

Dynamic Treatment Regimes and “SMART” Design

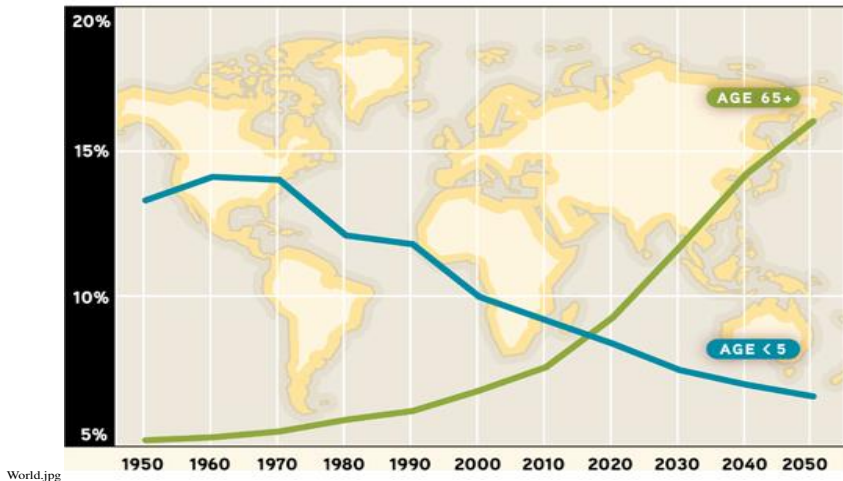
Bibhas Chakraborty

Duke-NUS Medical School and Department of Statistics & Applied Probability
National University of Singapore

bibhas.chakraborty@duke-nus.edu.sg

Institute for Mathematical Sciences, NUS
Singapore
February 15, 2019

The World is Aging!

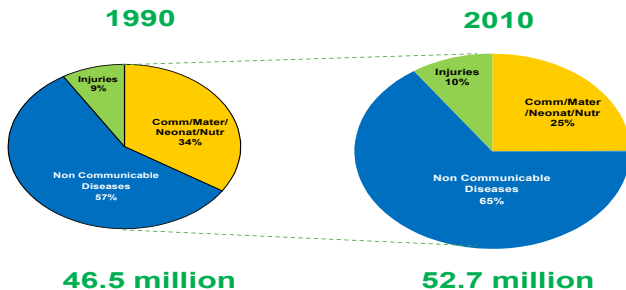


World.jpg

Source: U.S. Department of State Archive

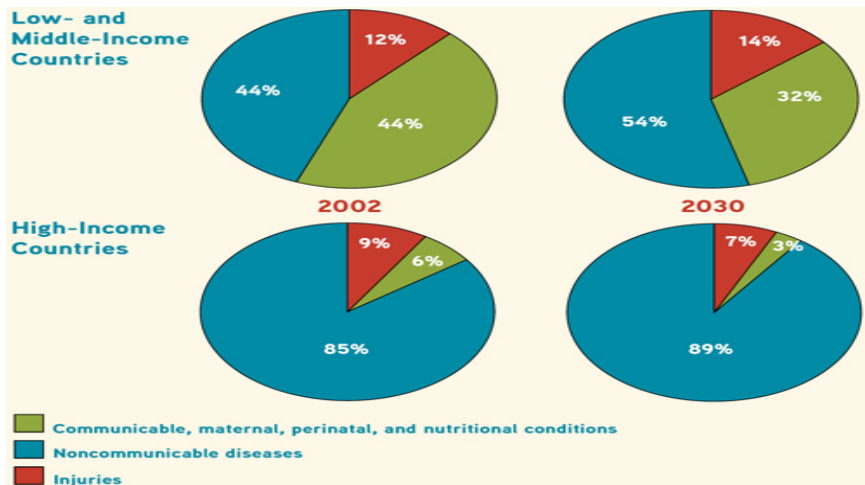
Global Shift in Causes of Death...

