



Workshop on Design of Experiments

April 30 – May 4 2018

CIRM, Marseille, France

Abstracts



Oral presentations

- Anthony C. **Atkinson**, *Experiments for determining non-isothermal kinetic rates*
- Rosemary A. **Bailey**, Peter J. Cameron, *Designs which allow each medical centre to treat only a limited number of cancer types with only a limited number of drugs*
- Julien **Bect**, François Bachoc, David Ginsbourger, *Uncertainty functionals and the greedy reduction of uncertainty*
- Stefanie G.M. **Biedermann**, Rana H. Khashab, Steven G. Gilmour, *Optimal designs for experiments with mixtures*
- Peng Ding, Tirthankar **Dasgupta**, *Randomization based perspectives of randomized block designs and a new test statistic for the Fisher randomization test*
- Valerii V. **Fedorov**, *Optimal designs for dose-response models with partially observed interim/hidden layers*
- Nancy **Flournoy**, Assaf Oron, *Statistical implications of informative dose allocation in binary regression*
- Yohann de Castro, Fabrice **Gamboa**, Didier Henrion, Roxana Hess, Jean-Bernard Lasserre, *Approximate optimal designs for multivariate polynomial regression*
- Bertrand **Gauthier**, *Sampling and spectral approximation*
- Alessandra **Giovagnoli**, Isabella Verdinielli, *Compound utility functions in Bayesian randomized adaptive designs*
- Ulrike **Grömping**, *An algorithm for generating good mixed level factorial designs*
- Radoslav **Harman**, *Computing D-optimal designs of experiments on finite spaces: a survey and comparison of algorithms*
- Ralf-Dieter **Hilgers**, *Evaluation of randomization procedures for clinical trial design optimization with various clinical trial layouts*
- Theodoros Papathanasiou, Anders Strathe, Rune Viig Overgaard, Trine Meldgaard Lund, Andrew C. **Hooker**, *Optimization of dose finding studies for fixed dose combinations using nonlinear mixed-effect models*
- Feifang **Hu**, Wei Ma, Yichen Qin, Yang Li, *Statistical inference of covariate-adaptive randomized studies*
- Volker **Kraft**, *The “When and why?” about definitive screening designs*
- Joachim **Kunert**, Johanna Mielke, *Efficient designs for the estimation of mixed and self carryover effects*

- Sergei L. **Leonov**, Valerii V. Fedorov, *Implementation of algorithms of optimal experimental design on a quantum computer*
- Yuanzhen He, C. Devon **Lin**, Fasheng Sun, *Recent development on design for computer experiment with mixed inputs*
- Dennis K.J. **Lin**, *Design of order-of-addition experiments*
- France **Mentré**, Florence Loingeville, Jeremy Seurat, Thu Thuy Nguyen, *Optimal designs for trials with discrete longitudinal data analyzed by nonlinear mixed effect models*
- Tobias **Mielke**, *Model-based design of dose-finding studies using longitudinal response modelling*
- David C. Woods, Werner G. **Müller**, *Copula-based robust optimal block designs*
- Victor **Picheny**, Tatiana Labopin-Richard, *Sequential design of experiments for estimating quantiles of black-box functions*
- Mohamed R. El Amri, Céline Helbert, Olivier Lepreux, Miguel Munoz Zuniga, Clémentine **Prieur**, Delphine Sinoquet, *Sampling issues for robust inversion*
- William F. **Rosenberger**, *Randomization-based inference: the forgotten component of randomized clinical trials*
- Guillaume **Sagnol**, *Using the S-Lemma to design robust experiments*
- Kirsten **Schorning**, Holger Dette, Katrin Kettelhake, Tilman Möller, *Optimal designs for enzyme inhibition kinetic models*
- Rainer **Schwabe**, *Simplify designs: reduction principles revisited*
- John **Stufken**, *Information-based optimal subdata selection*
- Yuanzhen He, Ching-Shui Cheng, Boxin **Tang**, *Second order saturated designs and strong orthogonal arrays*
- Dariusz **Uciński**, *Optimum experimental design for infinite dimensional inverse problems*
- HaiYing **Wang**, *Statistical inference based on optimal subdata*
- Sonja Surjanovic, William J. **Welch**, *Computer experiments with big n: has Gaussian process computation been tamed?*
- Weng Kee **Wong**, *Optimal experimental designs for complex or high dimensional statistical models*
- Henry P. **Wynn**, Hugo Maruri-Aguillar, *Hilbert series and polynomial models for Smolyak-type sparse grid designs*
- Xin Liu, Rong-Xian **Yue**, *Design admissibility, invariance and optimality in multiresponse linear models*
- Wei **Zheng**, *Optimal design of sampling survey for efficient parameter estimation*
- Anatoly **Zhigljavsky**, Luc Pronzato, *Energy functionals, minimizing measures and kernel herding*

Experiments for determining non-isothermal kinetic rates

ANTHONY C. ATKINSON

London School of Economics, London WC2A 2AE, UK

The optimal design of experiments for the nonlinear models arising in chemical kinetics was introduced by Box and Lucas [1]. One of their examples is first-order decay with rate a function of temperature. They consider a batch reaction which starts from known conditions. The experimental problem is to find the times, and temperatures, at which measurements are to be taken for best estimation of the parameters of the model. The experimental design so found consists of running the experiment on two batches for specified times at specified temperatures and then taking one reading of the concentrations on each batch. However, many industrial experiments afford the opportunity to take a series of readings as the reaction proceeds. The talk will investigate the properties of such designs.

The main example concerns design when there are two consecutive reactions with kinetics following the Arrhenius law, so the model has four parameters. The ‘Box and Lucas’ design requires four runs of the experiment at specific temperatures with a single measurement per batch at a specified time. With a series of readings the experimental requirements can be reduced to two runs at different constant temperatures. A third possibility is one run, in which the temperature follows a specified profile. Optimal designs under these strategies will be compared both for information per experimental run and for information per reading. These extreme assessments provide guidance when detailed information on costs is not available.

Throughout, as in Box and Lucas, the criterion is local D-optimality. The talk will show how the structure of the three classes of design depend upon the duration of the experiment. An extension of the generalized equivalence theorem of optimal experimental design provides insight into the properties of the calculated designs.

References:

[1] G.E.P. Box and H.L. Lucas, “Design of experiments in nonlinear situations”, *Biometrika*, **46**:77–90, 1959.

[Professor A.C. Atkinson; Department of Statistics, London School of Economics, London WC2A 2AE, UK]

[a.c.atkinson@lse.ac.uk]

Designs which allow each medical centre to treat only a limited number of cancer types with only a limited number of drugs

ROSEMARY A. BAILEY

PETER J. CAMERON

University of St Andrews, UK

In order to keep the protocol for a cancer clinical trial simple for each medical centre involved, it is proposed to limit each medical centre to only a few of the cancer types and only a few of the drugs. Let v_1 be the total number of cancer types, and v_2 the total number of drugs. At the workshop on *Design and Analysis of Experiments in Healthcare* at the Isaac Newton Institute, Cambridge, UK in 2015, Valerii Fedorov listed the following desirable properties.

- (a) All medical centres involve the same number, say k_1 , of cancer types, where $k_1 < v_1$.
- (b) All medical centres use the same number, say k_2 , of drugs, where $k_2 < v_2$.
- (c) Each pair of distinct cancer types are involved together at the same non-zero number, say λ_{11} , of medical centres.
- (d) Each pair of distinct drugs are used together at the same non-zero number, say λ_{22} , of medical centres.
- (e) Each drug is used on each type of cancer at the same number, say λ_{12} , of medical centres.

The first four conditions state that, considered separately, the designs for cancer types and drugs are balanced incomplete-block designs (a.k.a. 2-designs) with the medical centres as blocks. We propose calling a design that satisfies all five properties a *2-part 2-design*.

The parameters of a 2-part 2-design satisfy some equations, and also an inequality that generalizes both Fisher's inequality and Bose's inequality.

We give several constructions of 2-part 2-designs, then generalize them to m -part 2-designs.

[R. A. Bailey; School of Mathematics and Statistics, University of St Andrews, North Haugh, St Andrews, Fife KY16 9SS, UK]

[rab24@st-andrews.ac.uk — www-groups.mcs.st-andrews.ac.uk/~rab/]

Uncertainty functionals and the greedy reduction of uncertainty

JULIEN BECT

L2S, Gif-sur-Yvette, France

FRANÇOIS BACHOC

IMT, Toulouse, France

DAVID GINSBOURGER

Idiap, Martigny, Switzerland & IMSV, Bern, Switzerland

The idea of Stepwise Uncertainty Reduction (SUR) has appeared under various names and in various fields (psychophysics, computer vision, machine learning. . .) throughout the eighties and the nineties. More recently, starting with the work of E. Vazquez and co-authors [1, 2], it has been successfully applied to the sequential design of numerical experiments, in particular optimization and reliability analysis, based on Gaussian process priors. In a nutshell, a SUR sequential design greedily minimizes the expected value of some “measure of uncertainty” (e.g., the entropy or variance of some quantity of interest) in order to make it go to zero, hopefully as fast as possible. This talk will present recent results [3] about the almost sure consistency of some SUR sequential designs, in particular under Gaussian process priors, and discuss the properties of uncertainty functionals (i.e., the functionals used to compute quantitative measures of uncertainty from posterior distributions) that make such results possible.

References:

- [1] E. Vazquez and M. Piera-Martinez, “Estimation du volume des ensembles d’excursion d’un processus gaussien par krigeage intrinsèque”, 39^e Journées de Statistique, Angers, France, 2007.
- [2] J. Villemonteix, E. Vazquez and E. Walter, “An informational approach to the global optimization of expensive-to-evaluate functions”, *J. Glob. Optim.*, 44(4):509–534, 2009.
- [3] J. Bect, F. Bachoc and D. Ginsbourger, “A supermartingale approach to Gaussian process based sequential design of experiments”, arXiv:1608.01118v2, 2017.

[Julien Bect; L2S, CentraleSupélec, 3 rue Joliot-Curie, 91192 Gif-sur-Yvette cedex, France.]

[julien.bect@l2s.centralesupelec.fr — www.l2s.centralesupelec.fr/perso/julien.bect]

Optimal designs for experiments with mixtures

STEFANIE G.M. BIEDERMANN

RANA H. KHASHAB

University of Southampton, Southampton, UK

STEVEN G. GILMOUR

King's College London, London, UK

Experiments involving mixtures are conducted in a variety of areas, for example in food processing or in chemical research. The experimental region is constrained naturally, as the proportions of all ingredients have to sum to one. Additional constraints may arise when there are bounds on the proportions, for example a cake must contain a minimum percentage of flour to have the right texture and flavour. Khashab, Gilmour and Biedermann (2018) propose a new class of models to fit the data from such experiments, based on fractional polynomial models (Royston and Altman, 1994). In the talk, I will motivate this modelling approach, and will use a number of historical data sets to compare these models with various other models suggested in the literature. I will then present some optimal designs for these models, and will further discuss some general issues related to designing experiments for mixtures.

References:

- [1] R. H. Khashab, S. G. Gilmour and S. G. M. Biedermann, “Fractional polynomial models for constrained mixture experiments”, *Working paper*, 2018.
- [2] P. Royston and D. G. Altman, “Regression using fractional polynomials of continuous covariates”, *J. Roy. Statist. Soc.*, **C43**(3):429–467, 1994.

[Stefanie Biedermann; University of Southampton, Mathematical Sciences, Highfield Campus, SO17 1BJ, Southampton, UK]

[S.Biedermann@soton.ac.uk — <https://www.southampton.ac.uk/maths/about/staff/sb33.page>]

**Randomization based perspectives of randomized block designs and
a new test statistic for the Fisher randomization test**

PENG DING

University of California, Berkeley, USA

TIRTHANKAR DASGUPTA

Rutgers University, Piscataway, USA

Randomized complete block designs (RCBD) are extensions of matched pair designs when the number of treatments is greater than two. An early application and evaluation of the Fisher randomization test was done under RCBDs, and this is a classical topic that had sparked the famous Neyman-Fisher controversy in a meeting at the Royal Statistical Society in 1935 and still continues to intrigue the statistics community. A recent paper by Sabbaghi and Rubin (2014) has shed some light on this topic. In this paper, we extend the work of Sabbaghi and Rubin (2014) and our previous work on completely randomized designs by providing a more careful theoretical evaluation of the randomization test in an RCBD setting from an asymptotic perspective. Using the potential outcomes framework, we examine the behavior of the classical F statistic under Fishers sharp null hypothesis of no treatment effect on any experimental unit and Neymans null hypothesis of no average treatment effect. It is argued that using the F statistic in the Fisher randomization test under Neymans null does not necessarily yield the correct type-I error. We propose conducting the randomization test with a new Wald-type test statistic, which makes the test exact for Fishers sharp null and asymptotically conservative for Neymans null hypothesis.

