



International Conference
on Trends and Perspectives
in Linear Statistical Inference
LinStat 2026

Book of Abstracts

August 24 – 28, 2026
Poznań, Poland

Edited by

Katarzyna Filipiak

Institute of Mathematics

Poznań University of Technology, Poland

and

Monika Mokrzycka

Institute of Plant Genetics

Polish Academy of Sciences, Poznań, Poland

Contents

Part I. Introduction

Part II. Program

Part III. Invited Speakers

Smart estimation in high-dimensional sparse settings 31
S. Ejaz Ahmed

Pattern recovery by convex non-differentiable regularizers 32
Małgorzata Bogdan

On the combinatorics of generalized cumulants and k -statistics 33
Elvira Di Nardo

Green LIME: Improving AI explainability through design of experiments 35
Alexandra Stadler, Werner G. Müller, and Radoslav Harman

Equivalence testing: Bioequivalence and beyond 37
Thomas Mathew

Deviation tests for a high-dimensional mean 38
Jianfeng Yao

Part IV. Contributed Talks

Special Session: Copulas 41
Ostap Okhrin

Modeling and approximation of copulas for complex data structures using Cramér-von Mises statistic 42
Eckhard Liebscher

Dependence reduction by linear residualization 44
Jean-David Fermanian and Ostap Okhrin

Q-Tab: Quantized Tabular Data Generator 45
Julian Wustl, Philipp Haid, Yarema Okhrin, and Claudius Schnörr

Statistical information theory meets copulas	46
<i>Konstantinos (Kostas) Zografos</i>	
Special Session: Projection Pursuit	48
<i>Nicola Loperfido</i>	
Modeling incomplete compositional datasets	49
<i>Andriette Bekker, Jason Pillay, Cristina Tortora, and Antonio Punzo</i>	
Projection pursuit for detecting hidden structures	51
<i>Alessandro Berti and Nicola Loperfido</i>	
On the use of some unconventional projection pursuit indexes in cluster identification problems	52
<i>Claudio G. Borroni and Lucio De Capitani</i>	
Competing projections techniques of multivariate qual- ity data to monitor a production process	53
<i>Claudio G. Borroni, Manuela Cazzaro, and Paola M. Chiadini</i>	
On copula-based systems of quantile curves with appli- cation in outliers detection	54
<i>Luca Bagnato, Lucio De Capitani, and Antonio Punzo</i>	
Some remarks on kurtosis projection pursuit for clus- tering	55
<i>Cinzia Franceschini and Nicola Loperfido</i>	
A moment-based projection pursuit index	56
<i>Cinzia Franceschini and Nicola Loperfido</i>	
Modeling high-dimensional data with a multivariate Bernoulli distribution	57
<i>Mireille Boutin, Evzenie Coupkova, and Marco Morosin</i>	
A simple method for robust matrix completion with the- oretical recovery guarantees	58
<i>Eilon Vaknin Laufer, Pini Zilber, and Boaz Nadler</i>	
On nonstationary subspace analysis for multivariate ran- dom fields	59
<i>Jaakko Pere, Perttu Saarela, Sandra De Iaco, and Klaus Nordhausen</i>	

Stationary subspace analysis for spatial data	60
<i>Perttu Saarela, Klaus Nordhausen, Jaakko Pere, and Anne M. Ruiz</i>	
Optimal portfolio projections: Optimizing portfolio re- turns in the case of elliptically and skew-elliptically dis- tributed models	61
<i>Tomer Shushi</i>	
Special Session: New Trends in Robust Statistical Analysis – Theory and Application	62
<i>Agnieszka Wyłomańska</i>	
The Greenwood statistic	63
<i>Marek Arendarczyk, Tomasz J. Kozubowski, Anna K. Panorska, Katarzyna Skowronek, and Agnieszka Wyłomańska</i>	
When switching fractional Brownian motion becomes non-Gaussian	65
<i>Michał Balcerek, Adrian Pacheco-Pozo, Agnieszka Wyłomańska, and Diego Krapf</i>	
Kernel-based probabilistic path forecasting for electric- ity prices: empirical kernel calibration, scenario selec- tion and dynamic forecast updating	66
<i>Joanna Janczura and Andrzej Puć</i>	
Estimation methods of Matrix-valued Autoregressive model	68
<i>Alexander Lindner and Kamil Kołodziejcki</i>	
Coherent estimation of risk measures	69
<i>Marcin Pitera</i>	
The modified p-Greenwood statistic and its application to distinguishing Gaussian and near-Gaussian distribu- tions - platykurtic and leptokurtic cases	70
<i>Katarzyna Skowronek, Marek Arendarczyk, and Agnieszka Wyłomańska</i>	
Fractional lower-order covariance-based measures for cy- clostationary time series with heavy-tailed distributions: application to dependence testing and model order iden- tification	72
<i>Wojciech Żuławiński and Agnieszka Wyłomańska</i>	

Special Session: Inference for Non-Gaussian Distributions in Linear Models	74
<i>Krzysztof Podgórski</i>	
Moving random fields constructing and estimating motion in spatio-temporal random fields	75
<i>Anastassia Baxevari</i>	
Advances in nonparametric testing for overdispersed count data models	76
<i>Stefano Bonnini and Michela Borghesi</i>	
A class of continuous-time Gaussian mean-variance mixture state-space models	77
<i>Farrukh Javed and Krzysztof Podgórski</i>	
Distribution-free REML estimation of the AR(1) covariance structure	78
<i>Daniel Klein and Ivan Žežula</i>	
Testing patterned covariance matrices under quadratic subspace	79
<i>Daniel Klein, Mateusz John, and Yuli Liang</i>	
Test of the quadratic subspace under elliptically contoured distribution	80
<i>Daniel Klein, Malwina Mrowińska, and Ali Ali Taha</i>	
Covariance structure approximation under high-dimensionality with the use of improved Kullback-Leibler divergence ..	81
<i>Monika Mokrzycka, Jolanta Pielaszkiewicz, and Johan Alenlöv</i>	
Special Session: Machine Learning and Statistical Inference ...	82
<i>Tomasz Górecki</i>	
Matrix reduced-rank approximation through cross-validation	83
<i>Marisol García-Peña, Sergio Arciniegas-Alarcón, Johannes Forkman, and Diego Duarte</i>	
Positive unlabeled data and prior shift estimation	85
<i>Jan Mielniczuk, Wojciech Rejchel, and Paweł Teisseyre</i>	
Testing in matched-pairs functional data with missingness in one arm	86
<i>Marléne Baumeister, Łukasz Smaga, and Markus Pauly</i>	

Very simple and relatively precise mapping of the normal quantile function intended to generate normal pseudo-random numbers fastly	87
<i>Piotr Sulewski and Antoni Drapella</i>	
Special Session: Order Statistics	88
<i>Magdalena Szymkowiak</i>	
Discrete-time signatures of coherent systems: Properties, stochastic orderings and transformations.....	89
<i>Naraswamy Balakrishnan, He Yi, and Agnieszka Goroncy</i>	
Stochastic dominance between linear combination	91
<i>Tommaso Lando and Paulo Eduardo Oliveira</i>	
Preservation of IFR property for systems with two or three minimal paths of the same length	93
<i>Jagoda Papis and Magdalena Szymkowiak</i>	
Sharp bounds on L -moments	94
<i>Tomasz Rychlik and Magdalena Szymkowiak</i>	
ϕ -divergence as a measure of the predictability of reliability systems.....	95
<i>Narayanaswamy Balakrishnan, Maria Kateri, and Magdalena Szymkowiak</i>	
Special Session: Statistics in Applications	97
<i>Barbora Kessel</i>	
Beyond average speed: Functional data analysis of traffic calming effects	98
<i>Eva Fišerová</i>	
Quadratically constrained likelihood estimation for exponential regression under censoring.....	99
<i>Julianna Holmberg, Laila Hübbert, Dietrich von Rosen, and Martin Singull</i>	
Examining law students' intention to use GenAI systems for professional purposes using PLS modelling	100
<i>Vivien Kardos and Peter Kovacs</i>	
Linear-Gaussian Laubach–Williams state-space estimation and decomposition of the natural rate of interest – r^* : Evidence from Poland	101
<i>Paweł Kołbyko</i>	

Algorithmic support for judicial sentencing? Evidence from human smuggling cases	103
<i>Peter Kovacs, Krisztina Karsai, Zsanett Fantoly, Andor Gal, and Balint Kelemen</i>	
Statistical comparison of wavelet and Fourier series approximation quality for audio signals	104
<i>Mateusz John and Agnieszka Kwita</i>	
ANOVA: Is 100 years enough?	105
<i>Lynn R. LaMotte</i>	
Canonical variate analysis for two-way classification applied to a fertilization experiment with winter triticale .	106
<i>Alicja Lerczak, Dariusz Kayzer, Katarzyna Wielgusz, Dorota Czerwińska-Kayzer, and Renata Gaj</i>	
The resilience of random intercepts and slopes models to violations of the linearity assumption in the treatment group	108
<i>Im Chiew Ng, Katy Morgan, Jennifer Nicholas, and James Carpenter</i>	
Computing with characteristic functions: From numerical inversion to multivariate logistic goodness-of-fit	110
<i>Viktor Witkovský</i>	
Special Session: Experimental Designs	112
<i>Maryna Prus</i>	
A Bayesian updating framework for long-term multi-environment trial data in plant breeding	113
<i>Stephan Bark, Maryna Prus, and Hans-Peter Piepho</i>	
Selected properties of estimators in the model of weighing design	115
<i>Małgorzata Graczyk and Bronisław Ceranka</i>	
Estimation and design strategies for inter-individual differences in an overparameterized paired comparison model	116
<i>Heiko Großmann, Heinz Holling, Nadja Malevich, and Rainer Schwabe</i>	
Optimizing the allocation of trials to sub-regions in crop variety testing	117
<i>Maryna Prus</i>	
Session: Multivariate Statistics	118

The power of blocks: Matrix operators for statistics	119
<i>Katarzyna Filipiak</i>	
Testing hypotheses on covariance matrices within linear structure spaces	120
<i>Mateusz John and Adam Mieldzioc</i>	
Quasi-shrinkage estimation of block covariance matrix with structured off-diagonal blocks	122
<i>Monika Mokrzycka and Malwina Mrowińska</i>	
Convexity of the Moore–Penrose inverse	123
<i>Kenneth Nordström</i>	
The Safety Belt estimation method	124
<i>Katarzyna Filipiak and Dietrich von Rosen</i>	
Independence testing based on the covariance matrix . . .	125
<i>Adam Mieldzioc, Malwina Mrowińska, and Anna Szczepańska-Álvarez</i>	
Special Session: Memorial Session for Jerzy K. Baksalary (1944–2005)	126
<i>Simo Puntanen</i>	
On the Baksalary–Hauke condition aka the Rao–Mitra–Bhimasankaram relation	127
<i>Oskar Maria Baksalary</i>	
Remembering Jerzy K. Baksalary: My Mentor in Mathematics and Beyond	129
<i>Jan Hauke</i>	
Linear admissibility and sufficiency	130
<i>Augustyn Markiewicz</i>	
Jerzy Baksalary’s research: A personal note	131
<i>Thomas Mathew</i>	
Photographic reminiscences of Jerzy K. Baksalary (1944–2005)	132
<i>Simo Puntanen</i>	
Part V. List of Participants	
Index	145

Part I

Introduction

International Conference on Trends and Perspectives in Linear Statistical Inference, LinStat 2026, will be held from 24 till 28 August, 2026, at the Poznan University of Technology, Poznań, Poland. This is the follow-up of the LinStat series held in Będlewo, Poland (2008, 2012, 2018, 2021), in Tomar, Portugal (2010, 2022), in Linköping, Sweden (2014), Istanbul, Turkey (2016), and Poprad, Slovakia (2024).

The purpose of the meeting is to bring together researchers sharing an interest in a variety of aspects of statistics and its applications as well as matrix analysis and its applications to statistics, and offer them a possibility to discuss current developments in these subjects.

The work of young scientists has a special position in the LinStat 2026 to encourage and promote them. The best poster as well as the best talk will be chosen. Prizes will be awarded to graduate students or scientists with a recently completed Ph.D. Prize-winning works will be widely publicized and promoted by the conference.

Special sessions:

- Copulas
- Projection Pursuit
- New Trends in Robust Statistical Analysis - Theory and Application
- Inference for Non-Gaussian Distributions in Linear Models
- Machine Learning and Statistical Inference
- Order Statistics
- Statistics in Applications
- Experimental Designs,

as well as Multivariate Statistics session, will be held during the conference. Also, the Memorial Session Memorial Session for Jerzy K. Baksalary (1944-2005) will be organized.

The conference will include invited talks given by

- S. Ejaz Ahmed (Canada)
- Małgorzata Bogdan (Poland)
- Elvira Di Nardo (Italy)
- Werner Müller (Austria)
- Thomas Mathew (USA)
- Jianfeng Yao (China)

Organizers

- Institute of Mathematics, Poznań University of Technology, Poland
- Institute of Plant Genetics, Polish Academy of Sciences, Poznań, Poland
- Institute of Mathematics, P. J. Šafárik University, Košice, Slovakia
- Faculty of Mathematics and Computer Science, Adam Mickiewicz University, Poznań, Poland
- Department of Mathematical and Statistical Methods, Poznań University of Life Sciences, Poland

in partnership with the Foundation for the Development of the Poznan University of Technology.

Committees

The Scientific Committee for this Conference comprises

- Augustyn Markiewicz (Poland) - Chair
- Anthony C. Atkinson (UK)
- João Tiago Mexia (Portugal)
- Simo Puntanen (Finland)
- Dietrich von Rosen (Sweden)
- Müjgan Tez (Turkey)
- Götz Trenkler (Germany)
- Roman Zmysłony (Poland)
- Ivan Žežula (Slovakia)

The Organizing Committee comprises

- Katarzyna Filipiak (Poland) - Chair
- Daniel Klein (Slovakia)
- Monika Mokrzycka (Poland)
- Mateusz John (Poland)
- Malwina Mrowińska (Poland)
- Adam Mieldzioc (Poland)
- Tomasz Górecki (Poland)

Call for papers

We are pleased to announce special issue of *Statistical Papers* (Springer)

<https://link.springer.com/collections/fbcbicdafa>

devoted to LinStat 2026. It will include selected papers strongly correlated to the talks of the conference, with particular emphasize in linear models, including estimation, testing, and prediction; multivariate and high-dimensional data analysis; robustness of statistical procedures; extensions beyond linear settings; design and analysis of experiments; copulas; projection pursuit; and methodological developments involving matrix analysis in statistics.

All papers submitted to *Statistical Papers* must meet the publication standards of the journal and will be subject to normal refereeing procedure. The contributors must choose “S.I.: LinStat2026” as the manuscript type.

The deadline for paper submission: **31st of December 2026**

Guest Editors:

Katarzyna Filipiak and Daniel Klein

Part II

Program

Scientific Programme

Monday, August 24

08:00 **Registration**

09:00 - 09:10 **Opening**

Plenary Session

Chair: Dietrich von Rosen

09:10 - 09:55 **Jianfeng Yao**
Deviation tests for a high-dimensional mean

09:55 - 10:10 **Refreshment**

Session: Projection Pursuit

Chair: Tomer Shushi

10:10 - 10:40 **Andriette Bekker**
Modeling incomplete compositional datasets

10:40 - 11:10 **Jaakko Pere**
On nonstationary subspace analysis for multivariate random fields

11:10 - 11:30 **Marco Morosin**
Modeling high-dimensional data with a multivariate Bernoulli distribution

11:30 - 11:50 **Perttu Saarela**
Stationary subspace analysis for spatial data

11:50 - 12:20 **Coffee Break**

Session: New Trends in Robust Statistical Analysis - theory and application

Chair: Agnieszka Wyłomańska

12:20 - 12:50 **Michał Balcerek**
When switching fractional Brownian motion becomes non-Gaussian

12:50 - 13:20 **Marcin Pitera**
Coherent estimation of risk measures

13:20 - 13:50 **Joanna Janczura**
Kernel-based probabilistic path forecasting for electricity prices: empirical kernel calibration, scenario selection and dynamic forecast updating

14:00 **Lunch**

Plenary Session*Chair: Thomas Mathew*

- 15:00 - 15:45** **Werner Müller**
Green LIME: Improving AI explainability through design of experiments

15:45 - 16:15 Coffee Break**Session: Experimental Design***Chair: Maryna Prus*

- 16:15 - 16:40** **Nadja Malevich**
Estimation and design strategies for inter-individual differences in an overparameterized paired comparison model
- 16:40 - 17:00** **Małgorzata Graczyk**
Selected properties of estimators in the model of weighing design
- 17:00 - 17:20** **Stephan Bark**
A Bayesian updating framework for long-term multi-environment trial data in plant breeding
- 17:20 - 17:45** **Maryna Prus**
Optimal allocation of trials to sub-regions in multi-environment crop variety testing

17:45 - 17:50 Refreshment**Session: Inference for Non-Gaussian Distributions in Linear Models***Chair: Krzysztof Podgórski*

- 17:50 - 18:20** **Farrukh Javed**
A class of continuous-time Gaussian mean-variance mixture state-space models
- 18:20 - 18:40** **Yuli Liang**
Testing patterned covariance matrices under quadratic subspace
- 18:40 - 19:00** **Jolanta Pielaszkiewicz**
Covariance structure approximation under high-dimensionality with the use of improved Kullback-Leibler divergence

19:00 Welcome Reception

Scientific Programme

Tuesday, August 25

Plenary Session

Chair: Ivan Žežula

09:00 - 09:45 **Syed Ejaz Ahmed**
Smart estimation in high-dimensional sparse settings

09:45 - 10:00 **Refreshment**

Session: New Trends in Robust Statistical Analysis - theory and application

Chair: Agnieszka Wyłomańska

10:00 - 10:30 **Wojciech Żuławiński**
Fractional lower-order covariance-based measures for cyclostationary time series with heavy-tailed distributions: application to dependence testing and model order identification

10:30 - 11:00 **Marek Arendarczyk**
The Greenwood statistic

11:00 - 11:20 **Katarzyna Skowronek**
The modified p -Greenwood statistic and its application to distinguishing Gaussian and near-Gaussian distributions - platykurtic and leptokurtic cases

11:20 - 11:40 **Kamil Kołodziejski**
Estimation methods of matrix-valued autoregressive model

11:40 - 12:10 **Coffee Break**

Session: Machine Learning and Statistical Inference

Chair: Tomasz Górecki

12:10 - 12:30 **Johannes Forkman**
Matrix reduced-rank approximation through cross-validation

12:30 - 13:00 **Wojciech Rejchel**
Positive unlabeled data and prior shift estimation

13:00 - 13:30 **Łukasz Smaga**
Testing in matched-pairs functional data with missingness in one arm

13:30 - 14:00 **Piotr Sulewski**
Very simple and relatively precise mapping of the normal quantile function intended to generate normal pseudo-random numbers fastly

14:00 **Lunch**

Plenary Session

Chair: Werner Müller

15:00 - 15:45 **Thomas Mathew**
Equivalence testing: Bioequivalence and beyond

15:45 - 16:15 **Coffee Break**

Session: Statistics in applications

Chair: Viktor Witkovský

16:15 - 16:35 **Lynn Roy LaMotte**
ANOVA: is 100 years enough?