Towards **Chronic, Noncommunicable** Diseases



Source: Institute for Health Metrics and Evaluation, University of Washington

Similar Trend in Developing and Developed Countries



Source: U.S. Department of State Archive

Outline

- 1 Chronic Diseases, Personalized Medicine and Dynamic Treatment Regimes
- 2 SMART Design
 - Primary Analysis and Sample Size
- 3 SMART in Mobile Health (mHealth)
 - Non-Inferiority Testing
 - Adaptive SMART for Mobile Health
- 4 Estimation of Optimal DTRs
 - Q-learning
 - Analysis of Smoking Cessation Data: A Simple Case Study
- 5 Discussion

- 1 Chronic Diseases, Personalized Medicine and Dynamic Treatment Regimes
- 2 SMART Design
 - Primary Analysis and Sample Size
- 3 SMART in Mobile Health (mHealth)
 - Non-Inferiority Testing
 - Adaptive SMART for Mobile Health
- 4 Estimation of Optimal DTRs
 - Q-learning
 - Analysis of Smoking Cessation Data: A Simple Case Study
- 5 Discussion

Personalized Medicine for Chronic Diseases

Believed by many as the future of medicine ...



Source: <http://www.personalizedmedicine.com/>

Often refers to tailoring by **genetic** profile, but it's also common to personalize or stratify treatments to patients based on more “**macro**” level characteristics, some of which are **time-varying**

Personalized Medicine for Chronic Diseases

- Paradigm shift from “one size fits all” to individualized, patient-centric care
 - Can address inherent heterogeneity across patients
 - Can also address variability within patient, over time
 - Can increase patient compliance or adherence, thus increasing the chance of treatment success
 - Likely to reduce the overall cost of health care
- Overarching Methodological Questions:
 - How to decide on the optimal treatment for an individual patient?
 - How to make these treatment decisions evidence-based or data-driven?

Treatment Regimes or Regimens

- One perspective in personalized medicine: given an **individual patient's characteristics**, can we **identify a treatment**, among the available options, that is most likely to confer the most benefit?
- A **treatment regime** is a **rule** (or rules) that specifies **how to allocate treatment** based on an individual patient's characteristics
- If treatment is administered only once, then there is a **single** decision point, often the time of diagnosis
- In case of **chronic** diseases, treatments are typically administered at **multiple** decision points

Dynamic Treatment Regime(n)s

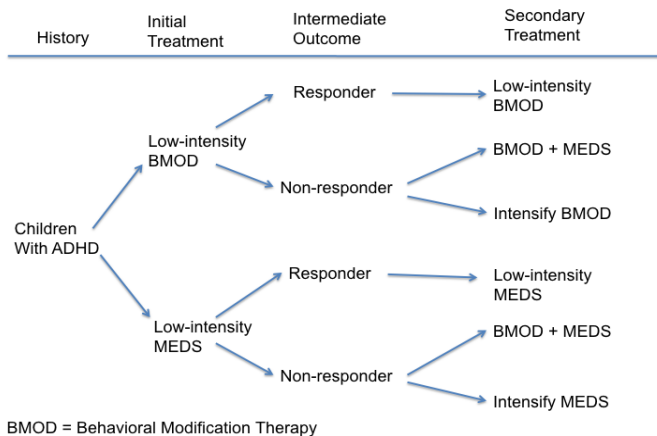
- A **dynamic treatment regime (DTR)** is a **sequence of decision rules**, one per decision point (or stage), that specify how to adapt the **type, dosage and timing of treatment** according to the ongoing information on an individual patient
 - Each decision rule takes a patient's treatment and covariate history as inputs, and outputs a recommended treatment
- Decision points can occur at regular intervals (e.g. yearly follow-up visits) or clinical events (e.g. remission, relapse, complication)

Dynamic Treatment Regimes (DTRs)

- DTRs offer a data-driven framework for **operationalizing** the adaptive clinical decision-making in a time-varying setting, and thereby potentially **improving** it
 - **Clinical decision support systems** for treating **chronic diseases**
- A subject's stage-specific treatment is not known at the start of a dynamic regime, since treatment depends on time-varying variables