References:

[1] A. Sabbaghi and D. B. Rubin, “Comments on the NeymanFisher Controversy and Its Consequences”, *Statist. Sci.*, **29**(2):267–284, 2014.

[2] P. Ding and T. Dasgupta, “A randomization-based perspective on analysis of variance: a test statistic robust to treatment effect heterogeneity”, *Biometrika*, **105**(1):31–44, 2018.

[Tirthankar Dasgupta; Rutgers University, Department of Statistics and Biostatistics, 110 Frelinghuysen Rd, Piscataway, NJ 08854]

[tirthankar.dasgupta — <https://dumptydg.wixsite.com/mysite>]

Optimal designs for dose-response models with partially observed interim/hidden layers

VALERII V. FEDOROV
ICON plc, North Wales, PA, USA

In dose selection/ranging trials a researcher knows the dose given, may/may not measure (PK stage or layer) certain characteristics, such as AUC, $C_{\max}$, or $T_{\max}$, and observes the response(s) to treatment, for example, efficacy and/or toxicity end-points (PD stage or layer). For every stage its own model can be suggested and outputs from the first one can be viewed as candidate inputs for the second stage model. In cases when outcomes of the first stage cannot be observed the setting reminds the neural networks modeling and that partially explains our terminology. In this presentation various designs will be compared and discussed.

[V.V. Fedorov; ICON plc, North Wales, PA, USA]
[Valerii.Fedorov@iconplc.com]

Statistical implications of informative dose allocation in binary regression

NANCY FLOURNOY

University of Missouri, Columbia MO, USA

ASSAF ORON

Institute for Disease Mapping, Bellevue WA, USA

In many fields such acute toxicity studies, Phase I cancer trials, sensory studies and psychometric testing, binary regression techniques are used to analyze data following informative dose allocation. We assume a binary response Y has a monotone positive response probability to a stimulus or treatment X , and we consider designs that sequentially select X values for new subjects in a way that concentrates treatments in a certain region of interest under the dose-response curve. We discuss how data analysis at the end of a study is affected by choosing the stimulus value for each subject sequentially according to some informative sampling rule.

Without loss of generality, we call a positive response a toxicity and the stimulus a dose. For simplicity, we restrict this talk to the case of a univariate treatment X and binary Y , and further assume that treatments are limited to a finite set $\{d_1, d_2, \dots, d_M\}$ of M values we call doses. Now suppose n subjects receive treatments that were sequentially selected (according to some rule using data from prior subjects) from the restricted set of M doses. Let N_m and T_m denote the number of subjects receiving treatment d_m and the number of toxicities observed on treatment d_m , respectively. Define $F_m \equiv P\{Y = 1|X = d_m\} = E[Y|X = d_m]$.

Then it is often said that the distribution of T_m given N_m is Binomial with parameters (F_m, N_m) . But taking N_m as fixed is not the same as conditioning on this random variable, and conditioning on informative dose assignments is not the same as conditioning on summary dose frequencies. In fact, the observed dose-specific toxicity rate, T_m/N_m , is biased for F_m . From first principals, we show unconditionally that

$$E\left[\frac{T_m}{N_m}\right] = F_m - \frac{\text{Cov}[T_m/N_m, N_m]}{E[N_m]}.$$

F_m is a first-order linear approximation to the expected dose-specific toxicity rate, $E[T_m/N_m]$. The observed toxicity rate is biased for F_m because adaptive allocations, by design, induce a correlation between toxicity rates and allocation frequencies.

This bias impacts inference procedures: Isotonic regression methods use dose-specific toxicity rates directly. Standard likelihood-based methods mask the bias by providing first-order linear approximations. We illustrate these biases using isotonic and likelihood-based regression methods in some well known (small sample size) adaptive methods including selected up-and-down designs, interval designs, and the continual reassessment method.

[Nancy Flournoy; University of Missouri]

[flournoy@missouri.edu — <http://web.missouri.edu/~flournoy/>]

Approximate optimal designs for multivariate polynomial regression

YOHANN DE CASTRO
IMO University of Orsay

FABRICE GAMBOA
IMT University of Toulouse

DIDIER HENRION
ROXANA HESS
JEAN-BERNARD LASSERRE
LAAS CNRS Toulouse

We introduce a new approach aiming at computing approximate optimal designs for multivariate polynomial regressions on compact (semi-algebraic) design spaces. We use the moment-sum-of-squares hierarchy of semidefinite programming problems to solve numerically the approximate optimal design problem. The geometry of the design is recovered via semidefinite programming duality theory. This work shows that the hierarchy converges to the approximate optimal design as the order of the hierarchy increases. Furthermore, we provide a dual certificate ensuring finite convergence of the hierarchy and showing that the approximate optimal design can be computed numerically with our method. As a byproduct, we revisit the equivalence theorem of the experimental design theory: it is linked to the Christoffel polynomial and it characterizes finite convergence of the moment-sum-of-square hierarchies

[Fabrice, Gamboa; IMT Université Paul Sabatier 118 Route de Narbonne 31000 Toulouse]
[fabrice.gamboa@math.univ-toulouse.fr — https://www.math.univ-toulouse.fr/~gamboa/newwab/Pages_Fabrice_Gamboa/Main_Page.html]

Sampling and spectral approximation

BERTRAND GAUTHIER

Cardiff University - School of Mathematics, Wales (UK)

We will give an overview of the results presented in [1]. This work addresses the problem of designing sparse quadratures for the approximation of integral operators related to symmetric positive-semidefinite kernels. We more specifically aim at obtaining quadratures leading to an accurate approximation of the main eigendirections of a given initial operator (i.e., the directions related to the largest eigenvalues). A particular attention is paid to the *quadrature-sparsification* problem, which consists in designing sparse quadratures with support included in a fixed finite set of points; this framework in particular encompasses the landmark-selection, or column-sampling, problem for the approximation of large-scale kernel matrices.

We assess the approximation error through the squared Hilbert-Schmidt norm of the difference between the initial and approximate operators, both operators being interpreted as operators acting on the reproducing kernel Hilbert space associated with the kernel considered; we refer to the underlying criterion as the *squared-kernel discrepancy* between the initial and approximate measures; the squared-kernel discrepancy can in addition be interpreted as a “weighted spectral sum-of-squared-errors-type criterion”.

For approximate measures with support included in a fixed finite set of points, the squared-kernel discrepancy can be expressed as a convex quadratic function; sparsity of the approximate measure can then be promoted through the introduction of an ℓ^1 -type *penalisation*, and the induced penalised squared-kernel-discrepancy minimisation problems then consist in convex quadratic minimisation problems. The so obtained quadratic programs can be interpreted as the Lagrange duals of *distorted one-class support-vector machines* (SVM) defined from the squared kernel, the initial measure and the penalisation term; the points selected through penalised squared-kernel-discrepancy minimisation thus correspond to the support vectors of these SVMs.

We pay a special attention to the approximation of the main eigenpairs of an initial operator induced by the eigendecomposition of an approximate operator; to assess the accuracy of an approximate eigendirection while estimating the associated approximate eigenvalue, we in particular rely on the notion of *geometric approximate eigenvalues*. Motivated by the invariance property of the spectral approximations induced by proportional approximate measures, and in order to derive bounds on the overall accuracy of these spectral approximations, we also introduce the notion of *conic squared-kernel discrepancy*.

Numerical strategies for solving large-scale penalised squared-kernel-discrepancy minimisation problems are discussed, and the efficiency of the approach is illustrated by a series of examples. In particular, the ability of the proposed methodology to lead to accurate approximations of the main eigenpairs of kernel matrices related to large-scale datasets is demonstrated.

References:

[1] B. Gauthier and J.A.K. Suykens, “Optimal quadrature-sparsification for integral operator approximation”, *preprint*, <https://hal.archives-ouvertes.fr/hal-01416786/>.

[Bertrand Gauthier; Senghennydd Road, Cardiff, CF24 4AG, United Kingdom]

[GauthierB@cardiff.ac.uk — <https://www.cardiff.ac.uk>]

Compound utility functions in Bayesian randomized adaptive designs

ALESSANDRA GIOVAGNOLI

University of Bologna, Bologna, Italy

ISABELLA VERDINELLI

Carnegie Mellon University, Pittsburgh, USA

Bayesian response-adaptive designs have become popular in the recent literature. They formalize the use of previous knowledge at the planning stage of the experiment, and also allow for a recursive update of the prior information. Randomization, although still debated in the Bayesian theory, also plays a role. Some of the approaches to experiments proposed in the literature combine frequentist and Bayesian frameworks. The inference is frequentist, but a prior probability on the parameters is used to help at the design stage. For example [1] presents ways to combine Bayesian models, utility functions and frequentist analyses in clinical trials. This is the approach of this paper, in which a binary response model on two treatments is considered, with two independent Beta priors on the success probabilities. The outlook is decisional: the treatments are randomized at each step by maximizing the updated expected utility. We propose a utility function that is a trade-off between the acquisition of scientific knowledge, through the experiment, and some ethical or utilitarian gain, typical in clinical trials. Thus our utility is a weighted average of those two quantities, as in [2], [3] and [4]. The double-adaptive design methods introduced in frequentist statistics are those in which the randomization probability depends at each step on the current allocation as well as the target. By extending this approach - in particular the Efficient Randomized Adaptive DEsign (ERADE) by [5] - to the Bayesian theory, we define a doubly-adaptive Bayesian design denoted as BRACE (Bayesian Randomized Adaptive Compound Efficient). The treatment allocation of the BRACE design is shown to converge to the one that maximises the utility function. Theoretical results are supported by numerical simulation studies that compare the behaviour of BRACE to other randomized Bayesian designs.

References:

- [1] S. Venz, G. Parmigiani, and L. Trippa, “Combining Bayesian experimental designs and frequentist data analyses: motivations and examples”, *Appl. Stoch. Model. Bus.*, **33**(3):302–313, 2017.
- [2] A. Baldi-Antognini and A. Giovagnoli, “Compound optimal allocation for individual and collective ethics in binary clinical trials”, *Biometrika*, **97**(4):935–946, 2010.
- [3] I. Verdinelli and J.B. Kadane, “Bayesian designs for maximizing information and outcome”, *J. Am. Statist. Ass.*, **87**:510–515, 1992.
- [4] O. Sverdlov and W.F. Rosenberger, “On recent advances in optimal allocation designs in clinical trials”, *J. Stat. Theor. Prac.*, **7**(4):753–773, 2013.
- [5] F. Hu, L.-X. Zhang, and X. He, “Efficient randomized-adaptive designs”, *Ann. Stat.*, **37**(5A):2543–2560, 2009.

[Alessandra Giovagnoli; University of Bologna, via Belle Arti 41, 40126, Bologna, Italy]
[alessandra.giovagnoli@unibo.it]

An algorithm for generating good mixed level factorial designs

ULRIKE GRÖMPING

Beuth University of Applied Sciences, Berlin, Germany

Mixed integer programming, implemented with R package **DoE.MIParray** using commercial optimizers Gurobi or Mosek, is applied for the creation of “good” mixed level factorial designs, where “good” refers to (possibly partial) generalized minimum aberration, as introduced in [7]. The algorithm is presented in [3]; it improves the algorithm of [1] by exploiting coding invariance results from [2], incorporating lower bounds from [4] and [5], and potentially reducing optimization to short word lengths only.

Usefulness of the algorithm is demonstrated on the biotechnological experiment presented in [6], whose (already quite reasonable) design could have been substantially improved.

References:

- [1] R. Fontana, “Generalized minimum aberration mixed-level orthogonal arrays: A general approach based on sequential integer quadratically constrained quadratic programming”, *Communications in Statistics — Theory and Methods*, **46**:4275–4284, 2017.
- [2] U. Grömping, “Coding Invariance in Factorial Linear Models and a New Tool for Assessing Combinatorial Equivalence of Factorial Designs”, *Journal of Statistical Planning and Inference*, **193**(1):1–14, 2018.
- [3] U. Grömping and R. Fontana, “An Algorithm for Generating Good Mixed Level Factorial Designs”, *Reports in Mathematics, Physics and Chemistry, Report 1/2018, Dep. II, Beuth University of Applied Sciences Berlin*, 1–23, 2018.
- [4] U. Grömping and H. Xu, “Generalized Resolution for Orthogonal Arrays”, *The Annals of Statistics*, **42**(3):918–939.
- [5] M.Q. Liu and D.K.J. Lin, “Construction of Optimal Mixed-Level Supersaturated Designs”, *Statistica Sinica* **19**:197–211, 2009.
- [6] N. Vasilev, Ch. Schmitz, U. Grömping, R., Fischer and S. Schillberg, “Assessment of Cultivation Factors that Affect Biomass and Geraniol Production in Transgenic Tobacco Cell Suspension Cultures”, *PLoS ONE*, **9**(8:e104620):1–7, 2014.
- [7] H. Xu, H. and C.F.J. Wu, “Generalized minimum aberration for asymmetrical fractional factorial designs”. *The Annals of Statistics*, **29**:1066–1077, 2001.