16:35 - 16:55 **Julianna Holmberg**
Quadratically constrained likelihood estimation for exponential regression under censoring

16:55 - 17:15 **Alicja Lerczak**
Canonical variate analysis for two-way classification applied to a fertilization experiment with winter triticale

17:15 - 17:30 **Refreshment**

Session: Memorial Session for Jerzy K. Baksalary (1944-2005)

Chair: Simo Puntanen

17:30 **Oskar Maria Baksalary**
On the Baksalary--Hauke condition aka the Rao--Mitra--Bhimasankaram relation

Jan Hauke
Remembering Jerzy K. Baksalary: My Mentor in Mathematics and Beyond

Augustyn Markiewicz
Linear admissibility and sufficiency

Thomas Mathew
Jerzy Baksalary's research: A personal note

Simo Puntanen
Photographic reminiscences of Jerzy K. Baksalary (1944-2005)

19:00 **Snack**

Scientific Programme

Wednesday, August 26

Session: Multivariate Statistics

Chair: Augustyn Markiewicz

- 09:00 - 09:20** **Dietrich von Rosen**
The Safety Belt estimation method
- 09:20 - 09:40** **Kenneth Nordström**
Convexity of the Moore–Penrose inverse
- 09:40 - 10:00** **Katarzyna Filipiak**
The power of blocks: Matrix operators for statistics

10:00 - 10:10 **Refreshment**

Session: Projection Pursuit

Chair: Jaakko Pere

- 10:10 - 10:40** **Claudio Borroni**
On the use of some unconventional projection pursuit indexes in cluster identification problems
- 10:40 - 11:10** **Manuela Cazzaro**
Competing projections techniques of multivariate quality data to monitor a production process
- 11:10 - 11:40** **Lucio De Capitani**
On copula-based systems of quantile curves with application in outliers detection

11:40 - 12:10 **Coffee Break**

Session: Copulas

Chair: Ostap Okhrin

- 12:10 - 12:40** **Konstantinos Zografos**
Statistical information theory meets copulas
- 12:40 - 13:10** **Eckhard Liebscher**
Modeling and approximation of copulas for complex data structures using Cramér-von Mises statistic
- 13:10 - 13:30** **Yarema Okhrin**
Q-Tab: Quantized tabular data generator
- 13:30 - 14:00** **Ostap Okhrin**
Dependence reduction by linear residualization

14:00 **Lunch**

Plenary Session*Chair: Jianfeng Yao*

15:00 - 15:45 **Elvira Di Nardo**
On the combinatorics of generalized cumulants and k-statistics

15:45 - 16:15 **Coffee Break**

Session: Statistics in applications*Chair: Eva Fišerová*

16:15 - 16:35 **Peter Kovacs**
Algorithmic support for judicial sentencing? Evidence from human smuggling cases

16:35 - 16:55 **Patryk Kołbyko**
Linear-Gaussian Laubach-Williams State-Space estimation and decomposition of the natural rate of interest - r^ : Evidence from Poland*

16:55 - 17:05 **Refreshment**

Session: Order Statistics*Chair: Magdalena Szymkowiak*

17:05 - 17:35 **Paulo Eduardo Oliveira**
Stochastic dominance between linear combination

17:35 - 18:05 **Agnieszka Goroncy**
Discrete-time signatures of coherent systems: properties, stochastic orderings and transformations

18:05 - 18:25 **Jagoda Papis**
Preservation of IFR property for systems with two or three minimal paths of the same length

20:00 **Conference Dinner**

Scientific Programme

Thursday, August 27

Plenary Session

Chair: Syed Ejaz Ahmed

09:00 - 09:45 **Małgorzata Bogdan**
Pattern recovery by convex non-differentiable regularizers

09:45 - 10:00 **Refreshment**

Session: Projection Pursuit

Chair: Claudio Borroni

10:00 - 10:30 **Alessandro Berti**
Projection pursuit for detecting hidden structures

10:30 - 11:00 **Boaz Nadler**
A simple method for robust matrix completion with theoretical recovery guarantees

11:00 - 11:30 **Tomer Shushi**
Optimal portfolio projections: Optimizing portfolio returns in the case of elliptically and skew-elliptically distributed models

11:30 - 12:00 **Coffee Break**

Session: Statistics in Applications

Chair: Barbora Kessel

12:00 - 12:30 **Viktor Witkovský**
Computing with characteristic functions: From numerical inversion to multivariate logistic goodness-of-fit

12:30 - 13:00 **Eva Fišerová**
Beyond average speed: Functional data analysis of traffic calming effects

13:00 - 13:20 **Im Chiew Ng**
The resilience of random intercepts and slopes models to violations of the linearity assumption in the treatment group

13:20 - 13:40 **Vivien Kardos**
Examining law students' intention to use GenAI systems for professional purposes using PLS modelling

13:40 - 14:00 **Agnieszka Kwita**
Statistical comparison of wavelet and Fourier series approximation quality for audio signals

14:00 **Lunch**

15:00 **Excursion**

Scientific Programme

Friday, August 28

Session: Order Statistics

Chair: Agnieszka Goroncy

- 09:00 - 09:30** **Tomasz Rychlik**
Sharp bounds on L-moments
- 09:30 - 10:00** **Magdalena Szymkowiak**
 Φ -divergence as a measure of the predictability of reliability systems

10:00 - 10:10 **Refreshment**

Session: Projection Pursuit

Chair: Tomer Shushi

- 10:10 - 10:40** **Cinzia Franceschini**
Some remarks on kurtosis projection pursuit for clustering
- 10:40 - 11:10** **Nicola Loperfido**
A moment-based projection pursuit index

11:10 - 11:40 **Coffee Break**

Session: Inference for Non-Gaussian Distributions in Linear Models

Chair: Krzysztof Podgórski

- 11:40 - 12:10** **Anastassia Baxevani**
Moving random fields constructing and estimating motion in spatio-temporal random fields
- 12:10 - 12:30** **Stefano Bonnini**
Advances in nonparametric testing for overdispersed count data models
- 12:30 - 12:50** **Malwina Mrowińska**
Test of the quadratic subspace under elliptically contoured distribution
- 12:50 - 13:10** **Daniel Klein**
Distribution-free REMLS estimation of the AR(1) covariance structure

13:10 **Lunch**

Session: Multivariate Statistics

Chair: Daniel Klein

14:00 - 14:20

Anna Szczepńska-Alvarez

Independence testing based on the covariance matrix

14:20 - 14:40

Mateusz John

Testing hypotheses on covariance matrices within linear structure spaces

14:40 - 15:00

Monika Mokrzycka

Quasi-shrinkage estimation of block covariance matrix with structured off-diagonal blocks

15:00

Closing

Part III

Invited Speakers

Smart estimation in high-dimensional sparse settings

S. Ejaz Ahmed

Brock University, Canada

Abstract

In this talk, we address estimation and prediction problems in regression models where the coefficients may be constrained to lie in a candidate subspace, considering both low- and high-dimensional settings. We propose a class of penalized, pretest, and shrinkage estimators, and establish their large-sample properties through the analysis of asymptotic distributional quadratic bias and risk.

Monte Carlo simulations are conducted to support the theoretical results and to assess the finite-sample performance of the proposed methods. In addition, real data applications-including comparisons with standard penalized techniques and machine learning-based approaches-illustrate the practical advantages of our methodology.

Our proposed strategies consistently outperform competing methods over meaningful regions of the parameter space induced by commonly assumed sparsity structures. Moreover, they continue to perform well even when the sparsity assumption is substantially violated, providing strong evidence of their robustness.

Pattern recovery by convex non-differentiable regularizers

Małgorzata Bogdan

University of Wrocław, Poland

Abstract

Convex non-differentiable regularization is usually associated with sparsity, but its scope is much broader. Many such penalties promote structured low-dimensional patterns in the parameter space, including sparsity, clustering, and fusion. In this talk, I will present a general framework in which the notion of pattern is characterized through the geometry of the subdifferential of the penalty. This perspective helps explain how regularization induces interpretable structure and why it can improve estimation, not only by reducing complexity, but also by shrinking the estimator toward a lower-dimensional pattern space.

Such shrinkage can lead to a substantial reduction in variance and, consequently, to improved estimation accuracy, especially when the underlying signal indeed possesses a simple latent structure. I will discuss theoretical results showing when these patterns can be consistently recovered and how pattern recovery differs from the more classical goals of estimation and prediction. The asymptotic analysis of the pattern convergence requires tools going beyond standard asymptotic theory and leads to conditions for pattern recovery that are closely related to irrepresentability-type phenomena. Examples based on fused LASSO and SLOPE will illustrate how convex regularization can recover meaningful structure beyond variable selection alone.

On the combinatorics of generalized cumulants and k -statistics

Elvira Di Nardo

University of Turin, Italy

Abstract

Generalized cumulants offer a flexible framework for the analysis of polynomial functions of random variables [3]. They can be viewed as intermediate objects between joint moments and joint cumulants, providing a structured way to compute the cumulants of such functions.

Within this framework, k -statistics and their extensions yield unbiased estimators of generalized cumulants. While the classical theory provides a well-established framework for ordinary cumulants and k -statistics, its extension to generalized cumulants presents additional challenges, primarily due to the complexity involved in enumerating complementary set partitions required for their computation. This complexity can limit computational efficiency and practical implementation in non-symbolic numerical environments.

In this talk, we present a novel combinatorial approach for generating complementary set partitions based on two-block partitions, leading to improved computational efficiency and enabling implementation in non-symbolic environments [2]. We also extend generalized cumulants to products of random variables indexed by multiset subdivisions, allowing for repetitions and powers.

A key ingredient is the use of multi-index partitions, which provide a tractable framework for handling multivariate cumulant formulas and for constructing multivariate polykays [1]. The results highlight the interplay between combinatorial techniques and statistical methodology, offering both theoretical insight and practical tools for cumulant-based analysis.

Keywords

Complementary set partitions, Multivariate polykays, Multi-index partitions, Combinatorial algorithms.

References

- [1] Di Nardo, E., Guarino, G. (2022). k statistics: Unbiased estimates of joint cumulant products from the multivariate Faà di Bruno's formula. *The R Journal*, 14, 208–228.

- [2] Di Nardo, E., Guarino, G. (2025). Efficient computation of complementary set partitions, with applications to an extension and estimation of generalized cumulants. *arXiv:2505.12706 [math.ST]*.
- [3] McCullagh, P. (2018). *Tensor Methods in Statistics: Monographs on Statistics and Applied Probability*. Chapman and Hall/CRC.

Green LIME: Improving AI explainability through design of experiments

Alexandra Stadler¹, Werner G. Müller¹,
and Radoslav Harman²

¹ Johannes Kepler University Linz, Austria

² Comenius University Bratislava, Slovakia

Abstract

In artificial intelligence (AI), the complexity of many models and processes surpasses human understanding, making it challenging to determine why a specific prediction is made. This lack of transparency is particularly problematic in critical fields like healthcare, where trust in a model’s predictions is paramount. As a result, the explainability of machine learning (ML) and other complex models has become a key area of focus. Efforts to improve model explainability often involve experimenting with AI systems and approximating their behavior through interpretable surrogate mechanisms. However, these procedures can be resource-intensive. Optimal design of experiments, which seeks to maximize the information obtained from a limited number of observations, offers promising methods for improving the efficiency of these explainability techniques.

To demonstrate this potential, we explore Local Interpretable Model-agnostic Explanations (LIME), a widely used method introduced by [1]. LIME provides explanations by generating new data points near the instance of interest and passing them through the model. While effective, this process can be computationally expensive, especially when predictions are costly or require many samples. LIME is highly versatile and can be applied to a wide range of models and datasets. In this work, we focus on models involving tabular data, regression tasks, and linear models as interpretable local approximations.

By utilizing optimal design of experiments’ techniques, we reduce the number of function evaluations of the complex model, thereby reducing the computational effort of LIME by a significant amount. We consider this modified version of LIME to be energy-efficient or “green”.

Keywords

Model interpretability, Optimal design of experiments, Explainable Artificial Intelligence, Post-hoc explanation, Local regression.

Acknowledgments

We acknowledge the use of ChatGPT (OpenAI) to assist in editing the English grammar and improving language clarity in the abstract. The content and ideas are solely the authors' own, unless cited or stated otherwise.

References

- [1] Ribeiro, M. T., Singh, S., Guestrin, C. (2016). “Why should I trust you?": Explaining the predictions of any classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, (pp. 1135–1144).

Equivalence testing: Bioequivalence and beyond

Thomas Mathew

University of Maryland, Baltimore County, USA

Abstract

In hypothesis testing problems, the null hypothesis is often a statement of equality. The two-sample t-test is a classic example where the null hypothesis states that two population means are equal. In some applications on the comparison of two means, it is more appropriate to test if the two means are equivalent; i.e., if they are close according to a specified threshold. Many applications call for testing equivalence, rather than equality, where the statement of equivalence is the alternative hypothesis so that equivalence is concluded by rejecting the null hypothesis. In the talk I will give a few applications where equivalence testing is required. These will include testing bioequivalence, testing the equivalence of several normal means relevant to the comparison of several treatments in a clinical trial, testing the equivalence of covariance matrices, comparing drug dissolution profiles, and equivalence testing for food safety assessment. All of the problems will be motivated and illustrated using practical applications.

Keywords

Equivalence of covariance matrices, Drug dissolution profiles, Food safety assessment, Likelihood-based criteria.

Deviation tests for a high-dimensional mean

Jianfeng Yao

The Chinese University of Hong Kong, Shenzhen, China

Abstract

This paper investigates testing for deviation of a high-dimensional mean vector $\boldsymbol{\mu}$. In contrast to the standard one-sample significance test of the form: $H_0^e : \boldsymbol{\mu} = \boldsymbol{\mu}_0$ versus $H_1^e : \boldsymbol{\mu} \neq \boldsymbol{\mu}_0$, we focus on testing the deviation $H_0 : \|\boldsymbol{\mu} - \boldsymbol{\mu}_0\|_2 \geq d_0$ versus $H_1 : \|\boldsymbol{\mu} - \boldsymbol{\mu}_0\|_2 < d_0$ for a prespecified length $d_0 > 0$. Constructing a valid test statistic for this problem is technically non-trivial. By applying the concept of positive and negative feedback processes from control theory, we propose a test statistic based on a two-armed bandit (TAB) process. The deviation test is also extended to the two-sample setting. Simulation experiments confirm a good performance of the tests in finite samples. Finally, a real data analysis demonstrates the practical significance of the proposed deviation tests.

This is a joint work with Zengjing Chen (Shandong University) and Ruihan Liu (The University of Hong Kong).

References

- [1] Chen, Z., Feng, S., Zhang, G. (2022). Strategy-driven limit theorems associated bandit problems. *arXiv preprint* arXiv:2204.04442.
- [2] Chen, Z., Yan, X., Zhang, G. (2023). Strategic two-sample test via the two-armed bandit process. *J. R. Stat. Soc. Ser. B. Stat. Methodol.*, 85, 1271–1298.

Part IV

Contributed Talks

Special Session: Copulas

Ostap Okhrin

TUD Dresden University of Technology, Germany

Abstract

Copulas have become a central tool in modern multivariate statistics, offering a flexible framework to model and analyze complex dependence structures beyond the limitations of classical correlation. By separating the marginal behavior of random variables from their dependence structure, copulas allow researchers to capture non-linear, asymmetric, and tail dependencies that often occur in real-world data. The session will cover both theoretical advances and practical applications of copula models. Topics include novel approaches to copula construction and estimation, dependence measures, inference procedures, and goodness-of-fit testing. Special emphasis will be placed on recent developments in high-dimensional settings, time-series contexts, and dynamic copula models.

Invited Speakers:

Eckhard Liebscher (Germany)
Ostap Okhrin (Germany)
Konstantinos Zografos (Greece)

Contributed Speakers:

Yarema Okhrin (Germany, Sweden)

Modeling and approximation of copulas for complex data structures using Cramér-von Mises statistic

Eckhard Liebscher

University of Applied Sciences Merseburg, Germany

Abstract

This talk addresses the recent challenges associated with analysing multivariate data. The objective is to approximate the copula of the random vector under consideration using an appropriate model. Goodness-of-fit tests often reveal that many null hypotheses are not rejected. This situation is due to the high complexity of the multivariate distribution and the limited size of the sample. In this talk, we present a method for overcoming this problem. Rather than seeking an exact fit, we search for a reasonable approximation of the sample copula. The aim is to find a copula from a parametric family or from families which approximates best the sample copula.

The specific challenges of complex data sets of a random vector $X = (X^{(1)}, \dots, X^{(d)})$ are:

- (i) the distribution of at least one $X^{(j)}$ has discrete components,
- (ii) asymmetry of the copula and different correlations,
- (iii) high-dimensional random vectors,
- (iv) copulas with special structure.

Regarding (i), the investigations are limited to positive random vectors, where zero values occur with positive probability (zero-inflated data). In this case we are not able to approximate the whole copula which is uniquely determined only on the codomain of the marginal distributions. We therefore work with a surrogate function $\tilde{C}$ for C .

Let C denote the copula of the sample. The copula C is modelled by a parametric family $\mathcal{M} = \{C_\theta\}_{\theta \in \Theta}$ of copulas. Θ is the compact parameter space. Against the above background, we assume that C does not belong to $\mathcal{M}$. We consider the Cramér-von-Mises divergence as a measure of discrepancy between the copula C and the model class $\mathcal{M}$. This has the advantage that copula densities, which can have complicated formulas, do not need to be evaluated. The parameter θ_0 of the best approximation is estimated by the approximate minimum distance estimator $\hat{\theta}_n$. Under certain regularity conditions, we have proved that $\hat{\theta}_n$ is a consistent estimator for θ_0 . Moreover we present an asymptotic normality result for $\hat{\theta}_n$. This was achieved by applying a central limit theorem for U -statistics. We show how comparisons of the approximation by several copula models can be done.

We provide a method for treating the distribution of high-dimensional random vectors by using product copulas which are mostly asymmetric. Finally, the use of the methods is illustrated by means of examples.

Keywords

Asymmetric copulas, Approximation of copulas, Zero-inflated data, High-dimensional random vectors.

References

- [1] Liebscher, E. (2009). Semiparametric estimation of the parameters of multivariate copulas. *Kybernetika*, 45, 972–991.
- [2] Liebscher, E. (2025). Selection of parametric copula models in the approximation of copulas using Cramér-von Mises divergence. In: Steland, A., Rafajłowicz, E., Parolya, N. (Eds.), *Stochastic Models, Statistics and Their Applications, TU Delft 2024* (pp. 19–31). Springer.

Dependence reduction by linear residualization

Jean-David Fermanian¹ and Ostap Okhrin²

¹ CREST, ENSAE Paris, France

² TUD Dresden University of Technology, Germany

Abstract

A standard way to reduce dependence between two random variables is to regress one on the other and work with the resulting residual. While this procedure removes linear correlation by construction, its effect on nonlinear dependence is much less clear. We study this question in a copula framework. Starting from a bivariate vector (U, V) with copula C , we consider monotone marginal transformations $Z_1 = H_1^{-1}(U)$ and $Z_2 = H_2^{-1}(V)$ and the residualized variable $\varepsilon_\rho = Z_2 - \rho Z_1$, $\rho \in [-1, 1]$. We derive the copula of (Z_1, ε_ρ) and investigate how residualization changes measures of dependence, with particular emphasis on Kendall's tau.

We show that the Kendall's tau of (Z_1, ε_ρ) is non-increasing in ρ , strictly so under mild regularity conditions. This yields conditions under which residualization moves dependence towards independence and provides bounds on the admissible residualization parameter. In the Gaussian case, choosing ρ equal to the correlation coefficient eliminates both Pearson correlation and Kendall dependence. We demonstrate, however, that this property is genuinely Gaussian: outside the Gaussian copula family, eliminating linear correlation need not eliminate rank dependence and can even reverse its sign. Explicit counterexamples based on mixtures of Gaussian copulas and asymmetric copulas illustrate this phenomenon. We further propose empirical estimators of the residualized copula and Kendall's tau and establish uniform consistency with respect to ρ . Finally, the framework is extended to jointly select the marginal transformations and the residualization parameter so as to minimize a chosen measure of dependence.

Keywords

Copula, Dependence reduction, Residualization, Kendall's tau, Nonlinear dependence, Empirical copula.

Q-Tab: Quantized Tabular Data Generator

Julian Wustl¹, Philipp Haid¹, Yarema Okhrin¹,
and Claudius Schnörr²

¹ University of Augsburg, Germany

² Munich University of Applied Sciences, Germany

Abstract

Codebook-based generators built on masked language model (MLM) transformers have become highly effective in text and vision, yet remain underused for tabular data. This is because codebooks typically act as information bottlenecks, whereas synthetic tabular generation requires a code space larger than the training sample, with additional codes trained to support new tabular rows. We address this gap with Q-Tab, a codebook-based tabular generator that uses lookup-free quantization (LFQ) with residual corruption to jointly tokenize numerical variables, categorical variables are directly one-hot tokenized. A BERT-style MLM captures dependencies in the token space and can then be sampled from. Corruption propagates reconstruction supervision across the numerical code space, but under joint encoder-decoder training induces a moving-target regression problem whose difficulty depends on the corruption structure. This motivates residual LFQ as the quantization mechanism, balancing broader supervision with locality. Q-Tab achieves state-of-the-art predictive utility and label prediction, while matching the distributional fidelity of diffusion-based generators.

Keywords

Tabular data, Generation, Lookup-free quantization.

