



ENBIS-26 ANNUAL CONFERENCE PROCEEDINGS



UNIVERSITÀ
DEGLI STUDI
FIRENZE

Dipartimento di Statistica,
Informatica, Applicazioni
“Giuseppe Parenti”

Eccellenza 2023-2027



ENBIS-26 gratefully acknowledges support of the following sponsors:



UNIVERSITÀ
DEGLI STUDI
FIRENZE

Dipartimento di Statistica,
Informatica, Applicazioni
"Giuseppe Parenti"

Eccellenza 2023-2027



ReDS
Rethinking
Data Science

Conference gold sponsor:

Minitab ®

General sponsor and conference gold sponsor:


Jmp®
**STATISTICAL
DISCOVERY**

ENBIS-26 ANNUAL CONFERENCE PROCEEDINGS

Proceedings contain all the abstracts accepted for the ENBIS-26 Annual Conference

Firenze, Italy

September 6-10, 2026

©ENBIS (www.enbis.org)

©2026 Dipartimento di Statistica, Informatica, Applicazioni “G. Parenti” – Università degli Studi di Firenze

All rights reserved

Published by: Dipartimento di Statistica, Informatica, Applicazioni “G. Parenti” – Università degli Studi di Firenze

Cover photo by:

Rossella Berni

EUROPEAN Network for Business and Industrial Statistics. Annual conference (26; 2026; Firenze)
ENBIS-26 Conference Proceedings: programme and abstracts of the 26th Annual Conference of the European Network for Business and Industrial Statistics (ENBIS), Firenze, Italy, September 6–10, 2026 / [photography by Rossella Berni]. – 1st ed. – Firenze: Dipartimento di Statistica, Informatica, Applicazioni “G. Parenti” – Università degli Studi di Firenze, 2026

ISBN 979-12-243434-1-7

Contents

Welcome to the 26th Annual Conference of ENBIS, the European Network for Business and Industrial Statistics	1
The Scientific Program and Organising Committees of ENBIS-26	3
ENBIS Executive Committee 2025-2027	3
ENBIS Permanent Office	3
Keynote Speakers	4
Detailed program	5
Abstracts	24
Useful Information	123
ENBIS-27 Conference Announcement	127

Welcome to The 26th annual conference of ENBIS, the European Network for Business and Industrial Statistics!

Dear participants of ENBIS-26 in Florence

On behalf of the Scientific Program Committee of ENBIS-26, it is my great honor and pleasure to welcome you to the 26th Annual Conference of the European Network for Business and Industrial Statistics host by the University of Florence at the Centro Didattico Morgagni in Florence (Italy) from September 6 to 10, 2026.

We are proud to present you with a program reflecting a strong and dynamic network in applied statistics, both from academia, industry & business, and demonstrating numerous and fruitful interfaces with other disciplines. With many talks and interactive sessions, covering a broad range of topics in statistics, we hope you will find it stimulating and relevant for your own interests.

We are also very proud of our plenary sessions. The opening keynote speaker is Laura Maria Sangalli, Politecnico di Milano (Italy), and the closing keynote speaker is Peter Bühlmann, ETH Zürich (Switzerland). Bradley Jones (EFFEX) is the 2026 George Box Medal recipient, and the Greenfield Challenge winner is Lluís Marco Almagro (Universitat Politècnica de Catalunya & BarcelonaTech). Their talks will take place on Tuesday afternoon. Two more speakers, whose presentations are scheduled in the Award Session on Monday afternoon, will receive ENBIS awards in 2026: Sophie Carr (Bays Consulting) for the Best Manager Award and : Fabio Centofanti (KU Leuven) for the Young Statistician Award.

Besides these keynote sessions, there will be 11 sessions' slots, each involving 4 or 5 sessions held in parallel. Among these, around 30 contributed sessions (more than 80 talks) will address the following topics:

- AI: Machine Learning and Predictive Analytics
- Data Analytics and Data Science: Case Studies
- Design of Experiments
- Measurement Uncertainty and Metrology
- Reliability
- Statistical / Stochastic Modelling and Statistical Computing
- Statistical Process Monitoring
- Statistics in Industry, Business and Finance
- Statistics in Pharma / Healthcare
- Teaching, Consulting and Knowledge Transfer in Statistics
- Time Series, Forecasting and Dynamic Systems
- Trustworthy and Explainable AI
- Uncertainty quantification and computer experiments

In addition, the conference hosts 19 invited sessions featuring more than 50 talks. Some are institutional and/or traditional ones, like the Software, Editor's Corner, Mathmet, INFORMS/QSR, itENBIS, frENBIS, ISEA, Young Statistician and JQT/QE/Technometrics. There are also new sessions reflecting the subject trends, for example the links between Statistics & Data Science/Artificial Intelligence, the topics of Uncertainty Quantification, Sensitivity Analysis &

Computer Experiments, and some scientific or technological fields as Climate Science and Digital Twins. A strong emphasis is finally made on Active/Interactive sessions, Generative AI tools and discussions on Teaching.

As usual, pre- and post-conference events of the ENBIS-26 conference will take place:

- On Sunday (6th of September) afternoon, the joint ECAS-ENBIS course on Adaptive Machine Learning for Time Series Forecasting,
- On Wednesday (10th of September) morning, the JMP Workshop: Human + AI in the Loop: AI-Supported Analysis and Learning in JMP,
- On Wednesday (10th of September) afternoon, Consultancy Skills: Storytelling with Data.

This year, as an exceptional additional event to celebrate the ENBIS hosting in Florence, a Satellite Event organized by the StEering Research Center will be held on Saturday 5th September, bringing high-level talks on Railway (Reliability, Safety and Sustainability), Reliability (Design of Experiments, Maintenance, Degradation Models), Tourism (Predictions and Sustainability) and Statistics for biomedical engineering.

Finally, a special issue of the Quality and Reliability Engineering International (QREI) Journal will feature (expected in 2027) selected papers presented in ENBIS-26. It will be edited by Véronique Maume-Deschamps (Université Claude Bernard Lyon 1, France), and Panagiotis Tsiamyrtzis (Athens University of Economics and Business, Athens, Greece). A formal Call for Papers will be issued after the conference.

I want to express my deepest thanks to the dedicated Program and Organizing Committee members, in particular Rossella Berni who has agreed to organize this conference in this wonderful city of Florence. Organizing this conference on an annual basis is essential to maintaining the vitality of the ENBIS community. I would like to take this opportunity to call for volunteers to embark on this adventure which, despite the administrative challenges, proves to be a source of valuable professional and personal insights for the organizers.

I am also grateful to the ENBIS Executive and Office members, who have helped get this conference rolling. Lara Kuhlmann de Canaviri is a highly efficient person who activates all the keys to solving many issues concerning the organization and the administration of the conference website. Finally, I am grateful to all colleagues and friends who have organized the invited sessions and/or have agreed to chair the different contributed sessions.

Last, but not least, I would like to express my greatest gratitude to all of you who will contribute to the success of ENBIS-26 with your talks and discussions between sessions and during the different social events. As you know, networking is the primary purpose of the ENBIS annual conference, and we must draw particular attention to involve in our discussions the newcomers.

Please enjoy the ENBIS-26 Conference and have a good time in Florence!

Bertrand Iooss

Chair of the Scientific Program Committee for ENBIS-26 Florence Conference

The Scientific Program and Organising Committees of ENBIS-26

ENBIS-26 Scientific Program Committee

Bertrand Iooss, Chair
 Rossella Berni
 Katy Klauenberg
 Sven Knoth
 Véronique Maume-Deschamps
 Alessandra Menafoglio
 Werner Müller
 Morten Bormann Nielsen
 Biagio Palumbo
 Peter Parker
 Winfried Theis
 Marina Vives-Mestres
 Margaux Zaffran

ENBIS-26 Organizing Committee

Rossella Berni, Chair
 Lucia Buzzigoli
 Francesca A. Giambona
 Alessandro Magrini
 Nedka D. Nikiforova
 Antonio Pievatolo
 Antonella Bonacorsi
 Renza Campagni
 Maria Nunzia Galdi
 Stefano Mariani
 Lorenzo Tofani

ENBIS Executive Committee 2025-2027

President	Jacqueline Asscher
President-Elect	Winfried Theis
Vice President	Bertrand Iooss
Vice President	Frøydis Bjerke
Vice President	Marina Vives-Mestres
Past President	Biagio Palumbo

ENBIS Permanent Office

Director	Sonja Kuhnt
Treasurer	Bernard Francq
Communication Officer	Christian Capezza
Webmaster	Morten Bormann Nielsen
Administrative Officer	Lara Ariane Saskia Kuhlmann de Canaviri

Address of the ENBIS Permanent Office: European Network for Business and Industry Statistics (ENBIS), Westzanerdijk 170, 1507AM Zaandam, The Netherlands.

Keynote Speakers

The conference will feature five keynote lectures delivered by leading experts, combining invited international speakers and distinguished ENBIS award recipients.

Invited Keynote Speakers

- Laura Maria Sangalli (Politecnico di Milano, Italy) – Laura Maria Sangalli's work focuses on spatial statistics and functional data analysis, with impactful applications in environmental and biomedical sciences.
Her lecture introduces physics-informed statistical learning methods that combine data with known physical principles, such as differential equations, to improve inference and decision-making in complex, high-dimensional, and structured applications. By integrating statistical modeling with physically grounded regularization, these approaches enhance interpretability and predictive performance across fields including environmental science, engineering, pharmacokinetics, and medical research.
- Peter Bühlmann (ETH Zürich, Switzerland) – Peter Bühlmann is internationally recognised for his contributions to high-dimensional statistics, machine learning, and causal inference, with wide-ranging applications across science and industry.
His lecture examines the challenge of applying statistical and machine learning models to new populations, particularly in digital health settings such as international ICU datasets. It presents Distributionally Robust Invariance Learning as a method for identifying stable patterns across environments and discusses the opportunities and limitations of emerging foundation models for improving generalization and adaptation.

ENBIS Award Keynote Speakers

- Bradley Jones – 2026 Box Medal – Bradley Jones is recognized for his outstanding contributions to the development and application of statistical methods in business and industry, particularly in design of experiments and industrial quality improvement.
His talk presents recent advances in Definitive Screening Designs (DSDs), extending their use beyond continuous and two-level categorical factors to include categorical factors with three or more levels. It introduces a general framework for constructing these enhanced designs and illustrates their application through practical examples.
- Sophie Carr (Bays Consulting, England) – 2026 Best Manager Award – Sophie Carr's keynote will explore the role of statistics and data science in addressing societal challenges, with a focus on communication, ethics, and interdisciplinary collaboration.
Her work highlights the impact of statistical thinking beyond academia, particularly in industry and decision-making contexts.
- Fabio Centofanti (KU Leuven, Belgium) – 2026 Young Statistician Award – Fabio Centofanti works focus on robust statistics, functional data analysis, and statistical process monitoring, addressing challenges in complex and high-dimensional data.
His talk explores statistical process monitoring methods for industrial profile data using functional data analysis, where profiles are treated as smooth functions observed over continuous domains. It presents a general framework for monitoring multivariate functional data and reviews recent advances addressing practical challenges such as outliers, partially observed profiles, adaptive detection strategies, and the integration of additional functional covariates.

Detailed Program

Sunday, September 6

Sun, September 6

1:45 –5:45 PM

ECAS-ENBIS Course: Adaptive Machine Learning for Time Series Forecasting

Location: Auditorium B | Instructor: Yannig Goude

Sun, September 6

3:15 –4:15 PM

Exec + Office meeting

Location: Room 107

Sun, September 6

4:30 –6:00 PM

Council Meeting

Location: Room 107

Monday, September 7

Mon, September 7

9:00 –9:30 AM

Opening Ceremony

Location: Auditorium B

Mon, September 7

9:30 –10:30 AM

Keynote

Location: Auditorium B | Chair: Biagio Palumbo

9:30 –10:30 AM **Physics-Informed Statistical Learning for Data-Driven Decision Making in Science and Industry**

Speaker: Laura M. Sangalli

Mon, September 7

10:55 –11:55 AM

AI: Machine Learning and Predictive Analytics

Location: Conference Room 103 | Chair: Nicolas Bousquet

10:55-11:15 AM **Data-Efficient Upscaling and Optimal Control of Fed-Batch Bioprocesses via Hybrid Modelling and Transfer Learning**

Speaker: Luca Riezzo

11:15 –11:35 AM **Local Constrained Bayesian Optimization for High-Dimensional Industrial Design**

Speaker: Jingzhe Jing

11:35 –11:55 AM **Extending Inverse Distance-based Exploration for Active Learning to Streaming: On-line Active Learning for Regression**

Speaker: Bernhard Spangl

Mon, September 7

10:55 –11:55 AM

Design of Experiments

Location: Auditorium B | Chair: Marco P. Seabra dos Reis

10:55-11:15 AM **Quality by Design: Data-Efficient Active Learning Approaches**

Speaker: Marco P. Seabra dos Reis

11:15 –11:35 AM **Fighting Antimicrobial Resistance, Functional Data Analysis, and the Future of Multi-factor Experiments**

Speaker: Phil Kay

11:35 –11:55 AM **Optimizing multi-arm clinical trials for personalized medicine using a genetic algorithm**

Speaker: Karla Cervantes

Mon, September 7

10:55 –11:55 AM

Statistical / Stochastic Modelling and Statistical Computing

Location: Conference Room 102 | Chair: Véronique Maume-Deschamps

10:55-11:15 AM **Least trimmed squares regression with missing values and cellwise outliers**

Speaker: Peter Rousseeuw

11:15 –11:35 AM **Characterization of multi-way binary tables with uniform margins and fixed correlations**

Speaker: Fabio Rapallo

11:35 –11:55 AM **Estimating Multivariate Generalized Gamma Convolutions via Kernel Stein Discrepancy**

Speaker: Léo Gonin

Mon, September 7

10:55 –11:55 AM

Statistical Process Monitoring**Location:** Conference Room 107 | **Chair:** Antonio Lepore**10:55-11:15 AM** **An entropy-based distribution-free approach for profile monitoring****Speaker:** Praise Obanya**11:15 –11:35 AM** **Nonparametric Statistical Process Control for Monitoring Structural Changes in Regional Drought Dynamics****Speaker:** Michele Scagliarini**11:35 –11:55 AM** **Linear Profile Monitoring of Functional Data****Speaker:** Davide Forcina**Mon, September 7**

10:55 –11:55 AM

Statistics in Pharma / Healthcare: Pharmaceutical Process Modelling and Manufacturing Analytics**Location:** Conference Room 106 | **Chair:** Marina Vives-Mestres**10:55-11:15 AM** **Estimating Stabilization Time in Continuous Manufacturing****Speaker:** Chellafe Ensoy-Musoro**11:15 –11:35 AM** **Incorporating Factor Variability into Risk Estimation in Pharmaceutical Manufacturing****Speaker:** Olympia Tumolva**Mon, September 7**

12:00 –1:00 PM

Design of Experiments**Location:** Auditorium B | **Chair:** Marco P. Seabra dos Reis**12:00 –12:20 PM** **Bayesian Optimisation under system drift****Speaker:** Stefanie Feiler**12:20 –12:40 PM** **Powerful Foldover Designs****Speaker:** Jonathan Stallrich**12:40 –1:00 PM** **Possible Pitfalls in Using Bayesian Optimization (BO) to Quickly Optimize the Levels of Factors in a Physical Experiment****Speaker:** Jacqueline Asscher**Mon, September 7**

12:00 –1:00 PM

Statistical / Stochastic Modelling and Statistical Computing**Location:** Conference Room 107 | **Chair:** Biagio Palumbo**12:00 –12:20 PM** **Tourist Mobility: A Markov Chain Approach Using Origin–Destination Survey Data****Speaker:** Francesca Atzori**12:20 –12:40 PM** **Adaptive Particle MCMC for non-linear battery state-space models****Speaker:** Edvinas Juozapaitis**12:40 –1:00 PM** **Robust graphical modeling under data contamination****Speaker:** Christian Capezza**Mon, September 7**

12:00 –1:00 PM

Statistical Process Monitoring**Location:** Conference Room 102 | **Chair:** Panagiotis Tsiamyrtzis**12:00 –12:20 PM** **When a Cusum Stops, What Confidence Is There that the Alarm Is Not False?****Speaker:** Moshe Pollak**12:20 –12:40 PM** **On the Design and Performance of a Control Chart for Monitoring Continuous data in (0,1) when Process Parameters are Unknown****Speaker:** Athanasios Rakitzis**12:40 –1:00 PM** **Some Stylized Facts of the Conditional Expected Delay (CED)****Speaker:** Sven Knoth

Mon, September 7

12:00 –1:00 PM

Statistics in Pharma / Healthcare: Statistical Decision Support in Healthcare and Biomedical Development**Location:** Conference Room 106 | **Chair:** Marina Vives-Mestres**12:00 –12:20 PM** Towards an Adaptive Framework for Longitudinal Survival Modeling in Clinical Decision Support: Case of the ISPY 1 Trial**Speaker:** Abdelaziz Berrado**12:20 –12:40 PM** Predicting Study Success in Diagnostic Test Evaluation: A Case Study Using Frequentist and Bayesian Approaches**Speaker:** Moreno Muñoz**Mon, September 7**

12:00 –1:00 PM

Truthworthy and explainable AI**Location:** Conference Room 103 | **Chair:** Nicolas Bousquet**12:00 –12:20 PM** Statistical Evaluation of CPU-Based Offline LLMs for Industrial Text Processing on Low-Cost Edge Hardware**Speaker:** Guido Moeser**12:20 –12:40 PM** Ups and Downs with AI and Old Data**Speaker:** Shirley Coleman**12:40 –1:00 PM** Towards the reliable use of metamodels in critical systems - Application to nuclear safety studies**Speaker:** Nicolas Bousquet, Bertrand Iooss**Mon, September 7**

2:00 –3:30 PM

Editors' Corner Session - INFORMS Journal on Data Science**Location:** Auditorium B | **Chair:** Yu Ding**2:00 –2:30 PM** Introduction to INFORMS Journal on Data Science (IJDS)**Speaker:** Yu Ding**2:30 –3:00 PM** Using Operational Data Analytics for Planning Decisions Under Uncertainty**Speaker:** Natarajan Gautam**3:00 –3:30 PM** Estimating Hidden Epidemic: A Bayesian Spatiotemporal Compartmental Modeling Approach**Speaker:** Kamran Paynabar**Mon, September 7**

2:00 –3:30 PM

Generative AI for Statistics: Practice, Productivity, and Prioritization**Location:** Conference Room 106 | **Chair:** Marina Vives-Mestres**2:00 –3:30 PM** Generative AI for Statistics: Practice, Productivity, and Prioritization**Speaker:** Christian Ritter, Morten Bormann Nielsen, Winfried Theis**Mon, September 7**

2:00 –3:30 PM

Statistics and data science in the technological field: current issues and new proposals**Location:** Conference Room 102 | **Chair:** Nedka Dechkova Nikiforova**2:00 –2:30 PM** ADeS: A Flexible Adaptive Design Strategy for Ethical and Covariate-Balanced Clinical Trials**Speaker:** Marco Novelli**2:30 –3:00 PM** Online Quality Monitoring in Selective Laser Melting via Image-Based Statistical Methods**Speaker:** Panagiotis Tsiamyrtzis**3:00 –3:30 PM** Fisher's Principles of Experimental Design and Why They Are Still Important**Speaker:** Geoff Vining

Mon, September 7

2:00 –3:30 PM

Young Statisticians Session**Location:** Conference Room 103 | **Chair:** Fabio Centofanti**2:00 –2:30 PM** **A Registration-free Approach for Shape and Color Monitoring of Functionally Graded Materials via 4D Point Clouds****Speaker:** Mariafrancesca Patalano**2:30 –3:00 PM** **On stochastic network trends in network time series****Speaker:** Amandine Pierrot**3:00 –3:30 PM** **Active Learning for Effect Screening in Manufacturing****Speaker:** Marcus Engsig**Mon, September 7**

4:00 –5:30 PM

Award session**Location:** Auditorium B | **Chair:** Winfried Theis**4:15 –4:45 PM** **The Irreducible Error: What Statistics and Management Have in Common****Speaker:** Sophie Carr**4:45 –5:15 PM** **Statistical Process Monitoring Based on Functional Data Analysis****Speaker:** Fabio Centofanti**Mon, September 7**

5:30 –6:30 PM

General Assembly**Location:** Auditorium B

Tuesday, September 8

Tue, September 8

9:00 –10:30 AM

Data science for climate and environment

Location: Conference Room 103 | **Chair:** Véronique Maume-Deschamps

9:00 –9:30 AM Model selection for extremal dependence structures using deep learning: Application to environmental data

Speaker: Pierre Ribereau

9:30 –10:00 AM Multivariate Discrete Generalized Pareto distributions: Theory, likelihood-free inference, and applications to drought risk assessment

Speaker: Samira Aka

10:00 –10:30 AM Concepts and methods for better predictions in climate and environmental science

Speaker: Georgia Papacharalampous

Tue, September 8

9:00 –10:30 AM

From Automation to Autonomy: New Frontiers of AI and Statistics for Data in Business and Industry

Location: Auditorium B | **Chair:** Christian Capezza, Mostafa Reisi Gahrooei

9:00 –10:30 AM From Automation to Autonomy: New Frontiers of AI and Statistics for Data in Business and Industry

Speaker: Bianca Maria Colosimo, Kamran Paynabar

Tue, September 8

9:00 –10:30 AM

ISBIS Session - Advances in Statistical Analysis for Large-Scale Dynamic Networks, Labour Market Outcomes, and Interpretable Ensemble Methods

Location: Conference Room 102 | **Chair:** Antonio Lepore

9:00 –9:30 AM Real-Time Change Detection in Large-Scale Dynamic Networks

Speaker: Nathaniel Stevens

9:30 –10:00 AM Heterogeneity in STEM Graduate Wages Across Italian Universities: A Weighted Mixed-Effects Analysis

Speaker: Andrea Carta

10:00 –10:30 AM Agreement over Association in Explainable Ensemble Trees

Speaker: Agostino Gnasso

Tue, September 8

9:00 –10:30 AM

Software Session

Location: Conference Room 106 | **Chair:** Volker Kraft

9:00 –9:30 AM Optimal design, analysis and optimization in the presence of random effects using Minitab DoE by Effex platform

Speaker: Jose Nunez Ares

9:30 –10:00 AM JASP for Quality Control: An Open-Source Software for Industrial Statistics

Speaker: Julius Pfadt, Don van den Bergh, Frantisek Bartos, Henrik Godmann, Jonas Petter, Johnny van Doorn, Simon Kucharsky, Eric-Jan Wagenmakers

10:00 –10:30 AM JMP Developments in JMP 19 and Coming in JMP 20

Speaker: Chris Gotwalt

Tue, September 8

10:55 –11:55 AM

Data Analytics and Data Science: Case Studies**Location:** Conference Room 103 | **Chair:** Morten Bormann Nielsen**10:55-11:15 AM** **Health Insurance Fraud Detection using Claim-Based Profiling****Speaker:** Sotiris Bersimis**11:15 –11:35 AM** **Simulating a key performance index for a family of future products****Speaker:** Carlo Leardi**11:35 –11:55 AM** **Dirty laundry –A data quality journey from ENBIS Active session to nationwide system redesign for industrial textile management****Speaker:** Morten Bormann Nielsen, Robert Heck

Tue, September 8

10:55 –11:55 AM

Reliability**Location:** Conference Room 107 | **Chair:** Antonio Pievatolo**10:55-11:15 AM** **Statistical Travel Time Inference Under Incomplete Data: A Queueing-Theoretic Approach****Speaker:** Dingyi Wang**11:15 –11:35 AM** **A scheme for the predictive replacement of lithium-ion batteries, using reinforcement learning****Speaker:** Antonio Pievatolo

Tue, September 8

10:55 –11:55 AM

Statistics in Industry, Business and Finance**Location:** Conference Room 102 | **Chair:** Jean-Michel Poggi**10:55-11:15 AM** **Lot Heterogeneity in Statistical Quality Control: An enhanced Optimization Approach for the Design of Rectifying Sampling Plans****Speaker:** Jaqueline Bojer**11:15 –11:35 AM** **Stacking Evidence across tests in R&D****Speaker:** Jan-Willem Bikker**11:35 –11:55 AM** **A Requirements-Driven Lifecycle Digital Twin Methodology for Innovative SMRs: Architecture Integration and Expected Benefits****Speaker:** Taecheol Park

Tue, September 8

10:55 –11:55 AM

Teaching, Consulting and Knowledge Transfer in Statistics**Location:** Conference Room 106 | **Chair:** Shirley Coleman**10:55-11:15 AM** **Bayesian Optimization as Design of Experiments: The Entropy Connection****Speaker:** Arno Strouwen**11:15 –11:35 AM** **A Practical Roadmap for Choosing Correct Statistical Tests, Based on Three Decades of Teaching Experience****Speaker:** Hossein Bevrani**11:35 –11:55 AM** **Plackett and Burman arrays Revisited****Speaker:** Jonathan Smyth-Renshaw

Tue, September 8

10:55 –11:55 AM

Uncertainty quantification and computer experiments**Location:** Auditorium B | **Chair:** Julien Pelamatti**10:55-11:15 AM** **Prediction of physical fields under linear constraints****Speaker:** Mahamat Hamdan Nassouradine**11:15 –11:35 AM** **Multifidelity Gaussian process regression for solving nonlinear partial differential equations****Speaker:** Fatima-Zahrae El-Boukkouri**11:35 –11:55 AM** **Taxonomy of statistical models for investigating order picking systems****Speaker:** Larissa Sander

Tue, September 8

12:00 –1:00 PM

Data Analytics and Data Science: Case Studies**Location:** Conference Room 103 | **Chair:** Andrea Ahlemeyer-Stubbe**12:00 –12:20 PM** An IoT-inspired concept for data-driven service provision and resource planning in the tourism sector, using a campsite as an example**Speaker:** Eva Scheideler**12:20 –12:40 PM** Multivariate Six Sigma for the Optimization of the Meat Roasting Process in the Ready-to-Eat Food Industry**Speaker:** Alberto Ferrer-Hermenegildo**12:40 –1:00 PM** Risk analysis of forest fires in Sicily, based on 2010-2023 observation period**Speaker:** Stefano Barone**Tue, September 8**

12:00 –1:00 PM

Reliability**Location:** Conference Room 107 | **Chair:** Antonio Pievatolo**12:00 –12:20 PM** Variable Selection under Cumulative Exposure Model for Time-to-Event Data with Time-Varying Covariates**Speaker:** Yueyao Wang**12:20 –12:40 PM** Reliability testing of repairable systems available in diverse configuration variants**Speaker:** Nikolaus Haselgruber**Tue, September 8**

12:00 –1:00 PM

Statistics in Industry, Business and Finance**Location:** Conference Room 102 | **Chair:** Jean-Michel Poggi**12:00 –12:20 PM** A Scalable Analytics Platform for Enterprise Level Quality Management**Speaker:** Manuel Martin**12:20 –12:40 PM** Adaptive soft sensor for product quality estimation in clinker production**Speaker:** Pierantonio Facco**12:40 –1:00 PM** The Lack of Impact of Lean Six Sigma in Textiles, Apparel and Other Basic Industries**Speaker:** A. Blanton Godfrey**Tue, September 8**

12:00 –1:00 PM

Teaching, Consulting and Knowledge Transfer in Statistics**Location:** Conference Room 106 | **Chair:** Marina Vives-Mestres**12:00 –12:20 PM** CROSSVALI: A Method for Prompting AI for Data Analysis**Speaker:** Anne Driscoll, Jennifer Van Mullekom**12:20 –12:40 PM** Bridging the Gap: Interactive Discovery as a Booster for Preparing Industry-Ready Graduates**Speaker:** Paolo Chiappa, Volker Kraft**12:40 –1:00 PM** Deep Adaptive Design for Model-Based DOE, Screening, and Bayesian Optimization**Speaker:** Arno Strouwen**Tue, September 8**

12:00 –1:00 PM

Uncertainty quantification and computer experiments**Location:** Auditorium B | **Chair:** Julien Pelamatti**12:00 –12:20 PM** Spectral Clustering for Detecting Stationary Subregions in Gaussian Process Regression**Speaker:** Théo Sylvestre**12:20 –12:40 PM** Multi-fidelity Gaussian processes for noisy outputs and non-nested experiment designs: a comparison between the recursive and non-recursive formulations**Speaker:** Nils Baillie**12:40 –1:00 PM** Integrating Global Sensitivity Analysis into Bayesian Optimization for Tactical Decision Support in Transport Logistics**Speaker:** Lara Kuhlmann de Canaviri

Tue, September 8

2:00 –3:30 PM

Design and modeling of computer experiments**Location:** Conference Room 102 | **Chair:** Bertrand Iooss**2:00 –2:30 PM** Random Designs for Quantisation in High Dimension**Speaker:** Luc Pronzato**2:30 –3:00 PM** Active Subspace Embedding for Parsimonious Gaussian Process Surrogates**Speaker:** Amandine Marrel**10:00 –10:30 AM** Multivariate sensitivity analysis for risk analysis with spatial outputs: a comparison exercise**Speaker:** Jérémy Rohmer**Tue, September 8**

2:00 –3:30 PM

Hands-on with physical experiments for teaching Design of Experiments**Location:** Conference Room 103 | **Chair:** Morten Bormann Nielsen, Christian Ritter**2:00 –3:30 PM** Hands-on with physical experiments for teaching Design of Experiments**Speaker:** Christian Ritter, Morten Bormann Nielsen**Tue, September 8**

2:00 –3:30 PM

ISEA Session - Statistical Engineering**Location:** Auditorium B | **Chair:** Inez Zwetsloot**2:00 –2:30 PM** XAI for signal diagnosis in SPM**Speaker:** Inez Zwetsloot**2:30 –3:00 PM** Assessing inter-rater reliability of LLMs with the probability of agreement**Speaker:** Nathaniel Stevens**10:00 –10:30 AM** A Cooperative Framework for Raw Material Acceptance: Integrating Supplier and Customer Perspectives via SMB-PLS**Speaker:** Alberto J. Ferrer-Riquelme**Tue, September 8**

2:00 –3:30 PM

Mathmet session - Digital twins for industrial machine vision systems and reference data generation**Location:** Conference Room 106 | **Chair:** Hichem Nouira**2:05 –2:25 PM** Digital Twin for optimal positioning of an industrial machine vision camera**Speaker:** Dario Pasin**2:25 –2:45 PM** Enhancing photogrammetric measurement systems through Gaussian Splatting**Speaker:** Mattia Trombini**2:45 –3:05 PM** Development of digital twins in metrology for a stereovision system**Speaker:** Katarina Josic**3:05 –3:30 PM** Generalized constraints on reference data generation for software verification**Speaker:** Louis-Ferdinand Lafon**Tue, September 8**

4:00 –5:00 PM

Box medal**Location:** Auditorium B | **Chair:** Winfried Theis**4:00 –5:00 PM** Enhancing Definitive Screening Designs with Multilevel Categorical Factors**Speaker:** Bradley Jones**Tue, September 8**

5:00 –5:10 PM

Greenfield challenge**Location:** Auditorium B | **Chair:** Winfried Theis**5:00 –5:10 PM** An Academic Year Dedicated to George E. P. Box at the Faculty of Mathematics and Statistics in Barcelona**Speaker:** Lluís Marco-Almagro

Wednesday, September 9

Wed, September 9

9:00 –10:30 AM

INFORMS Session - Quality Statistics Reliability (QSR)

Location: Auditorium B | **Chair:** Christian Capezza, Mostafa Reisi Gahrooei

9:00 –9:30 AM Response Grid Plots for Model-Agnostic Explainable AI

Speaker: Russell Barton

9:30 –10:00 AM FedCOT: Personalized Federated Transfer Learning With Conditional Optimal Transport for Manufacturing Predictive Modeling

Speaker: Jiaqi Qiu

10:00 –10:30 AM Spatial–Temporal Large Language Models for Time-Series Data

Speaker: Ziyue Li

Wed, September 9

9:00 –10:30 AM

Uncertainty quantification and sensitivity analysis

Location: Conference Room 103 | **Chair:** Véronique Maume-Deschamps

9:00 –9:30 AM Sensitivity analysis for sets: Application to pollutant concentration maps

Speaker: Celine Helbert

9:30 –10:00 AM Explicit functional ANOVA and new results in model and algorithm explainability and sensitivity analysis

Speaker: Nicolas Bousquet

10:00 –10:30 AM Sensitivity Measures for Continuous Actions

Speaker: Lea Friedli

Wed, September 9

9:00 –10:30 AM

itENBIS Session - Flexible and Scalable Models for Complex Data Structures

Location: Conference Room 106 | **Chair:** Amalia Vanacore

9:00 –9:30 AM Fuzzy jump models for soft and hard clustering of mixed-type multivariate time series