[Ulrike Grömping; Department II, Beuth University of Applied Sciences, Luxemburger Str. 10, D-13353 Berlin, Germany]

[groemping@bht-berlin.de — prof.beuth-hochschule.de/groemping]

Computing D-optimal designs of experiments on finite spaces: a survey and comparison of algorithms

RADOSLAV HARMAN
Comenius University, Bratislava, Slovakia

Computing D-optimal approximate designs of experiments on finite spaces is an important algorithmic problem, with applications in statistics, geometry, and elsewhere. While the first algorithms for D-optimality were developed half a century ago, significant improvements emerged only relatively recently. Modern algorithms allow solving problems with a million of design points in a few seconds.

In the talk, I will briefly discuss the basic principles of the classical and the state-of-the art algorithms. In particular, I will discuss the properties of a new, randomized batch-exchange method (see [1] for details). I will then compare the performance of the algorithms for problems of various sizes and structures, and suggest guidelines for selecting the best D-optimal design algorithm depending on a given situation.

References:

[1] R. Harman, L. Filová and P. Richtárik. A Randomized Exchange Algorithm for Computing Optimal Approximate Designs of Experiments. arXiv:1801.05661, 2018.

[Radoslav Harman; Faculty of Mathematics, Physics and Informatics, Comenius University in Bratislava, Mlynská dolina, 84248 Bratislava, Slovakia]

[harman@fmph.uniba.sk — www.iam.fmph.uniba.sk/ospm/Harman/]

Evaluation of randomization procedures for clinical trial design optimization with various clinical trial layouts

RALF-DIETER HILGERS

RWTH Aachen University, Aachen, Germany

Randomization is a key term in the definition of “randomized controlled trials”. Randomization is used to protect the clinical trial results against bias and justify the statistical model. There are a huge number of randomization procedure which show advantages and disadvantages with respect to various evaluation criteria. But no procedure performs best with respect to all criteria. Consequently research choose the randomization procedure not by a scientific evaluation. Recently Proschan [1] developed a model to describe the impact of bias on the type one error rate in clinical trials und by this connects the practical setting of a clinical trial via the randomization process with the level of evidence resulting from a clinical trial.

Recently Hilgers et al. [2] used this approach, generalizes the setting by inclusion of other bias types and proposed a comprehensive framework for evaluation of randomization procedures for clinical trial design optimization. The framework is exemplified for a single center clinical trial with a 2-arm parallel group design, using an 1:1 allocation ratio without interim analysis and no adaptation in the randomization process where the response is measured with a continuous normal endpoint to prove a superiority hypothesis. The framework includes the derivation of the distribution of the resulting t -test statistic under the model with misspecification and a model for selection and time trend bias.

Within this talk I will generalize the statistical model in two direction. I first will consider a multicenter clinical trial. I still consider a 2-arm parallel group design, using an 1:1 allocation ratio without interim analysis and no adaptation in the randomization process where the response is measured with a continuous normal endpoint to prove a superiority hypothesis. As statistical analysis technique I am using a stratified weighted t -test for statistical analysis. Then I propose a corresponding bias models for selection and time trend bias accounting for the multicenter nature of the trial. I will give the non-centrality parameters of the stratified weighted t -test statistic under the misspecification model which is doubly noncentral t . With this results it is possible to calculate the actual type I error rate for each allocation sequence resulting form the respective randomization procedure in the presence of selection and/or time trend bias. And finally it allows us depending on the amount of bias we would like to control, to select the best performing randomization procedure. The calculation is facilitated by using the R package RandomizeR [3]. I will present some numerical results.

References:

- [1] Proschan M. Influence of selection bias on type I error rate under random permuted block designs. *Statistica Sinica* 1994; 4: 219-231.
- [2] Hilgers RD, Uschner D, Rosenberger WF, Heussen N. ERDO - a framework to select an appropriate randomization procedure for clinical trials. *BMC Med Res Methodol.* 2017 Dec 4; **17(1)**:159. doi: 10.1186/s12874-017-0428-z.
- [3] Uschner D, Schindler D, Hilgers RD, Heussen N. randomizeR: An R Package for the Assessment and Implementation of Randomization in Clinical Trials. *Journal of Statistical Software* 2017 accepted.

[Ralf-Dieter Hilgers; Department of Medical Statistics, RWTH Aachen University, Pauwelsstrasse 19, D - 52074 Aachen]

[rhilgers@ukaachen.de — www.ideal.rwth-aachen.de]

Optimization of dose finding studies for fixed dose combinations using nonlinear mixed-effect models

THEODOROS PAPATHANASIOU
University of Copenhagen, Copenhagen, Denmark
Novo Nordisk A/S, Søborg, Denmark

ANDERS STRATHE
RUNE VIIG OVERGAARD
Novo Nordisk A/S, Søborg, Denmark

TRINE MELDGAARD LUND
University of Copenhagen, Copenhagen, Denmark.

ANDREW C. HOOKER
Uppsala University, Uppsala, Sweden

Drug therapies are increasingly becoming more targeted in their delivery to the body, which may improve the specificity of where and how a drug is active in the body, but may be suboptimal in a physiological system that has evolved to be regulated by a multiplicity of pathways. Thus, combinations of targeted treatments are being investigated for potentially higher clinical benefit, especially when the combined drugs act via synergistic interactions. The clinical development of combination treatments is particularly challenging, especially during the dose selection phase, where a vast range of possible combination doses exist. It has previously been shown that dose selection can be improved through the modeling of exposure-response (E-R) relationships of combinatory drug effects through population pharmacokinetic-pharmacodynamic (PKPD) drug interaction models [1]. However, as shown in [1], the study design is important in correctly characterizing these models used for dosing decisions. Traditionally, drug combination studies are conducted based on factorial designs and variations thereof (i.e. fractional factorial). While simple in their conception and construction, the choice of the investigated dose levels is often empirical. In this work we investigate how dose selection can be optimized in drug combination studies through the use of these population PKPD drug interaction models. We explore local optimal designs (D- and Ds-optimality) to maximize the precision of model parameters in a number of potential E-R surfaces. We also consider a compound criterion (D/V-optimality) to optimize the precision of model predictions in specific parts of the E-R surfaces [2]. Finally, for the most promising local designs, globally optimal design criteria are explored. We find the compound criterion designs investigated to be a promising way forward for combination therapy studies.

References:

- [1] Papathanasiou, P., Strathe, A., Hooker, A.C., Lund, T.M., and Overgaard, R.V., “Feasibility of Exposure-Response analyses for clinical dose-ranging studies of drug combinations”, *J. AAPS*, Accepted, 2018.
- [2] Miller, F., Guilbaud, O., and Dette, H. . “Optimal designs for estimating the interesting part of a dose-effect curve”. *Journal of Biopharmaceutical Statistics*, **17**(6), 1097-115, 2007.

[Andrew C. Hooker; Dept. of Pharmaceutical Biosciences, Uppsala University, Box 591, 751 24, Uppsala, Sweden]
[andrew.hooker@farmbio.uu.se — <http://katalog.uu.se/profile/?id=N4-631>]

Statistical inference of covariate-adaptive randomized studies

FEIFANG HU

George Washington University, Washington, DC, USA

WEI MA

Renmin University of China, Beijing, China

YICHEN QIN

University of Cincinnati, Cincinnati, USA

YANG LI

Renmin University of China, Beijing, China

Covariate-adjusted randomization is frequently used in comparative studies, such as clinical trials and causal inference. However, since the randomization inevitably uses the covariate information when forming balanced treatment assignments, the validity of classical statistical inference following such randomization is often unclear. In this talk, we derive the theoretical properties of statistical inference post general covariate-adjusted randomization under the linear model framework. More important, we explicitly unveil the relationship between covariate-adjusted and inference properties. We apply the proposed general theory to various randomization procedures including complete randomization (CR), re-randomization (RR), pairwise sequential randomization (PSR), and Atkinson's DA- optimality biased coin design (DA-BCD) and compare their performance analytically. We then proposed a new adjusted approach to obtain valid and more powerful tests. These results open a new door to understand and analyze comparative studies based on covariate-adjusted randomization. Simulation studies provide further evidence of the advantages of the proposed framework and theoretical results. This talk is based on joint research with Wei Ma, Yichen Qin and Yang Li.

References:

- [1] J. Hu, H. Zhu and F. Hu, "A Unified Family of Covariate-Adjusted Response-Adaptive Designs Based on Efficiency and Ethics". *Journal of the American Statistical Association*, **110**, 357–367, 2015.
- [2] Y. Hu and F. Hu, "Asymptotic Properties of Covariate-Adaptive Randomization" *Annals of Statistics*. **40** , 1794–1815, 2012.
- [3] W. Ma, F. Hu and L.X. Zhang, "Testing Hypotheses of Covariate-Adaptive Randomized Clinical Trials", *Journal of the American Statistical Association*. **110**, 669–680, 2015.

[Feifang Hu; Department of Statistics, George Washington University, 801 22nd St NW, Washington, DC, 20052]

[feifang@gwu.edu — <https://statistics.columbian.gwu.edu/feifang-hu>]

The “When and why?” about definitive screening designs

VOLKER KRAFT

SAS Institute - JMP Division, Heidelberg, Germany

For many good reasons, during recent years Definitive Screening Designs (DSD) have been embraced by many practitioners in industry and academia. In addition to modern design tools implementing DSDs, analysis techniques allow to estimate both main effects and second-order effects. Does it mean that this type of design became a new de-facto standard for designed experiments? One the one hand, you can learn a lot from DSDs at a minimal budget. One the other hand, limitations exist concerning the types of possible factors and the choice of models for data analysis. We demonstrate several design and analysis examples and discuss the limitations and advantages of this approach.

[Volker Kraft; SAS Institute - JMP Division]

[volker.kraft@jmp.com — www.jmp.com]

Efficient designs for the estimation of mixed and self carryover effects

JOACHIM KUNERT

Department of Statistics, TU Dortmund, Germany

JOHANNA MIELKE

Novartis Pharma AG, Basel, Switzerland

Biosimilars are copies of biological medicines that are developed by a competitor after the patent for the originator drug has expired. Extensive clinical trials are required to show therapeutic equivalence between the biosimilar and its reference product before a biosimilar can be sold on the market. However, even after more than 10 years of experience with biosimilars in Europe, there is still some uncertainty if the patients who are already taking the reference product can switch between the biosimilar and its reference product (see e.g. Ebbers et al, 2012). One convenient way to assess the impact of switches is the analysis of mixed and self carryover effects: if the products are switchable, there should not be any difference in the carryover effects. This paper determines a series of simple designs which are highly efficient for the comparison of the mixed and self carryover effects of two treatments.

Efficient designs for the estimation of direct effects have been determined by Kunert and Stufken (2008). It turns out that the determination of efficient designs for the estimation of carryover effects is harder, because the information matrix for the estimation of carryover effects is not completely symmetric, not even for the best designs.

References:

- [1] H.C. Ebbers, S.A. Crow, A.G. Vulto and H. Schellekens, “Interchangeability, immunogenicity and biosimilars”, *Nature Biotechnology*, **30**: 1186–1190, 2012.
- [2] J. Kunert and J. Stufken, “Optimal crossover designs for two treatments in the presence of mixed and self carryover effects”, *J. Amer. Statist. Assoc.*, **97**: 989–906, 2008.

[Joachim Kunert; Dept. of Statistics, TU Dortmund, 44221 Dortmund, Germany]

[joachim.kunert@udo.edu — www.statistik.tu-dortmund.de/kunert0.html]

Implementation of algorithms of optimal experimental design on a quantum computer

SERGEI L. LEONOV

ICON Clinical Research, North Wales, PA, USA

VALERII V. FEDOROV

ICON Clinical Research, North Wales, PA, USA

Iterative algorithms for the construction of optimal experimental designs for regression models and specifics of their implementation on a quantum computer are discussed. We present several examples of optimal model-based design problems originating in biopharmaceutical applications and show how to reformulate these problems as quadratic unconstrained binary optimization problems (QUBO) that can be solved on a D-Wave quantum annealer; see

<https://www.dwavesys.com/quantum-computing>

[Sergei Leonov; ICON Clinical Research, 2100 Pennbrook Parkway, North Wales, PA 19454, USA]
[Sergei.Leonov@iconplc.com]

Recent development on design for computer experiment with mixed inputs

YUANZHEN HE

School of Statistics, Beijing Normal University, Beijing, China

C. DEVON LIN

Department of Mathematics and Statistics, Queen's University, Kingston, Ontario, Canada

FASHENG SUN

KLAS and School of Mathematics and Statistics, Northeast Normal University, Changchun, Jilin, PR China

Computer experiments with qualitative and quantitative factors occur frequently in various applications in science, engineering and business. Marginally coupled designs were introduced to accommodate such experiments in a more efficient, less costly way. Some basic, general constructions of such designs are proposed. Further methods for improving projection and overall space-filling properties are introduced. When the designs for qualitative factors are two-level or multi-level orthogonal arrays, constructions based on subspace theory are proposed. The theoretical results on the proposed constructions are derived. For practical use, some constructed designs for two-level and three-level qualitative factors are tabulated.