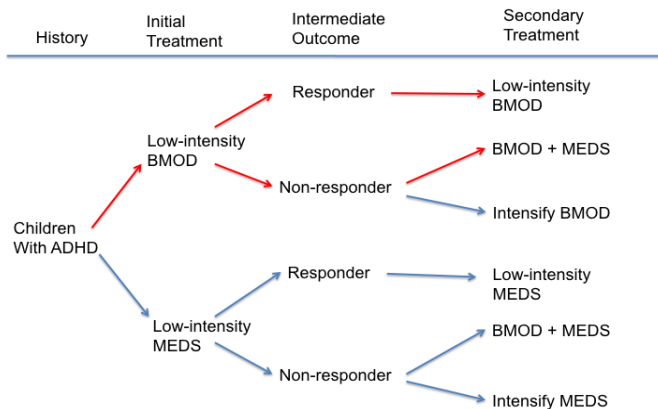
ADHD Example: Treatment Scenarios¹

ADHD: Attention Deficit Hyperactivity Disorder



¹Pelham WE and Fabiano GA (2008). Evidence-based psychosocial treatment for ADHD: An update. *Journal of Clinical Child and Adolescent Psychology*, 31, 184-214.

ADHD Example: One Simple DTR



“Give **Low-intensity BMOD** initially; if the subject **responds**, then continue **BMOD**, otherwise prescribe **BMOD + MEDS**”

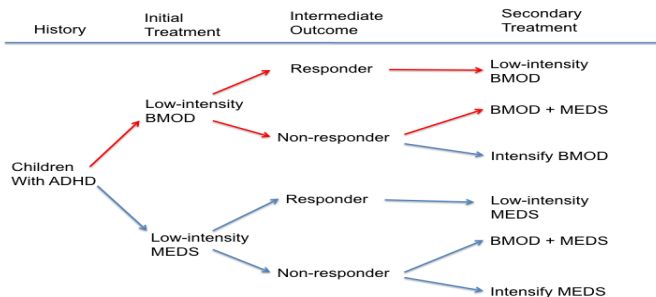
ADHD Example: One Not-so-simple DTR

- **Stage-1 Rule:** “If the **baseline level of impairment** is greater than a threshold (say, ψ), prescribe **MEDS**; otherwise prescribe **BMOD**”
- **Stage-2 Rule:** “If the subject is a **responder** to initial treatment, continue the same treatment; if **non-responder**, prescribe **BMOD + MEDS**”

How to specify ψ ?

Treatment Regime vs. Realized Treatment Experience

- Subjects following the **same DTR** can have **different realized treatment experiences**:
 - Subject 1 experiences “**Low-intensity BMOD**, followed by **response**, followed by **Low-intensity BMOD**”
 - Subject 2 experiences “**Low-intensity BMOD**, followed by **non-response**, followed by **BMOD + MEDS**”



Notation and Data Structure

It's a bit more than standard longitudinal data!

- K stages (or decision points) on a single patient:

$$O_1, A_1, \dots, O_K, A_K, O_{K+1}$$

O_j : Observation (pre-treatment) at the j -th stage

A_j : Treatment (action) at the j -th stage, $A_j \in \mathcal{A}_j$

H_j : History at the j -th stage, $H_j = \{O_1, A_1, \dots, O_{j-1}, A_{j-1}, O_j\}$

Y : Primary Outcome (assume larger is better, without loss of generality)

- A DTR is a sequence of decision rules:

$$d \equiv (d_1, \dots, d_K) \text{ with } d_j(h_j) \in \mathcal{A}_j$$

The Big Scientific Questions in DTR Research

- What would be the **mean outcome** if the population were to follow a particular pre-conceived DTR?
- How do the mean outcomes **compare** among two or more DTRs?
- What is the **optimal** DTR in terms of the mean outcome?
 - What individual information (**tailoring or prescriptive variables**) do we use to make these decisions?