Statistical information theory meets copulas

Konstantinos (Kostas) Zografos

University of Ioannina, Greece

Abstract

The aim of this talk is to present a link between statistical information theory with the theory of copulas. Measures of statistical information are omnipresent quantities, and they receive considerable attention in many fields. These measures are popular because of their diverse conceptual interpretations and their applications in solving practical problems across fields like probability, statistics, engineering, and economics. This talk will overview broad classes of measures of information which are defined by means of the probability density function. Properties of these measures will be briefly mentioned and applications of the said measures in several disciplines and frames, from estimation and testing of hypotheses to canonical correlation analysis, will be also briefly exposed. Recent reconsiderations of measures of information, defined in terms of the distribution function instead of the density, will be presented. This approach bridges statistical information theory and copula theory, enabling collaborative exploitation of both research areas to formulate, develop, and address problems in both fields. Some preliminary results in the light of extreme value copulas will be presented.

Keywords

Measures of information, Divergence, Extreme value copula, Pickand's dependence function.

References

- [1] Basu, A., Harris, I. R., Hjort, N. L. and Jones, M. C. (1998). Robust and efficient estimation by minimizing a density power divergence. *Biometrika*, 85, 549–559.
- [2] Csiszár, I. (1967). Information-type measures of difference of probability distributions and indirect observations. *Stud. Sci. Math. Hungar.*, 2, 299–318.
- [3] Ferentinos, K., Papaioannou, T. (1981). New parametric measures of information. *Inform. Control*, 51, 193–208.
- [4] Nelsen, R. B. (2006). *An Introduction to Copulas* (2nd ed). Springer.
- [5] Papaioannou, T. (1985). Measures of information. In: Kotz, S., Johnson, N.L. (Eds.), *Encyclopedia of Statistical Sciences* (pp. 391–397). John Wiley & Sons, New York.
- [6] Pardo, L. (2006). *Statistical Inference Based on Divergence Measures*. Boca Raton, FL. Chapman & Hall/CRC.

- [7] Rao, M., Chen, Y., Vemuri, B. C., Wang, F. (2004). Cumulative residual entropy: a new measure of information. *IEEE Trans. Inform. Theory*, 50, 1220–1228.
- [8] Zografos, K. (2023). On reconsidering entropies and divergences and their cumulative counterparts: Csiszár’s, DPD’s and Fisher’s type cumulative and survival measures. *Probab. Eng. Inf. Sci.*, 37, 294–321.
- [9] Zografos, K. (2025). On entropy and divergence type measures of bivariate extreme value copulas. *REVSTAT*. To appear.

Special Session: Projection Pursuit

Nicola Loperfido

University of Urbino “Carlo Bo”, Italy

Abstract

Projection pursuit is a multivariate statistical technique aimed at detecting interesting low-dimensional data projections. It looks for the data projection which maximizes the projection pursuit index, that is a measure of its interestingness. After an interesting projection is found, it is removed to facilitate the search for other interesting features. Projection pursuit addresses three major challenges of multivariate analysis: the curse of dimensionality, the presence of irrelevant features and the limitations of visual perception. Its applications have been hampered by computational, interpretative and inferential problems. Additional problems arise when data are high-dimensional, that is when there are more variables than units. This session outlines the main features of projection pursuit and its connections with other multivariate techniques. The theory is illustrated with both real and simulated datasets.

Invited Speakers:

Andriette Bekker (South Africa)
Alessandro Berti (Italy)
Claudio Borroni (Italy)
Manuela Cazzaro (Italy)
Lucio De Capitani (Italy)
Cinzia Franceschini (Italy)
Nicola Loperfido (Italy)
Marco Morosin (The Netherlands)
Boaz Nadler (Israel)
Jaakko Pere (Finland)
Perttu Saarela (Finland)
Tomer Shushi (Israel)

Modeling incomplete compositional datasets

Andriette Bekker¹, Jason Pillay¹, Cristina Tortora²,
and Antonio Punzo³

¹ University of Pretoria, South Africa

² San José State University, USA

³ University of Catania, Italy

Abstract

Incomplete compositional data analysis is hindered by the lack of likelihood-based methods that directly handle missing proportions on the simplex [1,2]. This topic introduces an estimation technique that operates directly on the original data, thereby accommodating missing values while preserving the interpretability of the variables within their natural domain [3,4]. Simulation studies demonstrate robust performance in parameter estimation and imputation across missingness mechanisms. To showcase its practicality, the proposed algorithm is applied to two real datasets exhibiting distinct missingness patterns. Analysis of the National Health and Nutrition Examination Survey dataset (with values censored and missing at random) focuses on proportions of mercury types in blood samples. Different types are subjected to differing levels of detection limits making more than 40% of the dataset unobserved [5,6]. Application to PM2.5 species data from the Air Quality System, supports a four-component mixture model, providing a descriptive profile of measured PM2.5 speciation across the United States.

Keywords

Censoring, Compositional, Clustering, Detection limit, Hidden values, Missing values, Mixture Models.

References

- [1] Van den Boogaart, K. G., Tolosana-Delgado, R., Bren, M. (2006). Concepts for handling zeroes and missing values in compositional data. In: *Proceedings of IAMG*, Vol. 6.
- [2] Schafer, J. L. (1997). *Analysis of Incomplete Multivariate Data*. Chapman and Hall/CRC.
- [3] McLachlan, G. J., Krishnan, T. (2008). *The EM Algorithm and Extensions*. John Wiley & Sons.
- [4] Favaro, S., Hadjicharalambous, G., Prünster, I. (2011). On a class of distributions on the simplex. *J. Stat. Plan. Inference*, 141(9), 2987–3004.
- [5] Heitjan, D. F., Rubin, D. B. (1991). Ignorability and coarse data. *Ann. Statist.*, 19, 2244–2253.

- [6] Seaman, S., Galati, J., Jackson, D., Carlin, J. (2013). What is meant by “Missing at random”? *Statist. Sci.*, *28*(2), 257–268.

Projection pursuit for detecting hidden structures

Alessandro Berti and Nicola Loperfido

University of Urbino "Carlo Bo", Italy

Abstract

The relationship between sport performance and the financial assets of football clubs has attracted increasing attention of both investors and academics. This paper investigates topic within the framework of Italian major league championship for the 2024/2025 season, regarded as one of the most competitive leagues at the continental level, albeit less financially endowed than those of other countries, such as the United Kingdom and Spain. From the statistical point of view, we show that principal components lead to unsatisfactory results, while projection pursuit detects both clusters and outliers. To a lesser degree, the same does invariant coordinate selection. From the subject-matter point of view, we show that financial variables explain about half of the sport performance, and that top Italian football teams are clearly separate from each other, with Inter Football Club being quite peculiar, consistently with common beliefs.

Keywords

Football, Kurtosis, Invariant coordinates selection, Principal component analysis, Projection pursuit.

On the use of some unconventional projection pursuit indexes in cluster identification problems

Claudio G. Borroni and Lucio De Capitani

University of Milano-Bicocca, Italy

Abstract

Despite kurtosis is often regarded as a univariate characteristic, its relevance in the analysis of multivariate data has been recognized as an important tool in the field of projection pursuit. The need to project a set of multivariate points onto a lower dimensional space (often the real line) struggles against the computational effort to work with infinitely many directions for such projections, but some “interesting” directions are possibly found by minimizing or maximizing the kurtosis of the projected data. The existing literature concentrated on the use of the standardized fourth moment as a pursuit index and on the possible simplification of the related computations. In this work we evaluate the use of different unconventional kurtosis measures. We underline that their strength rests in their decomposability, as long as in their balanced sensitivity for both the tails and the center of the distribution. In a simulation study involving bivariate data, the effectiveness of the unsupervised projection pursuit procedures built upon these unconventional kurtosis indexes, to reveal hidden clustering structures, is assessed. The results obtained are promising and will be further validated on higher-dimensional data in future works.

Keywords

Kurtosis, Dimension reduction, Clustering.

Competing projections techniques of multivariate quality data to monitor a production process

Claudio G. Borroni, Manuela Cazzaro,
and Paola M. Chiodini

University of Milano-Bicocca, Italy

Abstract

The change-point paradigm has been successfully applied to statistical process control by building many parametric and nonparametric charts for sequential monitoring. We concentrate here on a novel implementation of the change-point paradigm, based on dynamic windows, which forces comparisons only between samples of the same size. That implementation has been primarily applied to univariate control variables, despite many applications need to consider more than a single aspect of quality and to draw suitable information from their dependence structure. Thus, in this talk, we investigate the application of the dynamic-window approach to multivariate control variables. Essentially, we propose a preliminary reduction of the dimension of such variables and we evaluate different techniques of projection to the real line. We focus on the use of norms and inter-point distances. A comparison study is conducted to: i) identify such techniques which can provide stable results with respect to the actual dimension and the actual shape of the underlying distribution; ii) identify techniques which can provide a fast signal when some or all of the components of the underlying distribution undergo shifts in location or scale.

Keywords

Control charts, Dynamic-window approach, Dimension reduction, Inter-point distances.

On copula-based systems of quantile curves with application in outliers detection

Luca Bagnato¹, Lucio De Capitani²,
and Antonio Punzo³

¹ Catholic University of the Sacred Heart, Milano, Italy

² University of Milano-Bicocca, Milano, Italy

³ University of Catania, Catania, Italy

Abstract

Systems of curves representing quantiles of a bivariate distribution are introduced. Such systems consist of decreasing curves that partition the support of the distribution into nested subsets, each associated with a prescribed probability, which can be interpreted as the quantile level corresponding to the respective curve.

For continuous bivariate distributions, there exist uncountably many systems of quantile curves. Some approaches to defining such systems have been proposed, among others, by [1] and [2]. Here, we focus on a particular class of systems based on the copula associated with the underlying bivariate distribution. The use of the copula makes it possible to construct quantile curves in two steps. First, a system of quantile curves is identified on the support of the copula (the unit square). These curves are then mapped onto the support of the original bivariate distribution through the marginal distributions, yielding a system of copula-based quantile curves.

Some preliminary inferential results are also presented, allowing the construction of both pointwise and simultaneous confidence bands for the quantile curves.

Finally, a first attempt at using the proposed notion of quantile curves for outlier detection in bivariate data is presented.

Keywords

Multivariate distributions, Multivariate quantile, Copula function.

References

- [1] Belloni, A., Winkler, R. L. (2011). On multivariate quantiles under partial orders. *Ann Statist.*, 39(2), 1125–1179.
- [2] Coblenz, M., Dyckerhoff, R., Grothe, O. (2018). Nonparametric estimation of multivariate quantiles. *Environmetrics*, 29(2), e2488.

Some remarks on kurtosis projection pursuit for clustering

Cinzia Franceschini¹ and Nicola Loperfido²

¹ University of Sassari, Italy

² University of Urbino “Carlo Bo”, Italy

Abstract

In cluster analysis, data projections with either maximal or minimal kurtosis are used to gain a better insight into the cluster structure. The method, known as kurtosis projection pursuit for clustering, is theoretically motivated for samples generated from a mixture of two multivariate normal distributions with either equal or proportional covariances. When more than two clusters are present, both heuristic arguments and limited empirical findings suggest that the method tends to detect just two clusters. Theoretically wise, we provide conditions which support the use of kurtosis-optimizing projections in the presence of more than two clusters. Empirically wise, we successfully apply kurtosis minimizations to a real data set where more than two clusters are present. Our empirical findings suggest that kurtosis projection pursuit may accommodate for outliers and that may be used for variable selection in clustering.

Keywords

Dimension reduction, Finite mixtures, K-means, Kurtosis optimization, Principal components, Tandem analysis.

A moment-based projection pursuit index

Cinzia Franceschini¹ and Nicola Loperfido²

¹ University of Sassari, Italy

² University of Urbino “Carlo Bo”, Italy

Abstract

In projection pursuit, the projection index measures the interestingness of a given data projection. In moment-based projection pursuit, the projection index depends on some moments of the projection, such as its skewness or its kurtosis. We propose the skewness-to-kurtosis ratio as a projection pursuit index. It is particularly apt at detecting clusters and may be used to detect other departures from normality. The index might be derived as the solution of a generalized tensor eigenvector problem and is theoretically motivated when sampling from a location mixture of two multivariate normal distributions or from an extended multivariate skew-normal distribution. Properties and limitations of the proposed index are investigated with special emphasis on its connections with other moment-based indices.

Keywords

Finite mixture, Generalized tensor eigenpairs, Kurtosis, Projection pursuit, Skewness.

Modeling high-dimensional data with a multivariate Bernoulli distribution

Mireille Boutin¹, Evzenie Coupkova²,
and Marco Morosin¹

¹ Eindhoven University of Technology, The Netherlands

² Purdue University, West Lafayette, USA

Abstract

Recent experiments have uncovered that several high-dimensional datasets have a high probability of forming binary clusters after projecting the data on a random one-dimensional subspace. We present a probability model for the high-dimensional data that could explain this phenomenon. The model consists of a discrete skeleton formed by several Bernoulli random variables arranged as the vertices of a parallelotope, on which noise is added. While clusters in high dimension are difficult to observe or may not exist at all, the groupings of points drawn from this distribution can be easily identified. These groupings can be used in place of clustering, especially when the dataset is small, and their statistical significance can be tested. This structure allows for semantic grouping of datapoints based on different criteria and provides a binary, compressed representation of the data in which each bit represents the binary class for one criterion.

Keywords

High-dimensional data, Clustering, Small data, Data binarization, Random projection, Data model.

References

- [1] Boutin, M., Coupkova, E. (2026). A new model for natural groupings in high-dimensional data. In: Nielsen, F., Barbaresco, F. (Eds.), *Geometric Science of Information, GSI 2025*, Lecture Notes in Computer Science, vol. 16033. Springer.

A simple method for robust matrix completion with theoretical recovery guarantees

Eilon Vaknin Laufer, Pini Zilber, and Boaz Nadler

Weizmann Institute of Science, Rehovot, Israel

Abstract

Recovering a low rank matrix from a subset of its entries, is known as the matrix completion (MC) problem. When some of the observed entries are corrupted, this problem is known as robust matrix completion (RMC). Existing methods have several limitations: they require a relatively large number of observed entries; they may fail under overparameterization, when their assumed rank is higher than the correct one; and many of them fail to recover even mildly ill-conditioned matrices. In this talk, I'll describe simple iterative methods for both the MC and the RMC problems, that empirically overcome these limitations. Our approach to both the MC and RMC problems is a factorization-based iterative algorithm, based on a Gauss-Newton linearization, and in the latter case combined with removal of entries suspected to be outliers. We present theoretical guarantees for both methods, which are amongst the best known for factorization-based approaches.

Keywords

Matrix completion, Outliers, Low rank structure, Gauss-Newton.

References

- [1] Zilber, P., Nadler, B. (2022). GNMR: A provable one-line algorithm for low rank matrix recovery. *SIAM J. Math. Data Sci.*, 4(2), 909–934.
- [2] Vaknin Laufer, E., Nadler, B. (2026). RGNMR: A Gauss-Newton method for robust matrix completion with theoretical guarantees. *Adv. Neural Inf. Process. Syst.*, 38, 60146–60188.

On nonstationary subspace analysis for multivariate random fields

Jaakko Pere¹, Perttu Saarela², Sandra De Iaco³,
and Klaus Nordhausen²

¹ Aalto University, Finland

² University of Helsinki, Finland

³ University of Salento, Italy

Abstract

The analysis of multivariate spatiotemporal data can potentially be simplified by applying a blind source separation model for a random field. Here we consider a particular linear blind source separation model in which the unobserved multivariate latent field includes stationary and nonstationary components. For the presented model, we establish identifiability conditions. Furthermore, we consider basic asymptotics for a diagonalization-based estimator of subspace spans. The diagonalized matrix is constructed by comparing local spatial statistics to the corresponding global statistic.

Keywords

Blind source separation, Multivariate random field, Independent subspace analysis, Nonstationarity.

Stationary subspace analysis for spatial data

Perttu Saarela¹, Klaus Nordhausen¹, Jaakko Pere²,
and Anne M. Ruiz³

¹ University of Helsinki, Finland

² Aalto University, Finland

³ University Toulouse Capitole, France

Abstract

Stationary subspace analysis (SSA) is a blind source separation framework that decomposes linearly mixed multivariate data into stationary and nonstationary components. We extend SSA to spatially indexed data by introducing spatial stationary subspace analysis (spSSA), which explicitly accounts for spatial dependence. We propose three estimation procedures for the unmixing matrix based on first- and second-order spatial statistics. Each procedure targets a different type of nonstationarity and can be formulated as the solution to a generalized eigenvalue problem. To address situations where multiple forms of nonstationarity are present simultaneously, we combine the three procedures using approximate joint diagonalization. Simulation studies demonstrate that this combined approach yields superior separation performance. When the dimension of the nonstationary subspace is known, the proposed methods reliably recover the latent stationary and nonstationary components. However, determining this dimension remains a fundamental challenge in SSA, for which no generally accepted solution currently exists. Building on our estimation procedures, we propose a novel data augmentation approach to estimate the dimension of the nonstationary subspace and demonstrate its effectiveness through simulation studies. The proposed methodology is easily transferable to time series settings, making it of broader methodological interest.

Keywords

Blind source separation, Data augmentation, Dimension reduction, Multivariate spatial data, Nonstationarity, Order determination.

Optimal portfolio projections: Optimizing portfolio returns in the case of elliptically and skew-elliptically distributed models

Tomer Shushi

Ben-Gurion University of the Negev, Beer-Sheva, Israel

Abstract

The concept of optimal portfolio projection is introduced as a technique that projects the portfolio weight vector to a lower dimension, enabling explicit solution of the optimal portfolio selection problem for any given risk measure. We study the class of skew-elliptically distributed risks. We show that, by following the proposed procedure, we can obtain explicit optimal weights for such risks, with a dramatic reduction in the complexity of the optimization problem.

**Special Session:
New Trends in Robust Statistical Analysis –
Theory and Application**

Agnieszka Wyłomańska

Wrocław University of Science and Technology, Poland

Abstract

Recent research clearly shows that classic statistical approaches (e.g., those related to model parameter estimation) may be insufficient for describing phenomena observed in reality. There are many reasons for this, for example, the fact that in industrial data we observe impulsive behaviors, which indicate that the models describing the data are non-Gaussian, or in biological experiments (so-called single particle tracking), where we observe behaviors corresponding to models with random parameters. During this session, new statistical approaches and methods used to describe and analyze data with a complex structure will be presented. The theoretical results presented will be supported by real-world examples of data analysis from various fields.

Invited Speakers:

Marek Arendarczyk (Poland)
Michał Balcerek (Poland)
Joanna Janczura (Poland)
Marcin Pitera (Poland)
Wojciech Żuławiński (Poland)

Contributed Speakers:

Kamil Kołodziejowski (Poland)
Katarzyna Skowronek (Poland)

The Greenwood statistic

Marek Arendarczyk¹, Tomasz J. Kozubowski²,
Anna K. Panorska², Katarzyna Skowronek³,
and Agnieszka Wyłomańska³

¹ University of Wrocław, Poland

² University of Nevada, Reno, USA

³ Wrocław University of Science and Technology, Poland

Abstract

Let $X_1, X_2, \dots, X_n$ be an i.i.d. random sample from a distribution supported on $[0, \infty)$. The Greenwood statistic is defined by

$$T_n = \frac{\sum_{i=1}^n X_i^2}{(\sum_{i=1}^n X_i)^2}.$$

This talk reviews a series of results on the Greenwood statistic and its modifications. First, we discuss how star-shaped ordering of the underlying distributions leads to stochastic ordering of the Greenwood statistic, providing a basis for inference in heavy-tailed models [1]. In particular, for the generalized Pareto distribution, this approach leads to hypothesis tests for the tail index [2]. We then introduce a modified Greenwood statistic for symmetric distributions, with applications to Gaussianity testing and inference for symmetric α -stable and Student's t distributions [3,4]. Finally, we discuss multivariate extensions applied to sub-Gaussian, multivariate Pareto, and multivariate Student's t models [5]. These results illustrate the versatility of the Greenwood statistic and its modifications as simple tools for statistical inference in heavy-tailed and non-Gaussian settings.

Keywords

Stochastic orders, Clustering, Generalized Pareto distribution, Heavy-tailed distributions, Gaussianity testing.

References

- [1] Arendarczyk, M., Kozubowski, T. J., Panorska, A. K. (2022). The Greenwood statistic, stochastic dominance, clustering and heavy tails. *Scand. J. Stat.*, 49, 331–352.
- [2] Arendarczyk, M., Kozubowski, T. J., Panorska, A. K. (2023). A computational approach to confidence intervals and testing for generalized Pareto index using the Greenwood statistic. *REVSTAT*, 21(3), 367–388.