[C. Devon Lin; Department of Mathematics and Statistics, Queens University, Kingston, ON, K7L 3N6, Canada]

[devon.lin@queensu.ca — www.mast.queensu.ca/~cdlin]

Design of order-of-addition experiments

DENNIS K.J. LIN
Department of Statistics, The Pennsylvania State University, University Park,
PA, USA

In Fisher (1971), a lady was able to distinguish (by tasting) from whether the tea or the milk was first added to the cup. This is probably the first popular order of addition experiment. In general, there are m required components and we hope to determine the optimal sequence for adding these m components one after another. Knowing the optimal order of addition of components related in production is crucial. It is often unaffordable to test all the $m!$ treatments, and the design problem arises (note that when $m=10$, for example, $m!$ is about 3.5 million). We consider the model in which the response of a treatment depends on the pairwise orders of the components. The optimal design theory under this model is established, and the optimal values of the D-, A-, E-, and M:S:-criteria are derived. We identify a special constraint on the correlation structure of such designs. The closed-form construction of a class of optimal designs is obtained, with examples for illustration.

[Dennis K. J. Lin; 317 Thomas Building, University Park, PA 16802]
[dk15@psu.edu — <http://stat.psu.edu/people/dk15>]

**Optimal designs for trials with discrete longitudinal data analyzed
by nonlinear mixed effect models**

FRANCE MENTRÉ

FLORENCE LOINGEVILLE

JEREMY SEURAT

THU THUY NGUYEN

IAME, UMR 1137, INSERM and University Paris Diderot, Paris, France

Trials with discrete longitudinal data can be analyzed with nonlinear mixed effect models (NLMEM). To derive optimal designs or to compare several designs we can use the expected Fisher information matrix (FIM). We developed a new method to evaluate the FIM for NLMEM with discrete data based on Monte-Carlo Hamiltonian Monte-Carlo (MC/HMC) [1]. We implemented it in the R package MIXFIM using R-Stan for HMC sampling. This approach requires a priori knowledge on models and parameters to compute the FIM, leading to locally optimal designs. We extended this MC/HMC-based method to account for uncertainty in parameters and/or models. When introducing uncertainty on the population parameters, we evaluated the robust FIM as the expectation of the FIM computed by MC/HMC. To account for several candidate models, we used the compound D-optimality criterion [2, 3].

We illustrated this approach on two examples: one with repeated count data and one with repeated binary data. For the first example, we assumed a longitudinal Poisson count model, where the event rate parameter depends on dose through four candidate models. For the second example, we assumed a longitudinal binary data model, where the logit depends on time and treatment group through four candidate models. For each example, using MC/HMC to compute the FIM, we performed combinatorial optimization to define optimal doses or sampling times, respectively. First we derived D-optimal designs for each candidate model and then the CD-optimal design across the four models assuming equal a priori weights. We showed that misspecification of models could lead to low D-efficiencies and that CD-optimal designs provided a good compromise for the different candidate models. The proposed design strategy, based on MC/HMC and compound optimality theory, is a relevant approach which can be used to efficiently design longitudinal studies accounting for model uncertainty.

References:

- [1] M.K. Riviere, S. Ueckert and F. Mentré, “An MCMC method for the evaluation of Fisher information matrix for non-linear mixed effect models”, *Biostatistics*, **17**(4):735–750, 2016.
- [2] A. Atkinson, “DT-optimum designs for model discrimination and parameter estimation”, *Journal of Statistical Planning and Inference*, **138**(1):56–64, 2008.
- [3] T.T. Nguyen, H. Banech, M. Delaforge and N. Lenuzza, “DT-optimum designs for model discrimination and parameter estimation”, *Pharmaceutical Statistics*, **15**(2):165–177, 2016.

[France Mentré; IAME, UMR1137, UFR de Médecine, Université Paris Diderot, 16 rue Henri Huchard, 75018 Paris, France]
[france.mentre@inserm.fr]

Model-based design of dose-finding studies using longitudinal response modelling

TOBIAS MIELKE
Janssen-Cilag, Germany

Although regulatory guidance emphasizes the importance of dose-response modelling since more than 20 years, dose-finding in clinical drug development is still mainly based on pairwise comparisons to a control. Implications on study designs include small numbers of doses tested, to reduce multiplicity penalty. Only promising doses are then typically considered, increasing the likelihood of success for at least one of these. These designs are weak for dose-response modelling, where optimized designs would emphasize an allocation to doses on the steep part of the curve. Efficient methods for model-based testing against a flat dose-response have been published by Bornkamp et al.[1] and qualified by the EMA[2] in recent years. This has opened the door wider for the application of optimum experimental designs in drug development.

As dose-response relations are typically non-linear in some of the parameters, optimized study designs will be generally only locally optimal and might hence lose efficiency for a true, but unknown dose-response relation. Interim analyses allow to update the design using interim parameter estimates (Dragalin et al. [3]). The benefit of interim design adjustments depends on the amount of interim information available and hence the quality of the initial design. A true optimal design may only be worsened using interim design adjustments. Similarly, weak study designs will easily benefit from interim adjustments. In practice, initial designs will be neither optimal nor heavily suboptimal, such that an early interim analyses could lead to false design adjustments, whereas late interim analyses will only lead to limited benefit.

Patients are continuously recruited into dose-finding studies and followed-up for several months to describe the dose-response on the endpoint at a certain time, e.g. 12 weeks after start of treatment. Due to various operational reasons, a pause of recruitment to wait for interim data is not desired in these studies. This reduces the time window and available data for the conduct of interim analyses. Longitudinal modelling allows in this situation to increase information at the interim analyses, supporting interim decision making.

Dragalin[4] described adaptive designs in dose-finding using an integrated two-component model for longitudinal modelling of the dose-response introduced by Fu and Manner[5]. In this presentation, we extend the work from Dragalin to discuss timing of the interim analysis and the impact of modelling assumptions and information approximations on the efficiency of study designs.

References:

- [1] Bornkamp, B., Pinheiro, J., and Bretz, F. “MCPMod: An R Package for the Design and Analysis of Dose-Finding Studies” *J. of Statistical Software*, **29** 1–23, 2010.
- [2] Committee for Medicinal Products for Human Use (CHMP) “Qualification opinion of MCP-Mod as an efficient statistical methodology for model-based design and analysis of Phase II dose finding studies under model uncertainty”, European Medicines Agency, EMA/CHMP/SAWP/757052/2013, 2014.
- [3] Dragalin, V., Hsuan, F., Padmanabhan, S.K. “Adaptive designs for dose-finding studies based on sigmoid EMax model”, *J. Biopharm.Stat.*, **17** 1051–1070, 2007.
- [4] Dragalin, V., “Optimal Design of Experiments for Delayed Responses in Clinical Trials”, In: Uciniski, D., Atkinson, A., Patan, M. (eds.) *mODa 10 - Advances in Model-Oriented Design and Analysis*, 55–61, Springer, 2013.
- [5] Fu, H., Manner, D., “Bayesian adaptive dose-finding studies with delayed responses”, *J. Biopharm.Stat.*, **21** 888–901, 2011.

Copula-based robust optimal block designs

DAVID C. WOODS

University of Southampton, UK

WERNER G. MÜLLER

Johannes Kepler University Linz, Austria

Blocking is often used to reduce known variability in designed experiments by collecting together homogeneous experimental units. A common modelling assumption for such experiments is that responses from units within a block are dependent. Accounting for such dependencies in both the design of the experiment and the modelling of the resulting data when the response is not normally distributed can be challenging, particularly in terms of the computation required to find an optimal design. The application of copulas and marginal modelling provides a computationally efficient approach for estimating population-average treatment effects. Motivated by an experiment from materials testing, we develop and demonstrate designs with blocks of size two using copula models. Such designs are also important in applications ranging from microarray experiments to experiments on human eyes or limbs with naturally occurring blocks of size two. We present methodology for design selection, make comparisons to existing approaches in the literature and assess the robustness of the designs to modelling assumptions.

[Werner G. Müller; JKU Linz, Dept. of Applied Statistics, Altenberger Straße 69, 4040 Linz, Austria]

[werner.mueller@jku.at — www.jku.at/ifas]

Sequential design of experiments for estimating quantiles of black-box functions

VICTOR PICHENY

MIAT, Université de Toulouse, INRA, Castanet-Tolosan, France

TATIANA LABOPIN-RICHARD

Institut de Mathématiques de Toulouse, Université Paul Sabatier, Toulouse, France

We consider here the question of estimating a quantile of a random variable $g(X)$, where $g : \mathbb{X} \subset \mathbb{R}^d \rightarrow \mathbb{R}$ is a “black-box” function and X is a random vector of known distribution. Typically, g stands for a computer model and X for its uncertain multivariate input. We assume that g is accessible only through a limited dataset $\{g(x_1), \dots, g(x_n)\}$, which rules out a simple Monte-Carlo (MC) strategy. A classical approach in this scenario is to rely on metamodels: the observation set is used to build a fast-to-evaluate approximation $\hat{q}$ of g , on which a simple MC method can be conducted to estimate the quantile.

Designing the experiment $\{x_1, \dots, x_n\}$ is critical to maximize the quality of such estimate. We focus here on the Gaussian process (GP) metamodel, which has the advantage of being particularly well-suited for sequential sampling (i.e. adding observations one at a time, using the metamodel to guide the process). We build up from promising previous work [1,3] to present a novel procedure tailored to quantile estimation in the form of a classical Bayesian Optimization algorithm, with sequential maximization of an *acquisition function*. This function evaluates the value of a new candidate sample by estimating by how much it would affect the quantile estimates. A particular attention is given to the formulation of the function and implementation issues to keep it numerically tractable, even when the problem dimension increases. The proposed strategy is tested on several numerical examples with up to six uncertain parameters, showing that accurate estimators can be obtained using only small designs of experiments.

This abstract is excerpt from [2], which contains the full mathematical background, method description and numerical results.

References:

- [1] M. Jala, C. Lvy-Leduc, E. Moulines, E. Conil, and J. Wiart, “Sequential design of computer experiments for the assessment of fetus exposure to electromagnetic fields”, *Technometrics*, **58**(1): 30–42, 2016.
- [2] T. Labopin-Richard, V. Picheny, “Sequential design of experiments for estimating quantiles of black-box functions”, *Statistica Sinica*, 2017 (*to appear*).
- [3] J. Oakley, “Estimating percentiles of uncertain computer code outputs”, *J. Roy. Statist. Soc.*, **C53**(1), 83–93, 2004.

[Victor Picheny; MIAT, INRA, Auzeville, 31326 CASTANET TOLOSAN, FRANCE]

[victor.picheny@inra.fr — <https://sites.google.com/site/victorpicheny/>]

Sampling issues for robust inversion

MOHAMED R. EL AMRI^(*); CÉLINE HELBERT⁽⁺⁾; OLIVIER LEPREUX^(‡);
 MIGUEL MUNOZ ZUNIGA^(‡); CLÉMENTINE PRIEUR^(#); DELPHINE SINOQUET^(‡)
 (*) *Université Grenoble Alpes et IFP Énergies nouvelles*; (+) *École Centrale de*
Lyon; (‡) *IFP Énergies nouvelles, Solaize*; (#) *IFP Énergies nouvelles, Rueil*;
 (#) *Université Grenoble Alpes*

In the present talk, we consider a system evolving in an uncertain environment. That system is modeled by a numerical simulator f , whose inputs are of two types: a set of control variables $x \in \mathbb{X}$, and a set of uncertain variables $v \in \mathcal{V}$. More precisely, $f : \mathbb{X} \times \mathcal{V} \rightarrow \mathbb{R}_+$. Robust inversion consists in seeking the set of control variables $x \in \mathbb{X}$ such that $\sup_{v \in \mathcal{V}} f(x, v)$ is bounded by a prescribed threshold $c > 0$. Then, the difficulty of solving the robust inversion problem strongly depends on the uncertainty set $\mathcal{V}$. In our framework, the uncertainty set $\mathcal{V}$ is a functional space. We also assume that the uncertainty has a probabilistic description: the uncertainty is modeled by a random variable V valued in $\mathcal{V}$. We then consider the following stochastic inversion problem: we are seeking the set $\Gamma^* := \{x \in \mathbb{X}, g(x) := \mathbb{E}_V[f(x, V)] \leq c\}$. In our setting, the probability distribution of V is only known from a set of M realizations $\{v_1, \dots, v_M\}$.

