, mit dem sich bundesweit verfügbare Wahlbezirksergebnisse der Bundestagswahl 2025 mit kommunalen Geometrien, Zensusdaten und weiteren kleinräumigen Daten zu einer gemeinsamen Analysegrundlage verbinden lassen. Im Zentrum stehen drei gleichwertige Bausteine.

Erstens wird ein integriertes Wahlmodell vorgestellt, das Wahl-, Raum- und Sozialstrukturdaten zusammenführt, um regionale Muster des Wahlverhaltens und Veränderungen zwischen Wahlen sichtbar zu machen. Der Ansatz basiert auf einem hierarchischen bayesianischen Modell, das verschiedene räumliche Ebenen gemeinsam berücksichtigt.

Zweitens zeigt der Vortrag, wie sich Wahlbezirke mithilfe eines rekursiven Clustering-Ansatzes mit Gaussian Mixture Models zu datengetriebenen politischen Milieus verdichten lassen. So entstehen interpretierbare Raumtypen, die kleinräumige Muster von Wahlverhalten, Beteiligung und regionaler Prägung sichtbar machen.

Drittens wird gezeigt, wie sich Modell- und Clusterergebnisse in einer interaktiven Kartenanwendung räumlich darstellen, vergleichen und explorieren lassen. Dadurch werden komplexe kleinräumige Zusammenhänge visuell zugänglich und für Analyse, Interpretation und kommunale Praxis nutzbar gemacht.

92 Vergleich der Wahlstrukturanalysen der Bundestagswahlen 2025 und 2021 in München

Boris Fischer

Statistisches Amt, Landeshauptstadt München, Deutschland
boris.fischer@muenchen.de

SECTION: VDSSt (Session 1): Wahlanalyse und statistische Helfer für die Wahlorganisation

Am Beispiel der Bundestagswahl 2025 in München soll der Einfluss der Struktur eines Gebiets auf das Wahlergebnis analysiert werden. Die Ergebnisse dieser Analyse werden mit der Bundestagswahl 2021 verglichen, auf die dasselbe Modell angewendet wurde. Hierzu wird der Zusammenhang zwischen den errungenen Stimmenanteilen der Parteien bzw. der Wahlbeteiligung und verschiedenen sozioökonomischen Variablen mit Hilfe eines generalisierten additiven Modells geschätzt. Zur Darstellung der Ergebnisse dieses Modells werden die geschätzten marginalen Effekte für jede Partei berechnet und präsentiert, sowie mit den geschätzten marginalen Effekten der Bundestagswahl 2021 verglichen.

93 It's the social, economic and infrastructure, stupid! Erklärungsfaktoren auf Gemeindeebene und die Bundestagswahlergebnisse

Johannes Kiess

Universität Leipzig, Deutschland
johannes.kiess@uni-leipzig.de

SECTION: VDSSt (Session 1): Wahlanalyse und statistische Helfer für die Wahlorganisation

Dieser Beitrag analysiert, wie sich die Zweitstimmenergebnisse bei der Bundestagswahl 2025 auf Gemeindeebene unterscheiden und welche Faktoren diese Unterschiede erklären können. Zur Erklärung der parteispezifischen geografischen Verteilungen werden eine Reihe sozial-, wirtschafts- und infrastruktureller Faktoren herangezogen. Beruhend auf umfangreichen Vorarbeiten wird eine erprobte

Variablenauswahl weiterentwickelt, die neben naheliegenden erklärenden Faktoren wie der (Entwicklung der) Gemeindegröße, dem Steueraufkommen, der Sozialstruktur oder der Arbeitslosenquote

auch innovative Variablen wie die ÖPNV-Gütekategorie sowie die Wahlbeteiligung als Variable der politischen Kultur umfasst. Die Varianz der Zweitstimmenanteile von Union, Grünen, SPD, BSW, Linke und AfD lässt sich mit den Modellen sehr gut aufklären, für die bei der Wahl sehr schwach abscheidende FDP immerhin zu einem Teil. Die Untersuchung zeigt damit, dass die Ergebnisse der Bundestagswahl stark mit sozial-, wirtschafts- und infrastrukturellen Bedingungen in den Gemeinden zusammenhängen – jedoch für verschiedene Parteien auf sehr unterschiedliche Weise.

94 Vom Wahldatensatz zur Blaupause: Datengenerierung und -aufbereitung für Kommunalwahlen im Praxiseinsatz

Leonie Schröder

Amt für Statistik und Stadtforschung, Deutschland
leonie.schroeder@wiesbaden.de

SECTION: VDSSt (Session 1): Wahlanalyse und statistische Helfer für die Wahlorganisation

Für die Kommunalwahl 2026 wurde eine umfassende digitale Wahlanalyse entwickelt, welche die technische Aufbereitung der Stimmzetteldaten und die Aufbereitung sowie Auswertung der Wahlergebnisse umfasst. Der geplante Vortrag stellt die datengetriebenen Methoden und Werkzeuge in den Vordergrund. Ziel ist es, die zugrunde liegende Infrastruktur als skalierbares und übertragbares Modell für kommunale Wahlanalysen vorzustellen.

Im Mittelpunkt steht ein modular aufgebauter Analyse-Workflow, der von der Datensatzaufbereitung über die Generierung verschiedener Outputs bis hin zur Visualisierung reicht. Dazu zählen Tabellen, Grafiken und Karten ebenso wie spezialisierte Auswertungen zur Parteienlandschaft, Sitzverteilung sowie zur Stimmenanzahl von einzelnen Kandidierenden. Die einzelnen Code-Bausteine wurden so strukturiert, dass sie klar getrennte Aufgaben erfüllen und somit leicht nachvollziehbar, wiederverwendbar und anpassbar sind. Ein besonderer Fokus liegt auch auf der standardisierten und nachnutzbaren Bereitstellung maschinenlesbarer Daten einschließlich Metadaten über unsere Open Data Plattform.

Darüber hinaus beleuchtet der Vortrag die spezifischen Anforderungen des hessischen Kommunalwahlrechts, das sich in zentralen Punkten von Regelungen anderer Bundesländer unterscheidet. Diese Unterschiede haben direkte Auswirkungen auf die Datenmodellierung und erfordern eine flexible Anpassung der Analyse-Logik bei einer Übertragung auf andere Regionen.

95 Forschungsbefunde zu Absprüngen von Bewerberinnen und Bewerbern auf Ausbildungsplätze

Ute Leber, Barbara Schwengler

Institut für Arbeitsmarkt- und Berufsforschung, Deutschland
Barbara.Schwengler@iab.de

SECTION: Labour Markets and Social Security 2

Der Ausbildungsmarkt in Deutschland steht vor großen Herausforderungen. Seit Jahren bleibt ein hoher Anteil der betrieblichen Ausbildungsstellen unbesetzt. Ein Aspekt in diesem Zusammenhang ist das (kurzfristige) Abspringen von jungen Menschen, die sich auf Ausbildungsstellen bewerben, bzw. das Nicht-Erscheinen zum Ausbildungsbeginn, das unter dem Begriff „Ghosting“ bekannt ist.

In der Forschung wurden Nicht-Antritte von Ausbildungsverhältnissen bislang nur wenig untersucht. Die Befundlage ist entsprechend dünn und beschränkt sich weitestgehend auf die Häufigkeit, mit der einerseits Jugendliche angeben, die Ausbildung einfach nicht angetreten zu haben, und mit der andererseits Betriebe äußern, davon betroffen zu sein. Zu Absprüngen in anderer Form und zu anderen Zeitpunkten im Stellenbesetzungsprozess (z.B. Vertragslösungen vor Ausbildungsbeginn), den Gründen dafür sowie ihrer Relevanz für unbesetzt bleibende Ausbildungsstellen, mangelt es bislang allerdings für Deutschland an empirischer Evidenz.

Vor diesem Hintergrund untersuchen wir das Phänomen abspringender Bewerberinnen und Bewerber und angehender Auszubildender mit dem IAB-Betriebspanel, einer jährlich bundesweit durchgeführten Wiederholungsbefragung bei Betrieben aller Branchen und Größen. Ein Schwerpunkt in dieser Befragung sind Fragen zur betrieblichen Ausbildung, die regelmäßig auch Informationen zu unbesetzten Ausbildungsstellen erhebt. In den Wellen 2013, 2022 und 2023 wurden zudem die Gründe für die Nichtbesetzung erfragt, darunter auch Absprünge von Bewerberinnen und Bewerbern. Wie Ergebnisse des IAB-Betriebspanels zeigen, ist davon rund jeder vierte Betrieb mit unbesetzten Ausbildungsplätzen betroffen. Der Beitrag untersucht, welche Rolle Bewerberabsprünge im Zusammenhang von unbesetzten Ausbildungsstellen spielen und wie sich diese im Zeitverlauf entwickelt hat. Auch wird analysiert, welche betrieblichen Merkmale damit zusammenhängen, dass Ausbildungsplätze unbesetzt bleiben, weil Bewerberinnen und Bewerber sich umentschieden haben. Neben Betriebsgröße, Branche und Region betrachten wir dabei insbesondere auch solche Faktoren, die mit der Attraktivität eines Ausbildungsbetriebs zusammenhängen.

96 Monetäre Anreize im Sozialsystem und ihre Auswirkung auf die Beschäftigung von Grundsicherungsbeziehenden

Kerstin Bruckmeier, Böhme Fabian

Institut für Arbeitsmarkt- und Berufsforschung (IAB), Deutschland
kerstin.bruckmeier@iab.de

SECTION: Labour Markets and Social Security 2

Großzügige Sozialleistungen und hohe Transferenzugsraten bei Aufnahme einer Erwerbstätigkeit werden als ein Faktor für Langzeitarbeitslosigkeit in Deutschland diskutiert. Diese Studie untersucht den Einfluss der durch das Steuer-, Abgaben- und Sozialleistungssystem gesetzten finanziellen Anreize zur Aufnahme einer Erwerbstätigkeit auf den Übergang von arbeitslosen Grundsicherungsbeziehenden in eine Beschäftigung.

Datengrundlage ist eine 10%-Zufallsstichprobe von Grundsicherungsbeziehenden, basierend auf den in der „Stichprobe der integrierten Grundsicherungsbiografien“ (SIG) organisierten SGB-II-Verwaltungsdaten aus den Jahren 2016 bis 2019. In einem ersten Schritt wird die individuelle Partizipationsbelastung auf das potenzielle Erwerbseinkommen bei Aufnahme einer Beschäftigung als einem gängigen Maß für finanzielle Erwerbsanreize im Steuer- und Transfersystem ermittelt. Die Berechnung nutzt hypothetische Einkommen, die anhand des Open-Source-Mikrosimulationsmodells GETTSIM simuliert werden. Anders als in der Literatur üblich, nutzt die Mikrosimulation Verwaltungsdaten zu Einkommen und Haushaltsstruktur und minimiert damit auf Messfehlern in Surveydaten beruhende Simulationsungenauigkeiten.

Wir finden eine durchschnittliche Partizipationsbelastung bei Aufnahme einer Vollzeitbeschäftigung von etwa 75 Prozent, was im Vergleich zu Personen außerhalb des Sozialleistungssystems deutlich erhöht ist. Da die effektive Belastung stark durch die Anrechnungsregeln für Erwerbseinkommen im Sozialleistungssystem bestimmt wird, ist die Variation der Partizipationsbelastung über verschiedene soziodemografische Merkmale gering. Eine stärkere Variation zeigt sich hingegen für eine

Differenzierung nach Regionstypen entsprechend der regionalen Streuung der in der Grundsicherung übernommenen Kosten der Unterkunft.

Durch die Verknüpfung der Grundsicherungsdaten mit administrativen Daten zu sozialversicherungspflichtiger Beschäftigung können in einem zweiten Schritt individuelle Beschäftigungsaufnahmen beobachtet und modelliert werden. Regressionsanalysen weisen auf einen signifikanten negativen Zusammenhang zwischen der Höhe der Partizipationsbelastung und der Wahrscheinlichkeit, in eine sozialversicherungspflichtige Beschäftigung überzugehen, hin. Unsere Ergebnisse zeigen eine durchschnittliche Partizipationselastizität zwischen -0,1 und -0,2, was im Einklang mit der empirischen Literatur steht. Die Wahrscheinlichkeit der Aufnahme einer Vollzeitbeschäftigung aus der Grundsicherung heraus insgesamt jedoch gering ist, sollten die erzielbaren Beschäftigungseffekte erhöhter monetärer Anreize nicht überschätzt werden.

97 Arbeitsmarktintegration von Geflüchteten im Zeitverlauf: Erkenntnisse aus kohortenbasierten Verbleibsanalysen der Arbeitsmarktstatistik

Ehsan Vallizadeh, Anton Klaus

Statistik der BA, Deutschland
ehsan.vallizadeh@arbeitsagentur.de

SECTION: Labour Markets and Social Security 2

Die Arbeitsmarktintegration von Migrantinnen und Migranten wird häufig anhand von querschnittlichen Bestandsindikatoren beschrieben. Solche Ansätze erlauben jedoch nur eingeschränkt Aussagen über individuelle Integrationsverläufe und Dynamiken im Zeitverlauf. Der vorliegende Beitrag stellt einen kohortenbasierten Analyseansatz vor, der auf Statistikdaten der Bundesagentur für Arbeit basiert und Integrationsprozesse als langfristige Verläufe derselben Personengruppe abbildet.

Grundlage der Analyse ist die Bildung von Zugangskohorten erwerbsfähiger Leistungsberechtigter (ELB) im SGB II. Der erstmalige Zugang in den Regelleistungsbezug markiert dabei den Beginn der systematischen Beobachtung. Durch die Verknüpfung von Beschäftigten-, Leistungs- und Arbeitslosenstatistik können Statusverläufe über mehrere Jahre hinweg konsistent nachvollzogen werden. Exemplarisch werden Ergebnisse für die Jahresskohorte syrischer Staatsangehöriger dargestellt.

Die Ergebnisse zeigen, dass Arbeitsmarktintegration ein langfristiger und heterogener Prozess ist. Während zu Beginn Leistungsbezug und integrationsvorbereitende Statuslagen dominieren, steigt die Beschäftigung im Zeitverlauf kontinuierlich an. So erhöht sich die Zahl der Beschäftigten innerhalb der betrachteten Jahresskohorte 2016 von rund 20.000 Personen im Jahr 2017 auf etwa 111.000 Personen bis Ende 2024, während der Leistungsbezug deutlich zurückgeht. Gleichzeitig zeigen sich ausgeprägte Unterschiede zwischen Männern und Frauen: Männer integrieren sich deutlich schneller und häufiger in Beschäftigung, während Frauen häufiger in familienbezogenen oder qualifikationsorientierten Lebenslagen verbleiben.

Darüber hinaus diskutiert der Beitrag methodische Herausforderungen kohortenbasierter Integrationsanalysen, insbesondere selektive Ausfälle aus Statistikdaten sowie die Interpretation unterschiedlicher Bezugsgrößen bei Anteilswerten. Insgesamt zeigt der Beitrag, dass kohortenbasierte Verbleibsanalysen eine wichtige Ergänzung zu klassischen querschnittlichen Arbeitsmarktindikatoren darstellen und neue Perspektiven für eine dynamische Integrationsverläufe eröffnen.

98 Macht Pflege arm? Eine differenzierte Kausalanalyse mit Paneldaten des SOEP

Janosch Korell, Torsten Harms

DHBW Karlsruhe, Deutschland
harms@dhbw-karlsruhe.de

SECTION: Labour Markets and Social Security 2

Der mögliche Verzehr von vorhandenem Vermögen mit Eintritt in die Pflegebedürftigkeit ist sowohl individuell als auch gesellschaftlich von hoher Relevanz: Neben einer individuellen Angst vor Altersarmut durch Pflegekosten und dem vollständigen Aufbrauchen der finanziellen Lebensleistung in Form von angespartem Vermögen, stellt sich auch die gesellschaftliche Frage, ob und wie vorhandenes Vermögen adäquat an möglichen Pflegekosten beteiligt wird.

Auf Basis von Paneldaten des Sozio-oekonomischen Panels (SOEP) über einen Zeitraum von 15 Jahren wird untersucht, wie sich der Eintritt in die Pflegebedürftigkeit auf das Vermögen älterer Menschen in Deutschland auswirkt. Verglichen werden Personen, die im Beobachtungszeitraum mindestens einmal pflegebedürftig wurden, mit Personen, die im Untersuchungszeitraum nie Pflege benötigten.

Zur kausalen Identifikation wird eine Fixed-Effects-Panelregression mit Jahresdummies geschätzt, welche für unbeobachtete Heterogenität sowie gemeinsame Zeittrends (insb. die Finanzkrise und der Immobilienboom während des Untersuchungszeitraums) kontrolliert. Ergänzend werden ein First-Difference-Modell und eine Dummy- Wirkungskfunktion (Event-Study-Design) verwendet, um den zeitlichen Verlauf der Vermögensänderung relativ zum Pflegeeintritt abzubilden. Der Fokus liegt auf dem Average Treatment Effect on the Treated (ATT), also dem Vermögensseffekt ausschließlich für tatsächlich Pflegebedürftige.

Die Ergebnisse zeigen einen starken einmaligen Vermögensschock beim Eintritt in die Pflege: Der Fixed-Effects-Schätzer weist einen durchschnittlichen Vermögensrückgang von rund 68 Prozent über den gesamten Beobachtungszeitraum aus, der First-Difference-Schätzer quantifiziert dabei den unmittelbaren Verlust bei Pflegeeintritt auf etwa 62 Prozent. Die Dummy- Wirkungskfunktion bestätigt, dass der stärkste Effekt zum Zeitpunkt des Pflegeeintritts auftritt und anschließend weitgehend konstant bleibt. Etwa zehn Jahre nach Pflegeeintritt verliert der Effekt an Signifikanz, was auf eine selektive Sterblichkeit hindeutet. Insgesamt deutet die Analyse darauf hin, dass Pflege das Vermögen substantiell reduziert, aber in dieser vermögenden Teilpopulation nicht zwangsläufig zu unmittelbarer Armut führt. Ein Ausblick skizziert weitere Analysen zu Einkommen, Schulden und Sozialhilfeanspruch, um die Armutsrisiken pflegebedürftiger Personen umfassender zu bewerten.

99 Data Science for a Better World: Building a Safer Digital Environment

Stefano Maria Iacus

European Commission - Joint Research Centre, Spanien
stefano.iacus@ec.europa.eu

SECTION: Computational Statistics and Data Science (Session 1)

Digital platforms have become critical environments where information spreads, collective behaviours form, and societal risks can escalate rapidly and endogenously. This talk examines how statistics can contribute to a safer, more predictable, and more trusted online world. Several risks that arise naturally from technological features of today's digital ecosystems, such as end-to-end AI recommender systems, AI companions, age-verification challenges, online fraud, and other emerging harms, can be better understood and mitigated through robust data science pipelines.

Europe's Digital Services Act, particularly Article 40(12), now guarantees controlled data access for vetted researchers. This provision represents a major opportunity for the community of social scientist and statisticians: for the first time, independent scholars can systematically study platform dynamics, audit algorithmic systems, and generate evidence that contributes directly to safer, more transparent, and accountable digital environments. Therefore, this talk will also discuss how this new opportunity enables scientists to play a more active and impactful role in improving online safety.

Overall, the talk presents a unified vision: data science, supported by responsible governance and meaningful data access, provides essential capabilities for building digital ecosystems that are safer, more predictable, and worthy of public trust.

100 Identifying economic narratives in large text corpora

Kai-Robin Lange (1), Tobias Schmidt (1), Matthias Reccius (2), Henrik Müller (1), Michael Roos (2), Carsten Jentsch (1)

1: TU Dortmund, Deutschland; 2: Ruhr Universität Bochum, Deutschland
kai-robin.lange@tu-dortmund.de

SECTION: Computational Statistics and Data Science (Session 1)

As economic and political narratives spread rapidly across digital media platforms, it has become increasingly critical to automatically extract such narratives from a corpus to gain an understanding of which narratives are spread by whom and on which platform. Previous pipelines attempting to extract narratives from a corpus of documents often employ a mix of state-of-the-art natural language processing techniques, such as BERT, to tackle this task. While effective on foundational linguistic operations essential for narrative extraction, such models lack the deeper semantic understanding required to distinguish extracting economic narratives from merely conducting classic tasks like Semantic Role Labeling.

Instead of relying on complex model pipelines, we evaluate the benefits of Large Language Models (LLMs) to extract narratives in one prompt using their great language understanding capabilities. We discuss two approaches: In a computationally heavy approach, we apply a rigorous narrative definition and compare GPT-4o interpretations of every single document in a corpus to gold-standard narratives produced by expert annotators. The second approach focuses on decreasing

the computational demand of narrative analyses by pre-selecting documents of interest using a change-detection method based on the dynamic topic model RollingLDA. We discuss our findings and provide guidance for future work in economics and the social sciences that employs LLMs to pursue similar complex objectives.

101 Computing halfspace depth in $O(n^{d-1})$ via topological sweep

Rainer Dyckerhoff (1), Erik Mendros (2), Stanislav Nagy (2)

1: University of Cologne, Cologne, Germany; 2: Charles University, Prague, Czech Republic
rainer.dyckerhoff@statistik.uni-koeln.de

SECTION: Computational Statistics and Data Science (Session 1)

We present an algorithm for computing the exact halfspace depth of a single point in $\mathbb{R}^d$, $d \geq 3$. The algorithm achieves asymptotic running time $O(n^{d-1})$, improving upon the previously best known bound $O(n^{d-1} \log n)$. It has two main components. For $d = 3$, we employ gnomonic projection to map the data to a plane, apply point-line duality, and traverse the induced line arrangement by a topological sweep. For $d > 3$, we reduce the problem to the three-dimensional case by means of a suitable projection scheme. The talk will highlight the geometric ideas underlying the topological sweep, the use of duality, and the efficient traversal of the arrangement without explicit full

102 Schätzung und Bias-Anpassung regionaler alters-, geschlechts- und nationalitätsspezifischer Mortalitätsraten bei unvollständigen und fehlerhaften Informationen - Anwendung für die deutschen Bundesländer

Patrizio Vanella (1,2,3), Julian Ernst (4)

1: Abteilung für Gesundheitsberichtserstattung & Biometrie, aQua-Institut, Göttingen; 2: Lehrstuhl für Empirische Sozialforschung und Demographie, Universität Rostock; 3: Arbeitskreis Demografische Methoden, Deutsche Gesellschaft für Demographie; 4: Professur für Wirtschafts- und Sozialstatistik, Universität Trier
patrizio.vanella@qua-institut.de

SECTION: DGD (Session 1): Methodische Herausforderungen, Lösungsansätze und aktuelle Entwicklungen in der Regionaldemografie I

Regionale Mortalitätsraten sind zentrale Indikatoren für die Gesundheitsüberwachung und -planung; ihre belastbare Schätzung ist jedoch häufig durch Datenprobleme eingeschränkt. Wir entwickeln daher ein statistisches Rahmenwerk zur Schätzung und Bias-Adjustierung regionaler, alters-, geschlechts- und nationalitätsspezifischer Mortalitätsraten in Deutschland für den Zeitraum 1995 bis 2023. Ausgangspunkt ist die Beobachtung, dass insbesondere bei kleinräumigen Einheiten systematische Verzerrungen auftreten können: geringe Fallzahlen erhöhen die Volatilität und fehlerhafte Bevölkerungszahlen – verstärkt durch Interzensalperioden zwischen Volkszählungen – verfälschen die Expositionsbasis und damit die abgeleiteten Raten.

Zur Modellierung der Todesfälle verwenden die Autoren eine negative Binomial-Regression im Rahmen Generalisierter Linearer Modelle (GLM), um Überdispersion gegenüber Poisson-Ansätzen angemessen abzubilden. In das Modell fließen demografische Merkmale (Alter, Geschlecht, Nationalität), migrationsbezogene und regionale Einflussfaktoren, Zeittrends sowie spezifische Bias-Quellen

ein, darunter Fortschreibungsfehler und pandemiebezogene Effekte. Zusätzlich werden Interaktionsterme berücksichtigt, um heterogene Muster zwischen Bevölkerungsgruppen, Regionen und Zeitverläufen abzubilden. Damit werden sowohl langfristige Entwicklungen (z.B. steigende Lebenserwartung) als auch kurzfristige Schocks wie die COVID-19-Pandemie explizit modelliert.