The Big Statistical Questions

- ❶ What is the right kind of data for comparing two or more DTRs, or estimating optimal DTRs? What is the appropriate study design?
 - Sequential Multiple Assignment Randomized Trial (SMART)
- ❷ How can we compare pre-conceived, embedded DTRs?
 - primary analysis of SMART data
- ❸ How can we estimate the “optimal” DTR for a given patient?
 - secondary analysis of SMART data
 - e.g. Q-learning, a stagewise regression-based approach, originally developed in computer science

- 1 Chronic Diseases, Personalized Medicine and Dynamic Treatment Regimes
- 2 **SMART Design**
 - Primary Analysis and Sample Size
- 3 SMART in Mobile Health (mHealth)
 - Non-Inferiority Testing
 - Adaptive SMART for Mobile Health
- 4 Estimation of Optimal DTRs
 - Q-learning
 - Analysis of Smoking Cessation Data: A Simple Case Study
- 5 Discussion

DTR Discovery: Observational Data

- People extensively use observational longitudinal data (including **electronic health records**) – data collection is cheap!
- Usual concerns about observational data, e.g. **confounding** and other hidden **biases** (*Rubin, 1974; Rosenbaum, 1991*)
- Need **unverifiable assumptions** to make causal inference about treatment effects
- Analysis is more complex in general, and also in the DTR context (see, e.g. *Robins, 2004; Moodie, Chakraborty and Kramer, 2012*)

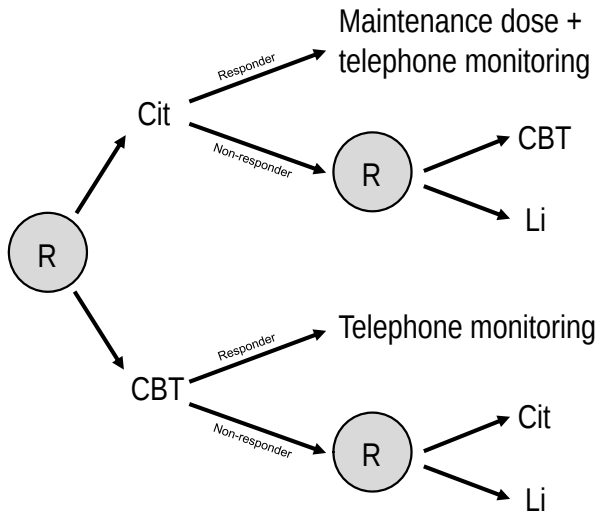
DTR Discovery: Experimental Data

- As is well-known, experimental data are generally better than observational data (of course, if you can afford it)
- Standard *randomized controlled trials* (RCTs) assess the effectiveness of a single dose-level of a single treatment, as compared to another
- Estimating the sequence of treatments that optimizes response in a longitudinal setting requires studying the elements in the sequence
- Can this be accomplished by a series of single-stage RCTs?

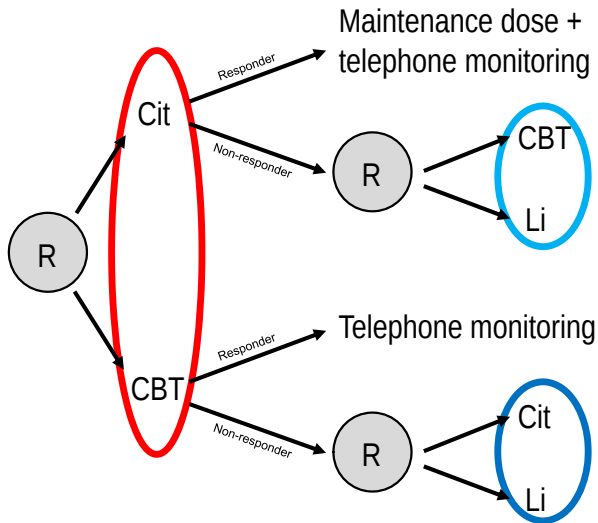
Example: Treating MDD

- Suppose we wish to compare both front-line and second-line treatment of *major depressive disorder* (MDD):
 - Front-line options: **citalopram (Cit)** or **cognitive behavioral therapy (CBT)**
 - Second-line options: treatment switch to **Cit**, **CBT**, or **Lithium (Li)**
 - All responders to first-line therapy will continue with maintenance and follow-up
- **Remember:** The goal is to find the **optimal treatment sequence**

