



Targeted Covariate-Adjusted Response-Adaptive LASSO-Based Randomized Controlled Trials

Antoine Chambaz, Mark J. van Der Laan, Wenjing Zheng

► To cite this version:

Antoine Chambaz, Mark J. van Der Laan, Wenjing Zheng. Targeted Covariate-Adjusted Response-Adaptive LASSO-Based Randomized Controlled Trials. Oleksandr Sverdlov. Modern Adaptive Randomized Clinical Trials: Statistical, Operational, and Regulatory Aspects, 2015, 978-1-48-223988-1. hal-00990761

HAL Id: hal-00990761

<https://hal.science/hal-00990761>

Submitted on 14 May 2014

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Targeted Covariate-Adjusted Response-Adaptive LASSO-Based Randomized Controlled Trials

A. Chambaz¹, M. J. van der Laan², W. Zheng²

¹Université Paris Ouest Nanterre

²University of California, Berkeley

May 14, 2014

Abstract

We present a new covariate-adjusted response-adaptive randomized controlled trial design and inferential procedure built on top of it. The procedure is targeted in the sense that (i) the sequence of randomization schemes is group-sequentially determined by targeting a user-specified optimal randomization design based on accruing data and, (ii) our estimator of the user-specified parameter of interest, seen as the value of a functional evaluated at the true, unknown distribution of the data, is targeted toward it by following the paradigm of targeted minimum loss estimation. We focus for clarity on the case that the parameter of interest is the marginal effect of a binary treatment and that the targeted optimal design is the Neyman allocation, in an effort to produce an estimator with smaller asymptotic variance. For clarity too, we consider the case that the estimator of the conditional outcome given treatment and baseline covariates, a key element of the procedure, is obtained by LASSO regression. Under mild assumptions, the resulting sequence of randomization schemes converges to a limiting design, and the TMLE estimator is consistent and asymptotically Gaussian. Its asymptotic variance can be estimated too. Thus we can build valid confidence intervals of given asymptotic levels. A simulation study confirms our theoretical results.

Keywords: covariate-adjusted response-adaptive (CARA) design, randomized controlled trial (RCT), least absolute shrinkage and selection operator (LASSO), targeted minimum loss estimation (TMLE)

1 Introduction

1.1 Overview

This technical report is devoted to the study of a so-called group-sequential CARA randomized controlled trial (RCT), with a particular focus on incorporating more flexible (*i.e.*, data-adaptive) techniques to model the response. A CARA RCT is Covariate-Adjusted: the treatment randomization schemes are allowed to be a function of the patients' pre-treatment covariates. In addition, a CARA RCT is Response-Adaptive: the investigators have the opportunity to adjust these schemes during the course of the trial based on accruing information, including previous responses, in order to meet some pre-specified objectives. In a group-sequential CARA RCT, the latter adjustments are made at interim time points given by sequential inclusion of blocks of c patients, where $c \geq 1$ is a pre-specified integer. We consider the case of $c = 1$ for simplicity of exposition, though the discussions generalize to any $c > 1$.

The trial protocol pre-specifies the observed data structure, scientific parameters of interest, analysis methods, and a criterion characterizing an optimal randomization scheme. Here, some baseline covariates and a primary outcome of interest are measured on each patient. We choose the marginal treatment effect of a binary treatment as our parameter of interest, ψ_0 . It is analyzed using targeted minimum

loss estimation (TMLE) on top of the so-called LASSO (least absolute shrinkage and selection operator) methodology (Tibshirani, 1996) that we choose to illustrate the application of data-adaptive techniques to model the response. The TMLE methodology was first introduced by (van der Laan and Rubin, 2006) in the independent identically distributed setting. Its extension to adaptive RCTs was considered in (van der Laan, 2008) and (Chambaz and van der Laan, 2013), upon which this technical report relies. The extension based on LASSO that we present here encompasses the parametric approach of (Chambaz and van der Laan, 2013) as a special case. For concreteness, we choose the so-called Neyman design as our optimal randomization scheme. The Neyman design minimizes the Cramér-Rao lower bound on the asymptotic variances of a large class of estimators of ψ_0 . The resulting Neyman allocation probabilities are evaluated conditionally on the baseline covariates. By targeting the Neyman design, we aim at improving the efficiency of the study, *i.e.*, at reaching a valid result using as few blocks of patients as possible. We emphasize that the results and procedures presented here are generally applicable to other parameters and optimal randomization schemes.

We show that, under mild conditions, the resulting TMLE estimator of ψ_0 is consistent and asymptotically normal regardless of the consistency of the LASSO estimator of the conditional expectation of the response given treatment and baseline covariates. Furthermore, the resulting targeted CARA design converges to a fixed limiting design, which equals the Neyman design if the LASSO estimator is consistent and if the Neyman design belongs to a user-supplied set of randomization schemes. The general framework that combines CARA RCTs with machine-learning techniques is presented in a separate article. Before we delve into the main contents, let us motivate our discussion with a bird’s eye view of the landscape of CARA designs.

1.2 Literature Review

Adaptive randomization has a long history that can be traced back to the 1930s. We refer to (Rosenberger, 1996, Rosenberger, Sverdlov, and Hu, 2012), (Hu and Rosenberger, 2006, Section 1.2) and (Jennison and Turnbull, 2000, Section 17.4) for a comprehensive historical perspective. Many articles are devoted to the study of response-adaptive randomizations, which select current treatment probabilities based on responses of previous patients, but not on the covariates of the current patients. We refer to (Hu and Rosenberger, 2006, Chambaz and van der Laan, 2011, Rosenberger et al., 2012) for a bibliography on that topic. In a heterogeneous population, however, it is often desirable to take into account the patients’ characteristics for treatment assignment. CARA randomization tackles the issue of heterogeneity by dynamically calculating the allocation probabilities based on previous responses and current and past values of certain covariates. Compared to the broader literature on response-adaptive randomization, the advances in CARA procedures are relatively recent, but growing steadily. Among the first approaches, (Rosenberger, Vidyashankar, and Agarwal, 2001, Bandyopadhyay and Biswas, 2001) considered randomization procedures defined as explicit functions of the conditional responses, which are modeled by generalized linear models. Though these procedures are not defined based on formal optimality criteria, their general goal is to allocate more patients to their corresponding “better” treatment arm. Atkinson and Biswas (2005) presented a biased-coin design with skewed allocation, which is determined by sequentially maximizing a function that combines the variance of the parameter estimate, based on a Gaussian linear model for the conditional response, and the conditional treatment effect given covariates. Up till here, very little work had been devoted to asymptotic properties of CARA designs. Subsequently, Zhang, Hu, Cheung, and Chan (2007), Zhang and Hu (2009) established the efficiency theory for CARA designs converging to any given target design, when the responses follow a generalized linear model, and proposed a covariate-adjusted doubly-adaptive biased coin design whose asymptotic variance achieves the efficiency bound. Chang and Park (2013) proposed a sequential estimation of CARA designs under generalized linear models for the response. This procedure allocates treatment based on the patients’ baseline covariates, accruing information and sequential estimates of the treatment effect and uses a stopping rule

that depends on the observed Fisher information. With regard to hypothesis testing, Shao, Yu, and Zhong (2010), Shao and Yu (2013) provided asymptotic results for valid tests under generalized linear models for the responses. Most recently, progress has also been made in CARA designs in the longitudinal settings, see for example (Biswas, Bhattacharya, and Park, 2014, Huang, Liu, and Hu, 2013, Sverdlov, Rosenberger, and Ryznik, 2013).

To tackle the issue of restrictive modeling assumptions, Chambaz and van der Laan (2013) proposed a TMLE analysis of a CARA design where the treatment allocation is conditional on a summary measure of the covariates that takes only finitely many values. Under such a framework, the treatment effect is defined nonparametrically, and the consistency and asymptotic normality of its estimator is robust to misspecification of the parametric working model for the response. However, assigning treatment based on such summary measures is perhaps too restrictive in real-life RCTs where response to treatment may be correlated with a large number of a patient’s baseline characteristics, some of which being continuous. Moreover, although a misspecified parametric working model for the response does not hinder the consistency of the treatment effect estimator, it may affect its efficiency and the convergence of the CARA design to the targeted optimal design.

In this technical report, we generalize the results of Chambaz and van der Laan (2013) to address the two issues mentioned above. We adopt a loss-based approach to the construction of more flexible CARA randomization schemes while exploiting data-adaptive estimators for the estimation of the response model, in search for greater efficiency through better variable adjustments and more accurate estimation of the variance of the estimator.