Die Anpassung verschiedener Modellspezifikationen wird mit Güte- und Effizienzmaßen verglichen. Dabei zeigt sich, dass eine regionale Modellierung auf Ebene der 16 Bundesländer deutlich bessere Ergebnisse liefert als ein globales Modell. Insgesamt belegen die Resultate, dass die Nutzung detaillierter Stratifizierungen sowie die explizite Berücksichtigung von Bias-Faktoren die Schätzgenauigkeit erheblich steigert. Das vorgeschlagene Framework ist flexibel auf andere Länder und räumliche Ebenen übertragbar und unterstützt eine verbesserte Mortalitätsüberwachung, evidenzbasierte Gesundheitsplanung sowie die Korrektur fehlerhafter Bevölkerungsdaten.

103 Sind Sie sicher? Eine Unsicherheitsanalyse am Beispiel regionaler Bevölkerungsprojektionen in Niedersachsen

Falk Voit (1), Patrizio Vanella (2)

1: Landesamt für Statistik Niedersachsen, Deutschland; 2: aQua-Institut, Deutsche Gesellschaft für Demographie
falk.voit@statistik.niedersachsen.de

SECTION: DGD (Session 1): Methodische Herausforderungen, Lösungsansätze und aktuelle Entwicklungen in der Regionaldemografie I

Regionale Bevölkerungsprojektionen stellen ein vielfach unverzichtbares Zahlenwerk für eine Reihe zentraler (planerischer) Entscheidungen zur Verfügung. Häufig wird von den Entscheidern in diesem Kontext nur ein Punktschätzer als relevant erachtet und Prognoseintervalle werden nicht selten vernachlässigt. Dabei sind gerade kleinräumige und regionale Projektionen mit Unsicherheiten verbunden. Doch wie groß sind diese Unsicherheiten?

Um die Unsicherheit und Genauigkeit kleinräumiger Bevölkerungsprojektionen besser einschätzen zu können, werden die Ergebnisse der 3. und 4. regionalisierten Bevölkerungsvorausberechnung (rBV) des Landesamtes für Statistik Niedersachsen (LSN) mit den tatsächlichen amtlichen Einwohnerzahlen verglichen. Während die 3. rBV am Basisjahr 2020 ansetzte, wurde die 4. rBV auf Grundlage der Einwohnerzahlen des Jahres 2022 berechnet. Folglich können jeweils vier bzw. sechs Jahre der Vorausberechnung mit der tatsächlichen Bevölkerungsentwicklung verglichen werden. Die rBV des LSN ist eine deterministische Bevölkerungsvorausberechnung nach Kohorten-Komponenten-Methode, welche die Unsicherheit der zukünftigen Bevölkerungsentwicklung mit verschiedenen Szenarien abzubilden versucht. Diese Szenarien spannen ein Intervall auf, in welchem sich die zukünftige Bevölkerungsentwicklung bewegen könnte. Mit Blick auf die aufgespannten Intervalle sollen folgende Fragen beantwortet werden:

- Wie deutlich weichen die Projektionen nach Hauptvariante von der tatsächlichen Einwohnerzahl ab?
- Wie stark wächst die Abweichung über die Zeit?
- Liegt die Abweichung innerhalb des Intervalls, das sich aus den Berechnungsvarianten ergibt?
- Waren spätere Berechnungen genauer mit Blick auf die tatsächliche Einwohnerzahl?
- Gibt es regionale Auffälligkeiten?
- Was sind Erklärungsansätze für mögliche Abweichungen? Hätten diese Aspekte schon im Vorfeld berücksichtigt werden können?

Im Nachgang wird aufgezeigt, wie die Abweichung vergangener Projektionen zu den tatsächlichen Einwohnerzahlen in die Modellierung der Unsicherheit zukünftiger Bevölkerungsprojektionen einbezogen werden könnten. Weiterhin gibt der Beitrag einen kurzen Ausblick, welche Vorteile

sich bei der Verwendung probabilistischer Verfahren im Hinblick auf die Modellierung bestehender Unsicherheiten ergeben, aber auch, welche fundamentalen Herausforderungen bestehen bleiben.

104 A Bayesian Approach to Municipal Cohort-Component Forecasting

Georg Wiegleb

Landeshauptstadt Magdeburg, Deutschland
georg.wiegleb@stadt.magdeburg.de

SECTION: DGD (Session 1): Methodische Herausforderungen, Lösungsansätze und aktuelle Entwicklungen in der Regionaldemografie I

Municipal population forecasting is particularly challenging when births, deaths, and migration flows must be estimated across many small strata defined by age, sex, and fine-grained spatial units such as urban districts. Event counts in these strata are not only often sparse and unevenly observed; demographic patterns can also vary substantially over time and remain highly sensitive to temporary shocks. This becomes especially visible during disruptive periods, when political contexts or crises produce pronounced shifts in in- and out-migration that temporarily reshape demographic patterns across population groups and spatial units.

This talk discusses how Bayesian modeling can support municipal cohort-component forecasting under these conditions. Separate models for births, deaths, and migration are linked through posterior predictive simulation, carrying uncertainty and characteristic temporal patterns forward into future population trajectories. Established demographic schedules provide reference structures that stabilize estimation when data is sparse, noisy, or potentially misleading, while preserving data-driven flexibility in the estimated patterns. For in-migration, a generative posterior predictive setup simulates synthetic inflow patterns sequentially, incorporating shock-adjusted dynamics and their consequences for future population composition. The approach is illustrated using administrative data from a German municipality.

105 NeDaMo – Kalibrierung von Mobilfunkdaten mithilfe einer Sondererhebung

Gloria Deetjen, Armando Häring

Statistisches Bundesamt, Deutschland
armando.haering@destatis.de

SECTION: Methodology of Statistical Surveys (Session 1)

Zahlreiche Forschungsarbeiten der letzten Jahre zur Erschließung von Mobilfunkdaten haben aufgezeigt, dass ohne relevante Zusatzinformationen eine qualitätsgesicherte Anwendung in der amtlichen Statistik nach heutigem Stand nicht umsetzbar ist.

Das vom Bundesministerium für Verkehr (BMV) im Rahmen der Innovationsinitiative „mFUND“ geförderte Projekt „Neue Daten für Mobilitätsanalysen“ (NeDaMo) strebt daher eine Kalibrierung von Mobilfunkdaten mithilfe einer maßgeschneiderten Erhebung nach §7 Bundesstatistikgesetz an. Erstmals werden dazu auf freiwilliger Basis bis zu 20.000 Haushalte zu ihrer Mobilfunknutzung und ihrem Mobilitätsverhalten befragt. Anschließend werden in Kombination mit weiteren Statistiken Verfahren zur statistisch validen Nutzung von Mobilfunkdaten entwickelt, um diese Datenquelle

zur qualitativ hochwertigen Nutzung im Rahmen des Bundesverkehrswege- und Mobilitätsplan zu erschließen. Gleichzeitig kommen im Projekt Daten gleich mehrerer Mobilfunknetzbetreiber zum Einsatz. Zuletzt werden qualitativ optimierte Verflechtungsmatrizen im Rahmen von Mikrosimulationen pilothaft für ausgewählte Regionen zur Prognose von Mobilitätsszenarien genutzt. Im Beitrag wird das Projektvorhaben vorgestellt sowie Zwischenergebnisse zur Konzeption und Durchführung der Sondererhebung nach §7-BstatG präsentiert.

106 Assessing the Quality of Survey Weights in Panel Studies: Limitations of Design Effects and Effective Sample Size

Anne Konrad

Leibniz-Institut für Bildungsverläufe, Deutschland
anne.konrad@ifbi.de

SECTION: Methodology of Statistical Surveys (Session 1)

Assessing the quality of survey weights is essential for valid statistical inference. This is particularly crucial in panel studies, where panel attrition and the multiplication of wave-specific weighting factors result in heterogeneous weight distributions. In practice, widely used quality measures are design effects and effective sample size. Both diagnostics are driven by weight variability and are commonly interpreted as indicators of efficiency loss due to unequal weighting. However, increased weight variability does not necessarily imply an increased variance of survey estimates. This is well illustrated by sampling designs such as probability-proportional-to-size sampling, where unequal weights may even lead to efficiency gains. Furthermore, survey weights are composed of multiple components, including design weights, nonresponse adjustment factors, and calibration factors. While the design weight reflects the sampling design, nonresponse adjustments and calibration factors are explicitly introduced to reduce nonresponse and coverage bias. In the presence of strong participation selectivity or binding calibration constraints, variable weights may therefore be both necessary and desirable.

In addition, the commonly used approximation of the design effect proposed by Kish relies on restrictive assumptions, including independence between weights and study variables and homoskedasticity. These assumptions are rarely met in practice and are particularly implausible in longitudinal panel studies, where participation and study variables are often endogenously related over time.

The aim of this paper is therefore to critically review the use of design effects and effective sample size as measures of survey weight quality. Using a Monte Carlo simulation study, we examine whether these diagnostics are informative indicators of key quality criteria such as bias and mean squared error or whether they systematically diverge under conditions typical of longitudinal panel studies.

107 Using Partial Least Squares Regression for Imputation of Multinomial Variables

Björn Rohr, Christian Bruch, Christof Wolf

GESIS - Leibniz Institute for Social Sciences, Deutschland
Bjoern.rohr@gesis.org

SECTION: Methodology of Statistical Surveys (Session 1)

Imputing multinomial variables with multiple imputations in social science surveys may require large imputation models with many parameters. The models should include all analysis variables to avoid bias. Furthermore, the imputation models must include variables needed to reduce nonresponse bias and might include variables correlated to the imputation variable to improve imputation quality. This might result in a large complexity of imputation models. In case of multinomial variables, this complexity can further increase with the number of categories. However, given the limited sample sizes in many surveys, common imputation methods such as multinomial logistic regression-based imputation or predictive mean matching (PMM) might include only a limited number of variables. For example, in multinomial logistic regression-based imputation, small sample sizes can lead to infinite estimates due to perfect prediction or limited precision. In other words, common imputation scenarios both require and suffer from large imputation models when the sample size is small. A potential solution for this problem might be to use partial least squares regression (PLS) prior to imputation, as PLS can reduce the dimensionality of the predictors while preserving information relevant to the imputation.

In a Monte Carlo simulation study, we evaluate three PLS-based methods for imputing multinomial variables. The first two methods combine an adaptation of PLS for multinomial variables (PLS-DA) with different PMM implementations for multinomial variables. In the third method, multiple correspondence analysis is used to transform all variables to metric variables, before imputing missing values with a combination of PLS and PMM. To compare these methods, the data and item non-response are simulated based on 87 variables from the European Values Study 2017 across various levels of overall non-response. In this presentation, we show initial results on how well the imputed datasets reproduce bivariate chi-square correlations between multinomial variables.

108 Microdata Linking in Official Statistics: Opportunities and Challenges

Hoda Mohamed

The Federal Statistical Office of Germany
hoda.mohamed@destatis.de

SECTION: Methodology of Statistical Surveys (Session 1)

The growing availability of administrative and statistical microdata provides an opportunity to modernize data production and reduce the burden on enterprises. Microdata linking (MDL) can leverage this infrastructure to enable statistical offices to reuse existing data sources more efficiently, minimize redundant survey collection, and improve data quality. By integrating information from business registries, administrative records, and survey data, microdata linking can reduce reporting obligations for enterprises while maintaining the richness of statistical outputs. Microdata linking also opens new analytical possibilities by combining complementary data sources, enabling the

creation of new indicators, improving coverage, and supporting more granular, policy-relevant analyses.

Drawing on applied examples, the presentation showcases different linking strategies and discusses implementation challenges. Two concrete cases illustrate the practical value of microdata linking: first, the replacement of an employment module in the Global Value Chains (GVC) survey through administrative data linkage — demonstrating the direct potential for response burden reduction for enterprises and second, the development of new indicators on globalization and GVC integration by combining complementary data sources, showcasing how linking expands the analytical scope of existing survey data. The presentation emphasizes microdata linking as an instrument for achieving more efficient, less burdensome, and more informative official statistics.

109 Characterizing M-Estimators using Canonical Divergences from Information Geometry

Paul Navas

TU Dortmund, Deutschland
paul.navas@tu-dortmund.de

SECTION: Statistical Theory and Methods (Session 1)

Recent work by Dimitriadis, Fissler and Ziegel establishes a fundamental connection between the theory of M-estimation and the theory of forecast evaluation by showing that model-consistent M-estimators are characterized through strictly consistent loss functions for elicitable functionals. In particular, classical M-estimation for conditional means, quantiles and related functionals can be understood through the geometry of strictly consistent scoring rules, such as Bregman losses and their associated identification functions.

The present paper extends this program from the dually flat setting underlying Bregman divergences to general statistical manifolds equipped with a dualistic information-geometric structure. Using the canonical divergences introduced by Amari and Ay, we show that strictly consistent loss functions arise canonically from the intrinsic geometry of the model space itself, rather than being externally specified. This yields a geometric construction of elicitable functionals directly from the dual affine structure of information geometry.

The resulting estimators are characterized as minimizers of expected canonical divergences and admit first-order conditions expressed through inverse exponential maps and geometric barycenters. In the self-dual Riemannian case, the construction reduces to squared geodesic distance and recovers Fréchet means and their empirical counterparts, thereby connecting the theory to modern non-Euclidean statistics in the sense of Panaretos and collaborators. In the dually flat case, the framework specializes to the classical correspondence between Bregman divergences, exponential families and elicitable expectations.

This perspective provides a unified geometric framework linking elicibility, Osband's principle, M-estimation, information geometry and Fréchet-type estimation. More broadly, it suggests that consistent loss functions should be viewed as manifestations of an underlying geometric divergence structure, thereby placing forecast evaluation and semiparametric estimation within a common differential-geometric theory.

110 Uncertainty Quantification in Forecast Comparisons

Marc-Oliver Pohle (1,2), **Tanja Zahn** (3), Sebastian Lerch (4,2)

1: University of Wuppertal, Germany; 2: Heidelberg Institute for Theoretical Studies, Germany; 3: Goethe University Frankfurt, Germany; 4: Marburg University, Germany
tzahn@wiwi.uni-frankfurt.de

SECTION: Statistical Theory and Methods (Session 1)

Skill scores, which measure the relative improvement of a forecasting method over a benchmark via consistent scoring functions and proper scoring rules, are a standard tool in forecast evaluation, yet their sampling uncertainty is rarely rigorously quantified. With modern forecasting applications being increasingly multivariate and involving evaluations across multiple horizons, variables, spatial locations, and forecasting methods, standard tools like the pairwise Diebold-Mariano forecast accuracy test or pointwise confidence intervals fail to account for the multiple comparison problem, leading to inflated Type I error rates and invalid joint inference. To address the lack of a coherent, statistically rigorous framework for quantifying uncertainty across these multi-dimensional evaluation problems, we introduce simultaneous confidence bands for expected scores and skill scores. Our framework provides a versatile tool for joint inference that is applicable to any forecast type from mean and quantile to full distributional forecasts. We develop a bootstrap implementation and show that our bands are valid under multivariate extensions of the classical Diebold-Mariano assumptions. We demonstrate the practical utility of the approach in two case studies by quantifying the benefits of time-varying parameter models for macroeconomic forecasting, and by comparing data-driven and physics-based models in probabilistic weather forecasting.

111 A robust test for equal predictive accuracy

Matei Demetrescu (1), Christoph Hanck (2), Yannick Hoga (2)

1: TU Dortmund University; 2: University of Duisburg-Essen
mdeme@statistik.tu-dortmund.de

SECTION: Statistical Theory and Methods (Session 1)

This paper studies what we call a "robust" Diebold–Mariano (DM) type test.

The unique feature of our test is that - even in the absence of any knowledge of the forecasting method - it is robust to estimation noise in the forecasts, i.e., size is kept irrespective of estimation effects induced by model fitting.

This feature is obtained by a test statistic that is based on rolling-window means whose length is a vanishing fraction of the total evaluation sample.

This leads to non-standard Gumbel limit laws.

Other desirable features of our test are that it is easily robustified against time-varying volatility, and that it deals naturally with time-varying differences in predictive ability under the alternative. Simulations demonstrate the benefits of our multiply robust implementation vis-a-vis several competitors.

An empirical application to forecasts for several variables, horizons, vintages and methods from the Survey of Professional Forecasters illustrates the relevance of the new approach, allowing us to identify forecasters with superior models.

Such conclusions are in fact impossible to infer by extant tests, since information on the models and estimation procedures behind the forecasts are typically proprietary and, hence, estimation effects cannot be factored out.

112 Sequential Comparison of Predictive Ability with Unbounded Score Differences

Timo Dimitriadis (1,3), **Jan-Lukas Wermuth** (1), Johanna Ziegel (2,3)

1: Goethe Universität Frankfurt, Deutschland; 2: Heidelberg Institut für Theoretische Studien; 3: ETH Zürich
lukaswermu@gmail.com

SECTION: Statistical Theory and Methods (Session 1)

Time series forecasting and its evaluation are inherently sequential. However, their classical comparison via the Diebold–Mariano test is valid only for fixed evaluation horizons and does not allow for continuous monitoring. We develop sequentially valid inference on forecast performance for general, potentially unbounded loss differentials in multi-step-ahead forecasting based on asymptotic confidence sequences. These confidence sequences provide asymptotic coverage guarantees uniformly over time, thereby enabling online monitoring and tracking of loss differentials while permitting data peeking. We document the favorable finite-sample properties of our procedure in simulations and demonstrate its practical relevance in an application to volatility forecasting.

113 Künstliche Intelligenz in der Kommunalverwaltung: Ein Überblick über den KI-Einsatz in den Kommunen

Anika Groß

KGSt, Deutschland
anika.gross@kgst.de

SECTION: VDSt (Session 2): KI - Einsatzrahmenbedingungen und Nutzen

Mit generativer KI wie ChatGPT, Claude oder Gemini hat die Debatte um Künstliche Intelligenz eine neue Ebene erreicht. Auch in deutschen Kommunalverwaltungen ist das Thema angekommen, nicht mehr als technologische Zukunftsmusik, sondern als konkrete Frage der Organisations- und Datenarbeit. Der Vortrag verortet KI im kommunalen Kontext, zeigt praxisnahe Beispiele und richtet den Blick gezielt auf die Schnittstelle zum Thema Daten.

Drei Themenbereiche stehen im Mittelpunkt. Erstens: Was ist KI eigentlich, und wie grenzt sie sich von determinierter Automation ab? Funktionsweisen und Systemtypen werden eingeordnet: von regelbasierten Algorithmen über klassisches maschinelles Lernen bis zu generativen Modellen. Zweitens: Wo entstehen Mehrwerte in der kommunalen Praxis? In den Bereichen operative Exzellenz, verbesserte Entscheidungsfindung sowie Nutzendenerlebnis und Inklusion werden praktische Beispiele aus Kommunen gezeigt. Drittens: Welche Rahmenbedingungen sind zu beachten? Hier geht es um einen kurzen und beispielhaften Überblick über die technischen, rechtlichen, ethischen, organisationalen und kulturellen Voraussetzungen verantwortungsvoller KI-Nutzung, eingebettet in kommunale Governance-Strukturen.

114 Agentic Coding am Beispiel GitHub Copilot. Ein Praxisbeispiel aus Wolfsburg

Jennifer Kreklow, Jan Lunge

Stadt Wolfsburg, Deutschland
jennifer.kreklow@stadt.wolfsburg.de

SECTION: VDSt (Session 2): KI - Einsatzrahmenbedingungen und Nutzen

KI-gestützte Coding-Assistenten wie GitHub Copilot sind inzwischen weit verbreitet und unterstützen Statistikerinnen und Statistiker durch Autovervollständigung und Chat-basierte Hilfe. Mit Agentic Coding entsteht nun eine neue Stufe: KI-Agenten, die eigenständig planen, Werkzeuge aufrufen, Ergebnisse prüfen und iterativ auf ein Ziel hinarbeiten – während der Mensch als Auftraggeber, Reviewer und Qualitätssicherer die Kontrolle behält (human-in-the-loop).

Dieser Vortrag führt in die Grundlagen des Agentic Coding ein und grenzt das Konzept von klassischen KI-Coding-Assistenten sowie dem unkontrollierten „Vibe Coding“ ab. Er stellt die Kernkonzepte vor und erklärt den typischen Agentic Loop, in dem ein Agent Aufgaben in Teilschritte zerlegt, Werkzeuge nutzt und sein Vorgehen reflektiert.

Dazu werden auch Sicherheitsaspekte und Best Practices behandelt: Wie gebe ich Leitplanken vor und setze Grenzen? Wie funktioniert das Model Context Protocol (MCP) als offener Standard für die Tool-Integration? Und wie können eigene Skills und spezialisierte Agents erstellt werden, um wiederkehrende Aufgaben zu automatisieren?

115 LLMs als Analysewerkzeug für offene Antworten bei Bürgerbefragungen

Annelie Heuser

Landeshauptstadt Wiesbaden, Deutschland
annelie.heuser@wiesbaden.de

SECTION: VDSt (Session 2): KI - Einsatzrahmenbedingungen und Nutzen

Offene Textantworten aus Bürgerumfragen enthalten wertvolles qualitatives Wissen, binden aber aufgrund des hohen manuellen Auswertungsaufwands viele Ressourcen und bleiben deshalb in der Praxis manchmal ungenutzt. Der Beitrag stellt eine KI-gestützte Pipeline zur automatisierten Inhaltsanalyse von offenen Bürgerantworten vor, die im Rahmen einer Online-Befragung zum Sicherheitsgefühl junger Menschen in Wiesbaden weiterentwickelt und erprobt wurde.

Eingebunden sind Large Language Models (LLMs) für aufeinander aufbauende Aufgaben: die induktive Generierung einer thematischen Taxonomie aus den Rohdaten, die anschließende Klassifikation der Antworten sowie weitere Analysen wie die automatisierte Extraktion repräsentativer O-Töne und konkreter Handlungsempfehlungen für die Stadtverwaltung. Ein besonderer Fokus lag dabei auf der Handhabung politisch sensibler Inhalte: Durch gezielte Prompt-Gestaltung werden Antworten konsequent nach ihrer konkreten Forderung kategorisiert, nicht nach politischem Subtext. Mehrfachkodierungen sind explizit vorgesehen.

Der Beitrag diskutiert die Grenzen des Ansatzes: Abhängigkeit von externen API-Diensten, Datenschutzanforderungen beim Einsatz kommunaler Umfragedaten sowie die Frage, wo manuelle Prüfung zwingend erforderlich bleibt.

Der Workflow ist quelloffen gestaltet und auf andere kommunale Befragungskontexte übertragbar.

116 Schuldenbremse und Investitionspaket - Wirkung der Grundgesetzänderung auf den Arbeitsmarkt

Jonas Krinitz (1), Christian Schneemann (2)

1: Gesellschaft Wirtschaftliche Strukturforchung, Deutschland; 2: Institut für Arbeitsmarkt- und Berufsforschung (IAB)
krinitz@gws-os.com

SECTION: Labour Markets and Social Security 1

Nachdem Geld für Infrastruktur und Militärausgaben nicht mehr knapp zu sein scheint, sind wir der Frage nachgegangen, welche neuen Flaschenhälse sich aus den Mehrinvestitionen ergeben und was es braucht, damit diese Investitionssummen ergebniswirksam Verwendung finden können. Hierzu haben wir mit dem QuBe-Modell den Brutto- und Nettoeffekt der Pakete gerechnet und diese so parametrisiert, dass die Wirkung auf dem Arbeitsmarkt kontinuierlich und, soweit es geht, realistisch umsetzbar ist. Die dabei aufkommenden Limitationen führen zu Annahmen bzw. Gelingensbedingungen, die die Umsetzbarkeit und wirtschaftliche Verträglichkeit erheblich beeinflussen. Über diese hinaus könnten durch flexibel strukturierte Rahmenbedingungen die Aufwuchseffekte nachhaltig in Wachstum umgemünzt werden und gleichzeitig die Preiseffekte abgemildert werden.

117 Demografie, Zuwanderung und Beschäftigungswachstum in Deutschland: Evidenz aus Arbeitsmarkt- und Decompositionsanalysen

Davit Adunts (2), Ehsan Vallizadeh (1), Anton Klaus (1)

1: Statistik der BA, Deutschland; 2: Institut für Arbeitsmarkt- und Berufsforschung
ehsan.vallizadeh@arbeitsagentur.de

SECTION: Labour Markets and Social Security 1

Der deutsche Arbeitsmarkt verzeichnete in den vergangenen Jahren trotz demografischer Alterung, strukturellem Wandel und konjunktureller Schwäche ein bemerkenswert robustes Beschäftigungswachstum. Der Beitrag bündelt die Ergebnisse zweier komplementärer Analysen: eines deskriptiven Arbeitsmarktberichts der Statistik der Bundesagentur für Arbeit zur Beschäftigungsentwicklung nach Staatsangehörigkeit sowie einer decomposition-basierten Analyse zu den Triebkräften des Beschäftigungswachstums in Deutschland zwischen 2017 und 2024. Gemeinsam ermöglichen beide Ansätze eine differenzierte Einordnung der Bedeutung von Zuwanderung, Erwerbsbeteiligung und demografischer Entwicklung für den deutschen Arbeitsmarkt.