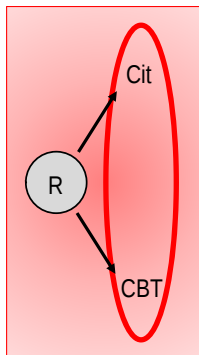
Example: Treating MDD



Example: Treating MDD

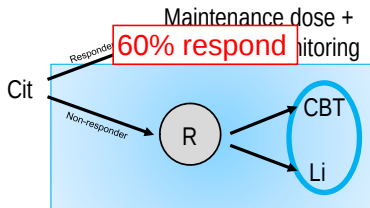


Front-line Treatment of MDD



- Suppose we observe 60% response with Cit, and only 50% with CBT
- Conclude: **Cit is the best front-line therapy**
- Now run another **one-stage trial** amongst Cit non-responders

Second-line Treatment of MDD

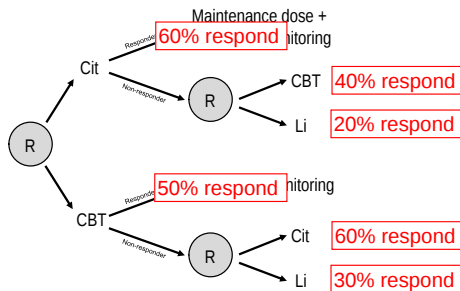


- We now observe 40% response to CBT and 20% to Li
- Conclude: **CBT is the best second-line therapy**
- Final treatment sequence: **“Cit, followed by CBT for non-responders”**

– Under this regimen, we expect to see 76% of patients respond

Delayed Effect

- What if initial treatment with CBT increases treatment adherence $\Rightarrow$ subsequent therapies more successful?



- Optimal DTR: “CBT, followed by Cit for non-responders”; 80% response expected

SMART!

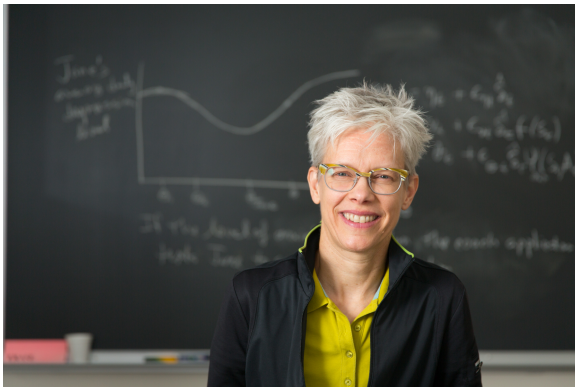
- Two single-stage trials would not have detected the best overall strategy for treatment
- Instead, we should have used a two-stage trial
- Such two-stage trials are known as **SMARTs** (*Lavori and Dawson, 2004; Murphy, 2005*), i.e.

Sequential Multiple Assignment Randomized Trials

SMART!

- **Susan Murphy** won the MacArthur Foundation “Genius Grant” in 2013 for inventing the SMART design:

<http://www.macfound.org/fellows/898/>



SMARTs, in general

- **Multi-stage** trials with a goal to inform the development of DTRs
- Same subjects participate **throughout** (they are followed through stages of treatment)
- Each stage corresponds to a treatment decision
- At each stage the patient is **randomized** to one of the available treatment options
- Treatment options at randomization may be **restricted** on ethical grounds, depending on intermediate outcome and/or treatment history

SMART: The Gains

- Ability to detect
 - delayed therapeutic effects (treatment interactions)
 - diagnostic effects
- More **generalizable** than standard RCT?
- Better **recruitment and retention** potential than standard RCT?
 - By virtue of the option to alter a non-functioning treatment while in the trial

SMART: The Costs

Unfortunately there is no free lunch!!