A Stepwise Uncertainty Reduction (SUR) strategy aims at constructing a sequence $x_1, x_2, \dots$ of evaluation points of g in such a way that the residual uncertainty about Γ^* given the information provided by the evaluation results is small. More precisely, SUR strategies are based on three main ideas [1]. The first (Bayesian) idea is to consider g as a sample path of a random process, which is assumed Gaussian for the sake of tractability. Doing so entails that any quantity depending on g is formally a random variable. The second idea is to introduce a measure of the uncertainty about the quantity of interest conditioned on the σ -algebra $\mathcal{A}_n$ generated by $\{(x_i, g(x_i)), 1 \leq i \leq n\}$. We will denote by $H_n(g)$ such a measure of uncertainty, which is an $\mathcal{A}_n$ -measurable random variable. In the context of robust inversion, its choice is based on the theory of random closed sets [2]. The third idea is to choose evaluation points sequentially in order to minimize at each step n the expected value of the future uncertainty $H_{n+1}(g)$ with respect to the random outcome of the new evaluation of g :

$$x_{n+1} = \arg \min_{x \in \mathbb{X}} J_n(x) := \mathbb{E}_n(H_{n+1}(g) \mid x_{n+1} = x)$$

where $\mathbb{E}_n(\cdot)$ stands for the conditional expectation $\mathbb{E}(\cdot \mid \mathcal{A}_n)$.

The key contribution of the present work is to adapt the aforementioned SUR strategy to our setting where g is defined as $g(x) = \mathbb{E}_V[f(x, V)]$ with V a random variable taking its values in a functional space $\mathcal{V}$. The expectation is estimated by an average on a set of m realizations of V among the M realizations $\{v_1, \dots, v_M\}$. We propose an adaptive strategy, sampling alternatively one point x in the control set $\mathbb{X}$ with the SUR strategy and m realizations of V in the set $\{v_1, \dots, v_M\}$ using an innovative space filling strategy.

Our new procedure will be tested on an analytical test case, as far as on an industrial application from the French Oil Institute (IFP Énergies nouvelles).

References:

- [1] C. Chevalier, J. Bect, D. Ginsbourger, E. Vazquez, V. Picheny and Y. Richet, “Fast parallel kriging-based stepwise uncertainty reduction with application to the identification of an excursion set”, *Technometrics*, **56**(4):455–465, 2014.
- [2] I.S. Molchanov, *Theory of random sets*, London: Springer **19**(2), 2005.

[Clémentine Prieur; Grenoble Alpes University]

[clementine.prieur@univ-grenoble-alpes.fr — <https://ljk.imag.fr/membres/Clementine.Prieur/>]

Randomization-based inference: the forgotten component of randomized clinical trials

WILLIAM F. ROSENBERGER
George Mason University, Fairfax, VA, USA

Randomization has been the hallmark of the clinical trial since Sir Bradford Hill introduced it in the 1948 streptomycin trial. An exploration of the early literature yields three rationales: (1) the incorporation of randomization provides unpredictability in treatment assignments, thereby mitigating selection bias; (2) randomization tends to ensure comparability in the treatment groups on known and unknown confounders (at least asymptotically); and (3) the act of randomization itself provides a basis for inference when random sampling is not conducted from a population model. Of these three, rationale (3) is often forgotten, ignored, or left untaught. And yet, since the dawn of statistics, it has been recognized that randomized experiments cannot use statistical techniques developed for random sampling from a population. Today, randomization is a rote exercise, scarcely considered in protocols or medical journal articles. Randomization was done by Excel is a standard sentence that serves to check the box that investigators specify how they conducted the randomization.

In this talk, we review the history of randomization as a basis for inference and describe how randomization-based inference can be used for virtually any outcome of interest in a clinical trial. We conclude that randomization matters!

[William F. Rosenberger; University Professor and Chairman, Department of Statistics, George Mason University, 4400 University Drive MS4A7, Fairfax, VA 22032 U.S.A.]
[wrosenbe@gmu.edu — statistics.gmu.edu]

Using the S-Lemma to design robust experiments

GUILLAUME SAGNOL

Technical University Berlin, Germany

We address the problem of designing optimal experiments when the data is uncertain. This situation includes –but is not limited to– the design of nonlinear experiments: here the information matrix depends on unknown parameters, so we have to rely on computationally expensive Bayesian or minimax design approaches. In particular, one difficulty to compute minimax designs is that, when a design ξ is given, the problem of finding the worst combination of parameters in some plausible region is an NP-hard optimization problem.

Building on work from [1-3] on robust estimators, we propose to design experiments that are robust to deviations of the design matrix from a nominal scenario. Rather than relying on a parametric model for the design matrix (the standard approach for nonlinear experiments), our uncertainty set is defined as a ball of design matrices relative to the spectral norm. By using the celebrated S-lemma [4], we formulate the problem of simultaneously computing a robust design and the corresponding robust linear estimator by semidefinite programming. This approach is a tractable alternative to minimax optimal designs, and seems appealing when the number of uncertain parameters is large.

References:

- [1] L. El Ghaoui and H. Le Bret, “Robust solutions to least-squares problems with uncertain data”, *SIAM Journal on matrix analysis and applications*, **18.4** (1997): 1035-1064.
- [2] G. Calafiore and L. El Ghaoui, “Robust maximum likelihood estimation in the linear model”, *Automatica*, **37.4** (2001): 573-580.
- [3] Y.C. Eldar, A. Ben-Tal and A. Nemirovski, “Robust mean-squared error estimation in the presence of model uncertainties”, *IEEE Transactions on Signal Processing*, **53.1** (2005): 168-181.
- [4] I. Plik and T. Terlaky, “A survey of the S-lemma”, *SIAM review*, **49.3** (2007): 371-418.

[Guillaume Sagnol; Institut für Mathematik, Sekr. MA 5-2, Technische Universität Berlin, Straße des 17. Juni 136, 10623 Berlin, Germany]

[sagnol@math.tu-berlin.de — <http://page.math.tu-berlin.de/~sagnol/>]

Optimal designs for enzyme inhibition kinetic models

KIRSTEN SCHORNING

HOLGER DETTE

KATRIN KETTELHAKE

TILMAN MÖLLER

Ruhr-University Bochum, Bochum, Germany

We consider a new method for determining optimal designs for enzyme inhibition kinetic models, which are used to model the influence of the concentration of a substrate and an inhibition on the velocity of a reaction. The approach uses a nonlinear transformation of the vector of predictors such that the model in the new coordinates is given by an incomplete response surface model. Although there exist no explicit solutions of the optimal design problem for incomplete response surface models so far, the corresponding design problem in the new coordinates is substantially more transparent, such that explicit or numerical solutions can be determined more easily. The designs for the original problem can finally be found by an inverse transformation of the optimal designs determined for the response surface model. We illustrate the method determining explicit solutions for the D -optimal design and for the optimal design problem for estimating the individual coefficients in a non-competitive enzyme inhibition kinetic model.

[Kirsten Schorning; Ruhr-University Bochum]

[`kirsten.schorning@rub.de`]

Simplify designs: reduction principles revisited

RAINER SCHWABE

Otto-von-Guericke-University, Magdeburg, Germany

Based on convex optimization the general theory of optimal design is well developed. However, in practice for every non-standard statistical situation an individual optimal solution still has to be computed which may be challenging in the case of high dimensions and/or nonlinear relationships.

While a diversity of algorithmic approaches is available ranging from steepest descent, multiplicative, and quasi-Newton to metaheuristic algorithms like particle swarm optimization involving high computational efforts, there may be still interest in analytical solutions or in reduction of the complexity of the problem to decrease the computational burden or to obtain exact benchmarks on the quality of competing designs.

Reduction principles in the construction of optimal designs we revisit here will be invariance and equivariance, majorization, and reduction to lower-dimensional problems. The applicability of these general concepts will be exhibited in a variety of non-standard experimental situations.

This is joint work with Fritjof Freise (Technical University Dortmund), Osama Idais, Eric Nyarko, Maryna Prus, Martin Radloff, Frank Röttger, Dennis Schmidt, and Marius Schmidt (all University of Magdeburg).

[Rainer Schwabe; Otto-von-Guericke-Universität, Universitätsplatz 2, D-39 106 Magdeburg, Germany]

[rainer.schwabe@ovgu.de — www.imst3.ovgu.de]

Information-based optimal subdata selection

JOHN STUFKEN

Arizona State University, Tempe, USA

The focus of this presentation is on the analysis of data with a very large number of cases, n , and a modest number of variables, p . The size of n , or the lack of access to a sufficiently powerful computing platform, might necessitate or suggest the use of subdata, which consists of only some of the cases. How should one select such subdata?

In the linear regression context, several subsampling methods have been proposed (cf. Meng et al., 2017) to obtain such subdata, along with appropriate estimation methods. Such methods have been referred to as algorithmic leveraging methods.

Also in the context of linear regression, Wang et al. (2018) proposed a deterministic method for subdata selection, referred to as Information-Based Optimal Subdata Selection (IBOSS). This method borrows ideas from design of experiments to select subdata that provides “maximum information”.

In this presentation, I will briefly introduce the different methods for subdata selection, with an emphasis on the IBOSS method. Selected results and comparisons from Wang et al. (2018) will be presented. I will conclude with a brief discussion of remaining challenges in this area of research.

References:

- [1] C. Meng, Y. Wang, X. Zhang, A. Mandal, W. Zhong and P. Ma, “Effective Statistical Methods for Big Data Analytics”, in *Handbook of Research on Applied Cybernetics and Systems Science*, eds. S. Saha, A. Mandal, A. Narasimhamurthy, V. Sarasvathi, S. Sangam, 280–299, 2017.
- [2] H. Wang, M. Yang and J. Stufken, “Information-Based Optimal Subdata Selection for Big Data Linear Regression”, *J. Americ. Statist. Assoc.*, to appear, 2018. (<https://doi.org/10.1080/01621459.2017.1408468>)

[John Stufken; Arizona State University, School of Mathematical and Statistical Sciences, PO Box 871804, Tempe, AZ 85287, USA]

[John.Stufken@asu.edu — <https://math.la.asu.edu/~jstufken/>]

Second order saturated designs and strong orthogonal arrays

YUANZHEN HE

Beijing Normal University, Beijing, China

CHING-SHUI CHENG

University of California, Berkeley, USA

BOXIN TANG

Simon Fraser University, Burnaby, Canada

Strong orthogonal arrays were recently introduced and studied as a class of space-filling designs for computer experiments. To enjoy the benefits of better space-filling properties, when compared to ordinary orthogonal arrays, strong orthogonal arrays need to have strength three or higher, which may require run sizes that are too large for experimenters to afford. To address this problem, we introduce a new class of arrays, called strong orthogonal arrays of strength two plus. These arrays, while being more economical than strong orthogonal arrays of strength three, still enjoy the better two-dimensional space-filling property of the latter. Among the many results we have obtained on the characterizations and construction of strong orthogonal arrays of strength two plus, worth special mention is their intimate connection with second order saturated designs.

[Boxin Tang; Simon Fraser University, Dept of Statistics and Acturial Science, Burnaby, BC V5A 1S6, Canada]

[boxint@sfu.ca — www.stat.sfu.ca/~boxint]

Optimum experimental design for infinite dimensional inverse problems

DARIUSZ UCİŃSKI
University of Zielona Góra, Zielona Góra, Poland

A great difficulty in parameter estimation of distributed parameter systems, i.e., systems described by partial differential equations, is the inability to observe the system states over the entire spatial domain. This leads to the question of where to locate sensors so that the information content of the resulting measurements be as high as possible. As ‘best’ sensor positions have to be determined prior to actual data collection, the choice of the appropriate optimality criterion becomes of paramount importance. Sensor location for parameter estimation usually follows the traditional approach of statistical experimental design and is based on various performance measures defined on the Fisher information matrix associated with the estimated parameters. An overview of this currently very active research area is contained in the monograph [1].

Over the last decade communications about sensor locations have continued to grow. A very prospective direction, which is stimulated by data assimilation for air quality monitoring, constitutes inclusion of the unknown initial state as an additional unknown parameter. The infinite dimensional nature of the resulting parameter space is inherently associated with the ill-posedness, which means that even low noise in the data may make the estimates extremely unstable. This generated interest in a Bayesian framework which quite naturally makes it possible to take account of prior statistical information of the unknown parameters and/or states. The relevant works, see, e.g., [2]–[4] are extremely auspicious, although at present the main impediment to this approach is the inordinately large scale of attendant computations.

The talk is aimed at characterizing the difficulties and challenges of this area of experimental design. Special attention will be paid to the definition of the covariance operator for the Gaussian prior as the inverse of an elliptic differential operator, the appropriate problem discretization and dimensionality reduction, evaluation of A- and D-optimality criteria for large matrices, and sparsity enforcing techniques to get a ‘continuous’ approximation to the binary optimization problem which the sensor selection from a predefined finite set of candidate locations actually is. The presentation will be complemented by illustrative numerical examples.