Die Ergebnisse zeigen, dass der Anstieg des Erwerbspersonenpotenzials in Deutschland in den vergangenen Jahren nahezu vollständig auf Zuwanderung zurückzuführen ist, während die Bevölkerung deutscher Staatsangehörigkeit im erwerbsfähigen Alter bereits deutlich rückläufig ist. Gleichzeitig hat sich die Beschäftigung ausländischer Staatsangehöriger – insbesondere von Drittstaatsangehörigen – dynamisch entwickelt. Die decomposition-basierte Analyse verdeutlicht, dass das Beschäftigungswachstum vor allem durch zwei Faktoren getragen wurde: erstens durch das Wachstum der Bevölkerung im erwerbsfähigen Alter und zweitens durch steigende Beschäftigungsquoten innerhalb verschiedener Bevölkerungsgruppen. Demografische Alterung wirkte dagegen

bereits beschäftigungsdämpfend. Positive Beiträge kamen insbesondere von Migrantinnen und Migranten, Frauen sowie älteren Beschäftigten.

Die deskriptiven Ergebnisse zeigen darüber hinaus erhebliche strukturelle Unterschiede nach Altersgruppen, Qualifikationsniveau, Berufssegmenten, Branchen und Regionen. Während deutsche Beschäftigte vor allem Beschäftigungszuwächse in höher qualifizierten Tätigkeiten verzeichnen, kompensieren ausländische Beschäftigte zunehmend Rückgänge bei Fachkraft- und Helfertätigkeiten. Besonders hoch ist ihr Beitrag in Engpassberufen, personenbezogenen Dienstleistungen sowie in ostdeutschen Regionen mit starkem demografischem Rückgang. Hinweise darauf, dass die steigende Beschäftigung ausländischer Staatsangehöriger zulasten deutscher Beschäftigung geht, ergeben sich aus den vorliegenden Daten nicht.

Insgesamt verdeutlichen die Ergebnisse, dass Zuwanderung und erfolgreiche Arbeitsmarktintegration zentrale Voraussetzungen für die Stabilisierung von Beschäftigung und Fachkräftesicherung in einer alternden Gesellschaft darstellen.

118 Wer arbeitet wie viel – und wie viel wäre gewünscht? Arbeitszeitrealitäten und -präferenzen im Geschlechtervergleich und nach Lebensform

Martina Rengers

Statistisches Bundesamt, Deutschland
martina.rengers@destatis.de

SECTION: Labour Markets and Social Security 1

Wer arbeitet wie viel – und wie viel wäre eigentlich gewünscht? Diese Frage steht im Zentrum einer Arbeitszeitdebatte, die in Deutschland derzeit so politisiert ist wie lange nicht mehr. Während die Teilzeitquote mit rund 40 Prozent einen historischen Höchststand erreicht hat, entzündet sich an ihr eine grundlegende Auseinandersetzung darüber, wie viel Arbeit eine Gesellschaft braucht.

Im politischen Diskurs wird dabei zunehmend zugespitzt argumentiert: Unter dem Schlagwort „Lifestyle-Teilzeit“ wird suggeriert, die Deutschen würden aus Gründen einer besseren Work-Life-Balance zu wenig arbeiten, obwohl mehr notwendig wäre. Forderungen, den Rechtsanspruch auf Teilzeit einzuschränken, verweisen auf Fachkräftemangel und wirtschaftlichen Druck. Gleichzeitig halten viele Stimmen dagegen, dass diese Diagnose an der Lebensrealität vorbeigeht: Teilzeit ist für viele kein „Lifestyle“, sondern eine Voraussetzung dafür, Erwerbsarbeit überhaupt leisten zu können – etwa wegen Kinderbetreuung, Pflege oder fehlender Arbeitszeitmodelle.

Gerade im Geschlechtervergleich zeigt sich, wie eng Arbeitszeitrealitäten mit sozialen Strukturen verknüpft sind. Frauen arbeiten deutlich häufiger in Teilzeit als Männer, oft nicht primär aus freier Präferenz, sondern im Kontext von Sorgearbeit und institutionellen Rahmenbedingungen. Gleichzeitig zeigen Befragungsergebnisse, wie beispielsweise der Mikrozensus, dass in vielen Fällen die tatsächlich geleisteten Arbeitszeiten nicht mit den gewünschten übereinstimmen: Ein Teil der Beschäftigten – insbesondere Frauen – würde gerne mehr arbeiten, während andere ihre Arbeitszeit reduzieren möchten.

Vor diesem Hintergrund gewinnt auch die Perspektive der Lebensformen an Bedeutung. Ob mit oder ohne Kinder, in Partnerschaft oder alleinlebend – Arbeitszeitwünsche sind eingebettet in unterschiedliche biografische und soziale Konstellationen. Die Frage nach „zu viel“ oder „zu wenig“ Arbeit lässt sich daher nicht pauschal beantworten, sondern bedarf einer differenzierten Betrachtung.

119 Zukunft der Gesundheitsberufe: Erste Ergebnisse des Fachkräftemonitorings für das BMG bis 2040

Anja Sonnenburg, Ines Thobe

Gesellschaft für Wirtschaftliche Strukturforschung mbH, Deutschland
sonnenburg@gws-os.com

SECTION: Labour Markets and Social Security 1

Der Bedarf an Fachkräften in den Gesundheitsberufen nimmt seit Jahren zu. Aufgrund demografischer Entwicklung, geänderter Verhaltensweisen und rechtlicher Gestaltung steigt dieser vermutlich auch in kommenden Jahren weiter an. Der Arbeitsmarkt wird hingegen durch einen Fachkräftewettbewerb geprägt sein. Daher stellt sich die Frage, wie eine qualitativ hochwertige Gesundheitsversorgung flächendeckend sichergestellt werden kann.

Um frühzeitig berufliche Passungsprobleme in den Berufen des Gesundheitswesens zu erkennen, hat das Bundesministerium für Gesundheit das Bundesinstitut für Berufsbildung, das Institut für Arbeitsmarkt und Berufsforschung und die Gesellschaft für Wirtschaftliche Strukturforschung mit dem Aufbau eines detaillierten langfristigen Monitorings von Angebot und Bedarf in den Gesundheitsberufen beauftragt.

Der Vortrag stellt die Ergebnisse der ersten Basisprojektion des Fachkräftemonitorings vor. Auf Grundlage einer umfangreichen Datenbasis wird der realisierte Bedarf an Erwerbstätigen sowie das Arbeitsangebot für 54 Berufe des Gesundheitswesens bis zum Jahr 2040 ermittelt. Die Zahl der zukünftig benötigten Erwerbstätigen wird mit den voraussichtlich für einen Beruf zur Verfügung stehenden Erwerbspersonen bilanziert, um Knappheiten zu erkennen.

Zentrale Ergebnisse sind ein Erwerbstätigenanstieg bis 2040 und auch die Zahl der Erwerbspersonen wird für die Berufe wachsen, obwohl sie für die Volkswirtschaft insgesamt schrumpft. Allerdings ist die Entwicklung nach Berufen heterogen. Bundesweit betrachtet verknappt sich in 23 der 54 Berufe das Arbeitsangebot im Verhältnis zum Bedarf, in 31 Berufen findet hingegen eine Ausweitung im Vergleich zum Bedarf statt.

Das höchste Erwerbstätigenwachstum und auch die höchsten Engpässe zeigen sich bei den Pflegeberufen. Berufsspezifisch lassen sich im Vortrag zudem die Entwicklungen in der Human- und Zahnmedizin, in den Therapie- und Heilkundeberufen, bei (zahn-)medizinischen Fachangestellten, in Berufen des Gesundheitshandwerks und in (medizinisch-)technischen Berufe sowie in Apothekenberufen hervorheben.

120 Scalable Estimation of Large Generalized Additive Models: Extended Fellner-Schall via Hutch++ and Conjugate Gradients

Monika Zimmermann (1), Claudia Collarin (2), Simon Wood (2), Florian Ziel (1)

1: Universität Duisburg-Essen; 2: University of Edinburgh
monika.zimmermann@uni-due.de

SECTION: Computational Statistics and Data Science (Session 2)

We address the computational bottleneck arising from formation and factorization of the penalized Hessian in large generalized additive models (GAMs) with high-dimensional parameter vectors. We combine the extended Fellner-Schall smoothing parameter update with stochastic trace estimation (Hutch++) and preconditioned conjugate gradients (PCG), thereby avoiding formation or Cholesky

factorization of the GAM penalized Hessian. The resulting procedure relies only on matrix–vector products, enables low memory-bandwidth parallelization and allows exploitation of model term sparsity with minimal infill. The performance of the approach is demonstrated on the NMMAPS respiratory mortality data with over one million observations and more than 20,000 coefficients, fitted in just over half an hour.

121 Online Regularized Generalized Linear Models for Vine Copulas

Christian Schulz (1), Christoph Hanck (2), Simon Hirsch (3), Florian Ziel (4)

1: Universität Duisburg-Essen; 2: Universität Duisburg-Essen; 3: Statkraft Trading GmbH; 4: Universität Duisburg-Essen
christian.schulz@vwl.uni-due.de

SECTION: Computational Statistics and Data Science (Session 2)

We propose an online estimation algorithm for vine copula models in which the dependence parameters are allowed to change with covariates. The algorithm updates the parameters sequentially as new observations become available. Existing approaches require full batch re-estimation using all available data at each update step, which makes them computationally prohibitive in large-scale forecasting applications and less suitable for settings in which dependence structures evolve over time. Simulation results show that the online procedure achieves forecasting accuracy comparable to that of the batch estimator, while reducing computation time by a factor of up to 17. We illustrate the method in a joint forecasting application to net electricity load across five regions of Great Britain. The algorithm is implemented in the Python package `ondil`.

122 Fast Training of Mixture-of-Experts for Time Series Forecasting via Expert Loss Integration

Btissame El Mahtout (1,2), Florian Ziel (1)

1: Duisburg-Essen University, Germany; 2: TU Dortmund, Germany
btissame.el@tu-dortmund.de

SECTION: Computational Statistics and Data Science (Session 2)

We propose a novel adaptive Mixture-of-Experts (MoE) framework designed for time series forecasting. The method incorporates expert-specific loss information directly into the training process to enhance expert specialization, improve regime differentiation, and increase utilization of diverse market signals. We introduce an online learning strategy that incrementally updates both the gating mechanism and expert parameters, thereby substantially reducing computational effort by avoiding repeated full model retraining. By combining expert-level loss integration and efficient online optimization, the proposed approach improves learning efficiency while preserving strong predictive accuracy. Empirical results across diverse economic, tourism, and energy datasets with varying frequencies demonstrate that the proposed approach outperforms both statistical methods and advanced machine learning models, such as Transformer and WaveNet, in terms of forecasting accuracy and computational efficiency.

123 QuBe - Projektion des zukünftigen Arbeitsangebots bis 2045

Michael Kalinowski

Bundesinstitut für Berufsbildung (BIBB), Deutschland
kalinowski@bibb.de

SECTION: DGD (Session 2): Methodische Herausforderungen, Lösungsansätze und aktuelle Entwicklungen in der Regionaldemografie II

Zur fundierten Projektion des zukünftigen Arbeitskräfteangebots nach Qualifikationsniveaus und Berufsfeldern ist eine detaillierte Modellierung des Bildungs- und Migrationsgeschehens notwendig. Im Rahmen der BIBB-IAB-Qualifikations- und Berufsprojektionen (QuBe) wurde hierzu ein Angebotsmodell entwickelt, das auf der Bevölkerungsprojektion des IAB basiert.

Das Modell unterscheidet zwischen Personen im Bildungssystem und jenen außerhalb und bildet Bildungswege, Abschlüsse, demografische Entwicklungen sowie Migration differenziert ab. Das sogenannte „berufliche Bildungssystem“ modelliert die Übergänge von Ausbildungseinrichtungen in den Arbeitsmarkt unter Berücksichtigung von Erfolgsquoten und Übergangswahrscheinlichkeiten. Die Migrationskomponente basiert primär auf Daten des Mikrozensus, wobei Qualifikationsstrukturen von Zu- und Fortzügen sowie Einbürgerungen modelliert werden.

Für die Projektion der Bestände im Bildungssystem erfolgt eine Vorausberechnung der Schülerinnen- und Schülerzahlen sowie der Studierenden. Die daraus resultierenden Absolventenzahlen werden unter Verwendung von Erfolgsquoten ermittelt.

Das Kohorten-Komponenten-Modell erlaubt eine differenzierte Betrachtung der demografischen und qualifikatorischen Entwicklung der Erwerbsbevölkerung, unterteilt nach Deutschen und Nicht-deutschen, Geschlecht, Alter, Beruf und Bildungsabschluss.

Insgesamt ermöglicht das Modell eine Darstellung der jährlichen Zu- und Abgänge am Arbeitsmarkt, einschließlich der Beiträge aus dem Bildungssystem und der Migration und liefert eine empirische Grundlage zur Bewertung zukünftiger Fachkräftepotenziale und unterstützt langfristige Planungen in Arbeitsmarkt-, Bildungs- und Migrationspolitik. Wegen der aktuellen demografischen und migratorischen Entwicklung fokussiert sich der Beitrag auf Entwicklung bis 2045: - Demografie - Qualifikationsspezifische Zu- und Abwanderung - Neuangebot auf dem Arbeitsmarkt und als Gegenstrom ausscheidende Personen aus dem Berufsleben

Diese vorgestellten Größen werden in den nächsten Dekaden den größten Einfluss auf den Arbeitsmarkt und die wirtschaftliche Entwicklung Deutschlands haben.

124 Analyzing the Effects of Demographic Changes on Regional Human Capital Stocks in Germany

Timon Hellwagner (1), Patrizio Vanella (2)

1: Institut für Arbeitsmarkt- und Berufsforschung, Deutschland; 2: aQua-Institut
timon.hellwagner@gmx.at

SECTION: DGD (Session 2): Methodische Herausforderungen, Lösungsansätze und aktuelle Entwicklungen in der Regionaldemografie II

Across countries, population ageing and decline are either unfolding or impending. Such large-scale and far-reaching changes raise justified questions about the accompanying societal and economic effects. While some contributions in the literature focus on the extent of overall labor supply as

the population is ageing and declining, other approaches address, for instance, the implications on the human capital stock. This often happens with a focus on the national level. However, the relevancy of those processes and questions are amplified by stark differences in the regional population structures, implying vastly different changes in regional labor supplies and capital stocks in the past, but also in the short- to medium-term future. From a policy perspective, sound knowledge of underlying drivers of past, and of scope and scale of future regional developments and differences is crucial. Yet, corresponding empirical evidence is often missing. We address this shortcoming and contribute to the literature by constructing regional human capital stocks for German regions and conducting a decomposition analysis to show how demographic changes have altered regional human capital patterns across Germany over the period 2011–2024. Moreover, by linking our decomposition to a deterministic population projection until 2035, we provide insights how these regional differences evolve in the near future.

125 Regionale Strukturen und Trends der Erwerbspersonenentwicklung in Deutschland 2022 bis 2045

Steffen Maretzke, Jana Hoymann, Claus Schlömer

Bundesinstitut für Bau-, Stadt- und Raumforschung im BBR (BBSR), Deutschland
steffen.maretzke@bbr.bund.de

SECTION: DGD (Session 2): Methodische Herausforderungen, Lösungsansätze und aktuelle Entwicklungen in der Regionaldemografie II

Das BBSR erarbeitet regelmäßig regionalisierte Prognosen, die für Politik und Planungen in Bund, Ländern und Kommunen eine wichtige Informationsgrundlage sind. Auf der Grundlage der Ergebnisse der BBSR-Bevölkerungsprognose und regional differenzierter Annahmen zur künftigen Entwicklung des Erwerbsverhaltens der Erwerbsfähigen wird auf der Ebene von Raumordnungsregionen (das Aggregat mehrerer Kreise) die Zahl der Erwerbspersonen für den Zeitraum 2022 bis 2045 vorausberechnet.

Das Besondere an dieser Prognose ist die Tatsache, dass sie für jede Region differenziert aufzeigt, wie die Wanderungen, die Altersstruktur der Erwerbsfähigen und deren Erwerbsverhalten die Entwicklung des Arbeitskräfteangebotes prägen. Der Vortrag wird zeigen, dass ohne Wanderungsgewinne bis 2045 keine Region Deutschlands mehr einen Anstieg des Arbeitskräfteangebotes zu erwarten hätte. Dass sich die Zahl der Erwerbspersonen in den besonders wachstumsstarken Regionen überdurchschnittlich stark erhöhen wird, resultiert dort vor allem aus den Wanderungsgewinnen wie aus der steigenden Erwerbsbeteiligung der Erwerbsfähigen. In Berlin ergibt sich diese positive Entwicklung zudem aus vorteilhaften Verschiebungen der Alters- und Geschlechterstruktur der Erwerbsfähigen zugunsten von Jahrgängen/Personen mit einer höheren Erwerbsbeteiligung. In den Regionen mit dem stärksten Rückgang des Arbeitskräfteangebots ist vor allem der unzureichende Ersatz der aus dem Erwerbsleben ausscheidenden durch jüngere Personen Erwerbspersonen für diese Entwicklung verantwortlich. Im Gegensatz zu den besonders wachstumsstarken Regionen können diese Verluste in den strukturschwächeren Regionen durch die anderen Einflussfaktoren nicht mehr kompensiert werden.

Ausgewählte Forschungsergebnisse dieser Prognose werden im Rahmen des angebotenen Vortrages präsentiert, inklusive erster Schlussfolgerungen für Politik und Praxis.

Literatur

Maretzke, S.; Hoymann, J.; Schlömer, C., 2024: Raumordnungsprognose 2045. Bevölkerungsprognose – aktualisiert anhand der Ergebnisse des Zensus 2022. Herausgeber: Bundesinstitut für Bau-, Stadt- und Raumforschung (BBSR). BBSR-Analysen KOMPAKT 13/2024. Bonn.

<https://doi.org/10.58007/q848-k680>

Maretske, Steffen; Hoymann, Jana; Schlömer, Claus, 2026: Raumordnungsprognose 2045. Erwerbspersonenprognose 2045. Herausgeber: Bundesinstitut für Bau-, Stadt- und Raumforschung (BBSR). BBSR-Analysen KOMPAKT02/2026. Bonn.
<https://doi.org/10.58007/hgvp-ck89>

126 Cost-Error Trade-Offs in Introducing Web in an Ongoing Telephone Labor Market Survey

Joseph Sakshaug (1), Jan Mackeben (2)

1: IAB & LMU-Munich, Germany; 2: Leibniz Institute for Educational Trajectories (LifBi)
joesaks@umich.edu

SECTION: Methodology of Statistical Surveys (Session 2)

Telephone surveys have historically been a popular form of data collection in labor market research and continue to be used to this day. Yet, telephone surveys are confronted with many challenges, including imperfect coverage of the target population, low response rates, risk of nonresponse bias, and rising data collection costs. To address these challenges, many telephone surveys have shifted to online and mixed-mode data collection to reduce costs and minimize the risk of coverage and nonresponse biases. However, empirical evaluations of the intended effects of introducing online and mixed-mode data collection in ongoing telephone surveys are lacking. We address this research gap by analyzing a telephone employee survey in Germany, the Linked Personnel Panel (LPP), which experimentally introduced a sequential web-to-telephone mixed-mode design in the refreshment samples of the 4th and 5th waves of the panel. By utilizing administrative data available for the sampled individuals with and without known telephone numbers, we estimate the before-and-after effects of introducing the web mode on coverage and nonresponse rates and biases. We show that the LPP was affected by known telephone number coverage bias for various employee subgroups prior to introducing the web mode, though many of these biases were partially offset by nonresponse bias. Introducing the web-to-telephone design improved the response rate but increased total selection bias, on average, compared to the standard telephone single-mode design. This result was driven by larger nonresponse bias in the web-to-telephone design and partial offsetting of coverage and nonresponse biases in the telephone single-mode design. Significant cost savings (up to 50% per respondent) were evident in the web-to-telephone design.

127 Participation and Selectivity in a Randomised Mixed-Mode Panel Design: Evidence from NEPS Starting Cohort 4

Michael Bergrab, Christian Assmann, Timo Gnambs

Leibniz-Institut für Bildungsverläufe, Deutschland
michael.bergrab@lifbi.de

SECTION: Methodology of Statistical Surveys (Session 2)

In langfristigen Panelstudien stellen Attrition und selektive Nichtteilnahme zentrale Herausforderungen für Inferenz, Vergleichbarkeit über Wellen und die effiziente Organisation von Feldarbeit dar. Dies gilt besonders, wenn etablierte Panels von klassischen intervieweradministrierten Erhebungsformen auf mixed-mode oder stärker selbstadministrierte Designs umgestellt werden. Der

Beitrag untersucht diese Frage anhand von Welle 17 der National Educational Panel Study, Starting Cohort 4. In dieser Welle wurden panelbereite Zielpersonen randomisiert einem von drei Erhebungsmodi zugewiesen: CAPI, CALVI oder webbasierter Kompetenztestung mit anschließendem CATI-Interview.

Wir analysieren Teilnahme als mehrstufigen Feldprozess. Die zentralen Outcomes sind gültige Kompetenztestdaten, der Start des Hauptinterviews und die Teilnahme an einer nachgelagerten CAWI-Befragung. Da die CAWI-Befragung nur für Personen freigeschaltet wurde, die das Hauptinterview begonnen hatten, wird der Teilnahmeprozess nicht als vollständig beobachteter trivariater Response-Vektor behandelt. Stattdessen verwenden wir eine zweistufige Modellstrategie: Kompetenztestung und Interviewstart werden gemeinsam mit einem bivariaten Probit-Modell analysiert, das korrelierte unbeobachtete Teilnahmebereitschaften zulässt. Die nachgelagerte CAWI-Teilnahme wird bedingt unter Interviewstartern modelliert.

Der Beitrag schätzt Effekte der randomisierten Moduszuweisung auf Test- und Interviewteilnahme, beschreibt die kumulative Selektivität entlang des Feldprozesses und untersucht, welche Respondentengruppen durch welche Modusbedingung erreicht oder verloren werden. Ergänzend werden realisierte Modi, Moduswechsel und Sensitivitätsanalysen betrachtet. Der Beitrag zeigt, wie randomisierte Mixed-Mode-Designs genutzt werden können, um nicht nur Response-Raten, sondern auch stage-specific selection und die Zusammensetzung realisierter Analysetichproben statistisch zu bewerten.

128 You've got Mail: Does sending Thank-you postcards increase response in a probability-based online panel?

Barbara Felderer (1), Anne Balz (2), **Julian Diefenbacher** (1), Jannis Kück (3), Aline Mohr (3), Tobias Rettig (2), Martin Spindler (4)

1: GESIS, Germany; 2: University of Mannheim, Germany; 3: University of Düsseldorf, Germany; 4: University of Hamburg, Germany
barbara.felderer@gesis.org

SECTION: Methodology of Statistical Surveys (Session 2)

Survey nonresponse is one of the major challenges to survey data quality. While various treatments, e.g., monetary incentives, have been tested to increase response rates, the evidence regarding differential effects of the treatments for different population groups is rather thin. For example, it is known that incentives work well overall, but it is unclear whether a less costly form of appreciation (like a postcard) would achieve the same or a greater effect for certain population groups or whether some groups do not need any treatment at all.

An experiment to increase survey response was fielded in mid-December 2025 among the newly recruited participants of the German Internet Panel (GIP). After the first panel wave, before the second panel wave, panelists were randomly assigned to four experimental groups:

- 1.) receiving a "Thank-you" postcard from the GIP team
- 2.) receiving a handwritten "Thank-you" postcard with the same text as 1.)
- 3.) receiving a postcard that states that they have been credited an extra of 5 Euro as a "Thank-you" for being in the panel
- 4.) control group

Based on the results of the random experiment, we derived an optimal treatment allocation rule via policy trees that maximizes expected response. The learned policy was then evaluated in a second experiment in the subsequent panel wave where a random half of the former control group was

assigned to treatments according to the learned policy and the other half was randomly assigned to treatments.

The presentation focuses on both the effects of the treatments (and their interactions with respondents' personal characteristics) on nonresponse in wave 2 from the first experiment as well as the evaluation of the optimal policy fielded in the second experiment.