1.3 Organization

The remainder of this technical report is organized as follows. In Section 2 we introduce our LASSO-based group-sequential CARA RCT design and the TMLE procedure built on top of it to infer the marginal treatment effect of a binary treatment. Section 3 is devoted to the presentation of results pertaining to the convergence of our targeted CARA design and to the asymptotics, consistency and central limit theorem, of the TMLE estimator. A simulation study is described and its results summarized in Section 4. The technical report closes on a discussion in Section 5.

2 Targeted CARA RCT using LASSO

In the introduction, we have outlined the motivation to use data-adaptive procedures to estimate the conditional response given treatment and covariates. For concreteness of the formal theoretical development, we consider here the LASSO estimator, which is a shrinkage and selection method for generalized regression models that optimizes a loss function of the regression coefficients subject to the constraint that the L^1 norm of the coefficient vector be upper-bounded by a given value. The parametric estimators considered in (Chambaz and van der Laan, 2013) are a special case of a LASSO estimator.

We begin by establishing the key features of the trial, namely, the parameter of interest, analysis method, and the optimal randomization scheme. Then, we describe the data generating process (including estimation of the response model using LASSO and adaptation of the randomization scheme) and the targeted maximum likelihood estimation procedure.

2.1 Observed Data Structure, Parameter of Interest and Optimal Design

Prior to data collection, the trial protocol notably specifies the observed data structure, parameter of interest, and the optimal randomization design to target, both expressed in terms of features of the true, unknown data-generating process in the population of interest. In this technical report we consider a

simple situation, judging by the definition of the data and our choice of parameter of interest. The range of application of the methods presented here extends beyond this limited yet instructive framework.

Sections 2.1.1, 2.1.2 and 2.1.3 are respectively devoted to the presentation and discussion of the observed data structure, parameter of interest, and optimal randomization design.

2.1.1 Observed Data Structure

The data structure O writes as $O \equiv (W, A, Y)$, where $W \in \mathcal{W}$ consists of the baseline covariates (some of which may be continuous), $A \in \mathcal{A} \equiv \{0, 1\}$ is the binary treatment of interest, and $Y \in \mathcal{Y}$ is the primary outcome of interest. We assume that the outcome space $\mathcal{O} \equiv \mathcal{W} \times \mathcal{A} \times \mathcal{Y}$ is bounded. Without loss of generality, we may then assume that $Y \in \mathcal{Y} \equiv (0, 1)$ is bounded away from 0 and 1.

Every distribution of O consists of three components. On one hand, the marginal distribution of W and the conditional distribution of Y given (A, W) form a couple which is given by nature. On the other hand, the conditional distribution of A given W , also know as (a.k.a.) randomization scheme, is controlled by the investigators of the RCT. To reflect this dichotomy, we denote $P_{Q,g}$ the distribution of O whose couple formed by the marginal distribution of W and the conditional distribution of Y given (A, W) equals Q and whose randomization scheme equals $g \in \mathcal{G}$, with $\mathcal{G}$ the set of all randomization schemes. For a given Q , we denote Q_W the related marginal distribution of W and Q_Y the related conditional expectation of Y given (A, W) . Moreover, we denote Q_0 the true couple in our population of interest, which is unknown to us, and we assume that this Q_0 does not vary during the whole duration of the RCT. Thus, for any Q and g , $P_{Q_0,g}$ is the true, partially unknown distribution of O when one relies on g , and $E_{P_{Q_0,g}}(Y|A, W) = Q_Y(A, W)$, $P_{Q_0,g}(A = 1|W) = g(1|W) = 1 - g(0|W)$ $P_{Q_0,g}$ -almost surely.

2.1.2 Parameter of Interest

The parameter of interest under consideration in this technical report is the marginal treatment effect on an additive scale:

$$\psi_0 \equiv E_{P_{Q_0,g}} \{Q_{Y,0}(1, W) - Q_{Y,0}(0, W)\} = \int (Q_{Y,0}(1, w) - Q_{Y,0}(0, w)) dQ_{W,0}(w),$$

which evidently depends on $P_{Q_0,g}$ only through Q_0 . Of particular interest in medical, epidemiological and social sciences research, this parameter can be interpreted causally under additional assumptions on the data-generating process (Pearl, 2000). Central to our approach is seeing ψ_0 as the value at any $P_{Q_0,g}$ of the mapping $\Psi : \mathcal{M} \rightarrow [-1, 1]$ characterized over the set $\mathcal{M}$ of all possible distributions of O by

$$\Psi(P_{Q,g}) \equiv E_{P_{Q,g}} \{Q_Y(1, W) - Q_Y(0, W)\} = \int (Q_Y(1, w) - Q_Y(0, w)) dQ_W(w). \quad (1)$$

The mapping Ψ enjoys a remarkable property: it is pathwise differentiable (think “smooth”) with an efficient influence curve (think “gradient”) which provides insight into the asymptotic properties of all regular and asymptotically linear (think “well-behaved”) estimators of $\Psi(P_{Q_0,g})$. The following lemma makes the latter statement more formal—we refer the reader to (Bickel, Klaassen, Ritov, and Wellner, 1998, van der Vaart, 1998, van der Laan and Robins, 2003) for definitions and proofs.

Lemma 1. *The mapping $\Psi : \mathcal{M} \rightarrow [-1, 1]$ is pathwise differentiable at every $P_{Q,g} \in \mathcal{M}$ with respect to (wrt) the maximal tangent space. Its efficient influence curve at $P_{Q,g}$ is $D^*(P_{Q,g})$ which satisfies $D^*(P_{Q,g})(O) = D_W^*(P_{Q,g})(W) + D_Y^*(Q, g)(O)$ with*

$$\begin{aligned} D_W^*(P_{Q,g})(W) &\equiv Q_Y(1, W) - Q_Y(0, W) - \Psi(P_{Q,g}), \\ D_Y^*(Q, g)(O) &\equiv \frac{2A - 1}{g(A|W)} (Y - Q_Y(A, W)). \end{aligned}$$

The variance $\text{Var}_{P_{Q,g}} D^*(P)(O)$ is a generalized Cramér-Rao lower bound for the asymptotic variance of any regular and asymptotically linear estimator of $\Psi(P_{Q,g})$ when sampling independently from $P_{Q,g}$. Moreover, if either $Q_Y = Q'_Y$ or $g = g'$ then $E_{P_{Q,g}} D^*(P_{Q',g'})(O) = 0$ implies $\Psi(P_{Q,g}) = \Psi(P_{Q',g'})$.

The last statement of Lemma 1, often referred to as a “double-robustness” property, shows that one can seek help from D^* to protect oneself against model misspecifications when estimating ψ_0 . This is especially relevant in our setting where we know precisely what is the randomization scheme g at play when one samples an observation from $P_{Q_0,g}$.

2.1.3 Optimal Design

Suppose our goal of adaptation is to reach a randomization scheme of higher efficiency, *i.e.*, to obtain a valid estimate of ψ_0 using as few blocks of patients as possible. By Lemma 1, the asymptotic variance of a regular, asymptotically linear estimator is lower-bounded by $\min_{g \in \mathcal{G}} \text{Var}_{P_{Q_0,g}} D^*(P_{Q_0,g})$. In this light, the *Neyman design* (Hu and Rosenberger, 2006)

$$g_0 \equiv \arg \min_{g \in \mathcal{G}} \text{Var}_{P_{Q_0,g}} D^*(P_{Q_0,g}) = \arg \min_{g \in \mathcal{G}} E_{P_{Q_0,g}} \frac{(Y - Q_{Y,0}(A, W))^2}{g^2(A|W)} \quad (2)$$

can be considered as an optimal randomization design (“optimal design” for short). Since its definition involves the unknown Q_0 , the optimal design g_0 is unknown too. It is readily seen that g_0 is characterized by $g_0(1|W) = \sigma_0(1, W) / (\sigma_0(1, W) + \sigma_0(0, W))$, where $\sigma_0^2(A, W)$ is the conditional variance of Y given (A, W) under Q_0 . It therefore appears that, under this randomization scheme, the treatment arm with higher probability for a patient with baseline covariates W is the one for which the conditional variance of the outcome is higher.

If we knew the optimal design then we could undertake the covariate-adjusted trial consisting in drawing independently observations from P_{Q_0,g_0} . The next task would be to build a regular, asymptotically linear estimator with asymptotic variance $\text{Var}_{P_{Q_0,g_0}} D^*(P_{Q_0,g_0})$ based on the resulting data. In the present situation, we are going to “target” g_0 at some pre-determined interim steps. By targeting g_0 we mean estimating g_0 based on past observations and relying on the resulting estimator to collect the next block of data. In addition to targeting g_0 , each interim analysis will also consist in building an adaptive, targeted, regular and asymptotically linear estimator of ψ_0 . The details of this procedure are presented in Section 2.2.

