

Dynamic Borrowing of External Controls in Clinical Trials*Xiaoting Chen, New York University*

In randomized clinical trials, sample size constraints arising from practical and ethical considerations often motivate borrowing external historical information, such as prior randomized studies or Electronic Health Record (EHR) data, to augment the control arm. In the first part of this talk, we review the existing dynamic borrowing methods, including power priors and hierarchical modeling frameworks that adaptively incorporate external historical information. While these approaches can substantially improve estimation efficiency and power, they may induce Type I error inflation under prior-data conflict and are less well-developed for longitudinal data. In the second half, motivated by a recent Ritlecitinib trial for Alopecia Areata, where extrapolating unobserved longitudinal endpoints from a 24-week control to a 36-week treatment target was ethically necessary, we propose a novel dynamic borrowing framework for continuous longitudinal outcomes under a linear mixed-effects (LME) model with slope constraints. Hope our research can stimulate further methodological development on dynamic borrowing that achieves efficiency gains while maintaining the Type I error control in clinical trials.

TabSODA: Tabular Diffusion based Imputation with Skip Pattern Detection and Ordinal Awareness*Yuyu Chen, New York University*

Missing data imputation in large-scale surveys faces two challenges that are not well handled by current tabular diffusion methods. First, structural skips, cells made inapplicable by questionnaire design, should not be imputed but are often conflated with item nonresponse. Second, ordinal responses encode ordered categories, yet most pipelines treat them as nominal levels through one-hot or analog-bit encodings. We introduce TabSODA (Tabular diffusion with Skip pattern detection and Ordinal Awareness), an Expectation-Maximization (EM)-based diffusion imputer built on the Elucidated Diffusion Model (EDM) framework. TabSODA propagates structural skips through the denoising loss and reverse-time sampler, and represents ordinal variables with cumulative-probit scalar latents while retaining analog-bit encodings for nominal variables. When a codebook skip mask is available, TabSODA uses it directly; otherwise, the TabSODA+SKIP variant estimates the mask from raw responses and questionnaire order using a CART-based skip-pattern miner. On Population Assessment of Tobacco and Health (PATH) study and the National Survey on Drug Use and Health (NSDUH), two nationally representative U.S. surveys, TabSODA reduces ordinal MACE by up to 23.7%(relative) and improves categorical accuracy by up to 9% over the strongest baseline across MCAR, MAR, and MNAR masking. The skip miner achieves near-perfect precision on both datasets, allowing TabSODA+SKIP to closely track the codebook-mask variant.

Matching is a widely used causal inference design that aims to approximate a randomized experiment using observational data by forming matched sets of treated and control units based on similarities in their covariates. Ideally, treated units are exactly matched with controls on these covariates, enabling randomization-based inference for treatment effects as in a randomized experiment, under the assumption of no unobserved covariates. However, inexact matching often occurs, leading to residual covariate imbalance. Previous matched studies have typically overlooked this issue and relied on conventional randomization-based inference, as long as some covariate balance criteria are met. Recent research has shown that this approach can introduce significant bias and proposed methods to correct for bias arising from inexact matching in randomization-based inference. These methods, however, are primarily focused on the constant treatment effect and its extensions (i.e., Fisher's sharp null) and do not apply to average treatment effects (i.e., Neyman's weak null). To address this gap, we introduce a new method, inverse post-matching probability weighting, for conducting randomization-based inference for average treatment effects under inexact matching. Our theoretical and simulation results indicate that, compared to conventional randomization-based inference methods, our approach significantly reduces bias and improves coverage rates in the presence of inexact matching.

Advancements in Transfer Learning -- Leveraging Heterogeneous Data for Better Predictions*Iris Zhang, New York University*

Transfer learning has emerged as a powerful framework for improving prediction in settings where the target dataset is limited but related information is available from multiple external sources. A central challenge is that these auxiliary datasets often differ substantially from the target population in terms of covariate distributions, feature spaces, or underlying data-generating mechanisms. Effectively leveraging such heterogeneous data while avoiding negative transfer remains an active area of statistical research. We propose Faster (Feature Alignment and Structured Transfer via Efficient Regularization), a two-stage framework for high-dimensional generalized linear models with partially overlapping feature spaces, in the regime where the total feature dimension may grow much faster than any individual sample size. We establish a high-probability guarantee for sparse recovery in the alignment step and derive an error bound for the final estimator. Simulation studies and real-data analysis demonstrate the effectiveness of this method.

Randomization-Based Inference for Average Treatment Effects in Inexactly Matched Observational Studies.*Jianan Zhu, New York University*