Speaker: Federico Cortese

9:30 –10:00 AM Variable selection and high dimension in multinomial models

Speaker: Marcella Niglio

10:00 –10:30 AM Small area models for count data: an area-level EFD model

Speaker: Serena Arima

Wed, September 9

9:30 –10:30 AM

AI: Machine Learning and Predictive Analytics

Location: Conference Room 102 | **Chair:** Francesca Bassi

9:30 –9:50 AM Belief Propagation for Early Sequential Feature Selection in Manufacturing

Speaker: Joana Martins

9:50 –10:10 AM Decision Reliability in Uplift Modeling: A Stability- and Value-Driven Framework for Customer Retention Policy Selection

Speaker: Massimo Pacella

10:10 –10:30 AM Wrapper-Based Variable Selection for Interacting Variables in Manufacturing

Speaker: Marcus Engsig

Wed, September 9

10:55 –11:55 AM

AI: Machine Learning and Predictive Analytics**Location:** Conference Room 102 | **Chair:** Riccardo Ceccato**10:55 –11:15 AM** **A Two-Dimensional View of Classification Difficulty: Local Ambiguity and Global Separability****Speaker:** Yariv N. Marmor**11:15 –11:35 AM** **Ranking Sets of Variables by Multivariate Association: a Non Parametric Combination Approach****Speaker:** Elena Barzizza**11:35 –11:55 AM** **A two-layer model for opinion spread in hypergraphs as a Markov chain and as the object of machine learning methods****Speaker:** András Zempléni**Wed, September 9**

10:55 –11:55 AM

Design of Experiments**Location:** Auditorium B | **Chair:** Luc Pronzato**10:55 –11:15 AM** **Gradient-Based Line Search for Continuous Optimal Designs in Spatial Statistics****Speaker:** Helmut Waldl**11:15 –11:35 AM** **Robust D-Optimal Designs for Ordinal Split-Plot Experiments via Surrogate Objective Functions****Speaker:** Marcus Perry**11:35 –11:55 AM** **Model Selection for Screening Experiments with Random Effects: An Adaptation of the SAMS Algorithm****Speaker:** Robin van der Haar**Wed, September 9**

10:55 –11:55 AM

Statistics in Industry, Business and Finance**Location:** Conference Room 103 | **Chair:** Andrea Ahlemeyer-Stubbe**10:55 –11:15 AM** **Understanding Tourism Demand: A Tree-Based Analysis of Territorial Drivers****Speaker:** Giulia Contu**11:15 –11:35 AM** **Gastronomic Tourism in Italy: Measuring Tourist Attitudes and Identifying Market Segments****Speaker:** Francesca Bassi**Wed, September 9**

10:55 –11:55 AM

Time Series, Forecasting and Dynamic Systems**Location:** Conference Room 106 | **Chair:** Ouassim Feliachi**10:55 –11:15 AM** **A Semi-Supervised Framework for Optimisation-Based Label Inference in RUL Prediction****Speaker:** Joseph Cupit**11:15 –11:35 AM** **Diffusion Models for Multivariate Probabilistic Net Load Forecasting in Transmission Grid Congestion Management****Speaker:** Maxime Ducourau**11:35 –11:55 AM** **INAR(1) Processes based on the Zipf-PSS distribution****Speaker:** Marta Perez-Casany**Wed, September 9**

12:00 –1:30 PM

Active Session - ENBIS-Live: Real Problems**Location:** Conference Room 103 | **Chair:** Jennifer Van Mullekom**12:00 –1:30 PM** **ENBIS-Live: Real Problems. Real Time. Real Stats.****Speaker:** Jennifer Van Mullekom

Wed, September 9

12:00 –1:30 PM

Bridging AI and Statistics**Location:** Conference Room 102 | **Chair:** Bart De Ketelaere**12:00 –12:30 PM** **Applied Statistics in the Era of Artificial Intelligence: A focus on AI Assurance****Speaker:** Jie Min**12:30 –1:00 PM** **Novel AI Solutions for Process Monitoring and Optimization****Speaker:** Bianca Maria Colosimo**1:00 –1:30 PM** **Discriminative Open Set Active Learning (DOSAL)****Speaker:** Bart De Ketelaere**Wed, September 9**

12:00 –1:30 PM

QT / QE / Technometrics Session**Location:** Auditorium B | **Chair:** Sven Knoth**12:00 –12:30 PM** **Constructing large OMARS designs by concatenating two definitive screening designs****Speaker:** Peter Goos**12:30 –1:00 PM** **The allure and the perils of summarizing process performance in a single index****Speaker:** Mahmut Onur Karaman**1:00 –1:30 PM** **Time series models for cultural heritage preservation: The case of the Brunelleschi's Dome****Speaker:** Fiammetta Menchetti**Wed, September 9**

12:00 –1:30 PM

frENBIS Session - Advances in machine learning and sensitivity analysis**Location:** Conference Room 106 | **Chair:** Jean-Michel Poggi**12:00 –12:30 PM** **Geographically Weighted Regression for Low-Cost Sensors Calibration****Speaker:** Jean-Michel Poggi**12:30 –1:00 PM** **Quantile oriented sensitivity analysis: why and how?****Speaker:** Véronique Maume-Deschamps**1:00 –1:30 PM** **Shapelet learning via sparse feature selection for interpretable time series classification****Speaker:** Julien Pelamatti**Wed, September 9**

2:30 –3:30 PM

AI: Machine Learning and Predictive Analytics**Location:** Conference Room 103 | **Chair:** Morten Bormann Nielsen**2:30 –2:50 PM** **Semi-Supervised PARAFAC2 Decomposition for Computational Phenotyping Using Electronic Health Records****Speaker:** Mostafa Reisi Gahrooei**2:50 –3:10 PM** **Sustainable specialty chemical design through multi-objective latent-variable model inversion****Speaker:** Pierantonio Facco**3:10 –3:30 PM** **Decomposition of Serially Dependent Data into Static and Dynamic Latent****Structures****Speaker:** Moritz Bauchrowitz**Wed, September 9**

2:30 –3:30 PM

Data Analytics and Data Science: Case Studies**Location:** Conference Room 107 | **Chair:** Elena Barzizza**2:30 –2:50 PM** **Nonparametric tests for multivariate stability data****Speaker:** Riccardo Ceccato**2:50 –3:10 PM** **Clustering for Detecting Careless Respondents: A Case Study****Speaker:** Silvia Villanova**3:10 –3:30 PM** **Preference mapping for product optimization: a case study****Speaker:** Alessandro Fanesi

Wed, September 9

2:30 –3:30 PM

Measurement Uncertainty and Metrology**Location:** Conference Room 106 | **Chair:** Hichem Nouria**2:30 –2:50 PM Bias and Multiplicity in Producer Risk Estimation under Simplified Conformity Assessment Models****Speaker:** Omar Jair Purata Sifuentes**2:50 –3:10 PM A study of prior sensitivity in random effects models for decisions in interlaboratory comparisons: credible intervals versus Bayesian evidence****Speaker:** Séverine Demeyer**3:10 –3:30 PM An Interactive Software Application for Gauge R&R Analysis****Speaker:** Mahmut Onur Karaman**Wed, September 9**

2:30 –3:10 PM

Statistical Process Monitoring**Location:** Auditorium B | **Chair:** Sven Knoth**2:30 –2:50 PM Self-Starting Control Charts for Few-Shot In Situ Monitoring in Customized Manufacturing****Speaker:** Bianca Maria Colosimo**2:50 –3:10 PM Practical challenges and statistical solutions for process monitoring and control in dairy production****Speaker:** Ewa Szymanska**Wed, September 9**

2:30 –3:30 PM

Statistics in Pharma / Healthcare: Statistical Modelling and Learning for Biomedical Systems**Location:** Conference Room 102 | **Chair:** Hichem Nouria**2:30 –2:50 PM Kinetic Models by Physics-Informed Statistical Learning Methods****Speaker:** Marco Galliani**2:50 –3:10 PM Bayesian network-based Tikhonov MRI reconstruction****Speaker:** Sébastien Marmin**3:10 –3:30 PM Statistics Can Go Farther****Speaker:** Dibyen Majumdar**Wed, September 9**

3:35 –4:35 PM

Keynote**Location:** Auditorium B | **Chair:** Bertrand Iooss**3:35 –4:35 PM Domain Generalization and Adaptation in Digital Health (and beyond)****Speaker:** Peter Bühlmann**Wed, September 9**

4:35 –4:45 PM

Closing Ceremony**Location:** Auditorium B**Wed, September 9**

5:00 –6:00 PM

ENBIS 2026 debrief meeting (Head LOC/POC Florence & Barcelona, Exec, Office)**Location:** Auditorium B

Thursday, September 10

Thu, September 10

9:00 AM – 1:00 PM

JMP Workshop: Human + AI in the Loop: AI-Supported Analysis and Learning in JMP

Location: Auditorium B | **Instructor:** Volker Kraft, Winfrid Theis

Thu, September 10

1:45 – 5:45 PM

Consultancy Skills: Storytelling with Data

Location: Auditorium B | **Instructor:** Winfried Theis, Morten Bormann Nielsen

Abstracts

ENBIS-Live: Real Problems. Real Time. Real Stats.	25
Data-Efficient Upscaling and Optimal Control of Fed-Batch Bioprocesses via Hybrid Modelling and Transfer Learning	26
Local Constrained Bayesian Optimization for High-Dimensional Industrial Design . . .	27
Extending Inverse Distance-based Exploration for Active Learning to Streaming: Online Active Learning for Regression	28
Belief Propagation for Early Sequential Feature Selection in Manufacturing	29
Decision Reliability in Uplift Modeling: A Stability- and Value-Driven Framework for Customer Retention Policy Selection	30
Wrapper-Based Variable Selection for Interacting Variables in Manufacturing	31
A Two-Dimensional View of Classification Difficulty: Local Ambiguity and Global Separability	32
Ranking Sets of Variables by Multivariate Association: a Non Parametric Combination Approach	33
A two-layer model for opinion spread in hypergraphs as a Markov chain and as the object of machine learning methods	33
Semi-Supervised PARAFAC2 Decomposition for Computational Phenotyping Using Electronic Health Records	34
Sustainable specialty chemical design through multi-objective latent-variable model inversion	35
Decomposition of Serially Dependent Data into Static and Dynamic Latent Structures . .	36
The Irreducible Error: What Statistics and Management Have in Common	36
Statistical Process Monitoring Based on Functional Data Analysis	37
Enhancing Definitive Screening Designs with Multilevel Categorical Factors	37
Applied Statistics in the Era of Artificial Intelligence: A focus on AI Assurance	38
Novel AI Solutions for Process Monitoring and Optimization	38
Discriminative Open Set Active Learning (DOSAL)	39

Health Insurance Fraud Detection using Claim-Based Profiling	39
Simulating a key performance index for a family of future products	40
Dirty laundry –A data quality journey from ENBIS Active session to nationwide system redesign for industrial textile management	40
An IoT-inspired concept for data-driven service provision and resource planning in the tourism sector, using a campsite as an example	41
Multivariate Six Sigma for the Optimization of the Meat Roasting Process in the Ready-to- Eat Food Industry	42
Risk analysis of forest fires in Sicily, based on 2010-2023 observation period	43
Nonparametric tests for multivariate stability data	44
Clustering for Detecting Careless Respondents: A Case Study	44
Preference mapping for product optimization: a case study	45
Model selection for extremal dependence structures using deep learning: Application to environmental data	46
Multivariate Discrete Generalized Pareto distributions: Theory, likelihood-free inference, and applications to drought risk assessment	46
Concepts and methods for better predictions in climate and environmental science	47
Random Designs for Quantisation in High Dimension	47
Active Subspace Embedding for Parsimonious Gaussian Process Surrogates	48
Multivariate sensitivity analysis for risk analysis with spatial outputs: a comparison exer- cise	49
Quality by Design: Data-Efficient Active Learning Approaches	50
Fighting Antimicrobial Resistance, Functional Data Analysis, and the Future of Multifactor Experiments	51
Optimizing multi-arm clinical trials for personalized medicine using a genetic algorithm	51
Bayesian Optimisation under system drift	52
Powerful Foldover Designs	52
Possible Pitfalls in Using Bayesian Optimization (BO) to Quickly Optimize the Levels of Factors in a Physical Experiment	53
Gradient-Based Line Search for Continuous Optimal Designs in Spatial Statistics	54
Robust D-Optimal Designs for Ordinal Split-Plot Experiments via Surrogate Objective Func- tions	54
Model Selection for Screening Experiments with Random Effects: An Adaptation of the SAMS Algorithm	55

Introduction to INFORMS Journal on Data Science (IJDS)	55
Using Operational Data Analytics for Planning Decisions Under Uncertainty	56
Estimating Hidden Epidemic: A Bayesian Spatiotemporal Compartmental Modeling Approach	56
Geographically Weighted Regression for Low-Cost Sensors Calibration	57
Quantile oriented sensitivity analysis: why and how?	57
Shapelet learning via sparse feature selection for interpretable time series classification	58
From Automation to Autonomy: New Frontiers of AI and Statistics for Data in Business and Industry	59
Generative AI for Statistics: Practice, Productivity, and Prioritization	60
An Academic Year Dedicated to George E. P. Box at the Faculty of Mathematics and Statistics in Barcelona	61
Hands-on with physical experiments for teaching Design of Experiments	62
Response Grid Plots for Model-Agnostic Explainable AI	62
FedCOT: Personalized Federated Transfer Learning With Conditional Optimal Transport for Manufacturing Predictive Modeling	63
Spatial–Temporal Large Language Models for Time-Series Data	63
Real-Time Change Detection in Large-Scale Dynamic Networks	64
Heterogeneity in STEM Graduate Wages Across Italian Universities: A Weighted Mixed-Effects Analysis	64
Agreement over Association in Explainable Ensemble Trees	65
XAI for signal diagnosis in SPM	65
Assessing inter-rater reliability of LLMs with the probability of agreement	66
A Cooperative Framework for Raw Material Acceptance: Integrating Supplier and Customer Perspectives via SMB-PLS	67
Fuzzy jump models for soft and hard clustering of mixed-type multivariate time series	68
Variable selection and high dimension in multinomial models	68
Small area models for count data: an area-level EFD model	69
Constructing large OMARS designs by concatenating two definitive screening designs	70
The allure and the perils of summarizing process performance in a single index	70
Time series models for cultural heritage preservation: The case of the Brunelleschi’s Dome	71

Physics-Informed Statistical Learning for Data-Driven Decision Making in Science and Industry	72
Domain Generalization and Adaptation in Digital Health (and beyond)	72
Digital Twin for optimal positioning of an industrial machine vision camera	73
Enhancing photogrammetric measurement systems through Gaussian Splatting	74
Development of digital twins in metrology for a stereovision system	75
Generalized constraints on reference data generation for software verification	76
Bias and Multiplicity in Producer Risk Estimation under Simplified Conformity Assessment Models	77
A study of prior sensitivity in random effects models for decisions in interlaboratory comparisons: credible intervals versus Bayesian evidence	78
An Interactive Software Application for Gauge R&R Analysis	79
Statistical Travel Time Inference Under Incomplete Data: A Queueing-Theoretic Approach	80
A scheme for the predictive replacement of lithium-ion batteries, using reinforcement learning	81
Variable Selection under Cumulative Exposure Model for Time-to-Event Data with Time-Varying Covariates	81
Reliability testing of repairable systems available in diverse configuration variants	82
Optimal design, analysis and optimization in the presence of random effects using Minitab DoE by Effex platform	82
JASP for Quality Control: An Open-Source Software for Industrial Statistics	83
JMP Developments in JMP 19 and Coming in JMP 20	83
Least trimmed squares regression with missing values and cellwise outliers	84
Characterization of multi-way binary tables with uniform margins and fixed correlations	84
Estimating Multivariate Generalized Gamma Convolutions via Kernel Stein Discrepancy	85
Tourist Mobility: A Markov Chain Approach Using Origin–Destination Survey Data	85
Adaptive Particle MCMC for non-linear battery state-space models	86
Robust graphical modeling under data contamination	87
An entropy-based distribution-free approach for profile monitoring	87
Nonparametric Statistical Process Control for Monitoring Structural Changes in Regional Drought Dynamics	88
Linear Profile Monitoring of Functional Data	88

When a Cusum Stops, What Confidence Is There that the Alarm Is Not False?	89
On the Design and Performance of a Control Chart for Monitoring Continuous data in (0,1) when Process Parameters are Unknown	89
Some Stylized Facts of the Conditional Expected Delay (CED)	90
Self-Starting Control Charts for Few-Shot In Situ Monitoring in Customized Manufacturing	90
Practical challenges and statistical solutions for process monitoring and control in dairy production	91
ADeS: A Flexible Adaptive Design Strategy for Ethical and Covariate-Balanced Clinical Trials	92
Online Quality Monitoring in Selective Laser Melting via Image-Based Statistical Methods	92
Fisher's Principles of Experimental Design and Why They Are Still Important	93
Lot Heterogeneity in Statistical Quality Control: An enhanced Optimization Approach for the Design of Rectifying Sampling Plans	93
Stacking Evidence across tests in R&D	94
A Requirements-Driven Lifecycle Digital Twin Methodology for Innovative SMRs: Archi- tecture Integration and Expected Benefits	95
A Scalable Analytics Platform for Enterprise Level Quality Management	96
Adaptive soft sensor for product quality estimation in clinker production	97
The Lack of Impact of Lean Six Sigma in Textiles, Apparel and Other Basic Industries . .	98
Understanding Tourism Demand: A Tree-Based Analysis of Territorial Drivers	99
Gastronomic Tourism in Italy: Measuring Tourist Attitudes and Identifying Market Seg- ments	99
Estimating Stabilization Time in Continuous Manufacturing	100
Incorporating Factor Variability into Risk Estimation in Pharmaceutical Manufacturing .	100
Multi-attribute modelling for mRNA specification setting	101
Towards an Adaptive Framework for Longitudinal Survival Modeling in Clinical Decision Support: Case of the ISPY 1 Trial	101
Predicting Study Success in Diagnostic Test Evaluation: A Case Study Using Frequentist and Bayesian Approaches	102
Kinetic Models by Physics-Informed Statistical Learning Methods	103
Bayesian network-based Tikhonov MRI reconstruction	104
Statistics Can Go Farther	104

Bayesian Optimization as Design of Experiments: The Entropy Connection	105
A Practical Roadmap for Choosing Correct Statistical Tests, Based on Three Decades of Teaching Experience	106
Plackett and Burman arrays Revisited	107
CROSSVALI: A Method for Prompting AI for Data Analysis	107
Bridging the Gap: Interactive Discovery as a Booster for Preparing Industry-Ready Graduates	108
Deep Adaptive Design for Model-Based DOE, Screening, and Bayesian Optimization . .	108
A Semi-Supervised Framework for Optimisation-Based Label Inference in RUL Prediction	109
Diffusion Models for Multivariate Probabilistic Net Load Forecasting in Transmission Grid Congestion Management	110
INAR(1) Processes based on the Zipf-PSS distribution	110
Statistical Evaluation of CPU-Based Offline LLMs for Industrial Text Processing on Low-Cost Edge Hardware	111
Ups and Downs with AI and Old Data	112
Towards the reliable use of metamodels in critical systems - Application to nuclear safety studies	113
Prediction of physical fields under linear constraints	114
Multifidelity Gaussian process regression for solving nonlinear partial differential equations	115
Taxonomy of statistical models for investigating order picking systems	116
Spectral Clustering for Detecting Stationary Subregions in Gaussian Process Regression	117
Multi-fidelity Gaussian processes for noisy outputs and non-nested experiment designs: a comparison between the recursive and non-recursive formulations	118
Integrating Global Sensitivity Analysis into Bayesian Optimization for Tactical Decision Support in Transport Logistics	119
Sensitivity analysis for sets: Application to pollutant concentration maps	119
Explicit functional ANOVA and new results in model and algorithm explainability and sensitivity analysis	120
Sensitivity Measures for Continuous Actions	120
A Registration-free Approach for Shape and Color Monitoring of Functionally Graded Materials via 4D Point Clouds	121
On stochastic network trends in network time series	121
Active Learning for Effect Screening in Manufacturing	122

Active Session - ENBIS-Live: Real Problems / 14

ENBIS-Live: Real Problems. Real Time. Real Stats.

Author: Jennifer Van Mullekom¹

¹ *Virginia Polytechnic Institute & State University*

Join us for one of the most dynamic and interactive sessions of the ENBIS conference —ENBIS-Live: Open Problem Solving in Action!

This is no ordinary talk. In this fast-paced, high-energy session, statisticians and data scientists roll up their sleeves to tackle real-world open problems —live and on the spot. Think of it as a collaborative brain trust powered by the collective wisdom and creativity of the ENBIS community. Here's how it works: Got a tricky problem? Volunteer to present it in 7 minutes. Outline the background, the data science or statistical angle, and where you're stuck. Need clarity? The audience asks their burning questions. Have ideas? The floor opens to contributions, suggestions, and insights from the crowd. Wrap-up: The problem owner reflects on the input, and the session host ties it all together.

Each problem gets 20–30 minutes of focused, expert attention. We'll feature 2–3 open problems —so come ready to think, question, suggest, and be inspired!

This year we would like to take it one step further to test how AI fares as a collaborative problem solver.

Got a challenge you'd like to throw into the ring or want to use your AI skills for one or more problems? Contact Jennifer Van Mullekom (vanmuljh@vt.edu) and be part of this one-of-a-kind problem-solving spectacle.

Bring your brains. Bring your curiosity. And let's solve some problems —ENBIS-style!

Keywords:

Collaboration, Problem Solving, Teamwork

AI: Machine Learning and Predictive Analytics / 22**Data-Efficient Upscaling and Optimal Control of Fed-Batch Bioprocesses via Hybrid Modelling and Transfer Learning****Author:** Luca Riezzo¹**Co-authors:** Dongda Zhang²; Harvey Al-Ramadhan³¹ *The University of Manchester*² *University of Manchester*³ *university of Manchester*

Optimisation and control of industrial fed-batch bioprocesses remains a complex and resource intensive task. Process development from laboratory to manufacturing scale requires extensive experimental trials to characterise system behaviour and identify optimal conditions, incurring significant costs. While advanced process control (APC) technologies have been widely applied to industry, their application is limited by the lack of data and predictive models at manufacturing scale. Within the last decade, hybrid modelling, combining mechanistic knowledge with data-driven techniques, has emerged as a promising approach to accelerate scale-up of complex processes. Furthermore, transfer learning has shown to significantly reduce data requirements for reliable predictions across scale. Therefore, to overcome the challenges of bioprocess development and optimisation at scale, this work proposes a unique predictive optimal control strategy that integrates two powerful modelling techniques: hybrid modelling and transfer learning. A hybrid model was trained on 2 datasets at laboratory scale and applied within a model predictive control framework to achieve a desired end-point product concentration by manipulating substrate feed flow within a fed-batch bioreactor at production scale. Two strategies were compared in this work: updating optimal control actions with and without transfer learning-based hybrid model update. Through comparison, the model update strategy demonstrated superior predictive performance however both strategies were successful at satisfying the end-point criteria. Overall, this work demonstrated feasibility to combine hybrid modelling with transfer learning to achieve accurate prediction and optimal control at scale and accelerate bioprocess development.

Keywords:

Machine Learning, Hybrid Modelling, Bioprocessing

AI: Machine Learning and Predictive Analytics / 39**Local Constrained Bayesian Optimization for High-Dimensional Industrial Design****Author:** Jingzhe Jing¹**Co-authors:** Zheyi Fan¹; Szu Hui Ng²; Qingpei Hu¹¹ *Academy of Mathematics and Systems Science, Chinese Academy of Sciences*² *National University of Singapore*

Bayesian optimization (BO) has become a cornerstone methodology for the data-efficient tuning of expensive black-box systems encountered throughout business and industrial practice, ranging from chemical process design and structural engineering to controller calibration and the configuration of large-scale machine learning pipelines. In most of these applications, the objective must be optimized subject to unknown and equally expensive-to-evaluate constraints reflecting safety thresholds, physical feasibility, or budgetary limits. While constrained BO (CBO) addresses this need in low dimensions, its performance degrades sharply as the number of design variables grows, since standard regret bounds scale exponentially with dimension. Existing high-dimensional remedies based on trust regions are prone to premature shrinkage when the descent direction is blocked by tight or complex constraints, a failure mode we observe repeatedly on industrial design tasks.

We propose Local Constrained Bayesian Optimization (LCBO), a framework that extends gradient-based local BO to constrained problems via a quadratic-penalty surrogate. LCBO alternates between rapid local descent and uncertainty-driven exploration. Under mild regularity conditions, we prove that the KKT residual converges at a rate depending only polynomially on the dimension for common kernels, which is a marked improvement over global CBO guarantees. Empirically, on synthetic benchmarks up to 100 dimensions, a 25-bar truss design, a 50-dimensional stepped cantilever beam, and a 102-dimensional MuJoCo control task, LCBO consistently outperforms state-of-the-art CBO baselines in sample efficiency, stability, and final solution quality, demonstrating its promise for real-world industrial optimization under constraints.

Keywords:

high dimensional optimization, Bayesian optimization, local search, convergence analysis

AI: Machine Learning and Predictive Analytics / 122

Extending Inverse Distance-based Exploration for Active Learning to Streaming: Online Active Learning for Regression

Author: Bernhard Spangl¹**Co-authors:** Wilhelm Otto²; Xiaojian Xu³¹ *BOKU University, Vienna*² *University of Applied Sciences Wiener Neustadt*³ *Brock University, St. Catharines*

In statistical and machine learning, efficient data acquisition is pivotal to model performance, particularly when labeled data are costly or time-intensive to obtain. This motivates active learning, in which the learning algorithm selectively queries maximally informative data points to accelerate training and improve predictive efficiency. While many active learning strategies consider query synthesis or pool-based sampling, we address the problem of online active learning in regression scenarios. In the stream-based scenario, unlabeled instances arrive continuously and must be either queried (labeled) or skipped on the fly. We introduce a new algorithm that adaptively selects which instances to label from a data stream by thresholding a suitably designed acquisition function. The method transfers and extends the Inverse Distance-based Exploration for Active Learning (IDEAL) principle, originally developed for pool-based settings, to the streaming context. This transfer preserves IDEAL's balance between exploitation of the current model structure and exploration to promote diversity in the feature space. We benchmark our method against state-of-the-art active learning strategies and against a passive baseline that labels incoming stream instances at random with a fixed probability, providing a clear reference for gains attributable to targeted querying. Performance is assessed through controlled numerical experiments on illustrative synthetic regression problems. We further demonstrate practical utility on a real chemometric data set.

Keywords:

online active learning, stream-based selective sampling, regression

AI: Machine Learning and Predictive Analytics / 28**Belief Propagation for Early Sequential Feature Selection in Manufacturing****Author:** Joana Martins¹**Co-authors:** Diogo Costa¹; Eugénio Rocha¹¹ *University of Aveiro*

Early sequential feature selection is crucial in manufacturing environments with heterogeneous sensors and tightly coupled process variables. In shop-floor applications, predicting End-of-Line test failures using data collected as early as possible is vital to enable timely corrective actions. However, classical feature selection and Explainable Artificial Intelligence methods (e.g., SHAP) evaluate variables globally or independently. Consequently, they often highlight temporally dispersed or late-stage variables as highly predictive, which prevents rapid early interventions.

To address this, we propose a novel Factor Graph topology and an associated Belief Propagation algorithm for early sequential feature selection on classification tasks. The factor graph encodes conditional relationships among candidate features and the target, computing “Influence Scores” to estimate each feature’s marginal. This space is then compressed into a highly informative subset of seven summary statistics.

Evaluated on a Bosch Thermotechnology heat pump production line, our approach restricted the analysis exclusively to an early “Instability Zone”. Our transparent, rule-based classification achieved a Macro F1-Score of 0.750, practically matching the performance of complex Machine Learning models (e.g., XGBoost and Random Forest) trained on the same early zone. Conversely, restricting complex models to the top 10 globally dominant features decreased performance (0.744). Ultimately, this approach eliminates the need for heavy Machine Learning inference on the shop floor and is adaptable to other continuous manufacturing processes. Furthermore, the proposed model is primed for future Digital Twin integration, projecting a massive time reduction of ≥ 8 hours per failed heat pump.

Keywords:

Sequential Feature Selection, Belief Propagation, Factor Graphs, Early Prediction, Predictive Analytics

AI: Machine Learning and Predictive Analytics / 27

Decision Reliability in Uplift Modeling: A Stability- and Value-Driven Framework for Customer Retention Policy Selection

Authors: Massimo Pacella¹; Gabriele Papadia²; Vincenzo Giliberti³

¹ *Department of Engineering for Innovation, Università del Salento- Lecce, ITALY*

² *Department of Engineering for Innovation, University of Salento, Lecce - ITALY*

³ *IN & OUT S.p.A. a Socio Unico Teleperformance S.E., Taranto - ITALY*

A stability-aware, value-driven framework for uplift policy selection is presented, applied to a telecommunications churn-retention dataset. The central argument is that reliable deployment of an uplift model requires more than optimizing a causal ranking metric: the targeted customer set must remain consistent across repeated training runs, and the selected policy must be economically sound and interpretable. Competing algorithmic families, including causal meta-learners, transformed-outcome approaches, four-quadrant models, and a conventional response baseline, are evaluated under a repeated-run experimental protocol that jointly quantifies economic effectiveness, run-to-run reliability of both targeted and persuadable customer sets via Jaccard-based consistency indices, and interpretability through surrogate CART distillation into compact Boolean decision rules. A Stability-Adjusted Revenue index synthesizes value and reliability into a single deployment-oriented criterion. Two complementary empirical regimes are examined: a semi-synthetic scenario with controlled treatment effects and an observational proxy scenario. Results reveal two significant findings: (i) the highest-yielding policy suffers from severe targeting instability, exposing a structural risk-return trade-off concealed by single-metric selection; and (ii) causal architectures do not universally dominate, as a conventional response model proves economically superior under confounded observational conditions. Developed within a regional research program and validated through an ongoing industrial implementation, this work extends the methodological framework toward practical, auditable deployment of causal machine learning.

Keywords:

Uplift Modeling; Decision Reliability; Algorithmic Stability

AI: Machine Learning and Predictive Analytics / 64

Wrapper-Based Variable Selection for Interacting Variables in Manufacturing**Author:** Marcus Engsig¹**Co-author:** Murat Kulahci²¹ *Technical University of Denmark*² *DTU*

In manufacturing, identifying variables that influence process outcomes is essential for control and optimization. While wrapper-based variable selection methods, such as Conditional Boruta by Rotari and Kulahci(2025), has been proposed, they remain susceptible to rejecting variables that only influence the outcome through interactions with other variables. Therefore, the purpose of the suggested method in this paper, is to extend the application to work for cases where variables are only important through interactions, we will call these “*interacting variables*”.

This paper proposes *WarmStart Conditional Boruta*, an extension of the Conditional Boruta method introduced by Rotari and Kulahci(2025). The Conditional Boruta method is a wrapper-based variable selection method that iteratively selects a subset of variables and evaluates their importance, to identify relevant variables. However, it can fail to identify *interacting variables*, as variables have a risk of being rejected in early iterations if they are not included in the same subset as their interacting partners, leading to a failure to observe the interaction effect.

The WarmStart Conditional Boruta method introduces a memorization loop in which variables accepted during one iteration are carried over as mandatory conditioning variables in subsequent iterations. This increases the chance of including *interacting variables* with their interacting partners within subsets, thereby improving the likelihood that *interacting variables* are correctly identified and retained.

The method is evaluated on a synthetic dataset with *interacting variables*, as well as on a real-world high-dimensional manufacturing dataset where conventional variable selection has been insufficient. Results demonstrate improved selection of *interacting variables* compared to Conditional Boruta and other baseline methods. The method is computationally intensive by nature, making it best suited for manufacturing settings where process understanding is crucial but experimentation is impossible due to time, cost or process constraints.

Keywords:

Variable Selection, Interactions, Manufacturing

AI: Machine Learning and Predictive Analytics / 61**A Two-Dimensional View of Classification Difficulty: Local Ambiguity and Global Separability****Authors:** Amalia Vanacore¹; Yariv N. Marmor²¹ *Department of Industrial Engineering University of Naples Federico II*² *BRAUDE - College of Engineering, Karmiel*

Classification performance is typically assessed empirically, yet its dependence on intrinsic data characteristics is not fully understood. In this study, we examine classification difficulty through two complementary dimensions: local class ambiguity measured in terms of instance hardness and global class separability captured by Silhouette score.

We generated synthetic datasets based on the `make_blobs` framework, where data complexity is controlled via cluster structure and dispersion. Instance hardness is quantified using 1-nearest-neighbor leave-one-out (LOOCV) error, corresponding to the N3 data complexity measure and capturing local overlap. In parallel, the Silhouette score—originating from cluster analysis—is used to quantify class separation and is interpreted as a proxy for global classification difficulty.

The individual and combined effects of local class ambiguity and global class separability on the accuracy of several classifiers have been investigated. Results show that local class ambiguity has a consistent negative effect across most classifiers, while global class separability has a positive but model-dependent effect. The interaction between local class ambiguity and global class separability is significant for several classifiers, indicating that performance depends on their interplay.