2.2 Data-Generating Mechanism and Estimation Procedures

Describing the data-generating mechanism amounts to presenting how we target the optimal design g_0 at each interim step, which involves the estimation of the conditional expectation $Q_{Y,0}$. We initiate the description in Section 2.2.1, describe a LASSO estimation procedure of $Q_{Y,0}$ in Section 2.2.2 and the related targeting procedure of g_0 in Section 2.2.3. By then, the data-generating mechanism is fully characterized by recursion.

2.2.1 Initiating the Data-Generating Mechanism

In the sequel, we denote $O_i \equiv (W_i, A_i, Y_i)$ the i th observation that we sample. The indexing reflects the time ordering of the data collection: $j < i$ implies that O_j was collected before or at the same time as O_i . For convenience, we let $\mathbf{O}_n \equiv (O_1, \dots, O_n)$ be the ordered vector of first n observations, with convention $O_0 \equiv \emptyset$. In the adaptive trial, the treatment A_i is drawn conditionally on W_i from the Bernoulli law with parameter $g_i(1|W_i)$, where the randomization scheme $g_i : \mathcal{A} \rightarrow [0, 1]$ depends on past observations $\mathbf{O}_{i-1}$. We set $\mathbf{g}_n \equiv (g_1, \dots, g_n)$, the ordered vector of first n randomization schemes. The data-generating

distribution of $\mathbf{O}_n$ is denoted $\mathbf{P}_{Q_0, \mathbf{g}_n}$. It is formally characterized by the following factorization of the density of $\mathbf{O}_n$ wrt the product of the dominating measures: for any $g \in \mathcal{G}$,

$$\mathbf{P}_{Q_0, \mathbf{g}_n}(\mathbf{O}_n) = \prod_{i=1}^n P_{Q_0, g_i}(O_i) = \prod_{i=1}^n Q_{W,0}(W_i) \times g_i(A_i|W_i) \times P_{Q_0, g}(Y_i|A_i, W_i).$$

Let g^b be the balanced randomization scheme, for which each arm is assigned with probability $1/2$ regardless of baseline covariates. For a pre-specified n_0 , we first draw n_0 independent observations $O_1, \dots, O_{n_0}$ from P_{Q_0, g^b} . At an interim point, suppose one has thus far drawn n observations $\mathbf{O}_n \sim \mathbf{P}_{Q_0, \mathbf{g}_n}$. An estimator of $Q_{Y,0}$ is obtained based on $\mathbf{O}_n$. The next randomization scheme g_{n+1} is defined using the latter estimator and $(\mathbf{O}_n, \mathbf{g}_n)$, then the $(n+1)$ th observation O_{n+1} is drawn from $P_{Q_0, g_{n+1}}$. We will describe the estimation of $Q_{Y,0}$ and construction of g_{n+1} in the two following sections.

2.2.2 LASSO Estimation of the Outcome's Conditional Expectation

Consider $\{b_n\}_{n \geq 1}$ and $\{d_n\}_{d \geq 1}$ two non-decreasing, possibly unbounded sequences over $\mathbb{R}_+$ and, for some $M > 0$ and every $n \geq 1$, introduce the subset

$$B_{M,n} \equiv \{\beta \in \ell^1 : \|\beta\|_1 \leq \min(b_n, M) \text{ and } \forall j \geq d_n, \beta^j = 0\} \quad (3)$$

of $\ell^1 \equiv \{\beta \in \mathbb{R}^{\mathbb{N}} : \sum_{j \in \mathbb{N}} |\beta^j| < \infty\}$. Let $\{\phi_j : j \in \mathbb{N}\}$ be a uniformly bounded set of functions from $\mathcal{A} \times \mathcal{W}$ to $\mathbb{R}$. Without loss of generality, we may assume that $\|\phi_j\|_\infty = 1$ for all $j \in \mathbb{N}$, where $\|\cdot\|_\infty$ denotes the supremum norm. For all $\beta \in \ell^1$, we denote $\Phi_\beta : \mathcal{A} \times \mathcal{W} \rightarrow \mathbb{R}$ the function such that $\Phi_\beta(A, W) \equiv \sum_{j \in \mathbb{N}} \beta^j \phi_j(A, W)$.

The construction of our LASSO estimators of $Q_{Y,0}$ relies on a working model $\mathcal{Q}_1$ and on a loss function L for $Q_{Y,0}$, both specified by the investigators. This means that $Q_{Y,0}$ is the minimizer of $Q_Y \mapsto P_{Q_0, g} L(Q_Y)$ over the set of all conditional expectations of Y given (A, W) , of which $\mathcal{Q}_1$ is a user-specified subset (the value of $g \in \mathcal{G}$ plays no role in this statement). For instance, they can take $\mathcal{Q}_1 \equiv \{Q_{Y, \beta} \equiv \Phi_\beta : \beta \in B_{M,n}\}$ with $M = 1$, and the least-square loss function L characterized over the latter by

$$L(Q_{Y, \beta})(O) \equiv (Y - Q_{Y, \beta}(A, W))^2. \quad (4)$$

They can also take $\mathcal{Q}_1 \equiv \{Q_{Y, \beta} \equiv \text{expit}(\Phi_\beta) : \beta \in B_{M,n}\}$ with M a deterministic upper-bound on $|\text{logit}(Y)|$ (recall that Y is assumed bounded away from 0 and 1), and the quasi negative-log-likelihood loss function L characterized over the latter by

$$-L(Q_{Y, \beta})(O) \equiv Y \log(Q_{Y, \beta}(A, W)) + (1 - Y) \log(1 - Q_{Y, \beta}(A, W)). \quad (5)$$

Note that in both cases, for all $\beta \in B_{M,n}$, $\|Q_{Y, \beta}\|_\infty$ is upper-bounded by a deterministic upper-bound on $|Y|$.

Recall that we have already drawn n observations $\mathbf{O}_n \sim \mathbf{P}_{Q_0, \mathbf{g}_n}$. Given a user-specified reference $g^r \in \mathcal{G}$ that is bounded away from 0 and 1, we estimate $Q_{Y,0}$ with Q_{Y, β_n} , where

$$\beta_n \in \arg \min_{\beta \in B_{M,n}} \frac{1}{n} \sum_{i=1}^n \left(L(Q_{Y, \beta})(O_i) \frac{g^r(A_i|W_i)}{g_i(A_i|W_i)} \right). \quad (6)$$

The above minimization with the constraint $\|\beta\|_1 \leq \min(b_n, M)$, see (3), can be rewritten as a minimization free of the latter constraint by adding a term of the form $\lambda_n \|\beta\|_1$ to the empirical criterion, where λ_n depends on b_n . This is the so-called LASSO procedure introduced by Tibshirani (1996) for the sake of obtaining estimators with fewer nonzero parameter values, thus effectively reducing the number of variables upon which the given solution is dependent. Note that when $d_n = d$ is held constant by choice, (6) should be interpreted as a standard parametric procedure rather than as a LASSO.

2.2.3 Adapting Towards the Optimal Design

We now turn to the construction of the next randomization scheme g_{n+1} .

Our optimal design minimizes $g \mapsto \text{Var}_{P_{Q_0, g}} D^*(P_{Q_0, g})$ over the class $\mathcal{G}$ of all randomization schemes, see (2). We adopt a loss-based approach, by defining g_{n+1} as the minimizer in g of an estimator of $\text{Var}_{P_{Q_0, g}} D^*(P_{Q_0, g})$ over a user-specified class of randomization schemes. This approach is applicable in the largest generality. In the case that W is discrete, or if one is willing to assign treatment only based on a discrete summary measure V of W , g_{n+1} can be defined explicitly as an estimator of the Neyman design based on Q_{Y, β_n} and observations $\mathbf{O}_n$; we refer the readers to (Chambaz and van der Laan, 2013) for details.

To proceed, we first note that, for all $g' \in \mathcal{G}$,

$$g_0 = \arg \min_{g \in \mathcal{G}} E_{P_{Q_0, g'}} \frac{(Y - Q_{Y, 0}(A, W))^2}{g(A|W)g'(A|W)}.$$

This equality teaches us that for the sake of estimating g_0 using observations drawn from $P_{Q_0, g'}$ we may consider the loss function L_{Q_Y} characterized over $\mathcal{G}$ by

$$L_{Q_Y}(g)(O) \equiv \frac{(Y - Q_Y(A, W))^2}{g(A|W)},$$

provided it is weighted by $1/g'(A|W)$. Note that this loss function is indexed by a given Q_Y .