129 Was Kommunalbefragungen zur Umfrageforschung beitragen: Methodische Befunde aus der BBSR-Kommunalbefragung „Kommunen im Wandel“

Daniel Meyer, Elisabeth Stürmer, Antonia Milbert

Bundesinstitut für Bau-, Stadt- und Raumforschung, Deutschland
daniel.meyer@bbr.bund.de

SECTION: Methodology of Statistical Surveys (Session 2)

Ob Mikrozensus, SOEP oder ALLBUS – Bevölkerungsbefragungen nehmen in der empirischen Sozialforschung eine zentrale Stellung ein. Entsprechend beruht ein Großteil des umfragemethodischen Wissens auf Befunden aus Personen- und Haushaltsbefragungen. Demgegenüber werden Organisationsbefragungen seltener durchgeführt und methodisch weniger umfassend reflektiert. Unser Beitrag setzt hier an und zeigt, welche methodischen Einsichten Kommunalbefragungen als spezifischer Typ von Organisationsbefragungen bereithalten. Als empirisches Anschauungsmaterial dient die BBSR-Kommunalbefragung „Kommunen im Wandel“, eine zwischen April und Juli 2024 durchgeführte Online-Vollerhebung aller 651 Städte und Gemeinden in den deutschen Kohleregionen. Der Beitrag untersucht zwei umfragemethodische Aspekte: Erstens analysieren wir anhand eines Split-Ballot-Experiments, wie Variationen der Skalenrichtung und Itemreihenfolge das Antwortverhalten der befragten Funktionsträgerinnen und Funktionsträger beeinflussen und inwieweit sich hierbei Unterschiede zwischen den adressierten Bürgermeisterinnen und Bürgermeistern und anderen Verwaltungsmitarbeitenden zeigen, an die die Beantwortung der Umfrage delegiert wurde. Zweitens untersuchen wir im Rahmen einer Non-Response-Analyse, wie die Teilnahmebereitschaft, der Teilnahmezeitpunkt und die Bearbeitungsdauer von den Ansprachen im Feldverlauf und von kommunalen Strukturmerkmalen – darunter die Gemeindegröße und die Zugehörigkeit zu einem Gemeindeverband – abhängen. Durch den Vergleich mit etablierten Befunden aus Personen- und Haushaltsbefragungen arbeiten wir methodische Besonderheiten kommunaler Organisationsbefragungen heraus. Abschließend diskutieren wir die Implikationen unserer Ergebnisse für zukünftige Kommunalbefragungen und leiten daraus konkrete Empfehlungen für die Fragebogenentwicklung, Feldsteuerung und Datenauswertung ab.

130 Spectral Mean Estimators of Time Series: Finite Sample Distributional Approximations

Andreas Anastasiou (1), Tobias Kley (2), Efstathios Paparoditis (1)

1: University of Cyprus, Cyprus; 2: University of Göttingen, Germany
tobias.kley@uni-goettingen.de

SECTION: Statistical Theory and Methods (Session 2)

We consider general spectral means, based on a finite signed spectral Borel measure, and derive finite sample bounds for the Wasserstein distance of the distributions of the corresponding sample estimators (integrated periodograms) to their Gaussian limits. Our results also apply to absolutely continuous and discretized versions of spectral means, based on a function φ . The bounds obtained are given explicitly, they are computable and are applicable under non-restrictive conditions regarding Fourier coefficients, and for a broad class of short-range dependent processes. They imply (close to) optimal rates for sufficiently smooth functions φ . To illustrate the versatility of our results we derive corresponding bounds for nonparametric spectral density estimators, where the function φ is allowed to vary with n . In the context of the latter statistic, existing results are extended and bounds for lag- and spectral-window density estimators are obtained, which are novel and interesting in themselves.

131 Multivariate Local Whittle Estimation for Multivariate Functional Time Series

Miao Yu, Philipp Sibbertsen, Yeliz Özer

Leibniz University Hannover, Germany
miao.yu@statistik.uni-hannover.de

SECTION: Statistical Theory and Methods (Session 2)

In this paper we establish consistency and asymptotic normality of the multivariate local Whittle estimator for multivariate functional curved time series. The proposed semi-parametric estimator is constructed by applying the local Whittle estimator method to periodograms formed from multivariate functional principal component scores, where each score is obtained as the inner product between the functional observation and an estimated leading multivariate eigenfunction. The methodology is examined using a Monte Carlo study based on multivariate functional long memory processes generated via a Davies-Harte spectral construction with dimension specific fractional parameters. The simulation results indicate that the proposed multivariate local Whittle estimator exhibits low bias and mean squared error in finite samples. In empirical application to high-frequency wind energy data, daily functional representations of wind speed, wind direction, and power production are constructed and analyzed using multivariate functional principal component analysis(MFPCA). The results show that the long memory behavior of the dominant latent dynamic is primarily driven by power production, as evidenced by the loading structure and the persistence properties of the leading principal component scores.

132 Testing the order of fractional integration when smooth deterministic trends are possibly present

Mustafa Kilinc (1), Michael Massmann (2,3)

1: University of Cologne, Germany; 2: WHU – Otto Beisheim School of Management; 3: Vrije Universiteit
m.kilinc@wiso.uni-koeln.de

SECTION: Statistical Theory and Methods (Session 2)

This paper introduces a test for fractional integration in a model that possibly contains smooth deterministic trends. We model the trend component using a Chebyshev polynomial and specify the short-run dynamics semi-parametrically, accommodating a broad class of possibly nonlinear processes, including those with conditional heteroskedasticity. We use a local Whittle approach for constructing a Lagrange multiplier test statistic and for constructing a frequency-domain information criterion for the selection of the order of the Chebyshev polynomial. We show that widely used time-domain information criteria are generally inconsistent for the true order, whereas our frequency-domain criterion remains robust under both short- and long-memory behaviour. Monte Carlo simulations and an empirical application to the UK Great Ratios support our theoretical findings.

133 Georeferenzierte Standortindikatoren: Kleinräumige Lageanalysen für kommunale Fragestellungen

Georg Wiegleb

Landeshauptstadt Magdeburg, Deutschland
georg.wiegleb@stadt.magdeburg.de

SECTION: VDSSt (Session 3): Bauen, Wohnen, Verkehr - Angewandte Geodatenanalysen

Viele kommunale Fragestellungen in den Bereichen Wohnen, Versorgung, Umwelt und Verkehr lassen sich nur begrenzt über Stadtteile oder Quartiere beantworten. Erreichbarkeit, Grün- und Freiraumversorgung, Umweltbelastungen oder Barrieren variieren innerhalb derselben Gebietseinheit häufig deutlich und hängen von der konkreten Lage einzelner Standorte ab.

Die dafür benötigten Informationen liegen in der Praxis selten direkt auf Adressniveau vor. Sie verteilen sich über sehr unterschiedliche räumliche Datentypen – Gebietsdaten, Rasterdaten, Punkt- und Netzwerkdaten – mit jeweils eigenen Eigenschaften und Auflösungen. Die methodische Aufgabe besteht darin, diese heterogenen Quellen auf eine gemeinsame, feinräumige Standort-Resolution zu bringen, sodass jeder Adresse ein konsistentes Bündel von Standortmerkmalen zugeordnet werden kann.

Der Vortrag zeigt, wie sich solche georeferenzierten Standortindikatoren in der Praxis aufbauen lassen. Je nach Datentyp und Fragestellung kommen unterschiedliche Verfahren zum Einsatz – von räumlichen Überlagerungen und Pufferanalysen über Dichte-, Glättungs- und Interpolationsverfahren bis zu netzwerkbasierter Erreichbarkeiten.

Auf Basis solcher Indikatorfelder lassen sich Standortprofile bilden und ähnliche Lagetypen identifizieren. So kann untersucht werden, welche Standortmerkmale gemeinsam auftreten und wo sich räumliche Häufungen ähnlicher Standortbedingungen zeigen.

134 Bautätigkeitsstatistik im Wandel: Datenqualität, Digitalisierung und Auswertungen nach Gemeindegrößenklassen auf Basis des tatsächlichen Ereignisdatums

Alexander Weigert, Kerstin Lange

Statistisches Bundesamt, Deutschland
alexander.weigert@destatis.de

SECTION: VDSt (Session 3): Bauen, Wohnen, Verkehr - Angewandte Geodatenanalysen

Die Bautätigkeitsstatistik verzeichnete zuletzt steigende Nachmeldequoten, welche die Aussagekraft der monatlichen Baugenehmigungs- und der jährlichen Baufertigstellungsstatistik einschränken. Qualitätsauswertungen belegen den Handlungsdruck zur Verbesserung der Gesamtprozesse. Neben der noch laufenden Implementierung digitalisierter Ende-zu-Ende Prozesse wird die bisherige Betrachtung nach dem Berichtszeitraum um eine Auswertung nach dem tatsächlichen Ereignisdatum ergänzt. Dies erhöht die Aussagekraft der konjunkturellen Analysen. Dieser Vortrag erläutert die Erhebungs- und Verarbeitungsprozesse der Bautätigkeitsstatistiken, gibt einen Einblick in die Digitalisierung der Statistik innerhalb der föderalen digitalen Bauantragsportale und zeigt erste konkrete Analysen des Wohnungsbaus in Deutschlands nach Gemeindegrößenklassen auf Basis des Ereigniszeitraumkonzepts.

135 Bautätigkeitsstatistik im Wandel: Neues Auswertungskonzept mit methodischen Schätzungen am aktuellen Rand

Kerstin Lange, Alexander Weigert

Statistisches Bundesamt, Deutschland
kerstin.lange@destatis.de

SECTION: VDSt (Session 3): Bauen, Wohnen, Verkehr - Angewandte Geodatenanalysen

Die Bautätigkeitsstatistik befindet sich im Wandel: Aufbauend auf den identifizierten Handlungsbedarfen und den aufgezeigten Digitalisierungspotenzialen wurden erste konkrete Schritte umgesetzt, um die Datenqualität zu verbessern. Im Mittelpunkt des Vortrags steht die Umstellung des Auswertungskonzepts von der bisherigen Betrachtung nach Berichtszeitraum auf eine Auswertung nach Ereigniszeitraum mit korrekter Periodenzuordnung der (nach-)gemeldeten Baugenehmigungen. Dadurch erfolgt die statistische Abbildung aktueller Bautätigkeiten präziser und zeitnäher. Dieser Beitrag analysiert die verspäteten Meldungseingänge der Baugenehmigungsstatistik, erörtert die methodische Entwicklung eines Schätzmodells und die Einführung eines neuen Revisionszyklus zur Verbesserung der Datenqualität; insbesondere am aktuellen Rand.

136 Schätzung kleinräumiger Armutsgefährdung in Deutschland mit GIS

Stefan Kaup, Silas Eichfuss, **Lizanne Hauser**

Esri Deutschland GmbH
s.kaup@esri.de

SECTION: VSt (Session 3): Bauen, Wohnen, Verkehr - Angewandte Geodatenanalysen

Armutsgefährdete Menschen in Deutschland haben nach Definition der Statistischen Ämter ein Einkommen von unterhalb 60% des Medians des Äquivalenzeinkommens der Gesamtbevölkerung. Die Sozialberichterstattung stellt diese Informationen für 1% der Bevölkerung in der Erhebung MZ-Kern dar. Die so erhobenen Daten lassen aggregierte Aussagen für regionale Quoten der Armutsgefährdung bis zu den Bundesländern und Regierungsbezirken zu. Für Darstellung auf einer tieferen regionalen Gliederung wie der Ebene der Kreise und kreisfreien Städte ist ein kleinräumiges Schätzverfahren notwendig.

Der Beitrag soll einen Ansatz sowie die Ergebnisse eines solchen Verfahren darstellen. Er greift dabei eine verwendete Methodik des Projektes TiPSE aus der ESPON Programmperiode 2013 auf. Das Verfahren simuliert die Verteilung der Armutsgefährdungsquote (Anteil der Personen in einer Region oder Gruppe mit einem als armutsgefährdend definierten Einkommen) auf der Zielebene über ein Regressionsverfahren auf Basis von Kovariablen auf 2 Ebenen. Die Datenquellen des Zensus 2011 sowie des EU-SILC 2012 werden dafür durch den aktuellen Stand des Zensus 2022 und Mikrozensus 2022 ersetzt.

Da hierbei der Raumbegriff und räumliche Abhängigkeiten eine Rolle spielen, zielt der Beitrag darauf ab, die Datenverarbeitung, Simulation und Darstellung ganzheitlich in einem GIS-System durchzuführen. Die Verwendung einer cloudbasierten zeitgemäßen GIS-Infrastruktur macht die Ergebnisse reproduzierbar. Mit der Nutzung moderner Analysen und Darstellungsweisen, entsteht ein praxisnahes Werkzeug zur evidenzbasierten Sozialpolitik, das analytische Tiefe, visuelle Verständlichkeit und operative Umsetzbarkeit verbindet.

137 Wie lassen sich forschende Unternehmen typisieren? Eine multivariate Analyse der FuE-Erhebung

Verena Leyendecker

SV Wissenschaftsstatistik, Deutschland
verena.leyendecker@stifterverband.de

SECTION: Economic, Social and Market Statistics

Im Jahr 2023 investierten deutsche Unternehmen über 90 Milliarden Euro in Forschung und Entwicklung (FuE). Ein wesentlicher Teil dieser Aufwendungen entfällt auf wenige große Unternehmen in technologieintensiven Kernbranchen, insbesondere der Automobilindustrie. Entsprechend konzentrieren sich viele Auswertungen der FuE-Erhebung auf diese forschungsintensiven Wirtschaftszweige und die Entwicklung ihrer FuE-Ausgaben.

Gleichzeitig ist die Grundgesamtheit der forschenden Unternehmen deutlich heterogener strukturiert: Kleine und mittlere Unternehmen (KMU) stellen rund 86 Prozent der in der FuE-Erhebung erfassten Unternehmen. Trotz dieser quantitativen Bedeutung ist weniger untersucht, durch welche weiteren strukturellen und organisatorischen Merkmale sich forschende Unternehmen charakterisieren lassen.

Insbesondere bleibt offen, inwiefern sich innerhalb der Kategorien Branche und Unternehmensgröße zusätzliche Unterschiede zwischen den Unternehmen zeigen.

Vor diesem Hintergrund analysiert der Beitrag auf Basis der FuE-Erhebung, welche Merkmale forschende Unternehmen über Branche und Unternehmensgröße hinaus typisieren. Hierzu werden ausgewählte Indikatoren der Unternehmens- und FuE-Struktur herangezogen, um mithilfe einer Clusteranalyse unterschiedliche Typen forschender Unternehmen zu identifizieren.

138 Forschung und Entwicklung in der Pandemie – Merkmale krisenresilienter Unternehmen

Jessica Ernst, Johannes Schmitt

SV gemeinnützige Gesellschaft für Wirtschaftsstatistik mbH, Deutschland
jessica.ernst@stifterverband.de

SECTION: Economic, Social and Market Statistics

Die Covid-19-Pandemie hat die Rahmenbedingungen für unternehmerische Forschung und Entwicklung (FuE) deutlich verändert. Während Eindämmungsmaßnahmen zu Unterbrechungen und Anpassungen in laufenden Prozessen führten, entwickelten sich gleichzeitig neue Forschungsbedarfe, etwa im pharmazeutischen Bereich. Insgesamt gingen die FuE-Aufwendungen in der Wirtschaft zunächst deutlich zurück, wobei sich einige Unternehmen in der Krise hinsichtlich ihrer FuE-Tätigkeit als resilient erwiesen oder ihre Aktivitäten sogar ausweiteten.

Der Beitrag analysiert ein Panel von FuE-tätigen Unternehmen ausgewählter Branchen im Zeitraum vor, während und nach der Pandemie. Ziel ist es, Erklärungsmuster für rückläufige, gleichbleibende oder steigende FuE-Tätigkeit während der Krise zu identifizieren. Ergänzend werden weitere Quellen wie Geschäftsberichte herangezogen, um zusätzliche Erklärungsmerkmale zu bestimmen, die anschließend mittels Qualitative Comparative Analysis (QCA) untersucht werden.

139 Fördert Demokratie die Innovationsfähigkeit eines Landes?

Andreas Kladroba

Stifterverband für die Deutsche Wissenschaft, Deutschland
andreas.kladroba@stifterverband.de

SECTION: Economic, Social and Market Statistics

Fragt man nach den innovativsten Ländern der Erde, fallen vielen sicher die Mitgliedsstaaten der Europäischen Union, die USA und Kanada, aber auch Japan, Südkorea oder vielleicht Australien ein. Dabei handelt es sich nicht nur um Länder mit hoher Wirtschaftskraft, sondern auch mit einer demokratischen Verfassung. Aber auch ein Land wie China mit einer eher autokratischen Staatsführung gehört inzwischen zu den innovativsten Volkswirtschaften der Welt. Welcher Zusammenhang zwischen der Staatsform und der Innovationsfähigkeit eines Landes ist also zu erkennen? Dieser Frage will der Vortrag mit Hilfe von Paneldaten zu Forschung und Entwicklung der OECD und der EU sowie dem Demokratieindex der Zeitschrift "The Economist" nachgehen.

140 Erweiterungsmöglichkeiten der Unternehmensstrukturstatistik zur Analyse der Qualifikationsnachfrage von Unternehmen

Tobias Maier (1), Dorothea Krabbe (1), Peter Dreuw (2)

1: Bundesinstitut für Berufsbildung, Deutschland; 2: Gesellschaft für Wirtschaftliche Strukturforschung
tobias.maier@bibb.de

SECTION: Economic, Social and Market Statistics

In den strukturellen Unternehmensstatistiken werden Struktur, Wirtschaftstätigkeiten und Leistungsfähigkeit von Unternehmen im Zeitverlauf abgebildet. Sie zeichnen sich durch einen detaillierten Erfassungsgrad hinsichtlich Größenklassen und Wirtschaftstätigkeiten aus und stehen auf nationaler wie europäischer Ebene zur Verfügung. Während Investitionen und Erträge gut darstellbar sind, bleiben Fragen zum Humankapital bislang weitgehend unberücksichtigt. Der Beitrag diskutiert deshalb Erweiterungsmöglichkeiten der Statistik zur Analyse von Qualifikationsnachfrage und Humankapitalstrukturen.

Grundlage der Unternehmensstrukturstatistik ist seit 2018 die Umsetzung des EU-Unternehmensbegriffs gemäß EWG Nr. 696/93. Rechtliche Einheiten werden mittels Imputation und Konsolidierung zu komplexen Unternehmen zusammengeführt, definiert als „kleinste Kombination rechtlicher Einheiten“, die eine organisatorische Einheit zur Produktion von Waren und Dienstleistungen bildet. Die Unternehmensgrößenklassen folgen der Empfehlung 2003/361 der EU-Kommission und kombinieren Beschäftigtenzahl und Umsatz. Sowohl die Unternehmensform als auch die Größenklasse sind in anderen Statistiken mit Informationen zur Qualifikationsstruktur nicht deckungsgleich abbildbar.

So basiert die Beschäftigungsstatistik der Bundesagentur für Arbeit auf sozialversicherungspflichtigen Meldungen. Ein Betrieb stellt dabei eine regional und wirtschaftsfachlich abgegrenzte Einheit mit einer Betriebsnummer dar, in der mindestens ein sozialversicherungspflichtiges oder geringfügiges Beschäftigungsverhältnis besteht. Der Mikrozensus erhebt dagegen Selbsteinschätzungen von Haushaltsmitgliedern zu Wirtschaftszweig und Betriebsgröße ihrer eigenen Tätigkeit bzw. des Arbeitgebers. In beiden offiziellen Statistiken lassen sich somit Wirtschaftszweig und Betriebsgröße ermitteln, allerdings ist bislang unklar, wie diese Einschätzungen zur Darstellung in der Unternehmensstrukturstatistik stehen.

Der Beitrag analysiert Abweichungen der Unternehmenszahlen nach Betriebsgröße und WZ08-Abteilungen zwischen Beschäftigtenhistorik und Mikrozensus, diskutiert deren Ursachen und bewertet die Konsequenzen für Daten-Matching und Analysen der Qualifikationsnachfrage. Konkret gehen wir der Frage nach, wie das Investitionsverhalten der Unternehmen mit der Veränderung der Anforderungsstruktur (Mikrozensus) und dem Ausbildungsverhalten der Betriebe (Beschäftigtenhistorik) zusammenhängt.

141 Von vorläufigen zu finalen Zahlen: Ein modellbasierter Ansatz zur Vervollständigung von Außenhandelsdaten

Heiko Limberg

Statistisches Bundesamt, Deutschland
heiko.limberg@destatis.de

SECTION: Computational Statistics and Data Science (Session 3)

Das Statistische Bundesamt veröffentlicht unter anderem monatliche Außenhandelsdaten zu Importen und Exporten mit Partnerländern weltweit. Diese Daten werden auf einer hohen Detailtiefe bereitgestellt und umfassen sowohl die Warenklassifikation als auch das jeweilige Partnerland. Die Überprüfung der Datenqualität auf dieser Ebene ist ein kontinuierlicher Prozess, um verlässliche Statistiken für Nutzende bereitzustellen. Eine zentrale Herausforderung besteht darin, teilweise vorliegende Werte innerhalb eines laufenden Monats zu validieren, obwohl weder der endgültige Gesamtwert noch der genaue Zeitpunkt und Umfang der eingehenden Meldungen mit Sicherheit bekannt sind.

Zur Bewältigung dieser Herausforderung wird die Entwicklung der kumulierten Meldungen mithilfe sogenannter Completion Curves modelliert, welche den zeitlichen Verlauf des Anteils bereits eingegangener Werte am späteren Gesamtwert beschreiben.

Im Rahmen eines Bayes'schen Ansatzes werden die partiellen Summen als Mischung historischer "Completion Curves" modelliert. Dadurch wird sowohl die Unsicherheit hinsichtlich der Form als auch der Dauer des Meldeprozesses berücksichtigt. Historische Verläufe werden anhand charakteristischer Eigenschaften, wie Meldegeschwindigkeit und Verteilungsmuster, gruppiert und zur Konstruktion informativer Prior-Verteilungen genutzt.

Dieser Ansatz ermöglicht es, den erwarteten täglichen Fortschritt der Datenvollständigkeit zu schätzen sowie den zu erwartenden Gesamtwert dynamisch anzupassen. Somit wird eine Grundlage geschaffen, um die Datenvollständigkeit besser zu bewerten und Auffälligkeiten frühzeitig zu erkennen.

142 Design und Evaluierung synthetisch generierter Mikrozensus-Campus Files für den Einsatz in der Hochschullehre

Amelie Plöger, Colin Schmidt, Sarah Redlich

Landesbetrieb Information und Technik NRW - Statistisches Landesamt
Amelie.Ploeger@it.nrw.de

SECTION: Computational Statistics and Data Science (Session 3)

Die Campus Files der Forschungsdatenzentren der amtlichen Statistik dienen in der Hochschullehre zur praxisnahen statistischen und sozialwissenschaftlichen Methodenausbildung von Studierenden. Die zur Erstellung von Campus Files verwendeten Anonymisierungskonzepte (Substichprobenziehung, Merkmalsausschluss, Merkmalsvergrößerung, Entfernung regionaler Ebenen) schränken die Einsetzbarkeit zur Methodenausbildung jedoch ein.

Synthetisch erstellte Datensätze bieten die Möglichkeit, wesentliche Strukturen der Originaldaten abzubilden, ohne den Informationsverlust herkömmlicher Anonymisierungsverfahren. Gleichzeitig wird die Anonymität einzelner Personen und Haushalte sichergestellt.

Der Beitrag befasst sich mit dem methodischen Vorgehen zur Erstellung von synthetischen Daten als Mikrozensus-Campus File. Unter Einsatz von binären Entscheidungsbäumen werden die Merkmale des Mikrozensus sequentiell modelliert und somit ein vollständig synthetischer Datensatz generiert. Als Trainingsdaten dienen faktisch anonymisierte Scientific-Use-Files. Im Vortrag werden außerdem die bisher aufgetretenen Herausforderungen und deren Umgang mit diesen dargestellt. Darüber hinaus wird mittels verschiedener statistischer Verfahren die Nähe der synthetischen Daten zu den Ausgangsdaten quantifiziert und das vermeintliche Aufdeckungsrisiko thematisiert. Abschließend werden die Möglichkeiten des Einsatzes in der Hochschullehre aufgezeigt.