References:

- [1] D. Uciński, “Optimal Measurement Methods for Distributed Parameter System Identification”, CRC Press, Boca Raton, FL, 2005.
- [2] A. Alexanderian, N. Petra, G. Stadler and O. Ghattas, “A-optimal design of experiments for infinite-dimensional Bayesian linear inverse problems with regularized ℓ_0 -sparsification,” *SIAM Journal on Scientific Computing*, **36**(5):A2122–A2148, 2014.
- [3] E. Haber, L. Horesh and L. Tenorio, “Numerical methods for the design of large-scale nonlinear discrete ill-posed inverse problems,” *Inverse Problems*, **26**(2):025002, 2010.
- [4] I.Y. Gejadze and V. Shutyaev, “On computation of the design function gradient for the sensor-location problem in variational data assimilation,” *SIAM Journal on Scientific Computing*, **34**(2):B127–B147, 2012.
- [5] M. Dashti and A.M. Stuart, “The Bayesian Approach to Inverse Problems,” In: R. Ghanem, D. Higdon and H. Owhadi (Eds.), “Handbook of Uncertainty Quantification,” Springer, 311–428, 2017.

[Dariusz Uciński; Institute of Control and Computation Engineering, University of Zielona Góra, ul. Prof. Z. Szafrana 2, 65–516 Zielona Góra, Poland]
[D.Ucinski@issi.uz.zgora.pl — www.staff.uz.zgora.pl/ducinski/]

Statistical inference based on optimal subdata

HAIYING WANG

University of Connecticut, Storrs, CT, USA

The Optimal Subsampling Method under the A-optimality Criterion (OSMAC) proposed in [1] samples more informative data points with higher probabilities. However, the original OSMAC estimator use inverse of optimal subsampling probabilities as weights in the likelihood function. This reduces contributions of more informative data points and the resultant estimator may lose efficiency. In this paper [2], we propose a more efficient estimator based on OSMAC subsample without weighting the likelihood function. Both asymptotic results and numerical results show that the new estimator is more efficient. In addition, our focus in this paper is inference for the true parameter, while [1] focuses on approximating the full data estimator. We also develop a new algorithm based on Poisson sampling, which does not require to approximate the optimal subsampling probabilities all at once. This is computationally advantageous when available random-access memory is not enough to hold the full data. Interestingly, asymptotic distributions also show that Poisson sampling produces more efficient estimator if the sampling rate, the ratio of the subsample size to the full data sample size, does not converge to zero. We also obtain the unconditional asymptotic distribution for the estimator based on Poisson sampling.

References:

- [1] Haiying Wang, Rong Zhu, and Ping Ma. Optimal subsampling for large sample logistic regression. *Journal of the American Statistical Association*, page <http://dx.doi.org/10.1080/01621459.2017.1292914>, 2017.
- [2] Haiying Wang. More Efficient Estimation for Logistic Regression with Optimal Subsample. <https://arxiv.org/abs/1802.02698>, 2018.

[HaiYing Wang; University of Connecticut, 215 Glenbrook Rd. U-4120, Storrs, CT 06269, USA]
[haiying.wang@uconn.edu]

Computer experiments with big n : has Gaussian process computation been tamed?

SONJA SURJANOVIC

WILLIAM J. WELCH

University of British Columbia, Vancouver, Canada

Computer models are routinely used in a wide variety of applications to replace or augment physical studies. Analogous to statistical design of physical experiments, a computer experiment is a planned set of runs of the computer code, varying the input variables. Because of its computational cost, the computer model is often replaced by a fast statistical surrogate to model the relationship between the inputs and the output(s). The surrogate is used for “what if” prediction, optimization, calibration, and other user objectives.

For decades Gaussian processes (GPs) have been the almost ubiquitous choice for the statistical surrogate (Sacks et al., 1989; Currin et al., 1991; O’Hagan, 1992). The GP itself is computationally expensive for large designs, however. Standard algorithms to train a GP using data from n computer model runs have $O(n^3)$ running time in every likelihood, and hence posterior, evaluation.

Several approaches have tackled the unfavourable computational complexity of GPs. To make analysis more efficient for arbitrary designs, Kaufman et al. (2011) employed sparse matrix techniques, and Gramacy and Apley (2015) introduced localized GPs. Alternatively, special designs can allow faster analysis: Gramacy and Lee (2008, 2009) combined sequential design and treed GPs, and Plumlee (2014) exploited the special structure of sparse-grid designs.

After presenting an overview of these methods, their performances will be compared in a variety of settings.

References:

- [1] Currin, C., Mitchell, T., Morris, M., and Ylvisaker, D. (1991), “Bayesian Prediction of Deterministic Functions, With Applications to the Design and Analysis of Computer Experiments”, *Journal of the American Statistical Association*, **86**, 953–963.
- [2] Gramacy, R. B. and Apley, D. W. (2015), “Local Gaussian Process Approximation for Large Computer Experiments”, *Journal of Computational and Graphical Statistics*, **24**, 561–578.
- [3] Gramacy, R. B. and Lee, H. K. H. (2008), “Bayesian Treed Gaussian Process Models With an Application to Computer Modeling”, *Journal of the American Statistical Association*, **103**, 1119–1130.
- [4] Gramacy, R. B. and Lee, H. K. H. (2009), “Adaptive Design and Analysis of Supercomputer Experiments”, *Technometrics*, **51**, 130–145.
- [5] Kaufman, C. G., Bingham, D., Habib, S., Heitmann, K., and Frieman, J. A. (2011), “Efficient Emulators of Computer Experiments Using Compactly Supported Correlation Functions, With an Application to Cosmology”, *Annals of Applied Statistics*, **5**, 2470–2492.
- [6] O’Hagan, A. (1992), “Some Bayesian Numerical Analysis”, in *Bayesian Statistics 4*, eds. Bernardo, J. M., Berger, J. O., Dawid, A. P., and Smith, A. F. M., Oxford University Press, pp. 345–363.
- [7] Plumlee, M. (2014), “Fast Prediction of Deterministic Functions Using Sparse Grid Experimental Designs”, *Journal of the American Statistical Association*, **109**, 1581–1591.
- [8] Sacks, J., Welch, W. J., Mitchell, T. J., and Wynn, H. P. (1989), “Design and Analysis of Computer Experiments”, *Statistical Science*, **4**, 409–423.

[William J. Welch; Department of Statistics, University of British Columbia, Vancouver, BC V6T 1Z4, Canada]

[will@stat.ubc.ca — www.stat.ubc.ca/~will]

Optimal experimental designs for complex or high dimensional statistical models

WENG KEE WONG

Department of Biostatistics, Fielding School of Public Health, UCLA, USA

Algorithms are practical ways to find optimal experimental designs. Most published work in the statistical literature concerns optimal design problems for models with a few factors or assume models are additive when there are multiple factors. With big data, there are increasingly high-dimensional design problems with many factors and many of the current algorithms may not work well. Nature-inspired meta-heuristic algorithms are general and powerful optimization tools that seem to have been under-utilized in statistical research. I describe their numerous advantages over current algorithms for finding efficient designs, and demonstrate how they can quickly find single or multiple-objective optimal designs for dose response studies and for tackling complex or high-dimensional optimal design problems in biostatistics.

[Weng Kee Wong; Department of Biostatistics, Fielding School of Public Health, University of California at Los Angeles, 405 Hilgard Avenue, Los Angeles, California 90095, USA]

[wk Wong@ucla.edu — <https://ph.ucla.edu/faculty/wong>]

Hilbert series and polynomial models for Smolyak-type sparse grid designs

HENRY P. WYNN

London School of Economics, UK

HUGO MARURI-AGUILLAR

Queen Mary University of London, UK

Sparse grid are used in the area of stochastic finite element solutions to differential equations to capture the input stochastic component, particularly with reference to polynomial quadrature and interpolation. In the case that the grids are nested they have a very special algebraic form as a union of tensor grids, or full factorial designs in the terminology of classical experimental design. This structure applies both to the grids themselves and to the polynomial basis of the underlying model via the minimal free resolution (MFR) of the associated ideal. The decomposition of the model itself is given by the associated Stanley representation. It also transpires that the inclusion-exclusion like formulae used for the underlying interpolation in the sparse grid theory follows from this decomposition and also that the complicated coefficients which multiply the separate tensor terms in the interpolation and quadrature formulae are derivable from the Betti coefficients of the MFR. These are held by the Hilbert series and indeed the whole interpolator has a Hilbert series representation.

A wide class of sparse grids have the nesting condition and hence give the above decomposition. Moreover, the aberration and Alexander duality theory in papers [1] and [2] go some way to explaining the interesting complementary and "hyperbolic" nature of some grids .

The methods can be thought of as a contribution to the modern theory of Uncertainty Quantification (UQ) and follow partly from paper [3]. But there is much work to be done in linking the optimality of sparse grids, for example the optimal spacing arising from the orthogonal polynomial theory, to the rather different criteria in optimal experimental design.

References:

- [1] Y. Bernstein, H. Maruri-Aguilar, S. Onn, E. Riccomagno, and H.P. Wynn, "Minimal average degree aberration and the state polytope for experimental designs." *Annals of the Institute of Statistical Mathematics*, **62** (4), 673-698, 2010.
- [2] H Maruri-Aguilar, E. Senz-de-Cabezn, E. and H. P. Wynn, H.P. "Alexander duality in experimental designs". *Annals of the Institute of Statistical Mathematics*, **65**(4), 667-686. 2013.
- [3] J. Ko, H.P. Wynn, H.P. "The algebraic method in quadrature for uncertainty quantification". *SIAM/ASA Journal on Uncertainty Quantification*, **4**(1), 331-357. 2016.

[Henry P. Wynn; Department of Statistics, London School of Economics, London WC2 2AE, UK]
[h.wynn@lse.ac.uk — <http://www.lse.ac.uk/Statistics/>]

Design admissibility, invariance and optimality in multiresponse linear models

XIN LIU

Donghua University, Shanghai, China

RONG-XIAN YUE

Shanghai Normal University, Shanghai, China

In many experimental situations, especially in engineering, pharmaceutical and biomedical as well as environmental research, there are more than one response to be measured for each unit. The multiresponse models play an important role in many areas of science. This work deals with optimal design problems for the multiresponse linear models. We focus on investigating the optimality, admissibility and invariance of approximate designs. Necessary and sufficient conditions are given for a design to be admissible and invariant. An Elfving's theorem for D-optimality is established for the multiresponse linear models. Several examples are given for illustration.

References:

- [1] H. Dette and V.B. Melas, “A note on the de la Garza phenomenon for locally optimal designs”, *Ann. Statist.*, **39**, 1266–1281, 2011.
- [2] H. Dette and K. Schorning, “Complete classes of designs for nonlinear regression models and principal representations of moment spaces”, *Ann. Statist.*, **41**, 1260–1267, 2013.
- [3] T. Holland-Letz, H. Dette and A. Pepelyshev, “A geometric characterization of optimal designs for regression models with correlated observations”, *J. Roy. Statist. Soc.*, **B73**, 239–252, 2011.
- [4] A.I. Khuri, “Multiresponse Surface Methodology”, In S. Ghosh and C. R. Rao, eds., *Handbook of Statistics*, Vol. 13, 1996.
- [5] O. Krafft and M. Schaefer, “D-optimal designs for a multivariate regression model”, *Multivar. Anal.*, **42**, 130–140, 1992.
- [6] X. Liu, R.-X. Yue and F.J. Hickernell, “Optimality criteria for multiresponse linear models based on predictive ellipsoids”, *Statist. Sinica*, **21**, 421–432, 2011.
- [7] X. Liu, R.-X. Yue and D.K.J. Lin, “Optimal design for prediction in multiresponse linear models based on rectangular confidence region”, *J. Statist. Plan. Inference*, **143**, 1954–1967, 2013.
- [8] F. Pukelsheim, *Optimal Design of Experiments*, Wiley, New York, 1993.
- [9] M. Yang, “On the de la Garza phenomenon”, *Ann. Statist.*, **38**, 2499–2524, 2010.
- [10] M. Yang and J. Stufken, “Identifying locally optimal designs for nonlinear models: A simple extension with profound consequences”, *Ann. Statist.*, **40**, 1665–1685, 2012.

[Rong-Xian Yue; Department of Mathematics, Shanghai Normal University, Shanghai 200234, China]

[yue2@shnu.edu.cn]

Optimal design of sampling survey for efficient parameter estimation

WEI ZHENG

University of Tennessee, Knoxville, USA

For many tasks of data analysis, we may only have the information of the explanatory variable and the evaluation of the response values are quite expensive. While it is impractical or too costly to obtain the responses of all units, a natural remedy is to judiciously select a good sample of units, for which the responses are to be evaluated. In this talk, I will introduce an algorithm with the following features *(i)* The statistical efficiency of any candidate sample can be evaluated without knowing the exact optimal sample; *(ii)* It can be applied to a very wide class of statistical models; *(iii)* It can be integrated with a broad class of information criteria; *(iv)* It is much faster than existing algorithms. *(v)* A geometric interpretation is adopted to theoretically justify the relaxation of the original combinatorial problem to continuous optimization problem.

In the end, I will propose some variates of the algorithm to solve some traditional design problems.