Recall that we have already drawn n observations $\mathbf{O}_n \sim P_{Q_0, \mathbf{g}_n}$ and estimated $Q_{Y, 0}$ with Q_{Y, β_n} . Now, given a class $\mathcal{G}_1 \subset \mathcal{G}$ of randomization schemes uniformly bounded away from 0 and 1, we define the next randomization scheme as

$$g_{n+1} \in \arg \min_{g \in \mathcal{G}_1} \frac{1}{n} \sum_{i=1}^n \frac{L_{Q_{Y, \beta_n}}(g)(O_i)}{g_i(A_i|W_i)} = \arg \min_{g \in \mathcal{G}_1} \frac{1}{n} \sum_{i=1}^n \frac{(Y - Q_{Y, \beta_n}(O_i))^2}{g(A_i|W_i)g_i(A_i|W_i)} \quad (7)$$

This completes the description of our data-generating mechanism.

2.3 Targeted Maximum Likelihood Estimation

Given n observations $\mathbf{O}_n \sim P_{Q_0, \mathbf{g}_n}$ and the estimator Q_{Y, β_n} of $Q_{Y, 0}$ defined in (6), we may carry out the estimation of the parameter of interest ψ_0 . We adopt the targeted minimum loss estimation methodology. In the setting of a covariate-adjusted RCT with fixed design, a TMLE estimator is unbiased and asymptotically Gaussian regardless of the specification of the working model used for the estimation of $Q_{Y, 0}$. It is known that unbiasedness and asymptotic normality still hold in the context of this technical report (CARA RCT for the estimation of ψ_0 based on copies of O), provided that the randomization schemes depend on W only through a summary measure taking finitely many values and that the working model used for the estimation of $Q_{Y, 0}$ be a simple linear model (this basically amounts to taking $d_n = d$ constant and $b_n = M$) in Section 2.2.2), see (Chambaz and van der Laan, 2013). Yet by relying on more flexible randomization schemes and on more adaptive estimators of $Q_{Y, 0}$ we may achieve a greater efficiency through better estimation of the optimality criteria that may facilitate adaptation toward the optimal design, better adjustment of the variables that may directly improve on the estimation of the parameter of interest, and a more accurate estimation of the variance of the estimator.

In a glimpse, the proposed strategy consists in targetedly fluctuating the initial estimator Q_{Y, β_n} by minimizing a pre-specified loss along a least favorable (wrt ψ_0) submodel through Q_{Y, β_n} , and then evaluating Ψ at the resulting updated estimator of Q_0 . Formally, consider the following one-dimensional parametric working model through Q_{Y, β_n} : for a given closed, bounded interval $\mathcal{E} \subset \mathbb{R}$ containing 0 in its interior,

$$\{Q_{Y, \beta_n}(\varepsilon) \equiv \text{expit}(\text{logit}(Q_{Y, \beta_n}) + \varepsilon H(g_n)) : \varepsilon \in \mathcal{E}\}, \quad (8)$$

with notation $H(g)(O) \equiv \frac{2A-1}{g(A|W)}$ for every $g \in \mathcal{G}_1$. This model passes through Q_{Y,β_n} at $\varepsilon = 0$ in such a way that $\frac{\partial}{\partial \varepsilon} L(Q_{Y,\beta_n}(\varepsilon))|_{\varepsilon=0} = D_Y^*(Q_{Y,\beta_n}, g_n)$. The optimal fluctuation parameter ε_n minimizes the weighted empirical risk along the working model:

$$\varepsilon_n \in \arg \min_{\varepsilon \in \mathcal{E}} \frac{1}{n} \sum_{i=1}^n L(Q_{Y,\beta_n}(\varepsilon))(O_i) \frac{g_n(A_i|W_i)}{g_i(A_i|W_i)}. \quad (9)$$

Set $Q_{Y,\beta_n}^* \equiv Q_{Y,\beta_n}(\varepsilon_n)$ then $Q_{\beta_n}^* \equiv (Q_{W,n}, Q_{Y,\beta_n}^*)$ where $Q_{W,n}$ is the empirical marginal distribution of the W . The TMLE estimator of ψ_0 is finally defined as

$$\psi_n^* \equiv \frac{1}{n} \sum_{i=1}^n Q_{Y,\beta_n}^*(1, W_i) - Q_{Y,\beta_n}^*(0, W_i).$$

It satisfies $\psi_n^* = \Psi(P_{Q_{\beta_n}^*, g})$ for any $g \in \mathcal{G}$.

3 Asymptotics

We first introduce further notation in Section 3.1 then we successively investigate the convergence of the targeted CARA design in Section 3.2 and the asymptotic behavior of the TMLE estimator in Section 3.3.

3.1 Notation

In general, given a known $g \in \mathcal{G}$ and an observation O drawn from $P_{Q_0, g}$, $Z \equiv g(A|W)$ is a deterministic function of g and O . Note that Z should be interpreted as a weight associated with O and will be used as such. Therefore, we can augment O with Z , i.e., substitute (O, Z) for O , while still denoting $(O, Z) \sim P_{Q_0, g}$. In particular, during the course of our trial, conditionally on $\mathbf{O}_{i-1}$, the randomization scheme g_i is known and we can substitute $(O_i, Z_i) = (O_i, g_i(A_i|W_i)) \sim P_{Q_0, g_i}$ for O_i drawn from P_{Q_0, g_i} . By uniform boundedness of $\mathcal{G}_1$, the inverse weights $1/g_i(A_i|W_i)$ are bounded.

The empirical distribution of $\mathbf{O}_n$ is denoted P_n . For a function $f : \mathcal{O} \times [0, 1] \rightarrow \mathbb{R}^d$, we will use the notation $P_n f \equiv n^{-1} \sum_{i=1}^n f(O_i, Z_i)$. Likewise, for any fixed $P_{Q, g} \in \mathcal{M}$, $P_{Q, g} f \equiv E_{P_{Q, g}} f(O, Z)$ and, for each $i = 1, \dots, n$, $P_{Q_0, g_i} f \equiv E_{Q_0, g_i} [f(O_i, Z_i) | \mathbf{O}_{i-1}]$, $\mathbf{P}_{Q_0, \mathbf{g}_n} f \equiv n^{-1} \sum_{i=1}^n E_{Q_0, g_i} [f(O_i, Z_i) | \mathbf{O}_{i-1}]$.

We endow the set $\{Q_Y : P_{Q, g} \in \mathcal{M}\}$ of all conditional expectations of Y given (A, W) under $P_{Q, g} \in \mathcal{M}$ with the norm $\|\cdot\|_{Y,0}$ characterized by

$$\|Q_Y - Q'_Y\|_{Y,0}^2 \equiv E_{P_{Q_0, g^r}} (Q_Y(A, W) - Q'_Y(A, W))^2.$$

Similarly, we endow the set $\mathcal{G}$ with the norm $\|\cdot\|_{A,0}$ characterized by

$$\|g - g'\|_{A,0}^2 \equiv E_{Q_{W,0}} (g(1|W) - g'(1|W))^2.$$

For any class $\mathcal{F}$ of functions equipped with a norm $\|\cdot\|$ and $\varepsilon > 0$, $N(\mathcal{F}, \|\cdot\|, \varepsilon)$ is the ε -bracketing of $\mathcal{F}$ wrt $\|\cdot\|$ and $J(1, \mathcal{F}, \|\cdot\|) \equiv \int_0^1 \sqrt{\log N(\mathcal{F}, \|\cdot\|, \varepsilon)} d\varepsilon$ is the corresponding bracketing entropy (evaluated at 1). Finally, the uniform norm of a real-valued operator Π on $\mathcal{F}$ is $\|\Pi\|_{\mathcal{F}} \equiv \sup_{f \in \mathcal{F}} |\Pi(f)|$.

3.2 Convergence of the Targeted CARA Design

Our first concern is the convergence of Q_{Y,β_n} , see (6).

Proposition 1 (Convergence of Q_{Y,β_n}). *Consider the following assumptions:*

A1. The conditional density under Q_0 of Y given (A, W) wrt some dominating measure is bounded away from 0.

A2. There exists a unique $\beta_0 \in \bigcup_{n \geq 1} B_{M,n}$ such that

$$\beta_0 \in \arg \min_{\beta \in \bigcup_{n \geq 1} B_n} P_{Q_0, g^r} L(Q_{Y, \beta}).$$

A3. It holds that $d_n = O(n^r)$ for some $r > 0$ and $\sup_{\beta \in B_{M,n}} |(P_n - \mathbf{P}_{Q_0, \mathbf{g}_n})L(Q_{Y, \beta})| = o_P(1)$.