143 Vom Prompt zum Prozessfluss: Reproduzierbare Statistik und KI mit SAS Viya, Open Source und KI-Assistenz

Ulrich Reincke

SAS Institute, Deutschland
Ulrich.Reincke@ger.sas.com

SECTION: Computational Statistics and Data Science (Session 3)

Statistische Analyseprozesse stehen zunehmend vor der Herausforderung, methodische Flexibilität, Reproduzierbarkeit, Governance und schnelle Umsetzung miteinander zu verbinden. Der Beitrag zeigt, wie SAS Viya als integrierte Analyseplattform klassische statistische Verfahren, moderne Data-Science-Workflows und Open-Source-Komponenten in einem nachvollziehbaren Prozessmodell zusammenführt.

Im Mittelpunkt stehen die neuen Möglichkeiten das auf GitHub bereitgestellte neue Prompt-Coding-Prozessfluss-Framework. Dieses Framework für assistierte Entwicklung, Analyse und Dokumentation kann lokal als Thin Client und Frontend zur Viya-Plattform genutzt werden und unterstützt eine Arbeitsweise vom Prompt über den generierten und geprüften Code bis zum fertig dokumentierten Prozessfluss. Open-Source-Integration wird im SAS Programmfluss über PROC PYTHON und PROC R ermöglicht.

Skizziert wird wie Analyseideen schneller operationalisiert, Implementierungsschritte transparenter dokumentiert und Endergebnisse besser nachvollziehbar gemacht werden können. Damit adressiert der Beitrag zwei Perspektiven zugleich: Für „Computational Statistics and Data Science“ steht die technische Integration von SAS, Python, R und KI-Assistenz im Vordergrund; für „Kompetenzentwicklung: Data Literacy und Statistik“ zeigt der Ansatz, wie statistisches Arbeiten durch erklärbare Prozessflüsse, Dokumentation und Prompt-Kompetenz besser vermittelbar wird.

144 Textklassifikation mit Machine Learning im ESS: Einblicke aus dem AIML4OS-Projekt

Ariane Lestrade

Statistisches Bundesamt, Deutschland
ariane.lestrade@destatis.de

SECTION: Methodology of Statistical Surveys (Session 3)

Die jüngsten Fortschritte in der Verarbeitung natürlicher Sprache und im maschinellen Lernen treiben zunehmend die Automatisierung zentraler Textklassifikationsaufgaben in der amtlichen Statistik voran, insbesondere bei der Kodierung von Freitexten in standardisierte Klassifikationen

wie Wirtschaftszweige, Berufe und Haushaltsausgaben (z. B. entsprechend NACE, ISCO, COICOP), und dies über zahlreiche nationale statistische Ämter hinweg. Im Rahmen des europäischen Projekts AIML4OS vereint das Arbeitspaket 10 „Text to Code“ (WP10) 13 nationale statistische Ämter (NSIs), um die Zusammenarbeit in diesen Themenbereichen zu stärken.

WP10 nutzte die besondere Gelegenheit, dass zahlreiche NSIs am selben Themenfeld arbeiten, und kombinierte Projektpräsentationen der beteiligten Institute mit einer systematischen Literaturrecherche, um einen umfassenden Überblick über interne Ansätze sowie aktuelle methodische Entwicklungen zu gewinnen und gemeinsame Herausforderungen organisationsübergreifend sichtbar zu machen. Dies führte zur Einrichtung von fünf thematischen Clustern innerhalb von WP10, die jeweils eine spezifische Herausforderung beim Einsatz von maschinellem Lernen für die Textklassifikation in der amtlichen Statistik adressieren: die Generierung synthetischer Daten zur Überwindung begrenzter Trainingsdaten, der Einsatz fortgeschrittener NLP-Modelle zur Bewältigung sprachlicher Komplexität, die Entwicklung hierarchischer Klassifikationsmethoden, der Aufbau eines MLOps-Frameworks für den Text-zu-Code-Kontext sowie die Anpassung von Modellen an sich weiterentwickelnde Klassifikationsstandards wie beispielsweise in dem Fall der NACE-Revision.

Ziel der WP10-Cluster ist es, die identifizierten methodischen Ansätze systematisch zu erproben und die daraus resultierenden Ergebnisse sowie den entwickelten Code offen bereitzustellen, um eine organisationsübergreifende, praxisnahe und breit zugängliche Open-Source-Sammlung von Lösungsansätzen zu schaffen, die NSIs als Ausgangspunkt und Orientierung für eigene Anwendungen dient.

145 Einsatz von LLMs zur Klassifikation der Wirtschaftszweige

Susanne Wegner

Statistisches Bundesamt, Deutschland
susanne.wegner@destatis.de

SECTION: Methodology of Statistical Surveys (Session 3)

Die Zuordnung von Freitextangaben zu einer standardisierten Wirtschaftszweigssystematik ist ein entscheidender Schritt für die einheitliche und vergleich-bare statistische Erfassung wirtschaftlicher Tätigkeiten in verschiedenen Statistiken. Traditionell wird diese Zuordnung manuell durchgeführt, wobei in den letzten Jahren zunehmend maschinelle Lernverfahren zur Unterstützung eingesetzt werden. Transformer-basierte Modelle, insbesondere solche, die auf großen Sprachmodellen (LLMs) basieren, eröffnen neue Möglichkeiten und bieten erhebliches Potenzial. Es wird aufgezeigt, wie LLMs spezifisch für die WZ-Klassifikation genutzt werden. Ein Anwendungsbereich ist die Datenvorverarbeitung, bei der LLMs zur Textzusammenfassung, Datenbereinigung oder zur Erstellung synthetischer Daten eingesetzt werden. Letzteres dient besonders dann zur Verbesserung von überwachten Lernverfahren, wenn nicht genügend Trainingsdaten zur Verfügung stehen, beispielsweise beim Umstieg von der Klassifikationsversion WZ08 auf die WZ25. Ein weiterer Schwerpunkt liegt auf der Klassifikation selbst. Hier werden Experimente mit feinabgestimmten Transformern und einem Retrieval- und Reranking-Ansatz präsentiert. Ein Vergleich dieser Methoden mit traditionellen maschinellen Lernverfahren zeigt das Potential dieser Verfahren auf.

146 Maschinelle Lernverfahren als Unterstützung für die Wirtschaftszweigschlüssel-Umstellung 2025 im Unternehmensregister URS

Richard Bündgens (1), Holger Leerhoff (2), Birgit Pech (2), Schulz Julian (3), Christian Borgs (1), Mödinger Patrizia (4), Dohn Nicole (4)

1: IT.NRW, Statistisches Landesamt, Deutschland; 2: Amt für Statistik Berlin-Brandenburg, Deutschland; 3: Landesamt für Statistik Niedersachsen, Deutschland; 4: Statistisches Bundesamt (Destatis), Deutschland
christian.borgs@it.nrw.de

SECTION: Methodology of Statistical Surveys (Session 3)

Entsprechend den europäischen Vorschriften musste die Klassifizierung der Unternehmen im statistischen Unternehmensregister (URS) in eine nationale Version der NACE Rev. 2.1, der WZ 2025 („Klassifikation der Wirtschaftszweige“, Ausgabe 2025), umgewandelt werden.

Die Umsetzung erfolgte größtenteils automatisiert. In Fällen, in denen die Entsprechung zwischen der neuen und der alten Klassifikation nicht eindeutig war (1-zu-n oder n-zu-m-Fälle) konnte die Umstellung nicht automatisiert werden, und es musste der wahrscheinlichste Fall ermittelt und anschließend manuell überprüft werden. Um die große Anzahl von 1-zu-n-Fällen – etwa 1,7 Millionen Unternehmen – zu überprüfen, wurden verschiedene Methoden eingesetzt. Mehrere technische Lösungen und automatisierte Prozesse unterstützten den Umstieg, darunter ein auf maschinellem Lernen (ML) basierendes Modell zur teilweisen Automatisierung der Nachqualifikationen, welches im Statistischen Verbund entwickelt wurde.

Verschiedene ML-Modelle wurden anhand bis zum Trainingszeitpunkt manuell überprüfter Einheiten trainiert und haben hervorragende Resultate erzielt. Das final eingesetzte ML-Modell berücksichtigt diverse Kennzahlen, wie den alten WZ2008-Schlüssel, Texte zur Beschreibung des Unternehmens, die Anzahl der Beschäftigten und den Umsatz.

Das beste Modell erzielte eine sehr hohe Genauigkeit von über 98%. Damit ließen sich von ca. 1,47 Mio. zu prüfenden Unternehmen im URS etwa 1,10 Mio. WZ-Schlüssel für Unternehmen durch das ML-Verfahren automatisiert umschlüsseln. Das entspricht einer Automatisierungsquote von ca. 76% aller übrigen Fälle.

Damit ergibt sich durch den Einsatz eines maschinellen Lernverfahrens eine erhebliche Arbeitersparnis sowie potenzielle Qualitätsgewinne bei Fällen, deren vollständige händische Prüfung aus Zeitgründen ursprünglich nicht vorgesehen war.

147 Optimierung der Datenverknüpfung mit überwachtem maschinellen Lernen: Random Forest

Murad Hajiyeu, Christopher Nölting

Statistisches Bundesamt, Deutschland
murad.hajiyeu@destatis.de

SECTION: Methodology of Statistical Surveys (Session 3)

Die Integration von Verwaltungsdaten in das statistische Unternehmensregister ist ein zentraler Bestandteil moderner amtlicher Statistik, wird jedoch häufig durch fehlende eindeutige Identifikatoren sowie Inkonsistenzen in Namen und Adressinformationen erschwert. Dieser Beitrag stellt

einen Ansatz des überwachten maschinellen Lernens vor, um die Qualität der Datenverknüpfung mithilfe eines Random-Forest-Modells zu verbessern.

Die Methode baut auf vorhandenen Exaktübereinstimmungen als Trainingsdaten auf und kombiniert normalisierte Textinformationen mit optimierten Zeichenkettenähnlichkeitsmaßen und zusätzlichen binären Variablen. Ein wesentliches Element ist die systematische Auswahl von Ähnlichkeitsmaßen für verschiedene Attribute sowie die Konstruktion eines ausgewogenen Datensatzes, der sowohl echte als auch falsche Übereinstimmungen enthält.

Der Ansatz integriert Datenvorverarbeitung, Feature-Engineering und modellbasierte Klassifikation in einen kohärenten Workflow. Durch die Erfassung von Wechselwirkungen zwischen mehreren Variablen bietet er eine flexiblere Alternative zu traditionellen regel- oder scoring-basierten Verknüpfungsmethoden – insbesondere bei unvollständigen oder inkonsistenten Daten.

Die Ergebnisse zeigen das praktische Potenzial des überwachten Lernens für die Datensatzverknüpfung in der amtlichen Statistik und bieten eine skalierbare und übertragbare Lösung zur Verbesserung von Datenintegrationsprozessen.

148 Learning Copulas via Distributional Regressions

Sven Pappert

Ruhr University Bochum, Germany
sven.pappert@rub.de

SECTION: Statistical Theory and Methods (Session 3)

Any multivariate distribution can be decomposed into its marginal distributions and its copula. Consequently copula-based modeling has become a popular method whenever multivariate distributions need to be modeled. Several methods to construct copulas have been proposed. Among the most popular ones are implicit copulas, Archimedean copulas and more recently non-parametric and deep learning based methods. All these methods model the copula or a joint distribution more or less directly. We propose an approach in which the conditional quantile functions associated to copulas are modeled. This approach is based on establishing a correspondence between regression models and the copula of the associated random variables. The resulting copula may also be understood as an implicit copula approach for regressions. For models with only one covariate the approach is straight-forward and known copulas, e.g. the Gaussian copula, can be reconstructed. The approach can readily be extended to more complex relationships, e.g. non-Gaussian dependence, non-linear dependence as well as distributional dependence. We formally define these copula-classes and provide examples. Furthermore, extensions to more covariates are discussed and statistical modeling is exemplified.

149 Time series quantiles without moments

Harry Haupt

Universität Passau, Deutschland

harry.haupt@uni-passau.de

SECTION: Statistical Theory and Methods (Session 3)

This paper studies time series quantile regression asymptotics in an UC model $y = d + u$ for nonstationary y (due to trend and/or season d) and weakly dependent u . The results do not require the existence of any moments of the underlying stochastic processes. A CLT based on moments of the quantile score function is discussed and contrasted to alternatives frequently used in Econometric Theory.

150 Model checks for copula regression

Philip Dörr, Holger Dette

Ruhr-Universität Bochum, Deutschland

philip.doerr@rub.de

SECTION: Statistical Theory and Methods (Session 3)

There is a great variety of statistical models expressing relations between response variables of interest and explanatory variables, ranging from classical conditional mean regression to fully distributional regression models. We are particularly interested in expressing regression models by means of copulas which are a valuable tool to separate marginal distributions and dependencies. New goodness-of-fit tests and new measures of deviation can be developed based on such copula representations. These tests are desirable since regression models often impose parametric or semiparametric assumptions to overcome the curse of dimensionality, running a risk of misspecification. We present a new goodness-of-fit test for the classical mean regression model. More importantly, we also introduce a new measure of deviation between the true regression function and the imposed parametric assumption. By self-normalization, we develop pivotal inference for this measure including tests for relevant hypotheses. These inference tools are illustrated via simulated and empirical data.

151 Practically Significant Differences between Conditional Distribution Functions

Holger Dette (1), Kathrin Möllenhoff (2), **Dominik Wied** (2)

1: Ruhr-Universität Bochum; 2: Universität zu Köln

dwied@uni-koeln.de

SECTION: Statistical Theory and Methods (Session 3)

In the framework of semiparametric distribution regression, we consider the problem of comparing the conditional distribution functions corresponding to two samples. In contrast to testing for exact equality, we are interested in the (null) hypothesis that the L^2 distance between the conditional

distribution functions does not exceed a certain threshold in absolute value. The consideration of these hypotheses is motivated by the observation that in applications, it is rare, and perhaps impossible, that a null hypothesis of exact equality is satisfied and that the real question of interest is to detect a practically significant deviation between the two conditional distribution functions. The consideration of a composite null hypothesis makes the testing problem challenging, and in this paper we develop a pivotal test for such hypotheses. Our approach is based on self-normalization and therefore requires neither the estimation of (complicated) variances nor bootstrap approximations. We derive the asymptotic distribution of the (appropriately normalized) test statistic and show consistency under local alternatives. A simulation study and an application to German SOEP data reveal the usefulness of the method.

152 Werkstattbericht zum Einsatz von generativer KI in den Volkswirtschaftlichen Gesamtrechnungen

Stefan Hauf, Christoph-Martin Mai

Statistisches Bundesamt, Deutschland
christoph-martin.mai@destatis.de

SECTION: National Accounts, Welfare Measurement

Die Volkswirtschaftlichen Gesamtrechnungen (VGR) stehen vor wachsenden Herausforderungen: Neue internationale Standards - wie das SNA 2025/ESVG 2030 - erhöhen die fachlichen Anforderungen, während personelle und zeitliche Ressourcen zunehmend begrenzt sind. Vor diesem Hintergrund gewinnt der Einsatz von generativer Künstlicher Intelligenz (KI) als Instrument zur Effizienzsteigerung und Zukunftssicherung der Statistikproduktion an Bedeutung.

Der Vortrag gibt einen praxisnahen Einblick in aktuelle KI-Pilotvorhaben in den VGR. Im Mittelpunkt steht die automatisierte Unterstützung fachlicher Prüfprozesse, am Beispiel der Sektorklassifizierung in der Staatsrechnung. Dabei wird untersucht, inwieweit KI-Modelle komplexe Abgrenzungsfragen, wie etwa zur institutionellen Kontrolle oder zur Marktbestimmung, strukturieren, vorbereiten und fachlich vorbereiten können. Zudem werden Erfahrungen aus der bereichsübergreifenden Zusammenarbeit präsentiert, etwa bei Datenaufbereitung, Modellbewertung und Qualitätssicherung. Ein weiterer Fokus liegt auf den Rahmenbedingungen des KI-Einsatzes in der amtlichen Statistik. Diskutiert werden Anforderungen an Datenschutz, statistische Geheimhaltung, technische Infrastruktur sowie die Nachvollziehbarkeit modellgestützter Entscheidungen. Hierbei wird auch der Einsatz von KIPITZ, der sicheren "Bundes-KI" im Statistischen Bundesamt, vorgestellt, die datenschutzkonform ist.

Der Vortrag richtet sich an Fachleute aus der amtlichen Statistik und der öffentlichen Verwaltung und lädt zum Austausch über den sinnvollen Einsatz von generativer KI in der Statistikproduktion ein. Er zeigt erste Erfahrungen, benennt Grenzen und Potenziale und bietet erste praktische Einblicke in die Umsetzung von KI-Projekten.

153 Navigating the value chain – from the top to the bottom and back again. A unified statistical framework for micro-level carbon accounting

Ulf von Kalckreuth

Deutsche Bundesbank, Deutschland
ulf.von-kalckreuth@bundesbank.de

SECTION: National Accounts, Welfare Measurement

This paper is on evaluating indirect emissions resulting from economic activity for the purposes of sustainable finance. Indirect emissions are somebody else's emissions, and this creates specific problems for reporting. Both upstream and downstream emissions may come from a multitude of counterparts, or from the counterparts of counterparts, and so on. In most cases, the counterpart entities do not have any reporting obligation to the company.

The aim of this paper is to present a unified analytical framework for measuring and estimating upstream and downstream emissions in a context of inter-company linkages. The framework clarifies specifically how downstream emissions in subsequent parts of the value chain can be conceptualised to enable analytically tractable solutions, in parallel to the well-known Leontief framework used for upstream emissions. A closed form solution will result if downstream emissions are allocated to intermediate products by their share in output value. With this intuitive allocation rule, the framework can readily be used to calculate industry level averages based on Input-Output and emission data from official statistics, which in turn may serve as proxies for company level carbon accounting, to get the analysis started and to fill partial information gaps.

The essay starts by reviewing the IO based framework for upstream indirect emissions. The framework is extended to encompass downstream emissions. The resulting methodology is used to compute both upward and downward emission intensities from FIGARO Input-Output data and emission statistics. Such data can be used as proxies for carbon accounting on the micro level. Ultimately, the potential of company level carbon accounting to deliver reliable product level information for downstream emissions is inspected.

154 Konjunkturstatistik neu denken – Navigation durch die Vielfalt der Nutzerbedarfe

Stefan Linz

Statistisches Bundesamt, Deutschland
stefan.linz@destatis.de

SECTION: National Accounts, Welfare Measurement

Die Unternehmensstatistiken haben sich über viele Jahrzehnte hinweg erweitert, um neue Informationsbedarfe zu bedienen. Dadurch ist ein komplexes, historisch gewachsenes System entstanden, dessen einzelne Statistiken nicht immer konsistent aufeinander abgestimmt sind. Der Reformansatz „System der Unternehmensstatistik“ (SysdU) setzt hier an, indem er das Gesamtsystem in den Blick nimmt und eine kohärentere Architektur sowie effizientere Produktionsprozesse anstrebt.

Eine grundlegende Neuordnung eines solchen Systems wirft jedoch zwangsläufig die Frage nach seiner inhaltlichen Ausrichtung auf: Wofür soll das System der Unternehmensstatistiken in erster

Linie Informationen bereitstellen? Die Beantwortung dieser Frage erfordert eine systematische Auseinandersetzung mit den Nutzerbedarfen. Diese sind jedoch vielfältig und teilweise heterogen: Je breiter der Kreis der Nutzerinnen und Nutzer betrachtet wird, desto größer wird die Bandbreite der Anforderungen an statistische Informationen.

Vor diesem Hintergrund besteht eine zentrale Herausforderung darin, diejenigen Bedarfe zu identifizieren, die für die grundlegende Ausrichtung des Systems maßgeblich sind. Darüber hinausgehende Anforderungen – etwa zu spezifischen Branchen, Detailfragen oder einzelnen Analyseperspektiven – sind für sich genommen ebenso relevant, eignen sich jedoch weniger als Leitlinie für die Gestaltung der Systemarchitektur. Sie lassen sich vielmehr gezielt aufgreifen, sobald die grundlegende Struktur des Systems definiert ist.

Der Beitrag nimmt die Konjunkturstatistiken als Teil der Unternehmensstatistiken in den Blick und diskutiert, welche grundlegenden Nutzerbedarfe für die amtliche Konjunkturbeobachtung prägend sind. Ziel ist es, einen konzeptionellen Rahmen zu skizzieren, der dabei helfen kann, die Vielfalt der Anforderungen zu strukturieren und damit eine Grundlage für die Weiterentwicklung der Konjunkturstatistik im Kontext des SysdU zu schaffen.

155 Verwendung des statistischen Unternehmensregisters für daten-basierte Analysen. Restriktionen und Lösungsansätze für ausgewählte Fragestellungen

Christian Salwiczek

Statistikamt Nord, Deutschland
Christian.salwiczek@statistik-nord.de

SECTION: National Accounts, Welfare Measurement

Grundsätzlich ermöglicht das Unternehmensregister (URS) vielfältige und umfassende Datenanalysen. Hierbei wirken sich die zahlreich integrierten Datenquellen und ein breites Spektrum an Merkmalen positiv aus. Für unterschiedliche Entitäten sind zeitreihenbasierte und georeferenzierte Analysen für Politik, Verwaltung, Forschung und weitere Nutzergruppen möglich.

In jüngerer Zeit stellten sich vermeintlich „einfache“ Fragestellungen als herausfordernd heraus. So ist dies bei der Frage der lokalen Versorgungssicherheit oder in Bezug auf die regionale Wärmeleitplanung der Fall gewesen.

Im Rahmen des Vortrags sollen ausgewählte, praxisorientierte Ansätze vorgestellt werden, welche trotz bestehender Restriktionen und Nicht-Datenverfügbarkeiten derartige Fragen adäquat bearbeitet werden können. Perspektivisch können hierdurch auch Impulse für eine kommende Wirtschaftsregister-Anwendung erhalten werden.

156 Code Generation for Open Data Statistics: Case Study on the Genesis Database

Valentin Schmidberger (2), Johannes Maucher (2), Julia Lapp (2), Albrecht Melan (1)

1: Statistisches Landesamt Baden-Württemberg; 2: Hochschule der Medien, Stuttgart
albrecht.melan@stala.bwl.de

SECTION: Poster Session

This work presents a method for enabling Large Language Models (LLMs) to accurately answer questions about complex statistical tables from open data sources, such as Germany's Genesis database. These tables often feature hierarchical structures embedded in visual formats (e.g., XLSX files with merged headers or denormalized CSV data), making them difficult for LLMs to process directly. To address this, we developed a two-step approach: First, tables are transformed into a unified JSON structure using deterministic heuristics that preserve semantic relationships by analyzing formatting cues like cell borders and indentation. Second, an LLM is fine-tuned to generate executable Python code that programmatically retrieves the requested information from the JSON data.

To train the model, we created a synthetic dataset of 13,000 question-code pairs covering diverse query types, including lookups, aggregations, and calculations. Each example was validated through automated execution to ensure correctness. The fine-tuned Microsoft Phi-4 model achieves an accuracy of 88.5% (peaking at 95.2% for XLSX files), matching the performance of leading proprietary models like GPT-4.1 while maintaining full data sovereignty through local processing. For resource-constrained environments, a smaller Llama-3.1-8B model delivers 78.6% accuracy with significantly lower computational requirements.

The system's architecture separates logical data retrieval (code generation) from result synthesis (natural language answers), enhancing transparency and adaptability. A prototype demonstrates practical usability, including multimodal input (text/voice) and output (speech synthesis). The results show that domain-specific fine-tuning enables medium-sized open-source models to rival large commercial LLMs in specialized table analysis tasks, offering a scalable solution for government agencies and data portals to make complex statistics more accessible.

157 Daten-Ethik als Element im digitalen Data-Literacy-Lehrprojekt

Christine Buchholz

Hochschule Bonn-Rhein-Sieg, Deutschland
christine.buchholz@h-brs.de

SECTION: Poster Session

Im Projekt „DAViD - Daten analysieren, visualisieren und deuten“ an der Hochschule Bonn-Rhein-Sieg entstanden digitale, freizugängliche Lehr-Lern-Materialien, die Basiskompetenzen im Umgang mit Daten vermitteln. 15 digitale E-Learning-Sessions mit jeweils 90-minütigem Umfang stehen als OER (Open Educational Resources) auf der Plattform ORCA.nrw zur Verfügung.

In diesem Projekt thematisiert eine Einheit den Bereich der Datenethik und wurde inhaltlich von Prof. Gert Scobel gestaltet. Studierende und interessierten Nutzer:innen erhalten eine Einführung in die Themenbereiche „Datenethik“ und „Ethik der Information“.

Im ersten Teil zur Datenethik wird der Begriff der Ethik erläutert und damit verbundene Fragestellungen vorgestellt. Dabei wird Grundwissen über Ethik und Moral vermittelt sowie die Fähigkeit thematisiert, gesamtgesellschaftliche wie wissenschaftliche Fragestellungen unter ethischen Gesichtspunkten zu würdigen.