- [3] Skowronek, K., Arendarczyk, M., Zimroz R., Wyłomańska, A. (2024). Modified Greenwood statistic and its application for statistical testing *J. Comput. Appl. Math.*, 116122.
- [4] Skowronek, K., Arendarczyk, M., Panorska, A. K., Kozubowski, T. J., Wyłomańska, A. (2025). Testing and estimation of the index of stability of univariate and bivariate symmetric α -stable distributions via modified Greenwood statistic. *J. Comput. Appl. Math.*, 116587.
- [5] Skowronek, K., Arendarczyk, M., Wyłomańska, A. (2026). Modified Greenwood statistic for multivariate Pareto and Student's t distributions in application to statistical testing. *Statist. Papers*, 67, 51.

When switching fractional Brownian motion becomes non-Gaussian

**Michał Balcerek^{1,2}, Adrian Pacheco-Pozo¹,
Agnieszka Wyłomańska², and Diego Krapf¹**

¹ Colorado State University, Fort Collins, USA

² Wrocław University of Science and Technology, Poland

Abstract

Anomalous diffusion is often observed in complex systems, including single-particle tracking experiments, crowded media, and financial time series. In this talk, we consider switching fractional Brownian motion, a model in which the diffusivity changes over time while the long-range temporal correlations of fractional Brownian motion are preserved. Although the underlying process is conditionally Gaussian, temporal heterogeneity may lead to non-Gaussian displacement distributions.

During the talk we analyze kurtosis as a simple diagnostic of non-Gaussianity for Markovian and non-Markovian switching mechanisms. We show that Markovian switching leads to asymptotic Gaussianity, whereas heavy-tailed dwell times may preserve non-Gaussianity over long timescales. We also compare kurtosis with distribution-based distances such as the Hellinger distance, showing that kurtosis captures the main deviations from Gaussianity.

The talk is based on the article [1].

Keywords

Switching fractional Brownian motion, Gaussianity, Anomalous diffusion, Heterogeneity, Kurtosis, Hellinger distance.

Acknowledgments

Michał Balcerek and Agnieszka Wyłomańska acknowledge the support from National Science Centre, Poland, via projects No. 2023/07/X/ST1/01139 (MB) and 2020/37/B/HS4/00120 (AW), respectively. Diego Krapf acknowledges the support from the National Science Foundation Grant 2102832.

References

- [1] Balcerek, M., Pacheco-Pozo, A., Wyłomańska A., Krapf, D. (2025). Evaluating Gaussianity of heterogeneous fractional Brownian motion. *J. Phys. A*, 58(17), 175001.

Kernel-based probabilistic path forecasting for electricity prices: empirical kernel calibration, scenario selection and dynamic forecast updating

Joanna Janczura and Andrzej Puć

Wrocław University of Science and Technology, Poland

Abstract

We propose a kernel-based framework for probabilistic path forecasting and illustrate its performance using intraday electricity prices as a challenging case study. The methodology builds on a corrected Support Vector Regression (cSVR) model, where information from an auxiliary forecast is incorporated directly into the kernel function. Kernel bandwidths are calibrated empirically from pairwise distance distributions, avoiding computationally demanding hyperparameter tuning. To obtain compact yet informative predictive ensembles, we introduce Support Vector Sorting (SVS), a model-driven scenario selection procedure that exploits the dual representation of the SVR model to rank trajectories according to their contribution to the prediction function. The final ensemble size is determined using a Wasserstein-distance stopping criterion, allowing a substantial reduction in the number of scenarios while preserving the essential characteristics of the predictive distribution. We further propose a dynamic updating mechanism based on scenario reweighting, allowing predictive distributions to adapt online to newly observed trajectory segments without retraining the model. An empirical study on the German intraday electricity market demonstrates that the proposed framework delivers informative point and probabilistic forecasts while substantially reducing ensemble complexity.

The talk is based on [1,2].

Keywords

Support vector regression, Probabilistic path forecasting, Scenario selection, Kernel methods, Weighted distribution.

References

- [1] Puć, A., Janczura, J. (2026). Corrected support vector regression for intraday point forecasting of prices in the continuous power market. *Int. J. Forecast.*, 42, 796–815.

- [2] Puć, A., Janczura, J. (2026). Scenario generation of intraday electricity price paths for optimal trading in continuous markets. Working paper. ArXiv:2605.13446.

Estimation methods of Matrix-valued Autoregressive model

Alexander Lindner¹ and Kamil Kołodziejski²

¹ Ulm University, Germany

² Lodz University of Technology, Poland

Abstract

In modern time series analysis, increasingly complex and high-dimensional data sets are often considered. To address the problem of modeling multivariate time series, the matrix-valued autoregressive (MAR) model was introduced in [1]. Estimation methods based on projection onto the space of Kronecker products, least squares (LSE), and maximum likelihood estimation (MLE) have already been developed. However, moment-based estimation methods have not yet been investigated.

In this talk, I introduce new estimation methods for the Matrix Autoregressive (MAR) model by extending the classical Yule-Walker equations to the matrix setting. I discuss the advantages of matrix-based formulations compared to vector autoregressive (VAR) models and outline the main similarities and differences between these approaches. MAR models preserve the inherent matrix structure of the data and require substantially fewer parameters than VAR models, making them efficient and interpretable tools for high-dimensional time series analysis.

Keywords

Matrix-valued time series, Multivariate time series, Estimation, Autoregressive.

Acknowledgements

This research has been completed while the second author was the Doctoral Candidate in the Interdisciplinary Doctoral School at the Lodz University of Technology, Poland.

References

- [1] Chen, R., Xiao, H., Yang, D. (2021). Autoregressive models for matrix-valued time series. *J. Econometrics*, 222(1), 539–560.
- [2] Tsay, R. S. (2024). Matrix-variate time series analysis: A brief review and some new developments. *Int. Stat. Rev.*, 92(2), 246–262.

Coherent estimation of risk measures

Marcin Pitera

Jagiellonian University, Poland

Abstract

We develop a statistical framework for risk estimation inspired by the axiomatic theory of risk measures and show that coherent risk estimators are characterized through robust dual representations linked to L -statistics. We also show on simulated and market data that theoretical coherence of a risk measure need not carry over to its estimator unless the proper weighting structure on estimation level is maintained.

The talk is based on a joint work with D. Jelito (UJ, Krakow), I. Cialenco (IIT, Chicago), and M. Aichele (ECB, Frankfurt).

Keywords

Estimation of risk measures, Coherent risk measure, Coherent risk estimator, Expected shortfall, Value at risk, L-estimator.

Acknowledgments

The author acknowledges support from the National Science Centre, Poland, via project 2024/53/B/HS4/00433.

References

- [1] Aichele, M., Cialenco, I., Jelito, D., Pitera, M. (2026). Coherent estimation of risk measures. *J. Financ. Econom.* Forthcoming.

The modified p -Greenwood statistic and its application to distinguishing Gaussian and near-Gaussian distributions - platykurtic and leptokurtic cases

Katarzyna Skowronek¹, Marek Arendarczyk²,
and Agnieszka Wyłomańska¹

¹ Wrocław University of Science and Technology, Poland

² University of Wrocław, Poland

Abstract

This paper introduces the modified p -Greenwood statistic, a novel generalization of the classical Greenwood statistic designed for goodness-of-fit testing. While traditional methodologies predominantly focus on exponentiality or specific heavy-tailed distributions, we extend this framework to a comprehensive range of distribution families, encompassing both leptokurtic and platykurtic classes. A primary theoretical contribution of this work is the characterization of the stochastic ordering of the statistic within analyzed classes, which provides the foundation for a formal Gaussianity testing framework. Furthermore, we develop a computational procedure for identifying the optimal parameter p , ensuring maximum statistical power when discriminating between Gaussian distributions and near-Gaussian alternatives within the analyzed classes. The proposed methodology is validated via extensive Monte Carlo simulations, demonstrating superior performance compared to established high-efficiency benchmarks, particularly in small-sample regimes. The practical utility of the approach is demonstrated through an application in condition monitoring area.

Keywords

Greenwood statistic, Modified Greenwood statistic, Platykurtic distribution, Stochastic monotonicity, Star-shaped order, Statistical testing, Monte Carlo simulations, Machine diagnostic signal.

Acknowledgments

The work of AW is supported by National Center of Science under Weave-Unisono project No. 2025/07/Y/ST8/00070 “Advanced signal processing techniques for cyclostationary modelling in Gaussian and non-Gaussian noisy environment - detection of cyclic sources, estimation, optimisation of algorithms and validation in the context of fault identification”. The work of M.A.

was supported by project BPIDUB.17.2024 under the program 'Excellence initiative - research university' for the years 2020-2026 at the University of Wrocław.

References

- [1] Greenwood, M. (1946). The statistical study of infectious diseases. *J. R. Stat. Soc.*, 109 (2), 85–110.
- [2] Arendarczyk, M., Kozubowski, T. J., Panorska, A. K. (2022). The Greenwood statistic, stochastic dominance, clustering and heavy tails. *Scand. J. Stat.*, 49, 331–352.
- [3] Skowronek, K., Arendarczyk, M., Zimroz R., Wyłomańska, A. (2024). Modified Greenwood statistic and its application for statistical testing *J. Comput. Appl. Math.*, 116122.
- [4] Skowronek, K., Arendarczyk, M., Wyłomańska, A. (2026). Modified Greenwood statistic for multivariate Pareto and Student's t distributions in application to statistical testing. *Statist. Papers*, 67, 51.
- [5] Skowronek, K., Arendarczyk, M., Wyłomańska (2026). The modified p -Greenwood statistic and its application to distinguishing Gaussian and near-Gaussian distributions - platykurtic and leptokurtic cases. In preparation.

Fractional lower-order covariance-based measures for cyclostationary time series with heavy-tailed distributions: application to dependence testing and model order identification

Wojciech Żuławiński and Agnieszka Wyłomańska

Wrocław University of Science and Technology, Poland

Abstract

This presentation introduces new methods for the analysis of cyclostationary time series with infinite variance. Traditional cyclostationary analysis, based on periodically correlated (PC) processes, relies on the autocovariance function (ACVF). However, the ACVF is not suitable for data exhibiting a heavy-tailed distribution, particularly with infinite variance. Thus, we propose a novel framework for the analysis of cyclostationary time series with heavy-tailed distribution, utilizing the fractional lower-order covariance (FLOC) as an alternative to covariance. This leads to the introduction of two new autodependence measures: the periodic fractional lower-order autocorrelation function (peFLOACF) and the periodic fractional lower-order partial autocorrelation function (peFLOPACF). These measures generalize the classical periodic autocorrelation function (peACF) and periodic partial autocorrelation function (pePACF), offering robust tools for analyzing infinite-variance processes. Two practical applications of the proposed measures are explored: a portmanteau test for testing dependence in cyclostationary series and a method for order identification in periodic autoregressive (PAR) and periodic moving average (PMA) models with infinite variance. Both applications demonstrate the potential of new tools, with simulations validating their efficiency. The methodology is further illustrated through the analysis of real-world air pollution data, which showcases its practical utility. The results indicate that the proposed measures based on FLOC provide reliable and efficient techniques for analyzing cyclostationary processes with heavy-tailed distributions.

Keywords

Cyclostationary time series, Heavy-tailed distribution, Fractional lower-order covariance, Testing dependence, PARMA model, Order identification.

Acknowledgments

The work is supported by the National Center of Science under the Sheng2 project No. UMO-2021/40/Q/ST8/ 00024 “NonGauMech - New methods of processing non-stationary signals (identification, segmentation, extraction, modeling) with non-Gaussian characteristics to monitor complex mechanical structures”.

The authors thank Prof. Valderio Reisen for sharing the real data analyzed in this work.

Special Session: Inference for Non-Gaussian Distributions in Linear Models

Krzysztof Podgórski

Lund University, Sweden

Abstract

Non-Gaussian features in data can be seen either as a curse or a blessing, depending on one's perspective. For those who value the elegance of the Gaussian paradigm and its well-established theoretical results, the disruptive effects of non-Gaussianity can be a source of frustration, challenging long-standing frameworks. Yet for others, these same departures from normality offer opportunities: by accounting for non-Gaussian structures, one can develop more efficient inference procedures and more accurate predictive methods, moving beyond the familiar – if somewhat predictable – Gaussian path. This session explores the world beyond the Gaussian paradigm by presenting new theoretical foundations and innovative inferential methodologies designed to address non-Gaussian data.

Invited Speakers:

Anastassia Baxevani (Cyprus)
Farrukh Javed (Sweden)

Contributed Speakers:

Stefano Bonnini (Italy)
Daniel Klein (Slovakia)
Yuli Liang (China)
Malwina Mrowińska (Poland)
Jolanta Pielaszkiewicz (Sweden)

Moving random fields constructing and estimating motion in spatio-temporal random fields

Anastassia Baxevani

University of Cyprus, Cyprus

Abstract

Many environmental processes - precipitation, ocean waves, wind fields - are not merely spatially correlated but move: weather systems drift, grow, and decay coherently in time. Capturing this motion requires more than a spatial random field evolving pointwise in time; it requires a model with a genuine underlying velocity. This talk presents a framework for constructing and estimating such moving random fields. We first formalize what it means for a random field to move, using excursion sets, crossing contours, and a generalized Rice formula that characterizes the distribution of the field's velocity at a level crossing. Building on this, we construct Gaussian random fields driven by a deterministic flow: first advected by a constant velocity, then generalized to a spatially and temporally varying velocity field via a flow of diffeomorphisms driving a stochastic integral. A key theorem shows that, under isotropic innovations, the field's own random velocity distribution is centered exactly on the deterministic flow used to construct it - confirming the construction behaves as intended. We then discuss two complementary strategies for estimating the velocity field from data: a method-of-moments approach exploiting structure in the spatio-temporal covariance, and a state-space formulation solved via particle filtering. We close by showing that the resulting motion signature is directly visible in the estimated covariance of real precipitation data.

Advances in nonparametric testing for overdispersed count data models

Stefano Bonnini and Michela Borghesi

University of Ferrara, Italy

Abstract

In generalized linear models with a discrete dependent variable representing count data, Poisson regression is most often used. Consequently, the goodness-of-fit of the model is tested using the likelihood ratio test, which is based on the assumption that the response follows the Poisson distribution. However, in many practical applications, this assumption is unrealistic. For example, when the variance is significantly different from the mean, as in the case of overdispersed data, i.e. when the variance is much larger than the mean, the assumption that the underlying distribution is the Poisson law is not plausible [2].

In this work, we present recent developments in a new nonparametric method for testing the validity of a model for count data. This innovative approach relies on a suitable combination of permutation tests on the significance of individual model coefficients, which are analyzed jointly [1]. The proposed method does not assume any distribution for the response variable and is therefore appropriate for overdispersed count data.

Keywords

Overdispersed count data, Poisson regression, Generalized linear model, Combined permutation test.

References

- [1] Bonnini, S., Borghesi, M. (2022). Relationship between mental health and socio-economic, demographic and environmental factors in the COVID-19 lockdown period - a multivariate regression analysis. *Mathematics*, 10(18), 3237.
- [2] Winkelmann, R. (2008). *Econometric Analysis of Count Data* Springer.

A class of continuous-time Gaussian mean-variance mixture state-space models

Farrukh Javed and Krzysztof Podgórski

Lund University, Sweden

Abstract

The generalized inverse Gaussian (GIG) distribution is central to the construction of generalized hyperbolic (GH) distributions, widely used in Gaussian mean-variance mixtures for modeling skewness and heavy tails. However, the class of GIG distributions, which are infinitely divisible (ID), in general, is not closed under convolution, which limits its applicability in continuous-time and Lévy-based modeling frameworks. This paper introduces an alternative, although related, class of mixing ID distributions, denoted by IGG, defined as a convex combination of inverse Gaussian and gamma distributions. The IGG class has as many parameters as GIG, with an additional parameter controlling the mixing proportion between its components. Subordinating Brownian motion with drift to an IGG Lévy process then yields a new family of Gaussian mixtures, referred to as Gaussian-IGG (NIGG) distributions, which contain the normal-inverse Gaussian and generalized asymmetric Laplace distributions as special cases. This framework also clarifies a confusion about the role of scale and infinite divisibility in models such as NIG, and provides a flexible alternative for constructing state-space models.

Distribution-free REMLS estimation of the AR(1) covariance structure

Daniel Klein and Ivan Žežula

P. J. Šafárik University, Košice, Slovakia

Abstract

The paper propose a restricted expected multivariate least squares (REMLS) approach for estimating the parameters of the first-order autoregressive covariance structure in multivariate linear models. The resulting estimators are available in closed form, satisfy the admissible parameter space without additional modifications, and require only mild distributional assumptions. Exact finite-sample moments are established for the scale estimator, whereas higher-order bias and variance approximations are derived for the correlation estimator. Particular attention is paid to elliptical distributions, where the influence of non-Gaussianity on the finite-sample properties is characterized explicitly through the fourth-order moments. Under Gaussian errors, the proposed scale estimator coincides with the maximum likelihood estimator. Monte Carlo experiments indicate that the REMLS estimators achieve finite-sample performance comparable to, and in several settings superior to, maximum likelihood estimation while remaining applicable under substantially weaker distributional assumptions.

Keywords

Restricted expected multivariate least squares, Autoregressive covariance structure, Covariance parameter estimation, Multivariate linear model, Elliptical distributions.

Acknowledgments

The research was supported by the Slovak Research and Development Agency under the Contract No. APVV-21-0369, and grant VEGA No. 1/0585/24.

Testing patterned covariance matrices under quadratic subspace

Daniel Klein¹, Mateusz John², and Yuli Liang³

¹ Institute of Mathematics, P. J. Šafárik University in Košice, Košice, Slovakia

² Institute of Mathematics, Poznań University of Technology, Poznań, Poland

³ Guangxi Normal University, Guilin, China

Abstract

Although real-world phenomena are inherently complex, empirical data frequently exhibit structured dependence patterns arising from nested or grouped observations. This article addresses the problem of testing whether a population scale matrix belongs to a commutative quadratic subspace. Under a multivariate elliptically contoured model, we derive both the likelihood ratio test and the Rao score test statistics, and determine the exact distribution of the likelihood ratio test statistic. The performance of the proposed tests is evaluated through simulation studies, and their practical utility is illustrated with two real-data examples. In addition, we consider a high-dimensional setting in which the number of distinct eigenvalues remains fixed as the dimension grows.

Keywords

Beta random variables, Commutative quadratic subspace, Doubly multivariate model, Elliptical contoured distribution, Likelihood ratio test, Rao score test.

References

- [1] Filipiak, K., John, M., Liang, Y. (2024). Testing covariance structures belonging to a quadratic subspace under a doubly multivariate model. *TEST*, 33, 847–876.
- [2] Seely, J. (1971). Quadratic subspaces and completeness. *Ann. Math. Statist.*, 42, 710–721.
- [3] Olkin, I. (1973). Testing and estimation for structures which are circularly symmetric in blocks. Report no. OLK NSF 67, Stanford University.
- [4] Sperling, M. R., et al. (1990). Subcortical metabolic alterations in partial epilepsy. *Epilepsia*, 31, 145–155.
- [5] Worsley, K., et al. (1991). A linear spatial correlation model, with applications to PET. *J. Am. Stat. Assoc.*, 86, 55–67.

Test of the quadratic subspace under elliptically contoured distribution

Daniel Klein¹, Malwina Mrowińska², and Ali Ali Taha¹

¹ P. J. Šafárik University, Košice, Slovakia

² Poznan University of Technology, Poland

Abstract

We study tests for verifying whether a covariance matrix belongs to a quadratic subspace generated by block-repeated Kronecker products up to permutation, assuming that the observation matrix follows a matrix-variate elliptically contoured distribution. Under the null hypothesis, the covariance matrix is represented as a sum of Kronecker-structured components embedded through orthogonal projections. We derive the likelihood ratio test statistic and characterize its finite-sample null distribution through the characteristic function of $-n \log \Lambda$. In addition, we derive the modified Rao score test statistic and obtain its Wishart representation under the null hypothesis. The performance of both tests is compared by Monte Carlo simulations, including their convergence under the null hypothesis to the limiting chi-square distribution and their empirical power.

Keywords

Elliptically contoured distribution, Covariance structure, Quadratic subspace, Rao score test, Likelihood ratio test.

Acknowledgments

The author was supported by the project no. 0213/SBAD/0122 and by the Slovak Research and Development Agency under the contract no. APVV-21-0369 and by grant VEGA No. 1/0585/24.