We identify distinct patterns: some classifiers (e.g., kNN, decision trees) are mainly sensitive to local class ambiguity, others (e.g., logistic regression) to global class separability, while hybrid models (e.g., SVM and boosting methods) depend on both.

Keywords:

Classification difficulty, Local class ambiguity, Class separability

AI: Machine Learning and Predictive Analytics / 31**Ranking Sets of Variables by Multivariate Association: a Non Parametric Combination Approach****Authors:** Rosa Arboretti¹; Elena Barzizza¹; Riccardo Ceccato¹; Giacomo Vezzosi¹¹ *University of Padova*

In many applied contexts, organizations need to evaluate and compare sets of explanatory variables in terms of their association with a Key Performance Indicator (KPI) of interest. This problem frequently arises in industrial and marketing applications, where companies seek to identify which groups of product characteristics or drivers are most strongly related to outcomes such as customer satisfaction or perceived performance.

This work proposes a methodological framework to rank clusters of explanatory variables according to the strength of their multivariate association with a response variable. For each cluster, the association between the outcome and the corresponding set of drivers is assessed through a permutation-based approach. Specifically, the method relies on the NonParametric Combination (NPC) methodology, which enables the joint evaluation of multiple partial tests measuring the association between the response and each variable within the cluster.

The partial permutation tests are combined within the NPC framework to obtain a global statistic that summarizes the overall strength of association for each cluster. These global measures are then used to derive a ranking of the clusters, thereby identifying which sets of drivers exhibit stronger relationships with the response variable.

The methodology provides a flexible and robust tool for applications such as customer satisfaction analysis and product development, supporting the identification of the most influential groups of drivers on the KPI of interest.

Keywords:

Association, Multivariate, Permutation, Ranking

AI: Machine Learning and Predictive Analytics / 98**A two-layer model for opinion spread in hypergraphs as a Markov chain and as the object of machine learning methods****Author:** András Zempléni¹**Co-authors:** Csegő Balázs Kolok ; Damján Tárkányi ; Villő Csiszár ; Ágnes Backhausz¹ *Eötvös Loránd University, Budapest*

In studies of opinion spread, peer pressure is often modeled through interactions of more than two individuals (higher-order interactions). We introduce a two-layer random hypergraph model, in which households and workplaces form the layers and hyperedges represent individual households and workplaces. Within this structure, individuals may react when their opinion is in the minority within their groups. The process evolves through stochastic steps: individuals can either change their opinion, or quit their workplace and join another one in which their opinion belongs to the majority. The model can be considered as a Markov chain, and its absorbing states are investigated. The effects of the parameters governing opinion change and workplace switching are described via computer simulations. The performance of different statistical and machine learning methods — namely linear regression, XGBoost, convolutional neural network and LSTM - in estimating these parameters is investigated. It turns out that in most of the cases, especially for shorter observation periods LSTM is the best, while in many cases XGBoost is also a strong contender. It is also an important observation that the information required for accurate estimation depends on the strength of the peer pressure effect.

Keywords:

hypergraph, opinion spread, parameter estimation

AI: Machine Learning and Predictive Analytics / 115

Semi-Supervised PARAFAC2 Decomposition for Computational Phenotyping Using Electronic Health Records

Author: Elif Konyar¹**Co-author:** Mostafa Reisi Gahrooei²¹ *Georgia Tech*² *University of Florida*

Computational phenotyping uses data mining methods to extract clusters of clinical descriptors, known as phenotypes, from electronic health records (EHR). Tensor factorization methods are very effective in extracting meaningful patterns and have become popular in computational phenotyping. Nevertheless, these techniques mainly focus on regular tensors and are used in a fully unsupervised manner. EHR data is often represented by irregular tensors (due to varying hospital visits) and contains some label information that is useful in extracting meaningful phenotypes. While techniques to decompose irregular tensors in an unsupervised manner have been developed, methods to integrate label information in such decomposition are lacking. In this work, we propose a semi-supervised PARAFAC2 decomposition model to extract meaningful patterns from irregular EHR data by incorporating label information available for a subset of instances. Experiments on a synthetic data set and a case study based on MIMIC-IV data show the superiority of our approach over the benchmarks. This novel computational phenotyping method can potentially facilitate medical decision-making in many healthcare applications.

Keywords:

Computational phenotyping; Tensor Decomposition

AI: Machine Learning and Predictive Analytics / 110**Sustainable specialty chemical design through multi-objective latent-variable model inversion****Authors:** Marco Brendolan¹; Michele Olivi²; Massimiliano Barolo¹; Pierantonio Facco¹¹ *University of Padova*² *Sirca S.p.A. –Resins and Coatings*

Product formulation in the specialty chemicals industry requires balancing product quality, cost, and environmental impact. This work proposes a data-driven framework for sustainable formulation design based on latent-variable model inversion combined with multi-objective optimization. Partial Least Squares (PLS) models are built to relate raw-material properties and compositions to product quality within a reduced latent space. The inversion of the PLS model is then formulated as a multi-objective mixed-integer nonlinear optimization problem. The approach simultaneously optimizes three competing objectives: (i) achieving a desired product quality target; (ii) minimizing formulation cost; and (iii) maximizing environmental sustainability. The economic objective minimizes formulations unit cost combining polymer and component costs weighted by mass fractions, using proxy costs for confidentiality reasons. Sustainability is quantified through an index derived from hazard classifications using a decision-tree methodology based on the Globally Harmonised System classification and labelling of chemicals.

The methodology is applied to the formulation of unsaturated polyester resins, where multiple raw materials (polymers, reactants, and additives) must be selected and blended under compositional and statistical feasibility constraints.

The optimization generates a set of feasible non-dominated solutions, enabling the exploration of trade-offs among objectives. Results highlight that: higher mechanical performance is associated with increased cost and reduced sustainability, while more sustainable formulations tend to be more expensive and may slightly reduce product quality. Furthermore, the proposed methodology supports informed decision-making by quantifying these trade-offs and identifying alternative formulations that maintain good product performance while improving cost and environmental impact.

Keywords:

AI, product formulation, PLS model inversion, resins

AI: Machine Learning and Predictive Analytics / 117**Decomposition of Serially Dependent Data into Static and Dynamic Latent Structures****Author:** Moritz Bauchrowitz¹**Co-authors:** Pau Cabaneros²; Peter Westergaard Jakobsen²; Tobias Eifler³; Murat Kulahci⁴¹ *PhD Student*² *Novo Nordisk A/S*³ *Technical University of Denmark*⁴ *DTU*

In many industrial settings, data is collected over time causing a serial dependence among the observations. Many chemometric methods, such as Principal Component Analysis (PCA), function under the assumption of time independence. This assumption is violated for most industrial data, creating challenges for both descriptive modelling as well as fault detection.

Dynamic PCA (DPCA), which employs lagged augmentations of the data matrix to capture serial relationships, has been proposed for multivariate statistical process control (MSPC). An integral step in DPCA is the definition of the lag structure, for which multiple algorithms have been proposed. However, none of the lag selection algorithms inherently limits the number of lags, so the dimensionality of the resulting DPCA model may be inflated. Furthermore, DPCA leads to latent directions that mix static and dynamic relationships, which may complicate interpretation.

In this presentation, we propose formulating vector autoregressive models as latent directions to achieve a latent space that consists of both static and dynamic components but strictly separates them. This simplifies interpretation and limits the maximum number of latent directions. We apply the methodology to data generated from the Tennessee Eastman Process simulator.

Keywords:

Autocorrelation, Time Series, Multivariate Data

Award session / 20**The Irreducible Error: What Statistics and Management Have in Common****Author:** Sophie Carr¹¹ *Bays Consulting*

The truth is that some error, no matter how hard we try, simply can't be modelled away. That irreducible error that stubbornly remains, no matter how time we have spent selecting predictors, or agonising over parameter tuning. Accepting that there will always be some randomness in statistics goes a long way to helping manage a technical team.

In this talk, Sophie will draw on her own experience and lessons learned from mentors, friends and colleagues to argue the most important lesson in management (and statistics) is to embrace uncertainty and act wisely in its presence. In statistics, we learn to distinguish noise from signal, not to eliminate it. When managing a team, the aim should be to create conditions where people can confidently deliver their best work.

To achieve this, how can managers ensure that they resist the temptation to overfit measurable performance for what matters most (trust, motivation, and safety)? Perhaps more crucially, how can a manager have the honesty to say when their model is wrong and they need to refine their approach? Throughout, Sophie will ask and challenge if the habits and traits that make a great statistician, are also those that make a great statistician.

As a Bayesian statistician, Sophie would like you to know that no p-values were harmed in the preparation of her talk and significance is not guaranteed.

Keywords:

management, uncertainty, decision

Award session / 62

Statistical Process Monitoring Based on Functional Data Analysis

Author: Fabio Centofanti¹¹ *KU Leuven*

In modern industrial settings, advanced acquisition systems allow for the collection of data in the form of profiles, that is, functional relationships linking responses to explanatory variables. In this context, statistical process monitoring (SPM) aims to assess the stability of profiles over time in order to detect unexpected behavior. This talk focuses on SPM methods that model profiles as functional data, that is, smooth functions defined over a continuous domain, and apply functional data analysis (FDA) tools to address limitations of traditional monitoring techniques. A reference framework for monitoring multivariate functional data is first presented. The talk then offers a focused survey of several recent FDA-based profile monitoring methods that extend this framework to address common challenges encountered in real-world applications. These include approaches that integrate additional functional covariates to enhance detection power, a robust method designed to accommodate outlying observations, a real-time monitoring technique for partially observed profiles, and adaptive strategies that target the characteristics of the out-of-control distribution. These methods are implemented in the R package *funcharts*, available on CRAN.

This presentation is based on the article

Centofanti, F. (2026). Statistical Process Monitoring Based on Functional Data Analysis. *Wiley Interdisciplinary Reviews: Computational Statistics*, 18(1), e70057. <https://doi.org/10.1002/wics.70057>

Keywords:

Anomaly detection; Control charts; Profile Monitoring

Box medal / 78

Enhancing Definitive Screening Designs with Multilevel Categorical Factors

Author: Bradley Jones¹¹ *EFFEX*

When first introduced, Definitive Screening Designs (DSDs) supported continuous factors only. Later developments added two-level categorical factors and flexible blocking of DSDs to the mix. However, until now, support for categorical factors at three or more levels has been elusive. This talk provides a general framework for doing so, along with examples of enhanced DSDs having multiple-level factors.

Keywords:

design

Bridging AI and Statistics / 82**Applied Statistics in the Era of Artificial Intelligence: A focus on AI Assurance**

Authors: Jie Min¹; Xinyi Song^{None}; Simin Zheng^{None}; Caleb King^{None}; Xinwei Deng^{None}; Yili Hong^{None}

¹ *University of South Florida*

The advent of artificial intelligence (AI) technologies has significantly changed many domains, including applied statistics. This talk explores the evolving role of applied statistics in the AI era, drawing from our experiences in engineering statistics, focusing on how statistics can be employed to study the properties of AI models and enhance AI systems, especially in AI assurance.

We briefly review the role of applied statistics in traditional areas, then focus on the relationship between statistics and AI, with particular emphasis on AI assurance. We explore how applied statistics contributes to key aspects of AI reliability by introducing the “SMART” statistical framework for AI reliability research. This framework consists of five components, including the structure of the system, metrics of reliability, analysis of failure causes, reliability assessment, and test planning, with the overarching goal of ensuring that an AI system can perform its designed functionality for the intended period. Several concrete examples from our recent research are presented, including the reliability analysis of autonomous vehicles, out-of-distribution detection, and AI test planning. Together, these examples demonstrate that statistical principles provide a rigorous foundation for evaluating and improving AI systems. The talk concludes with reflections on the future role of statisticians in an increasingly AI-driven world.

Keywords:

AI reliability; AI robustness

Bridging AI and Statistics / 89**Novel AI Solutions for Process Monitoring and Optimization**

Author: Bianca Maria Colosimo¹

¹ *Politecnico di Milano*

The landscape of industrial data mining is rapidly evolving, opening new opportunities to transform production systems from merely instrumented to truly intelligent. This talk explores innovative AI-driven solutions that go beyond traditional data-driven approaches, encompassing generative AI—including GANs and diffusion models—and transfer learning, to support a new generation of smart manufacturing systems. Through selected examples drawn from additive manufacturing and bioprinting, we discuss how these methodologies can enable adaptive process monitoring and optimization, moving production systems toward greater autonomy.

Keywords:

AI, GEN AI, Statistical process monitoring, adaptive optimization, Additive manufacturing, bioprinting

Bridging AI and Statistics / 116**Discriminative Open Set Active Learning (DOSAL)****Authors:** Bart De Ketelaere¹; Laura Carolina Martinez Esmeral²; Wouter Saeys^{None}¹ *Catholic University of Leuven*² *KU Leuven*

Active learning is commonly framed in the artificial intelligence community as a strategy to reduce labeling costs by selectively querying informative samples for model improvement, yet its conceptual roots are closely aligned with classical ideas from statistical design of experiments, particularly sequential design and design augmentation. In this work, we position active learning as a modern instantiation of experimental design under evolving data representations, where the design space itself is learned rather than predefined. We introduce DOSAL, a discriminative open set active learning framework that bridges statistical design principles with deep learning-based representation learning. Starting from a convolutional neural network (CNN) trained on known classes, learned latent representations are treated as empirical feature spaces in which sample utility can be quantified. DOSAL combines novelty, measured through distances in feature space, and surprisal, derived from discriminative probability estimates, into a unified acquisition strategy that parallels multi-criterion optimality in experimental design. This strategy is used to iteratively augment the training set whilst accounting for open set conditions, where previously unseen classes may be present. In this way, DOSAL provides an explicit and operational connection between classical design-based sample selection and modern neural representations, enabling discovery-oriented learning driven by design-inspired criteria rather than purely heuristic querying. The working principles of DOSAL will be outlined during the presentation, and its empirical behavior relative to baseline methods are demonstrated on a benchmark dataset.

Keywords:

active learning, design of experiments, AI

Data Analytics and Data Science: Case Studies / 143**Health Insurance Fraud Detection using Claim-Based Profiling****Author:** Sotiris Bersimis¹**Co-authors:** Charalampos-Panagiotis Michelakis²; Panagiotis Biris³; Polychronis Economou³¹ *University of Piraeus, Greece*² *University of Piraeus, Department of Business Administration*³ *University of Patras*

Medical claim expenses are inherently compositional, as fraud-relevant patterns often emerge from the relative allocation of costs across categories rather than from total expenditure alone. We propose a claim-level fraud screening framework based on compositional profiling, using the Aitchison distance to compare new claims with a historical reference distribution. Statistical significance is assessed via bootstrap resampling. Simulation results under a multivariate normal setting demonstrate effective false positive control, increased sensitivity to meaningful profile deviations, and robustness to scale-only changes. The framework offers an efficient, interpretable, and practically relevant approach to anomaly detection in healthcare expenditure data.

Keywords:

Fraud detection; Healthcare claims; Medical expenses; Claim-level profiling; Compositional data analysis; Aitchison distance; Anomaly detection; Bootstrap;

Data Analytics and Data Science: Case Studies / 38**Simulating a key performance index for a family of future products****Author:** Carlo Leardi¹¹ *tetra pak packaging solutions*

A complex system is currently under validation as implemented in its initial instantiation. The technical bet regards more than doubling a key performance at parity of the other ones. The preliminary estimation has been performed by simulation in the concept's exploration phase by risk reduction by Fault Tree Analysis. The current studies are devoted to allowing the estimation of the probability of success of a full family of products by combining the system structure and flows by Design Structure Matrices with Bayesian propagation. The proposal for a presentation regards the rationales, the opportunities and the consequences of the methodological estimation evolution. The potentialities, the limits of the potential alternatives are proposed for sharing and discussion.

Keywords:

KPI modelling FTA DSM Bayesian propagation

Data Analytics and Data Science: Case Studies / 97**Dirty laundry –A data quality journey from ENBIS Active session to nationwide system redesign for industrial textile management****Authors:** Morten Bormann Nielsen¹; Robert Heck¹¹ *Danish Technological Institute*

At ENBIS-24 in Leuven, we brought an emerging challenge to the ENBIS Active Session: a large industrial laundry operator managing millions of textile items across multiple sites wanted to understand and extend textile lifespans as part of a circular economy strategy using existing operational data on textile discarding events. The key advice we received —to not trust the data from the outset —proved transformative.

Rather than proceeding directly to lifetime modeling, our project team invested substantial effort in exploratory data visualization and knowledge building through interviews with operational staff. This exploration revealed that the registration system fundamentally conflated two distinct decisions: whether to invoice the customer and why the textile was actually discarded. Staff routinely selected codes to achieve correct billing rather than to document quality defects, rendering the data unusable for lifetime analysis.

This discovery led to the co-development of a new process to cleanly separate billing decisions from quality assessments. The workflow was iteratively refined through workshops, pilot testing, and continuous dialogue with operational staff across the organization, before being deployed nationwide across all sites.

This talk traces the full journey: from seeking advice at the ENBIS Active Session, through the project that uncovered and solved a fundamental data quality problem, to a validated nationwide system —demonstrating how statistical thinking and early skepticism toward data quality can reshape an entire organizational data infrastructure.

Keywords:

Data Quality, Statistical Engineering, Exploratory Data Analysis

Data Analytics and Data Science: Case Studies / 72**An IoT-inspired concept for data-driven service provision and resource planning in the tourism sector, using a campsite as an example****Author:** Eva Scheideler¹**Co-author:** Andrea Ahlemeyer-Stubbe²¹ *Technische Hochschule Ostwestfalen-Lippe*² *Ahlemeyer-Stubbe*

The digitalisation of tourism facilities means that these facilities have access to a wide range of IoT-like data sources. This article presents a conceptual approach that describes how such heterogeneous data streams can be used to systematically improve service offerings, resource planning and operational decisions through targeted short-term, medium-term and long-term forecasts. The analysis is based on the example of a medium-sized campsite, which provides data on the electricity consumption of individual pitches, water consumption, access information for various facilities, occupancy figures, and guests' movement and usage profiles.

The key idea is that, taken together, this data creates a digital representation of the business. This representation makes it possible to identify patterns in infrastructure utilisation, guest behaviour and the demand placed on technical systems. Linking this information creates a conceptual framework that supports predictions for three key areas of application: Firstly, service offerings can be dynamically adapted to expected usage patterns, for example through flexible opening hours or the targeted allocation of resources. Secondly, areas of high wear and tear can be identified from occupancy and usage data, enabling proactive planning of maintenance measures. Thirdly, the analysis of movement and access data forms the basis for demand-driven staff deployment planning in high-traffic areas.

The focus of the presentation lies in the area of data integration and the establishment of the data structures required for the respective analyses; in this context, knowledge of data generation and process rules within the business environment is of paramount importance.

Keywords:

Heterogeneous data streams, Data integration, domain knowledge

Data Analytics and Data Science: Case Studies / 112**Multivariate Six Sigma for the Optimization of the Meat Roasting Process in the Ready-to-Eat Food Industry****Authors:** Alberto Ferrer-Hermenegildo¹; Joan Borràs-Ferrís²**Co-authors:** José Manuel Prats-Montalbán³; Borja Galdón-Navarro¹; Alberto J. Ferrer-Riquelme³¹ *Kensight Solutions S.L.*² *Kensight Solutions, S.L.*³ *Universitat Politècnica de València, Spain*

This talk presents a Six Sigma project developed in a ready-to-eat food company aimed at optimizing a meat roasting process while balancing food safety, product appearance, juiciness, and production yield.

Following the DMAIC methodology, historical data analysis, Measurement System Analysis (Gage R&R), and Root Cause Analysis tools were initially applied to understand process variability and identify potential sources of performance loss. A Design of Experiments was subsequently planned and executed based on expert knowledge and process understanding.

In addition to the experimental factors, several process and contextual covariates were collected during experimentation. To overcome the limitations of traditional univariate approaches in complex industrial environments, latent variable-based multivariate techniques such as Principal Component Analysis were incorporated into the Six Sigma statistical toolkit. These techniques allowed the evaluation of hidden relationships and potential confounding structures between process covariates and experimental effects before DOE interpretation.

The proposed multivariate approach provided a more reliable understanding of process behavior and supported the identification and implementation of improved operating conditions. This work reinforces how the integration of latent variable methods into the DMAIC methodology, the so-called Multivariate Six Sigma, leads to a powerful process improvement framework for Industry 4.0 environments.

References:

Ferrer, A. (2021). Multivariate six sigma: A key improvement strategy in industry 4.0. *Quality Engineering*, 33(4), 758–763. <https://doi.org/10.1080/08982112.2021.1957481>

García-Carrión, S., Pozueta, L., & Ferrer, A. (2026). Enhancing Six Sigma with latent variable models: An industrial application of multivariate Six Sigma in the automotive sector. *Quality Engineering*, 1–15. <https://doi.org/10.1080/08982112.2026.2626827>

Keywords:

Multivariate Six Sigma, DMAIC, PCA

Data Analytics and Data Science: Case Studies / 120**Risk analysis of forest fires in Sicily, based on 2010-2023 observation period****Author:** Stefano Barone¹**Co-authors:** Antonino Paladino ; Santo Orlando ¹¹ *University of Palermo***Introduction.**

Forest fires are complex phenomena causing significant damage to the environment and human health, habitat destruction, soil erosion, greenhouse gas emissions, and biodiversity loss. They are increasing globally, with extreme events becoming more frequent and destructive. Understanding their root causes and influencing factors is crucial.

Methods.

This work focuses on analyzing data of forest fires that occurred in the period 2010-2023 in Sicily, a big island with special orographic characteristics and substantial agricultural and forestry-pastoral activities. The methods concern a careful extraction of data by using QGIS software and official databases, and their appropriate statistical analysis.

Results.

A suitable definition of forest fire risk is formulated, and a risk ranking of Municipalities is obtained. Risk factors and their significance on risk are determined via multiple regression analysis with Box-Cox transformed dependent variable.

Conclusions.

The work shows an optimal balancing between ecological perspective and operational risk management. Forest fire data collection empowerment is highlighted, such as fire-starting location and total damage caused by each fire event. The study allows optimally distributing the Regional budget for forest fire prevention among the 390 Sicilian municipalities.

Keywords:

Environmental statistics; Forest fire risk management; Spatial statistics

Data Analytics and Data Science: Case Studies / 32**Nonparametric tests for multivariate stability data**

Authors: Riccardo Ceccato¹; Marta Disegna¹; Alessandro Fanesi¹; Daniele Pennisi²

¹ *University of Padova*

² *Stat4Value S.r.l.*

Stability studies are commonly conducted to evaluate how product characteristics evolve during storage. In many industrial applications, several quality attributes are measured repeatedly over time for multiple products, generating multivariate longitudinal datasets. A key objective in these studies is to compare products in terms of their stability and identify those exhibiting more stable behavior during storage.

This work investigates nonparametric permutation-based approaches for the analysis of multivariate stability data collected at different time points. To properly account for the longitudinal structure of the data, a restricted permutation strategy is adopted. This approach preserves the temporal dependence of the measurements while allowing valid inference on differences between products.

Partial permutation tests are performed for each quality attribute, and the resulting p-values are combined using the Nonparametric Combination (NPC) methodology to obtain a global assessment of product differences across the multivariate responses. The proposed framework enables the construction of a ranking of products according to their stability while avoiding distributional assumptions and remaining suitable for small sample sizes or heterogeneous variables.

The methodology provides a flexible and robust tool for the analysis of multivariate stability studies and supports product comparison and decision-making in industrial applications.

Keywords:

Stability, Nonparametric statistics, Permutation tests

Data Analytics and Data Science: Case Studies / 42**Clustering for Detecting Careless Respondents: A Case Study**

Authors: Elena Barzizza¹; Marta Disegna¹; Sara Naziri²; Silvia Villanova¹

¹ *University of Padova*

² *Stat4Value*

The presence of careless respondents represents a well-known threat to the quality of survey data. Respondents who provide inattentive or random answers can distort statistical analyses, reduce measurement reliability, and bias substantive conclusions. A variety of indicators have been proposed in the literature to detect such respondents, including response pattern measures such as longstring indices, intra-individual response variability (IRV), and multivariate distance measures such as Mahalanobis distance. However, these indicators are often applied individually and require arbitrary thresholds, which may limit their effectiveness in complex datasets.

This work proposes a data-driven framework for identifying careless respondents that integrates several classical diagnostic indicators with clustering algorithms. First, multiple indicators capturing different aspects of response behavior are computed for each respondent, including measures of response consistency, variability, and multivariate outlyingness. These indicators are then jointly analyzed using unsupervised clustering techniques to identify groups of respondents with similar response patterns. In this way, respondents exhibiting atypical combinations of diagnostic indicators can be detected without relying on predefined cutoffs.

The proposed approach allows the simultaneous exploitation of complementary information contained in multiple quality indicators and provides a flexible and scalable tool for detecting potentially careless respondents. An empirical study based on survey data illustrates how the integration of traditional indicators with clustering methods can improve the identification of problematic response patterns and support more reliable data cleaning procedures in questionnaire-based research.

Keywords:

careless responding, machine learning, smart data

Data Analytics and Data Science: Case Studies / 45**Preference mapping for product optimization: a case study**

Authors: Rosa Arboretti¹; Taha Asghari^{None}; Alessandro Fanesi²; Silvia Villanova²

¹ *Department of Civil, Environmental and Architectural Engineering, University of Padova, Via Marzolo, 9, Padova, 35131, Italy*

² *Department of Management and Engineering, University of Padova, Stradella San Nicola, 3, Vicenza, 36100, Italy*

Understanding how technical product characteristics translate into consumer perception remains a key challenge in product development. This study presents a case study in which preference mapping techniques are used to explore the relationship between laboratory-based technical measurements and consumer evaluations.

A set of products was characterized through a series of objective instrumental measurements collected under controlled laboratory conditions. In parallel, consumer data were collected to assess perceived performance and overall product preference. Preference mapping was then applied to model the relationship between the technical variables and the consumer responses, enabling the identification of the technical drivers associated with higher consumer liking.

The resulting maps provide a visual and analytical representation of how products are positioned in the space defined by instrumental attributes and consumer preference. The analysis highlights which technical characteristics are most strongly associated with favorable consumer evaluations and illustrates how preference mapping can be used to bridge the gap between laboratory measurements and consumer perception.

The case study demonstrates how combining technical data and consumer research through preference mapping can support product optimization and guide product development toward configurations that better meet consumer expectations.

Keywords:

Preference mapping, consumer perception, Product development

Data science for climate and environment / 80

Model selection for extremal dependence structures using deep learning: Application to environmental data

Authors: Manaf Ahmed¹; Pierre Ribereau²; Véronique Maume-Deschamps³

¹ *University of Mosul*

² *Université Claude Bernard Lyon 1*

³ *Institut Camille Jordan, Université Claude Bernard Lyon 1*

Although the CLIC-based model selection approach is widely used to identify spatial extreme models, the complexity of the associated statistical inference limits the reliability of this criterion. In addition, the strong spatial dependence in small or moderate regions may lead to substantial overlap among the spatial extremes models. This potential overlap increases the risk of model misidentification. In this paper, we exploit the ability of Convolutional Neural Networks (CNNs) to extract spatial patterns in order to develop a CNN-based model selection framework. The proposed approach evaluates how well the dependence structure observed in the data matches the dependence patterns implied by competing models. Two identification strategies are considered. In the first strategy, both the max-stable model and its associated covariance function are identified simultaneously by a single CNN in a one-step procedure. In the second strategy, model identification is performed hierarchically. First, a CNN identifies the class of max-stable model, and then additional CNNs are trained for each model to determine the corresponding covariance function. The performance of the two strategies is evaluated through an extensive simulation study designed to reproduce the spatial dependence structure of 2-m air temperature data over Iraq, where strong dependence and model overlap are observed. The results demonstrate that the proposed CNN-based approach provides an effective alternative for model selection in spatial extremes.

Keywords:

Statistics, Spatial statistics, Extreme value theory

Data science for climate and environment / 81

Multivariate Discrete Generalized Pareto distributions: Theory, likelihood-free inference, and applications to drought risk assessment

Author: Samira Aka¹

¹ *Square Management*

Understanding extreme environmental phenomena is crucial for risk management in a changing climate. In particular, dry spells, defined as consecutive days without precipitation, play a key role in drought dynamics, with direct impacts on agriculture, water resources, and insurance systems. Dry spell lengths are inherently discrete and often exhibit complex dependence structures across locations. Classical extreme value models, designed for continuous data, are therefore not well suited to such settings.

We introduce the multivariate discrete generalized Pareto distribution (MDGPD), a probabilistic framework tailored to model multivariate exceedances with integer-valued support. This model extends extreme value theory to discrete settings while preserving key tail properties. Inference for MDGPD models is challenging due to the intractability of the likelihood. To address this, we develop likelihood-free inference procedures combining neural Bayes estimation with Wasserstein-based discrepancies.

We illustrate the methodology on multivariate dry spell lengths derived from daily precipitation records in Switzerland, highlighting its relevance for environmental risk modeling.

Keywords:

Extreme Value Theory, Discrete extremes, Simulation-based Inference

Data science for climate and environment / 26

Concepts and methods for better predictions in climate and environmental science**Author:** Georgia Papacharalampous¹**Co-author:** Marco Borga¹¹ *University of Padova*

Skillful predictions in climate and environmental science are essential for planning operations, assessing risks, guiding adaptation strategies, and building resilience. This talk synthesizes key concepts and methods for enhancing predictive performance in these domains, with a particular emphasis on predictive uncertainty estimation, extreme event prediction, and the role of big datasets and large-scale benchmarking in this context. Along the way, we discuss core strengths, limitations and open challenges across the concepts and methods in predictive settings. Rather than diving deeply into any single method, we focus on their interconnections and practical trade-offs. We explore both established and envisioned applications in areas such as hydrological forecasting and bias correction of climate forecasts and satellite data. To foster the discussion, the talk concludes by reflecting on the transferability of the analyzed concepts and methods to fields such as energy demand forecasting, and on how such a transfer can be effectively achieved.

Keywords:

Big data, extreme events, predictive uncertainty

Design and modeling of computer experiments / 29

Random Designs for Quantisation in High Dimension**Author:** Luc Pronzato¹**Co-author:** Anatoly Zhigljavsky²¹ *CNRS*² *Cardiff University*

Incremental design for computer experiments traditionally relies on space-filling or uniformity criteria, but these become impractical in high dimensions. The greedy minimisation of the L_s -mean quantisation error (or distortion) offers a valuable alternative, though it is also computationally intractable for large d .

This talk focuses on random designs composed of i.i.d. points sampled from a specific distribution. For large d , observing the behaviour predicted by Zador's theorem requires an impractically large sample size n , growing super-exponentially with d . We address this challenge by analysing the quantisation problem for spherically symmetric distributions. Our results show that for moderate n random quantisers uniformly distributed on a sphere of suitable radius R achieve exceptional performance. The expected distortion, computed exactly via a triple integral, allows numerical optimisation of R . Using extreme-value theory, we also derive approximations for R , revealing that, when n grows with d and $d \rightarrow \infty$, R may converge to zero or approach a limiting value R_∞ independent of s , depending on the growth rate of n .

While spherically symmetric distributions may seem restrictive, they provide a starting point for further applications, such as quantising the uniform measure on the hypercube $X = [-1, 1]^d$. For large d , the uniform measure on X can be approximated by a spherically symmetric distribution, and one can consider random quantisers distributed according to a product measure, which can itself be approximated by a spherically symmetric distribution. Preliminary results suggest that quantisers distributed on the vertices of a smaller hypercube exhibit promising performance.

Keywords:

space-filling design, quantisation, distortion, spherically symmetric distribution, Zador's theorem, extreme-value theory

Design and modeling of computer experiments / 16

Active Subspace Embedding for Parsimonious Gaussian Process Surrogates**Author:** Amandine MARREL¹**Co-author:** Bertrand Iooss²¹ CEA² EDF R&D

Uncertainty quantification for complex physical systems often relies on computationally expensive numerical simulators. When execution times limit the number of feasible runs, **surrogate modeling** becomes essential for tasks such as sensitivity analysis, design optimization, and safety assessment.

Gaussian process regression (GPR) is a leading

surrogate due to its uncertainty quantification capabilities and flexibility. However, high input dimensionality—common in industrial applications—poses significant challenges: increased computational cost, deteriorated prediction accuracy, and numerical instability in covariance estimation.

We propose a dimension reduction methodology combining **statistical screening and active subspace (AS) identification to enable efficient and parsimonious GPR metamodeling in large-dimensional contexts**. The approach proceeds in three stages: (1) initial variable screening via independence tests based on Hilbert-Schmidt Independence Criterion (HSIC) reduces dimensionality to a tractable scale (~20 variables), (2) a medium-dimensional GPR is fitted to estimate gradient information, from which the gradient covariance matrix and its active subspace decomposition are computed, and (3) a final parsimonious low-dimensional GPR is constructed on the identified active directions. Key methodological contributions include optimal dimension selection via cross-validation adapted to the given-data context, and covariance modeling recommendations tailored to each stage: high-regularity kernels for stable gradient estimation, and flexible kernel choices for final surrogate construction.