Im zweiten Teil zur Informationsethik geht es um Fragen der Ethik in Verbindung mit der Digitalisierung der Gesellschaft. Dabei befasst sich die Informationsethik mit Argumenten, Werten, Handlungen und deren Auswirkungen, sofern diese durch und mit Informationen, Informationsverarbeitung und Daten verbunden sind.

In diesem Beitrag erfolgt zum einen eine kurze Vorstellung der Inhalte, zum anderen wird die didaktische Einbettung sowie die curriculare Verankerung an Beispielen aus der Hochschulpraxis vorgestellt. Die Nutzung über ORCA.nrw ermöglicht eine Implementierung im Blended-Learning-Format, wobei die digitalen Inhalte vorbereitend für die Präsenzlehre genutzt werden können und dann mit den Studierenden im Kontext der jeweiligen Fachdisziplin eingebettet, diskutiert und vertieft werden können.

158 Forschungsstrategie des Statistischen Bundesamtes: Kooperation mit der Wissenschaft

Lisa Carius-Munz, Jannek Mühlhan

Statistisches Bundesamt, Deutschland
lisa.carius-munz@destatis.de

SECTION: Poster Session

Im Kontext der methodischen Weiterentwicklung statistischer Erhebungen ist die systematische Integration wissenschaftlicher Innovationen von zentraler Bedeutung. Dies betrifft insbesondere den Umgang mit neuen Datenquellen, nicht-probabilistischen Stichprobenverfahren und technologischen Entwicklungen, die eine enge Verzahnung von Praxis und Forschung erfordern. Das Statistische Bundesamt verfolgt daher eine strategische Kooperationspolitik mit der Wissenschaft, die verschiedene Formen der Zusammenarbeit umfasst. Institutionell verankert ist bereits der Statistische Beirat, der als beratendes Gremium Wissenschaft, Nutzer und Produzenten amtlicher Statistik zusammenbringt. Gemeinsam berufene Professuren mit Hochschulen zu ausgewählten methodischen Fragestellungen der amtlichen Statistik sind mittelfristig vorgesehen. Darüber hinaus gibt es Überlegungen zur Einrichtung eines eigenständigen Begleitgremiums vergleichbar zu anderen nationalen Statistikämtern wie Istat und CBS oder anderen Forschungseinrichtungen, die die Modernisierung der Statistikproduktion auf Basis aktueller wissenschaftlicher Erkenntnisse begleiten soll.

Etablierte Kooperationsformate umfassen gemeinsame Forschungsprojekte – gefördert durch Bundesbehörden, Forschungsgemeinschaften und die Europäische Union – sowie die Organisation wissenschaftlicher Konferenzen, darunter das jährliche Wissenschaftliche Kolloquium und die zweijährliche Wissenschaftliche Tagung zu Erhebungsmethoden mit den Markt- und Sozialforschungsverbänden ADM und ASI. Die Fachzeitschrift WISTA fungiert als wissenschaftliches Magazin zur Publikation fortgeschrittener Methodendokumentation oder wissenschaftlicher Analysen. Ein weiterer Schwerpunkt der Forschungsstrategie liegt in der Nachwuchsförderung. Das Statistische Bundesamt engagiert sich im EMOS-Programm, vergibt den Wissenschaftlichen Nachwuchspreis, unterstützt Abschlussarbeiten mit Bezug zur amtlichen Statistik, bietet Praktika an und hat Lehraufträge an Hochschulen inne.

Der Beitrag diskutiert die bisherigen Ergebnisse dieser Kooperationsstrategie, benennt bestehende Herausforderungen und zeigt offene Fragen auf dem Weg zu einer systematisch verankerten Forschungspartnerschaft zwischen amtlicher Statistik und Wissenschaft.

159 Kernergebnisse des Projektes "Freiraumverluste 2030"

Nadine Blätgen (1), Jana Hoymann (1), Sylvie Dugay (1), Fabian Dosch (1), Florian Bernardt (2), Philipp Ulrich (2), Martin Behnisch (3), Hannah Poglitsch (3), Jens-Martin Gutsche (4)

1: Bundesinstitut für Bau-, Stadt- und Raumforschung (BBSR); 2: Gesellschaft für wirtschaftliche Strukturforchung (GWS); 3: Leibnitz Institut für ökologische Raumentwicklung; 4: Gertz Gutsche Rümenapp
nadine.blaetgen@bbr.bund.de

SECTION: Poster Session

Das Poster präsentiert zentrale Ergebnisse des Forschungsvorhabens Freiraumflächen 2030 – Quantitative Abschätzung von Änderungen in der Flächennutzung. Ziel des Vorhabens ist die Erstellung einer evidenzbasierten Prognose zum Freiraumflächenverlust in Deutschland bis 2030 und 2050. Grundlage bilden qualitative und quantitative Modellierungen der künftigen Flächennutzungsentwicklung.

160 Kleinräumige Entwicklungen in Städten analysieren: Datensatz „Innerstädtische Raubeobachtung“

Judith Kaschowitz, Cornelia Müller, Dorothee Winkler

Bundesinstitut für Bau-, Stadt- und Raumforschung (BBSR)
cornelia.mueller@bbr.bund.de

SECTION: Poster Session

Der Datensatz der Innerstädtischen Raubeobachtung (IRB) ist eine kleinräumige, jährlich aktualisierte kommunalstatistische Datenbasis mit rund 400 Merkmalen zu Bevölkerung, Haushalten, Wanderungen, Beschäftigung und Wohnungsbestand auf Stadtteilebene. Die IRB ist ein Kooperationsprojekt zwischen dem BBSR und 55 deutschen Großstädten. Seit 2002 stellen die Städte auf Basis eines abgestimmten Merkmalskatalogs jährlich kleinräumige Daten auf Stadtteilebene bereit. Eine Verschneidung mit weiteren georeferenzierten Daten (z.B. Rasterdaten, Umfragedaten) ist möglich. Der Wert der IRB als in Deutschland einmalige Datenbasis wird zunehmend in der Wissenschaft erkannt; die Daten können für Masterarbeiten und wissenschaftliche Projekte angefragt werden. Das Poster stellt die IRB vor und zeigt Forschungsergebnisse auf Basis des Datensatzes (z. B. zu sozialer Ungleichheit und demografischen Entwicklungen). Ziel des Posterbeitrags ist es, Studierende und Forschende auf den Datensatz aufmerksam zu machen und durch BBSR-Expertinnen eine Vor-Ort-Beratung zur Datennutzung anzubieten.

161 Modeling Long Memory Cyclical Trends in the Cenozoic

Tomás del Barrio Castro (1), Álvaro Escribano (2), **Yeliz Özer** (3), Philipp Sibbertsen (3)

1: Universitat de les Illes Balears; 2: Universidad Carlos III de Madrid; 3: Leibniz Universität Hannover, Deutschland
oezer@statistik.uni-hannover.de

SECTION: Poster Session

Long paleoclimate time series combine strong persistence, multiple orbital cycles, and regime shifts, which complicates the analysis of coupling and predictability. We analyze Cenozoic variability using the CENOGRID dataset, in a regime-based framework with segments defined by major climate-state transitions, combining deterministic decomposition, cyclical fractional cointegration, and forecasting. Within each segment we estimate and remove linear trends, orbital forcing variables, and harmonic cycles identified via GARMA-based filtering. We then test for cyclical fractional cointegration at shared orbital frequencies to assess cointegration between the proxies and orbital variables. The results reveal regime dependence: the 405 kyr eccentricity cycle band repeatedly cointegrates with both proxies across different climate states. Finally, regime-specific VAR(2) models are fitted to the residuals and in-sample forecasts reported. Additionally, we generate a 100 kyr out-of-sample projection based on the Icehouse specific dynamics. Forecast performance varies across regimes, highlighting that non-stationarity and regime-specific dynamics constrain predictability.

162 Probabilistic Modeling and Forecasting of Household Energy Consumption

Marco Jeschke, **Leia Betting**, Roland Fried

Technische Universität Dortmund, Deutschland
leia.betting@tu-dortmund.de

SECTION: Poster Session

Our aim is to model electricity consumption at the household level and to deliver probabilistic forecasts. To this end, we explore hidden Markov models and periodic autoregressive (PAR) models, and, in particular, seek to model the influence of various covariates on the model parameters. Among other things, the first approach leads to an extension of the classical Markov model, as previous observations can influence the emission parameters as covariates (autoregressive hidden Markov model). Concerning the second approach, while PAR models typically assume normally distributed innovations, empirical residuals often exhibit skewness and heavy tails. To address this limitation, we apply different techniques to handle non-normality in the residuals and thus enable the construction of more realistic prediction intervals while maintaining the simplicity and interpretability of PAR models.

163 Skalierbares Record Linkage für administrative Bevölkerungsdaten: Erfahrungen aus einem Methodentest

René Kremer

Destatis, Deutschland
rene.kremer@destatis.de

SECTION: Poster Session

Die Verknüpfung administrativer Bevölkerungsdaten gewinnt für die amtliche Statistik zunehmend an Bedeutung. In Deutschland ist sie jedoch aufgrund dezentraler Datenstrukturen, heterogener Datenformate, fehlender eindeutiger Identifikatoren sowie der großen Zahl potenzieller Datensatzvergleiche besonders herausfordernd. Der vorliegende Beitrag stellt zentrale methodische Überlegungen und praktische Erkenntnisse aus einem Methodentest des Statistischen Bundesamtes vor, der der Vorbereitung künftiger zensusbezogener Verknüpfungsaufgaben dient.

Im Mittelpunkt stehen wesentliche Vorbereitungsschritte, darunter die Standardisierung und Normalisierung von Variablen sowie die Auswahl geeigneter Verknüpfungsstrategien. Dabei werden exakte und deterministische Verfahren dem probabilistischen Record Linkage gegenübergestellt, dessen Stärken sich insbesondere im Umgang mit Tippfehlern, fehlenden Werten und Inkonsistenzen zwischen Registern zeigen. Besonderes Augenmerk gilt dem Einsatz von Splink, einer Open-Source-Python-Bibliothek, die das Fellegi-Sunter-Modell in skalierbaren Umgebungen wie Spark implementiert. Die Unterstützung von Blocking, ähnlichkeitsbasierten Vergleichen und eine Match-Gewichtung macht Splink zu einem vielversprechenden Werkzeug für die Verknüpfung umfangreicher administrativer Datenbestände.

Erste Ergebnisse des Methodentests deuten darauf hin, dass probabilistisches Record Linkage einen flexiblen und skalierbaren Rahmen für die Identifikation plausibler Übereinstimmungen bietet. Damit wird das Potenzial moderner Verknüpfungsmethoden für die amtliche Bevölkerungsstatistik in komplexen administrativen Datenlandschaften deutlich.

164 The Tail the Subsidy Doesn't Reach: Distributional Evidence on Extreme Energy Burdens in the US

Louise Drozdek (1), Franziska Dorn (1), Simone Maxand (2)

1: University of Duisburg-Essen, Deutschland; 2: Europa-Universität Viadrina
louise.drozdek@stud.uni-due.de

SECTION: Poster Session

Energy poverty constitutes a structural and intersectional dimension of inequality in the United States, yet its extreme manifestations remain systematically obscured by conventional mean-based measurement approaches that fail to capture the heterogeneity of at-risk households. While the literature increasingly acknowledges this heterogeneity, existing studies remain largely limited to descriptive assessments, leaving the systematic drivers of distributional inequality empirically unquantified. To address this gap, this study applies distribution-sensitive GAMLSS regression models to the full distribution of the relative energy burden, simultaneously modeling location, dispersion, and tail parameters as functions of key covariates. Drawing on data from the 2019 Panel Study of Income Dynamics, differences are examined by gender, ethnicity, race, household composition, and region, including states with distinct climate and price regimes such

as Massachusetts and Texas. Results indicate substantially higher burdens and a greater likelihood of extreme energy costs among Black and Hispanic households, particularly in female-headed households and colder regions. Quantile ratio measures reveal that heterogeneity is primarily driven by upper-tail stretching rather than lower-tail dispersion. A tail-risk criterion based on the 75th percentile illustrates meaningful variation in burden intensity within the subsidy-eligible population, while governmental energy subsidies exhibit markedly heterogeneous effects across the burden distribution, only partially mitigating structural and intersectional inequalities among vulnerable households.

165 Verbessern europäische Regelungen den Zugang zu Privately Held Data? Mobilfunk- und Plattformdaten im Kontext der revidierten Verordnung (EG) Nr. 223/2009

Maren Koehlmann, Markus Zwick

Statistisches Bundesamt, Deutschland
maren.koehlmann@destatis.de

SECTION: Poster Session

Die Revision der EU-Statistikverordnung als Verordnung (EU) 2024/3018 des Parlaments und des Rates zur Änderung der Verordnung (EG) Nr. 223/2009 über europäische Statistiken zielt darauf ab, das europäische Statistiksysteem an neue Datenrealitäten anzupassen. Ein zentraler Aspekt der aktuellen Diskussion basierend auf der Revision, ist der verstärkte Zugang zu Privately Held Data (PHD), also Datenbeständen, die von privaten Unternehmen generiert und gehalten werden. Solche Daten – etwa aus digitalen Plattformen, Mobilitätsdiensten oder Online-Transaktionen – bieten das Potenzial, klassische statistische Erhebungen komplementär zu erweitern, indem sie unter geeigneten methodischen und rechtlichen Rahmenbedingungen eine effiziente, zeitnahe und granulare Gewinnung belastbarer Datengrundlagen ermöglichen.

Der Beitrag beleuchtet die bisherigen Erfahrungen mit der Nutzung von PHD im Rahmen amtlicher Statistik sowie bestehende Pilotprojekte und Kooperationen zwischen statistischen Ämtern und privaten Datenhaltern. Zudem wird diskutiert, in welchen statistischen Domänen PHD künftig einen besonderen Mehrwert bieten können und welche Herausforderungen sich für die praktische Umsetzung aus der Revision der Verordnung (EG) Nr. 223/2009 über europäische Statistiken ergeben.

166 WISTA Wirtschaft und Statistik

Mareike Muth

Statistisches Bundesamt, Deutschland
mareike.muth@destatis.de

SECTION: Poster Session

WISTA als bedeutendste wissenschaftliche Publikation des Statistischen Bundesamtes dient der Methodendokumentation der amtlichen Statistik. WISTA ermöglicht die wissenschaftliche Diskussion, indem auch Wissenschaftlerinnen und Wissenschaftler aus den Statistischen Ämtern der Länder, Forschungsinstituten oder dem Hochschulbereich in WISTA veröffentlichen können. Mit einem Poster auf der Statistischen Woche soll die Bekanntheit von WISTA beim Fachpublikum gesteigert und sowohl um neue Lesende als auch Autorinnen und Autoren geworben werden. Hierfür stellt das Poster das Wissenschaftsmagazin, den redaktionellen Service und die Veröffentlichungskanäle vor.

167 Record Linkage powered by LLMs: Mikrodatenverknüpfungen mit Transformer-basierten Sprachmodellen

Igor Franjić, Sara Schiesberg

Statistisches Bundesamt, Deutschland
igor.franjic@destatis.de

SECTION: Methodology of Statistical Surveys (Session 4)

Transformer-basierte Sprachmodelle wie BERT oder GPT erhalten derzeit immer größere Aufmerksamkeit und Verbreitung in der amtlichen Statistik. Repräsentative Sprachmodelle ermöglichen es, die Bedeutung von Wörtern, Sätzen und längeren Texten numerisch darzustellen und so Aufgaben wie Textklassifikation, semantische Suche oder Named-entity recognition elegant zu lösen, ohne die Texte vorher aufwendig standardisieren zu müssen.

Der Vortrag stellt die Potenziale Transformer-basierter repräsentativer Sprachmodelle zur Mikrodatenverknüpfung vor. Anhand eines Testdatensatzes aus 300 handkodierten Unternehmen wird die Performance Transformer-basierter Mikrodatenverknüpfung mit und ohne vorheriges Finetuning mit etablierten Methoden des Maschinellen Lernens und regelbasierten Ansätzen verglichen. Die Ergebnisse zeigen, dass auch ohne Finetuning des zugrundeliegenden Sprachmodells Resultate erzielt werden können, die bestehenden Machine-Learning-Verfahren ebenbürtig sind.

168 Knochensägen, Katheter und Ibuprofen: Produktklassifikation mit LLMs im Krankenhausbereich

Yannik Buhl

Statistisches Bundesamt, Deutschland
yannik.buhl@destatis.de

SECTION: Methodology of Statistical Surveys (Session 4)

Large Language Models (LLMs) entwickeln sich innerhalb des Natural Language Processing zunehmend zu einer der besten Methoden für Textklassifikation. Die Vorteile liegen auf der Hand: Die Modelle bringen ein gutes Kontextverständnis mit, sind bereits vortrainiert und benötigen wenig bis keine Textaufbereitung. Aus diesem Grund untersucht der Bereich Gesundheitsstatistiken im Statistischen Bundesamt im Rahmen eines Projektes derzeit den Einsatz von LLMs, welches die Klassifizierung von Produktlisten aus Krankenhauswarenwirtschaftssystemen zum Ziel hat. Dabei sollen die Produkte, für die nur wenig Zusatzinformationen verfügbar sind, in eine eigens entwickelte hierarchische Struktur einsortiert werden. Getestet werden klassische BERT-basierte Modelle, aber auch Ansätze der Retrieval Augmented Generation (RAG), in deren Rahmen auch agentische Ansätze denkbar sind, in denen z.B. ein Sprachmodell mit klassischem Retrieval kombiniert wird. Eine weitere Herausforderung ist die Notwendigkeit, für weiterführende Berechnungen u.a. Mengeninformationen bzw. Packungsgrößen aus den Produktbeschreibungen zu extrahieren. Erste Ergebnisse zeigen, dass LLM-basierte Ansätze deutlich bessere Ergebnisse erzielen als klassische Machine-Learning-Algorithmen. Jedoch ist in vielen Fällen ein Finetuning der genutzten Modelle unerlässlich, was wiederum Aufwand in Form von Trainingsdaten erzeugt.

169 Kontextbasierte Klassifikation in den freiwilligen Haushaltserhebungen: Ein agentenbasierter GenAI-Workflow

Adrian Montag

Statistisches Bundesamt, Deutschland
adrian.montag@gmail.com

SECTION: Methodology of Statistical Surveys (Session 4)

In freiwilligen Haushaltserhebungen müssen Konsumausgaben nach dem hierarchischen COICOP/SEA-System klassifiziert werden. Eine Automatisierung der Klassifikation dieser Konsumausgaben stellt in der Praxis eine erhebliche Herausforderung dar, da von Haushalten erfasste Produktbeschreibungen oft stark abgekürzt oder innerhalb des Klassifikationssystem in mehr als eine Kategorie eingeordnet werden können. Klassische Machine-Learning-Methoden und String-Matching-Verfahren bieten zwar einen großen Mehrwert, stoßen bei diesen Mehrdeutigkeiten jedoch an ihre Grenzen.

Um diese Herausforderungen zu überwinden, wird das Potenzial von GenAI-basierten agentischen Systemen zur Lösung mehrdeutiger Klassifikationsfälle untersucht. Das System modelliert den Entscheidungsprozess menschlicher Kodierer, die Mehrdeutigkeiten häufig durch die Betrachtung ähnlicher Ausgaben oder des Haushaltskontexts auflösen. Der vorgestellte Workflow integriert eine Vektordatenbank für den initialen Abruf historischer Daten und nutzt KI-Agenten, um mehrdeutige Ergebnisse zu analysieren und bei Bedarf spezifische Entscheidungsregeln und Kontextvariablen des Haushalts zur finalen Entscheidungsfindung heranzuziehen.

Der Vortrag evaluiert die Klassifikationsgenauigkeit (F1-Score) der GenAI-Agenten im Vergleich zu klassischen ML-Ansätzen. Zudem werden die Auswirkungen verschiedener Agenten-Muster (z. B. ReAct), Modell Größen sowie automatische Prompt-Optimierungsstrategien (GEPA) auf Robustheit, Fairness und Kosteneffizienz beleuchtet. Abschließend wird gezeigt, inwiefern sich der entwickelte, quelloffene Workflow als abstraktes Modell auf andere hierarchische Klassifikationssysteme in der amtlichen Statistik übertragen lässt.

170 Neuronale Nomenklatur: Retrieval-then-Reasoning für die Außenhandelsstatistik

Sara Schiesberg

Statistisches Bundesamt, Deutschland
sara.schiesberg@destatis.de

SECTION: Methodology of Statistical Surveys (Session 4)

Der Vortrag widmet sich der automatisierten Klassifikation von Freitextbeschreibungen in die Kombinierte Nomenklatur. Dabei handelt es sich um ein hochkomplexes System mit knapp 10.000 Klassen. Dieses System bildet die Grundlage der Außenhandelsstatistik, die den internationalen Warenverkehr nach Warenklassen aufschlüsselt.

Die automatisierte Klassifikation wird durch einen Retrieval-then-Reasoning-Ansatz umgesetzt. Die zugrunde liegende Architektur ist modular aufgebaut und kombiniert repräsentative und generative Sprachmodelle zur Identifikation relevanter Klassen sowie zur finalen Entscheidungsfindung. Im Rahmen des Vortrags liegt ein besonderer Fokus auf der generativen Komponente, die aufgrund ihres hohen Ressourcenbedarfs eine zentrale Herausforderung darstellt. Untersucht wird insbesondere, inwieweit die Modellgröße die Performanz beeinflusst und welche Effekte durch FineTuning erzielt werden können.

Der Beitrag gibt Einblicke in die praktische Nutzung großer Sprachmodelle in der amtlichen Statistik und zeigt deren Potenzial als leistungsfähige, zugleich aber anspruchsvolle Erweiterung bestehender Verfahren.

171 Asymptotic Limits for Component-wise L2 Boosting: Theory, Derivation, and Simulation Evidence

Daria Ovsyannikova

TU Dortmund, Germany
daria.ovsyannikova@tu-dortmund.de

SECTION: Statistical Theory and Methods (Session 4)

This paper develops asymptotic inference theory for L2-Boosting estimators in linear regression. In the joint-update setting, the boosted estimator is shown to be exactly Gaussian at any finite sample size, and valid root-n inference is established under explicit conditions on the stopping time schedule. The high-dimensional regime where p/n converges to a positive constant is explored: in this setting OLS either does not exist or becomes highly variable, and the regularization provided by early stopping is essential rather than optional. The implicit regularization of boosting is connected to ridge regression via spectral shrinkage, providing a bridge to the well-developed ridge inference literature. For the component-wise setting, the asymptotic distribution of the one-step

block estimator is derived in closed form: when both blocks are equally informative, the limiting distribution is a mixture of skew-normal distributions, driven by the correlation between block selection noise and within-block estimation noise. A consequence of bivariate normal symmetry is that symmetric two-sided confidence intervals retain correct coverage despite non-Gaussianity. However, one-sided inference is systematically distorted — the naive one-sided test rejects at roughly 9.4% rather than the nominal 5%. Block-specific conditional quantiles are derived to restore valid directional inference. Monte Carlo simulations support the theoretical results throughout and illustrate the practical relevance of the corrections.

172 Distributional Random Forests in the context of hierarchical data

Johanna Einhorn, Timo Schmid

Otto-Friedrich-Universität Bamberg, Deutschland
johanna.einhorn@uni-bamberg.de

SECTION: Statistical Theory and Methods (Session 4)

Random forests are widely used non-parametric methods for regression and classification, but standard implementations focus on conditional expectations. Distributional Random Forests extend this framework by enabling the estimation of full conditional distributions and thereby allow inference on a broad class of distributional functionals for potentially multivariate response variables.

In many applications, observations exhibit a hierarchical structure, where units are nested within groups. Standard Distributional Random Forest approaches typically neglect such dependencies, which may lead to biased results and an inadequate representation of group-level variation.

We propose the Adjusted Distributional Random Forest (ADRF), a non-parametric method for estimating conditional distributions under hierarchical structure. The approach modifies predictions from Distributional Random Forests by incorporating group-level information through a probabilistic adjustment. This yields conditional distributions that account for both covariates and group membership, while retaining the flexibility of the underlying machine learning model, including its non-parametric nature and its ability to handle multivariate response variables.

Furthermore, we develop an approach to assess prediction accuracy and introduce diagnostic measures to quantify the impact of group-level effects on the estimated distributions. The proposed method provides a systematic extension of the Distributional Random Forest to clustered data settings and enables the analysis of distributional heterogeneity across groups.