Covariance structure approximation under high-dimensionality with the use of improved Kullback-Leibler divergence

Monika Mokrzycka¹, Jolanta Pielaszkiewicz²,
and Johan Alenlöv²

¹ Institute of Plant Genetics, Polish Academy of Sciences, Poznań, Poland

² Linköping University, Sweden

Abstract

Covariance approximation and structure identification become challenging in high-dimensional settings when the sample covariance matrix is singular. We propose to utilize an improved Kullback-Leibler divergence framework for comparing and approximating Kronecker structured covariance matrices for matrix-variate elliptically contoured distributions. The approach jointly estimates a regularization parameter and the underlying covariance matrix by positive-definite separable covariance structures. Method applies number of penalty functions representing different prior distributions on underlying parameter. Simulation studies show that the proposed method reliably identifies covariance structures as well as approximate structural parameters in high-dimensional settings. Applications to food-habit, barley, and gene-expression data illustrate its practical usefulness across longitudinal, repeated-measurement, and phylogenetic settings. The proposed framework provides a flexible approach to regularized covariance approximation and structure identification when conventional methods are hindered by high dimensionality and singularity problems.

Keywords

Regularization, Elliptical distribution, Kronecker product, Covariance matrix, Kullback-Leibler divergence.

Acknowledgments

This research was funded by the NAWA Bekker program under project no. BPN/BEK/2023/1/00142 (M. Mokrzycka).

Special Session: Machine Learning and Statistical Inference

Tomasz Górecki

Adam Mickiewicz University, Poznań, Poland

Abstract

The session focuses on current trends at the intersection of machine learning and classical statistical inference. It covers topics such as model interpretability, uncertainty quantification, and extensions of linear methods for analyzing large and complex data sets. The aim is to highlight how statistical principles contribute to the development of reliable and transparent AI models.

Invited Speakers:

Wojciech Rejchel (Poland)
Łukasz Smaga (Poland)
Piotr Sulewski (Poland)

Contributed Speakers:

Johannes Forkman (Sweden)

Matrix reduced-rank approximation through cross-validation

Marisol García-Peña¹, Sergio Arciniegas-Alarcón²,
Johannes Forkman³, and Diego Duarte⁴

¹ Pontificia Universidad Javeriana, Bogotá, Colombia

² Universidade Federal da Bahia, Salvador, Brasil

³ Swedish University of Agricultural Sciences, Uppsala, Sweden

⁴ Universidad de La Sabana, Chía, Colombia

Abstract

Low-rank approximation of data matrices is used in many statistical methods, including principal component analysis, and is commonly obtained through singular value decomposition (SVD) by retaining the leading singular values and their associated principal components. A key challenge is determining the optimal rank of the approximation. Cross-validation for rank selection was pioneered by Wold [8] and Eastment and Krzanowski [4,7]. Several alternative approaches have been proposed, see for example [2,3,6]. More recently, Arciniegas-Alarcón et al. [1] introduced a leave-one-out cross-validation procedure in which each observation is omitted in turn, imputed, and subsequently predicted using an SVD-based low-rank approximation. The optimal rank is then selected as the one that minimises the prediction error in a least-squares sense.

This study extends the method of Arciniegas-Alarcón et al. [1] in two respects. First, leave-group-out cross-validation is adopted in place of leave-one-out cross-validation since the latter may be asymptotically inconsistent [5] and thus could have lower predictive ability. Second, several alternative imputation methods are evaluated. The proposed approaches are assessed through simulation studies and applications to real genotype-by-environment data matrices. The results indicate that the most effective strategy combines an initial imputation based on either projection onto the model plane or an approximate Bayesian bootstrap with subsequent fitting of an SVD-derived multiplicative model.

Keywords

Low-rank approximation, Missing data imputation, Principal component analysis, Rank selection.

Acknowledgments

The first author thanks the Pontificia Universidad Javeriana for the support through project 21243. The fourth author thanks the Faculty of Engineering.

The High Performance Computing Center ZINE was employed the simulation study.

References

- [1] Arciniegas-Alarcón, S., García-Peña M., Krzanowski, W. J., Rengifo, C. (2024). Cross-validation to select the optimum rank for a reduced-rank approximation to multivariate data. *J. Crop Improv.*, 38, 344–367.
- [2] Bro, R., Kjeldahl, K., Smilde, A. K., Kiers, H. A. (2008). Cross-validation of component models: a critical look at current methods. *Anal. Bioanal. Chem.*, 390, 1241–51.
- [3] Camacho, J., Ferrer, A. (2014). Cross-validation in PCA models with the element-wise k-fold (ekf) algorithm: practical aspects. *Chemom. Intell. Lab. Syst.*, 131, 37–50.
- [4] Eastment, H. T., Krzanowski, W. J. (1982). Cross-validators choice of the number of components from a principal component analysis. *Technometrics*, 24, 73–77.
- [5] Eshghi, P. (2014). Dimensionality choice in principal component analysis via cross-validators methods. *Chemom. Intell. Lab. Syst.*, 130, 6–13.
- [6] Josse, J., Husson, F. (2012). Selecting the number of components in principal component analysis using cross-validation approximations. *Comput. Stat. Data Anal.*, 56, 1869–1879.
- [7] Krzanowski, W.J., Kline, P. (1995). Cross-validation for choosing the number of important components in principal component analysis. *Multivariate Behav. Res.*, 30, 149–165.
- [8] Wold, S. (1978). Cross-validators estimation of the number of components in factor and principal components models. *Technometrics*, 20, 397–405.

Positive unlabeled data and prior shift estimation

Jan Mielniczuk^{1,2}, Wojciech Rejchel³,
and Paweł Teisseyre^{1,2}

¹ Institute of Computer Science, Polish Academy of Sciences, Warsaw, Poland

² Warsaw University of Technology, Poland

³ Nicolaus Copernicus University, Toruń, Poland

Abstract

We study estimation of a class prior for unlabeled target samples which possibly differs from that of the source population. Moreover, it is assumed that the source data is partially observable: only samples from the positive class and from the whole population are available (PU learning scenario) [1]. We introduce a novel direct estimator of the class prior which avoids estimation of posterior probability in both populations and has a simple geometric interpretation. It is based on a distribution matching technique together with kernel embedding in a Reproducing Kernel Hilbert Space and is obtained as an explicit solution to an optimisation task.

In [2], we establish its asymptotic consistency, as well as an explicit, nonasymptotic bound on its deviation from the unknown prior, which is computable in practice. We also study finite sample behaviour for synthetic and real data and show that the proposal works consistently on par or better than its competitors.

Keywords

Binary classification, Partially observed data, Distribution shift.

References

- [1] Bekker, J., Davis, J. (2020). Learning from positive and unlabeled data: a survey. *Mach. Learn.*, 109, 719–760.
- [2] Mielniczuk, J., Rejchel, W., Teisseyre, P. (2026). Prior shift estimation for positive unlabeled data through the lens of Kernel embedding. In: *Proceedings of the 29th International Conference on Artificial Intelligence and Statistics (AISTATS)*.

Testing in matched-pairs functional data with missingness in one arm

Marléne Baumeister¹, Łukasz Smaga²,
and Markus Pauly^{1,3}

TU Dortmund University, Germany

Adam Mickiewicz University, Poznań, Poland

Research Center Trustworthy Data Science and Security, UA Ruhr, Germany

Abstract

We consider the matched-pairs testing problem for functional data when some observations from the second measurement are missing. Such incomplete designs arise, for example, in longitudinal medical studies when not all patients participate in a follow-up examination. Our motivating application concerns functional measurements collected from patients with multiple sclerosis at consecutive visits. We propose test statistics that use all available data without discarding incomplete pairs or imputing the missing functional observations. We establish the asymptotic properties of the proposed statistics. Since their null asymptotic distributions are complicated and depend on unknown parameters, direct asymptotic calibration is of limited practical use. We therefore develop parametric bootstrap procedures and show that the resulting tests are asymptotically valid and consistent. Their finite-sample performance is investigated through simulation studies, and the methodology is illustrated using the motivating medical data. The proposed approach provides uncertainty-aware inference for complex, incompletely observed functional data.

Keywords

Bootstrap method, Functional data, Missing data, Repeated measurements.

Acknowledgments

The work of M. Baumeister and M. Pauly were supported by the German Research Foundation project, Grant no. 565106208. The work of Ł. Smaga was funded by the National Science Centre, Poland under the Weave-UNISONO (grant no. 2025/07/Y/ST6/00001).

Very simple and relatively precise mapping of the normal quantile function intended to generate normal pseudo-random numbers fastly

Piotr Sulewski¹ and Antoni Drapella²

¹ Pomeranian University in Slupsk, Poland

² Professor emeritus, Poland

Abstract

This paper puts forward both simple and precise method of mapping of the normal quantile function. Our method aims to fastly generate normal pseudo-random numbers (NPRNs). We use the method of Cholesky (LU decomposition). In order to get in touch with accuracy of interpolation, we define two discrepancy measures namely local and total ones. We compare our very simple and relatively precise mapping of the normal quantile function with other methods of mapping the Gaussian quantile function. Cubic splines are applied to shorten time of generating NPRNs. We compare our fast generator with the generator implemented in R software and realized using the “**rnorm**” function. Products of generators in question are compared to each other in terms of goodness-of-fit tests statistics. In addition, trend tests and seasonality test were performed. The authors also propose to use order statistics to test the generator quality. All the calculations are performed in the R software. Proposed methods are also implemented in the VBA (Excel) and Mathcad. Proposed NPRN generator is a very interesting alternative to the commonly used NPRN generator from the R software.

Keywords

Normal distribution, Quantile function, Pseudo-random number generator, Monte Carlo simulation.

References

- [1] Eidous, O. M., Al-Rawwash, M. Y. (2025). Approximations for standard normal distribution function and its invertible. *J. Algorithms Comput. Technol.*, 19, 174.
- [2] Thomas, D. B., Luk, W., Leong, P. H. W., Villasenor, J. D. (2007). Gaussian random number generators. *ACM Comput. Surv.*, 39(4), 1–38.
- [3] Sulewski, P. (2019). Comparison of normal random number generators. *Pol. Stat.*, 64(7), 5–31.
- [4] Wichura, M. J. (1988). Algorithm AS 241: The percentage points of the normal distribution. *J. R. Stat. Soc. Ser. C Appl. Stat.*, 37(3), 477–484.

Special Session: Order Statistics

Magdalena Szymkowiak

Poznan University of Technology, Poland

Abstract

Order Statistics special session, focuses on the theoretical advancements and practical applications of ordered random variables in both continuous and discrete cases. The session is dedicated to all aspects of ordered data, including: distribution theory and related probability models, reliability theory and survival analysis, statistical inference for ordered and censored data, as well as interesting related applications.

Invited Speakers:

Agnieszka Goroncy (Poland)
Paulo Eduardo Oliveira (Portugal)
Tomasz Rychlik (Poland)
Magdalena Szymkowiak (Poland)

Contributed Speakers:

Jagoda Papis (Poland)

Discrete-time signatures of coherent systems: Properties, stochastic orderings and transformations

Naraswamy Balakrishnan¹, He Yi²,
and Agnieszka Goroncy³

¹ McMaster University, Canada

² Beijing University of Chemical Technology, China

³ Nicolaus Copernicus University in Toruń, Poland

Abstract

In the study of coherent systems in reliability theory, the concept of the system signature, originally introduced by Samaniego [1] and subsequently developed by Kochar, Mukerjee, and Samaniego [2], has become a fundamental tool for analysing the reliability properties of systems, assessing their structural efficiency, and developing inferential methods based on system lifetime data. Several extensions of the signature concept have subsequently been proposed in the literature, including the survival signature, the extended signature, the joint signature, and the ordered system signature. Previous research has focused predominantly on coherent systems composed of components with continuous lifetimes. In practice, however, such systems often consist of components with discrete lifetimes. This situation arises, for example, when the system is monitored only at discrete time points or when it operates in cycles and the observations represent the number of successfully completed cycles before failure occurs. In [3], we consider coherent systems composed of independent components with discrete lifetimes and introduce the concept of a discrete-time signature. We discuss its fundamental properties, establish stochastic ordering results, and derive transformation formulas that facilitate comparisons between the signatures of systems with different numbers of components. In addition, we present several examples illustrating the theoretical results.

Keywords

Reliability system, Discrete-time signatures.

References

- [1] Samaniego, F. J. (1985). On closure of the IFR class under formation of coherent systems. *IEEE Trans. Reliab.*, R-34, 69–72.

- [2] Kochar, S., Mukerjee, H., Samaniego, F. J. (1999). The “signature” of a coherent system and its application to comparisons among systems. *Nav. Res. Logist.*, *46:5*, 507–523.
- [3] Balakrishnan, N., Yi, H., Goroncy, A. (2026). On the notion of discrete-time signature and some associated properties and results. *Probab. Eng. Inf. Sci.*, *40:3*, 355–375.

Stochastic dominance between linear combination

Tommaso Lando¹ and Paulo Eduardo Oliveira²

Department of Economics, University of Bergamo, via dei Caniana 2, 24127,
Bergamo, Italy

CMUC, Department of Mathematics, University of Coimbra, Portugal

Abstract

Convex combinations of independent and identically distributed random variables without a finite mean can behave in a strikingly different way from the finite-mean case. First results proved the stochastic dominance of convex combinations with respect to the underlying distribution [2], and fostered the interest in enlarging the class of distributions for which the dominance holds and in finding tractable conditions implying this dominance [4,1]. The more general comparison between arbitrary combinations of variables was first discussed in [3], identifying a large class of distribution and a suitable majorization relation between the combination coefficients such that the stochastic dominance still holds.

We discuss a relaxation of the class of distributions that implies the dominance between convex combinations. Moreover, for the specific case of sample means, we discuss some distribution free results that also allow to derive dominance relations. From the application point of view, we propose a Kolmogorov-Smirnov one-sample type of test to decide upon this dominance, characterising its asymptotic properties and providing some insight on its numerical performance.

Keywords

Bootstrap, Empirical process, Infinite mean, Majorization, Nonparametric test, Sample mean, Stochastic order.

Acknowledgments

P.E.O. acknowledges financial support by the Centre for Mathematics of the University of Coimbra (CMUC, <https://doi.org/10.54499/UID/00324/2025>) under the Portuguese Foundation for Science and Technology (FCT), Grants UID/00324/2025 and UID/PRR/00324/2025.

References

- [1] Arab, I., Lando, T., Oliveira, P. E. (2025). Convex combinations of random variables stochastically dominate the parent for a new class of heavy tailed distributions. *Electron. Commun. Probab.*, 30, 1–11.

- [2] Chen, Y., Embrechts, P., Wang, R. (2025). Technical note – an unexpected stochastic dominance: Pareto distributions, dependence, and diversification. *Oper. Res.*, 73, 1336–1344.
- [3] Chen, Y., Hu, T., Shneer, S., Zou, Z. (2025). Stochastic dominance for linear combinations of infinite-mean risks. *arXiv:2505.01739 [math.PR]*.
- [4] Müller, A. (2025). Some remarks on the effect of risk sharing and diversification for infinite mean risks. *ASTIN Bulletin, Published online*, 1–10.

Preservation of IFR property for systems with two or three minimal paths of the same length

Jagoda Papis and Magdalena Szymkowiak

Poznań University of Technology, Poland

Abstract

In this work we study increasing failure rate property (IFR property) preservation appearing in reliability theory. That means that if components have this aging property, then the system inherits it. We investigate whether coherent systems with the same length of minimal paths preserve IFR property. Such studies support both the design of new systems and the maintenance of existing ones. They contribute to the development of reliable systems while helping to keep implementation and operating costs reasonable.

We prove that coherent systems with two or three minimal paths of the same length indeed inherit mentioned property. Proves of this statement were not trivial. Our approach needed more precise description of such systems, which can be seen in this work.

These developments extend existing results on the topic of aging properties preservation under the construction of systems and is another approach on proving hypothesis that all systems with the same length of minimal paths preserve IFR property, which was presented in [3].

Keywords

Reliability theory, Coherent systems, Aging classes, IFR property.

Acknowledgments

The first author has been partially supported by PUT under project no. 0213/SBAD/0124 and the second author by PUT under project no. 0211/SBAD/0126.

References

- [1] Navarro, J. (2022). *Introduction to System Reliability Theory*. Springer.
- [2] Navarro, J., Rychlik, T., Szymkowiak, M. (2025). Key distributions in the preservation of aging classes under the construction of systems. *J. Comput. Appl. Math.*, 469, 116650.
- [3] Papis, J., Szymkowiak, M. (2026). Preservation of IFR property under the construction of systems. *Ric. Mat.*, 75, 241–253.
- [4] Rychlik, T., Szymkowiak, M. (2022). Signature conditions for distributional properties of system lifetimes if coherent lifetimes are i.i.d. exponential. *IEEE Trans. Reliab.*, 71(2), 590–602.

Sharp bounds on L -moments

Tomasz Rychlik¹ and Magdalena Szymkowiak²

¹ Institute of Mathematics, Polish Academy of Sciences, Poland

² Poznan University of Technology, Poland

Abstract

The expectations of specific linear combinations of order statistics from i.i.d. samples with a finite expectation parent distribution are called L -moments and denoted as λ_r , $r = 1, 2, \dots$. They were introduced by Hosking [1]. The L -moments uniquely characterize the distribution, and admit to explicitly recover its quantile function. The first L -moments λ_r , $r = 1, 2, 3, 4$, describe basic shape properties of the distribution: location, dispersion, skewness, and peakedness, respectively. We present sharp lower and upper bounds on L -moments expressed in the scale units based on central absolute moments of the baseline distribution, and the second L -moment being a half of the Gini mean difference. We also determine the distributions attaining these bounds. The talk is based on papers [2] and [3].

Keywords

L -moment, Sharp bound, Attainability condition.

References

- [1] Hosking, J. R. M. (1990). L -moments: analysis and estimation of distributions using linear combinations of order statistics. *J. R. Stat. Soc. Ser. B Stat. Methodol.*, 52, 105–124.
- [2] Rychlik, T., Szymkowiak, M. (2026). Extreme values of scaled L -moments. *Statist. Probab. Lett.* 231, 110626.
- [3] Rychlik, T., Szymkowiak, M. (2026). Optimal bounds for L -moments. Under review.

ϕ -divergence as a measure of the predictability of reliability systems

Narayanaswamy Balakrishnan¹, Maria Kateri²,
and Magdalena Szymkowiak³

¹ McMaster University, Canada

² RWTH Aachen University, Germany

³ Poznan University of Technology, Poland

Abstract

We extend the Jensen–Shannon divergence to a broader class of divergence measures, referred to as the Jensen ϕ -divergence and denoted by JD_ϕ . An important parametric subfamily of JD_ϕ is the Jensen–Cressie–Read power divergence JD_λ , where the parameter λ provides a flexible scaling mechanism. The proposed divergences quantify predictability of reliability systems, with smaller divergence values indicating higher predictability.

In general, the divergence of a system depends both on the lifetime distribution of its homogeneous components and on the system structure, characterized by its system signature (see, e.g., [2]). To illustrate the practical applicability of the Jensen ϕ -divergence, we examine JD_λ for systems with three components and compare the results with those obtained from the Jensen–Shannon divergence introduced by Asadi et al. [1] and preferable system criterion proposed by Toomaj et al. [3].

Keywords

Reliability systems, Lifetime distribution, Signature, Divergence, Predictability measure.

Acknowledgments

The author has been partially supported by PUT under Grant no. 0211/SBAD/0126.

References

- [1] Asadi, M., Ebrahimi, N., Soofi, E. S., Zohrevand, Y. (2016). Jensen-Shannon information of the coherent system lifetime. *Reliab. Eng. Syst. Saf.*, 156, 244–255.
- [2] Rychlik, T., Szymkowiak, M. (2022). Signature conditions for distributional properties of system lifetimes if component lifetimes are iid exponential. *IEEE Trans. Reliab.*, 71, 590–602.

- [3] Toomaj, A., Chahkandi, M., Balakrishnan, N. (2021). On the information properties of working used systems using dynamic signature. *Appl. Stoch. Models Bus. Ind.*, 37, 318–341.

Special Session: Statistics in Applications

Barbora Kessel

University of Gothenburg, Sweden

Abstract

Statistical methods play a crucial role in understanding complex systems and driving data-informed decisions across a wide range of disciplines. This special session aims to showcase innovative applications of statistics in real-world contexts, including but not limited to, science, engineering, economics, healthcare, environmental studies, and social sciences. We invite contributions that demonstrate how statistical thinking and methodology can address practical challenges, uncover patterns in data, and support decision-making under uncertainty. Both methodological advancements motivated by applications and case studies highlighting impactful use of statistical tools are welcome.