We illustrate the methodology on a **nuclear safety application**: peak cladding temperature prediction during intermediate-break loss-of-coolant accidents (IB-LOCA) using the a thermal-hydraulic code. Starting from several dozens of uncertain input parameters, our approach achieves a four-fold dimension reduction yielding a parsimonious surrogate while preserving accuracy and reliable predictive intervals.

Keywords:

Uncertainty quantification, Computer experiments, Surrogate model, Gaussian process regression, Active subspaces.

Design and modeling of computer experiments / 74**Multivariate sensitivity analysis for risk analysis with spatial outputs: a comparison exercise****Author:** Jérémy ROHMER¹**Co-author:** Corinna Perchtold ¹¹ BRGM

Maps play a key role in many applications to facilitate decision-making for risk analysis, such as in assessing natural hazards, soil pollution, water quality, and so on. Performing global sensitivity analysis in a spatial context can benefit from multiple approaches adapted to multivariate outputs, namely: 1. the combination of variance-based sensitivity indices and functional principal component analysis (PCA); 2. the use of the Hilbert-Schmidt Independence Criterion (HSIC) with kernels adapted to functional outputs; 3. a recently-developed moment-independent method grounded in the theory of optimal transport (OT). In this communication, we aim to compare these different approaches by considering three types of criteria. Firstly, we analyse their practical implementation, that is, their ability to handle complex types of output data and to estimate measures of importance directly from a dataset containing a limited number of samples. Secondly, we analyse their interpretability in relation to the objective of risk analysis. Finally, we examine their robustness with respect to their parameterisation by analysing the stability of the results in the face of different modelling choices, more specifically the amount of variance retained in PCA, the definition of the kernel for the HSIC method, and the estimator used to calculate the Wasserstein distance for the OT method. This comparative exercise is based on a simple synthetic function, namely the Campbell2D function, as well as on real-world environmental studies for assessing soil pollution and groundwater quality.

Keywords:

Uncertainty; High-dimensional output; Importance measure

Design of Experiments / 129

Quality by Design: Data-Efficient Active Learning Approaches**Author:** Marco P. Seabra dos Reis¹**Co-author:** Daniel Alexandre Vidinha Batista²¹ *Department of Chemical Engineering, University of Coimbra*² *University of Coimbra*

As referred in last year's abstract by the same authors, "the landscape of the pharmaceutical industry is evolving". Such a process continues, with more efforts being devoted to developing data-efficient methodologies for exploring operational spaces of increasing dimensionality that may be composed of continuous, categorical, and mixture factors. From what was (and still is) a science-based discipline, more awareness exists nowadays of the opportunities arising from exploring data-driven methodologies to conduct various key activities, namely to find the range of "optimal" conditions to operate a process or the best formulation for a pharmaceutical product. In this regard, two classes of active learning (AL) approaches are increasingly regarded as the most competitive: Statistical Design of Experiments (DOE) and Bayesian Optimisation (BO). Each one has emerged from a different scientific community –applied statistics and machine learning, respectively–, and has been conquering the confidence of supporters who, at the same time, build a diminished vision of the other class. DOE supporters tend to find BO lacking a sound theoretical structure and optimality guarantees, whereas BO proponents view the DOE approach as outdated and overly constrained by assumptions.

Ultimately, "the proof of the pudding is in the eating". Therefore, in this work, we present results on the use of both classes of methods in the sequential search for the best conditions, under a limited budget of experiments. Different systems are considered, and Monte Carlo simulations conducted to gather robust information on scenarios where one methodology is expected to outperform the other with high confidence. Drawing such an operational map of AL methods would, we argue, be a more useful outcome for practitioners than the entrenched defense of each class.

Acknowledgements

This work was funded by the CInTech project - Technological Hub for Innovation, Translation and Industrialisation of Complex Injectable Drugs -, under reference no. C644865576-00000005, co-financed by Componente C5 - Capitalização e Inovação Empresarial integrada na Dimensão Resiliência do Plano de Recuperação e Resiliência (PRR), through the NextGenerationEU fund. Authors also acknowledge support from CERES –Chemical Engineering and Renewable Resources for Sustainability Research Center, funded by FCT –Fundação para a Ciência e Tecnologia (UID/00102/2025), PRR –Recovery and Resilience Program, of the Portuguese Republic (UID/PRR/00102/2025), Equipar+2 (UID/PRR2/00102/2025).

Keywords:

Quality by Design; Sequential Design of Experiments; Bayesian Optimization; Active Learning

Design of Experiments / 44

Fighting Antimicrobial Resistance, Functional Data Analysis, and the Future of Multifactor Experiments**Author:** Phil Kay¹¹ *JMP Statistical Discovery*

Antibiotic resistance is one of the greatest health threats facing the world. The University of Oxford are pioneering a novel approach by targeting the RecBCD enzyme —a key regulator of DNA repair in bacteria. Central to this was the development of a robust, high-throughput biochemical assay, which meant navigating a complex 11-factor space.

Initially limited by manual pipetting and single-variable experimentation, we adopted design and analysis of multidimensional experiments and automated liquid handling hardware. This enabled execution of hundreds of experimental runs at low volumes, dramatically increasing throughput and insight, and transforming assay development from a weeks-long process into a single day of experimentation.

Functional Data Analysis was instrumental in modeling enzyme kinetics, allowing us to visualize and interpret complex time-course data with ease. This enabled rapid optimization of assay conditions and a deeper understanding of the biochemical dynamics at play.

This project exemplifies the cutting edge of biological experimentation, where automation, high-dimensional design, and advanced analytics converge to accelerate discovery. The synergy between DOE and lab automation not only enabled us to develop a scalable assay for drug screening but also points to a future where biological research is faster, more reproducible, and more insightful than ever before.

Keywords:

functional data analysis, design of experiments, laboratory automation

Design of Experiments / 136

Optimizing multi-arm clinical trials for personalized medicine using a genetic algorithm**Author:** Alan Vazquez¹**Co-authors:** Karla Cervantes¹; Weng-Kee Wong²¹ *Tecnologico de Monterrey*² *University of California, Los Angeles*

Personalized medicine aims to improve treatment decisions using patient-specific covariates. In diseases with heterogeneous treatment responses, estimating treatment-covariate interactions is essential for identifying effective therapies across patient subgroups. Multi-arm clinical trials provide an efficient framework for evaluating several treatments simultaneously; however, the design problem becomes increasingly challenging as the numbers of treatments and covariates increase. In this work, we propose a statistical criterion for evaluating multi-arm trial designs based on interaction estimation across all potential subject covariates, including both continuous and categorical variables. To address the resulting combinatorial optimization problem, we develop a genetic algorithm that efficiently searches for statistically efficient treatment assignments. The proposed approach generates efficient designs by minimizing the maximum subject-covariate variance across treatment groups, thereby reducing uncertainty in treatment assignment under the individualized treatment rule considered. Extensive numerical experiments, including a real clinical trial application, demonstrate that the proposed algorithm consistently outperforms existing methods, yielding more efficient multi-arm trial designs. The proposed methodology provides a flexible and scalable framework for designing multi-arm clinical trials in personalized medicine.

Keywords:

Personalized medicine, Multi-arm clinical trials, Genetic algorithm, Individualized treatment rule

Design of Experiments / 86**Bayesian Optimisation under system drift****Author:** Stefanie Feiler¹¹ *FHNW School of Life Sciences*

The classical approach to DoE, as shaped by Fisher, now reaches back almost 100 years. Bayesian optimisation, or “active learning”, is now often presented as a more modern alternative. As an iterative method, it selects each new experimental run based on the information currently available. This means that randomisation, which is one of the central aspects in “classical” DoE, is inherently impossible. This becomes problematic in the presence of a system drift, i.e. a trend in the measurements over time.

Using simulated data with varying degrees of system drift (and different noise levels), we assess how the two approaches behave under such non-ideal conditions. Can these effects safely be ignored or do they present a genuine problem?

Keywords:

Bayesian Optimisation, Design of Experiments, System Drift

Design of Experiments / 87**Powerful Foldover Designs****Author:** Jonathan Stallrich¹¹ *North Carolina State University*

The foldover technique for screening designs is well known to guarantee zero aliasing of the main effect estimators with respect to two factor interactions and quadratic effects. It is a key feature of many popular response surface designs, including central composite designs, definitive screening designs, and most orthogonal, minimally-aliased response surface designs. In this paper, we show the foldover technique is even more powerful, because it produces degrees of freedom for a variance estimator that is independent of model selection. These degrees of freedom are characterized as either pure error or fake factor degrees of freedom. A fast design construction algorithm is presented that minimizes the expected confidence interval criterion to maximize the power of screening main effects. An augmented design and analysis method is also presented to avoid having too many degrees of freedom for estimating variance and to improve model selection performance for second order models. Simulation studies show our new designs are at least as good as traditional designs when effect sparsity and hierarchy hold, but do significantly better when these effect principles do not hold. A real data example is given for a 20-run experiment where optimization of ethylene concentration is performed by manipulating eight process parameters.

Keywords:

Response surface methodology, screening experiments, optimal design

Design of Experiments / 125**Possible Pitfalls in Using Bayesian Optimization (BO) to Quickly Optimize the Levels of Factors in a Physical Experiment****Authors:** Jacqueline Asscher¹; Sonja Kuhnt²¹ *Kinneret College, Technion*² *Dortmund University of Applied Sciences and Arts*

BO methods have recently been advocated as a newly accessible, straightforward, hands-off alternative to design of experiments (DOE) to efficiently optimize the levels of the factors in physical experiments. Physical experiments have random variation between replicated runs.

In the typical hands-on DOE approach to process optimization, an initial choice of factors and their ranges is made and an experimental design (classical or optimal) that enables an initial model to be fit is chosen. A significant model is fitted and used to optimize the process. Often a series of experiments are run, with the selection of factors and/or factor ranges adjusted. We assume (and check) that a simple linear or quadratic model plus random error will give a useful local approximation to the true relationship between the factors and the response/s.

The BO approach comes from the world of computer experiments, where both the nature of the problem and the solution strategy are very different. In this world, complex models are considered, random variation is non-existent or negligible, and runs are obtained one at a time or in small groups. An algorithm is used to efficiently explore a fixed design space with prescribed factors and ranges. Kriging models are automatically fitted and updated behind the scenes.

We consider possible pitfalls when using BO tools to optimize a real process with non-negligible random variation between runs. Can we start with only a handful of initial runs? Will variation cause the BO algorithm to fail? Will we know if it does fail?

Keywords:

DOE, BO, optimization

Design of Experiments / 56

Gradient-Based Line Search for Continuous Optimal Designs in Spatial Statistics**Author:** Helmut Waldl¹¹ *Johannes Kepler University Linz*

In spatial statistics, optimal designs select sampling locations such that a specific optimality criterion - for instance, maximizing the precision of parameter estimation or prediction accuracy - is satisfied. This task is particularly challenging in high-dimensional design spaces. Frequently, space-filling designs are used as an alternative; however, these perform well only under certain idealized conditions and fail to cover the area effectively as the dimensionality of the input space increases.

Search methods, such as the well-known sequential Wynn algorithm or the Fedorov exchange algorithm, usually add or exchange sampling locations from a finite set of candidate points, often relying on a trial-and-error approach. In the context of these algorithms, the term "gradient" is sometimes used misleadingly to denote the marginal gain in the optimality criterion resulting from a point exchange.

In this talk, the gradient is used in its genuine mathematical sense within an exchange algorithm. The presented method is compatible with continuous and differentiable optimality criteria, such as GV-optimality (Waldl & Müller, 2023). It computes the gradient of the criterion function with respect to the coordinates of the sampling locations. This gradient is then utilized within a line-search method as a descent direction to minimize the objective function. From a theoretical perspective, line-search methods are less susceptible to the curse of dimensionality than exchange algorithms that rely on a finite set of candidate points.

This strategy considerably reduces the number of criterion function evaluations, thereby identifying locally optimal designs faster than conventional exchange algorithms. Computer simulations in low-dimensional settings demonstrate interesting and promising preliminary results.

Keywords:

optimal design, exchange algorithms, gradient

Design of Experiments / 92

Robust D-Optimal Designs for Ordinal Split-Plot Experiments via Surrogate Objective Functions**Author:** Marcus Perry¹¹ *University of Alabama*

In industrial split-plot experiments, traditional algebraic block generators often force complete aliasing between critical sub-plot interactions and whole-plot blocks, trapping vulnerable effects within the high-variance whole-plot error stratum and severely reducing experimental power. While algorithmic D-optimal designs can mitigate this loss through partial confounding, generating such designs for ordinal responses has remained computationally intractable. Evaluating the exact Fisher Information Matrix (FIM) requires an $O(K^n)$ combinatorial enumeration, further compounded by an $O(2^n)$ dimensional penalty in exact Archimedean copula models. To address this structural limitation, we propose two efficient surrogate objective functions: a conditional GLMM surrogate that collapses the expected FIM into localized univariate expectations, and a Pairwise Composite Likelihood (PCL) copula surrogate that restricts the dependence structure to bivariate components. These surrogates reduce computational burden by several orders of magnitude while recovering designs that are identical or nearly identical to those obtained under exact evaluation. This enables the construction of parameter-robust split-plot designs that avoid sacrificing critical interactions to the whole-plot stratum, providing a computationally tractable approach for generating structurally robust, precision-efficient experiments.

Keywords:

Ordinal data, Restricted randomization, D-Optimal designs

Design of Experiments / 145**Model Selection for Screening Experiments with Random Effects: An Adaptation of the SAMS Algorithm**

Authors: Robin van der Haar¹; Peter Goos¹; Bart De Ketelaere²; Eric Schoen³; Mathias Born¹

¹ *KU Leuven*

² *Catholic University of Leuven*

³ *KU Leuven, Belgium*

Model selection in screening experiments is challenging when data exhibit random effects arising from split-plot structures, batch variation, or other sources of non-independent errors. Standard approaches such as stepwise selection, LASSO, and mixed-integer optimisation (MIO) typically assume independent errors or provide limited support for the correlation structures commonly encountered in industrial experiments, restricting their applicability in practice.

We present an adaptation of the Simulated Annealing Model Search algorithm (SAMS; Wolters and Bingham, 2011) that accounts for correlation induced by random effects, enabling computationally efficient model search while retaining generalized least squares or GLS-based inference for final model selection. The adapted algorithm retains the main strengths of SAMS: exploration of a large collection of well-fitting models, visualisation of effect aliasing via raster plots, and identification of active effects using an entropy-based criterion that quantifies how consistently terms appear across the model set. The method is implemented in the open-source Python package PyOptEx.

We evaluate the approach through a simulation study covering a range of experimental designs and variance structures, including a real split-plot type of screening experiment from potato fry production. We also apply the method to the data from the potato fry experiment to demonstrate practical applicability. Results show robust recovery of active effects across a range of designs and variance structures. Additional simulations investigate the performance of SAMS under strong random effects, where methods that ignore correlation structure, including MIO and backward selection, show systematic degradation in performance.

Keywords:

Screening experiments; random effects; split-plot designs

Editors' Corner Session - INFORMS Journal on Data Science / 7**Introduction to INFORMS Journal on Data Science (IJDS)**

Author: Yu Ding¹

Co-author: Antonio Lepore²

¹ *Georgia Tech*

² *Università degli Studi di Napoli Federico II - Dept. of Industrial Engineering*

In this talk, the Editors of IJDS will explain the mission and editorial structure of IJDS. We will also present some statistics and the recent/ongoing special issue. Time will be allocated for engaging questions and discussions.

Keywords:

IJDS, Editorial Process, Feature Articles

Editors' Corner Session - INFORMS Journal on Data Science / 8

Using Operational Data Analytics for Planning Decisions Under Uncertainty

Author: Natarajan Gautam¹

¹ *Syracuse University*

We consider a stochastic decision-making system with unknown parameters that need to be estimated to make appropriate decisions. We take the standard approach of exploring first and then exploiting. We start with a stylized model but present numerous applications in restaurant bookings, bike-share replenishments, customized order-fulfilment, air traffic control, virtual queueing systems, and inventory orders. In all these problems the underlying parameters of the stochastic process need to be learnt to make a decision. The approach of first estimating the parameters and then solving an optimization problem assuming those parameters are accurate does not work well. Even methods that use Bayesian analysis and bootstrapped simulations are not as effective. We discuss a promising approach called operational data analytics (ODA) where we optimally scale the estimated metrics and show that it results in much more accurate decisions. We present both simulated results where an oracle knows the underlying parameters as well as real data from bike-sharing to illustrate the effectiveness of the ODA approach.

Keywords:

operational data analytics; data-driven decision making under uncertainty; exploration and exploitation

Editors' Corner Session - INFORMS Journal on Data Science / 9

Estimating Hidden Epidemic: A Bayesian Spatiotemporal Compartmental Modeling Approach

Authors: Che-Yi Liao¹; Kamran Paynabar²; Lance Waller³; Peiliang Bai⁴

¹ *Georgia Tech*

² *School of Industrial and Systems Engineering*

³ *Emory University*

⁴ *JP Morgan*

Efforts to mitigate public health crises have been complicated by unreported cases and the ever-changing trends of those monitored health events across geographic regions and socioeconomic cultures. To resolve both challenges, we propose a Bayesian spatiotemporal susceptible-exposed-infected-recovered-removed (BayST-SEIRD) framework that builds the hidden effects of neighboring communities, local features, and the reporting rates into its transmission mechanism. To alleviate the computational burdens embedded in a fully Bayesian algorithm, we propose an alternating approach that learns the compartmental structure and the spatial effects separately. With a simulation study, we show that this algorithm can accurately retrieve our designed system. Then, we apply BayST-SEIRD to model the coronavirus disease 2019 (COVID-19) dynamics in the metropolitan Atlanta area. We observe that most counties' reporting rates were below 10% of the projected total infected population and that age and educational level are negatively correlated with the exposing rate, suggesting the needs for stronger incentives for COVID-19 testing and quarantine among the younger population. Importantly, BayST-SEIRD facilitates the reconstruction of actual case counts of the monitored subject among neighboring communities, which is critical to designing impactful public health policy interventions.

Keywords:

BayST-SEIRD, Spatiotemporal model, dynamics process

frENBIS Session - Advances in machine learning and sensitivity analysis / 2

Geographically Weighted Regression for Low-Cost Sensors Calibration

Author: Jean-Michel Poggi¹

Co-authors: Bruno Portier²; Emma Thulliez²

¹ *University of Paris-Saclay*

² *INSA Rouen Normandie*

Low-cost sensors are a new tool for improving air quality maps, which are of major interest in the current era of high-resolution, urban-scale air quality monitoring. These sensors require calibration using reference analyzers. A variety of strategies can be employed, ranging from individual pointwise calibration models to network calibration models. Here, we propose using geographically weighted regression (GWR) as an alternative method for calibrating or correcting micro-sensors. GWR is a local spatial statistical technique that can be used to model real-world phenomena by incorporating spatial nonlinearities. The additional flexibility that GWR provides, in the form of spatially varying coefficients, makes it possible to model the fact that the coefficients of a micro-sensor's calibration model change according to its position. We present a comprehensive approach, covering the selection of learning and test sets, the choice of spatial window, and the final evaluation using a spatial cross-validation scheme. The calibration results for Nitrogen dioxide (NO₂) are provided alongside some remarks about the estimated GWR model and the spatial content of the estimated coefficients. This study was carried out using the publicly available SenseURCity dataset in Antwerp. This dataset is especially relevant since it comprises 9 reference stations and 34 micro-sensors, all of which were collocated and deployed within the city.

More details can be found here: <https://revstat.ine.pt/index.php/REVSTAT/article/view/1012>

Keywords:

Low-cost sensors , Geographically Weighted Regression, Sensors network calibration

frENBIS Session - Advances in machine learning and sensitivity analysis / 88

Quantile oriented sensitivity analysis: why and how?

Authors: Véronique Maume-Deschamps¹; Kevin Elie-dit-Cosaque^{None}; Clémentine Prieur²; Ri Wang^{None}

¹ *Institut Camille Jordan, Université Claude Bernard Lyon 1*

² *université Grenoble Alpes*

Quantile-oriented sensitivity analysis allows to quantify uncertainty around quantiles, at different levels, while sensitivity analysis is often focused on deviation around mean (as it involves variances). We will consider quantile-oriented sensitivity indices (QOSA) and quantile-oriented Shapley effects (QOSE). We will present their relevance on some analytical examples, show how to estimate them and present some concrete examples.

Keywords:

quantile-oriented sensitivity analysis, quantile-oriented Shapley effects, projected forest

frENBIS Session - Advances in machine learning and sensitivity analysis / 91

Shapelet learning via sparse feature selection for interpretable time series classification

Authors: Julien Pelamatti¹; Nicolas Bousquet²; Ibrahim Seydi¹

¹ EDF R&D

² EDF

Time-series classification faces recurring challenges, including high dimensionality, autocorrelation, and the difficulty of identifying features that capture essential dynamics across temporal scales and phase shifts. We address these issues through shapelet decomposition, a technique that extracts shape-based features from time series while preserving both temporal and frequency information. The approach represents a dataset through minimal distances to representative patterns, thereby improving interpretability by highlighting the patterns most relevant to classification decisions.

In this framework, selecting the most informative shapelets, or equivalently an effective projection basis, is critical for both performance and interpretability. Building on prior work that optimizes shapelet values to maximize classifier accuracy, we propose a formulation based on sparse feature selection. Rather than jointly optimizing shapelets as classifier parameters, our approach seeks a data transformation that produces covariates that are highly predictive of the class label while being weakly dependent, thereby reducing redundancy and improving classifier stability and interpretability.

To this end, shapelets are determined by optimizing a HSIC-Lasso criterion on the transformed dataset. HSIC-Lasso is a nonlinear feature-selection method for high-dimensional settings that identifies a sparse subset of features that are highly informative about the response while exhibiting very limited or null redundancy. To make the optimization tractable, the non-smooth minimum-distance operator between shapelets and time series is replaced by a soft-minimum approximation, yielding a differentiable loss function.

The proposed method is validated on both synthetic and real-world datasets, demonstrating improvements in classification performance, scalability, and interpretability.

Keywords:

interpretable machine learning, time series classification, shapelet transformation, shapelet learning, sparse feature selection

From Automation to Autonomy: New Frontiers of AI and Statistics for Data in Business and Industry / 123**From Automation to Autonomy: New Frontiers of AI and Statistics for Data in Business and Industry**

Authors: Bianca Maria Colosimo¹; Kamran Paynabar²

¹ *Politecnico di Milano*

² *School of Industrial and Systems Engineering*

The rapid evolution of Artificial Intelligence is transforming how data is generated, analyzed, and leveraged across business environments, industrial systems, and organizational processes. Moving beyond the traditional Industry 4.0 paradigm centered on automation, new AI technologies are opening the way toward increasingly autonomous data-driven systems capable of supporting complex decision-making, knowledge generation, and adaptive operations.

This panel session brings together internationally recognized experts to explore the evolving intersection of artificial intelligence (AI), generative AI, and statistical methodologies in modern business and industrial contexts. The discussion will focus on how statistical thinking continues to provide essential foundations for AI-driven approaches, enabling robust, interpretable, reliable, and scalable solutions across a wide range of applications.

Particular attention will be devoted to emerging developments such as large language models (LLMs), synthetic data generation, foundation models, and generative AI tools, and to the ways these technologies are reshaping data management, analytics, and operational processes within supply chains, business functions, and industrial applications. Panelists will discuss both opportunities and challenges related to data quality, model interpretability, governance, human-AI interaction, and the deployment of AI systems in real-world environments.

The session will also examine how these technological advances may influence future research methodologies and industrial innovation, reflecting on the evolving role of statisticians, data scientists, and domain experts in bridging methodological rigor with practical implementation. By fostering a multidisciplinary and collective discussion, the panel aims to provide insights into the transition “from automation to autonomy” and the new frontiers emerging at the intersection of AI, statistics, business, and industry.

Keywords:

Artificial Intelligence, Generative AI, Statistics for Business and Industry, Synthetic Data, Large Language Models, Model Interpretability

Generative AI for Statistics: Practice, Productivity, and Prioritization / 90**Generative AI for Statistics: Practice, Productivity, and Prioritization**

Authors: Christian Ritter¹; Morten Bormann Nielsen²; Winfried Theis³

¹ *Ritter and Danielson Consulting*

² *Danish Technological Institute*

³ *Kufuu Consultancy*

This session offers a practical and thought-provoking exploration of how generative AI and large language models are transforming the daily work of statisticians, data scientists, and educators.

We start with a concise “kaleidoscope” of real examples illustrating what modern general-purpose AI tools can achieve in practice, with demonstrations that show how complex tasks can now be approached through simple prompts. Examples include interactive maps, image-based data interpretation and data visualization, highlighting AI’s growing ability to extract insights from non-tabular data and its implications for both teaching and applied work.

The second contribution brings a personal perspective on using generative AI as a “virtual team.” Drawing on experience from industry and entrepreneurship, it shows how AI can accelerate tasks such as generating R code, building interfaces, and creating simulations for exploratory purposes. These capabilities enable rapid prototyping, while also requiring careful validation and critical thinking to ensure reliable outcomes.

The session concludes with an organizational viewpoint on prioritizing AI initiatives. It addresses how to distinguish meaningful applications from low-impact experimentation, avoid tool proliferation, and implement governance approaches that balance decentralization with coordinated, value-driven development.

An open discussion will follow, inviting participants to engage, share experiences, and reflect on the role of AI in practice.

Keywords:

Generative AI, Large Language Models (LLMs), Applied Statistics, AI in Practice, AI Governance

Greenfield challenge / 149**An Academic Year Dedicated to George E. P. Box at the Faculty of Mathematics and Statistics in Barcelona****Author:** Lluís Marco-Almagro¹**Co-authors:** Ariel Duarte López²; Ignasi Puig de Dou²; Lourdes Rodero³; Marta Perez-Casany²; Pedro Delicado⁴; Pere Grima⁴; Xavier Tort-Martorell⁵; Josep Ginebra⁴¹ *UPC Universitat Politècnica de Catalunya | BarcelonaTech*² *Universitat Politècnica de Catalunya*³ *Technical University of Catalonia (UPC)*⁴ *Universitat Politècnica de Catalunya*⁵ *UPC Universitat Politècnica de Catalunya, Barcelona Tech*

Every year, the Faculty of Mathematics and Statistics at Universitat Politècnica de Catalunya · BarcelonaTech (UPC) dedicates the academic year to a prominent scholar. Traditionally, the chosen figure has been a mathematician, reflecting the Faculty's stronger tendency to recognise mathematicians rather than statisticians.

Nevertheless, our proposal to dedicate the 2023–2024 academic year to George E. P. Box, marking the tenth anniversary of his death, was accepted. The initiative of devoting an academic year to a particular figure often has a limited impact, as only people who are already interested in that person's field tend to participate in the activities. Our approach was precisely the opposite: we prepared activities and events that could attract not only the usual members of the Faculty—mathematicians and statisticians—but also people from other disciplines and from both academia and industry.

The initiative aimed to promote the importance of collaboration between statisticians and experts from other fields within a multidisciplinary framework, and to highlight the role of statistics as an essential partner in scientific progress. We think this approach would have pleased both George E. P. Box and Tony Greenfield. In fact, Tony Greenfield taught a short course at the Faculty for many years.

Each month, a quote from Box was displayed in the Faculty's main hall. The first quote was deliberately provocative and helped establish the theme of the entire year: "Why do we aspire to be second-rate mathematicians when we can be first-rate scientists?" This quote was also used as the title of the opening talk, delivered by Geoff Vining.

Activities held during the year included two talks, by Geoff Vining and Daniel Peña; a design of experiments competition; a series of debates among teachers on controversial topics; and audiovisual materials celebrating Box's legacy (<https://fme.upc.edu/ca/la-facultat/activitats-fme-personalitats-del-curs/personalitat-del-curs/2023-2024-Box/Audiovisuals>).

Keywords:

Hands-on with physical experiments for teaching Design of Experiments / 37**Hands-on with physical experiments for teaching Design of Experiments****Authors:** Christian Ritter¹; Morten Bormann Nielsen²¹ *Ritter and Danielson Consulting*² *Danish Technological Institute*

Design of Experiments (DOE) is powerful but rarely intuitive. What is wrong with poking around in design space? Why not vary one factor at a time? The mathematics answers clearly, but the classroom often doesn't.

Physical experiments —where participants can see, touch, and interact with a real system —bring DOE concepts to life. They make abstract ideas like screening, response surface optimization, and measurement variability concrete in ways simulations cannot. But teaching with them requires experience and calm: real systems have their own lives and don't only do what we want them to.

In this 90-minute active session, participants work in small groups at stations built around different physical experiments, each illustrating some aspect of DOE: screening designs, response surface methodology, measurement systems analysis, or mixture/sensory experimentation. Rather than rotating through all stations, each participant engages deeply with one or two, guided by structured exercises. A reporter at each station captures observations on a prepared checklist, and a plenary discussion closes the session.

The goal is twofold. First, participants experience firsthand running a physical DOE exercise feels like —the engagement, the learning, and the challenge of real systems pushing back. Second, the session builds a shared resource: a standardized booklet of physical experiment systems for DOE teaching, tagged by topic, time, and suitability, hosted as a living document on the ENBIS website. Whether you teach DOE in universities or in industry, this session will equip you with practical tools and inspiration to make your teaching more effective —and more fun!

Keywords:

DOE, teaching

INFORMS Session - Quality Statistics Reliability (QSR) / 70**Response Grid Plots for Model-Agnostic Explainable AI****Author:** Russell Barton¹¹ *Pennsylvania State University*

Machine learning (ML) models are used to provide predictive characterization of response functions based on observed or simulated data. Insight on the nature of the approximation to the underlying response function is essential for the model output to be trusted and used by decision makers, a key part of explainable AI. When a response function is well-behaved, methods that identify marginal effects and interactions for important input variables can provide adequate insight on its nature. Alternatively, interpretable ML models can provide an ensemble of simple models to decompose a complex response into interpretable elements. This talk introduces a model-generated but model-independent visual display of functional information that provides direct insight, not interpreted through model form, model coefficients, or numerical characterizations such as those produced by sensitivity analysis. This is shown to complement existing xAI approaches. Further, the method can be applied directly to system response data if the data are collected from a factorial (grid) design. The value of response grid plots (RGPs) is shown through several examples.

Keywords:

Interpretable AI, Explainable AI, Graphical Methods

INFORMS Session - Quality Statistics Reliability (QSR) / 77

FedCOT: Personalized Federated Transfer Learning With Conditional Optimal Transport for Manufacturing Predictive Modeling

Authors: Anyi Li¹; Jia Liu²

¹ Auburn University

² University of Florida

Effective predictive modeling in large-scale manufacturing is hampered by the isolated and limited data from individual organizations, collected from costly experiments and various inspections. Collaboration across organizations can handle these limitations, but it faces two main challenges: privacy concerns over organizations and heterogeneous features from varied sensing and inspection capabilities. Federated learning (FL) offers a solution by allowing organizations to collaboratively train a predictive model without sharing raw data, but standard FL struggles with the problem of feature heterogeneity. To address these challenges, we propose a personalized federated transfer learning framework with conditional optimal transport (FedCOT). FedCOT enables “target” organizations with limited features to benefit from “source” organizations with sufficient features in prediction performance while keeping data privacy through a central server. Each organization learns a personalized encoder-regressor structure by alternating optimization: the encoder maps heterogeneous inputs into a shared latent space, and the regressor predicts responses from the latent representations. Target organizations align their latent representations’ structure and corresponding responses with the source organizations’ information via COT. We evaluate FedCOT through simulations and a manufacturing case study on fatigue life prediction of additive-manufactured parts. Our case study demonstrates that FedCOT achieves the latent space where target organizations are highly aligned with source organizations, leading to a significant 34.99% improvement in prediction performance over the baseline method. Additionally, we provide theoretical guarantees on OT gradients and predictive consistency.

Keywords:

Federated learning, feature heterogeneity, additive manufacturing.

INFORMS Session - Quality Statistics Reliability (QSR) / 127

Spatial–Temporal Large Language Models for Time-Series Data

Author: Ziyue Li¹

¹ Technical University of Munich

Since 2023, large language models (LLMs) have begun to reshape the landscape of time-series analysis. In this talk, we present our latest research, insights, and perspectives on using LLMs to model time-series data—both as a standalone modality and in combination with other spatial or contextual information/modality. We will explore three key questions:

- (1) What are spatiotemporal LLMs?
- (2) How can they be applied effectively to time-series data?
- (3) Why do they succeed—or fail—in practical scenarios?