- More **expensive** than a single-stage RCT
 - Longer follow-up
 - Need more participants **if the trial is powered to compare whole sequences of treatments** – but there exist less expensive alternatives to power SMARTs!
- More **complex methods** for planning and analysis: requires an experienced statistician, or one with time to learn new methods
- May require **additional work** to get funded: still relatively new, unfamiliar

SMART vs. Crossover Trial Designs

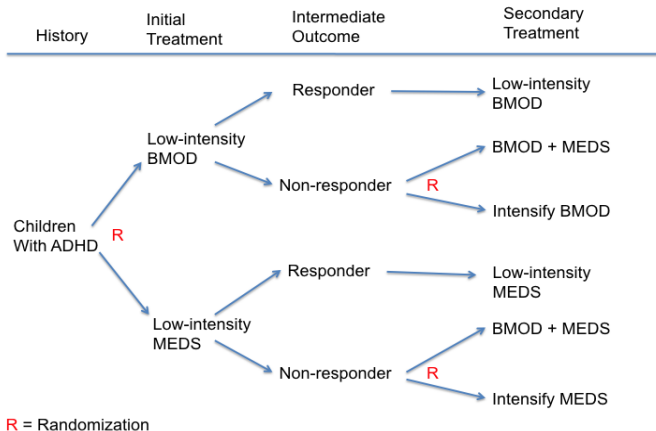
- Operationally, SMARTs look similar to **crossover trials**
- However, conceptually they are very different
 - Unlike SMART, Crossover trials aim to **evaluate stand-alone treatments**, not DTRs
 - Unlike a crossover trial, treatment **allocation in a SMART is typically adaptive to a subject's intermediate outcome**
 - Crossover trials attempt to “wash out” the “carry-over” effect while SMARTs attempt to **capture** it

SMART vs. Usual Adaptive Trial Designs

- SMARTs are different from usual **adaptive trials**
 - Within-subject vs. between-subject adaptation
- In an (outcome-) adaptive trial, the **randomization probabilities can change during the course of the trial** depending on the relative success of already-recruited patients on each of the treatments
 - Adaptive trials typically aim to reduce the number of trial participants exposed to the inferior treatment
 - Adaptive trials are most useful for examining questions where the outcome follows fairly quickly after treatment and/or the recruitment process is (relatively) slow
 - SMARTs are **fixed trial designs** – while randomization probabilities may depend on covariates, those probabilities remain constant over the course of the trial

Does anyone actually use SMARTs?

ADHD Trial (PI: Pelham; see *Lei et al.*, 2012, for design details)



Primary Outcome: Teacher-rated Impairment Rating Scale (TIRS)

Other Examples of SMART Studies

- **Schizophrenia:** CATIE (*Schneider et al., 2001*)
- **Depression:** STAR*D (*Rush et al., 2003*)
- **Prostate Cancer:** Trials at MD Anderson Cancer Center (e.g., *Thall et al., 2000*)
- **Leukemia:** CALGB Protocol 8923 (e.g., *Stone et al., 1995; Wahed and Tsiatis, 2004*)
- **Smoking Cessation:** Project Quit (*Strecher et al., 2008; Chakraborty et al., 2010*)
- **Alcohol Dependence:** *Oslin et al.* (see, e.g., *Lei et al., 2012*)
- Many recent examples available at the Methodology Center, PennState University website:

<http://methodology.psu.edu/ra/smart/projects>

SMART Design Principles

Primary and Secondary Hypotheses

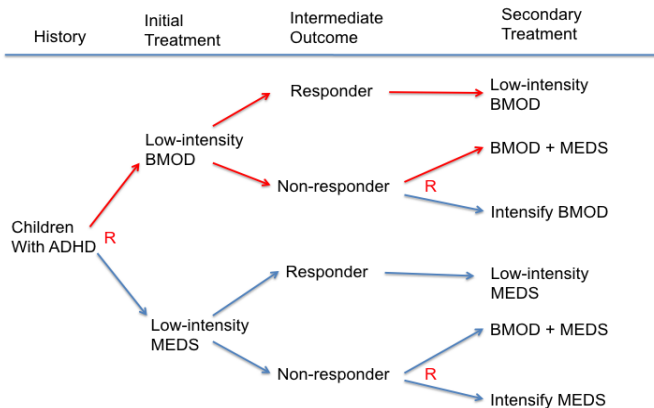
- Choose scientifically important primary hypotheses that also aid in developing DTRs
 - Power trial to address these hypotheses
- Choose secondary hypotheses that further develop the DTR, and use randomization to eliminate confounding
 - Trial is **not necessarily powered** to address these hypotheses
 - Still better than **post hoc observational analyses**
 - Underpowered randomizations can be viewed as **pilot studies** for future full-blown comparisons
- From a funding perspective, it may be more feasible to design a standard two-arm trial as primary goal, and add the secondary randomizations and DTR estimation as a secondary goal

SMART Design: Tailoring

- At each stage, restrict the class of treatments only by ethical, feasibility or strong scientific considerations
- Use a low-dimension summary (e.g. **responder status**) instead of all intermediate outcomes to determine randomization and restrict class of next treatments
 - This will often be a key **tailoring variable**
 - The definition must be concrete, e.g., determined using specific measure(s) at a specific time
 - Generally, this variable should be regularly available in clinical practice, so that tailoring with this variable is feasible
 - In mental illness studies, feasibility considerations may force investigators to use **preference** in this low dimensional summary (e.g. STAR*D)
- Collect intermediate outcomes potentially useful in further tailoring of treatment at the analysis stage

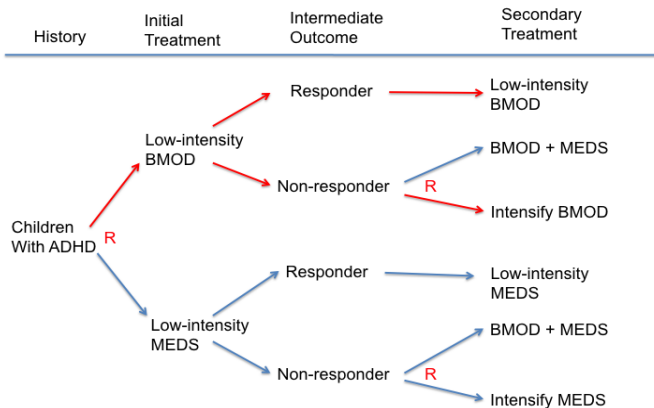
SMART Design: Embedded DTRs

Embedded DTR-1 (d_1)



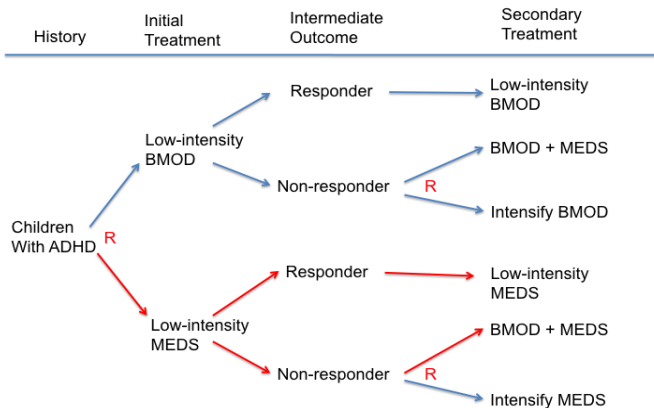
SMART Design: Embedded DTRs

Embedded DTR-2 (d_2)