Under **A1–A3**, $\|Q_{Y, \beta_n} - Q_{Y, \beta_0}\|_{Y,0} \rightarrow 0$ in probability.

In summary, Q_{Y, β_n} converges to Q_{Y, β_0} in probability for the norm $\|\cdot\|_{Y,0}$ if such a limit exists and if the dimension of the LASSO parameters grows polynomially wrt the sample size. Note that the limit Q_{Y, β_0} depends on the user-supplied reference design g^r . Based on (Dümbgen, Van De Geer, Veraar, and Wellner, 2010) we show that **A3** holds for instance when L is given by (4) and $\mathcal{Q}_1 = \{\Phi_\beta : \beta \in B_{1,n}\}$. The proof of Proposition 1 relies on empirical process theory for martingales, chaining arguments and tools developed by van Handel (2011). It requires that the sequence $\{g_n\}_{n \geq 1}$ of randomization schemes be uniformly bounded away from 0 and 1, which is guaranteed by specification of the user-supplied set $\mathcal{G}_1$.

We now turn to the convergence of the targeted CARA design $\{g_n\}_{n \geq 1}$, see (7), toward a fixed, limiting design $g_0^* \in \mathcal{G}_1$.

Proposition 2 (Convergence of the targeted CARA Design). *Consider the setup of Proposition 1 and the following additional assumptions:*

A4. There exists a unique $g_0^* \in \mathcal{G}_1$ such that

$$g_0^* \in \arg \min_{g \in \mathcal{G}_1} P_{Q_0, g^r} \frac{L_{Q_{Y, \beta_0}}(g)}{g^r}. \quad (10)$$

A5. The class $1/\mathcal{G}_1 \equiv \{1/g : g \in \mathcal{G}_1\}$ satisfies the finite entropy condition $J(1, 1/\mathcal{G}_1, \|\cdot\|_{A,0}) < \infty$.

Under **A1–A5**, $\|g_n(1|W) - g_0^*(1|W)\|_{A,0} \rightarrow 0$ in probability.

We have already emphasized that through the choice of $\mathcal{G}_1$, the investigators of the RCT benefit from a great flexibility in treatment allocation. The main constraint on $\mathcal{G}_1$ is **A4**, a condition on the complexity/richness of the class. We refer the reader to (van der Vaart, 1998, Examples 19.7–19.11, Lemma 19.15) for typical examples. They notably include “well-behaved” parametric classes and VC classes. In particular, $\mathcal{G}_1$ can consist of randomization schemes such that the allocation probabilities only depend on W through a discrete summary measure of it, as considered in (Chambaz and van der Laan, 2013).

The limiting randomization scheme g_0^* depends on the user-supplied reference design g^r only through Q_{Y, β_0} : replacing g^r with any $g \in \mathcal{G}$ in (10) does not alter the definition of g_0^* . Furthermore, g_0^* can be interpreted as the most optimal design in $\mathcal{G}_1$ given the limiting conditional outcome model Q_{Y, β_0} :

$$g_0^* \in \arg \min_{g \in \mathcal{G}_1} \text{Var}_{P_{Q_0, g}} D_Y^*(Q_{Y, \beta_0}, g) = \arg \min_{g \in \mathcal{G}_1} \left\{ \text{Var}_{P_{Q_0, g}} D_Y^*(Q_0, g) + P_{Q_0, g} \frac{(Q_{Y,0} - Q_{Y, \beta_0})^2}{g^2} \right\}.$$

Comparing the above equality with (2) yields that $g_0^* = g_0$, the Neyman design, whenever $Q_{Y, \beta_0} = Q_{Y,0}$. In general, g_0^* minimizes an objective function writing as the sum of the Cramér-Rao lower bound and a second-order term residual. This underscores the motivation for using a flexible estimator in estimating $Q_{Y,0}$: by minimizing this second-order residual of the limiting conditional outcome model, we are closer to adapting toward the desired optimal design.

3.3 Consistency and Central Limit Theorem

As with the initial LASSO estimators of the conditional outcome, we are firstly concerned with the convergence of the updated estimators Q_{Y,β_n}^* :

Proposition 3 (Consistency). *Consider the setups of Propositions 1 and 2 and the follow additional assumption:*

A6. *There exists a unique $\varepsilon_0 \in \mathcal{E}$ such that*

$$\varepsilon_0 \in \arg \min_{\varepsilon \in \mathcal{E}} P_{Q_0, g_0^*} L(Q_{Y, \beta_0}(\varepsilon)).$$

*Assume that **A1–A6** are met and define $Q_{Y, \beta_0}^* \equiv Q_{Y, \beta_0}(\varepsilon_0)$. It holds that $\|Q_{Y, \beta_n}^* - Q_{Y, \beta_0}^*\|_{Y, 0} \rightarrow 0$ in probability. Moreover, ψ_n^* consistently estimates ψ_0 .*

If $Q_{Y, \beta_0} = Q_{Y, 0}$ then $\varepsilon_0 = 0$: the updating procedure preserves the consistency of the initial estimator $\Psi(P_{Q_{Y, \beta_n}, g})$ for any $g \in \mathcal{G}$. More importantly, Proposition 3 guarantees that even if $Q_{Y, \beta_0} \neq Q_{Y, 0}$ then ψ_n^* still consistently estimates ψ_0 , by double-robustness. Nonetheless, the convergence of the updated estimators Q_{Y, β_n}^* (to the truth or otherwise) is crucial for studying the asymptotic behavior of ψ_n^* .

Proposition 4 (Central Limit Theorem for ψ_n^*). *Consider the setups of Propositions 1, 2 and 3 and the following additional assumption:*

A7. *For any deterministic function F , $F(O) = 0$ P_{Q_0, g_0^*} -almost surely implies that $F = 0$.*

*Assume that **A1–A7** are met. For both $\beta = \beta_0$ and $\beta = \beta_n$, introduce $d_{Y, \beta}^*$ and $q_{Y, \beta}^*$ characterized by*

$$\begin{aligned} d_{Y, \beta}^*(O, Z) &\equiv \frac{2A - 1}{Z} \left(Y - Q_{Y, \beta}^*(A, W) \right), \\ q_{Y, \beta}^* &\equiv Q_{Y, \beta}^*(1, W) - Q_{Y, \beta}^*(0, W) \end{aligned}$$

and, for any $g \in \mathcal{G}$,

$$\Sigma_n \equiv \frac{1}{n} \sum_{i=1}^n \left(d_{Y, \beta_n}^*(O_i, Z_i) + D_W^*(P_{Q_{\beta_n}, g}^*(W_i)) \right)^2. \quad (11)$$

Then $(\Sigma_n/n)^{-1/2}(\psi_n^ - \psi_0)$ converges in distribution to the standard normal distribution.*

The asymptotic results in Proposition 4 underpin the statistical analysis of the proposed targeted CARA RCT. In particular, denoting $\xi_{1-\alpha/2}$ the $(1 - \alpha/2)$ -quantile of the standard normal distribution, $[\psi_n^* \pm \xi_{1-\alpha/2}(\Sigma_n/n)^{1/2}]$ is a confidence interval of asymptotic level $(1 - \alpha)$.

4 Simulation Study

We present here the results of a simulation study of the performances of the targeted procedure exposed in the previous sections.

4.1 Simulation Scheme

We rely on the same simulation scheme as in (Chambaz and van der Laan, 2013). For completeness, let us recall that Q_0 is such that:

- the baseline covariate W equals (U, V) , where U and V are independently drawn with U uniformly distributed on $[0, 1]$ and $Q_{W,0}(V = 1) = 1/2$, $Q_{W,0}(V = 2) = 1/3$, $Q_{W,0}(V = 3) = 1/6$;
- the conditional distribution of Y given (A, W) is the Gamma distribution with conditional mean

$$Q_{Y,0}(A, W) = 2U^2 + 2U + 1 + \left(AV + \frac{1-A}{1+V} \right)$$

and conditional variance

$$\sigma_0^2(Y|A, W) = \left(U + A(1+V) + \frac{1-A}{1+V} \right)^2.$$

The marginal treatment effect on an additive scale satisfies $\psi_0 = \frac{91}{72} \simeq 1.264$.