Keywords:

Large Language Model, Time-series Data

ISBIS Session - Advances in Statistical Analysis for Large-Scale Dynamic Networks, Labour Market Outcomes, and Interpretable Ensemble Methods / 108

Real-Time Change Detection in Large-Scale Dynamic Networks

Author: Nathaniel Stevens¹

¹ *University of Waterloo*

In this talk we present a real-time change detection method for monitoring large, dynamic networks with community structure. We model the propensity for communication within and between communities to incorporate the structure of the underlying network. Our focus on communities makes our method scalable to large-scale networks and we use a window-based approach to accommodate network dynamics and a changing node set over time. We monitor deviations from the underlying model, where unexpectedly large deviations indicate potential changes. We benchmark our method using a simulation study with networks of 10,000 nodes and demonstrate its flexibility and detection accuracy. We apply the proposed method to a Reddit network defined by discussions on the r/WallStreetBets subreddit about the stock GME, which experienced one of the most notorious short squeezes in market history as retail investors took on hedge funds to drive GameStop's stock price sky-high. Our method is able to identify changes in the network well before the short squeeze.

Keywords:

anomaly detection, multivariate control chart, network monitoring

ISBIS Session - Advances in Statistical Analysis for Large-Scale Dynamic Networks, Labour Market Outcomes, and Interpretable Ensemble Methods / 93

Heterogeneity in STEM Graduate Wages Across Italian Universities: A Weighted Mixed-Effects Analysis

Author: Andrea Carta¹

Co-authors: Francesca Atzori¹; Giulia Contu²; Luca Frigau²

¹ *university of Cagliari*

² *University of Cagliari*

This study examines wage outcomes at the cohort level among STEM graduates in Italy, using data from the AlmaLaurea surveys covering 60 institutions over the period 2008-2023. The unit of analysis is a graduate cohort sharing the same university, degree level, and disciplinary category, observed one, three, or five years after completing their studies. Given the nested structure of the data and the heteroscedasticity arising from differences in cohort size, estimation relies on a weighted linear mixed-effects model incorporating random intercepts for both institution and survey year.

The findings reveal significant predictors of mean cohort remuneration. Geographic relocation for work is positively linked to earnings, whereas a greater share of women within a programme is negatively associated with cohort-average wages, capturing a composition effect at the programme level. Graduates holding a master's degree and those specialising in Computer Science and ICT or Industrial and Information Engineering command higher wages, while cohorts in Architecture and Civil Engineering earn considerably less relative to the Scientific baseline. Time since graduation emerges as one of the strongest predictors, with average wages rising markedly across survey intervals. A clear territorial divide is apparent, with cohorts based in the North earning substantially more than their Southern counterparts.

The cohort-level perspective treats average wages as synthetic indicators of the labour market returns associated with specific educational paths. Goodness of fit is evaluated via the correlation between fitted and observed cohort means and a size-weighted root mean squared error. The core findings hold across four robustness checks.

Keywords:

Returns to Education, STEM, Wages, Mixed-Effects Models, Italy, regression

ISBIS Session - Advances in Statistical Analysis for Large-Scale Dynamic Networks, Labour Market Outcomes, and Interpretable Ensemble Methods / 101

Agreement over Association in Explainable Ensemble Trees

Author: Agostino Gnasso¹

Co-authors: Carmela Iorio¹; Massimo Aria¹

¹ *University of Naples Federico II*

Ensemble methods such as Random Forests achieve strong predictive accuracy but at the cost of interpretability. Explainable Ensemble Trees (E2Tree) address this trade-off by constructing a single interpretable tree that approximates the co-occurrence structure induced by the ensemble. The quality of this approximation matters: when interpretability is invoked for regulatory or scientific purposes, a poor reconstruction does not merely underperform: it actively misleads.

Existing validation approaches for E2Tree rely on the Mantel test, which measures the association between proximity matrices. We argue that this answers the wrong question. The issue parallels the classical distinction between correlation and concordance in method comparison: two matrices can be perfectly correlated while differing substantially in absolute terms. What E2Tree validation requires is a scale-sensitive measure of agreement.

We propose a family of divergence and similarity measures for this task. The centrepiece is the Normalized Loss of Interpretability (nLoI), a statistic rooted in the Cressie–Read power divergence family, whose key feature is a decomposition into within-node and between-node components. This identifies not only how much reconstruction quality is lost, but where and why, a diagnostic capability unavailable from correlation-based approaches. Four complementary measures complete the family: Hellinger distance, weighted Root Mean Squared Error, the RV coefficient, and the Structural Similarity Index, each targeting a distinct facet of matrix agreement.

A unified permutation testing framework based on simultaneous row/column permutation provides valid inference for all measures. Monte Carlo simulations confirm correct Type-I error control and adequate power; empirical results on benchmark datasets illustrate the framework's utility.

Keywords:

Ensemble Models, Interpretability, Agreement Measures

ISEA Session - Statistical Engineering / 11

XAI for signal diagnosis in SPM

Authors: Inez Zwetsloot¹; Jiaqi Qiu¹

¹ *University of Amsterdam*

Artificial Intelligence (AI) has shown become very popular as modelling strategy within statistical process monitoring (SPM), particularly in detecting abnormal process behaviours. However, for existing AI-based SPM methods, diagnosing features associated with signal remains challenging, as traditional diagnosis methods are not directly applicable. This lack of diagnosis makes it difficult to make an out-of-control action plan and take appropriate actions once a signal is detected, and thus impedes the AI-based SPM methods from being applied in practice. Explainable AI (XAI) offers a promising framework for addressing this limitation by providing feature relevance information for the model outputs, which can help identify the features related to abnormal process behaviour. This work proposes a general framework for combining XAI with AI-based monitoring. Simulation studies and a real-world case study show the effectiveness of the proposed method.

Keywords:

Signal diagnosis, XAI, statistical process monitoring

ISEA Session - Statistical Engineering / 107**Assessing inter-rater reliability of LLMs with the probability of agreement****Author:** Nathaniel Stevens¹¹ *University of Waterloo*

Inter-rater reliability, the quantification of agreement between individuals who assign scores to the same phenomenon, is an important consideration in all fields for which data drives decision-making (e.g., business and industry, healthcare, social and behavioural sciences, education, etc.). Traditionally, the raters scoring the phenomenon have been human beings. With the proliferation of AI, a natural question arises: can a large language model (LLM) perform this task as well as humans? Central to this question is the assessment of inter-rater reliability of LLMs relative to humans. In this talk, we describe one such problem in the ed-tech space, where the goal is to establish the reliability of LLM-evaluation of educational material. In particular, we describe an end-to-end framework implemented at a prominent ed-tech company in which agreement studies are designed and analyzed to compare LLM raters with human raters via the probability of agreement.

Keywords:

agreement, reliability, concordance

ISEA Session - Statistical Engineering / 111

A Cooperative Framework for Raw Material Acceptance: Integrating Supplier and Customer Perspectives via SMB-PLS**Author:** Alberto J. Ferrer-Riquelme¹**Co-authors:** Joan Borràs-Ferrís²; Carl Duchesne³¹ *Universitat Politècnica de València, Spain*² *Kensight Solutions, S.L.*³ *Université Laval, Canada*

The increasing digitalization of manufacturing processes is transforming the relationship between raw material suppliers and customers. Rather than defining rigid specifications that often lead to unnecessary rejection of raw material lots, Industry 4.0 opens the door to more collaborative and knowledge-driven strategies. In this context, multivariate raw material specifications should not be understood as static acceptance limits, but as dynamic regions that can be jointly improved by both actors to guarantee final product quality while facilitating raw material acceptance.

This work proposes an integrated framework in which supplier and customer cooperate to improve the feasibility of raw material acceptance through Sequential Multi-Block Partial Least Squares (SMB-PLS). On the supplier side, the methodology provides a systematic diagnosis of assignable causes affecting the capability of the supplied raw material, identifying which properties and variability sources should be modified to better satisfy the customer requirements [1]. Simultaneously, from the customer perspective, the approach identifies process variables that can be manipulated to enlarge the feasible raw materials specifications without compromising the quality attributes of the final product [2].

The novelty of the proposal lies in combining both perspectives into a unified cooperative strategy where supplier and customer iteratively adapt to each other instead of operating under a restrictive pass/fail paradigm. The sequential structure of SMB-PLS is especially suitable for this purpose as it explicitly accounts for the ordered relationships between raw material properties, process conditions and final product quality. The proposed methodology is illustrated through a real industrial case study concerning cheese production from milk.

Keywords:

Sequential Multi-Block Partial Least Squares, Design Space, Raw materials specifications, supplier and customer.

itENBIS Session - Flexible and Scalable Models for Complex Data Structures / 34

Fuzzy jump models for soft and hard clustering of mixed-type multivariate time series

Author: Federico Cortese¹

Co-authors: Antonio Pievatolo²; Elisa Maria Alessi²

¹ *University of Milan*

² *CNR-IMATI*

Statistical jump models have been recently introduced to detect persistent regimes by clustering temporal features while discouraging frequent regime changes. However, they rely on hard clustering and therefore do not account for uncertainty in state assignments.

In this work, we propose a fuzzy extension of the statistical jump model that incorporates uncertainty in cluster membership. Leveraging the similarities with the fuzzy *c*-means framework, the proposed *fuzzy jump model* sequentially estimates time-varying state probabilities. The approach is flexible, as it encompasses both soft and hard clustering through a fuzziness parameter and naturally accommodates multivariate time series of mixed type.

Through extensive simulation studies, we show that the proposed method accurately recovers the latent state distribution and outperforms competing approaches in scenarios with high assignment uncertainty. We further illustrate its practical relevance on real data from celestial mechanics, addressing the identification of co-orbital regimes in the three-body problem, with implications for asteroid dynamics and space mission design.

Keywords:

co-orbital motion, mixture models, regime-switching models, time series analysis, unsupervised learning

itENBIS Session - Flexible and Scalable Models for Complex Data Structures / 52

Variable selection and high dimension in multinomial models

Authors: Marcella Niglio¹; Maria Iannario²; Marialuisa Restaino¹

¹ *University of Salerno*

² *University of Naples Federico II*

Identifying predictors associated with specific response categories in multinomial logistic regression is a challenging task. It is furthermore complex in a high-dimensional setting, where the number of covariates is higher than the number of units. To address the variable selection in high dimensional domain and in the presence of multinomial models with unordered responses, we propose a two-step ranking-based approach for category-specific variable selection. The method relies on marginal multinomial regressions, in which each covariate is separately regressed on the response variable, thereby enabling the identification of predictors relevant to individual categories.

The comparison of the proposed approach with standard penalized techniques demonstrates the parsimony of the final model (after variable selection) and the accuracy of its predictive performance. A further advantage of the proposed approach is the computational efficiency and scalability, which make the method well-suited for models involving a large number of predictors and response categories.

Keywords:

Multinomial regression model, high dimension, variable selection

itENBIS Session - Flexible and Scalable Models for Complex Data Structures / 76**Small area models for count data: an area-level EFD model****Author:** Serena Arima¹**Co-author:** Silvia Polettini²¹ *University of Salento*² *University of Rome "La Sapienza"*

In this contribution, we focus on small-area compositional data. These data are defined as vectors whose elements are strictly positive and sum to one (e.g., proportions). Compositional data arise in various fields, including medicine, economics, psychology, and environmetrics. They are defined on the D -part simplex (S^D) and require complex techniques for proper analysis.

A traditional approach involves applying log-ratio transformations, which map the D -part simplex onto a $(D-1)$ -dimensional real space. However, this approach has several limitations: parameter estimates are interpretable only in the transformed space, and issues such as skewness, heteroscedasticity, non-normality, and outliers may bias inference.

An alternative solution for regression with compositional responses is the Dirichlet model. It is often implemented using a multinomial logit function to link the response mean vector to covariates, allowing for a straightforward interpretation of regression coefficients in terms of log-odds ratios. Nevertheless, the Dirichlet model has significant drawbacks: it uses only a single parameter to describe the entire variance-covariance matrix (the precision parameter), offers limited flexibility, and implies several forms of simplicial independence. Moreover, it always yields negative covariances, making it unable to model many relevant phenomena (e.g., heavy-tailed or multimodal responses). We therefore propose using the Extended Flexible Dirichlet (EFD), a structured mixture of Dirichlet components. Unlike general Dirichlet mixture models, the parameters of each EFD component are strictly linked to one another, providing greater flexibility. We employ the EFD for modelling small-area data: we reparameterize the model in terms of mean and precision and incorporate covariates to account for design effects. We estimate the model within a fully Bayesian framework and assess its performance using simulated data.

Keywords:

EFD, counts, Bayesian models

JQT / QE / Technometrics Session / 132

Constructing large OMARS designs by concatenating two definitive screening designs

Author: Peter Goos¹**Co-authors:** Alan Vazquez²; Eric Schoen³¹ *KU Leuven*² *Tecnologico de Monterrey*³ *KU Leuven, Belgium*

Orthogonal minimally aliased response surface or OMARS designs permit the study of quantitative factors at three levels using an economical number of runs. In these designs, the linear effects of the factors are neither aliased with each other nor with the quadratic effects and the two-factor interactions. Complete catalogs of OMARS designs with up to five factors have been obtained using an enumeration algorithm. However, the algorithm is computationally demanding for designs with many factors and runs. To overcome this issue, we propose a construction method for large OMARS designs that concatenates two definitive screening designs and improves the statistical features of its parent designs. The concatenation employs an algorithm that minimizes the aliasing among the second-order effects using foldover techniques and column permutations for one of the parent designs. We study the properties of the new OMARS designs and compare them with alternative designs in the literature.

Keywords:

orthogonal design, screening design, response surface design

JQT / QE / Technometrics Session / 131

The allure and the perils of summarizing process performance in a single index

Author: Murat Kulahci¹**Co-author:** Mahmut Onur Karaman²¹ *DTU*² *Hacettepe University*

Industrial statisticians are well acquainted with the temptation to simplify the analysis of process performance through summary statistics. One of the most prominent examples is process capability analysis, in which a process's ability to produce items within specification limits is expressed as the ratio between the specification range and the natural variability of the process. The latter is typically estimated from observed data so as to cover a specified proportion of the underlying distribution, resulting in point estimates of the "true" capability indices that are inherently subject to uncertainty. Furthermore, the calculation of these indices often relies on assumptions regarding the characteristics of the process data. In our collaborations with industry, we have found violations of the normality assumption to be among the most common challenges encountered in practice. In this presentation, we discuss several approaches to capability analysis for processes in which the data cannot reasonably be assumed to follow a normal distribution.

Over the years, we have also observed growing interest in multivariate capability indices as an increasing number of quality characteristics are collected and jointly used to assess process performance. We compare several multivariate capability indices proposed in the literature and discuss their differences and limitations using both simulated data and a case study drawn from industrial practice.

Keywords:

Point Estimates, Capability Analysis

JQT / QE / Technometrics Session / 106

Time series models for cultural heritage preservation: The case of the Brunelleschi's Dome

Authors: Bruno Bertaccini¹; Fabrizio Cipollini¹; Fiammetta Menchetti¹; Silvia Bacci¹

¹ *Università di Firenze*

Structural Health Monitoring (SHM) of historical heritage is a crucial challenge for preserving humanity's cultural assets. Increasingly, monuments are equipped with multi-sensor monitoring systems that continuously collect large volumes of data over extended periods. These data require the application of appropriate statistical methods to provide meaningful insights into the structural health of such monuments. In this context, we explore the potential of time series models, specifically ARIMA and Structural VAR models, which until now have been used relatively little in the context of static SHM. As a case study, we focus on the Cathedral of Santa Maria del Fiore (Florence). Since the late 15th century, cracks have appeared in its Dome, prompting the initiation of systematic monitoring activities. In 1987, an extensive electronic monitoring system was installed, including numerous deformometers that record crack width variations several times a day. Using ARIMA and Structural VAR models, we investigate the interconnections among cracks and their dynamic responses to exogenous thermal shocks. Our findings reveal a negative relationship between crack width and humidity, consistent patterns between cracks and masonry temperature across the Dome, and novel evidence of differential responses between even and odd webs, as well as between the inner and outer domes.

Keywords:

ARIMA, Impulse-response analysis, Sensor data, Structural health monitoring, Structural VAR

Keynote / 84**Physics-Informed Statistical Learning for Data-Driven Decision Making in Science and Industry****Author:** Laura M. Sangalli¹¹ *Politecnico di Milano*

The increasing availability and complexity of data are transforming decision-making processes across science, industry, and engineering. Modern datasets are often high-dimensional, heterogeneous, and structured over space and time, and are collected on domains with complex geometries, including environmental domains, biological structures, and engineering systems. In many applications, the underlying phenomena are governed by known physical mechanisms, while multiple data sources provide complementary information. In other cases, such as pharmacokinetics, data may be limited, making robust inference challenging without integrating prior knowledge. In these settings, standard statistical approaches may fail to fully exploit available information.

This lecture presents a class of physics-informed statistical learning methods that integrate data with physical knowledge to support more reliable inference and decision-making. These approaches extend nonparametric regression and classical statistical learning models by incorporating regularization based on differential operators, including problem-specific partial differential equations. The resulting framework enables the analysis of spatial, spatio-temporal, and functional data over complex domains, and provides a natural approach to kinetic modeling in industrial applications.

By combining statistical rigor with physically grounded modeling, these methods improve interpretability, enhance predictive accuracy, and support the integration of heterogeneous data sources. Their practical value will be highlighted by illustrative applications from environmental monitoring, engineering design, and medical and pharmaceutical research, including stability studies and emerging measurement technologies such as spatial transcriptomics. The lecture concludes by outlining key challenges and future directions, emphasizing the role of physics-informed statistical learning as a unifying framework for advancing data-driven decision-making and scientific discovery.

Keywords:

high-dimensional and complex data; nonparametric regression

Keynote / 128**Domain Generalization and Adaptation in Digital Health (and beyond)****Author:** Peter Bühlmann¹¹ *ETH Zurich, Seminar for Statistics*

Statistical models and machine learning algorithms are often deployed in populations that differ from those on which they were trained, a challenge that is particularly acute in digital health. We discuss domain generalization and adaptation for a large-scale database from multiple countries with intensive care unit (ICU) data. We introduce Distributionally Robust Invariance Learning as an approach to exploiting stable structure across environments, and conclude with a brief discussion of the potential and limitations of a novel foundation model in this context.

Keywords:

machine learning

Mathmet session - Digital twins for industrial machine vision systems and reference data generation / 60

Digital Twin for optimal positioning of an industrial machine vision camera

Author: Dario Pasin¹

Co-authors: Sofia Catalucci ¹; Enrico Savio ¹

¹ *Università degli Studi di Padova*

In industrial machine vision, camera positioning is traditionally a manual, iterative, trial-and-error process. Even if sufficient accuracy can be reached, this leads to prolonged downtime during initial installation and maintenance, especially for inspection tasks where the camera must be positioned at a precise location, orientation, and working distance. In addition, the operator-dependent nature of the conventional manual camera setup cannot often guarantee the reproducibility of machine vision inspection results. This is also crucial for setups in which multiple cameras work in parallel, performing the same task to speed up batch inspection, resulting in the possibility of slightly different outcomes between different cameras on the same object.

This work presents a Digital Twin that assists the operator in the camera positioning process. A calibration target is used both to perform camera calibration and to detect the current camera position, allowing feedback to the operator for fine-tuning the position of the camera.

An industrial case study on pharmaceutical inspection machines, a field subject to strict regulatory and acceptance rules, is presented. The machines need accurate positioning of the multiple cameras involved in the different inspection tasks, at different times: first assembly on the shop floor, re-assembly after shipping parts to the customer and after maintenance. In this context, manufacturers need the imaging configuration to be reproducible quickly and accurately. The proposed Digital Twin addresses this challenge providing continuous feedback on camera setup and is paired with a specifically designed calibration target, enabling more reproducible acquisition conditions.

Keywords:

Digital Twin, Camera Positioning, Machine Vision

Mathmet session - Digital twins for industrial machine vision systems and reference data generation / 65

Enhancing photogrammetric measurement systems through Gaussian Splatting

Authors: Mattia Trombini¹; Giacomo Maculotti¹; Gianfranco Genta¹; Maurizio Galetto¹

¹ *Politecnico di Torino*

Photogrammetric measurement systems are widely adopted in manufacturing since they combine accurate geometric reconstruction, operational flexibility and relatively low implementation costs, while also supporting scalable and repeatable measurement workflows. However, photogrammetry is constrained by the initial image acquisition, since any extension of the dataset requires additional physical acquisitions. Recent advances in radiance-field methods, particularly Gaussian Splatting (GS), enable the creation of high-quality and fully navigable 3D scenes, allowing the extraction of renders (synthetic images) that can be integrated into the photogrammetric workflow.

This work investigates two potential benefits of integrating renders into conventional photogrammetry. First, it addresses data augmentation through synthetic views, as renders may represent a practical strategy for densifying the input image dataset without further physical acquisition. Second, the availability of a navigable scene enables the virtual capture of images from novel and controlled camera positions within a fully digital environment. In this sense, this second objective corresponds to the construction of a digital twin of the photogrammetric measurement system, with renders intended to mirror physical image acquisition while providing access to a theoretically unlimited range of virtual camera positions.

Preliminary results show that the progressive replacement of real images with renders introduces a systematic effect in the measurement. This study also includes the analysis of residual form errors, e.g., planarity and sphericity. It further examines the compensation of the observed systematic component through the introduction of a controlled perturbation term in the collinearity equation.

Keywords:

Photogrammetry, Gaussian Splatting, Digital Twin

Mathmet session - Digital twins for industrial machine vision systems and reference data generation / 21

Development of digital twins in metrology for a stereovision system

Author: Katarina Josic¹

Co-authors: Ladji Fofana¹; Louis-Ferdinand Lafon²; Joffray Guillory²; Romain Brault³; Nabil Anwer⁴; Olivier Bruneau⁴; Hichem Nouira²

¹ *LNE/CNAM/LURPA/CETIM*

² *LNE/CNAM*

³ *CETIM*

⁴ *LURPA*

Machine vision systems are important in Industry 4.0 as they allow fast automated inspection and quality control. Traceable metrology for machine vision systems is critical for the digital transformation of the Industry 4.0 objectives defined by the EU Green Deal. Nevertheless, these systems currently lack well-defined uncertainty frameworks and calibration techniques. For contactless 3D scanning of large volume mechanical parts and form error assessment, a structured light stereovision system is developed based on the active stereovision principle. The system is mounted on a positioning industrial robot with 6 rotational axes fixed on one additional translational guiding stage. A number of traceable high accuracy multilateration systems were also used for the validation of the extrinsic parameters. A digital twin architecture is presented for automating the calibration process and optimizing the measurement strategy based on environmental factors surrounding stereovision systems, assuring traceability and increasing accuracy. Digital Twin in metrology is proposed as a digital model of a measurement process connected to the physical system by a closed loop, then providing the associated measurement uncertainty for a given measured value traceable to the meter unit definition. They enable virtual testing and validation by linking the digital and physical domains, paving the way for more effective, and efficient adoption of stereovision technologies in industrial settings. The proposed framework is centred around metrology, where system positioning decisions are based on the lowest measurement uncertainty.

The project (23IND08 DI-Vision) has received funding from the European Partnership on Metrology, co-financed from the European Union's Horizon Europe Research and Innovation Programme and by the Participating States.

Keywords:

Digital twin, machine vision, measurement uncertainty evaluation

Mathmet session - Digital twins for industrial machine vision systems and reference data generation / 73

Generalized constraints on reference data generation for software verification

Author: Louis-Ferdinand LAFON¹

Co-authors: Hichem Nouira ¹; Katarina Josic ²; Ladji Fofana ²; Nabil Anwer ³

¹ *LNE*

² *LNE/CNAM/LURPA/CETIM*

³ *LURPA*

Vision-based systems in industrial applications involve a wide range of software, including fitting, association, and cloud-to-cloud registration. Software verification is required to guarantee the accuracy of estimated parameters. Verification typically relies on realistic datasets generated using ray casting to sample points on the predefined surface, followed by the addition of random deviations from the surface to simulate measurement noise. However, random and unconstrained deviations introduce a statistical bias, as the predefined fitting, association or registration parameters do not necessarily correspond to the optimal parameters of the underlying problem.

This paper proposes a generalized method for reference data generation that ensures unbiased verification. The approach is applicable to any differentiable parametric surface representation and to many least-squares problems. This is achieved through inequality constraints based on curvature and outlier bounds, combined with a positive semi-definite Hessian condition and a null gradient equality constraint.

To implement these constraints, a symbolic programming paradigm is adopted for automatically computing gradients, Hessians, and curvature-related quantities from the problem definition. The generation of deviations satisfying both equality and inequality constraints is formulated as a quadratic programming problem and is solved using an efficient operator splitting solver.

The proposed method produces unbiased reference datasets, ensuring metrological traceability to the reference parameters. Furthermore, the method can be generalised straightforwardly by modifying the surface parametrization and distance definition, making it accessible to non-expert users. In addition, the efficient solver leads to practical data generation for a wide range of vision and metrology software.

Keywords:

Reference data, software verification, machine vision

Measurement Uncertainty and Metrology / 53**Bias and Multiplicity in Producer Risk Estimation under Simplified Conformity Assessment Models****Author:** Omar Jair Purata Sifuentes¹¹ *Universidad de Guanajuato*

This work investigates how accounting for multiple sources of input measurement uncertainty affects the estimation of global risk in manufacturing processes. For example, conformity assessment for Non-Automatic Weighing Instruments involves several quantities of interest, one of which is the producer risk—the probability of incorrectly rejecting a conforming instrument—which is strongly dependent on how measurement uncertainty is treated. In current industrial practice, risk estimation is sometimes based solely on compliance with Maximum Permissible Error (MPE) limits using standard weights. Such an approach neglects additional sources of uncertainty and may distort both producers' and consumers' global risk evaluations. In this study, we contrast this simplified framework with more comprehensive scenarios that explicitly account for standard weight uncertainty and the implementation of guard bands. The results demonstrate that omitting these elements leads to systematic bias. Moreover, the analysis reveals multiple input configurations that yield identical global producer risk, indicating non-uniqueness with respect to process centring. In contrast, no such multiplicity is observed for the global consumer risk. These findings highlight the need for more rigorous uncertainty modelling in conformity assessment to ensure reliable risk quantification.

Keywords:

producer global risk, input multiplicity, conformity assessment

A study of prior sensitivity in random effects models for decisions in interlaboratory comparisons: credible intervals versus Bayesian evidence

Authors: Gerd Wübbeler¹; Séverine Demeyer²

¹ PTB

² Laboratoire National de Métrologie et d'Essais

Random effects models are widely used in interlaboratory comparisons to estimate between-laboratory variability τ and assess degrees of equivalence of laboratories [1]. In this context, decisions are often based on 95% credible intervals [2], while Bayesian hypothesis testing provides an alternative probabilistic framework based on Bayes factors for assessing laboratory effects [3].

This work compares these two decision paradigms, focusing on the sensitivity of conclusions to the prior specification of the between-laboratory standard deviation. We investigate how credible interval based decisions and Bayes factor based hypothesis testing behave under different weakly informative priors on τ , including Half-Cauchy prior, inverse Chi-square and data-informed scale choices.

We further assess the robustness of credible interval and Bayes factor conclusions using posterior predictive simulations [4]. This predictive additional step to standard analysis allows us to examine whether decisions remain stable under replicated data scenarios, and to identify cases where apparent evidence may not be supported by predictive behavior.

Results from simulation studies and real interlaboratory data from CCQM-K53 [5] show that credible interval decisions are relatively stable with respect to prior choices, whereas Bayes factors exhibit strong sensitivity to the prior on τ . Posterior predictive checks provide additional insight to the robustness of the decisions. These findings emphasize the importance of combining inferential and predictive perspectives when making decisions in metrological applications such as laboratory equivalence assessments.

[1] Toman B, Possolo, A, Laboratory effects models for interlaboratory comparisons, *Accred Qual Assur* (2009) 14:553–563

[2] CCQM-KCWG/01, Guidelines for the CCQM KCWG on the review of CCQM CMCs for inclusion in the key comparison database, 2020.

[3] Wübbeler G, Bodnar O, Elster C. Bayesian hypothesis testing for key comparisons. *Metrologia*. 2016:1131-8.

[4] Raghu N Kacker RK Alistair Forbes, Sommer KD. Bayesian posterior predictive p-value of statistical consistency in interlaboratory evaluations. *Metrologia*. 2008:512-23.

[5] Lee J, Lee JB, Moon DM, Kim JS, van der Veen AMH, Besley L, et al. Final report on international key comparison CCQM-K53: Oxygen in nitrogen. *Metrologia*. 2010;47

Keywords:

Bayesian hypothesis testing, Bayes factors, posterior predictive distribution, degrees of equivalence, prior knowledge

Measurement Uncertainty and Metrology / 135

An Interactive Software Application for Gauge R&R Analysis**Author:** Mahmut Onur Karaman¹**Co-author:** Murat Caner Testik¹¹ *Hacettepe University*

Measurement System Analysis is a fundamental element in quality improvement initiatives in manufacturing and is commonly conducted by Gauge Repeatability and Reproducibility (Gauge R&R) studies. However, most of the widely used software tools for Gauge R&R analysis confine the analyst to a restricted design in which the analysis is performed by ANOVA. In such analyses, generally part and operator factors are considered as sources of variability. However, industrial applications sometimes have a different structure and extend beyond this standard template. For instance, measurement variance can be decomposed into other sources of variability, have a nested structure, or include higher order interactions. Forcing such designs into a standard two-factor crossed model can be unsuitable.

This study presents an interactive software application developed to allow practitioners specify different measurement system designs without requiring programming and provide variance component estimates together with Gauge R&R metrics. The user identifies the factors affecting the measurements and classifies each factor by its source. The interface then enables the configuration of nested relationships and interactions. Based on the user's selections, the application builds the corresponding effects model and estimates the variance components using the chosen method, either the classical ANOVA approach or REML. Unlike ANOVA, REML also enables the analysis of unbalanced and incomplete measurement data. A range of diagnostic graphics including main-effect plots, interaction plots, and variance-contribution plots are provided as the visual assessment of each variance component.

Keywords:

Measurement System Analysis, Gauge R&R, Statistical Quality Control, Measurement Uncertainty

Reliability / 35

Statistical Travel Time Inference Under Incomplete Data: A Queueing-Theoretic Approach

Author: Dingyi Wang¹**Co-author:** Qingpei Hu¹¹ *Academy of Mathematics and Systems Science, Chinese Academy of Sciences*

Travel time reliability (TTR) is an important issue in transportation systems, significantly influencing individual decision-making and aggregate travel demand. The inherent uncertainty in travel time is crucial for various applications, which requires the estimation of entire travel time distribution instead of merely the expected travel time. This work proposes a statistical framework for estimating travel time distributions using the Mt/G/∞ queueing model, specifically tailored for scenarios where only interval-censored data are available. A comprehensive joint likelihood function is derived for incomplete data that considers the observed arrival counts and departure counts within each time interval. The model assumes a log-normal distribution for travel times, which effectively captures the characteristic right-skewness and non-negativity of empirical traffic data. The performance of the proposed framework is validated using the NGSIM US101 dataset, a high-fidelity trajectory database from a southbound freeway segment in Los Angeles. Our results demonstrate that the Queueing-Theoretic method can estimate the travel time distribution effectively for incomplete datasets. The findings suggest that the proposed methodology offers a powerful, cost-effective approach for estimating travel time distributions.

Similar approach could be applied for the recovery time inference as well, when the transportation systems operation is interrupted by natural disasters, traffic accidents, or failure of critical equipment, which would be helpful to assess the resilience of the transportation systems.

Keywords:

Travel Time Reliability; Mt/G/∞ Queue; Interval-censored Data.

Reliability / 71

A scheme for the predictive replacement of lithium-ion batteries, using reinforcement learning

Author: Antonio Pievatolo¹

Co-authors: Elisa Varini²; Giovanni Meccariello³

¹ *CNR-IMATI*

² *CNR IMATI*

³ *CNR STEMS*

Based on previous results we obtained, it is possible to establish a replacement decision policy for every pair of charge/discharge cycle index and (discretized) state of health of a lithium-ion battery, using Monte Carlo reinforcement learning to optimize a state-action value function. This function balances the value of using the battery for as long as possible against the loss due to outages caused by its unserviceability. The capacity degradation process of the battery has been described via a Gaussian Process model, which is also learned as new capacity observations are collected. This model is fundamental for the correct computation of the expected value function, because the probability distribution of the future state of health is required for this task.

We expand on this work by applying degradation models with better predictive performance. We also examine uncertainty around the optimal policy and compare our method with more conventional replacement policies.

Acknowledgment: This work was produced with funding from the Italian Ministry of University and Research, under the PRIN 2022 call, assigned with Decree No. 20428 adopted on 06/11/2024. Project: 2022WBN75S - E3DM - Experimental Design and Maintenance, a Decision-Making approach driven by Degradation Models.