173 A universal time series model (for continuous data)

Malte Jahn

Hochschule Merseburg, Deutschland
malte.jahn@hs-merseburg.de

SECTION: Statistical Theory and Methods (Session 4)

A novel time series framework is proposed which addresses all relevant empirical properties of a time series, making it an essentially universal model. More specifically, the dynamics in all conditional moments of a suitable continuous or discrete distribution are modelled jointly and without the need to make restrictive assumptions about the functional form of the link functions.

Furthermore, all considered explanatory variables are allowed to exhibit nonlinear and potentially time-varying effects on the conditional moments. This can be achieved by employing a simple feedforward neural network with a single hidden layer and an output for each conditional moment (parameter). In contrast to many (deep) neural network approaches, the proposed model is stochastically interpretable and allows for the calculation of standard errors, and in particular, confidence intervals. Many conventional time series frameworks such as (integer-valued) GARCH can be interpreted as simplified special cases of the proposed model. Several empirical applications are presented to illustrate the capabilities and the implementation.

174 High dimensional item selection via penalized multinomial logistic

Jennifer Simota, Anne Leucht

Otto-Friedrich-Universität Bamberg, Deutschland
jennifer.simota@uni-bamberg.de

SECTION: Statistical Theory and Methods (Session 4)

Survey instruments frequently contain a large number of items, which tends to increase respondent burden and the prevalence of missing values. To address this, we propose a method for data-driven item selection based on multinomial logistic regression with a network lasso penalty. The network penalty encourages the fusion of items with similar effects by shrinking their coefficient vectors toward one another. Thereby, we effectively identify redundant items that can be removed from the questionnaire.

Exploiting the separability of the resulting objective, we employ an inexact version of the Alternating Direction Method of Multipliers (ADMM) as an efficient solver and establish its convergence to the global optimum.

To illustrate our approach, we apply it to a linguistic survey examining the tendencies of Maltese English speakers toward British or American English varieties, where questionnaire reduction is of direct practical relevance. The proposed method yields a substantially reduced item set while maintaining the inferential structure of the original survey.

175 Evaluierung der Haushaltegenerierung im Zensus 2022

Johannes Baumstark (1), Michael Fürnrohr (2), Adrian Reichert (1), Judith Schwitulla (1)

1: Bayerisches Landesamt für Statistik; 2: Otto-Friedrich-Universität Bamberg
johannes.baumstark@statistik.bayern.de

SECTION: Methodology of Statistical Surveys (Session 5)

Im Zensus 2022 wurde wie im Zensus 2011 ein registergestütztes Verfahren in Deutschland genutzt, welches auf eine klassische Vollerhebung verzichtet. Zur Verknüpfung der Daten der Gebäude- und Wohnungszählung mit den Personendaten sowie zur gleichzeitigen Erzeugung von Haushalten kam in beiden Zensus die Haushaltegenerierung zum Einsatz. Die Funktionsweise der Haushaltegenerierung wurde bislang nicht unabhängig von ihren Eingangsdaten empirisch überprüft, was mit dieser Arbeit erstmals geleistet werden konnte. Dazu wurden die Zensusstichprobendaten 2022 als Untersuchungsgrundlage verwendet, indem die erhobenen Haushalte aus der Stichprobe mit auf denselben Daten generierten Haushalten verglichen wurden. Als Gütekriterium wurde u.a.

ein Distanzmaß verwendet, welches den Abstand zwischen den erhobenen und den generierten Haushalten an einer Anschrift quantifiziert. Die Ergebnisse zeigen, dass die Haushaltegenerierung unter Idealbedingungen, d.h. bei annähernd perfekten Eingangsdaten, die primärstatistisch erhobenen Haushalte ebenfalls nahezu vollständig korrekt abbilden kann. Auch unter Realbedingungen, wie sie im Zensus 2022 vorzufinden waren, konnte die Haushaltegenerierung die Haushalts- und Familienstrukturen zuverlässig abbilden. Die Haushaltegenerierung kommt als bewährtes Verfahren aller Voraussicht nach in der nächsten Zensusrunde 2031 wieder zum Einsatz.

176 Eine methodische Kritik des Zensus 2022

Rainer Schnell (1), Rainer Lenz (2), Rolf Schmidt (1)

1: Universität Duisburg-Essen; 2: TH Köln
rainer.schnell@uni-due.de

SECTION: Methodology of Statistical Surveys (Session 5)

Die Ergebnisse des Zensus 2022 wichen in zahlreichen Gemeinden so stark von den Melderegistern ab, dass mehrere Hundert Gemeinden Widerspruch gegen die Ergebnisse des Zensus eingelegten. Ein Zusammenschluss von ca. 290 Gemeinden hat die Autoren mit einem Gutachten zu den Verfahren des Zensus 2022 beauftragt. Das Statistische Bundesamt hat keine einheitliche Dokumentation des Zensus 2022 vorgelegt, es gab keine Begleitforschung und mit Ausnahme weniger Arbeiten zur sogenannten „Präzisionszielfunktion“ keine wissenschaftlichen Arbeiten speziell zur Methodik des Zensus 2022. Im Rahmen des Gutachtens wurde die Literatur vollständig gesichtet, Dutzende Anfragen an die statistischen Ämter gestellt und schließlich Anfragen nach dem Informationsfreiheitsgesetz gestellt. Zusätzlich wurden auf der Basis anonymisierter Mikrodaten Simulationen, insbesondere zum Hochrechnungsverfahren, durchgeführt. Aufgrund der anhängigen Klagen wurde das ca. 100 Seiten lange Gutachten bislang nicht veröffentlicht. Der Vortrag wird die Kernpunkte des Gutachtens erstmals öffentlich vorstellen.

177 Die Illusion der Automatisierung des Zensus

Rainer Schnell, Severin V. Weiland

Universität Duisburg-Essen, Deutschland
severin.weiland@uni-due.de

SECTION: Methodology of Statistical Surveys (Session 5)

Zurzeit arbeitet das Statistische Bundesamt an der Umstellung des Zensusverfahrens, von einem registergestützten Verfahren hin zu einem vollständigen Registerzensus. Im Vergleich zum registergestützten Verfahren, bei dem eine 10%-Stichprobe der Bevölkerung befragt wurde, um Fehler insbesondere in den Einwohnermelderegistern zu korrigieren, soll das neue Verfahren nur noch auf Basis von Registern funktionieren. Die Register sollen dabei anhand einer eindeutigen Identifikationsnummer (IDNr) verknüpft werden. Um die IDNr jedoch in ein vorhandenes Register einzuführen, muss dieses Register mit einem anderen Register mittels Record-Linkage verknüpft werden. Insbesondere bei dezentralen Registern, wie den Melderegistern, muss darüber hinaus sichergestellt werden, dass das konsolidierte Register frei von Duplikaten ist.

Das Ziel des Projekts ist, die einzelnen Aspekte des Registerzensus systematisch zu simulieren. Der Vortrag behandelt die Simulation des zentralen Schritts der Mehrfachfallprüfung – die Deduplizierung des konsolidierten Melderegisters. Hierzu wurde zunächst eine synthetische Bevölkerung

der deutschen Bevölkerung auf Basis der Ergebnisse des Zensus 2011 simuliert. Der Datensatz umfasst ca. 80 Mio. geolokalisierte Fälle. Bei der Simulation wurden die Haushaltsstrukturen sowie die Region, das Alter und die Nationalität bei den Namensverteilungen und Geburtsdaten berücksichtigt. Bisherige Simulationen vernachlässigen diese Abhängigkeiten immer, so dass die Leistungsfähigkeit automatischer Verfahren überschätzt wird.

Aus der synthetischen Bevölkerung wurde ein konsolidiertes Melderegister abgeleitet und anschließend mit Fehlern auf der Basis empirischer Verteilungen aus administrativen Daten behaftet. Dabei wurden unterschiedliche Annahmen über die Datenqualität getroffen. Die Deduplizierung zeigt, dass ein großer Teil der Duplikate durch kein Verfahren gefunden werden kann. Die Ergebnisse zeigen besonders, dass die für den Zensus verwendeten Kernmerkmale häufig zur zweifelsfreien Identifikation nicht ausreichen, insbesondere wenn Datenfehler vorliegen.

178 Simultaneous Inference Bands for Autocorrelations

Uwe Hassler (1), **Marc-Oliver Pohle** (2,3), Tanja Zahn (1)

1: Goethe Universität Frankfurt, Deutschland; 2: Bergische Universität Wuppertal, Deutschland; 3: Heidelberger Institut für Theoretische Studien, Deutschland
pohle@uni-wuppertal.de

SECTION: Statistical Theory and Methods (Session 5)

Sample autocorrelograms typically come with significance bands (non-rejection regions) for the null hypothesis of no temporal correlation. These bands have two shortcomings. First, they build on pointwise intervals and suffer from joint undercoverage (overrejection) under the null hypothesis. Second, if this null is clearly violated one would rather prefer to see confidence bands to quantify estimation uncertainty. We propose and discuss both simultaneous significance bands and simultaneous confidence bands for time series and series of regression residuals. They are as easy to construct as their pointwise counterparts and at the same time provide an intuitive and visual quantification of sampling uncertainty as well as valid statistical inference. For regression residuals, we show that for static regressions the asymptotic variances underlying the construction of the bands are the same as those for observed time series, and for dynamic regressions (with lagged endogenous regressors) we show how they need to be adjusted. We study theoretical properties of simultaneous significance bands and two types of simultaneous confidence bands (sup-t and Bonferroni) and analyse their finite-sample performance in a simulation study. Finally, we illustrate the use of the bands in an application to monthly US inflation and residuals from Phillips curve regressions.

179 Novel Stein-type Characterizations of Bivariate Count Distributions with Applications

Shaochen Wang (1), **Christian Weiß** (2)

1: South China University of Technology, Guangzhou, P. R. China; 2: Helmut-Schmidt-Universität, Hamburg, Deutschland
weissc@hsu-hh.de

SECTION: Statistical Theory and Methods (Session 5)

The derivation and application of Stein identities have received considerable research interest in recent years, especially for continuous or discrete-univariate distributions. In this paper, we

complement the existing literature by deriving and investigating Stein-type characterizations for the three most common types of bivariate count distributions, namely the bivariate Poisson, binomial, and negative-binomial distribution. Then, we demonstrate the practical relevance of these novel Stein identities by a couple of applications, namely the deduction of sophisticated moment expressions, of flexible goodness-of-fit tests, and of novel tests for the symmetry of bivariate count distributions. The paper concludes with an analysis of real-world data examples.

References:

Kocherlakota, S. and Kocherlakota, K. (2014). Bivariate discrete distributions. In N. Balakrishnan et al. (eds): Wiley StatsRef: Statistics Reference Online, stat00972.

Sudheesh, K.K. and Tibiletti, L. (2012). Moment identity for discrete random variable and its applications. *Statistics*, 46(6):767–775.

Wang, S. and Weiß, C.H. (2026). Novel Stein-type characterizations of bivariate count distributions with applications. *arXiv preprint* 2602.23775.

Weiß, C.H. and Aleksandrov, B. (2022). Computing (bivariate) Poisson moments using Stein–Chen identities. *The American Statistician*, 76(1):10–15.

180 The Combined INAR Model for Count Random Fields

Angelika Silbernagel, Christian H. Weiß

Helmut-Schmidt-Universität, Deutschland
silbernagel@hsu-hh.de

SECTION: Statistical Theory and Methods (Session 5)

We propose the combined integer-valued autoregressive (CINAR) random field model for count data observed on a regular lattice. The model can be viewed as a planar analogue of the CINAR time series model, as well as an extension of the first-order INAR random field model to a higher-order dependence structure.

We study its fundamental stochastic properties and consider methods for parameter estimation. The theoretical results are complemented by simulations and an application to real data.

181 Bootstrap and Minimum Distance Methods for Spatio-Temporal Counting Processes

Mirko Alexander Jakubzik

Technische Universität Dortmund, Deutschland
mirko.jakubzik@tu-dortmund.de

SECTION: Statistical Theory and Methods (Session 5)

Spatio-temporal counting processes arise naturally in a wide range of applications such as epidemiology, finance, and information propagation. We study multivariate counting processes defined on simple directed networks, where each node represents one component of the process and events may trigger further events along network edges. Important examples include self-exciting Hawkes-type processes, which generate temporal and spatial dependence through local excitation effects.

Our focus lies on parametric models for the conditional cumulative intensity, the Doob-Meyer compensator, of such counting processes. Compensating a counting process by its conditional

cumulative intensity at the true parameter yields a martingale. Using i.i.d. repetitions of the network process, we construct pooled martingales that serve as our basis for inference. Evaluating squared martingale increments at the jump times of the counting processes and aggregating them over the observation horizon leads to Cramér–von Mises type statistics whose minimization enables minimum distance estimation. While these estimators are known to be consistent and asymptotically normal under correct specification, they often behave poorly under model misspecification.

Motivated by these limitations, we investigate bootstrap-based inference procedures for network-valued counting processes, with particular emphasis on wild bootstrap methods adapted to the martingale difference structure. Our aim is to improve finite-sample calibration for goodness-of-fit testing and confidence procedures in strongly dependent settings. We compare several resampling approaches and study their performance relative to classical likelihood-based methods and standard asymptotic approximations. Particular attention is paid to the behavior of the proposed procedures under model misspecification and deviations from the assumed excitation structure. The resulting framework combines martingale methods, minimum distance estimation, and resampling techniques for statistical inference in spatial-temporal counting processes on networks.

Social Programme:

For activities of the social programme, prior registration via ConfTool is necessary. Please note the respective registration deadline and the maximum number of participants. For late registration please send an e-mail to statistische-woche@dstatg.de.

- **Monday, Sep 08, 2026:**

Casual Get-Together – Bootshaus Essen

– 6:30 pm

Meeting point: Freiherr-vom-Stein-Str. 204a, 45133 Essen

Reservation only. Participants cover their own expenses.

- **Tuesday, Sep 09, 2026:**

Guided Tours at Zollverein Coal Mine and Coking Plant – 7 pm

Duration: 1.5 hours

Über Kohle und Kumpel

Meeting point: Gelsenkirchener Str. 181, 45309 Essen

Language: German

Grubenlicht und Wetterzug

Meeting point: Gelsenkirchener Str. 181, 45309 Essen

Language: German

Von Kohle, Koks und harter Arbeit

Meeting point: Arendahls Wiese, 45141 Essen

Language: German

- **Wednesday, Sep 10, 2025:**

Conference Dinner – Weststadthalle Essen – 6 pm

Meeting point: Thea-Leymann-Str. 23, 45127 Essen

Participation fee: free of charge for participants of the Statistical Week 2026. Drinks are covered only until 9 pm.

Rahmenprogramm:

Für Aktivitäten des Rahmenprogramms ist eine vorherige Anmeldung über das ConfTool notwendig. Bitte beachten Sie den jeweiligen Anmeldeschluss bzw. die maximale Teilnehmerzahl. Zur nachträglichen Anmeldung senden Sie bitte eine E-Mail an statistische-woche@dstatg.de.

- **Montag, 08.09.2026:**

Gemütliches Beisammensein – Bootshaus Essen – 18:30 Uhr

Treffpunkt: Freiherr-vom-Stein-Str. 204a, 45133 Essen

Nur Reservierung. Teilnehmende tragen ihre Kosten selbst.

- **Dienstag, 09.09.2026:**

Führung Zeche Zollverein – 19:00 Uhr

Dauer: 1.5 Stunden

Über Kohle und Kumpel

Treffpunkt: Gelsenkirchener Str. 181, 45309 Essen

Sprache: Deutsch

Grubenlicht und Wetterzug

Treffpunkt: Gelsenkirchener Str. 181, 45309 Essen

Sprache: Deutsch

Von Kohle, Koks und harter Arbeit

Treffpunkt: Arendahls Wiese, 45141 Essen

Sprache: Deutsch

- **Mittwoch, 10.09.2026:**

Conference Dinner – Weststadthalle Essen – 18:00 Uhr

Treffpunkt: Thea-Leymann-Str. 23, 45127 Essen

Teilnahmegebühr: für Teilnehmende der Statistischen Woche 2026 kostenfrei. Getränke sind nur bis 21:00 Uhr inkludiert.

Participants / Teilnehmende

Adunts, Davit; Institute for Employment Research (IAB) of the German Federal Employment Agency (BA); Davit.Adunts@iab.de
Akman, Muhammet; IT.NRW; muhammet.akman@it.nrw.de
Alfken, Christoph; IT.NRW - Landesamt für Statistik Nordrhein-Westfalen; christoph.alfken@it.nrw.de
Andrä, Diana; Stadt Dortmund, Dortmunder Statistik; dandrae@stadtdo.de
Balczynski, Oliver; IT.NRW; oliver.balczynski@it.nrw.de
Barthel, Edwin; TU Dortmund Fakultät Statistik; barthel@statistik.tu-dortmund.de
Bauer, Lukas; University of Freiburg; lukas.bauer@vwl.uni-freiburg.de
Baumstark, Johannes; Bayerisches Landesamt für Statistik; johannes.baumstark@statistik.bayern.de
Becker, Angelika; Verband der Chemischen Industrie e.V.; becker@vci.de
Berend, Lukas; University of Hagen; lukas.berend@fernuni-hagen.de
Berger, Theo; Hochschule Hannover; theo.berger@hs-hannover.de
Bergrab, Michael; Leibniz-Institut für Bildungsverläufe; michael.bergrab@lifbi.de
Bettels, Timm; Stadt Wolfsburg; timm.bettels@stadt.wolfsburg.de
Betting, Leia; Technische Universität Dortmund; leia.betting@tu-dortmund.de
Bilgic, Daniel; Bayerisches Landesamt für Statistik; daniel.bilgic@statistik.bayern.de
Birke, Melanie; Universität Bayreuth; melanie.birke@uni-bayreuth.de
Blätgen, Nadine; BBSR; nadine.blaetgen@bbr.bund.de
Bleninger, Sara; Bayerisches Landesamt für Statistik; sara.bleninger@statistik.bayern.de
Bonny, Carlotta; Regionalverband Ruhr; fortbildungen@rvr.ruhr
Boswijk, Peter; University of Amsterdam; h.p.boswijk@uva.nl
Braakmann, Albert; Statistisches Bundesamt a.D.; alwies15@gmail.com
Brand, Ruth; Statistisches Bundesamt; ruth.brand@destatis.de
Brenzel, Hanna; Statistisches Bundesamt; Hanna.Brenzel@destatis.de
Bruckmeier, Kerstin; Institut für Arbeitsmarkt- und Berufsforschung (IAB); kerstin.bruckmeier@iab.de
Buchholz, Christine; Hochschule Bonn-Rhein-Sieg; christine.buchholz@h-brs.de
Buhl, Yannik; Statistisches Bundesamt; yannik.buhl@destatis.de
Bündgens, Richard; Information und Technik Nordrhein-Westfalen; richard.buendgens@it.nrw.de
Buscher, Lioba; Statistisches Landesamt des Freistaates Sachsens; Lioba.Buscher@statistik.sachsen.de
Buschner, Andrea; Bayerisches Landesamt für Statistik; andrea.buschner@statistik.bayern.de
Camehl, Annika; Erasmus University Rotterdam; camehl@ese.eur.nl
Carius-Munz, Lisa; Statistisches Bundesamt; lisa.carius-munz@destatis.de
Cekaj, Enjeda; University of Augsburg; enjeda.cekaj@uni-a.de
Chlumsky, Jürgen; Statistisches Bundesamt (ehem.); chlumsky.harttmann@t-online.de
Chomon, Sultan Mahmud; University of Duisburg-Essen; sultan.chomon@uni-due.de
Chrysopoulou Tseva, Kleio; Humboldt-Universität zu Berlin; kleio.chrysopoulou.tseva@hu-berlin.de
Cramer, Katharina; IT.NRW; katharina.cramer@it.nrw.de
Czudaj, Robert L.; TU Freiberg; robert-lukas.czudaj@vwl.tu-freiberg.de
Demetrescu, Matei; TU Dortmund; mdeme@statistik.tu-dortmund.de
Deschermeier, Philipp; Institut der deutschen Wirtschaft; deschermeier@iwkoeln.de
Diefenbacher, Julian; University of Hamburg & GESIS – Leibniz Institute for the Social Sciences; julian.diefenbacher@gesis.org
Dimitriadis, Timo; Goethe Universität Frankfurt; Dimitriadis@econ.uni-frankfurt.de
Dittrich, Stefan; Statistisches Bundesamt; stefan.dittrich@destatis.de
Dörr, Philip; Ruhr-Universität Bochum; philip.doerr@rub.de
Dröge, Lennart; Universität Augsburg; lennart.droege@uni-a.de
Drozdek, Louise; University of Duisburg-Essen; louise.drozdek@stud.uni-due.de
Dumpert, Florian; Statistisches Bundesamt; florian.dumpert@destatis.de

Dyckerhoff, Rainer; Universität zu Köln; rainer.dyckerhoff@statistik.uni-koeln.de
Dymarz, Martin; Stadt Ratingen; martin.dymarz@ratingen.de
Dzikowski, Daniel; TU Dortmund; daniel.dzikowski@tu-dortmund.de
Eichfuss, Silas; Esri Deutschland GmbH; s.eichfuss@esri.de
Eilers, Marie; WZB; marie.eilers@ratswd.de
Einhorn, Johanna; Universität Bamberg; johanna.einhorn@uni-bamberg.de
El Mahtout, Btissame; Duisburg-Essen University and TU Dortmund; btissame.el@tu-dortmund.de
Elsner, Svenja; SV Wissenschaftsstatistik gGmbH; svenja.elsner@stifterverband.de
Ernst, Jessica; SV gemeinnützige Gesellschaft für Wirtschaftsstatistik mbH; jessica.ernst@stifterverband.de
Fessler, Pirmin; Oesterreichische Nationalbank; pirmin.fessler@oenb.at
Fioravanti, Antonio; TU Dortmund; antonio.fioravanti@tu-dortmund.de
Fischer, Boris; Statistisches Amt, Landeshauptstadt München; boris.fischer@muenchen.de
Forster, Michael; ehemals IT NRW; family_forster@t-online.de
Franjić, Igor; Statistisches Bundesamt; igor.franjic@destatis.de
Frese, Maria; IT.NRW; Maria.Frese@it.nrw.de
Fried, Roland; TU Dortmund; fried@statistik.tu-dortmund.de
Fries, Jan; Amt für Statistik und Stadtforschung Wiesbaden; dr.jan.fries@wiesbaden.de
Fröhlich, Maximilian; Statistisches Amt für Hamburg und Schleswig-Holstein; maximilian.froehlich@statistik-nord.de
Fürrnrohr, Michael; Universität Bamberg; michael.fuerrnrohr@statistik.bayern.de
Germakowsky, Gerd; Privat; gerd.germakowsky@online.de
Göhlich, Maximilian; SV Wissenschaftsstatistik gGmbH; maximilian.goehlich@stifterverband.de
Golosnoy, Vasyli; Ruhr-Universität Bochum; vasyli.golosnoy@rub.de
Gondrom, Sabine; go2data; contact@go2data.de
Görz, Sheila; Technische Universität Dortmund; sheila.goerz@tu-dortmund.de
Greger, Tim; Universität Bayreuth; tim.greger@uni-bayreuth.de
Groh, Ariane; Stadt Bochum; agroh@bochum.de
Große, Franziska; Landesamt für Statistik Niedersachsen; franziska.grosse@statistik.niedersachsen.de
Grüne, Anna; IT.NRW; anna.gruene@it.nrw.de
Guljanov, Gaygysyz; GEFA BANK GmbH; gguljanow40@gmail.com
Haid, Philipp; Universität Augsburg; philipp1.haid@uni-a.de
Hajiyev, Murad; Statistisches Bundesamt; murad.hajiyev@destatis.de
Halbleib, Roxana; University of Freiburg; roxana.halbleib@vwl.uni-freiburg.de
Hanck, Christoph; University of Duisburg-Essen; christoph.hanck@vwl.uni-due.de
Harms, Torsten; DHBW Karlsruhe; harms@dhbw-karlsruhe.de
Harris, Emily; Statistikstelle Stadt Neuss; Emily.Harris@stadt.neuss.de
Hauf, Stefan; Statistisches Bundesamt; stefan.hauf@destatis.de
Haupt, Harry; Universität Passau; harry.haupt@uni-passau.de
Hauser, Lizanne; Esri Deutschland; l.hauser@esri.de
Heiden, Sebastian; DStatG; heiden@dstatg.de
Heinrichs, Florian; FH Aachen; f.heinrichs@fh-aachen.de
Helmers, Viola; RWI - Leibniz Institute for Economic Research; vhelmers@rwi-essen.de
Hense, Svenja; Stadt Bochum; shense@bochum.de
Herkner, Thomas; Privat; thomasherker@outlook.com
Heuser, Annelie; Landeshauptstadt Wiesbaden; annelie.heuser@wiesbaden.de
Hildenbrand, Andreas; IHK Rhein-Neckar; andreas.hildenbrand@rhein-neckar.ihk24.de
Hoffmann, Hanna; Statistisches Landesamt Nordrhein-Westfalen; hanna.hoffmann@it.nrw.de
Hofmann, Bernd; Bundesagentur für Arbeit; Bernd.Hofmann5@arbeitsagentur.de
Holzendorf, Volker; Stadtverwaltung Jena; volker.holzendorf@jena.de
Hoymann, Jana; Bundesinstitut für Bau-, Stadt- und Raumforschung; jana.hoymann@bbr.bund.de