We target the optimal designs corresponding to eight parametric working models $\mathcal{G}_{11}, \dots, \mathcal{G}_{18}$ that we present in Table 1.

working model	parametric form	dimension	optimal variance
$\mathcal{G}_{11}$	θ_0	1	18.50
$\mathcal{G}_{12}$	$\sum_{v=1}^3 \theta_v \mathbf{1}\{V = v\}$	3	18.18
$\mathcal{G}_{13}$	$\theta_0 + \theta_1 U$	2	18.37
$\mathcal{G}_{14}$	$\sum_{v=1}^3 \theta_v \mathbf{1}\{V = v\} + \theta_4 U$	4	18.05
$\mathcal{G}_{15}$	$\theta_0 + \sum_{v=1}^3 \theta_v \mathbf{1}\{V = v\} U$	4	18.12
$\mathcal{G}_{16}$	$\sum_{v=1}^3 \theta_v \mathbf{1}\{V = v\} + \theta_4 U + \sum_{v=2}^3 \theta_{3+v} \mathbf{1}\{V = v\} U$	6	18.01
$\mathcal{G}_{17}$	$\theta_0 + \sum_{v=1}^3 \theta_v \mathbf{1}\{V = v\} U + \sum_{v=1}^3 \theta_{4+v} \mathbf{1}\{V = v\} U^2$	7	18.36
$\mathcal{G}_{18}$	$\sum_{v=1}^3 \theta_v \mathbf{1}\{V = v\} + \theta_4 U + \theta_5 U^2 + \sum_{v=2}^3 \theta_{4+v} \mathbf{1}\{V = v\} U + \sum_{v=2}^3 \theta_{6+v} \mathbf{1}\{V = v\} U^2$	9	18.03

Table 1: **Parametric working models** $\mathcal{G}_{1k}$ ($k = 1, \dots, 8$). In the second column, we report the parametric forms of $\logit((g_\theta(W) - \delta)/(1 - 2\delta))$ for generic elements $g_\theta \in \mathcal{G}_{1k}$ ($k = 1, \dots, 8$). We set $\delta = 10^{-2}$. In the third column, we give the dimensions of the models. In the fourth column, we report the numerical values of $\arg\min_{g \in \mathcal{G}_{1k}} \text{Var}_{P_{Q_{0,g}}} D^*(P_{Q_{0,g}})(O)$ ($k = 1, \dots, 8$), with precision 10^{-2} . Recall that $\text{Var}_{P_{Q_{0,g^b}}} D^*(P_{Q_{0,g}})(O) = 23.87$, with precision 10^{-2} .

In addition to the latter parametric working models, we consider eight statistical procedures for the estimation of the conditional expectation $Q_{Y,0}$. Four of them consist in parametric estimation on small-dimensional models $\mathcal{Q}_{11}, \dots, \mathcal{Q}_{14}$. In contrast, the four others rely on moderate-dimensional parametric models, ℓ^1 -penalization and cross-validation to select the best regularization parameter. We denote $\mathcal{Q}_{15}, \dots, \mathcal{Q}_{18}$ these “machine-learning”, as opposed to “parametric”, procedures/models, which embody the LASSO estimating procedure of Section 2.2.2. We summarize in Table 2 what are $\mathcal{Q}_{11}, \dots, \mathcal{Q}_{18}$. All procedures involve the logistic loss, even though the support of the marginal distribution of Y under P_0 is $\mathbb{R}_+$, not $[0, 1]$. In fact, given a sample $O_1, \dots, O_n$, we first scale $Y_1, \dots, Y_n$ to $[0, 1]$, then regress the scaled outcomes on (A, W) based on the logistic loss and one procedure among $\mathcal{Q}_{11}, \dots, \mathcal{Q}_{18}$, then scale back the resulting conditional expectation to the original range of the observed outcomes.

Set $B = 1000$ and let $n = (250, 500, 750, 1000, 1250, 1500, 1750, 2000, 2250, 2500)$ be a sequence of sample sizes. For each combination $(k, l) \in \{1, \dots, 8\}^2$, we repeatedly simulate $B = 1000$ times a targeted CARA RCT based on $\mathcal{G}_{1k}$ and $\mathcal{Q}_{1l}$, performing an update of the randomization scheme and the computation of the TMLE of ψ_0 at every intermediate sample size n_i ($1 \leq i \leq 10$), which we denote $\psi_{n_i, klb}^*$. The simulations are mutually independent. The associated 95%-confidence intervals $\mathcal{I}_{n_i, klb}$ rely on estimated

	working model	parametric form	dimension
parametric	$\mathcal{Q}_{11}$	$\sum_{v=1}^3 \theta_v \mathbf{1}\{V=v\} + \theta_4 U + \theta_5 A$	5
	$\mathcal{Q}_{12}$	$\theta_0 + A (\theta_1 U + \sum_{v=2}^3 \theta_v \mathbf{1}\{V=v\})$ $+ (1-A) (\theta_4 U + \sum_{v=2}^3 \theta_{3+v} \mathbf{1}\{V=v\})$	7
	$\mathcal{Q}_{13}$	$A (\sum_{v=1}^3 \theta_v \mathbf{1}\{V=v\} + \theta_4 U)$ $+ (1-A) (\sum_{v=1}^3 \theta_{4+v} \mathbf{1}\{V=v\} + \theta_8 U)$	8
	$\mathcal{Q}_{14}$	$A (\sum_{v=1}^3 \theta_v \mathbf{1}\{V=v\} + \theta_4 U + \theta_5 U^2)$ $+ (1-A) (\sum_{v=1}^3 \theta_{5+v} \mathbf{1}\{V=v\} + \theta_9 U + \theta_{10} U^2)$	10
LASSO	$\mathcal{Q}_{15}$	$A (\sum_{v=1}^3 \theta_v \mathbf{1}\{V=v\} + \theta_4 U + \theta_5 U^2)$ $+ (1-A) (\sum_{v=1}^3 \theta_{5+v} \mathbf{1}\{V=v\} + \theta_9 U + \theta_{10} U^2)$	10
	$\mathcal{Q}_{16}$	$A (\sum_{v=1}^3 \theta_v \mathbf{1}\{V=v\} + \sum_{l=1}^5 \theta_{3+l} U^l)$ $+ (1-A) (\sum_{v=1}^3 \theta_{8+v} \mathbf{1}\{V=v\} + \sum_{l=1}^5 \theta_{11+l} U^l)$	16
	$\mathcal{Q}_{17}$	$A (\sum_{v=1}^3 \theta_v \mathbf{1}\{V=v\} + \sum_{l=1}^{10} \theta_{3+l} U^l)$ $+ (1-A) (\sum_{v=1}^3 \theta_{13+v} \mathbf{1}\{V=v\} + \sum_{l=1}^{10} \theta_{16+l} U^l)$	26
	$\mathcal{Q}_{18}$	$A (\sum_{v=1}^3 \theta_v \mathbf{1}\{V=v\} + \sum_{l=1}^{20} \theta_{3+l} U^l)$ $+ (1-A) (\sum_{v=1}^3 \theta_{23+v} \mathbf{1}\{V=v\} + \sum_{l=1}^{20} \theta_{26+l} U^l)$	46

Table 2: **Working models $\mathcal{Q}_{1k}(k=1, \dots, 8)$ for the conditional expectation $Q_{Y,0}$.** In the second column, we report the parametric form of $\text{logit}((q_\theta(A, W) - \delta)/(1 - 2\delta))$ for generic elements $q_\theta \in \mathcal{Q}_{1k}$ ($k = 1, \dots, 8$). We set $\delta = 10^{-2}$. In the third column, we give the dimensions of the models. All working models are exploited in combination with the quasi negative-log-likelihood loss function (5). Models $\mathcal{Q}_{11}, \mathcal{Q}_{12}, \mathcal{Q}_{13}, \mathcal{Q}_{14}$ are straightforwardly fitted by relying on the R function `glm`. Models $\mathcal{Q}_{15}, \mathcal{Q}_{16}, \mathcal{Q}_{17}, \mathcal{Q}_{18}$ are LASSO-fitted by relying on the R function `glmnet`.

variances of the TMLE as given in (11). For each combination (k, l) and intermediate sample size n_i , we compute the empirical variance of the corresponding TMLE

$$\hat{S}_{n_i,kl} = \frac{1}{B} \sum_{b=1}^B \psi_{n_i,klb}^{*2} - \left(\frac{1}{B} \sum_{b=1}^B \psi_{n_i,klb} \right)^2$$

and the empirical coverage of the corresponding confidence interval

$$\hat{C}_{n_i,kl} = \frac{1}{B} \sum_{b=1}^B \mathbf{1}\{\psi_0 \in \mathcal{J}_{n_i,klb}\}.$$

The simulation study is conducted using R (R Core Team, 2014) and the package `glmnet` (Friedman, Hastie, and Tibshirani, 2010).