Invited Speakers:

Eva Fišerová (Czech Republic)
Viktor Witkovský (Slovakia)

Contributed Speakers:

Julianna Holmberg (Sweden)
Vivien Kardos (Hungary)
Peter Kovacs (Hungary)
Patrik Kołbyko (Poland)
Agnieszka Kwita (Poland)
Lynn Roy LaMotte (USA)
Alicja Lerczak (Poland)
Im Chiew Ng (UK)

Beyond average speed: Functional data analysis of traffic calming effects

Eva Fišerová

Palacký University Olomouc, Czech Republic

Abstract

Traditional evaluations of traffic calming measures typically rely on speeds recorded at selected locations or on summary measures for entire road sections, potentially obscuring spatial variation in drivers' responses. This contribution presents the statistical framework for evaluating changes in continuous speed profiles using floating car data and functional data analysis [1,3].

Individual GPS trajectories are represented as smooth speed functions over distance. Mean speed profiles observed before and after the installation of traffic calming measures are compared using a global permutation test and interval-wise testing [2]. The latter identifies statistically significant road sections with statistically significant road sections while controlling the interval-wise error rate. These sections define zones of influence, within which the magnitude of the change is quantified by aggregated and pointwise relative effects.

The methodology is illustrated using real-world data from Spain and Czechia. At a Spanish urban gateway, significant speed reductions of approximately 6% and 9% were identified over road sections exceeding 500 meters. In contrast, no significant changes were found at eight Czech locations equipped with driver feedback signs. The proposed approach moves the evaluation beyond average speeds by determining not only whether speed profiles differ, but also where the differences occur and how their magnitude varies along the road.

Keywords

Functional data analysis, Speed profiles, Floating car data, Interval-wise testing, Traffic calming.

References

- [1] Elgner, J., Ambros, J., Muñoz, M. A., Valentová, V., Fišerová, E. (2026). Assessing the speed impact of urban traffic calming devices using functional data analysis. *European Transport Res. Rev.*, 18, 14.
- [2] Pini, A., Vantini, S. (2017). Interval-wise testing for functional data. *J. Non-parametr. Stat.*, 29, (2), 407–424.
- [3] Ramsay, J., Silverman, B. W. (2005). *Functional data analysis*. Springer, New York.

Quadratically constrained likelihood estimation for exponential regression under censoring

Julianna Holmberg¹, Laila Hübbert^{1,2},
Dietrich von Rosen¹, and Martin Singull¹

¹ Linköping University, Sweden

² Region Östergötland, Department of Cardiology, Norrköping, Sweden

Abstract

Parametric survival models are commonly estimated using unconstrained maximum likelihood, even when prior knowledge suggests that the effects of covariates on the hazard should satisfy structural constraints. In this presentation, we consider exponential regression under right-censoring with a quadratic constraint, leading to maximum likelihood estimation over a restricted parameter space. A simulation study demonstrates the bias-variance trade-off induced by the constraint and compares the estimator with the standard maximum likelihood estimator and ridge-type penalised estimators. An application to real survival data illustrates the proposed method.

Keywords

Exponential regression, Right-censored survival data, Constrained maximum likelihood.

Acknowledgments

Julianna Holmberg was supported by the Research School in Interdisciplinary Mathematics.

Examining law students' intention to use GenAI systems for professional purposes using PLS modelling

Vivien Kardos and Peter Kovacs

University of Szeged, Hungary

Abstract

Generative artificial intelligence systems (GenAI systems), and within them large language models (LLMs), have become increasingly important in recent years. Their rapid development and growing societal impact also have significant implications for the legal profession and legal education. In 2026, we conducted a survey among 162 primarily first-year law students at the Faculty of Law and Political Sciences of the University of Szeged, Hungary, to examine their attitudes, patterns, and use purposes related to the use of GenAI systems. The study focused on how frequently and for what purposes law students use different GenAI systems; the extent to which they are reluctant to use these tools in general, academic, and legal-professional contexts; how willing and able they are to develop their competences in these areas; and how they assess their own digital competences. The relationships among these factors were analysed using partial least squares structural equation modelling (PLS-SEM), a method that combines factor analysis and regression modelling in a simultaneous analytical framework. The presentation reports the main findings of this empirical study. The results indicate that GenAI systems are playing an increasingly important role in students' learning processes. The main modelling results show that the intention to use GenAI systems for legal studies and future legal work is most strongly and negatively influenced by reluctance towards the use of GenAI systems. At the same time, the frequency of GenAI systems use and openness towards these systems have positive, although moderate, effects on usage intention. Perceived digital competence does not appear to have a direct overall effect on usage intention; however, this result reflects the combined effect of a moderate direct negative relationship and a moderate indirect positive relationship mediated by openness towards GenAI systems.

Keywords

Generative AI systems, Legal education, PLS-SEM.

Linear-Gaussian Laubach–Williams state-space estimation and decomposition of the natural rate of interest – r^* : Evidence from Poland

Paweł Kołbyko

Maria Curie-Skłodowska University in Lublin, Poland

Abstract

The aim of the study is to obtain a time-varying estimate of Poland’s natural rate of interest (NRI) $-r^*$ and to provide a structural decomposition that separates trend growth (g) from other persistent determinants (z) within a semi-structural Laubach–Williams state-space framework. Quarterly data for 2010Q1–2025Q1 are analysed, with parameters estimated by Gaussian maximum likelihood and latent states inferred via Kalman filtering and smoothing in a linear-Gaussian state-space model. The measurement system links centred log industrial production, core inflation, a lagged real-rate input for the IS relation, and an activity-growth measurement equation that anchors the trend-growth state and improves identification of long-run versus cyclical movements. The ex-ante real policy rate is constructed as the nominal policy proxy minus survey-based expected inflation quantified via the Carlson–Parkin mapping. The resulting r^* exhibits pronounced regime variation, remaining moderately positive in the early 2010s, declining into negative territory around 2019–2021, and rising sharply after the pandemic, with a minimum of -1.34% (2020Q3), a maximum of 3.07% (2023Q4), and a sample mean of 0.76% (annualised). Uncertainty is reported using 95% intervals derived from smoothed-state covariances, and robustness checks indicate that the main turning points are preserved under calibration-variance (Q/H) perturbations and selected model-class extensions, while policy-rule closure generates the largest deviations. The results imply that monetary-policy stance diagnostics based on the natural-rate gap should be communicated jointly with uncertainty bands and sensitivity analysis, as inferred levels are specification- and end-sample-sensitive even when regime shifts remain stable. The contribution is a Poland-focused, identification-oriented LW-SSM implementation with an explicit decomposition and a transparent robustness protocol suitable for policy-relevant inference.

Keywords

Natural rate of interest (r^*), Linear-Gaussian inference, Kalman filter, Monetary policy stance.

References

- [1] Laubach, T., Williams, J. C. (2003). Measuring the natural rate of interest. *Rev. Econ. Statist.*, 85(4), 1063–1070.
- [2] Holston, K., Laubach, T., Williams, J. C. (2017). Measuring the natural rate of interest: International trends and determinants. *J. Int. Econ.*, 108, S59–S75.
- [3] Bielecki, M., Błażejowska, A., Brzoza-Brzezina, M., Kuziemska-Pawlak, K., Szafrński, G. (2023). Estimates and projections of the natural rate of interest for Poland and the euro area (NBP Working Paper No. 364). Narodowy Bank Polski.
- [4] Mésonnier, J.-S., Renne, J.-P. (2007). A time-varying “natural” rate of interest for the euro area. *European Econ. Rev.*, 51(7), 1768–1784.
- [5] Orphanides, A., van Norden, S. (2002). The unreliability of output-gap estimates in real time. *Rev. Econ. Statist.*, 84(4), 569–583.
- [6] Lewis, K. F., Vazquez-Grande, F. (2019). Measuring the natural rate of interest: A note on transitory shocks. *J. Appl. Econom.*, 34(6), 1096–1118.
- [7] Brand, C., Mazelis, F. (2019). Taylor-rule consistent estimates of the natural rate of interest (ECB Working Paper Series No. 2257). European Central Bank.
- [8] Benigno, G., Hofmann, B., Nuño, G., Sandri, D. (2024). Quo vadis, r^* ? The natural rate of interest after the pandemic. *BIS Quarterly Rev.*, 17–30.
- [9] Buncic, D. (2024). Econometric issues in the estimation of the natural rate of interest. *Econ. Modelling*, 132, 106641.
- [10] Juselius, M. (2025). Navigating the sea of natural real interest rate estimates (BoF Economics Review No. 3/2025). Bank of Finland.
- [11] Carlson, J. A., Parkin, M. (1975). Inflation expectations. *Economica*, 42(166), 123–138.
- [12] Durbin, J., Koopman, S. J. (2012). *Time Series Analysis by State Space Methods* (2nd ed.). Oxford University Press.

Algorithmic support for judicial sentencing? Evidence from human smuggling cases

Peter Kovacs, Krisztina Karsai, Zsanett Fantoly,
Andor Gal, and Balint Kelemen

University of Szeged, Hungary

Abstract

It has often been raised whether algorithmic support for judicial sentencing can be implemented. We examined this question in a relatively homogeneous case type, namely human smuggling cases, where the Hungarian Criminal Code explicitly identifies the aggravating and mitigating factors that may be taken into account, while not specifying their relative weights. In the empirical analysis, we examined 545 case files with respect to the legal classification of the offence under the Criminal Code. The results of the statistical analysis suggest that judicial practice cannot necessarily be regarded as consistent. Of the nine aggravating factors examined, only three were associated with significantly longer sentences, while of the seven mitigating factors analysed, only three were linked to a statistically significant reduction in sentence length. Although the number of mitigating factors showed a significant but weak correlation with the length of imprisonment, no such significant relationship could be detected for aggravating factors. Since the imposed sentence must be justifiable, we applied HC3-corrected multiple linear regression models and mixed models to estimate the length of imprisonment imposed. Our results show that aggravating and mitigating factors alone have low explanatory power ($R^2 = 0,32$); including the statutory midpoint of the sentencing range in the models improved predictive performance ($R^2 = 0,59$), but also reduced the relative importance of aggravating and mitigating circumstances in determining sentence length. The following factors had a significant effect: the number of persons transported, continuation of the offence, lapse of time, juvenile offender status, and no prior criminal record. In addition, the findings indicate that different judicial panels did not assess the weight of individual factors in the same way, which makes algorithmisation more difficult.

Keywords

Linear regression models, Mixed models.

Statistical comparison of wavelet and Fourier series approximation quality for audio signals

Mateusz John¹ and Agnieszka Kwita²

¹ Poznań University of Technology, Poland

² Wrocław University of Science and Technology, Poland

Abstract

A statistical procedure for comparing the quality of audio signal approximation obtained using the discrete wavelet transform (Daubechies wavelets dbN, $N = 1, \dots, 10$) and Fourier series approximation, estimated by the least squares method, is presented. The audio signal is sampled either uniformly or adaptively, with node density in the adaptive scheme proportional to the local variability of the signal amplitude. This makes it possible to assess the influence of the sampling scheme on reconstruction accuracy. Optionally, the wavelet detail coefficients may be denoised by soft or hard thresholding prior to reconstruction.

Since reconstruction errors obtained for different wavelets arise from the same observation vector, they are treated as dependent samples. Differences between wavelets are assessed using a nonparametric test for repeated measures, followed by post-hoc pairwise comparisons for dependent samples, with a correction for multiple testing and the computation of effect size measures. Additionally, a compression analysis based on M -term approximation in the wavelet basis is performed to determine which wavelet achieves a given reconstruction error level with the fewest retained coefficients.

The entire procedure is implemented in a custom Python application (built using `PyWavelets`, `scipy.stats`, `numpy`), enabling audio signal acquisition, sampling, approximation, and full statistical analysis together with result visualization.

Keywords

Signal approximation, Wavelet transform, Fourier series, Friedman test, Wilcoxon test, Wavelet coefficient compression.

ANOVA: Is 100 years enough?

Lynn R. LaMotte

LSU Health - New Orleans, USA

Abstract

R. A. Fisher introduced analysis of variance for balanced designs in his 1925 book *Statistics for Research Workers*. ANOVA quickly became the most widely used statistical method, particularly in agronomic disciplines.

It partitions total variation into identifiable parts that together account for all possible differences among the cell means. For unbalanced designs, though, no such consensus has been reached in the intervening 100 years, and disputatious colloquy continues to this day. A 2003 paper by Øyvind Langsrud, for example, titled “ANOVA for unbalanced data: Use Type II instead of Type III sums of squares,” has been cited more than 600 times.

The purpose of this presentation is to propose a strategy for unbalanced designs that parallels classical ANOVA for balanced designs. It is based on sums of squares that test exclusively those classical ANOVA effects that are estimable, which in some cases is none, together with other contrasts up to the degrees of freedom available for all differences among the means of the non-empty cells. Techniques are illustrated to express and identify the contrasts that are tested.

Keywords

Unbalanced designs, Extended effects, Empty cells.

Canonical variate analysis for two-way classification applied to a fertilization experiment with winter triticale

Alicja Lerczak¹, Dariusz Kayzer¹, Katarzyna Wielgusz²,
Dorota Czerwińska-Kayzer¹, and Renata Gaj¹

¹ Poznań University of Life Sciences, Poland

² Institute of Natural Fibres and Medicinal Plants, National Research Institute, Poland

Abstract

Canonical variate analysis, a technique applied in multilinear modelling, searches for transformations of the original variables into an orthogonal space in which the between-group variance is maximized. Our study shows how, in a linear model subject to two-way classification, singular value decomposition [2] can be combined with linear combinations [1] of the original variables. Transformations into the same orthogonal space, on the one hand, reveal the relationships between the variables and the experimental objects and, on the other hand, show the dispersion of the experimental units.

Canonical variate analysis was applied to analyze winter triticale yield-related data collected over several years under different fertilizer treatments [3]. The analyzed data consist of the yield of winter triticale, micronutrient uptake, and biomass assessed based on plants collected during the heading stage. The application of canonical variate analysis to the data allows the identification of relationships between parameters of winter triticale and experimental factors. This leads to a much deeper and broader interpretation of the agricultural experimental results.

Keywords

Canonical variate analysis, Singular value decomposition, Linear combination, Two-way classification, Winter triticale.

Acknowledgments

The publication was financed by the Polish Minister of Science and Higher Education as part of the Strategy of the Poznań University of Life Sciences for 2024–2026 in the field of improving scientific research and development work in priority research areas (PREIDUB project).

References

- [1] Campbell, N. A., Atchley, W. R. (1981). The geometry of canonical variate analysis. *Syst. Zool.* 30, 268–280.
- [2] Lejeune, M., Caliński, T. (2000). Canonical analysis applied to multivariate analysis of variance. *J. Multivar. Anal.* 72, 100–119.
- [3] Lerczak, A., Prałat, T., Sychalski, M., Kayzer, D., Kukawka, R., Gaj, R. (2025). Impact of hybrid fertilization on winter triticales yield and its stability based on SVD analysis. *Sustainability*, 17, 11385.

The resilience of random intercepts and slopes models to violations of the linearity assumption in the treatment group

Im Chiew Ng¹, Katy Morgan¹, Jennifer Nicholas¹,
and James Carpenter^{1,2}

¹ London School of Hygiene and Tropical Medicine, UK

² University College London, UK

Abstract

When planning a clinical trial with a continuous outcome, researchers may know that the trajectory in the control group will be approximately linear. However, they may be less sure of the trajectory in the treatment group. If this is also linear, one possible analysis method is the random intercepts and slopes (RIS) model, which estimates the mean difference in slopes between the treatment and control groups. However, if it is not linear, the assumptions of the RIS do not hold. Other methods which do not assume linearity, such as a mixed model for repeated measures (MMRM) may perform better as they do not assume linearity but estimates the mean difference between groups. We conducted a simulation study to assess the performance of the RIS and other analysis models that relax the linearity assumption in the setting of a linear control-group trajectory and non-linear treatment-group trajectory. We considered a wide range of scenarios, including seven treatment-group trajectories, several between to within-participant variance ratios, and ranging the spacing of follow-up visits.

We found that RIS is relatively robust to small amounts of non-linearity in treatment group trajectories when the between-participant variance is large compared to within-participant, with only small amounts of bias (calculated from the mean difference between groups at final visit) and maintaining statistical power at the expected level. However, in other scenarios we found the RIS had a large amount of bias. We found a two-stage method that first tests for non-linearity before choosing either the RIS or MMRM performs moderately well across a wide range of scenarios. We also present analytic results that support our simulations and consider the implications of these results in the design of clinical trials by applying these methods to real trial data.

Keywords

Random intercepts and slopes, Linearity assumption, Randomised clinical trial, Mixed model with repeated measures, Statistical power.

Acknowledgments

The author was supported by funding from the National Institute of Health Research under the Pre-Doctoral Fellowship (NIHR304900)

Computing with characteristic functions: From numerical inversion to multivariate logistic goodness-of-fit

Viktor Witkovský

Slovak Academy of Sciences, Bratislava, Slovakia

Abstract

Characteristic functions provide a natural computational interface between probability models and statistical inference, especially when densities or distribution functions are unavailable in closed form or are numerically cumbersome. This lecture first gives a brief overview of characteristic-function computation and numerical inversion as implemented in the MATLAB CharFunTool [3], including numerical algorithms for univariate and bivariate distributional calculations [1].

These computational tools provide a natural basis for extending characteristic-function methods from probability-distribution evaluation to statistical inference based on empirical characteristic functions.

The main part is devoted to a recently developed goodness-of-fit test for multivariate logistic distributions [2], illustrating how the computational methodology underlying CharFunTool can be extended from distributional calculations to statistical inference. After affine standardization of the observations, the test compares the empirical characteristic function with the theoretical characteristic function of the standard multivariate logistic model through a weighted L^2 distance. Its implementation combines numerical characteristic-function calculations and Monte Carlo methods with an analytical reduction of the defining multidimensional integral.

Radial symmetry and Fourier/Hankel-transform techniques lead to an explicit and computationally efficient statistic, while characteristic-function algorithms developed within the CharFunTool framework support multivariate distributional calculations, simulation, and visualization. The procedure is affine invariant, its asymptotic behaviour is established under the null hypothesis and contiguous alternatives, and finite-sample critical values are obtained by Monte Carlo calibration.

Simulation studies show accurate size and strong power relative to multivariate Kolmogorov–Smirnov, energy, and chi-square procedures, particularly under heavy-tailed and higher-dimensional alternatives. Applications to bivariate educational test scores and four-dimensional Iris measurements illustrate practical model validation. The presentation emphasizes a broader message: characteristic functions can serve not only as theoretical descriptors of distributions, but also as an effective computational route from probabilistic

modelling and numerical distribution theory to implementable multivariate statistical inference.

Keywords

Characteristic functions, CharFunTool, Numerical inversion, Empirical characteristic function, Goodness-of-fit, Multivariate logistic distribution.

Acknowledgments

The results on multivariate logistic goodness-of-fit are joint work with Božidar V. Popović and Andjela Mijanović. This research was supported by the Ministry of Education, Science and Innovation of Montenegro; the Ministry of Education, Research and Development and Youth of the Slovak Republic; and the Slovak Research and Development Agency under the bilateral project SK-MNE-25-0025, and by the Scientific Grant Agency of the Ministry of Education, Research, Development and Youth of the Slovak Republic and the Slovak Academy of Sciences under projects VEGA 2/0120/24 and VEGA 2/0094/26.

References

- [1] Mijanović, A., Popović, B. V., Witkovský, V. (2023). A numerical inversion of the bivariate characteristic function. *Appl. Math. Comput.*, 443, 127807.
- [2] Popović, B. V., Mijanović, A. and Witkovský, V. (2026). A goodness-of-fit test for multivariate logistic distributions. *Statistics*.
DOI: 10.1080/02331888.2026.2624510.
- [3] Witkovský, V. (2020). *CharFunTool: The Characteristic Functions Toolbox (MATLAB)*.
GitHub repository: <https://github.com/witkovsky/CharFunTool/>.

Special Session: Experimental Designs

Maryna Prus

University of Hohenheim, Germany

Abstract

Experimental design is a fundamental area of statistics that focuses on planning studies to collect data efficiently and draw reliable conclusions. This session is devoted to model-based design of experiments, where statistical models are used to select the most informative experimental settings before data are collected. The contributions cover a range of applications, including agricultural science and psychology, and present both analytical solutions and computational methods for constructing optimal designs. The proposed methods are illustrated using real data examples.

