

### Selection Bias in Experimental Results: A Proxy Pattern-Mixture Model for Randomized Trial Generalizability

*Rebecca Andridge, The Ohio State University College of Public Health*

Randomized controlled trials (RCTs) are the gold standard for assessing interventions, yet most use convenience samples rather than random samples from a target population. If unobserved differences between participants and non-participants include factors related to treatment effectiveness ("effect modifiers"), trial results may fail to generalize, contributing to the "efficacy-effectiveness gap". In this talk, I introduce a novel sensitivity analysis framework to assess the impact of such unmeasured factors on treatment effect estimates called the Proxy Pattern-Mixture Model in the context of RCTs (RCT-PPMM). This framework that treats this selection bias as a fundamental problem of incomplete or missing data because outcomes for the entire target population are not observed. The RCT-PPMM relies on two bounded sensitivity parameters that capture the deviation from sample selection at random and that can be varied systematically to determine how robust trial results are to a departure from ignorable sample selection. A critical advantage is that it requires only summary-level baseline covariate data for the target population, rather than individual-level data for nonparticipants. I illustrate the use of the method using a yoga intervention RCT for breast cancer survivors, showing how conclusions may shift under plausible selection biases. The approach offers a practical and interpretable tool for evaluating generalizability, particularly when individual-level data on nonpart

matters. We focus on two related problems: testing mutual independence among several variables and testing independence between two random vectors in the presence of missing observations. Particular attention is given to the impact of different missingness mechanisms, such as MCAR and MAR, on test statistics, asymptotic behaviour, and finite-sample performance. We compare several approaches, including complete-case procedures, weighting-based methods, and imputation strategies, and provide practical recommendations for different missing-data scenarios.

### On the k-sample problem in the presence of random right censoring

*Anke Steyn, Centre for Business Mathematics and Informatics, North-West University, Potchefstroom, South Africa*

Across different disciplines and fields, a question that often arise is whether two or more samples of time-to-event data are realised from the same distribution. The problem is complicated by the presence of random right censoring. We discuss a new test based on empirical characteristic functions. The proposed test statistic is a weighted L2 type distance between the empirical characteristic functions of the observed samples. In the presence of censoring, the empirical characteristic function is based on the Kaplan-Meier estimator. We investigate the asymptotic properties of the test, and consider the finite-sample performance of the test for different censoring proportions, and  $k=2,3$  and 5.

### A comprehensive evaluation of missing data imputation methods

*Krystyna Grzesiak, University of Wrocław, Faculty of Mathematics and Computer Science*

Missing values are ubiquitous in many real-world datasets. While imputation methods are typically evaluated through pointwise errors, the primary goal of imputation is to recover the underlying data distribution. In this work, we present a comprehensive benchmark of modern imputation methods, comparing their ability to reconstruct the original data distribution across a wide range of datasets and missingness mechanisms.

### Independence Testing with Incomplete Data: When Missingness Matters

*Bojana Milošević, Faculty of Mathematics, University of Belgrade*

High-dimensional data are increasingly common in modern applications, including genomics, biomedical studies, and environmental monitoring. In such settings, testing independence is already challenging due to the large number of variables, complex dependence structures, and limited sample sizes. These difficulties become even more pronounced when the data are incomplete. Missingness may substantially affect both the validity and the power of independence tests, while naive imputation may distort the underlying dependence structure.

In this talk, we discuss independence testing for high-dimensional incomplete data and examine when and how missingness