4.2 Discussion of the Results

4.2.1 Coverage

We propose an evaluation of the coverage performances based on testing. For every $(k, l) \in \{1, \dots, 8\}^2$ and n_i ($1 \leq i \leq 10$), the statistic $B \times \hat{C}_{n_i,kl}$ follows a Binomial distribution with parameter $(B, \pi_{n_i,kl})$ for some $\pi_{n_i,kl} \in [0, 1]$. Denote $\hat{p}_{n_i,kl}^{95}$ the exact p -value of the one-sided binomial test of $H_{n_i,kl}^{95} : \pi_{n_i,kl} \geq 95\%$ against “ $\pi_{n_i,kl} < 95\%$ ”. Under $H_{n_i,kl}^{95}$, $\hat{p}_{n_i,kl}^{95}$ is drawn from the uniform distribution on $[0, 1]$.

For every n_i ($1 \leq i \leq 10$), we carry out one-sample Kolmogorov-Smirnov tests of the null stating that the common law of $\{\hat{p}_{n_i,kl}^{95} : 1 \leq k \leq 8, l \in \mathcal{L}\}$ ($\mathcal{L} \subset \{1, \dots, 8\}$) is the uniform distribution on $[0, 1]$ against

n_i	250	500	750	1000	1250	1500	1750	2000	2250	2500
$\bigcap_{1 \leq k \leq 8} \bigcap_{1 \leq l \leq 8} H_{n_i,kl}^{95}$	< 0.001	< 0.001	0.011	0.003	0.011	0.006	0.110	0.362	0.003	0.059
$\bigcap_{1 \leq k \leq 8} H_{n_i,kl}^{95}$	< 0.001	0.015	0.023	< 0.001	0.151	0.034	0.025	0.080	0.281	0.414
$\bigcap_{1 \leq k \leq 8} \bigcap_{5 \leq l \leq 8} H_{n_i,kl}^{95}$	< 0.001	< 0.001	0.175	0.567	0.004	0.037	0.785	0.804	0.004	0.072
$\bigcap_{1 \leq k \leq 8} H_{n_i,kl}^{94}$	0.028	0.999	0.999	0.999	0.999	0.999	0.999	0.999	0.999	0.999

Table 3: **Evaluating the coverage performances based on testing.** The first row gives p -values of Kolmogorov-Smirnov tests of the null consisting of the intersection of all $H_{n_i,kl}^{95}$. The second and third rows give p -values of Kolmogorov-Smirnov tests of the nulls consisting in the intersections of all $H_{n_i,kl}^{95}$ based on parametric procedures (second row) and of all $H_{n_i,kl}^{95}$ based on LASSO procedures (third row). The fourth row gives p -values of Kolmogorov-Smirnov tests of the null consisting of the intersection of all $H_{n_i,kl}^{94}$.

the alternative that the common law is stochastically smaller than the uniform distribution on $[0, 1]$. Rejecting the null for its alternative indicates a defective coverage. The p -values of four such Kolmogorov-Smirnov tests are reported in Table 3. The first row corresponds to the choice $\mathcal{L} = \{1, \dots, 8\}$. It teaches us that the expected 95%-coverage is generally not guaranteed. One may wonder if the same conclusion holds when focusing in turn on the parametric procedures (set $\mathcal{L} = \{1, \dots, 4\}$) or on the LASSO procedures (set $\mathcal{L} = \{5, \dots, 8\}$). Inspecting the second and third rows of Table 3 does not reveal an interesting pattern. One may now wonder to what extent the 95%-coverage is deficient. To answer this question, we proceed similarly. We denote $\hat{p}_{n_i,kl}^{94}$ the exact p -value of the one-sided binomial test of $H_{n_i,kl}^{94}$: “ $\pi_{n_i,kl} \geq 94\%$ ” against “ $\pi_{n_i,kl} < 94\%$ ”. Under $H_{n_i,kl}^{94}$, $\hat{p}_{n_i,kl}^{94}$ is drawn from the uniform distribution on $[0, 1]$. For every n_i ($1 \leq i \leq 10$), we carry out a one-sample Kolmogorov-Smirnov test of the null stating that the common law of $\{\hat{p}_{n_i,kl}^{94} : 1 \leq k \leq 8, 1 \leq l \leq 8\}$ is the uniform distribution on $[0, 1]$ against the alternative that the common law is stochastically smaller than the uniform distribution on $[0, 1]$. The p -values of these tests are reported in the fourth row of Table 3. The conclusion is clear and satisfactory: even if the 95%-confidence intervals fail to guarantee the wished coverage, one can safely consider them as valid 94%-confidence intervals.

4.2.2 Standard Deviation

Here we investigate how the targeted CARA RCT behaves in terms of standard deviation of the produced estimators. As in the previous subsection, the investigation relies on testing. For every $(k, l) \in \{1, \dots, 8\}^2$ and n_i ($1 \leq i \leq 10$), we first compute the statistic

$$T_{n_i,kl} = \frac{\frac{1}{B} \sum_{b=1}^B (\Sigma_{n_i,klb})^{1/2} - (\hat{S}_{n_i,kl})^{1/2}}{\left(\frac{1}{B} \sum_{b=1}^B \Sigma_{n_i,klb} - \left(\frac{1}{B} \sum_{b=1}^B (\Sigma_{n_i,klb})^{1/2} \right)^2 \right)^{1/2}},$$

where $\Sigma_{n_i,klb}$ is the estimated variance of the TMLE produced at intermediate sample size n_i by the b th simulated targeted CARA RCT based on $\mathcal{G}_{1k}$ and $\mathcal{Q}_{1l}$, see (11). Thus, $T_{n_i,kl}$ sheds some light on the estimation of the standard deviation of the TMLE $\psi_{n_i}^*$ at sample size n_i by $(\Sigma_{n_i}/n)^{1/2}$ for the targeted CARA RCT based on $\mathcal{G}_{1k}$ and $\mathcal{Q}_{1l}$.

For every n_i ($1 \leq i \leq 10$), we perform a Lilliefors test of normality based on the sample $\{T_{n_i,kl} : 1 \leq k \leq 8, l \in \mathcal{L}\}$ with $\mathcal{L} = \{1, \dots, 8\}$. The p -values of these tests are reported in Table 4. They teach us that there is no stark evidence of non-normality across the ten intermediate sample sizes. This first conclusion justifies the next step: for every n_i ($1 \leq i \leq 10$), we perform a one-sided Student test of “ $\mu_{n_i} \geq 0$ ” against “ $\mu_{n_i} < 0$ ”, where μ_{n_i} denotes the mean of the common distribution of $\{T_{n_i,kl} : 1 \leq k \leq 8, l \in \mathcal{L}\}$ with $\mathcal{L} = \{1, \dots, 8\}$. The p -values of these tests are reported in the two first rows of Table 4. Adjusting for multiple testing in terms of the Benjamini and Yekutieli procedure for controlling the false discovery rate at the 5% level, we conclude that estimating the variance as in (11) is over-optimistic at least for intermediate sample sizes smaller than or equal to $n_3 = 750$. One may wonder if the same conclusions hold when focusing in turn on the parametric procedures (set $\mathcal{L} = \{1, \dots, 4\}$) or on the LASSO procedures (set $\mathcal{L} = \{5, \dots, 8\}$). Inspecting separately the third and fourth rows of Table 4 on one hand then the fifth and sixth rows on the other hand leads to the conclusion that estimating the variance as in (11) is over-optimistic only for intermediate sample sizes smaller than or equal to $n_2 = 500$, still adjusting for multiple testing in terms of the Benjamini and Yekutieli procedure for controlling the false discovery rate at the 5% level.