Keywords:

Degradation forecasting, reinforcement learning, lithium-ion batteries

Reliability / 30

Variable Selection under Cumulative Exposure Model for Time-to-Event Data with Time-Varying Covariates

Authors: Jiayi Lian¹; Xinwei Deng²; Yili Hong²; Yueyao Wang³

¹ *Wells Fargo*

² *Virginia Tech*

³ *Zhejiang Gongshang University*

In various industrial applications, sensors are widely used to collect the signals for predicting the lifetime of product units or systems.

From a modeling perspective, the signal from each sensor can be considered as a time-varying covariate and the lifetime of units can be considered as the response. In the literature, cumulative exposure models are used to link the lifetime response with time-varying covariates. It is important to identify the useful sensors and estimate their effect for predicting the lifetime of units, which leads to a variable selection and estimation problem. In this work, we consider variable selection under the cumulative exposure model for time-to-event data with time-varying covariates. Specifically, we propose to use a penalized likelihood method to select informative covariates and handle collinearity and redundancy among the sensor data, such that the model estimation and prediction accuracy can be greatly improved.

The proposed approach is demonstrated using both simulations and the NASA jet engine dataset.

Keywords:

Adaptive Elastic-net; Dynamic Covariate; Lifetime Regression; Reliability; Sensor Data

Reliability / 47**Reliability testing of repairable systems available in diverse configuration variants****Author:** Nikolaus Haselgruber¹¹ *CIS Consulting in Industrial Statistics GmbH*

Reliability testing is one of the last and most expensive steps in the development process of a complex technical repairable system. It ensures a certain level of reliability prior to market release and provides the basis for estimation of expected maintenance and warranty costs. Depending on system complexity and diversity of available configurations, current industrial practice is to select some critical variants and carry out the testing for this choice without consideration of transferability of test contributions and lack of testing for less critical variants.

This talk discusses a strategy for empirical reliability testing of repairable systems that are available in diverse configurations. It shows how and to which extent test contributions for different configuration variants can be combined efficiently and considered in stochastic reliability models. An example from the automotive industry illustrates the practical applicability.

Keywords:

reliability, testing, repairable system

Software Session / 69**Optimal design, analysis and optimization in the presence of random effects using Minitab DoE by Effex platform****Author:** Jose Nunez Ares¹¹ *Minitab*

Minitab DoE by Effex platform has been expanded to handle random blocks and complex split-plot structures with up to five levels of difficult-to-change factors. In this talk, we will first explain how random factors are considered when generating an optimal design. Then, we will explain how to assess the trade-off between run size and the quality of competing optimal design candidates. Next, we will present the all-subset model selection algorithm, which has been adapted to handle random factor structures, and we will detail our mixed model estimation procedure. Finally, we will demonstrate how to optimize one or more responses simultaneously using interactive graphs and analytical methods. The presented design generation, analysis, and optimization procedures in the presence of random effects will be illustrated using relevant industrial examples.

Keywords:

Design of Experiments, Random Effects, Split-plot designs

Software Session / 49

JASP for Quality Control: An Open-Source Software for Industrial Statistics

Authors: Julius Pfadt¹; Don van den Bergh^{None}; Frantisek Bartos^{None}; Henrik Godmann^{None}; Jonas Petter^{None}; Johnny van Doorn^{None}; Simon Kucharsky^{None}; Eric-Jan Wagenmakers^{None}

¹ *University of Amsterdam*

Quality control methods such as measurement systems analysis, control charts, capability studies, and design of experiments are central to modern manufacturing and increasingly used in service industries. However, many established software solutions (e.g., Minitab, JMP) are costly or require substantial technical expertise (e.g., R, Python). In this presentation, we introduce the Quality Control module of JASP, a free and open-source statistical software package with a graphical user interface that does not require programming. The module supports a broad range of methods used in industrial practice, including Type 1 gauge studies, Gauge R&R, linearity studies, attributes agreement analysis, classical (Shewhart) and advanced control charts (e.g., EWMA, CUSUM), process capability analysis for normal and non-normal data, and design of experiments (factorial and response surface designs). Because results update in real time as users adjust settings, JASP allows rapid exploration of analyses, assumptions, and model choices. Beyond quality control, JASP includes a broad range of other statistical tools and modules, including general statistics, Bayesian inference, machine learning, reliability, and time series analysis, among many others, making it a flexible platform for the broader analytical needs of quality professionals. We demonstrate how JASP can serve as an accessible and transparent alternative for quality engineers, Six Sigma practitioners, and other industrial statisticians.

Keywords:

Statistical process control, measurement systems analysis, design of experiments, open-source software

Software Session / 58

JMP Developments in JMP 19 and Coming in JMP 20

Author: Chris Gotwalt¹

¹ *JMP Division of SAS Institute*

JMP continues to develop powerful capabilities for statisticians and data scientists in industry. In the session we will demonstrate capabilities in JMP and JMP Pro 19 for Bayesian Optimization of multiple responses, and a new Causal Inference platform that makes establishing causality from observational data easily accessible to non-statistician researchers. We will also give a preview of new developments coming in JMP 20, including Mixed Model Variable Selection, Bayesian Optimization for split plot experiments, as well as other highlights.

Keywords:

JMP, Bayesian Optimization, Mixed Models

Statistical / Stochastic Modelling and Statistical Computing / 24

Least trimmed squares regression with missing values and cell-wise outliers

Authors: Jakob Raymaekers¹; Peter Rousseeuw²

¹ *University of Antwerp, Belgium*

² *University of Leuven*

Regression is the workhorse of statistics, and is often faced with real data that contain outliers. When these are casewise outliers, that is, cases that are entirely wrong or belong to a different population, the issue can be remedied by existing casewise robust regression methods. It is another matter when cellwise outliers occur, that is, suspicious individual entries in the data matrix containing the regressors and the response. We propose a new regression method that is robust to both casewise and cellwise outliers, and handles missing values as well. Its construction allows for skewed distributions. We show that it obeys the first breakdown result for cellwise robust regression. It is also the first such method that is geared to making robust out-of-sample predictions. Its performance is studied by simulation, and it is illustrated on a substantial real dataset.

Keywords:

Casewise outliers; Outlying cells; Prediction; Robust regression; Symmetrization.

Statistical / Stochastic Modelling and Statistical Computing / 59

Characterization of multi-way binary tables with uniform margins and fixed correlations

Authors: Roberto Fontana^{None}; Elisa Perrone¹; Fabio Rapallo²

¹ *Eindhoven University of Technology*

² *University of Genova*

In many applied settings involving binary variables, practitioners typically rely on pairwise measures of dependence, such as correlations or agreement indices. However, when more than two variables are involved, these quantities do not uniquely determine the joint distribution. Instead, they define a family of admissible distributions that share the same pairwise structure while potentially differing substantially in their higher-order interactions.

In this work, we introduce a geometric framework to characterize the full set of joint distributions with fixed pairwise dependence and uniform margins, using the framework of discrete copulas, a recently introduced approach to modeling multivariate dependence. We show that this admissible set forms a convex polytope whose structure can be explicitly analyzed. In particular, we investigate its symmetries and identify its extremal elements, which represent limiting configurations of higher-order dependence consistent with the observed pairwise information.

We highlight the practical relevance of this framework through two motivating examples drawn from medical and psychometric studies, where only partial dependence information is available. These examples illustrate how different admissible joint distributions may lead to substantially different conclusions, despite agreeing on all pairwise measures.

By enabling a systematic exploration of the full admissible set, our approach may be useful for applications in simulation and missing data imputation, where accounting for multiple compatible dependence structures can be important.

Keywords:

Contingency tables, Discrete copulas, Odds ratios

Statistical / Stochastic Modelling and Statistical Computing / 85

Estimating Multivariate Generalized Gamma Convolutions via Kernel Stein Discrepancy

Authors: Esterina Masiello^{None}; Léo Gonin¹; Véronique Maume-Deschamps²

¹ *Institut Camille Jordan*

² *Institut Camille Jordan, Université Claude Bernard Lyon 1*

This work focuses on the estimation of multivariate generalized gamma convolutions (MGGC), a class of distributions widely used in risk modeling for which no closed-form density is available. In practice, only their characteristic functions are known, which makes standard estimation methods such as maximum likelihood inapplicable. To overcome this difficulty, we adopt an RKHS-based approach and use the Kernel Stein Discrepancy (KSD) as an estimation criterion. More precisely, we develop a method to identify a Stein operator for multivariate MGGC from their characteristic function only, leading to a tractable expression of the associated Stein kernel and KSD. We also highlight the link between the Stein operator, the underlying subordinator, and the Lévy measure of GGC distributions. Finally, we illustrate the relevance of the proposed approach through several numerical experiments.

Keywords:

Generalized Gamma Convolutions -RKHS - Stein Operator

Statistical / Stochastic Modelling and Statistical Computing / 113

Tourist Mobility: A Markov Chain Approach Using Origin–Destination Survey Data

Author: Francesca Atzori¹

¹ *university of Cagliari*

This study develops a data-driven Markov chain framework to analyse tourist mobility patterns using empirical origin–destination data collected through surveys at a tourism information point. The dataset records both the municipality visited immediately prior to the survey and the subsequent intended destination, enabling the estimation of transition probability matrices that govern the stochastic evolution of tourist flows. To capture behavioural heterogeneity, separate transition matrices are constructed for two age groups (15–35 and 36–55), and Monte Carlo simulations are performed to examine long-run visitation distributions. The results reveal significant differences in destination preferences across age cohorts and show that the survey location primarily functions as a transit node rather than a final destination. From an economic perspective, the findings provide insights into the spatial allocation of tourism demand and the connectivity structure of local destinations. Identifying high-probability transitions and persistent visitation patterns can support more effective destination management, targeted marketing strategies, and improved allocation of local resources. More broadly, the study demonstrates how partial mobility data can be integrated into a stochastic modelling framework to extract statistically robust and economically meaningful information, offering a flexible tool for the analysis and planning of tourism systems.

Keywords:

Tourist mobility, Markov chains, Tourism flows, Destination management

Adaptive Particle MCMC for non-linear battery state-space models

Author: Edvinas Juozapaitis¹

Co-author: Robertas Alzbutas¹

¹ *Kaunas University of Technology*

State-space models have become core tools in industry as the basis of digital twin technology, enabling online system state monitoring. Advanced Bayesian methods, such as Particle Markov Chain Monte Carlo (PMCMC), may be used for state and parameter inference in non-linear state-space models. The approach combines particle filtering to approximate the hidden state posterior distribution and a Metropolis-Hastings sampler for model parameter proposal. An adaptive PMCMC framework is presented, consisting of multiple stages: first phase sees effective exploration of the parameter-space, achieved by inflating filter hyperparameters which govern particle diversity; during refinement phases posteriors approach stationary distributions as proposal acceptance rate and data subsampling step size are gradually decreased.

Presented methodology is applied to an equivalent circuit model of a lithium-ion battery for assessment of State of Charge and State of Health –battery characteristics useful for predictive maintenance and remaining useful life estimation. Versatility of adaptive PMCMC across different battery designs and charging conditions is confirmed via experiments using real battery cycling data. A comparative assessment against an electrochemical model, typical for lithium-ion battery simulation, reveals significant trade-offs. While the electrochemical model boasts higher interpretability through physical parameters, it suffers from numerical instability and disproportionate sensitivity to inputs. On the other hand, the proposed adaptive PMCMC framework features superior robustness and comprehensive uncertainty quantification at a cost of longer computation times. However, the inherent structure of presented methodology lends itself to parallelization for efficient computing and long-term system degradation monitoring through Bayesian updating.

Keywords:

state-space models; Bayesian inference; lithium-ion batteries

Statistical / Stochastic Modelling and Statistical Computing / 118

Robust graphical modeling under data contamination

Authors: Christian Capezza¹; Antonio Lepore²; Biagio Palumbo³

¹ *Department of Industrial Engineering, University of Naples "Federico II"*

² *Università degli Studi di Napoli Federico II - Dept. of Industrial Engineering*

³ *Università di Napoli Federico II*

High-dimensional data generated by modern multi-sensor systems call for statistical methods able to capture complex dependency structures. Graphical models are a popular tool for this purpose, as they represent conditional relationships between variables through a network. However, classical estimation techniques can be severely affected by the presence of outliers. Traditional contamination models primarily assume rowwise outliers (entirely anomalous observations), whereas cellwise outliers (individually corrupted data entries) are more challenging, as their effects can propagate across variables and severely affect high-dimensional estimation. This work presents a comprehensive evaluation of robust graphical model estimators under both types of contamination. Motivated by the observed limitations, we introduce a novel estimation strategy that integrates robust covariance estimation with sparse precision matrix recovery via adaptive graphical lasso and stability-driven model selection. An extensive simulation study demonstrates that the proposed robust framework substantially improves both structural graph recovery and precision matrix estimation compared to classical methods, particularly under cellwise contamination. An application to a case study on monitoring hourly traffic flow data further illustrates the practical advantages of the method, which provides more interpretable and stable network structures with fewer spurious connections, enabling more effective anomaly detection.

Keywords:

High-dimensional inference, Cellwise outliers, Robust covariance estimation

Statistical Process Monitoring / 99

An entropy-based distribution-free approach for profile monitoring

Author: Praise Obanya¹

Co-authors: Roelof Coetzer¹; Shawn Liebenberg¹

¹ *North-West University*

Profile monitoring is a branch of Statistical Process Monitoring (SPM) that uses statistical methods to identify irregularities in process data. The data is characterized by a profile, or response curve, observed over a given time interval. Profile monitoring consists of two main phases: the first involves defining an in-control (IC) profile, and the second focuses on comparing subsequent profiles with the IC profile to detect departures from normal behavior. In this study, permutation entropy (PE) is proposed as a nonparametric approach for profile monitoring. A general IC profile is obtained using B-spline fitting, while multiple IC profile replications are generated through nonparametric residual bootstrapping. The PE values of the IC profiles are then computed to establish upper and lower control limits for monitoring future profiles. The performance of the proposed entropy-based method is evaluated using simulated data, and the results demonstrate that PE is effective in identifying out-of-control process profiles.

Keywords:

B-spline fitting, control charts, non-parametric residual bootstrapping, permutation entropy

Statistical Process Monitoring / 40

Nonparametric Statistical Process Control for Monitoring Structural Changes in Regional Drought Dynamics

Author: Michele Scagliarini¹

¹ *University of Bologna*

This study develops a statistical process control framework for monitoring drought as a stochastic process characterized by frequency, duration, and severity. The analysis focuses on the Emilia-Romagna region (Italy) and relies on a spatially weighted SPEI-12 index, ensuring a robust and representative aggregation of regional climatic conditions. The main contribution from a statistical process control perspective is the adoption of nonparametric Time-Between-Event and Amplitude (TBEA) control charts, which jointly monitor inter-event times (frequency) and event magnitudes (severity). Two complementary monitoring schemes are implemented: (i) a distribution-free EWMA control chart, designed to detect small and persistent shifts, and (ii) a change-point control chart based on the Kolmogorov–Smirnov statistic, aimed at identifying abrupt distributional changes. A classical Phase I and Phase II structure is employed to estimate in-control parameters and subsequently detect out-of-control signals. The empirical results highlight significant out-of-control sequences in recent years (2017–2018 and 2022–2023) and multiple structural breakpoints, with a clear deterioration in drought conditions starting in the early 2000s. Methodologically, the integration of EWMA and change-point charts enhances diagnostic capability: the former captures gradual trends and persistence, while the latter identifies discrete regime shifts. Overall, the proposed framework extends SPC methodologies to environmental and climate applications, providing an effective tool for continuous monitoring and evidence-based decision-making in drought risk management.

Keywords:

Change Point Detection; Non-Parametric Methods; Regional Climate Analysis; Time-Between-Event and Amplitude

Statistical Process Monitoring / 68

Linear Profile Monitoring of Functional Data

Authors: Antonio Lepore¹; Davide Forcina²; Kamran Paynabar³

¹ *Università degli Studi di Napoli Federico II - Dept. of Industrial Engineering*

² *Università degli Studi di Napoli Federico II*

³ *School of Industrial and Systems Engineering*

Linear profile monitoring assesses the stability of a process described by a linear relationship between a scalar response variable and multiple explanatory variables. When both the response and explanatory variables are functions, this translates into tracking the stability of the underlying functional linear model (FLM). However, unlike the scalar setting, where batches of data points are collected at each sampling stage to construct a profile, this setting provides only a single paired functional observation at each stage, making stagewise estimation of the regression coefficient functions at each sampling stage ill-posed.

To address this challenge, this article proposes a framework for the online monitoring of the relationship between a functional response and multiple functional covariates. An extensive Monte Carlo simulation study compares the performance of the proposed method with a state-of-the-art method, and a case study on monitoring the relationship between lambda-signals and engine state values illustrates its practical applicability.

Keywords:

Profile Monitoring, Functional Data Analysis, Statistical Process Control

Statistical Process Monitoring / 1

When a Cusum Stops, What Confidence Is There that the Alarm Is Not False?**Author:** Moshe Pollak¹¹ *hebrew university of jerusalem*

In the framework of the Cusum procedure, the evolution of a false alarm has a well-understood stochastic behavior. So, if observations preceding an alarm were to exhibit a behavior that is significantly different, there would be reason to reject the hypothesis that the alarm is false.

We develop a test of this difference. The method is applied to detecting a change in a Covid-19 context involving a possible increase of a mean and in a context involving a possible increase in the probability of a rare event.

Keywords:

control chart; p-value; Covid-19

Statistical Process Monitoring / 6

On the Design and Performance of a Control Chart for Monitoring Continuous data in (0,1) when Process Parameters are Unknown**Author:** Athanasios Rakitzis¹**Co-authors:** Nirpeksh Kumar²; Subhabrata Chakraborti³¹ *University of Piraeus, Department of Statistics and Insurance Science*² *Department of Statistics, Institute of Science, Banaras Hindu University, Varanasi*³ *Department of Information Systems, Statistics and Management Science Culverhouse College of Commerce, University of Alabama*

In this work, we consider monitoring continuous data in the unit interval and investigate the statistical design and performance of a two-sided Shewhart chart when the process parameters are unknown. The most common distribution assumed for such data is the Beta distribution. Although control charts based on the Beta distribution have been studied by several authors, the case of estimated parameters, which is the most practical case, has not been considered in much detail in literature. The chart's performance is investigated in a Monte Carlo study, and empirical rules are provided regarding the size of the Phase I sample. Also, we explore the effectiveness of possible adjustments to the control limits of the chart, which take into account the size of the available Phase I sample data. The performance of the chart is also investigated for several out-of-control situations. The results show that for Phase I samples of small to moderate size, practitioners need to choose between guaranteed in-control performance or improved out-of-control performance. A numerical example based on real data is also provided.

Acknowledgement: This work has been partly supported by the University of Piraeus Research Center.

Keywords:

Beta distribution, Conditional Average Run Length, Maximum Likelihood Estimation

Statistical Process Monitoring / 95

Some Stylized Facts of the Conditional Expected Delay (CED)**Author:** Sven Knoth¹¹ *Helmut Schmidt University Hamburg, Germany*

The popular zero-state average run length (ARL) is just the mean of the random run length, which is the core element of a control chart. However, more appropriate measures for evaluating the detection power make use of the conditional expected delay (CED), which is the mean of the detection delay for a given change point position $\tau = 1, 2, \dots$ under the condition that no false alarm was triggered. Originally, it was only a prerequisite for building sophisticated delay measures. Aiming to optimize the latter, interesting patterns of the CED series were found. While the more theoretical strand of monitoring research struggles with some open optimality problems, did the more applied one not recognize the CED stylized facts so far. This talk tries to close the gap. More importantly, it emphasizes and affirms the importance of an appropriate CED analysis.

Keywords:

ARL, CED, control chart

Statistical Process Monitoring / 124

Self-Starting Control Charts for Few-Shot In Situ Monitoring in Customized Manufacturing**Authors:** Bianca Maria Colosimo¹; Giovanna Capizzi²; Marco Grasso³¹ *Politecnico di Milano*² *Department of Statistical Sciences*³ *Politecnico di Milano, Department of Mechanical Engineering*

Additive manufacturing processes are increasingly characterized by high customization, small batch sizes, and limited availability of historical data, making traditional statistical process control approaches difficult to apply. This work proposes a self-starting monitoring framework for few-shot additive manufacturing environments, enabling effective process monitoring from the earliest production stages without requiring large calibration datasets.

The methodology focuses on the *in situ* monitoring of powder bed fusion processes through the layer-wise analysis of the maximum geometrical deviation between reconstructed and nominal geometries. Since the monitored statistic follows an extreme-value behaviour, the proposed control scheme is developed under a Gumbel-distributed setting, extending self-starting control chart approaches beyond the standard Gaussian assumption.

A further contribution of the work is the introduction of a strategy to estimate the process transient phase and identify the transition toward steady-state conditions, improving monitoring effectiveness during process start-up. Results from a real industrial case study in additive manufacturing demonstrate the effectiveness of the proposed framework for online detection of geometrical defects using layerwise images.

Keywords:

Self-starting control chart, Additive Manufacturing, In situ monitoring, Image data mining

Statistical Process Monitoring / 144

Practical challenges and statistical solutions for process monitoring and control in dairy production**Author:** Ewa Szymanska¹**Co-authors:** Daniel Florea¹; Edward Suyata¹; Diaa Al Mohamad¹; Huong Ho Thanh¹; Yulianto Hartono¹¹ *FrieslandCampina*

Industrial manufacturing processes are complex, often including many known and unknown factors like various raw materials and their properties, multiple production lines, hundreds of process variables and different product quality parameters. Using data analysis to understand, optimize and monitor and control several parts of these processes is a common practice with big proven value. However, not all data-driven projects deliver fast and long-lasting results. It is because of challenges like limited data traceability, low data quality and ever-changing process conditions, raw materials and industrial practices.

In this presentation we would like to discuss few challenges and statistical approaches used on recent project focusing on monitoring and controlling production process for one of dairy products. First challenge was using parallel production lines with similar but not identical settings in the production of different product batches and sometimes using two different lines to produce the same one batch. Another challenge was complex relationships between product quality parameters for fresh products and shelf-life products that need to be modelled and understood. Moreover, using different types of raw materials for the same product with limited knowledge on relationship of their properties with final product quality was also challenging. Historical dataset including data on raw material, process and product was collected and analyzed with different methods and approaches including correlation analysis and PLS regression models. Results are compared and method's performance evaluated and statistically validated to select the best strategy and provide the best insights on critical-to-quality process and raw material variables for process monitoring and control.

Keywords:

industry 4.0., critical-to-quality parameters, process control

Statistics and data science in the technological field: current issues and new proposals / 18

ADeS: A Flexible Adaptive Design Strategy for Ethical and Covariate-Balanced Clinical Trials

Author: Marco Novelli¹

¹ *University of Bologna*

Adaptive randomization in clinical trials often requires balancing competing goals: improving patient benefit, preserving statistical efficiency, and maintaining adequate randomness in treatment assignment. We propose the Adaptive Design Strategy (ADeS), a flexible group-sequential framework that unifies covariate-adaptive (CA), response-adaptive (RA), covariate-adjusted response-adaptive (CARA), and hybrid RA+CA/CARA+CA designs within a single objective-based formulation. At each interim step, ADeS selects treatment allocations for the incoming patient group by minimizing a composite criterion that combines ethical or response-adaptive targets with covariate balance. The optimization is performed through a simulated annealing engine, while an acceptance-randomization step preserves a controlled level of randomness in the implemented assignments. The framework accommodates multiple treatments, different outcome types, and arbitrary baseline covariates. In particular, predictive components can be specified through either parametric or non-parametric models; in our implementation, Bayesian Additive Regression Trees are used to capture complex treatment-covariate interactions. For finite stratified covariates, we establish strong consistency of the resulting stratified estimators. Extensive simulation studies with homogeneous and heterogeneous treatment effects show that ADeS achieves a favorable trade-off between ethical allocation and inferential performance, while remaining more flexible than existing adaptive procedures.

Keywords:

Biomarkers; Covariate-Adaptive procedures; Covariate-Adjusted Response-Adaptive designs; Inferential efficiency; Precision medicine; Treatment by covariate interactions

Statistics and data science in the technological field: current issues and new proposals / 134

Online Quality Monitoring in Selective Laser Melting via Image-Based Statistical Methods

Author: Panagiotis Tsiamyrtzis¹

¹ *Politecnico di Milano*

A major challenge in Additive Manufacturing (AM) is the development of reliable in-situ and online quality monitoring methodologies. Visible and infrared cameras can provide near real-time image data that can be exploited for anomaly detection through Statistical Process Control and Monitoring (SPC/M) methods.

This work investigates image-based monitoring methods for Selective Laser Melting (SLM) processes, aiming to detect shifts from the in-control (IC) to the out-of-control (OOC) state. Two approaches are compared: a partial first-order stochastic dominance methodology and generalized multilinear models for sufficient dimension reduction with tensor-valued predictors. In addition, a hybrid approach combining elements of both methodologies is proposed.

The methods are evaluated using simulated datasets generated from images of a real SLM process, with emphasis on monitoring performance and sensitivity to training sample size. The results highlight the potential of statistically grounded, data-efficient image monitoring methods for next-generation smart manufacturing systems.

Keywords:

Non-Parametric, Sufficient Dimension Reduction, Statistical Process Control and Monitoring (SPC/M)

Statistics and data science in the technological field: current issues and new proposals / 54**Fisher's Principles of Experimental Design and Why They Are Still Important****Author:** Geoff Vining¹¹ *Virginia Tech Statistics Department*

Sir Ronald Fisher was one of the giants in both the fields of statistics and genetics. His seminal work was done at Rothamsted Experimental Station outside of London.

Fisher was both a statistician and a scientist. He was well trained in mathematics, but in the final analysis, he was a scientist who understood how to perform proper data analysis. Fisher strongly rejected the Neyman-Pearson approach based on the strict distributional assumptions required to perform hypothesis testing. Fisher understood that the assumption that the data are independent and identically distributed does not exist for real-world data.

This talk summarizes his three basic principles for conducting experiments best understood as randomization, replication, and local control of error (often referred to as "blocking"). This talk presents why these principles are extremely important through simple examples, especially randomization. Fisher understood that real data are always correlated, often highly correlated. Randomization provided a mechanism for making fair comparisons among the various treatments under study.

This talk then discusses how to adapt these principles for a large data universe, something that did not exist when Fisher was alive, much less when he was at Rothamsted. In the current big data universe, all of the data are highly correlated. His fundamental approach, properly modified, provides a formal basis for the analysis of complex data that corrects for the complexity of the models used to explain the data.

Keywords:

mathematical statistics versus scientific data analysis, deterministic data, complex model identification

Statistics in Industry, Business and Finance / 57**Lot Heterogeneity in Statistical Quality Control: An enhanced Optimization Approach for the Design of Rectifying Sampling Plans****Author:** Jaqueline Bojer¹**Co-authors:** Horst Lewitschnig²; Vera Hofer¹¹ *University of Graz*² *Infineon Technologies Austria AG*

In rectifying sampling inspection, a lot is subjected to full inspection if the number of defects in a random sample exceeds a predefined acceptance criterion. Traditional models assume a constant probability of being defective p for all items within a lot. In contrast to this, we consider heterogeneous lots in which individual items can have different probabilities of being defective, to better reflect practical situations such as when items are produced on different machines or handled by different operators. Moreover, the discussed model accounts for possible dependencies among the individual defect probabilities by incorporating correlations.

We analyze the key performance metrics, Average Outgoing Quality (AOQ) and Average Total Inspection (ATI), in heterogeneous settings and under different correlation scenarios. Based on the findings, we propose an enhanced optimization approach for the design of rectifying sampling plans.

Keywords:

Rectifying sampling inspection, Heterogeneous lots, Sampling plan design

Stacking Evidence across tests in R&D

Author: Jan-Willem Bikker^{None}

CQM is a consultancy company with over four decades of experience in industrial R&D projects. One of its long-standing customers has developed consumer products for many years and seeks to reduce the test effort and improve decision making in development projects for a certain class of products. In these development projects, different types of tests are performed on prototype designs, from A (cheap) to D (expensive). The relatively cheap tests A, B assess the prototype early on in the projects, and the more expensive tests C, D involve a trained panel of assessors for confirmation. CQM co-develops with the R&D department a method for stacking evidence, which allows predictions of outcomes of expensive tests (C or D) conditional on observed test results, typically the cheaper ones. These predictions help in deciding to stop the project or form a prior in a Bayesian analysis of future test D. As a consequence, the approach is expected to reduce overall costs for the expensive tests, and has motivated substantial investment in its development.

The model is based on historical development projects, where typically the 4-vectors of test results (A,B,C,D) have missing entries in historical dataset. The approach uses Bayesian inference for a multivariate model of test results (A,B,C,D), capturing conditional dependencies as in Bayesian networks, and incorporating measurement models akin to structural equation models. Strong priors based on expert opinion are needed to complement scarce data, but the multivariate nature poses challenges. In addition to the technical modelling aspects, I will discuss practical learnings from employing Bayesian analysis in an R&D organisation, including communication, prior specification, and the gradual development of statistical intuition.

Keywords:

Bayesian statistics, Bayesian network, SEM

A Requirements-Driven Lifecycle Digital Twin Methodology for Innovative SMRs: Architecture Integration and Expected Benefits

Author: Jihyeon Kim¹

Co-author: Taecheol Park¹

¹ *Innovative Small Modular Reactor Development Agency*

Innovative Small Modular Reactors (i-SMRs) introduce fundamentally different design characteristics compared with conventional large nuclear power plants, including integral reactor configurations, compact containment, multi-module deployment, extended fuel cycles, and highly automated operation. In particular, i-SMRs are designed for coordinated multi-module operation, requiring integrated monitoring and control of multiple reactor modules to achieve safe and efficient plant operation. Under the Industry 4.0 paradigm, these unique operational and structural characteristics make it essential to adopt a digital twin framework built upon a structured Requirements Management (RM) system and system architecture.

To address these challenges, this study proposes a lifecycle-oriented methodology for integrating digital twin technology into i-SMRs. Unlike conventional nuclear power plants, where digital twins are mainly introduced during the operational phase, the proposed methodology establishes a requirements-driven digital twin architecture from the early design phase and continuously utilizes it throughout construction, commissioning, operation, and maintenance. By incorporating structured requirements engineering and domain element models from inception, the framework links physics-based simulation models with real-time operational data to establish a highly reliable, unified virtual representation of the plant. This enables high-precision design verification, multi-module operational support, operator decision support, and predictive maintenance across the entire plant lifecycle.

To support key milestones of the i-SMR project—securing Standard Design Approval (SDA) by 2028 and completing the First-Of-A-Kind (FOAK) plant by 2035—this study presents a methodology for integrating a structured Requirements Management (RM) system and system architecture into the digital twin framework. By directly embedding technical requirements and regulatory guidelines into the virtual plant from inception, the proposed approach enables rigorous design verification and validation (V&V), minimizes engineering and licensing risks, and ensures long-term operational reliability across the plant lifecycle.

Keywords:

i-SMR, Digital Twin

Statistics in Industry, Business and Finance / 46**A Scalable Analytics Platform for Enterprise Level Quality Management****Author:** Manuel Martin¹**Co-authors:** Chiang Leo ; Wang Zhenyu¹ *Lubrizol*

Statistical Quality Control (SQC) plays a critical role in modern manufacturing processes by enabling early detection of process deviations, quantification of variability, and consistent delivery of products that meet customer specifications. To keep up with evolving customer demands and market dynamics, manufacturing networks grow in scale and complexity. Classic SQC workflows struggle to provide timely, consistent, and actionable insights across sites, materials, and suppliers etc. This creates a need for scalable digital platforms that embed rigorous statistical methods into everyday quality management.

This paper presents an enterprise level analytics platform designed to modernize SQC and process capability analysis across raw materials, intermediates, and finished products. This SQC platform integrates quality control data into a centralized Power BI environment, enabling standardized yet flexible analysis of process capability indices (Cpk), control charts, and statistical process control (SPC) rule violations. The platform supports granular slicing by plant, material, characteristic, and supplier, while maintaining a consistent analytical framework aligned with corporate quality metrics.

The system has been successfully applied to routine quality monitoring, cross site benchmarking, and identification of improvement opportunities, supporting both operational teams and leadership decision making. By modernizing SQC workflows and embedding advanced analytics into daily quality management, the presented SQC system demonstrates how scalable digital platforms can enhance process understanding, consistency, and continuous improvement in large scale manufacturing environments.

The next steps involve optimizing data flow to boost refresh rates, bringing the tool towards real-time operation to alert QC managers for process deviations.

Keywords:

SQC QM

Adaptive soft sensor for product quality estimation in clinker production

Authors: Mihnea Stefan¹; Fabrizio Bezzo¹; Wilson Ricardo Leal da Silva²; Pierantonio Facco¹

¹ *University of Padova*

² *Fuller Technologies, Green Innovation*

The real-time estimation of key quality variables remains a critical challenge in industrial environments due to the limited availability of direct measurements and the presence of complex, dynamic process behavior. This work proposes an adaptive soft-sensing framework for the estimation of cement quality in clinker production, where quality indicators are traditionally measured through costly and infrequent laboratory analyses. The proposed framework builds upon Partial Least Squares (PLS) regression, extending it to address nonlinearity, process dynamics, and non-stationarity through a combination of recursive updating, lagged-variable modelling, and local learning strategies.