Huang, Wenzhen; Universität Duisburg-Essen; wenzhen.huang@uni-due.de
Hützen, Dijana; Landesbetrieb Information und Technik NRW - Statistisches Landesamt; dijana.huetzen@it.nrw.de
Jäger, Solveigh; BDI; s.jaeger@bdi.eu
Jahn, Malte; Hochschule Merseburg; malte.jahn@hs-merseburg.de
Jakubzik, Mirko Alexander; Technische Universität Dortmund; mirko.jakubzik@tu-dortmund.de
Jentsch, Carsten; TU Dortmund University; jentsch@statistik.tu-dortmund.de
Jeschke, Marco; Technische Universität Dortmund; marco.jeschke@tu-dortmund.de
Johannssen, Arne; Harz University of Applied Sciences; ajohannssen@hs-harz.de
Kalinowski, Michael; Bundesinstitut für Berufsbildung (BIBB); kalinowski@bibb.de
Kanistras, Henrike; Stadt Witten; henrike.kanistras@stadt-witten.de
Kaschowitz, Judith; Bundesinstitut für Bau-, Stadt- und Raumforschung (BBSR); judith.kaschowitz@bbr.bund.de
Kateri, Maria; RWTH Aachen University; maria.kateri@rwth-aachen.de
Kauermann, Göran; Ludwig-Maximilians-Universität München; goeran.kauermann@LMU.DE
Kaveriappa-von Ramin, Beate; KOSIS-Gemeinschaft SIKURS; beate.kaveriappa@stadt.nuernberg.de
Keller, Manuel; Stadt Dortmund; mkeller@stadtdo.de
Kern, Christoph; LMU Munich; Christoph.Kern@stat.uni-muenchen.de
Kiess, Johannes; Universität Leipzig; johannes.kiess@uni-leipzig.de
Kilinc, Mustafa; University of Cologne; m.kilinc@wiso.uni-koeln.de
Kladroba, Andreas; Stifterverband für die Deutsche Wissenschaft; andreas.kladroba@stifterverband.de
Klapheck, Jochen; Stadt Hagen; jochen.klapheck@stadt-hagen.de
Klein, Stella; Bayerisches Landesamt für Statistik; Stella.Klein@statistik.bayern.de
Kley, Tobias; University of Göttingen; tobias.kley@uni-goettingen.de
Klinke, Sigbert; Humboldt-Universität zu Berlin; sigbert@wiwi.hu-berlin.de
Koehlmann, Maren; Statistisches Bundesamt; maren.koehlmann@destatis.de
Konrad, Anne; Leibniz-Institut für Bildungsverläufe; anne.konrad@lifbi.de
Korell, Janosch; DHBW KA; janosch.korell@dhbw-karlsruhe.de
Krabbe, Dorothea; Bundesinstitut für Berufsbildung; Dorothea.Krabbe@bibb.de
Kraus, Petra; Hauptverband der Deutschen Bauindustrie; petra.kraus@bauindustrie.de
Kreklow, Jennifer; Stadt Wolfsburg; jennifer.kreklow@stadt.wolfsburg.de
Kremer, René; Destatis; rene.kremer@destatis.de
Kück, Jannis; Heinrich-Heine-Universität Düsseldorf; kueck@dice.hhu.de
Kudjawu, Alexandra Fafali; Europa-Universität Viadrina; kudjawu@europa-uni.de
Kuhn, Sarah; Institut für Arbeitsmarkt- und Berufsforschung; sarah.kuhn@iab.de
Kürbis, Ilka; LH München; ilka@icara.de
Lange, Kai-Robin; TU Dortmund; kai-robin.lange@tu-dortmund.de
Lange, Tatjana; HS Merseburg; tanja.28.lange@gmail.com
Laudy, Daniel; FOM; daniel.laudy@bcw-gruppe.de
Lehmann, Niklas Valentin; TU Bergakademie Freiberg; niklasl.2306@gmail.com
Lenz, Jürgen; Kreis Mettmann; juergen.lenz@kreis-mettmann.de
Lestrade, Ariane; Destatis; ariane.lestrade@destatis.de
Leukert, Valerie; Bayerisches Landesamt für Statistik; DyanneValerie.Leukert@statistik.bayern.de
Leyendecker, Verena; SV Wissenschaftsstatistik gGmbH; verena.leyendecker@stifterverband.de
Liesenfeld, Roman; Universität zu Köln; liesenfeld@statistik.uni-koeln.de
Limberg, Heiko; Statistisches Bundesamt; heiko.limberg@destatis.de
Linz, Stefan; Statistisches Bundesamt; stefan.linz@destatis.de
Löffler, Natascha; IT.NRW; natascha.loeffler@it.nrw.de
Ludsteck, Johannes; IAB; johannes.ludsteck@online.de
Lunge, Jan; Stadt Wolfsburg; jan.lunge@stadt.wolfsburg.de
Mai, Christoph-Martin; Statistisches Bundesamt; christoph-martin.mai@destatis.de

Maier, Tobias; Bundesinstitut für Berufsbildung; tobias.maier@bibb.de
Mancini, Cecilia; University of Verona; cecilia.mancini@univr.it
Mannigel, Alice; Statistisches Amt für Hamburg und Schleswig-Holstein; alice.mannigel@statistik-nord.de
Maretzke, Steffen; BBSR; steffen.maretzke@bbr.bund.de
Martinez Hernandez, Catalina; European Central Bank; Catalina.Martinez_Hernandez@ecb.europa.eu
Mattes, Mara; Heinrich-Heine-Universität Düsseldorf; mattes@dice.hhu.de
Maxand, Simone; European University Viadrina; maxand@europa-uni.de
Meer, Uwe; Stadt Wolfsburg; uwe.meer@stadt.wolfsburg.de
Meyer, Daniel; Bundesinstitut für Bau-, Stadt- und Raumforschung; daniel.meyer@bbr.bund.de
Mireille, Fangueng; Leibniz University of Hannover; Mireille.Fangueng@ikg.uni-hannover.de
Missong, Martin; Universität Bremen; missong@uni-bremen.de
Mohamed, Hoda; Statistisches Bundesamt; hoda.mohamed@destatis.de
Moreira Lara, Ignacio; Universität Duisburg Essen; ignacio.moreira-lara@tu-dortmund.de
Moritz, Steffen; Statistisches Bundesamt; steffen.moritz@destatis.de
Mosler, Karl; Universität zu Köln; kmosler@uni-koeln.de
Mühlhan, Jannek; Statistisches Bundesamt; jannek.muehlhan@destatis.de
Müller, Christine; TU Dortmund; cmueller@statistik.tu-dortmund.de
Musta, Eni; University of Amsterdam; e.musta@uva.nl
Muth, Mareike; Statistisches Bundesamt; mareike.muth@destatis.de
Nagel, Jule; Stadt Wolfsburg; jule.nagel@stadt.wolfsburg.de
Neubert, Agatha; Stadt Witten; agatha.neubert@stadt-witten.de
Neumann, Uwe; RWI-Leibniz-Institut für Wirtschaftsforschung; uwe.neumann@rwi-essen.de
Noack, Marcel; IT.NRW; marcel.noack@it.nrw.de
Nölting, Christopher; Statistisches Bundesamt; christopher.noelting@destatis.de
Nüsing, Linus; Universität Konstanz; linus.nuesing@uni-konstanz.de
Ochsmann, Andrea; Stadt Recklinghausen; andrea.ochsmann@recklinghausen.de
Oess, Eva-Maria; University of Cologne; oess@wiso.uni-koeln.de
Okhrin, Yarema; University of Augsburg; yarema.okhrin@wiwi.uni-augsburg.de
Orlowski, Miriam; Landesamt für Statistik Bayern; miriam.orldowski@statistik.bayern.de
Otto, Philipp; University of Glasgow; philipp.otto@glasgow.ac.uk
Overmann, Gina Lena; IT.NRW; ginalena.overmann@it.nrw.de
Ovsyannikova, Daria; TU Dortmund; daria.ovsyannikova@tu-dortmund.de
Özer, Yeliz; Leibniz Universität Hannover; oezer@statistik.uni-hannover.de
Pappert, Sven; Ruhr University Bochum; sven.pappert@rub.de
Peters, Jonas; ETH Zurich; jonas.peters@stat.math.ethz.ch
Petruk, Viktoriia; Europa-Universität Viadrina; petruk@europa-uni.de
Piffer, Michele; Bank of England, KCL; m.b.piffer@gmail.com
Pinson, Pierre; Imperial College London; p.pinson@imperial.ac.uk
Pinto, Rebecca; Bürgeramt, Statistik und Wahlen, Stadt Frankfurt; rebecca.pinto@stadt-frankfurt.de
Plagens, Manfred; Stadt Würzburg; manfred.plagens@stadt.wuerzburg.de
Plöger, Amelie; Landesbetrieb Information und Technik NRW; Amelie.Ploeger@it.nrw.de
Pohle, Marc-Oliver; Bergische Universität Wuppertal; pohle@uni-wuppertal.de
Prüser, Jan; EBS University; J.Pruesser@gmx.net
Puschmann, Timo; Universität Duisburg-Essen; timo.puschmann@vwl.uni-due.de
Raatz, Philip; RWI - Leibniz-Institut für Wirtschaftsforschung; philip.raatz@rwi-essen.de
Rabenau, Kai; Manage Now GmbH; kai.rabenau@manage-now.de
Rabenseifner, Jan; Universität Hamburg; jan.rabenseifner@uni-hamburg.de
Rauwolf, Diana; RWTH Aachen; rauwolf@isw.rwth-aachen.de
Redlich, Sarah; Landesbetrieb Information und Technik NRW - Statistisches Landesamt; sa-

rah.redlich@it.nrw.de

Reichle, Hannes; Humboldt Universität zu Berlin; hannes.reichle31@gmail.com

Reincke, Ulrich; SAS Institute; Ulrich.Reincke@ger.sas.com

Rendtel, Ulrich; Freie Universität Berlin; ulrich.rendtel@fu-berlin.de

Rengers, Martina; Statistisches Bundesamt; martina.rengers@destatis.de

Rhoades, Julia; Stadt Ratingen; julia.rhoades@ratingen.de

Rigbers, Anke; Statistisches Landesamt Baden-Württemberg; anke.rigbers@stala.bwl.de

Rohde, Johannes; IT.NRW - Statistisches Landesamt; johannes.rohde@it.nrw.de

Rohr, Björn; GESIS - Leibniz Institute for Social Sciences; Bjoern.rohr@gesis.org

Rottwinkel, Mareen; IT.NRW; mareen.rottwinkel@it.nrw.de

Ruß-Obajtek, Uwe; Statistisches Landesamt Baden-Württemberg; uwe.russ-obajtek@stala.bwl.de

Salwiczek, Christian; Statistikamt Nord; Christian.salwiczek@statistik-nord.de

Savvidou, Paraskevi; Universitätsklinikum Essen (AöR); paraskevi.savvidou@uk-essen.de

Schäfer, Josef; Privat; jugschaefer@aol.com

Schaumburg, Julia; Vrije Universiteit Amsterdam; j.schaumburg@vu.nl

Schepers, Marianne; Statistisches Bundesamt; marianne.schepers@destatis.de

Schiesberg, Sara; Statistisches Bundesamt; sara.schiesberg@destatis.de

Schmandt, Marco; TU Berlin; m.schmandt@tu-berlin.de

Schmid, Wolfgang; Europa-Universität; milesle@icloud.com

Schmidt, Rolf; Universität Duisburg-Essen; schmidt.Ob@gmx.de

Schmidt, Fabian; TU Dortmund University; f.schmidt@statistik.tu-dortmund.de

Schmidt, Colin; Information und Technik Nordrhein-Westfalen; colin.schmidt@it.nrw.de

Schmidtke, Kerstin; IT.NRW - Statistisches Landesamt NRW; kerstin.schmidtke@it.nrw.de

Schmiedt, Anja; OTH Regensburg; anja.schmiedt@oth-regensburg.de

Schmitt, Johannes; SV gemeinnützige Gesellschaft für Wirtschaftsstatistik mbH; johannes.schmitt@stifterverband.org

Schmitz, Theresa M. A.; Heinrich Heine University Düsseldorf; theresa.schmitz@hhu.de

Schneider, Nicholas-Sebastian; Statistisches Bundesamt; nicholas-sebastian.schneider@destatis.de

Schneider, Gaby; Universität Frankfurt; schneider@math.uni-frankfurt.de

Schnurbus, Joachim; Universität Passau; joachim.schnurbus@uni-passau.de

Schröder, Leonie; Amt für Statistik und Stadtforschung; leonie.schroeder@wiesbaden.de

Schüller, Katharina; STAT-UP Statistical Consulting & Data Science GmbH; katharina.schueller@stat-up.com

Schüller, Daniela; Stadtverwaltung Koblenz; Daniela.Schueller@stadt.koblenz.de

Schultz, Andrea; Stadt Leipzig - Amt für Statistik und Wahlen; andrea.schultz@leipzig.de

Schulz, Christian; Universität Duisburg-Essen; christian.schulz@vwl.uni-due.de

Schulz, Julian; Landesamt für Statistik Niedersachsen; julian.schulz@statistik.niedersachsen.de

Schulze, Kathrin; IT.NRW; kathrin.schulze@it.nrw.de

Schürmann, Daniel; Stadt Dortmund; dstatg@dostat.de

Schüssler, Rainer Alexander; Aalborg University Business School; raineralexanderschuessler@gmail.com

Schütz, Martin; SAS Institute GmbH; martin.schuetz@sas.com

Schweitzer, Michael; TU Dortmund University; michael.schweitzer@tu-dortmund.de

Schwengler, Barbara; Institut für Arbeitsmarkt- und Berufsforschung; Barbara.Schwengler@iab.de

Seidel, Gerald; c/o Bundesagentur für Arbeit; gerald-seidel@web.de

Seifert, Ingo; Stadt Leipzig, Amt für Statistik und Wahlen; ingo.seifert@leipzig.de

Sibbertsen, Philipp; Leibniz Universität Hannover; sibbertsen@statistik.uni-hannover.de

Sieger, Lisa; Universität Duisburg-Essen; lisa.sieger@uni-due.de

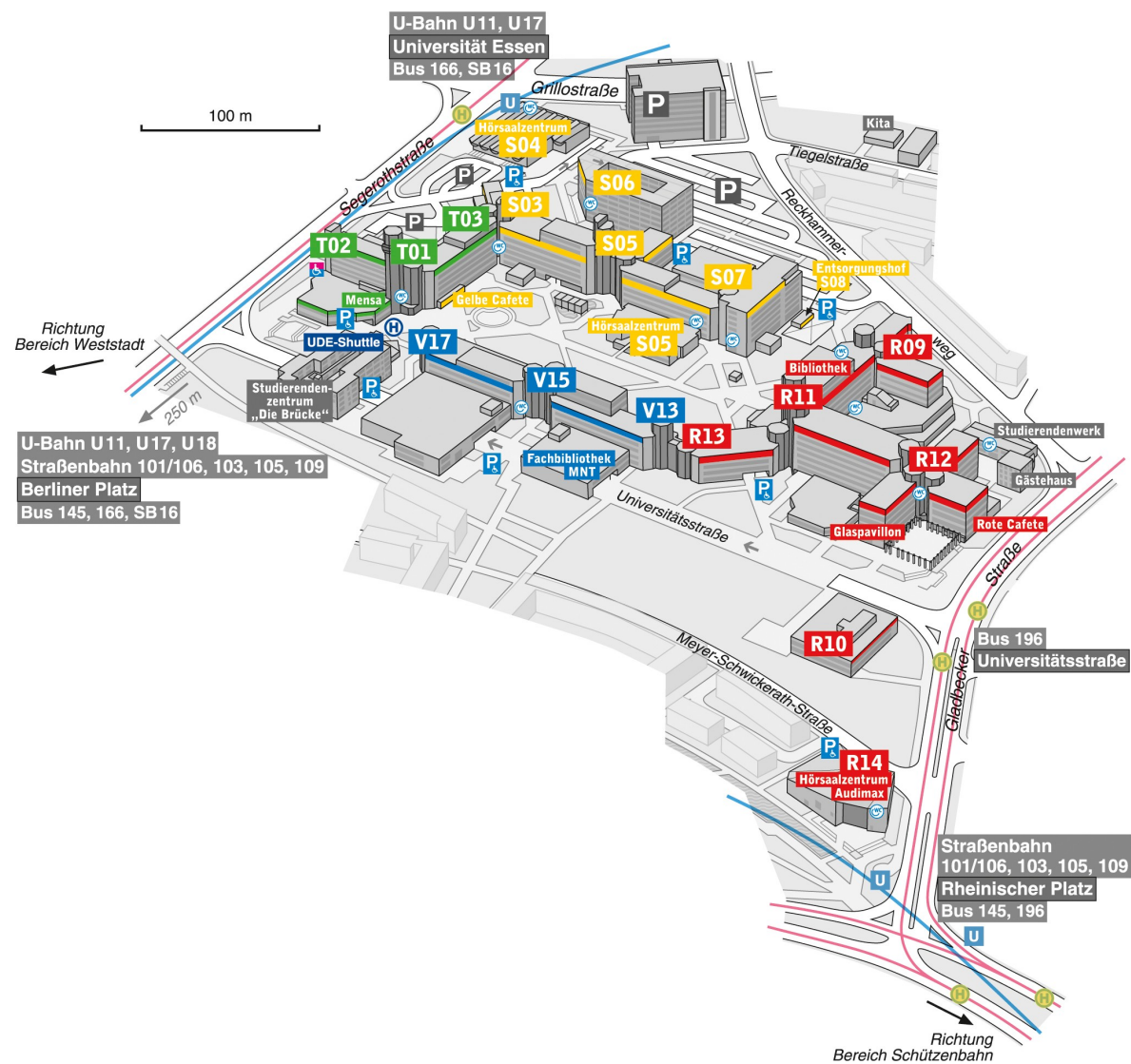
Simota, Jennifer; Otto-Friedrich-Universität Bamberg; jennifer.simota@uni-bamberg.de

Sonnenburg, Anja; Gesellschaft für Wirtschaftliche Strukturforschung mbH; sonnenburg@gws-os.com

Speich, Wolf-Dietmar; Statistisches Landesamt des Freistaates Sachsen; wolf-dietmar.speich@statistik.sachsen.de

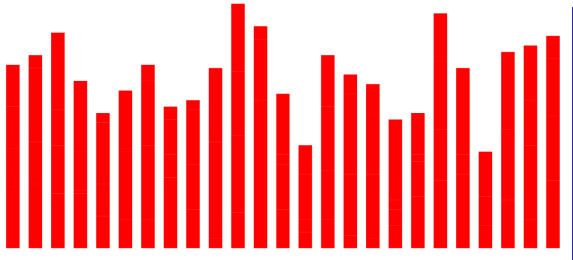
Spies, Julian; Universität zu Köln; julian.spies@wiso.uni-koeln.de
Spieß, C. Katharina; Bundesinstitut für Bevölkerungsforschung (BiB); Direktorin@bib.bund.de
Steland, Ansgar; RWTH Aachen University; steland@stochastik.rwth-aachen.de
Stender, Axel; Stadt Moers; axel.stender@moers.de
Strebel, Alex; Physikalisch Technische Bundesanstalt Berlin; alex.strebel@ptb.de
Stürmer, Elisabeth; BBSR; Elisabeth.Stuermer@bbr.bund.de
Suffert, Martin; manage now GmbH; martin.suffert@manage-now.de
Thees, Nicole; Stadt Trier; nicole.thees@trier.de
Thien-Seitz, Uta; Statistisches Amt der LH München; uta.thien@muenchen.de
Thobe, Ines; GWS Osnabrück; thobe@gws-os.com
Tillmann, Isabelle; Kreis Gütersloh; i.tillmann@kreis-guetersloh.de
Vallizadeh, Ehsan; Statistik der BA; ehsan.vallizadeh@arbeitsagentur.de
Vance, Colin; RWI - Leibniz Institute for Economic Research; vance@rwi-essen.de
Vance, Liam; Independent Researcher; liam.jvance@web.de
Vanella, Patrizio; aQua-Institut; patrizio.vanella@aqua-institut.de
Verneuer-Emre, Lena; Stadt Münster; verneuer-emre@stadt-muenster.de
Voit, Falk; Landesamt für Statistik Niedersachsen; falk.voit@statistik.niedersachsen.de
von Eschwege, Katja; Statistisches Bundesamt; katja.eschwege@destatis.de
von Kalckreuth, Ulf; Deutsche Bundesbank; ulf.von-kalckreuth@bundesbank.de
Vorgrimler, Daniel; Statistisches Bundesamt; daniel.vorgrimler@destatis.de
Wagner, Stephan; Statistikamt Nord; stephan.wagner@statistik-nord.de
Waschipky, Martin; Amt für Statistik und Wahlen, Stadt Leipzig; martin.waschipky@leipzig.de
Waßenberg, Sebastian; IT.NRW - Statistisches Landesamt; sebastian.wassenberg@it.nrw.de
Weber, Nina; Sonares; n.weber@sonares.org
Wegner, Susanne; Statistisches Bundesamt; susanne.wegner@destatis.de
Weiand, Severin V.; Universität Duisburg-Essen; severin.weiand@uni-due.de
Weigert, Alexander; Statistisches Bundesamt; alexander.weigert@destatis.de
Weinert, Henrike; TU Dortmund; henrike.weinert@tu-dortmund.de
Weiß, Christian; Helmut-Schmidt-Universität; weissc@hsu-hh.de
Weitzel, Nils; TU Dortmund; weitzel@statistik.tu-dortmund.de
Wendler, Martin; Otto-von-Guericke Universität; martin.wendler@ovgu.de
Wermuth, Jan-Lukas; Goethe Universität Frankfurt; lukaswermu@gmail.com
Wied, Dominik; Universität zu Köln; dwied@uni-koeln.de
Wiegleb, Georg; Landeshauptstadt Magdeburg; georg.wiegleb@stadt.magdeburg.de
Windmann, Michael; LMU München; michael.windmann@lmu.de
Winkler, Dorothee; Bundesinstitut für Bau-, Stadt- und Raumforschung (BBSR); dorothee.winkler@bbr.bund.de
Wübker, Ansgar; Hochschule Harz; awuebker@hs-harz.de
Wurtz, Christian; Technische Universität Dortmund; christian.wurtz@tu-dortmund.de
Yu, Miao; Leibniz University Hannover; miao.yu@statistik.uni-hannover.de
Zemann, Christian; Statistik der Bundesagentur für Arbeit; Christian.Zemann@arbeitsagentur.de
Zessin, Elisabeth; BBSR; Elisabeth.Zessin@bbr.bund.de
Zimmermann, Monika; Universität Duisburg-Essen; monika.zimmermann@uni-due.de
Zöphel, Moritz; Humboldt-Universität zu Berlin; moritz.zoephel@student.hu-berlin.de
Zukovska, Jekaterina; Humboldt Universität zu Berlin; zukovskj@hu-berlin.de
Zwick, Markus; Statistisches Bundesamt; markus.zwick@destatis.de

Building Plans / Gebäudepläne



Die Raumnummern ergeben sich wie folgt:

R	11	T	03	C	20
Gebäude (aussen)	Eingang	Gebäude (innen)	Stock	Gang	Raum



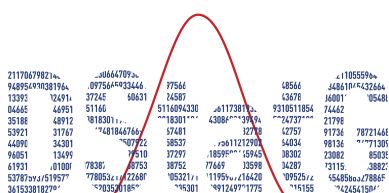
Essen

Statistische Woche 2026

08.-11. September



Alle Informationen zur Statistischen Woche 2026
für Android, IOS und als Web App unter
statistische-woche.de/app



Deutsche Statistische Gesellschaft

UNIVERSITÄT
DUISBURG
ESSEN

Offen im Denken