[Wei Zheng; University of Tennessee, Knoxville, USA]

[wzheng9@utk.edu]

Energy functionals, minimizing measures and kernel herding

ANATOLY ZHIGLJAVSKY

School of Mathematics, Cardiff University, Cardiff, UK

LUC PRONZATO

Laboratoire I3S, Université Côte d'Azur, CNRS, Sophia Antipolis, France

This is a survey of some recent results dealing with energy functionals which are

$$\Phi(\mu) = \int_{\mathcal{X}} \int_{\mathcal{X}} K(x, y) \mu(dx) \mu(dy),$$

where μ is a signed measure on $\mathcal{X}$. The main mathematical concepts involved are RKHS and integrally strictly conditionally positive functions and kernels. We make an emphasis on a relation between the problems of construction of the signed measure minimizing the energy and the construction of the BLUE for the location scale model. This has several interesting consequences, especially for smooth kernels.

In the second part of the talk we consider some known algorithms of kernel herding and observe that these algorithms are particular instances of common algorithms of construction of optimal designs.

At the end of the talk we discuss extensions of some classical results on positive definite kernels to the case of singular kernels whose values at the diagonal are not defined; here some results obtained by Tomos Phillips (PhD student at Cardiff University) are essential.

Part of the discussion is based on a recent work of H.Dette, A.Pepelyshev and the first coauthor, see [1]. Other key papers surveyed in the talk are [2-4].

References:

- [1] H. Dette, A. Pepelyshev and A. Zhigljavsky, “The BLUE in continuous-time regression models with correlated errors”, *Annals of Statistics* (submitted), arXiv preprint arXiv:1611.09804.
- [2] B. Sriperumbudur, A. Gretton, K. Fukumizu, B. Schölkopf, and G. Lanckriet, “Hilbert space embeddings and metrics on probability measures”, *Journal of Machine Learning Research*, **11**:1517–1561, 2010.
- [3] S. Damelin, F. Hickernell, D. Ragozin, and X. Zeng, “On energy, discrepancy and group invariant measures on measurable subsets of Euclidean space”, *J. Fourier Anal. Appl.*, **16**:813–839, 2010.
- [4] T.R.L. Phillips, K.M. Schmidt and A. Zhigljavsky, “Extension of the Schoenberg theorem to integrally conditionally positive definite functions”, submitted to *Mathematika*.

[Anatoly Zhigljavsky; School of Mathematics, Cardiff University, CF24 4AG, UK]
[Zhigljavskyaa@cardiff.ac.uk — <http://ssa.cf.ac.uk/zhigljavsky/>]



Posters

- Bertrand **Iooss**, Michaël Baudin, Anne-Laure Popelin, Anne Dutfoy, *OpenTURNS: an open source uncertainty engineering software*
- Asya **Metelkina**, Luc Pronzato, *Information-regret compromise in covariate-adaptive treatment allocation*
- Jérémy Seurat, Thu Thuy **Nguyen**, France Mentré, *Robust designs accounting for model uncertainty in longitudinal studies with binary outcomes*
- Nedka D. **Nikiforova**, Rossella Berni, Jesús Fernando López-Fidalgo, *Optimal heterogeneous choice designs for correlated choice preferences*
- Nicolas **Parisey**, Melen Leclerc, Solenn Stoeckel, Katarzyna Adamczyk, Jacques Baudry, *Some considerations on optimal experimental design for landscape ecology*
- Tomos R.L. **Phillips**, *Extension of the Schoenberg theorem to integrally conditionally positive definite functions*
- Maryna **Prus**, *Various optimality criteria for the prediction of individual response curves*
- Thomas Kahle, Frank **Röttger**, Rainer Schwabe, *Geometry of parameter regions for optimal designs in the Bradley-Terry-model*
- Yevgen **Ryezniuk**, Oleksandr Sverdlov, Andrew C. Hooker, *Implementing optimal designs for dose-response studies through adaptive randomization for a small population group*
- Chenlu **Shi**, Boxin Tang, *Selection of strong orthogonal arrays of strength two*

OpenTURNS: an open source uncertainty engineering software

BERTRAND IOOSS, MICHAËL BAUDIN, ANNE-LAURE POPELIN
EDF R&D, Chatou, France

ANNE DUTFOY
EDF R&D, Saclay, France

The needs to assess robust performances for complex systems and to answer tighter regulatory processes (security, safety, environmental control, health impacts, etc.) have led to the emergence of a new industrial simulation challenge: to take uncertainties into account when dealing with complex numerical simulation frameworks [1]. Many attempts at treating uncertainty in industrial applications have involved different mathematical approaches and standards (sometimes domain-specific): metrology, structural reliability, variational approaches, design of experiments, global sensitivity analysis, machine learning approaches, etc. However, facing the questioning of their certification authorities in an increasing number of different domains, a generic methodology has emerged from the joint effort of several industrial companies and academic institutions [2]. The specific organizational challenges attached are transparency (with respect to safety authorities), genericity (multi-applicative domain issues), modularity (easy integration from the open-source community), multi-accessibility (different levels of use and users) and industrial computing capabilities.

As no software was fully answering the mathematical techniques and organizational challenges mentioned above, EDF R&D, Airbus Group, and Phimeca Engineering started a collaboration at the beginning of 2005, joined by IMACS in 2014, for the development of an open-source software platform dedicated to uncertainty propagation by probabilistic methods, named OpenTURNS for open-source treatment of uncertainty, Risk N Statistics [3]. Based on a probabilistic representation of the model input variables, this tool aims to integrate all the useful algorithms for the probabilistic modeling (including the dependence issues), the statistical sampling tools (including various numerical designs of experiments), the uncertainty quantification (UQ) and sensitivity analysis steps, the metamodeling and parameter calibration techniques.

OpenTURNS is supported by its core team and its user community via the website www.openturns.org. At EDF, OpenTURNS is the reference software on UQ issues, for methodological dissemination in the business units. Numerous efforts have been made for its integration into the various computing environments. It is thus integrated within the Salome platform [4], allowing to couple many field-physics models. Since 2017, OpenTURNS has facilitated the coupling of any system model, thanks to the use of the Functional Mock-Up Interface (FMI) standard, an API widely used by most modeling languages, such as Modelica.

References:

- [1] R. Smith, “Uncertainty quantification”, SIAM, 2014.
- [2] A. Pisanisi and A. Dutfoy, “An industrial viewpoint on uncertainty quantification in simulation: stakes, methods, tools, examples”, In: *A. Dienstfrey and R. Boisvert (eds.), Uncertainty Quantification in Scientific Computing*, IFIP Advances in Information and Communication Technology, vol. 377, pp. 2745. Springer, Berlin, 2012.
- [3] M. Baudin, A. Dutfoy, B. Iooss and A-L. Popelin, “Open TURNS: An industrial software for uncertainty quantification in simulation”, In: *R. Ghanem, D. Higdon and H. Owhadi (eds), Handbook of uncertainty quantification*, Springer, 2017.
- [4] OPEN CASCADE SAS, “Salome: the open source integration platform for numerical simulation”, <http://www.salome-platform.org>, 2006.

Information-regret compromise in covariate-adaptive treatment allocation

ASYA METELKINA

LUC PRONZATO

Laboratoire I3S, Université Côte d’Azur, CNRS, Sophia Antipolis, France

We consider the construction of adaptive designs for ethical treatment allocation in (phase III) clinical trials. The usual objective of clinical trials concerns statistical inference about the effects of treatments as functions of covariates, through the estimation of each treatment model. In practice, this objective must be balanced with individual ethics that requires to favor allocation of the best treatment for each patient enrolled in the trial. To explicitly account for these conflicting objectives, we introduce a compromise criterion which is a convex combination of (i) an information criterion (concave function of Fisher information matrix) reflecting the precision of statistical inference and of (ii) an additive cost associated to the allocation of poorest treatment. The presentation is based on the paper [1].

The concept of design measures is natural for investigating asymptotic properties of designs, when the allocation strategy can be constructed to target the desired limiting measure. Within the framework of approximate design theory, the determination of an allocation measure that maximizes the compromise criterion forms a compound design problem. We show that when covariates are i.i.d. with a probability measure μ , its solution possesses some similarities with the construction of optimal design measures bounded by μ and is characterized through an equivalence theorem. The form of our compromise criterion insures that the optimal measure maximizes the information criterion under a bound constraint on ethical cost.

To target the optimal measure, we construct adaptive sequential designs, where the patients enter the study sequentially and current allocation depends on covariates of the current individual and possibly on covariates, allocations and responses of previous individuals. We introduce an oracle covariate-adaptive sequential allocation strategy that converges to the optimal measure and derive its asymptotic properties. In general, this sequential allocation is deterministic. When randomization is needed, an alternative optimal allocation strategy is proposed, where for each subject we use random balanced allocation with probability β and the predictable rule with probability $(1 - \beta)$.

The (randomized) oracle strategy requires the knowledge of the distribution of covariates μ ; moreover, its construction is complicated for trials involving non-scalar covariates. These difficulties can be avoided by using a covariate-adaptive allocation rule based on empirical allocation measures, which can be shown to converge to the optimal allocation. The target measure is locally optimal since its construction depends on the parameters of treatment models. A covariate-adjusted response-adaptive version of this allocation rule is proposed, which uses the current ML parameter estimates and targets the optimal allocation measure for the true unknown model parameters. The obtained optimal designs can be used as benchmarks for other, more usual, allocation methods. Several illustrative examples are provided and a comparison is made with recent results of the literature on CARA design.

References:

[1] A. Metelkina and , L. Pronzato, “Information-regret compromise in covariate-adaptive treatment allocation”, *The Annals of Statistics*, **45**(5):2046–2073, 2017.

[Metelkina Asya; Laboratoire I3S, UNS-CNRS, bât. Euclide Les Algorithmes, BP121, 2000 route des Lucioles, F06903 Sophia Antipolis Cedex, France]

[metelkin@i3s.unice.fr — <http://www.i3s.unice.fr/~metelkin/>]

Robust designs accounting for model uncertainty in longitudinal studies with binary outcomes

JÉRÉMY SEURAT
THU THUY NGUYEN

FRANCE MENTRÉ
IAME, INSERM, UMR 1137, University Paris Diderot, Paris, France

Objectives: Nonlinear mixed effect models are widely used for the analysis of longitudinal data obtained during clinical trials. To design these studies, a method evaluating the expected Fisher Information Matrix (FIM), without any linearization, was proposed based on Monte-Carlo and Hamiltonian Monte-Carlo (MC/HMC) and implemented in the R package MIXFIM [1]. This approach however requires a priori knowledge of the model, which may lead to non-informative designs if the guessed model is inaccurate. We aimed to propose a robust design approach to account for model uncertainty and to ensure a compromise between the overall precision of estimation and the power of the Wald test to detect a covariate effect. We illustrated and evaluated by simulations the proposed approach through an example of designing a longitudinal trial with binary outcomes.

Methods: First, to optimize designs given a predefined model, different optimality criteria based on the FIM evaluated by MC/HMC were computed: the D-optimality to account for the whole set of parameters, the D_S -optimality for a subset of parameters of interest, and the DD_S -optimality for a compromise between the D- and D_S -optimality. Then, to account for model uncertainty, we assumed a set of candidate models with their respective weights and we computed robust designs across these models using compound CD -, CD_S - and CDD_S -optimality [2]. These methods were applied to design a study with two treatment groups, using a logistic model for repeated binary responses. Four candidate models describing the evolution of the logit-probability of the response over time (0 to 12 months) were defined: M1 linear, M2 log-linear, M3 quadratic and M4 exponential models. We performed combinatorial optimization to obtain sparse sampling times which were optimal for each model separately or optimal over the four models. Using the FIM for a given design, we also predicted the average power of the Wald test to detect a significant treatment effect over the four models. Clinical trial simulations were then used to evaluate the performances of the CDD_S -optimal design ξ_{CDD_S} vs. the DD_S -optimal design for a given model ($\xi_{DD_{S1}}$, for M1) vs. the equi-spaced design ξ_{ES} . For that we simulated 500 datasets under each model and analyzed them using SAEM algorithm in the software MONOLIX2016R1.

Results: Misspecification of models led to designs with D-efficiencies as low as 64.6%. The compound criteria provided efficient robust designs across the four models, with D-efficiencies always above 80%. With the designs ξ_{ES} , $\xi_{DD_{S1}}$, and ξ_{CDD_S} , we predicted respectively 358, 320 and 274 subjects needed to achieve an average power of 0.9 over the four models. The simulation study confirmed that, for the same number of subjects, the robust design ξ_{CDD_S} is more informative than $\xi_{DD_{S1}}$ and ξ_{ES} , giving acceptable estimation errors and good power for the four models.

Conclusion: The proposed design strategy based on MC/HMC and compound optimality theory, is a relevant approach which can be used to efficiently design longitudinal studies. This approach accounts for model uncertainty and ensures a balance between the overall precision of estimation and the power of the Wald test to detect a covariate effect.