In particular, two complementary adaptive modelling strategies are investigated. The first is based on a Quasi-Ensemble PLS approach, in which multiple models with different hyperparameter configurations are combined to enhance estimation accuracy against model uncertainty. The second strategy proposes an autonomous soft sensor capable of self-adapting to time-varying plant conditions by integrating recursive PLS modelling with Bayesian optimization for the real-time tuning of hyperparameters, enabling continuous adaptation while preserving model interpretability and computational tractability.

The methodologies are validated on industrial data from cement production plants, demonstrating accurate predictive ability and robustness compared to state-of-the-art soft sensors. Furthermore, the proposed framework offers a general and scalable solution for adaptive soft-sensing in complex industrial systems, with potential applications beyond the cement industry.

Keywords:

soft sensor, virtual sensor, PLS, cement

The Lack of Impact of Lean Six Sigma in Textiles, Apparel and Other Basic Industries

Author: A Blanton Godfrey¹

Co-author: Alireza Vahedi Fakrh¹

¹ *North Carolina State University*

When we began researching the impact of Lean Six Sigma in the textile and apparel industries, we set out to identify success factors, understand differences in implementation strategies, and explore how applications varied by business type. We expected implementation to be strongest in segments with demanding customer requirements (for example aerospace, automotive, and medical textiles), moderate in technical operations (spinning, weaving, and dyeing and finishing), and perhaps weakest in apparel and household furnishings segments where style, fashion, and price dominate over quality.

Our findings largely confirmed these expectations, but the reality was more troubling than we anticipated. Textile companies face strong pressure from two very different types of powerful outside forces. On one side, customers in highly regulated industries such as aerospace, automotive, and medical impose their own strict quality standards—including FDA regulations and pharmacopoeia requirements—and textile suppliers must comply with these standards to keep their business. On the other side, large fashion retailers squeeze apparel and household furnishings suppliers by cutting prices and shortening lead times without offering any support for improvement. The automotive industry works well with Lean Six Sigma because production is standardized and suppliers are treated as long-term partners. Textile and apparel manufacturing is the opposite—products change constantly, fast fashion leaves no time for improvement, and suppliers are chosen on price rather than trust or collaboration. Many companies in this sector are small and medium-sized enterprises.

While we did find isolated successes, the overwhelming pattern was one of partial adoption, short-lived efforts, and disappointing returns. We propose a set of revised implementation strategies—ones that account for the unique structural, cultural, financial, and educational realities of this industry—with the potential to deliver far more meaningful and lasting results.

Looking ahead, the integration of artificial intelligence with Lean Six Sigma represents both a significant opportunity and an urgent strategic challenge for the textile industry. AI-powered tools have the potential to accelerate and democratize key LSS capabilities—enabling real-time process monitoring, predictive quality control, automated root cause analysis, and data-driven decision-making at a scale and speed that traditional statistical methods alone cannot achieve. For industries historically hampered by weak data infrastructure and limited analytical capability, this convergence could lower the barrier to meaningful LSS adoption. However, realizing this potential requires deliberate preparation: companies must begin building the digital foundations, data literacy, and cross-functional capability needed to operate in an AI-augmented improvement environment. This is not simply a technology investment—it demands a new generation of improvement professionals who are equally fluent in lean thinking, statistical reasoning, and AI-enabled analytics. The industry that fails to prepare for this evolution risks falling further behind, while those that embrace it strategically may finally unlock the sustained LSS impact that has so far remained out of reach.

Keywords:

Lean Six Sigma

Statistics in Industry, Business and Finance / 121

Understanding Tourism Demand: A Tree-Based Analysis of Territorial Drivers

Authors: Andrea Carta^{None}; Francesca Atzori¹; Marco Ortu¹**Co-author:** Giulia Contu¹¹ *University of Cagliari*

We investigate the territorial and structural drivers of tourism arrivals, with the aim of supporting more effective destination planning. Identifying the factors that attract tourists is essential for designing policies that balance development and sustainability across regions.

Our analysis applies a Multivariate Regression Tree (MRT), a flexible and interpretable method that allows to uncover both key determinants of tourism demand and homogeneous groups of destinations. Compared to traditional regression approaches, MRT captures complex interactions and non-linear relationships, making it particularly suitable for heterogeneous territorial contexts.

The empirical study focuses on 307 municipalities in Sardinia (Italy). Tourist arrivals, disaggregated by nationality, are considered as multiple response variables. Explanatory variables include geographic characteristics, demographic structure, and transport accessibility.

Results indicate that the density of accommodation facilities is the main driver of tourism attractiveness. Additionally, a clear coastal–inland divide emerges: coastal areas benefit from a stronger combination of accommodation supply and accessibility, leading to higher tourist concentration. Inland areas, in contrast, show more limited performance, even when other favorable characteristics are present.

These findings highlight how interactions between infrastructure and territorial features shape tourism demand across different international markets. The study illustrates the value of tree-based methods for applied statistical analysis in tourism and suggests their broader applicability to other regional planning contexts.

Keywords:

Multivariate Regression Tree, Tourism Arrivals, Destination Attractiveness

Statistics in Industry, Business and Finance / 142

Gastronomic Tourism in Italy: Measuring Tourist Attitudes and Identifying Market Segments

Author: Francesca Bassi¹¹ *University of Padova*

Gastronomic tourism has become an increasingly important component of contemporary tourism demand, particularly in Mediterranean countries such as Italy, where food culture and local traditions are deeply embedded in territorial identity. This study investigates the importance assigned to gastronomy during travel and identifies different profiles of gastronomic tourists in Italy. The analysis combines two complementary methodological approaches. First, a composite indicator was developed to measure the overall involvement with gastronomic tourism and its main dimensions. Second, a mixture confirmatory factor analysis model was estimated to identify tourist segments characterized by different attitudes towards gastronomy. The results show that gastronomy represents a relevant determinant of destination choice and travel satisfaction for Italian tourists. Segments differ substantially in terms of motivations, perceptions, and socio-demographic characteristics. The findings provide important managerial and policy implications for destination marketing and sustainable tourism development.

Keywords:

gastronomic tourism; tourist segmentation; composite indicators; mixture confirmatory factor analysis; tourist behavior; food tourism; Italy

Statistics in Pharma / Healthcare: Pharmaceutical Process Modelling and Manufacturing Analytics / 25

Estimating Stabilization Time in Continuous Manufacturing

Author: Chellafe Ensoy-Musoro¹

¹ *Johnson & Johnson*

This work develops a statistical framework to estimate stabilization time for core tablet batches produced by continuous roller compaction. Using simulated data from multiple production runs of a representative product, the project focuses on characterizing concentration stability during start-up and identifying an optimal start-up duration to guarantee product quality and consistency. The current approach employs a two-component segmented linear model (broken-stick) implemented via the `gslns` package in R to separate a transient phase of rapid concentration change from a subsequent stable phase. Although this model captures overall process behavior, practical application has revealed limitations: stabilization-time estimates and their confidence intervals are sometimes unacceptably large, reducing interpretability for operational decision-making. These issues likely arise from sparse transient-phase sampling, within-batch variability, and model inflexibility for complex dynamics. To address these shortcomings, alternative techniques will be investigated with a focus on robustness and precision of estimated change times.

The two-component linear (broken stick) model is given as:

$$Y_i = \beta_1 + \beta_2 \left(\left(1 - \frac{1}{1 + \exp(-\gamma(\text{time}_i - \tau))} \right) \text{time}_i + \left(\frac{1}{1 + \exp(-\gamma(\text{time}_i - \tau))} \right) \tau \right) + \epsilon_i$$

Where: Y_i is the concentration at time i , γ is the numeric smoothness of transition between linear model components. Higher value gives a sharper changepoint. β_1, β_2, τ are the parameters of the model which defines the components of interest (see below), $\epsilon_i \sim N(0, \sigma^2)$ is the residual error with σ^2 variance.

Keywords:

stabilization, broken-stick model, continuous manufacturing

Statistics in Pharma / Healthcare: Pharmaceutical Process Modelling and Manufacturing Analytics / 55

Incorporating Factor Variability into Risk Estimation in Pharmaceutical Manufacturing

Author: Olympia Tumolva¹

¹ *Johnson & Johnson*

In pharmaceutical manufacturing, process optimization and control are critical for ensuring consistent drug product quality and regulatory compliance. In real-world applications, certain study factors are treated as fixed and maintained constant to standardize production conditions. However, it is sometimes inevitable for other factors to exhibit variability, introducing uncertainties that can affect the final drug product. This study presents a statistical framework designed to operate under conditions where variability in the factors, are present. The focus is on incorporating this variability into risk estimation, specifically in evaluating the probability of failing to meet the specifications of a critical quality attribute of the drug product. This study offers valuable insights for practitioners in the pharmaceutical industry, aiming to enhance product quality and ensure compliance with regulatory standards.

Keywords:

risk estimation, variability, pharmaceutical manufacturing,

Statistics in Pharma / Healthcare: Pharmaceutical Process Modelling and Manufacturing Analytics / 141

Multi-attribute modelling for mRNA specification setting

Author: Marco Mariti¹

¹ GSK

A key challenge for mRNA vaccine and medicines development is represented by mRNA degradation under normal refrigerated storage condition (2-8°C). Product evolution over time is primarily driven by mRNA Integrity degradation, occurring mainly through chemical hydrolysis.

mRNA molecules are known to be particularly susceptible to hydrolysis under alkaline conditions, where degradation via backbone hydrolysis is accelerated. Consequently, understanding the impact of pH on degradation kinetics is critical for defining an appropriate control strategy.

A systematic investigation of pH effects on Integrity evolution was performed, to identify appropriate pH specification at release, assessing its impact both at release and in stability.

An enhanced Sestak-Berggren model was developed, including both temperature and pH effects to predict mRNA integrity degradation over time. A simulation approach was then developed to establish a quantitative linkage between pH and Integrity specifications at release and to compute pH impact on Integrity risk of out of specification at end of shelf-life. The results provide a scientifically justified framework to support definition of pH and Integrity product specifications.

Keywords:

Stability modelling, Specification setting, Simulation

Statistics in Pharma / Healthcare: Statistical Decision Support in Healthcare and Biomedical Development / 130

Towards an Adaptive Framework for Longitudinal Survival Modeling in Clinical Decision Support: Case of the ISPY 1 Trial

Author: Abdelaziz Berrado¹

¹ Mohammed V University, EMI

Clinical surveillance of cancer patients is necessary to ensure early detection of recurrence after curative treatment and to monitor patient progression. Although clinical guidelines commonly recommend fixed surveillance schedules for all patients, static intervals between follow-up visits may not be compatible with individual disease progression, increasing the risk of delayed recurrence detection for some cases as it may be burdensome for healthcare resources and patients for others.

Motivated by the idea of personalizing follow-up schedules based on each patient's prognosis, we introduce a generic framework for an intelligent cancer post-treatment follow-up. The framework enables the prediction of time to recurrence using patient-specific clinical longitudinal data along with survival modeling, which can be used to customize follow-up schedules for each patient.

The framework can enhance medical visits planning by incorporating evidence-based follow-up suggestions aligned with each patient's clinical development, thereby reducing under-surveillance risk, optimizing resource use and easing patient burden.

We illustrate the application of the proposed framework to a cohort of patients with locally advanced breast cancer undergoing neoadjuvant chemotherapy from the I-SPY 1 trial, a multicenter prospective

study conducted across nine institutions and publicly available.

The approach, while validated for breast cancer, applies to other cancer types and diverse data modalities. Its application can support doctors in predicting relapse time, establishing optimal follow-up intervals, and enabling prompt treatment decisions, so improving the efficiency, efficacy, and tailoring of post-treatment care.

Keywords:

Longitudinal survival modeling, customized Patient follow-up, ISPY 1 trial

Statistics in Pharma / Healthcare: Statistical Decision Support in Healthcare and Biomedical Development / 147

Predicting Study Success in Diagnostic Test Evaluation: A Case Study Using Frequentist and Bayesian Approaches

Author: Alejandro Moreno Muñoz¹

Co-author: Ignasi Puig de Dou²

¹ *Werfen*

² *Universitat Politècnica de Catalunya*

In diagnostic test evaluation, agreement between a candidate method and a reference method is sometimes assessed using a composite comparator method when no gold standard is available. A composite comparator combines the results of multiple assays to derive a final adjudicated classification of the patients.

This work addresses the prediction of study success after a partial testing of an initial subset of the study population. Using a case study of an immunoassay In Vitro Diagnostic (IVD) test evaluation, we leverage the results of the initial testing to predict the outcomes in the remaining samples and quantify the probability of meeting predefined agreement criteria for the positive and negative percent of agreement (PPA/NPA).

Two complementary frameworks are considered: a frequentist approach based on confidence intervals for transition probabilities, and a Bayesian approach using beta-binomial posterior predictive distributions under different prior assumptions. Sensitivity analyses compare pessimistic, likely, and optimistic scenarios for the two frameworks. Additionally, a total probability framework computes the posterior predictive probability of achieving target performance criteria across all possible outcome combinations.

The methodology is illustrated for additional study settings to evaluate the robustness of the approaches, assessing the impact of sampling design, prior assumptions, and uncertainty quantification on decision-making in diagnostic development.

Keywords:

Diagnostic test evaluation, decision-making under uncertainty, Bayesian inference, frequentist methods, sensitivity analysis

Statistics in Pharma / Healthcare: Statistical Modelling and Learning for Biomedical Systems
/ 96

Kinetic Models by Physics-Informed Statistical Learning Methods

Author: Marco Galliani¹

Co-authors: Bernard Francq²; Laura Sangalli¹

¹ *Politecnico di Milano*

² *GSK*

Accelerated stability studies are a fundamental tool in the development of pharmaceutical and vaccine products, enabling the prediction of long-term stability from short-term experiments conducted under stressed conditions. In this context, kinetic models—such as the Šesták–Berggren formulation combined with Arrhenius-type temperature dependence—are widely used to describe degradation processes. However, reliably estimating such models from limited, noisy data remains a challenging problem.

This work proposes a physics-informed statistical learning framework that integrates mechanistic knowledge, expressed through ordinary differential equations (ODEs), with experimental observations. The approach is formulated as a regularised regression problem, where a data-fitting term is combined with a penalty enforcing consistency with a parametrized dynamical system. This formulation enables a flexible balance between data-driven modelling and adherence to physical laws, improving robustness in data-scarce settings.

To address the resulting estimation problem, we adopt a hierarchical strategy that separates the reconstruction of the system trajectory from the estimation of the kinetic parameters, allowing for efficient computation. The proposed estimator is a flexible tool that allows accurate estimation of ODE parameters when the ODE model is well-specified, while still guaranteeing good predictive performance in the case of a misspecified regularizing model.

The methodology is validated through simulation studies and applied to real data from accelerated stability experiments on vaccine antigenicity decay. Results demonstrate improved predictive performance and more reliable parameter estimation, highlighting the potential of physics-informed approaches for kinetic modelling in pharmaceutical applications.

Keywords:

physics-informed nonparametric regression, inverse problems, kinetic modeling

Statistics in Pharma / Healthcare: Statistical Modelling and Learning for Biomedical Systems
/ 126

Bayesian network-based Tikhonov MRI reconstruction

Authors: Sébastien Marmin¹; Christoph Kolbitsch²; Andreas Kofler²

¹ *Laboratoire national de métrologie et d'essais*

² *PTB*

Uncertainty quantification is essential for assessing the reliability of MRI reconstructions. The Network-based Tikhonov reconstruction method was demonstrated to produce excellent results in accelerated multicoils settings. However, in challenging low-field settings, where noise is high and scanners are single coil, this scheme requires a careful evaluation of its reconstruction uncertainty. In this work, we study the Bayesian network-based Tikhonov reconstruction scheme and derive an upper bound on the posterior expected reconstruction error. Starting from the posterior mean-squared error, we separate the error into a variance term, given by the posterior covariance trace, and a bias term measuring the discrepancy between the posterior mean and the true image. By introducing a fixed network-based image prior and a worst-case prior mismatch, we obtain a tractable bound that depends on the posterior variance, the deviation of the reconstruction from the prior, and a bounded prior-to-truth error. We further derive pixelwise bounds, providing uncertainty map estimates.

We will present further experiments showing that for these challenging reconstruction conditions, the CNN-based Tikhonov regularization, yielding a linear reconstruction scheme, is too simple and relies too heavily on the prior, motivating the need to use more sophisticated reconstruction schemes like plug-and-play methods.

This analysis offers a simple theoretical framework and an analytical error bound for interpreting reconstruction quality in low-field MRI. It highlights the trade-off between data consistency and prior dependence for the reconstruction and its uncertainty.

Keywords:

Uncertainty quantification, MRI reconstructions, imaging

Statistics in Pharma / Healthcare: Statistical Modelling and Learning for Biomedical Systems
/ 139

Statistics Can Go Farther

Author: Dibyen Majumdar¹

¹ *University of Illinois at Chicago*

A pre-clinical trial for the treatment of cancer, with data analyzed by a traditional statistical method with random coefficient model, showed the significance of the treatment based on purified bacterial redox protein. A data analysis performed later, based on quantile regression, went beyond. It demonstrated the remarkable specificity of the effectiveness of experimental treatment.

Keywords:

Pre-clinical trial, Traditional regression, Quantile regression

Bayesian Optimization as Design of Experiments: The Entropy Connection

Author: Arno Strouwen¹

¹ *Strouwen Statistics; PumasAI; KULeuven*

Bayesian optimization (BO) and classical design of experiments (DOE) are rarely taught together. DOE courses cover factorial designs, response surface methodology, and optimality criteria like D and I-optimality. BO courses cover Gaussian processes and acquisition functions like expected improvement. The two communities use different notation, different software, and publish in different journals. Yet both address a closely related problem: how to allocate a limited experimental budget to learn what matters most.

This talk grew out of the Experimental Design course in the Master of Statistics and Data Science at KU Leuven. It starts from a classical DOE lecture and extends it with modern Bayesian optimization material, showing how both fit within a common experimental-design framework. The key message is that every experimental design involves three choices: a model for the response, an objective that defines what “good” means, and a rule for choosing additional runs. In familiar DOE settings, that rule may produce an initial design and then augment it to improve parameter estimation or prediction, often using polynomial models and criteria such as D- and I-optimality. BO makes different choices within the same template: a Gaussian process prior instead of a polynomial model, an objective focused on the optimum rather than all coefficients, and an adaptive rule for adding runs as data are collected.

The bridge between the two is entropy. D-optimality maximizes $\log \det(\mathbf{X}^\top \mathbf{X})$, which for Gaussian linear models is equivalent to maximizing the information gained about the regression coefficients. Information-theoretic BO acquisition functions such as Max-value Entropy Search do the same, but for the location or value of the optimum rather than for the model parameters. The principle is shared; what changes is what we want information *about*.

This perspective has practical value for teaching and consulting. It gives a DOE-trained audience a principled way to understand what is new in BO and what is not: the surrogate model, the objective, and the rule used to augment the design change, but the underlying logic of designing informative experiments remains. Framing BO in this way helps students and practitioners extend familiar DOE ideas to modern model-based optimization, rather than treating BO as a separate black-box machine-learning technique.

Keywords:

Bayesian optimization, design of experiments, entropy, Gaussian processes, information theory

Teaching, Consulting and Knowledge Transfer in Statistics / 104

A Practical Roadmap for Choosing Correct Statistical Tests, Based on Three Decades of Teaching Experience**Author:** Hossein Bevrani¹¹ *University of Kurdistan*

Throughout my years of research and university teaching, as well as advising master's and doctoral theses in applied fields such as economics, management, biology, geology, and agriculture, I have noticed that students and researchers often face difficulties in selecting appropriate statistical methods to validate their hypotheses. They may either choose an inappropriate method or fail to consider its underlying assumptions. This issue applies not only to complex models but also to conventional statistical methods.

Selecting the correct statistical test is critical, especially when examining relationships between variables or choosing between parametric and nonparametric approaches. Based on my thirty years of teaching experience, this presentation provides a clear, practical roadmap for making this choice correctly.

We begin by distinguishing between parametric and nonparametric methods, highlighting necessary conditions including normality, homogeneity of variances, sample size, and measurement scale. For relationships between two variables, we categorize relevant tests such as Pearson/Spearman correlation, t-test, Mann–Whitney, ANOVA, Kruskal–Wallis, Wilcoxon and chi-square test.

A novel contribution of this presentation is a set of operational flowcharts and decision algorithms that guide researchers step-by-step toward the correct test. These tools are built from practical scenarios frequently encountered in applied statistics. Each algorithm explicitly checks assumptions, data types, and research questions before recommending a specific test.

Attendees will leave this session having observed a comprehensive dashboard of relevant statistical tests simultaneously, while learning how to select the appropriate test for their specific research context.

Keywords:

Hypothesis testing, Parametric vs. nonparametric methods, Applied statistics

Teaching, Consulting and Knowledge Transfer in Statistics / 146

Plackett and Burman arrays Revisited**Author:** Jonathan Smyth-Renshaw¹¹ *Jonathan Smyth-Renshaw & Associates Ltd*

Plackett-Burman designs are experimental designs presented in 1946 by Robin L. Plackett and J. P. Burman while working in the British Ministry of Supply. Their goal was to find experimental designs for investigating the dependence of some measured quantity on a number of independent variables (factors), each taking L levels, in such a way as to minimize the variance of the estimates of these dependencies using a limited number of experiments. Interactions between the factors were considered negligible.

To this day I believe many experimenters overlook the use of these designs as they don't require a large amount of complex analysis.

During my time teaching Design of Experiment (DoE) I have focussed on the use of the Plackett and Burman (PB) 8 run experiment with the possibility to change 7 factors with 2 levels each. I have also started to study using 2 PB 8 run arrays with the same factors to generate data for an inner array and an outer array. I wish to describe the approach and demonstrate the benefits. I will also detail how domain knowledge and Statistical Process Control can further enhance learning.

I will use the helicopter experiment commonly used to test DoE as an example to demonstrate how domain knowledge, SPC and Plackett and Burman designs can be used in a practical and easy to follow manner.

This talk will focus on practical applications and delivering business solutions with minimal cost. I have found no other reference to my approach.

Keywords:

Teaching, Consulting and Knowledge Transfer in Statistics / 13

CROSSVALI: A Method for Prompting AI for Data Analysis**Author:** Jennifer Van Mullekom¹**Co-author:** Anne Driscoll²¹ *Virginia Polytechnic Institute & State University*² *Virginia Tech*

Over the past few years, we have been interested in answering the question “Can ChatGPT Think Like a Statistician?” The answer is “Yes, but...”. Our data analysis prompting experience has prompted us (pun intended!) to create a framework for obtaining appropriate data analyses from artificial intelligence called CROSSVALI (Context, Refined Questions, Options, Specificity, Scrutiny, Verify & Validate, Ask Questions & Chain of Thought, Look Over, Interpret). When communicating in a collaborative data analysis team, the team shares the burden of communication. When collaborating with AI, the prompter takes on the majority of the burden of communication. The CROSS portion of the framework is input focused to help the prompter create a well-defined, well-communicated prompt with the necessary elements to translate the domain question into a statistical one. The VALI portion is human in the loop focused to ensure that we comply with ethical and responsible use of AI. In this talk we will cover the details of the framework in the context of quality applications for both typical LLMs and agents using skills (e.g. Claude Code and associated skills). At the end of the talk, attendees should be able to apply the framework and instruct teammates or students on the elements of the prompting framework in order to obtain more appropriate data analyses.

Keywords:

AI, Prompting, Data Analysis

Teaching, Consulting and Knowledge Transfer in Statistics / 67

Bridging the Gap: Interactive Discovery as a Booster for Preparing Industry-Ready Graduates**Authors:** Paolo Chiappa¹; Volker Kraft¹¹ JMP

The GAISE recommendations (2025) are shaping best practices in teaching statistics, shifting the focus toward statistical thinking as a holistic investigative process. Despite related academic advancements, a persistent skills gap remains in preparing graduates for the complex unstructured problem-solving requirements of modern industry. This session proposes a framework where teaching methodology and tools work in tandem to transform statistical learning into an interactive discovery process.

We will emphasize examples of teaching best practices that prioritize an exploratory mindset over rote procedural verification. For instance, the presentation showcases how dynamic visualizations like the JMP Profiler enable students to intuitively navigate multi-factor relationships and communicate insights effectively. Furthermore, we highlight the pedagogical value of DOE simulations for experimental design, allowing learners to experience industrial complexity in a risk-free environment. The role of AI support is also discussed as a booster to lower technical barriers and accelerate conceptual mastery.

The session concludes with success stories of students who have “hit the ground running” in industrial roles. We will welcome an open discussion on the primary friction points in collaborative projects, illustrating how the right analytical environment acts as a booster for the statistical mindset required in Industry 4.0.

Keywords:

Skills gap, interactive discovery, AI support

Teaching, Consulting and Knowledge Transfer in Statistics / 148

Deep Adaptive Design for Model-Based DOE, Screening, and Bayesian Optimization**Author:** Arno Strouwen¹¹ Strouwen Statistics; PumasAI; KULeuven

Sequential experiments that adapt to incoming data are more efficient than static designs, but the required posterior inference and design optimization between steps are usually too expensive to run online. Deep Adaptive Design [DAD, Foster et al., 2021] sidesteps this by training a neural network policy online that maps any experimental history to the next design point in a single forward pass. We apply DAD to three problems central to industrial statistics: model-based design of experiments, factorial screening, and Bayesian optimization. On a Monod bioreactor, the adaptive policy reduces posterior RMSE by about 30% relative to the best static design [Strouwen and Micluș-Câmpeanu, 2026]. We describe ongoing work on adaptive screening. On the multimodal 2D Rastrigin benchmark, the learned policy achieves roughly an order-of-magnitude lower simple regret than the best tested classical baseline.

Keywords:

Time Series, Forecasting and Dynamic Systems / 43

A Semi-Supervised Framework for Optimisation-Based Label Inference in RUL Prediction**Author:** Joseph Cupit¹**Co-authors:** Dongda Zhang²; Elizabeth Dickinson ; James Birbeck¹ *The University of Manchester*² *University of Manchester*

Industrial systems generate large volumes of operational data, enabling predictive maintenance strategies to reduce unplanned downtime and costs. Over the past decade, machine learning (ML) models have been widely used for predicting equipment degradation. However, their effectiveness is constrained by the scarcity of high-quality labels, as industrial datasets remain largely unlabelled. Existing labelling approaches typically rely on expert domain knowledge or simplified degradation assumptions, which can introduce bias and limit generalisability.

This work proposes a novel semi-supervised framework for inferring key performance indicators, including Remaining Useful Life (RUL), from limited labelled data. The approach formulates label estimation as an optimisation problem, where candidate labels are iteratively updated to minimise prediction error. Long Short-Term Memory (LSTM) neural networks are constructed within each iteration to evaluate solution quality and guide the optimisation process, enabling progressive refinement of inferred labels while improving predictive accuracy.

The framework is evaluated on two established predictive maintenance datasets. Initial results indicate that it can effectively infer the RUL for large amounts of previously unlabelled data and achieves high predictive accuracy on unseen batches. Furthermore, by extracting additional process information from the newly labelled batch data, the framework expands the effective training dataset and enables continuous refinement of the LSTM networks. This allows them to learn a more comprehensive representation of system behaviour and supports the development of more robust and transferable predictive maintenance models.

Keywords:

Semi-Supervised Learning, Label Inference, Predictive Maintenance

Time Series, Forecasting and Dynamic Systems / 51

Diffusion Models for Multivariate Probabilistic Net Load Forecasting in Transmission Grid Congestion Management**Authors:** Maxime Ducourau¹; Ouassim Feliachi²¹ EPFL² RTE

Short-term congestion risk assessment in transmission grids is still largely based on deterministic load flow computations from point forecasts. We investigate multivariate probabilistic net load forecasting across substations as a way to better quantify the probability of congestion events, which arise from correlated forecast errors across substations. This is particularly challenging with growing renewable penetration, which induces high-dimensional joint distributions with complex, often strongly non-Gaussian behaviour and strong spatial dependence. In this context, we develop and assess a denoising diffusion probabilistic model designed to generate joint predictive distributions with calibrated univariate behaviour and realistic spatial dependence across substations.

Our study focuses on three-hour-ahead forecasting at 115 substations on the French extra-high-voltage grid. We compare this approach with alternative strategies for constructing multivariate predictive distributions, combining different marginal models —including empirical residual-based methods and quantile regression —with different dependence structures such as Gaussian and empirical copulas. The comparison is intended to clarify whether diffusion-based generative modelling yields practical benefits over simpler statistical approaches in this setting.

Performance is evaluated using multivariate proper scoring rules together with operational metrics relevant for congestion risk. We also discuss permutation-invariant model designs to handle unseen substations and hierarchical extensions to scale probabilistic forecasting to the roughly 5,000 substations of the French grid.

Keywords:

diffusion models, power systems, probabilistic forecasting

Time Series, Forecasting and Dynamic Systems / 79

INAR(1) Processes based on the Zipf-PSS distribution**Authors:** Ariel Duarte-López¹; Manuel Gonzalez Scotto²; Marta Perez-Casany³¹ Technical University of Catalonia² Técnico de Lisboa³ Universitat Politècnica de Catalunya

The Zipf-PSS distribution is a Poisson Stopped-Sum with a Zipf distribution as secondary distribution. In this work, we consider two INAR(1) processes: The Zipf-PSS-INAR(1) innovations process, whose innovations follow a Zipf-PSS distribution, and the Zipf-PSS-INAR(1) marginal process, whose stationary marginal distribution is Zipf-PSS. Working with the marginal process is more complex than working with the innovation process, because it requires to compute the unknown distribution of the innovations. Nevertheless, the distribution of the innovation has a notable feature: it depends on the survival parameter (a larger survival parameter implies less immigration). This property is appealing from an applied perspective, and it is never achieved in the INAR(1) processes which are defined specifying the innovation distribution. A practical parameter interpretation of the two processes is provided, and their performance fitting real time series is compared with the Poisson INAR(1) and the NB-INAR(1) processes.

Keywords:

INAR(1), Zipf, Poisson-Stopped-Sum

Truthworthy and explainable AI / 19

Statistical Evaluation of CPU-Based Offline LLMs for Industrial Text Processing on Low-Cost Edge Hardware**Author:** Guido Moeser¹**Co-authors:** Andrea Ahlemeyer-Stubbe²; Igor Dub³; Rainer Bumm⁴¹ *masem research institute*² *Data Mining & More*³ *Wiesbaden Smart City Department*⁴ *PHOENIX*

Recent advances in generative AI have enabled powerful language models for industrial applications. However, most solutions rely on cloud-based infrastructures or GPU-accelerated environments, which raise concerns regarding data privacy, latency, and operational cost—particularly in industrial settings dealing with sensitive internal documents.

In this study, we investigate the feasibility of deploying offline large language models (LLMs) on CPU-only edge hardware, such as standard notebooks and low-cost mini PCs. In particular, we evaluate the performance of highly quantized models, including emerging architectures such as BitNet, within a real-world industrial use case.

The application scenario is based on a production system developed using n8n, where incoming customer emails are processed locally without external data transfer. Two core tasks are considered:

1. Structured information extraction (classification task): Extraction of machine identifiers and request types from customer emails.
2. Text summarization and interpretation (generation task): Generation of concise summaries and actionable insights from unstructured text.

For the classification task, performance is evaluated using classical statistical metrics derived from the confusion matrix, including precision, recall, and F1-score. For the generative task, we apply G-Eval-based metrics to assess dimensions such as correctness, completeness, and consistency. Beyond output quality, we introduce a comprehensive set of system-level performance indicators, including:

- processing latency
- tokens generated per second
- total token count per task
- CPU utilization and memory footprint

These metrics are analyzed under different deployment configurations, comparing execution on a high-performance notebook (64 GB RAM) and a low-cost edge device (Intel N95 mini PC, 12 GB RAM).

Our results demonstrate that CPU-based offline LLMs can achieve competitive performance for industrial text processing tasks, while significantly reducing infrastructure complexity and enabling fully local data processing. The study highlights that, with appropriate model selection and evaluation metrics, cost-efficient edge deployments without GPUs are a viable alternative for industrial AI applications.

Keywords:

Edge Computing, Quantized Language Models (BitNet), Offline-LLMs

Truthworthy and explainable AI / 23**Ups and Downs with AI and Old Data****Author:** Shirley Coleman¹**Co-authors:** Andrea Ahlemeyer-Stubbe²; Robert Shiel³; Darren Evans³¹ *Visiting Fellow, Newcastle University*² *Ahlemeyer-Stubbe*³ *Newcastle University*

Farming is vital business. Agricultural experiments have long been carried out on crops including sugar, wheat, potatoes and grass. The second oldest grassland experiment in the UK has been in continuous action at Newcastle University's Cockle Park farm in Northumberland since 1897. (The oldest dates from 1856 at Rothamsted Research Lab.)