Invited Speakers:

Stephan Bark (Germany)
Nadja Malevich (Germany)
Maryna Prus (Germany)

Contributed Speakers:

Małgorzata Graczyk (Poland)

A Bayesian updating framework for long-term multi-environment trial data in plant breeding

Stephan Bark, Maryna Prus, and Hans-Peter Piepho

University of Hohenheim, Germany

Abstract

Multi-environment trial (METs) are essential in plant breeding to assess genotype performance across diverse climatic zones ([2]). Consequently, the key question of the design regarding the optimal allocation of a fixed total number of trials across climatic zones arises. An optimal allocation allows for more precise estimates of genotype performance within each zone, taking into account local environmental conditions ([3]). Their A-optimal design calculation starts with a linear mixed model (LMM) analysis estimated via residual (restricted) maximum likelihood (REML). Some variance components may be estimated as exactly zero when signals are weak, unintentionally removing random effects and reducing accuracy.

A fully Bayesian reformulation calculated via Markov chain Monte Carlo (MCMC) methods uses priors on variance components to keep them positive, yielding more realistic estimates ([6]). This has been addressed in this study ([1]) applied to the MET data of [4]. The corresponding Bayesian Conditional Mixed Model (BCMM) reformulation incorporates historical prior information and quantifies uncertainty in the design. To the best of the available knowledge, this study comprises one of the first serious attempts to objectively inform priors in the context of historical MET data subdivided into L successive multi-year data windows. Following the terminology of [5], this approach is referred to as the Bayesian updating approach.

It is crucial that the structural assumptions of the variance parameter priors are maintained throughout the windows to remain within the Bayesian mixed model framework. Posterior Gibbs samples of variance components are used to approximate prior distributions for the next successive data window via maximum likelihood estimation.

Our results indicate that Bayesian methods can enhance the quality of variance component estimates and consequently of optimal design strategies. The findings provide insights into improving the efficiency and accuracy in MET data analysis for more reliable genotype selection. Key challenges include prior assumptions and residual variance heterogeneity, which can pose convergence issues in Bayesian posterior sampling. Future research should reformulate design criteria to address variance heterogeneity, explore alternative prior specifications, and enhance sampling strategies. Moreover, extending

the modeling to include environmental covariates could lead to a remarkable gain in accuracy.

Keywords

Bayesian linear mixed model, Multi-environment trial, Residual (restricted) maximum likelihood, Markov chain Monte Carlo, Optimal design.

Acknowledgments

We are grateful to the Bangladesh Rice Research Institute (BRRI) for generating, maintaining, and generously providing access to the long-term multi-environment trial data used in this study. The exploration of our statistical method would not have been possible without all the efforts of BRRI scientists, field staff, and technical personnel!

We thank Reinhard Meister for suggesting a fully Bayesian analysis of historical MET data at a seminar held by Hans-Peter Piepho at the Berliner Kolloquium “Statistische Methoden in der empirischen Forschung” in 1998. It has taken quite some time to get there, but better late than never!

References

- [1] Bark, S., Malik, W. A., Prus, M., Piepho, H.-P., Schmid, V. (2026). A Bayesian updating framework for long-term multi-environment trial data in plant breeding. Submitted manuscript, arXiv:2604.16203.
- [2] Kleinknecht, K., Möhring, J., Singh, K.-P., Zaidi, P.-H., Atlin, G.-N., Piepho, H.-P. (2013). Comparison of the performance of best linear unbiased estimation and best linear unbiased prediction of genotype effects from zoned Indian maize data. *Crop Science*, 53, 1384–1391.
- [3] Prus, M., Piepho, H.-P. (2026). Optimizing the allocation of trials to sub-regions in crop variety testing with multiple years and locations. *J. Agric. Biol. Environ. Stat.*, 31, 287–307.
- [4] Rahman, N. M. F., Malik, W. A., Kabir, M. S., Baten, M. A., Hossain, M. I., Rudra Paul, D. N. et al. (2023). 50 years of rice breeding in Bangladesh: genetic yield trends. *Theoretical and Applied Genetics* 136, 18.
- [5] Sorensen, D. and Gianola, D. (2002). *Likelihood, Bayesian, and MCMC Methods in Quantitative Genetics*. Springer, New York.
- [6] Studnicki, M., Piepho, H.-P. (2024). Hierarchical modelling of variance components makes analysis of resolvable incomplete block designs more efficient. *Theoretical and Applied Genetics* 137, 134.

Selected properties of estimators in the model of weighing design

Małgorzata Graczyk and Bronisław Ceranka

Poznań University of Life Sciences, Poland

Abstract

Several aspects of the determination of optimal weighing designs are investigated. A methodological approach is proposed for identifying the best design under the D-optimality criterion. The analysis is conducted under the assumption that the measurements exhibit equal positive correlations and have identical variances. Related aspects of this problem have also been addressed in previous studies: [1], [2].

Keywords

D-optimality, weighing design.

References

- [1] Ceranka, B., Graczyk, M. (2018). Highly D-efficient weighing designs for an even number of objects. *REVSTAT*, 16(4), 475–486.
- [2] Graczyk M., Ceranka B. (2025). A study of the efficiency of chemical balance weighing designs. *Biom. Letters*, 62(1), 91–100.

Estimation and design strategies for inter-individual differences in an overparameterized paired comparison model

Heiko Großmann¹, Heinz Holling², Nadja Malevich^{1,2},
and Rainer Schwabe¹

¹ Otto von Guericke University Magdeburg, Germany

² University of Münster, Germany

Abstract

Motivated by psychological applications, we consider an overparameterized general linear paired comparison model with alternative-specific intercepts and respondent-specific vectors of trait parameters. We are interested in estimating contrasts of the trait parameter vectors. We study this problem both in the original model and in a transformed Gauss-Markov model arising from taking differences between response vectors. We show that the best linear unbiased estimator of a contrast in the original model coincides with the GLS and OLS estimators of the parameter vector in the transformed model. As an application, we discuss an optimal design problem that can be solved more easily in the transformed model setting.

Keywords

Paired comparison experiment, Linear model, Contrast estimation, Optimal design.

Acknowledgments

This research was funded by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) under grant GR 6301/2-1|HO 1286/20-1|SCHW 531/18-1.

Optimizing the allocation of trials to sub-regions in crop variety testing

Maryna Prus

University of Hohenheim, Germany

Abstract

New crop varieties are extensively tested in multi-environment trials to provide a basis for recommendations to farmers. When the target population of environments is large, it is often advantageous to divide it into sub-regions. If the same set of genotypes is tested across all sub-regions, a linear mixed model (LMM) can be fitted with random genotype-within-sub-region effects. The first analytical results on optimizing the allocation of trials to sub-regions were presented in [1] and [2]. In [1] the focus was on single-year experiments. In [2] multi-year experiments were investigated, for which the same allocations of trials (designs) were assumed for all years. In practice, however, the number of genotypes and the total number of locations often vary from year to year. Consequently, it is more realistic to allow different designs in different years. In the present work, a general LMM is considered in which allocations of trial may differ across years. The proposed analytical solution, formulated in terms of design criteria and optimality conditions, enables an easy and fast computation of optimal designs. The obtained results are illustrated using real-data examples.

Keywords

Best linear unbiased predictor, Experimental design, Linear mixed model, Multi-environment trials, Optimal design.

Acknowledgments

This work has been funded by the DFG-project 533007746.

References

- [1] Prus, M., Piepho, H.-P. (2021). Optimizing the allocation of trials to sub-regions in multi-environment crop variety testing. *J. Agric. Biol. Environ. Stat.*, 26, 267–288.
- [2] Prus, M., Piepho, H.-P. (2026). Optimizing the allocation of trials to sub-regions in crop variety testing with multiple years and locations. *J. Agric. Biol. Environ. Stat.*, 31, 287–307.

Session: Multivariate Statistics

Contributed Speakers:

Katarzyna Filipiak (Poland)

Mateusz John (Poland)

Monika Mokrzycka (Poland)

Kenneth Nordström (Finland)

Anna Szczepańska-Álvarez (Poland)

Dietrich von Rosen (Sweden)

The power of blocks: Matrix operators for statistics

Katarzyna Filipiak

Poznan University of Technology, Poland

Abstract

Partitioned matrices naturally appears in multivariate models, especially when particular covariance structure is assumed. In this talk we present, how two block-versions of standard matrix operator, trace of a matrix, can be used to simplify estimation of covariance matrix.

Keywords

Block-trace, Partial-trace, Projection, Estimation of partitioned matrix.

Acknowledgments

The author was supported by the PUT project no. 0213/SBAD/0122.

Testing hypotheses on covariance matrices within linear structure spaces

Mateusz John and Adam Mieldzioc

Poznań University of Technology, Poland

Abstract

The problem of testing whether the true covariance matrix belongs to a given non-quadratic linear space of symmetric matrices is considered. This distinction matters because, as shown by [1] and [2], the maximum likelihood estimator (MLE) reduces to an orthogonal projection of the sample covariance matrix only in the quadratic case; for non-quadratic subspaces, no closed form solution exists, and numerical optimization is required. To overcome this difficulty, two alternative estimators are proposed: the restricted maximum likelihood (RML) estimator, obtained via numerical optimization, and a shrinkage structured (SH) estimator and use them in place of the standard MLE within the classical likelihood ratio test (LRT) and Rao score test (RST).

Our main theoretical contribution establishes that the LRT statistics based on the RML and SH estimators retain the same asymptotic chi-squared distribution as their classical MLE-based counterparts. A simulation study is conducted to examine the finite sample behavior of both the LRT and RST statistics: for the LRT, this includes verifying the rate of convergence to the limiting distribution, while for the RST, the chi-squared distribution is used as an empirical reference. The study further compares the power of the tests and the computation times associated with each estimator.

The proposed methodology is further illustrated through an application to real data, where model selection criteria are used to determine the most appropriate covariance structure under each method.

Keywords

Covariance matrix, Testing hypotheses, Linear structures, Toeplitz matrix, Shrinkage estimator, Restricted maximum likelihood estimator.

Acknowledgments

Mateusz John was supported by Statutory Activities no. 0213/SBAD/0121.

References

- [1] Filipiak, K., John, M., Markiewicz, A. (2020). Comments on maximum likelihood estimation and projections under multivariate statistical models. In:

- Holgersson, T., Singull, M. (Eds.), *Recent Developments in Multivariate and Random Matrix Analysis* (pp. 51–66). Springer.
- [2] Zehna, P. W. (1966). Invariance of maximum likelihood estimators. *Ann. Math. Stat.*, 37(3), 744–744.

Quasi-shrinkage estimation of block covariance matrix with structured off-diagonal blocks

Monika Mokrzycka¹ and Malwina Mrowińska²

¹ Institute of Plant Genetics PAS, Poznań Poland

² Poznań University of Technology, Poland

Abstract

The block covariance matrix partitioned into four sub-blocks, with unstructured diagonal blocks and structured cross-covariance matrices is assumed. The structure of the off-diagonal block corresponds to the appropriate part of autoregression of order one structure. The least squares estimation method for such a structured covariance matrix as well as its improvement by shrinkage method are presented.

Keywords

Block covariance matrix, Autoregression of order one, Least squares estimation, Shrinkage method.

Convexity of the Moore–Penrose inverse

Kenneth Nordström

University of Helsinki, Finland

Abstract

In this talk I shall review convexity properties of the Moore–Penrose inverse matrix function. Convexity refers here to the partial (Löwner) ordering of Hermitian matrices. The convexity result for positive definite matrices was no doubt known to Karl Löwner and his students in the 1930s (and perhaps more broadly), but appeared in print later due to its use in problems of optimum experimental design and multivariate probability (Chebyshev-type) inequalities. Motivated in part by such optimality problems, various extensions to nonnegative definite matrices using generalized inverses have also been derived; see Section 1 of [1] for a survey of such convexity results for positive definite as well as nonnegative definite matrices. Also, in [1] the question was raised whether in the convexity inequality (with $\bar{\lambda} := 1 - \lambda$)

$$(\lambda A + \bar{\lambda} B)^+ \leq \lambda A^+ + \bar{\lambda} B^+$$

the inequality holding for a single $\lambda \in]0, 1[$ is enough to guarantee its validity for all $\lambda \in]0, 1[$. In [2] this was shown to be the case, the proof relying on a general result on "t-convexity" based on Rode's abstract version of the classical (Helly–)Hahn–Banach Theorem. Time permitting I will present alternative simpler proofs which avoid the use of this general result.

Keywords

Jensen convexity, Midpoint convexity, Rational convexity.

References

- [1] Nordström, K. (2011). Convexity of the inverse and Moore–Penrose inverse. *Linear Algebra Appl.* 434, 1489–1512.
- [2] Nordström, K. (2018). A note on the convexity of the Moore–Penrose inverse. *Linear Algebra Appl.* 538, 143–148.

The Safety Belt estimation method

Katarzyna Filipiak¹ and Dietrich von Rosen²

¹ Poznań University of Technology, Poland

² Linköping University, Sweden

Abstract

The main goal of this talk is to determine maximum likelihood estimators in the classical multivariate linear model as well as the growth curve model, assuming known prior information introduced via inequality restrictions on the parameters of the model. The restrictions are in the form of quadratic inequalities. Methods from convex optimization theory play a fundamental role in determining the estimators. A characteristic of the new estimators, called Safety Belt estimators, is that depending on the observed data, there are two alternative solutions to the likelihood equations. The result are illustrated using real-world data examples. Some similarities of Safety Belt estimation method with Ridge penalization of estimators under univariate models are also presented.

Keywords

Safety Belt estimation, Maximum likelihood estimation, Inequality constraints, Ridge regression, MANOVA, Growth curve model.

References

- [1] Filipiak, K., von Rosen, D., Singull, M. (2026). Maximum likelihood estimation under inequality constraints in univariate linear models. In: Wood, D.R., Etheridge, A.M., de Gier, J., Joshi, N. (Eds.), *2024 MATRIX Annals* (pp. 149–154). Springer.
- [2] Filipiak, K., von Rosen, D., Rejchel, W., Singull, M. (2026). The Safety Belt estimator under multivariate linear models with inequality constraints. *J. Statist. Plann. Inference* *241*, 106335.
- [3] Filipiak, K., von Rosen, D., and Singull, M. (2026). The Safety Belt estimator applied to the growth curve model with inequality constraints. *Submitted*.

Independence testing based on the covariance matrix

Adam Mieldzioc¹, Malwina Mrowińska¹,
and Anna Szczepańska-Álvarez²

¹ Poznań University of Technology, Poland

² Poznań University of Life Sciences, Poland

Abstract

The problem of testing the independence between groups of variables has a long history in multivariate statistics. One of the earliest contributions was made by Wilks (1935) ([3]), who studied the problem of testing the independence of two groups of variables. Since then, the theory has been extended, particularly in the context of high-dimensional data. We consider two groups of jointly normally distributed random variables. By combining both groups into a single random vector, the corresponding covariance matrix has a block structure, where the diagonal blocks describe the covariance structure within each group and the off-diagonal blocks represent the dependence between the groups.

We formulate likelihood ratio tests for assessing the separability of one block of the covariance matrix and, subsequently, the independence between two groups of variables under the assumption of a separable covariance structure. The finite sample properties of the proposed tests are evaluated through an extensive simulation study.

Keywords

Block covariance structure, Independence testing, Separable covariance matrix, Likelihood ratio test.

References

- [1] Wilks, S. S. (1935). Test criteria for statistical hypotheses involving several variables. *J. Amer. Statist. Assoc.*, 30, 549–560.

**Special Session:
Memorial Session for Jerzy K. Baksalary
(1944-2005)**

Simo Puntanen

Tampere University, Finland

Abstract

Jerzy K. Baksalary worked intensively on the theory of linear models, in particular, aspects of linear algebra. In this special session his colleagues and collaborators comment on his life and work.

Invited Speakers:

Oskar Maria Baksalary (Poland)
Jan Hauke (Poland)
Augustyn Markiewicz (Poland)
Thomas Mathew (USA)
Simo Puntanen (Finland)

On the Baksalary–Hauke condition aka the Rao–Mitra–Bhimasankaram relation

Oskar Maria Baksalary

Adam Mickiewicz University, Poznań, Poland

Abstract

I honor the great teachers and they live in my work and they dance invisibly in the margins of my prose.

Donald Patrick Conroy to Barbra Streisand, quoted from [5].

Jerzy K. Baksalary was one of Oskar Maria Baksalary’s great teachers, likely the one. Oskar trusts that Jerzy lives in all his works (though he doubts Jerzy tries any dancing). A striking example of Oskar’s work on which Jerzy had a tremendous impact is provided by paper [2], resulting from Oskar’s collaboration with Dennis S. Bernstein. The problem considered therein is rooted in two prior articles, namely [1] written by Jerzy jointly with Jan Hauke and [4] co-authored by C. Radhakrishna Rao, Sujit Kumar Mitra, and P. Bhimasankaram. Both these papers pay attention to the identity of the form $KK^*K = KL^*K$, where K and L are complex matrices of the same size. The identity was in [3] called “the Baksalary–Hauke condition”, whereas in [2] it was referred to as “the Rao–Mitra–Bhimasankaram relation”. The talk, based on [2], will shed light on the properties of the identity, which deal inter alia with the notions of intransitivity, a well-known feature associated with several research areas, and strong antisymmetry, whose meaning was coined only in [2]. Various personal recollections concerned with Jerzy will be provided in the talk as well.

Keywords

Jerzy K. Baksalary (1944-2005), Intransitivity, Strong antisymmetry, Matrix partial order.

References

- [1] Baksalary, J. K., Hauke, J. (1990). A further algebraic version of Cochran’s theorem and matrix partial orderings. *Linear Algebra Appl.*, 127, 157–169.
- [2] Baksalary, O. M., Bernstein, D. (2025). The Rao–Mitra–Bhimasankaram relation is strongly antisymmetric. *Linear Algebra Appl.*, 710, 80–94.

- [3] Baksalary, O. M., Trenkler, G. (2021). A partial ordering approach to characterize properties of a pair of orthogonal projectors. *Indian J. Pure Appl. Math.*, 52, 323–334.
- [4] Rao, C. R., Mitra, S. K., Bhimasankaram, P. (1972). Determination of a matrix by its subclasses of generalized inverses. *Sankhyā Ser. A* 34, 5–8.
- [5] Streisand, B. (2023). *My name is Barbra*. Viking Press.

Remembering Jerzy K. Baksalary: My Mentor in Mathematics and Beyond

Jan Hauke

Adam Mickiewicz University, Poznań, Poland

Abstract

The vast majority of researchers, regardless of their field of study, hold their teachers in the highest esteem, because they remain present in their scholarly work long after their joint research has ended. The legacy of great mentors continue to live on through the theorems, proofs, and research questions pursued by their students. My presentation has been prepared in this spirit. My own scientific achievements are modest compared with those of my mentor, Professor Jerzy K. Baksalary. However, whatever I have managed to accomplish is a direct result of knowledge, inspiration, and example that I received from him. This influence extends beyond my research work, to my teaching activities and many other spheres of my life.

In my presentation, I will trace the history of my long-standing collaboration with Jerzy K. Baksalary from 1979 to 2005, highlighting selected aspects that were particularly significant. This period includes interruptions caused by his involvement in the first Solidarity movement and its consequences, as well as the years when he served as Rector of the Zielona Gora Pedagogical University. I will also describe the background and circumstances that led to the publication of two of our joint papers, [1] and [2].

Keywords

Jerzy K. Baksalary, Matrix equation, Matrix partial ordering

References

- [1] Baksalary, J.K., Hauke, J., Kala, R. (1980). Nonnegative definite solutions to some matrix equations occurring in distribution theory of quadratic forms. *Sankhya, Ser. A*, 42, 283–291.
- [2] Baksalary, J.K., Hauke, J., Liu, X., Liu, S. (2004). Relationships between partial orders of matrices and their powers. *Linear Algebra Appl.*, 379, 277– 287.

Linear admissibility and sufficiency

Augustyn Markiewicz

Poznań University of Life Sciences, Poland

Abstract

The problem of admissible linear estimation under the Gauss-Markov model was extensively studied by Rao in [7]. Rao's results were developed in [3] and [2] under a possibly singular model. However, the linear sufficiency was studied and characterized in [1] and [4]. The subject of linear sufficiency and admissibility was considered in [5] and [6].

The purpose of this paper is to present J. K. Baksalary's contribution to the theory, along with further complementary results.

Keywords

Admissibility, Linear sufficiency.