References:

- [1] M-K. Riviere, S. Ueckert and F. Mentré, “An MCMC method for the evaluation of the Fisher information matrix for non-linear mixed effect models.”, *Biostat. Oxf. Engl.*, **17**:737–50, 2016.
- [2] A.C. Atkinson, A.N Donev and R.D. Tobias, “Optimum Experimental Designs, with SAS.”, *Oxford, New York: Oxford University Press*, 2009.

[Jérémy Seurat and Thu Thuy Nguyen; IAME, INSERM, UMR 1137, University Paris Diderot]
[jeremy.seurat@inserm.fr – thu-thuy.nguyen@inserm.fr]

Optimal heterogeneous choice designs for correlated choice preferences

NEDKA D. NIKIFOROVA

ROSSELLA BERNI

Department of Statistics, Computer Science, Applications "G. Parenti", University of Florence, Florence, Italy

JESÚS FERNANDO LÓPEZ-FIDALGO

Universidad de Navarra, ICS, Unidad de Estadística, Pamplona, Spain

In this work we propose an innovative approach for the construction of heterogeneous choice designs for correlated choice preferences. Differently from existing researches in the choice experiment literature that usually employs the exact design framework, we build optimal heterogeneous choice designs based on approximate design theory under the Panel Mixed Logit model that explicitly takes into account the fact that the responses given by the same respondent are correlated. The approach we have developed allows us for obtaining optimal heterogeneous choice designs composed by groups of choice-sets to be administered to a proportion of subject according to the optimal weights.

We demonstrate the efficiency of our proposal through an application to a real case study that concerns the analysis of the consumers' preferences for sustainable coffees integrating a choice experiment with consumers sensory tests. To this end, we develop our proposal under a compound design criterion (Wynn, 1970; Atkinson and Bogacka, 1997; Atkinson et al., 2007) in order to address the following two main issues: i) an efficient estimation of the attributes of the choice experiment, and ii) detection of the effect of the sensory assesment scores obtained through a guided tasting session. We present the estimation results related to the proposed optimal heterogeneous choice design that are very satisfactory by also confirming the validity of our innovative approach.

References:

- [1] H.P. Wynn, "The sequential generation of D-optimal experimental designs", *The Annals of Mathematical Statistics*, **41**:1055–1064, 1970.
- [2] A.C. Atkinson and B. Bogacka, "Compound D - and D_s - Optimum Designs for Determining the Order of a Chemical Reaction", *Technometrics*, **39**(4):347–356, 1997.
- [3] A.C. Atkinson, A.N. Donev and R.D. Tobias, "Optimum Experimental Designs, with SAS", *Oxford U.K.: Oxford University Press*, 2007.

[Nedka D. Nikiforova; Department of Statistics, Computer Science, Applications "G. Parenti", University of Florence, Florence, Italy]
[nikiforova@disia.unifi.it]

Some considerations on optimal experimental design for landscape ecology

NICOLAS PARISEY

MELEN LECLERC

SOLENN STOECKEL

UMR INRA 1349 IGEPP, 35650 Le Rheu, France

KATARZYNA ADAMCZYK

Unité MaIAGE, INRA, Domaine de Vilvert, 78352 Jouy-en-Josas, France

JACQUES BAUDRY

INRA, UMR BAGAP, CS 84215, 35000 Rennes, France

Population dynamics in heterogeneous landscapes can be investigated using parsimonious reaction-diffusion models whose parameters can be inferred from geolocalised population measurements and remote sensing data [1]. Such methodology is now well established for the study of beneficial insects and pest populations in agricultural landscapes [2]. But the cost of landscape-scale experiments remains the main bottleneck against a rich interaction between empirical and modelling approaches.

In this work, we are interested in reaction-diffusion models with a known (or independently estimated) initial condition equation, a non-linear state equation, an observation equation related to the cumulated population density at a given position over a target period and a data model for count or presence/absence.

We advocate that practitioners would greatly benefit from using optimal experimental design [3,4]. We illustrate this by constructing an on-average D_s -optimal exact design for a previously published model [2]. Our desired perspectives include relying on sequential design for parameter estimation and model discrimination [5] and, from a more applied point of view, to use optimal design for landscape genetics as well [6].

References:

- [1] S. Soubeyrand, L. Roques, “Parameter estimation for reaction-diffusion models of biological invasions”, *Popul. Ecol.*, **56**, 427–434, 2014.
- [2] N. Parisey, Y. Bourhis, L. Roques, S. Soubeyrand, B. Ricci and S. Poggi, “Rearranging agricultural landscapes towards habitat quality optimisation: In silico application to pest regulation”, *Ecol. Com.*, **28**, 113–122, 2016.
- [3] A.C. Atkinson, V.V. Fedorov, A.M. Herzberg, R. Zhang, “Elemental information matrices and optimal experimental design for generalized regression models”, *J. Stat. Plan. Infer.*, **144**, 81–91, 2014.
- [4] D. Ucinski, “Optimal Measurement Methods for Distributed Parameter System Identification”, CRC Press, 2004.
- [5] C. Tommasi, “Optimal designs for both model discrimination and parameter estimation”, *J. Stat. Plan. Infer.*, **139**, 4123–4132, 2009.
- [6] M. Hainy, W.G. Mller, H.P. Wynn, “Approximate Bayesian Computation Design (ABCD), an Introduction”, MODa 10 Advances in Model-Oriented Design and Analysis, Springer International Publishing, Heidelberg, 135143, 2013.

[N. Parisey; INRA, UMR 1349 IGEPP, 35650 Le Rheu, France]

[nicolas.parisey@inra.fr — www6.rennes.inra.fr/igepp_eng/Research-teams/Demecology]

Extension of the Schoenberg theorem to integrally conditionally positive definite functions

TOMOS R.L. PHILLIPS

School of Mathematics, Cardiff University, Cardiff, UK

The celebrated Schoenberg theorem establishes a relation between positive definite and conditionally positive definite functions. In this paper, we consider the classes of real-valued functions $P(J)$ and $CP(J)$, which are positive definite and respectively, conditionally positive definite, with respect to a given class of test functions J . For suitably chosen J , the classes $P(J)$ and $CP(J)$ contain classically positive definite (respectively, conditionally positive definite) functions, as well as functions which are singular at the origin. The main result of the paper is a generalization of Schoenberg's theorem to such function classes. We provide many examples of integrally positive definite functions and integrally conditionally positive definite functions with singularity at the origin.

The main application of the integrally conditionally positive definite functions is associated with the functional (energy)

$$\Phi_f(\mu) = \int \int f(x - y) d\mu(x) d\mu(y)$$

on the space of finite signed measures μ : if f is integrally conditionally positive definite function then the functional $\Phi_f(\cdot)$ is convex on the space of signed measures μ such that their energy is finite.

We also describe an algorithm of numerical construction of signed measure minimizing the energy functional $\Phi_f(\cdot)$ in the case of functions f with singularity.

[Tomos R.L. Phillips; School of Mathematics, Cardiff University, CF24 4AG, UK]

[PhillipsT5@cardiff.ac.uk]

Various optimality criteria for the prediction of individual response curves

MARYNA PRUS

Otto von Guericke University, Magdeburg, Germany

The subject of this work is random coefficients regression models, where observational units (individuals) are assumed to come from same population and differ from each other by individual random parameters. Analytical results for optimal designs for commonly used design criteria as linear and determinant criteria for the prediction in these models are presented in [1] and [2] (see also [3]). A practical approach for computation of optimal approximate and exact designs for linear criteria was discussed in [4]. Here we consider optimal designs for some particular Kiefer-criteria as well as G-criterion for the prediction of individual response curves.

References:

- [1] Prus, M.; Schwabe, R., “Interpolation and extrapolation in random coefficient regression models: Optimal design for prediction”, In C. H. Müller, J. Kunert, A. C. Atkinson (eds.): “mODa 11 - Advances in Model-Oriented Design and Analysis”, Springer, 209–216, 2016.
- [2] Prus, M., Schwabe, R., “Optimal designs for the prediction of individual parameters in hierarchical models”, *Journal of the Royal Statistical Society B*, **78**, 175–191, 2016.
- [3] Prus, M., “Optimal Designs for the Prediction in Hierarchical Random Coefficient Regression Models”, Ph.D. thesis, Otto-von-Guericke University, Magdeburg, 2015.
- [4] Harman, R. and Prus, M., “Computing optimal experimental designs with respect to a compound Bayes risk criterion”, *Statistics and Probability Letters*, **137**, 135-141, 2018.

[Maryna Prus; Otto von Guericke University, FMA, IMST, PF 4120, 39016 Magdeburg, Germany]
[maryna.prus@ovgu.de — www.math.ovgu.de]

Geometry of parameter regions for optimal designs in the Bradley-Terry-model

THOMAS KAHLE

Otto-von-Guericke-Universität Magdeburg, Germany

FRANK RÖTTGER

Otto-von-Guericke-Universität Magdeburg, Germany

RAINER SCHWABE

Otto-von-Guericke-Universität Magdeburg, Germany

Optimal design theory for nonlinear regression studies local optimality on a given design space. We identify the Bradley-Terry paired comparison model with graph representations and prove for an arbitrary number of parameters, that every saturated D-optimal design is displayed as a path in its graph representation. Via this path property we give a complete description of the optimality regions of saturated designs. Furthermore, we exemplify the unsaturated D-optimal designs with full support for 4 parameters by finding representations of the semi-algebraic sets given by the Kiefer-Wolfowitz equivalence theorem. This leads to formulas for the weights of the optimal designs being rational functions in the intensities, hence the parameters. This extends the results of Graßhoff and Schwabe in [1] for the one-way layout to the Bradley-Terry paired comparison model with 4 parameters.

References:

[1] Graßhoff U., Schwabe R. "Optimal design for the Bradley–Terry paired comparison model" *Statistical Methods and Applications*, v. 17, n. 3, p. 275–289, 2008.

[Frank Röttger; Otto-von-Guericke-Universität Magdeburg, Germany]

[frank.roettger@ovgu.de — www.imst3.ovgu.de]

Implementing optimal designs for dose-response studies through adaptive randomization for a small population group

YEVGEN RYEZNIK

*Department of Mathematics, Department of Pharmaceutical Biosciences,
Uppsala University, Uppsala, Sweden*

OLEKSANDR SVERDLOV

*Early Development Biostatistics – Translational Medicine,
Novartis Pharmaceutical Corporation, East Hanover, NJ, USA*

ANDREW C. HOOKER

*Department of Pharmaceutical Biosciences,
Uppsala University, Uppsala, Sweden*

A problem of implementing the D-optimal design for dose-response studies with censored time-to-event outcomes is addressed. Particularly, we consider a quadratic dose-response model for log-transformed Weibull event times that are subject to right censoring. D-optimal designs for such a problem depend on the true model parameters and the amount of censoring in the model. In practice, such designs can be implemented adaptively, when dose assignments are made according to updated knowledge of the dose-response curve at interim analysis [1]. It is also essential that treatment allocation involves randomization – to mitigate various experimental biases and enable valid statistical inference at the end of the trial. In this work, we perform a comparison of several randomization procedures that can be used for implementing D-optimal designs for dose-response studies with time-to-event outcomes with small to moderate sample sizes. We compare them in terms of balance, randomness, estimation efficiency, and impact on bias and uncertainty of parameter estimates. We consider single stage, two-stage, and multi-stage adaptive designs. Scenarios with chronological and selection bias are explored as well.

References:

[1] Y. Ryznik, O. Sverdlov and A.C. Hooker “Adaptive Optimal Designs for Dose-Finding Studies with Time-to-Event Outcomes”, *AAPS J.*, **20**(2):24, 2017.

[Yevgen Ryznik; Department of Mathematics, Uppsala University, A14133 Lgerhyddsvgen 1, Hus 1, 6 och 7, 75601, Uppsala, Sweden]

[yevgen.ryznik@math.uu.se — <http://katalog.uu.se/empinfo/?id=N14-1199>]

Selection of strong orthogonal arrays of strength two

CHENLU SHI

BOXIN TANG

Simon Fraser University, Burnaby, Canada

Strong orthogonal arrays, a class of space-filling designs, used for computer experiments were introduced and studied by He and Tang (2013). For these arrays, to enjoy better space-filling properties than comparable ordinary orthogonal arrays, they need to be of strength three or higher. But for some studies, such arrays of strength three or higher often require large run sizes, which may be too expensive for experimenters to afford. This draws our attention back to strong orthogonal arrays of strength two. In this paper we investigate those arrays that maximize the number of subarrays that retain the two-dimensional properties of strong orthogonal arrays of strength three.

References:

[1] Y. He and B. Tang, “Strong orthogonal arrays and associated Latin hypercubes for computer experiments”, *Biometrika*, **100**(1):254–260, 2013.

[Chenlu Shi; Simon Fraser University, Burnaby, British Columbia V5A 1S6, Canada]
[chenlus@sfu.ca — <https://www.stat.sfu.ca/>]