Over the years data on grass (hay) yield, fertiliser treatments, soil structure, grass composition and the weather have been meticulously recorded in handwritten notes, spreadsheets and pdf files. Data was analysed in 1960 by Pawson and in 1980 by Coleman. Further analysis has been piecemeal and hampered by the disparate sources of data.

Using Artificial Intelligence (AI) to convert a photo of handwritten data or a pdf file into a comprehensive spreadsheet has been transformative. It has motivated work on collating this valuable data into a definitive resource available for analysis by soil scientists, climatologists and statisticians.

The full set of yield data from 1897 to 2025 has now been reanalysed replicating the basic analysis carried out in 1980 and extending it with newer methods of assessing correlations between plots and detecting carry-over effects and cycles in the yields. The wider project aim is to collate the data into a definitive, settled dataset that could be made available more widely via an open-source website.

The analytical results for hay yields are presented in this paper and clearly show dramatic changes over time, and relationships between the 14 experimental plots receiving different nutrient regimes. This is the upside of AI.

The downside of AI is that it does not solve the challenges in collating disparate data into a sound resource. Administrative decisions have to be made and recorded. AI is instrumental but still needs significant input from personnel with sound domain knowledge. Data is extremely valuable, the cost in time, land use and labour is enormous for 128 years' worth of data. AI unleashes the opportunity to maximise the information from experiments and improve quality and efficiency in this agricultural business. We discuss these issues in the paper.

Keywords:

agricultural business, time series, treatment effects

Truthworthy and explainable AI / 150

Towards the reliable use of metamodels in critical systems - Application to nuclear safety studies

Authors: Nicolas Bousquet¹; Eric Chojnacki²; Fabrice Fouet²; Bertrand Iooss³; Mathieu Segond⁴¹ EDF² ASNR³ EDF R&D⁴ Framatome

In many engineering studies, computer codes are increasingly used to understand, model, and predict physical phenomena. However, the numerical models underlying these codes can be computationally expensive, severely limiting the number of simulations that can be performed. A widely adopted solution consists in replacing the expensive computer code with a computationally efficient mathematical approximation, referred to as a surrogate (or meta-model). Surrogates may rely on a broad range of supervised learning techniques, including polynomial regression, Gaussian processes, random forests, and neural networks. Trained on a set of numerical simulations, they should accurately reproduce the code outputs over the input domain of interest while maintaining strong predictive performance at unseen points.

This talk addresses the challenges involved in validating surrogates intended for use in safety-critical systems, whose failure could have catastrophic consequences, with a particular focus on nuclear safety applications. The work was initially conducted by a working group dedicated to assessing the reliability of thermal-hydraulic passive systems.

We first examine the fundamental differences between traditional physics-based numerical models and data-driven surrogates constructed using machine-learning algorithms. Rigorous validation of these tools is essential to establish their reliability and suitability for critical applications. We also discuss connections with recent developments in trustworthy artificial intelligence.

Drawing inspiration from the standardized Verification, Validation, and Uncertainty Quantification process applied to scientific computing tools in nuclear safety studies, we propose a methodological framework structured around four confidence dimensions: conformity, robustness, explainability, and transparency. Each dimension must be assessed using appropriately selected methods and criteria. A case study concerning the reliability of thermal-hydraulic passive systems illustrates the relevance and practical implementation of this framework.

Keywords:

Uncertainty quantification and computer experiments / 33**Prediction of physical fields under linear constraints**

Authors: Clément Gauchy¹; Mahamat Hamdan Nassouradine¹; Pierre-Emmanuel Angeli²; Sébastien Da Veiga³

¹ *Université Paris-Saclay, CEA, Service de Génie Logiciel pour la Simulation, France*

² *Université Paris-Saclay, CEA, Service de Thermohydraulique et de Mécaniques des Fluides, France*

³ *ENSAI, CREST, F-35000 Rennes, France*

Metamodeling is a fundamental approach for approximating computationally expensive numerical simulations in engineering applications, such as uncertainty propagation and sensitivity analysis. In this work, we address the simultaneous prediction of multiple high-dimensional physical fields governed by linear equality constraints, a setting that arises naturally in problems involving conservation laws or incompressibility conditions. Gaussian Process (GP) regression is a natural tool for this task given its effectiveness in small sample regimes, but it faces two intertwined challenges: managing the high dimensionality of the discretized output fields, and enforcing the physical constraint in predictions. A common strategy to handle linear constraints consists in deducing one output from the others via the constraint relation, and building an unconstrained surrogate on the remaining fields. Through an extensive empirical benchmark, we show that this deductive approach suffers from a sensitivity to the arbitrary choice of which output to deduce, affecting both predictive accuracy and uncertainty quantification. To address these limitations, we propose first a specific PCA procedure for multi-field data, coined row-wise PCA, which has the interesting property of preserving the constraint in the latent space. Since standard PCA strategies for multi-field data (field-wise, column-wise) do not preserve such constraints, we investigate theoretically the optimality of the row-wise choice and provide conditions under which it incurs a negligible reconstruction cost. In a second step, we consider a linearly-constrained multi-output GP approach based on a specific kernel parametrization. The proposed framework is validated on a population dynamics problem and on an industrial computational fluid dynamics application, which involves the prediction of Reynolds stress tensor components under the incompressibility constraint. We demonstrate competitive predictive performance while guaranteeing strict constraint satisfaction.

Keywords:

Gaussian process, PCA, linear constraints

Uncertainty quantification and computer experiments / 41

Multifidelity Gaussian process regression for solving nonlinear partial differential equations**Authors:** Fatima-Zahrae EL-BOUKKOURI¹; Josselin Garnier²; Olivier Roustant¹¹ *INSA Toulouse / IMT*² *CMAP*

Kernel-based methods provide a principled alternative to classical numerical solvers for nonlinear partial differential equations (PDEs), especially in mesh-free settings with built-in regularization and uncertainty quantification. Traditional discretization techniques such as finite differences or finite elements can become computationally demanding for nonlinear or multiscale problems. In contrast, the variational framework of Y. Chen et al. formulates PDE solving as a constrained minimization problem in a reproducing kernel Hilbert space (RKHS), with a natural Gaussian-process (GP) interpretation.

In this setting, the solution is defined as the minimizer of the RKHS norm subject to PDE and boundary constraints at collocation points. This formulation is equivalent to the maximum a posteriori estimator of a GP conditioned on the constraints, making the choice of kernel central to solution quality, regularity, and stability.

We introduce a multifidelity physics-informed methodology for constructing RKHSs adapted to PDE resolution. We assume access to low-fidelity simulations and sparse high-fidelity observations. From the low-fidelity observations, we estimate an empirical covariance and approximate it by a smooth kernel within a parametric family. High-fidelity information is incorporated using the autoregressive cokriging model of Kennedy and O'Hagan :

$$Y_H = \rho Y_L + Y_d.$$

The resulting RKHS is used directly in the PDE-constrained problem. We also extend the framework to non-centered GPs using a cokriging-informed mean. Numerical results on Burgers equations show reduced sensitivity to hyperparameters compared with single-fidelity approaches, while preserving the differentiability required by kernel-based solvers.

Keywords:

kernel learning, multifidelity, physics-informed learning

Uncertainty quantification and computer experiments / 100**Taxonomy of statistical models for investigating order picking systems****Authors:** Larissa Sander¹; Sonja Kuhnt²¹ *Fachhochschule Dortmund*² *Dortmund University of Applied Sciences and Arts*

Order picking systems are known as complex logistics systems, where uncertainties, e.g., caused by randomness of incoming orders or delay of supplies, are present. Investigating the relationship between input variables and key performance indicators (KPIs) in common types of order picking systems, described by reference models, is aimed. Input variables are for example system load, batch size or picking strategy and typical KPIs are throughput or utilization rate.

Design and Analysis of Computer Experiments (DACE) is applied in the following steps: A discrete-event simulation model of a reference model is generating simulation output data based on a design of experiment (DoE). The data is used to build statistical models, i.e., metamodels, with uncertainty bands for depicting the influence of the input variables on the KPIs.

A taxonomy of classes of statistical models is developed to be used for fitting statistical models to reference models. Goal is the assignment of specific classes of statistical models to reference models taking different requirements from both the simulation model and the type of order picking system into account. Aspects to be kept in mind are the types of input variables (continuous, discrete, categorical), the distributional assumptions of the KPIs, the types of dependencies (linear, curvilinear, exponential, interactions, ...), the possibility to calculate a prediction interval with specific characteristics as well as dependencies between observations in case of replications. First statistical analysis of some reference models are presented.

Keywords:

Design and Analysis of Computer Experiments, order picking system, statistical metamodeling

Uncertainty quantification and computer experiments / 48

Spectral Clustering for Detecting Stationary Subregions in Gaussian Process Regression

Author: Théo Sylvestre¹**Co-authors:** Amandine MARREL²; Guillaume Damblin²; Raksmy Nop²¹ *CEA Paris Saclay*² *CEA*

Complex physical systems are often modeled using high-fidelity simulation codes. However, their inherent complexity makes each simulation computationally expensive, which severely limits their direct use in tasks such as uncertainty quantification. To overcome this, a widely adopted solution is to approximate the simulator using a Gaussian Process Regression (GPR) surrogate model.

However, when the underlying simulation exhibits heterogeneous behaviour over the input domain, such as in two-phase fluid mechanics where flow regime may vary, a standard GPR model with a stationary covariance function may struggle to accurately represent it. In such cases, more sophisticated non-stationary GPR models should be considered. Several approaches have been proposed in the literature, some of which rely on firstly identifying stationary subregions through clustering methods applied to input–output pairs.

These methods typically group points based solely on proximity in the joint input–output space, without explicitly incorporating information about what characterises a shared stationary regime.

To address this limitation, we propose to perform the clustering step using local features as variances and gradients rather than raw outputs, thereby making the clustering process more consistent with the notion of stationarity. In addition, we propose the use of a spectral clustering approach to better account for variations in cluster density induced by local features, while promoting spatial contiguity within the input domain.

The proposed methodology demonstrates encouraging results in moderate dimension for identifying stationary subregions compared to standard approaches, despite the increase in computational complexity.

Keywords:

Non-stationary Gaussian process regression, Spectral clustering

Uncertainty quantification and computer experiments / 50

Multi-fidelity Gaussian processes for noisy outputs and non-nested experiment designs: a comparison between the recursive and non-recursive formulations**Author:** Nils Baillie¹**Co-authors:** Baptiste Kerleguer²; Cyril Feau¹; Josselin Garnier³¹ *Université Paris-Saclay, CEA, Service d'Etudes Mécaniques et Thermiques, 91191 Gif-sur-Yvette, France*² *CEA, DAM, DIF, F-91297, Arpajon, France*³ *CMAP, CNRS, École polytechnique, Institut Polytechnique de Paris, 91120 Palaiseau, France*

Surrogate models provide fast approximations of computationally expensive simulations (or experiments) and are trained using a limited set of observations generated by these codes. In the multi-fidelity framework, we assume the availability of two computer models with different levels of cost and accuracy. The high-fidelity model z_H provides the most accurate predictions but is also the most expensive to evaluate. In contrast, the low-fidelity model z_L is significantly cheaper to compute but offers lower accuracy. We focus in this work on the auto-regressive model which supposes a linear relation between the two codes: $\forall \mathbf{x} \in \mathbb{R}^D, z_H(\mathbf{x}) = \rho(\mathbf{x}) \cdot z_L(\mathbf{x}) + \delta_H(\mathbf{x})$, where ρ is the scaling factor and the functions z_L and δ_H are approximated with Gaussian processes.

This model was initially developed by Kennedy and O'Hagan and improved afterwards by Le Gratiet and Garnier with the more computationally efficient recursive formulation. However, these works rely on two important assumptions: first, that the observed outputs are deterministic; and second, that the experimental designs are nested, meaning that each high-fidelity input point coincides with a low-fidelity input point. Under these assumptions, the optimization of the model parameters is significantly simplified.

We generalize the recursive formulation to the case of noisy outputs and non-nested designs. An alternative optimization approach based on the expectation-maximization algorithm is compared to the direct maximum likelihood estimation for the initial, non-recursive formulation. We apply both formulations of the auto-regressive model to several cases of varying difficulty and show that the proposed approach achieves faster training times for large low-fidelity datasets.

Keywords:

Surrogate modeling, Auto-regressive co-kriging, Expectation-Maximization

Uncertainty quantification and computer experiments / 105

Integrating Global Sensitivity Analysis into Bayesian Optimization for Tactical Decision Support in Transport Logistics**Authors:** Lara Kuhlmann de Canaviri¹; Sonja Kuhnt²¹ *Fachhochschule Dortmund*² *Dortmund University of Applied Sciences and Arts*

Transport logistics facilities such as less-than-truckload terminals require fast and robust tactical decisions under uncertainty, for example regarding task scheduling, resource allocation, and terminal configuration. Since detailed simulation experiments are computationally expensive, surrogate-assisted optimization methods provide an important basis for decision support.

This contribution presents and compares strategies for integrating global sensitivity analysis into Bayesian optimization for simulation-based decision support in transport logistics. Gaussian process surrogate models are combined with sensitivity measures to guide sequential optimization, including variable screening, candidate generation, soft search space reduction, and sensitivity-informed acquisition strategies.

The approaches are evaluated in a multi-objective logistics setting involving throughput, waiting times, resource utilization, and process efficiency. The comparison focuses on optimization performance, computational efficiency, and interpretability, demonstrating the potential of sensitivity-guided Bayesian optimization for more targeted tactical decision-making.

Keywords:

Bayesian Optimization, Global Sensitivity Analysis, Transport Logistics Simulation

Uncertainty quantification and sensitivity analysis / 10

Sensitivity analysis for sets: Application to pollutant concentration maps**Author:** Noé Fellmann^{None}**Co-authors:** Adrien Spagnol ; Celine Helbert ¹; Christophette Blanchet ; Delphine Sinoquet ; Mathis Pasquier¹ *Centrale Lyon - ICJ*

We are motivated by the field of air quality control, where one goal is to quantify the impact of uncertain inputs such as meteorological conditions and traffic parameters on pollutant dispersion maps. Sensitivity analysis is one answer, but the majority of sensitivity analysis methods are designed to deal with scalar or vector outputs and are badly suited to an output space of maps. To address this problem, we propose a generic approach to sensitivity analysis of set-valued models. This approach can be applied to the case of maps. We propose and study three different types of sensitivity indices. The first ones are inspired by Sobol' indices but adapted to sets based on the theory of random sets. The second ones adapt universal indices defined for a general metric output space. The last set of indices uses kernel-based sensitivity indices adapted to sets. The proposed methods are implemented and tested to perform an uncertainty analysis for a toy excursion set problem and for time-averaged concentration maps of pollutants in an urban environment.

Keywords:

kernel-based sensitivity indices, random sets, HSIC

Uncertainty quantification and sensitivity analysis / 75

Explicit functional ANOVA and new results in model and algorithm explainability and sensitivity analysis**Authors:** Baptiste Ferrere¹; Fabrice Gamboa²; Jean-Michel Loubes²; Joseph Muré³; Nicolas Bousquet³¹ EDF / IMT² IMT³ EDF

We present a unified perspective on explicit functional ANOVA as a principled decomposition framework for black-box models, bridging explainability, sensitivity analysis, and algorithmic understanding. We derive an exact closed-form functional ANOVA for categorical inputs, valid under arbitrary dependence structures and even on sparse or non-rectangular supports, thereby removing a major limitation of standard ANOVA-based methods outside the independent setting. It extends other new results showing that classical Fourier analysis on the Boolean hypercube is in fact a particular case of Hoeffding functional decomposition under the uniform product measure. These results yield tractable decompositions that preserve hierarchical orthogonality, recover the classical orthogonal ANOVA in regular regimes, and naturally induce generalized SHAP-like feature attributions in machine learning. Beyond local and global explainability, they provide new tools to analyze interaction structure, detect spurious effects induced by encoding or dependence, compare models through their decomposed mechanisms, and study stability, robustness, and mismatch between algorithms and data distributions. They also suggest broader avenues for sensitivity analysis under dependence and for model compression through low-order effect dictionaries. Finally, these advances point toward continuous-input settings, notably because modern tabular architectures increasingly process numerical variables through learned embeddings or tokenizations, making explicit ANOVA on intermediate representations a promising route toward continuous-domain extensions.

Keywords:

Functional ANOVA, Hoeffding, SHAP

Uncertainty quantification and sensitivity analysis / 15

Sensitivity Measures for Continuous Actions**Author:** Lea Friedli¹**Co-authors:** Daniel Straub¹; Julius Weidinger¹¹ Technical University of Munich

The value of information (VOI) is a decision sensitivity measure that quantifies the expected improvement in decision quality when uncertainty in selected inputs is removed. Unlike many other sensitivity measures, the VOI provides not only a relative ranking of factors but also an absolute metric of decision quality. Despite this, its use has been limited, particularly to decision problems with discrete alternatives. Our work addresses this gap by developing an approach for settings with continuous decisions, which arise in various fields, including medicine (e.g., determining dosage intensity), environmental management (e.g., setting water release levels in reservoir management), and engineering design (e.g., choosing dimensions of structural elements). In such contexts, instead of selecting among a small number of discrete alternatives, decision makers must determine an action that varies over a continuous domain. Since the optimal decision depends on uncertain inputs, the VOI indicates which uncertainties are most valuable to reduce. To address the computational challenges of computing the VOI for continuous actions, we develop a method based on smoothing techniques that approximate optimal decisions from a set of samples. Overall, our results in engineering test cases show good performance for the problems investigated. The main remaining challenges for practical use are selecting appropriate hyperparameters and the relatively high number of samples required.

Keywords:

Value of information; decision sensitivity; smoothing techniques

Young Statisticians Session / 3

A Registration-free Approach for Shape and Color Monitoring of Functionally Graded Materials via 4D Point Clouds

Author: Mariafrancesca Patalano¹

Co-authors: Giovanna Capizzi²; Kamran Paynabar³

¹ *University of Padua*

² *Department of Statistical Sciences*

³ *School of Industrial and Systems Engineering*

Recent advances in additive manufacturing enable the fabrication of complex parts with intricate geometries and spatially-varying material composition. Data fusion integrates point cloud data with chromatic attributes, yielding 4D point clouds, a rich representation that jointly encodes shape and material information. We introduce a registration-free framework for jointly monitoring shape and surface color via 4D point clouds. The proposed approach leverages the Laplace-Beltrami operator to capture intrinsic spectral features. A combined monitoring scheme is developed to detect shape deformations and color anomalies, complemented by a spatially-aware post-signal diagnostic procedure to determine the source of change and localize color anomalies. Crucially, neither component requires point cloud registration or mesh reconstruction, thereby eliminating error-prone and computationally expensive pre-processing steps. The performance of the proposed framework is assessed through a Monte Carlo simulation study and a case study.

Keywords:

Statistical Process Monitoring; Post-signal diagnostic; Laplace-Beltrami operator

Young Statisticians Session / 66

On stochastic network trends in network time series

Author: Amandine PIERROT¹

Co-authors: Guy NASON²; Matthew NUNES¹

¹ *University of Bath*

² *Imperial College London*

In many applications of interest, multivariate time series data feature trend behaviors. Yet, trends that may affect multivariate stochastic processes are still largely dealt with in a univariate manner. Calling on differencing and co-integration concepts for univariate time series, we introduce stochastic trends for multivariate data, with particular focus on trends that are constrained by an underlying network. When introduced into an auto-regressive time series model, stochastic network trends allow practitioners to fit models that assume the component time series to move together, can discriminate between what comes from the network and what is only influenced by the past and whose sparsity is a priori enforced through the network.

Such stochastic network trends embed contemporaneous effects in a matrix, and estimating this matrix through an ordinary least squares approach leads to inconsistent estimators. Hence, we propose to estimate this trend matrix using maximum likelihood estimation and transform a network-constrained optimization problem into an unconstrained one. We show that the objective function for this problem is strongly convex in the trend parameters and propose an efficient algorithm for estimation based on block coordinate descent. We show that this algorithm converges to a stationary point, and that the corresponding estimators are consistent.

Keywords:

multivariate time series; auto-regressive models; networks

Young Statisticians Session / 63**Active Learning for Effect Screening in Manufacturing****Author:** Marcus Engsig¹**Co-authors:** Bart De Ketelaere ²; Murat Kulahci ³¹ *Technical University of Denmark*² *Catholic University of Leuven*³ *DTU*

This paper addresses effect screening from observational data where sampling is constrained by cost, time, or process limitations, with a main focus on manufacturing applications. We propose a novel active learning strategy that introduces principles from optimal experimental design (A- and D-optimality) and combines it with an optimization for multicollinearity using Variance Inflation Factors (VIF), making the approach suitable for observational manufacturing data with strong dependency structures and feedback loops.

We aim to ultimately obtain prescriptive models consisting of the selected effects, for control and optimization of the manufacturing processes.

We evaluate the performance of the proposed effect selection strategy across multiple methods, including Lasso, Pearson correlation, Boruta, and a Bayesian approach. The results indicate that although the strategy can improve screening efficiency, it can also lead to the selection of the irrelevant variables that are strongly correlated with truly relevant variables. As a result, correlated but non-relevant variables may be retained, while relevant variables may be excluded when multicollinearity is high.

These findings highlight an important trade-off between screening ability and predictive performance in constrained industrial sampling, and challenges the purpose of sampling.

Keywords:

Active Learning, Optimal Experimental Designs, Manufacturing

Useful Information

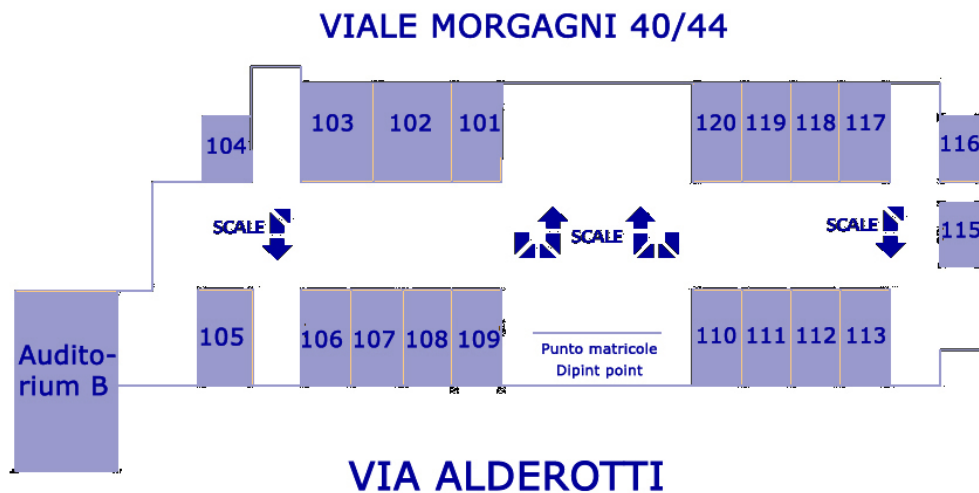
Congress venue

Centro Didattico Morgagni
Viale G. B. Morgagni 40
50134 Firenze



CDM PIANO 1

Scuola di Scienze della Salute Umana
Centro Didattico Morgagni, viale Morgagni 40/44



Registration and ENBIS Secretary Desk:

- Room #108 (1st floor)

Sunday, September 6 from 2:00 PM to 5:00 PM

Monday, September 7 from 8:15 AM to 4:00 PM

Tuesday, September 8 from 8:15 AM to 12:00 AM

Wednesday, September 9 from 8:15 AM to 8:45 AM

Conference sessions:

- Auditorium B (1st floor)
- Conference Room #102 (1st floor)
- Conference Room #103 (1st floor)
- Conference Room #106 (1st floor)
- Conference Room #107 (1st floor)

Coffee Breaks & Lunches:

- 1st Floor Hall (close to the conference rooms and room #108)

Group photo:

- Please make sure to be at the Centro Didattico Morgagni for the group photo on Tuesday, September 8 at 1:45 PM

Assistance to Conference Participants

Conference assistants, wearing yellow badges, are ready to help ENBIS-26 participants.

Wireless Network

- UNIFI provides access to eduroam.
- External guests from institutions that are not part of an eduroam service, can access eduroam with username W500005 and password xad40388. For more details, use this QR code:



Coming to Florence

Florence has a small international airport (Amerigo Vespucci), but there are also convenient alternatives nearby:

- **Pisa Airport (PSA)**
Pisa is very close to Florence, and frequent trains connect Pisa Airport directly to Santa Maria Novella Station in the city center. Trains are reliable and usually on time, making this a practical choice for many travelers.
- **Bologna Airport (BLQ)**
Located about 100 km from Florence, Bologna Airport is another excellent option. High-speed trains link Bologna to Florence in just under 40 minutes, ensuring a smooth journey.

Both Pisa and Bologna airports are well connected and often provide more international flight options than Florence itself. No matter which airport you choose, reaching Florence is straightforward and convenient.

Reaching the venue

The conference venue is located in northern Florence and is easily accessible by tram:

- From Santa Maria Novella railway station: take the T1 from Alamanni–Stazione towards Careggi Ospedale and get off at Morgagni–Università (approximately 15 minutes).
- From Amerigo Vespucci Airport (FLR): take the T2 towards San Marco, change to the T1 towards Careggi Ospedale at Alamanni–Stazione, and get off at Morgagni–Università (approximately 40 minutes in total).

Tram tickets

A single tram ticket costs 2.00 Euros and is valid for 90 minutes from validation. Tickets can be purchased and validated on board using the contactless *Tip Tap* service, with a credit, debit, or prepaid card, or an enabled smartphone or smartwatch. When changing tram lines, passengers

must tap the same card or device again; no additional fare will be charged within the 90-minute validity period.

Car parks

The nearest public car parks to the conference venue are:

- **Careggi–CTO**, Viale Gaetano Pieraccini 4 (approximately 200 m from the conference venue)
- **Pieraccini–Meyer**, Viale Gaetano Pieraccini 24 (approximately 1.3 km from the conference venue)

Parking fees apply.

Emergency

112 is the European emergency number you can dial free of charge from fixed and mobile phones everywhere in the EU. It will get you straight through to the emergency services – police, ambulance, fire brigade.

Electricity

Electricity in Italy is a 220-240 Volts, 50 Hz system. The power sockets are of type F and L. Please bring an appropriate adapter with you.

Disabled Access

All the lecture rooms in the Centro Didattico Morgagni (CDM) are accessible by wheelchair. Search “Viale Morgagni, 44, Firenze” on OpenStreetMap to obtain the correct position of the entrance of the CDM.

Social events

1. Welcome Reception

- Monday 7th September 2026, 19:30-22:00 at Braumeister (Via Madonna della Tosse 12/r) <https://braumeisterfirenze.com/>
- The venue is reachable by tram: take the T1 from Morgagni–Università towards Villa Costanza, change to the T2 towards San Marco at Fortezza, and get off at Parterre. From there, Braumeister is a short walk away (approximately 25 minutes in total).
- It is also reachable by car. The parking is very close to the venue (Parking Parterre).

2. Conference Dinner

- Tuesday 8th September 2026, 19.45:23.30 at Istituto degli Innocenti, Piazza SS. ma Annunziata 12
 - 19.45:20.30 Cocktail (Cortile delle Donne)
 - 20.30:23.00 Sitting Gala Dinner (Salone Brunelleschi)
 - 23.00:23.30 Closing Gala Dinner (Cortile delle Donne)
- The venue will be closed to external public at 20:30. You are asked to arrive on time for the sitting dinner.

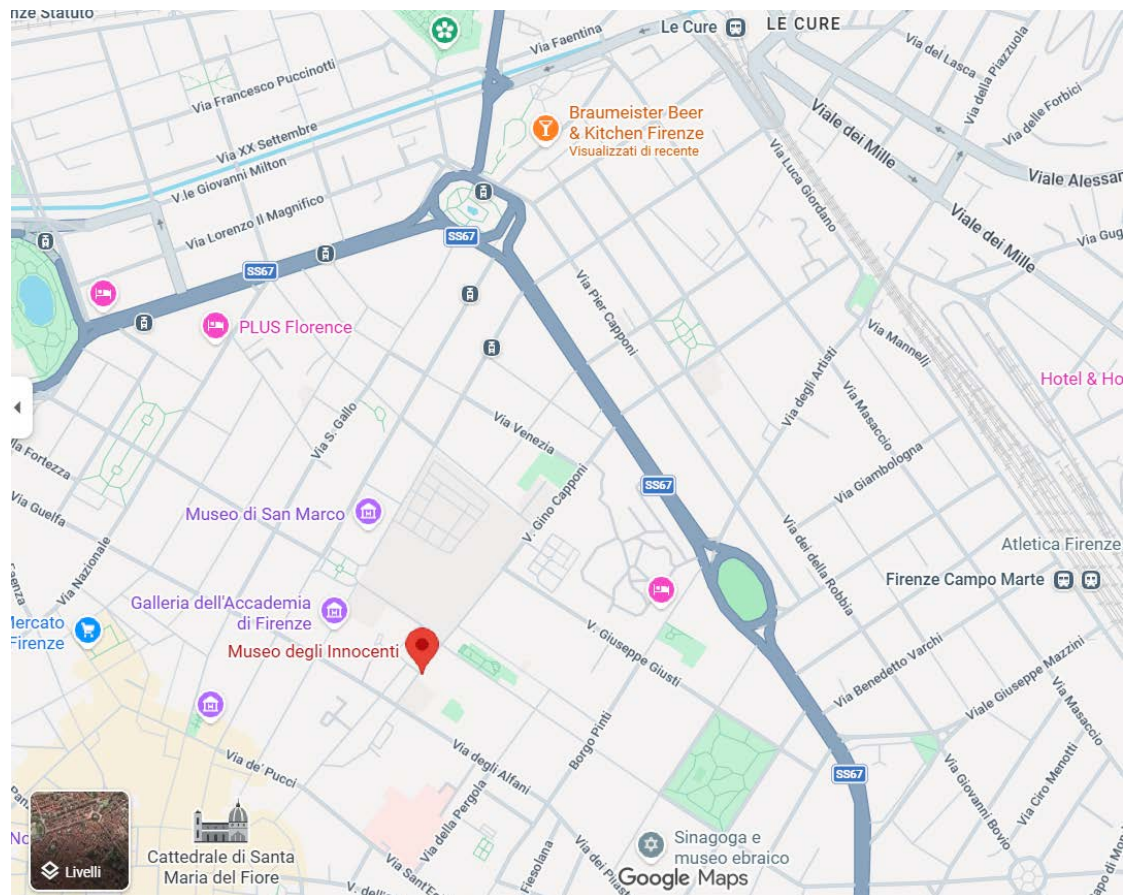
- The venue is reachable by tram: take the T1 from Morgagni–Università towards Villa Costanza, change to the T2 towards San Marco at Fortezza, and get off at San Marco–Università. From there, Istituto degli Innocenti is approximately a 7-minute walk away (approximately 30 minutes in total).
- A wonderful view of Florence is available on the rooftop of the building (Caffè del Verone).



With its famous loggia on Piazza SS. Annunziata, this monumental complex, home of the Institute, has been acknowledged as one of the earliest examples of Italian Renaissance architecture. The 'factory' designed by Brunelleschi was conceived with such far-sightedness that it fulfilled, also thanks to subsequent extensions, the various functions that are still housed there today. It is almost a miniature city, with internal and external passageways, façades, porticoes, courtyards, loggias and galleries. A city organically integrated into the urban context.

<https://www.istitutodeglinnocenti.it/en/architectural-heritage>

<https://www.istitutodeglinnocenti.it/en/node/37952>



ENBIS-27 Conference Announcement



Photos: Logan Armstrong and Lief Peng / Unsplash

The 27th Annual Conference of the European Network for Business and Industrial Statistics (ENBIS) will take place in Barcelona, Spain, from Wednesday, 1 September, to Friday, 3 September 2027.

The conference programme will feature distinguished keynote speakers, invited and contributed sessions, workshops, and panel discussions. As usual, pre- and post-conference courses are planned for Tuesday, 31 August, and Saturday, 4 September, respectively.

The conference will be hosted by the **Universitat Politècnica de Catalunya · BarcelonaTech (UPC)** and held at the **Barcelona School of Industrial Engineering (ETSEIB)**. The venue is conveniently located and easily accessible from across the city by metro, tram, and bus.

All plenary and parallel sessions will take place on the ground floor, in rooms located close to one another and fully air-conditioned.

For up-to-date information about ENBIS-27, please visit www.enbis.org.

To contact the Organising Committee, please write to enbis27@event.upc.edu.



ENBIS-26 gratefully acknowledges support of the following sponsors:



UNIVERSITÀ
DEGLI STUDI
FIRENZE

Dipartimento di Statistica,
Informatica, Applicazioni
"Giuseppe Parenti"

Eccellenza 2023-2027



ReDS
Rethinking
Data Science

Conference gold sponsor:

Minitab ®

General sponsor and conference gold sponsor:


jmp®
**STATISTICAL
DISCOVERY**

ISBN 979-12-243-4341-7



9 791224 343417