References

- [1] Baksalary, J. K., Kala, R. (1981). Linear transformations preserving best linear unbiased estimators in a general Gauss-Markoff model. *Ann. Statist.*, 9, 913–916.
- [2] Baksalary, J. K., Markiewicz, A., Rao, C. R. (1995). Admissible linear estimation in the general Gauss-Markov model with respect to an arbitrary quadratic risk function. *J. Statist. Plann. Inference*, 44, 341–347.
- [3] Baksalary, J. K., Rao, C. R., Markiewicz, A. (1992). A study of the “natural restrictions” on the estimation problems in the singular Gauss-Markov model. *J. Statist. Plann. Inference*, 31, 335–351.
- [4] Drygas, H. (1983). Sufficiency and completeness in the general Gauss-Markov model. *Sankhya, Ser. A*, 45, 88–98.
- [5] Markiewicz, A. (1996). Characterization of general ridge estimators. *Stat. Probab. Lett.*, 27, 145–148.
- [6] Markiewicz, A., Puntanen, S. (2009). Admissibility and linear sufficiency in linear model with nuisance parameters. *Stat. Pap.*, 50, 847–854.
- [7] Rao, C. R. (1976). Estimation of parameters in a linear model. *Ann. Statist.*, 4, 1023–1037.

Jerzy Baksalary's research: A personal note

Thomas Mathew

University of Maryland, Baltimore County, USA

Abstract

I will reminisce on some of the early work of Jerzy Baksalary that has influenced me and made an impact on my own research. I will also comment on some of his results, though simple, that I have found rather striking.

Photographic reminiscences of Jerzy K. Baksalary (1944–2005)

Simo Puntanen

Tampere University, Finland

Abstract

There are moments, Jeeves, when one asks oneself,

"Do matrices matter?"

Jeeves: "The mood will pass, sir."

Transformed conversation created by P.G. Wodehouse.

On 18 December 1981, 45 years ago, I wrote, with my colleague Erkki Liski, a letter to Jerzy K. Baksalary and Radosław Kala; below is an excerpt.

"Thank you for very interesting research papers you have published in different journals. We have read them with a great interest. Our own interest lies to a great extent in the same topics and this is why we would like to contact you. We enclose some technical reports and a doctoral dissertation and we would be very glad for any comments. Also we would be very happy if you could send us your papers. By the way, we think that we have at least one common friend: Professor George P.H. Styan, who positively referred to your articles. He visited our department first in January 1976 and second time in July this year. We felt him a very inspiring person."

On 3 June 1984, Erkki Liski and I stepped out of the train at the Poznań Railway Station, and were friendly welcomed by Barbara Bogacka, who kindly guided us to the conference site: *International Statistical Conference in Linear Inference*, held in Poznań, 4–9 June 1984. That was the first time I had an opportunity to see live performances of the Polish Linear Model Stars.

... Oh my holy goodness ... all these memories, they really push me thinking of the Polish-Nordic Party at Jan Hauke's Penthouse ... Thursday, 7 June 1984. The attached two pictures of Jerzy were taken there.



On 14 June 1984, I sent a thank-you letter to Jerzy. In a nutshell it said: *All in all—what a superb Party!*

Last time I met Jerzy in 2004, in Będlewo, near Poznań, at the IWMS meeting. In the photograph below Jerzy is posing for the group photo with Ingram Olkin, T.W. Anderson, Dorothy Anderson, Imbi Traat, Tõnu Kollo, Dietrich von Rosen, Tatjana von Rosen, etc.

In my talk I'll go through some further JKB-memories from 1984–2004.



Keywords

Jerzy K. Baksalary (1944–2005), Agricultural University of Poznań, Tampere University, LinStat conference series, IWMS conference series, Linear statistical models, Linear algebra and its statistical applications.

References

- [1] Atkinson, A. C., Bogacka, B. (2015). A conversation with Professor Tadeusz Caliński. *Statistical Science*, 30, 432–442.
- [2] Baksalary, O. M., Styan, G. P. H., eds. (2005). Some comments on the life and publications of Jerzy K. Baksalary (1944–2005). *Linear Algebra Appl.*, 410, 3–53.
- [3] Hauke, J., Markiewicz, A. (2013). Comments on the development of Linear Algebra in Poznań, Poland. *Image*, 50, 16–17.
- [4] Hauke, J., Markiewicz, A., Puntanen, S. (2020). Beyond the MatTriad conferences. *App. Math.*, 65, 547–556.
- [5] Isotalo, J., Puntanen, S., Styan, G. P. H. (2008). Decomposing matrices with Jerzy K. Baksalary. *Discuss. Math. Probab. Stat.*, 28, 91–111.
- [6] Szulc, T. (2022). Jerzy K. Baksalary (1944–2005). In: Krzyśko, M. (Ed.), *Polish Statisticians. Bibliographical Notes* (pp. 7–13). Główny Urząd Statystyczny, Warszawa.

Part V

List of Participants

Participants

1. **S. Ejaz Ahmed**
Department of Mathematics and Statistics
Brock University, St. Catharines, Canada
sahmed5@brocku.ca
2. **Marek Arendarczyk**
Mathematical Institute
Univeristy of Wrocław, Wrocław, Poland
marek.arendarczyk@math.uni.wroc.pl
3. **Oskar Maria Baksalary**
Institute of Spintronics and Quantum Information
Faculty of Physics and Astronomy,
Adam Mickiewicz University, Poznań, Poland
Oskar.Baksalary@amu.edu.pl
4. **Michał Balcerek**
Department of Applied Mathematics
Faculty of Pure and Applied Mathematics
Wrocław University of Science and Technology, Wrocław, Poland
michal.balcerek@pwr.edu.pl
5. **Stephan Bark**
Institute of Crop Science
University of Hohenheim, Stuttgart, Germany
stephan-bark@web.de
6. **Anastassia Baxevani**
Department of Mathematics and Statistics
University of Cyprus, Nicosia, Cyprus
baxevani@ucy.ac.cy
7. **Andriette Bekker**
Department of Statistics
Faculty of Natural and Agricultural Sciences
University of Pretoria, Pretoria, South Africa
andriette.bekker@up.ac.za
8. **Alessandro Berti**
Department of Economics, Society, and Politics
Urbino University “Carlo Bo”, Urbino, Italy
alessandro.berti@uniurb.it
9. **Małgorzata Bogdan**
Institute of Mathematics
University of Wrocław, Wrocław, Poland
Malgorzata.Bogdan@uw.edu.pl

10. **Stefano Bonnini**
Department of Economics and Management
University of Ferrara, Ferrara, Italy
stefano.bonnini@unife.it
11. **Claudio Giovanni Borroni**
Department of Statistics and Quantitative Methods
University of Milano-Bicocca, Milan, Italy
claudio.borroni@unimib.it
12. **Manuela Cazzaro**
Department of Statistics and Quantitative Methods
University of Milano-Bicocca, Milan, Italy
manuela.cazzaro@unimib.it
13. **Lucio De Capitani**
Department of Statistics and Quantitative Methods
University of Milano-Bicocca, Milan, Italy
lucio.decapitani1@unimib.it
14. **Elvira Di Nardo**
Department of Mathematics
University of Turin, Turin, Italy
elvira.dinardo@unito.it
15. **Katarzyna Filipiak**
Institute of Mathematics
Poznan University of Technology, Poznań, Poland
katarzyna.filipiak@put.poznan.pl
16. **Eva Fišerová**
Department of Mathematical Analysis and Applications of Mathematics
Palacký University Olomouc, Olomouc, Czech Republic
eva.fiserova@upol.cz
17. **Johannes Forkman**
Department of Crop Production Ecology
Swedish University of Agricultural Sciences, Uppsala, Sweden
johannes.forkman@slu.se
18. **Cinzia Franceschini**
Department of Economics and Business
University of Sassari, Sassari, Italy
cfranceschini@uniss.it
19. **Agnieszka Goroncy**
Department of Mathematical Statistics and Data Mining
Faculty of Mathematics and Computer Science
Nicolaus Copernicus University, Toruń, Poland
gemini@mat.umk.pl

20. **Tomasz Górecki**
Faculty of Mathematics and Computer Science
Adam Mickiewicz University, Poznań, Poland
tomasz.gorecki@amu.edu.pl
21. **Małgorzata Graczyk**
Department of Mathematical and Statistical Methods
Poznań University of Life Sciences, Poznań, Poland
malgorzata.graczyk@up.poznan.pl
22. **Jan Hauke**
Institute of Socio-Economic Geography and Spatial Management
Adam Mickiewicz University, Poznań, Poland
jhauke@amu.edu.pl
23. **Julianna Holmberg**
Department of Mathematics
Linköping University, Linköping, Sweden
julianna.lise.holmberg@liu.se
24. **Joanna Janczura**
Faculty of Pure and Applied Mathematics
Wrocław University of Science and Technology, Wrocław, Poland
joanna.janczura@pwr.edu.pl
25. **Farrukh Javed**
Statistics School of Economics and Management
Lund University, Lund, Sweden
Farrukh.Javed@stat.lu.se
26. **Mateusz John**
Institute of Mathematics
Poznan University of Technology, Poznań, Poland
mateusz.john@put.poznan.pl
27. **Vivien Kardos**
University of Szeged, Szeged, Hungary
kardos.vivien.kata@szte.hu
28. **Barbora Kessel**
School of Public Health and Community Medicine
Sahlgrenska Academy
University of Gothenburg, Gothenburg, Sweden
barbora.kessel@gu.se
29. **Daniel Klein**
Institute of Mathematics
P. J. Šafárik University, Košice, Slovakia
daniel.klein@upjs.sk

30. **Patryk Kołbyko**
Doctoral School of Social Sciences, Economics and Finance
Maria Curie-Skłodowska University, Lublin, Poland
patryk.kolbyko2000@gmail.com
31. **Kamil Kołodziejski**
Institute of Mathematics
Lodz University of Technology, Łódź, Poland
kamil.kolodziejski@dokt.p.lodz.pl
32. **Peter Kovacs**
Department of Statistics and Demography
Faculty of Economics and Business Administration
University of Szeged, Szeged, Hungary
pepe@eco.u-szeged.hu
33. **Agnieszka Kwita**
Wrocław University of Science and Technology, Wrocław, Poland
294764@student.pwr.edu.pl
34. **Lynn Roy LaMotte**
Biostatistics Program
School of Public Health
Louisiana State University, New Orleans, Louisiana, USA
lrlamo@icloud.com
35. **Alicja Lerczak**
Department of Mathematical and Statistical Methods
Poznań University of Life Sciences, Poznań, Poland
alicja.lerczak@up.poznan.pl
36. **Yuli Liang**
School of Mathematics and Statistics
Guangxi Normal University, Guangxi, China
yuli.liang@gxnu.edu.cn
37. **Eckhard Liebscher**
Faculty of Engineering and Natural Sciences
University of Applied Sciences, Merseburg, Germany
eckhard.liebscher@hs-merseburg.de
38. **Nicola Loperfido**
Department of Economics, Society, and Politics
Urbino University “Carlo Bo”, Urbino, Italy
nicola.loperfido@uniurb.it
39. **Nadja Malevich**
University of Münster, Münster, Germany
nadja.malevich@uni-muenster.de

40. **Augustyn Markiewicz**
 Department of Mathematical and Statistical Methods
 Poznań University of Life Sciences, Poznań, Poland
augustyn.markiewicz@up.poznan.pl
41. **Thomas Mathew**
 Department of Mathematics and Statistics
 University of Maryland, Baltimore County, Maryland, USA
mathew@umbc.edu
42. **Adam Mieldzioc**
 Institute of Mathematics
 Poznan University of Technology, Poznań, Poland
adam.mieldzioc@put.poznan.pl
43. **Monika Mokrzycka**
 Institute of Plant Genetics
 Polish Academy of Sciences, Poznań, Poland
mmok@igr.poznan.pl
44. **Marco Morosin**
 Department of Mathematics and Computer Science
 Eindhoven University of Technology, Eindhoven, The Netherlands
m.morosin@tue.nl
45. **Malwina Mrowińska**
 Institute of Mathematics
 Poznan University of Technology, Poznań, Poland
malwina.mrowinska@put.poznan.pl
46. **Lev Muchnik**
 Hebrew University of Jerusalem, Jerusalem, Israel
levmuchnik@gmail.com
47. **Werner Müller**
 Department of Applied Statistics
 Johannes Kepler University Linz, Linz, Austria
werner.mueller@jku.at
48. **Boaz Nadler**
 Department of Computer Science and Applied Mathematics
 Weizmann Institute of Science, Rehovot, Israel
boaz.nadler@weizmann.ac.il
49. **Im Chiew Ng**
 Medical Statistics Department
 London School of Hygiene & Tropical Medicine, London, UK
lshin11@lshtm.ac.uk
50. **Kenneth Nordström**
 Department of Mathematics and Statistics
 University of Helsinki, Helsinki, Finland
kenneth.nordstrom@helsinki.fi

51. **Ostap Okhrin**
Chair of Applied Statistics
Dresden University of Technology, Dresden, Germany
ostap.okhrin@tu-dresden.de
52. **Yarema Okhrin**
Department of Statistics and Data Science
Faculty of Business and Economics
University of Augsburg, Augsburg, Germany
Linnaeus University, Växjö, Sweden
yarema.okhrin@uni-a.de
53. **Paulo Eduardo Oliveira** Department of Mathematics
Univeristy of Coimbra, Coimbra, Portugal
paulo@mat.uc.pt
54. **Jagoda Papis** Institute of Mathematics
Poznan University of Technology, Poznań, Poland
jagoda.papis@put.poznan.pl
55. **Jaakko Pere**
Department of Information and Service Management
Aalto University, Espoo, Finland
jaakko.pere@aalto.fi
56. **Jolanta Pielaszkiewicz**
Department of Computer and Information Science
Linköping University, Linköping, Sweden
jolanta.pielaszkiewicz@liu.se
57. **Marcin Pitera**
Institute of Mathematics
Jagiellonian University, Cracow, Poland
marcin.pitera@uj.edu.pl
58. **Krzysztof Podgórski**
Department of Statistics
Lund University, Lund, Sweden
Krzysztof.Podgorski@stat.lu.se
59. **Maryna Prus**
Biostatistics Unit
University of Hohenheim, Stuttgart, Germany
maryna.prus@uni-hohenheim.de
60. **Simo Puntanen**
Tampere University, Tampere, Finland
simo.puntanen@tuni.fi
61. **Wojciech Rejchel**
Nicolaus Copernicus University, Toruń, Poland
wrejchel@gmail.com

- 62. **Tomasz Rychlik**
Institute of Mathematics
Polish Academy of Sciences, Warsaw, Poland
trychlik@impan.pl
- 63. **Perttu Saarela**
University of Helsinki, Helsinki, Finland
perttu.saarela@helsinki.fi
- 64. **Tomer Shushi**
Ben-Gurion University of the Negev, Beersheba, Israel
tomershu@bgu.ac.il
- 65. **Katarzyna Skowronek**
Wrocław University of Science and Technology, Wrocław, Poland
katarzyna.skowronek@pwr.edu.pl
- 66. **Łukasz Smaga**
Faculty of Mathematics and Computer Science
Adam Mickiewicz University, Poznań, Poland
ls@amu.edu.pl
- 67. **Piotr Sulewski**
Institute of Exact and Technical Sciences
Pomeranian University, Słupsk, Poland
piotr.sulewski@upsl.edu.pl
- 68. **Anna Szczepańska-Álvarez**
Department of Mathematical and Statistical Methods
Poznań University of Life Sciences, Poznań, Poland
anna.szczepanska-alvarez@up.poznan.pl
- 69. **Magdalena Szymkowiak**
Institute of Automation and Robotics
Poznan University of Technology, Poznań, Poland
magdalena.szymkowiak@put.poznan.pl
- 70. **Svend Vendelbo Nielsen**
Danish Technological Institute, Taastrup, Denmark
sven@teknologisk.dk
- 71. **Dietrich von Rosen**
Linköping University, Linköping, Sweden
dietrich.von.rosen@liu.se
- 72. **Viktor Witkovský**
Institute of Measurement Science
Slovak Academy of Sciences, Bratislava, Slovakia
witkovsky@savba.sk

- 73. **Agnieszka Wylomańska**
Department of Applied Mathematics
Faculty of Pure and Applied Mathematics
Wrocław University of Science and Technology, Wrocław, Poland
agnieszka.wylomanska@pwr.edu.pl
- 74. **Jianfeng Yao**
School of Data Science
The Chinese University of Hong Kong, Shenzhen, China
jeffyyao@cuhk.edu.cn
- 75. **Ivan Žežula**
Institute of Mathematics
P. J. Šafárik University, Košice, Slovakia
ivan.zezula@upjs.sk
- 76. **Konstantinos Zografos**
Department of Mathematics
University of Ioannina, Ioannina, Greece
kzograf@uoi.gr
- 77. **Wojciech Żuławiński**
Wrocław University of Science and Technology, Wrocław, Poland
wojciech.zulawinski@pwr.edu.pl

Index

- Ahmed, S.E., *31*
Alenlöv, J., *81*
Ali Taha, A., *80*
Arciniegas-Alarcón, S., *83*
Arendarczyk, M., *63, 70*
- Bagnato, L., *54*
Baksalary, O.M., *127*
Balakrishnan, N., *89, 95*
Balcerek, M., *65*
Bark, S., *113*
Baumeister, M., *86*
Baxevari, A., *75*
Bekker, A., *49*
Berti, A., *51*
Bogdan, M., *32*
Bonnini, S., *76*
Borghesi, M., *76*
Borroni, C.G., *52, 53*
Boutin, M., *57*
- Carpenter, J., *108*
Cazzaro, M., *53*
Ceranka, B., *115*
Chiodini, P.M., *53*
Coupkovva, E., *57*
Czerwińska-Kayzer, D., *106*
- De Capitani, L., *52, 54*
De Iaco, S., *59*
Di Nardo, E., *33*
Drapella, A., *87*
Duarte, D., *83*
- Fantoly, Z., *103*
Fermanian, J.-D., *44*
Fišerová, E., *98*
Filipiak, K., *119, 124*
Forkman, J., *83*
Franceschini, C., *55, 56*
- Górecki, T., *82*
Gaj, R., *106*
Gal, A., *103*
García-Peña, M., *83*
- Goroncy, A., *89*
Graczyk, M., *115*
Großmann, H., *116*
- Hübbert, L., *99*
Haid, P., *45*
Harman, R., *35*
Hauke, J., *129*
Holling, H., *116*
Holmberg, J., *99*
- Janczura, J., *66*
Javed, F., *77*
John, M., *79, 104, 120*
- Kardos, V., *100*
Karsai, K., *103*
Kateri, M., *95*
Kayzer, D., *106*
Kelemen, B., *103*
Kessel, B., *97*
Klein, D., *78–80*
Koľbyko, P., *101*
Kołodziejski, K., *68*
Kovacs, P., *100, 103*
Kozubowski, T.J., *63*
Krapf, D., *65*
Kwita, A., *104*
- LaMotte, L.R., *105*
Lando, T., *91*
Lerczak, A., *106*
Liang, Y., *79*
Liebscher, E., *42*
Lindner, A., *68*
Loperfido, N., *48, 51, 55, 56*
- Müller, W.G., *35*
Malevich, N., *116*
Markiewicz, A., *130*
Mathew, T., *37, 131*
Mieldzioc, A., *120, 125*
Mielniczuk, J., *85*
Mokrzycka, M., *81, 122*
Morgan, K., *108*

- Morosin, M., 57
Mrowińska, M., 80, 122, 125

Nadler, B., 58
Ng, I.C., 108
Nicholas, J., 108
Nordhausen, K., 59, 60
Nordström, K., 123

Okhrin, O., 41, 44
Okhrin, Y., 45
Oliveira, P.E., 91

Pacheco-Pozo, A., 65
Panorska, A.K., 63
Papis, J., 93
Pauly, M., 86
Pere, J., 59, 60
Pielaszkiewicz, J., 81
Piepho, H.-P., 113
Pillay, J., 49
Pitera, M., 69
Podgórski, K., 74, 77
Prus, M., 112, 113, 117
Puć, A., 66
Puntanen, S., 126, 132
Punzo, A., 49, 54

Rejchel, W., 85
von Rosen, D., 99, 124
Ruiz, A.M., 60

Rychlik, T., 94

Saarela, P., 59, 60
Schnörr, C., 45
Schwabe, R., 116
Shushi, T., 61
Singull, M., 99
Skowronek, K., 63, 70
Smaga, Ł., 86
Stadler, A., 35
Sulewski, P., 87
Szczepańska-Álvarez, A., 125
Szymkowiak, M., 88, 93–95

Teisseyre, P., 85
Tortora, C., 49

Vaknin Laufer, E., 58

Wielgusz, K., 106
Witkovský, V., 110
Wustl, J., 45
Wyłomańska, A., 62, 63, 65, 70, 72

Yao, J., 38
Yi, H., 89

Žežula, I., 78
Zilber, P., 58
Zografos, K., 46
Żuławiński, W., 72