The gap between the conclusions reached when considering all procedures or the parametric and LASSO ones separately may be simply explained by a loss of power due to the reduction of sample size (64 versus 32), or by subtle differences induced by the nature of $\mathcal{Q}_l$. In any case, in light of Section 4.2.1, the under-estimation of the true variance based on (11) is necessarily slight at most.

n_i	250	500	750	1000	1250	1500	1750	2000	2250	2500
Lilliefors	0.670	0.330	0.866	0.033	0.538	0.837	0.133	0.528	0.466	0.022
Student	< 0.001	< 0.001	0.002	0.006	0.008	0.012	0.007	0.008	0.044	0.038
Lilliefors	0.755	0.043	0.270	0.021	0.543	0.620	0.206	0.172	0.685	0.206
Student	< 0.001	< 0.001	0.013	0.026	0.025	0.026	0.021	0.036	0.226	0.420
Lilliefors	0.561	0.894	0.864	0.517	0.500	0.314	0.251	0.783	0.971	0.283
Student	< 0.001	< 0.001	0.044	0.059	0.087	0.116	0.084	0.063	0.050	0.011

Table 4: **Investigating the targeted CARA RCT in terms of standard deviation of the produced estimators.** In the first row we report the p -values of the Lilliefors tests of normality of the sample $\{T_{n_i,kl} : 1 \leq k, l \leq 8\}$ ($1 \leq i \leq 10$). In the second row, we report the p -values of the Student tests of “ $\mu_{n_i} \geq 0$ ” against “ $\mu_{n_i} < 0$ ”, where μ_{n_i} denotes the mean of the common distribution of $\{T_{n_i,kl} : 1 \leq k, l \leq 8\}$. In the third and fourth rows (fifth and sixth rows, respectively), we report the p -values of the same Lilliefors and Student tests based on the samples $\{T_{n_i,kl} : 1 \leq k \leq 8, 1 \leq l \leq 4\}$ corresponding to parametric procedures (on the samples $\{T_{n_i,kl} : 1 \leq k \leq 8, 5 \leq l \leq 8\}$ corresponding to LASSO procedures, respectively).

5 Discussion

We have presented in this technical report a new group-sequential CARA RCT design and inferential procedure built on top of it. The procedure is targeted in the sense that (i) the sequence of randomization schemes is group-sequentially determined by targeting a user-specified optimal randomization design based on accruing data and, (ii) our estimator of the user-specified parameter of interest, seen as the value of a functional evaluated at the true, unknown distribution of the data, is targeted toward it by following the paradigm of targeted minimum loss estimation. We focused for clarity on the case that the parameter of interest is the marginal effect of a binary treatment and that the targeted optimal design is the Neyman allocation, in an effort to produce an estimator with smaller asymptotic variance, but our methodology extends beyond this instructive framework. For clarity too, we considered the case that the estimator

of the conditional outcome given treatment and baseline covariates, a key element of the procedure, is obtained by LASSO regression, although our methodology can hinge on a wide class of data-adaptive estimators. Under mild assumptions, the resulting sequence of randomization schemes converges to a limiting design, and the TMLE estimator is consistent and asymptotically Gaussian, with an asymptotic variance that we can estimate too. Thus we can build valid confidence intervals of given asymptotic levels. A simulation study confirms our theoretical results. Across 64 different choices of pairs of working models and 10 intermediate sample sizes ranging from 250 to 2500, there is no empirical evidence that our 95%-confidence intervals do not provide at least 94%-coverage, based on 1000 independent replications. In addition, in the same framework, there is no empirical evidence that our estimators of the variances of the TMLE estimators are over-optimistic for sample sizes larger than 500, adjusting for multiple testing in terms of the Benjamini and Yekutieli procedure for controlling the false discovery rate at the 5% level. For smaller sample sizes, the under-estimation is slight at most.

We will soon make available a R package to allow interested readers to test the procedure. The proofs of our theoretical results will be presented in a forthcoming article. In this article, we will also describe and study a more general targeted CARA RCT design and related TMLE methodology involving possibly aggressive data-adaptive/machine-learning procedures and not only LASSO regression. In the future, we will also consider alternative strategies to randomly assign successive patients to the treatment arms in such a way that the overall empirical conditional distribution of treatment given baseline covariates be as close as possible to the current best estimator of the targeted optimal design. This will require both new theoretical developments and simulation studies.

Acknowledgements.

Antoine Chambaz acknowledges the support of the French Agence Nationale de la Recherche (ANR), under grant ANR-13-BS01-0005 (project SPADRO). Mark van der Laan was supported by NIH grant R01 AI074345-06A1. Wenjing Zheng acknowledges the support of the Embassy of France in the United States through a Chateaubriand Fellowship–Science, Technology, Engineering, and Mathematics (STEM).

References

- A. C. Atkinson and A. Biswas. Adaptive biased-coin designs for skewing the allocation proportion in clinical trials with normal responses. *Stat. Med.*, 24(16):2477–2492, 2005.
- U. Bandyopadhyay and A. Biswas. Adaptive designs for normal responses with prognostic factors. *Biometrika*, 88(2):409–419, 2001.
- P. J. Bickel, C. A. J. Klaassen, Y. Ritov, and J. A. Wellner. *Efficient and adaptive estimation for semi-parametric models*. Springer-Verlag, New York, 1998. Reprint of the 1993 original.
- A. Biswas, R. Bhattacharya, and E. Park. On a class of optimal covariate-adjusted response-adaptive designs for survival outcomes. *Statistical methods in medical research*, 2014.
- A. Chambaz and M. J. van der Laan. Targeting the optimal design in randomized clinical trials with binary outcomes and no covariate: Simulation study. *Int. J. Biostat.*, 7(1), 2011.
- A. Chambaz and M. J. van der Laan. Inference in targeted group sequential covariate-adjusted randomized clinical trials. *Scandinavian Journal of Statistics*, 41(1):104–140, 2013.
- Y. I. Chang and E. Park. Sequential estimation for covariate-adjusted response-adaptive designs. *J. Korean Statistical Society*, 42(1):105–116, 2013.

- L. Dümbgen, S. A. Van De Geer, M. C. Veraar, and J. A. Wellner. Nemirovski's inequalities revisited. *The American mathematical monthly: the official journal of the Mathematical Association of America*, 117(2):138, 2010.
- J. Friedman, T. Hastie, and R. Tibshirani. Regularization paths for generalized linear models via coordinate descent. *Journal of Statistical Software*, 2010.
- F. Hu and W. F. Rosenberger. *The theory of response-adaptive randomization in clinical trials*, volume 525. John Wiley & Sons, 2006.
- T. Huang, Z. Liu, and F. Hu. Longitudinal covariate-adjusted response-adaptive randomized designs. *J. Statistical Planning and Inference*, 143(10):1816–1827, 2013.
- C. Jennison and B. W. Turnbull. *Group Sequential Methods with Applications to Clinical Trials*. Chapman & Hall/CRC, Boca Raton, FL, 2000.
- J. Pearl. *Causality*. Cambridge University Press, Cambridge, 2000. Models, reasoning, and inference.
- R Core Team. *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria, 2014.
- W. F. Rosenberger. New directions in adaptive designs. *Statist. Sci.*, 11:137–149, 1996.
- W. F. Rosenberger, A. N. Vidyashankar, and D. K. Agarwal. Covariate-adjusted response-adaptive designs for binary response. *Journal of biopharmaceutical statistics*, 11(4):227–236, 2001.
- W. F. Rosenberger, O. Sverdlov, and F. Hu. Adaptive randomization for clinical trials. *J Biopharm Stat*, 22(4):719–36, 2012.
- J. Shao and X. Yu. Validity of tests under covariate-adaptive biased coin randomization and generalized linear models. *Biometrics*, 69(4):960–969, 2013.
- J. Shao, X. Yu, and B. Zhong. A theory for testing hypotheses under covariate-adaptive randomization. *Biometrika*, 97(2):347–360, 2010.
- O. Sverdlov, W. F. Rosenberger, and Y. Ryznik. Utility of covariate-adjusted response-adaptive randomization in survival trials. *Statistics in Biopharmaceutical Research*, 5(1):38–53, 2013.
- R. Tibshirani. Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society. Series B (Methodological)*, pages 267–288, 1996.
- M. J. van der Laan. The construction and analysis of adaptive group sequential designs. Technical report 232, Division of Biostatistics, University of California, Berkeley, March 2008.
- M. J. van der Laan and J. M. Robins. *Unified methods for censored longitudinal data and causality*. Springer Series in Statistics. Springer-Verlag, New York, 2003.
- M. J. van der Laan and D. Rubin. Targeted maximum likelihood learning. *Int. J. Biostat.*, 2(1), 2006.
- A. W. van der Vaart. *Asymptotic statistics*, volume 3 of *Cambridge Series in Statistical and Probabilistic Mathematics*. Cambridge University Press, Cambridge, 1998.
- R. van Handel. On the minimal penalty for markov order estimation. *Probability Theory and Related Fields*, 150:709–738, 2011.

L-X. Zhang and F-F. Hu. A new family of covariate-adjusted response-adaptive designs and their properties. *Appl. Math. J. Chinese Univ. Ser. B*, 24(1):1–13, 2009.

L-X. Zhang, F. Hu, S. H. Cheung, and W. S. Chan. Asymptotic properties of covariate-adjusted response-adaptive designs. *Ann. Statist.*, 35(3):1166–1182, 2007.