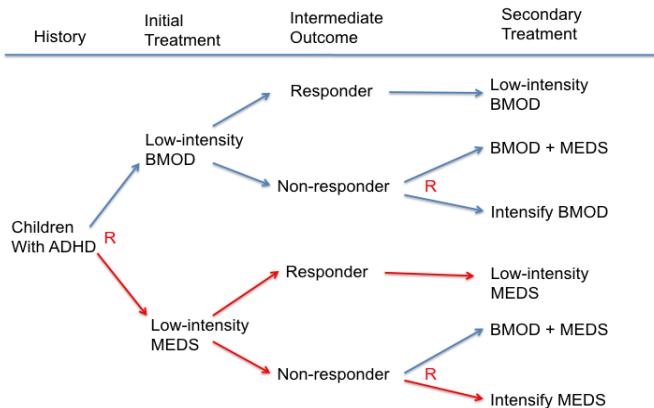
SMART Design: Embedded DTRs

Embedded DTR-3 (d_3)



SMART Design: Embedded DTRs

Embedded DTR-4 (d_4)



SMART Design: Simplified Data Structure

- $(O_{1i}, A_{1i}, O_{2i}, A_{2i}, Y_i)$ for $i = 1, \dots, N$, where
 - $O_1 = \emptyset$
 - $A_1 \in \{BMOD, MEDS\}$
 - $O_2 = R$ is the response indicator, where $R = 1$ denotes responder and $R = 0$ denotes non-responder
 - $A_2 \in \{BMOD, BMOD^+, MEDS, MEDS^+, BMOD + MEDS\}$
 - Y is the **continuous** end-of-trial primary outcome
 - N is the total sample size of the trial

SMART Design: Embedded DTRs (Formally)

- The 4 embedded DTRs in the above SMART are:

$$d_1 = \left(BMOD, BMOD^R (BMOD + MEDS)^{1-R} \right)$$

$$d_2 = \left(BMOD, BMOD^R (BMOD^+)^{1-R} \right)$$

$$d_3 = \left(MEDS, MEDS^R (BMOD + MEDS)^{1-R} \right)$$

$$d_4 = \left(MEDS, MEDS^R (MEDS^+)^{1-R} \right)$$

- Key Questions:

- How can we measure the performance of an embedded DTR?
- How can we compare two embedded DTRs in the SMART?

Regime Mean or Value Function: The Key Estimand

- What would be the mean outcome if the entire population follows a DTR d ? $\Rightarrow$ Call this the “Value” of d and denote by $\mu_d = E_d(Y)$
- Recall the general data structure, K stages (or decision points) on a single patient: $O_1, A_1, \dots, O_K, A_K, O_{K+1}$
- If $\pi_j(a_j|h_j)$ ’s are used to allocate the treatments, then the data likelihood (or, joint distribution) P_π is:

$$f_1(o_1)\pi_1(a_1|o_1) \prod_{j=2}^K f_j(o_j|h_{j-1}, a_{j-1})\pi_j(a_j|h_j)f_{K+1}(o_{K+1}|h_K, a_K)$$

- What is the likelihood of the data under an arbitrary set of deterministic decision rules $d_j(h_j)$ ’s for allocating treatments? Call it P_d :

$$f_1(o_1)\mathbb{I}\{a_1 = d_1(o_1)\} \prod_{j=2}^K f_j(o_j|h_{j-1}, a_{j-1})\mathbb{I}\{a_j = d_j(h_j)\}f_{K+1}(o_{K+1}|h_K, a_K)$$

Regime Mean or Value Function: The Key Estimand

- The notation $\mu_d = E_d(Y)$ actually means expectation with respect to the distribution P_d even though the data were actually generated by P_π
- Computing $E_d(Y)$ non-parametrically involves a **change of probability measure**, under the assumption that P_d is **absolutely continuous** or **feasible** with respect to P_π
 - Any data trajectory that can result by implementing d must also have positive probability of occurring under the generative distribution π
- Under the feasibility assumption, $E_d(Y) = \int Y dP_d = \int Y \left(\frac{dP_d}{dP_\pi} \right) dP_\pi$ where $\frac{dP_d}{dP_\pi}$ is a version of the **Radon-Nikodym derivative** and is given by the ratio of the two likelihoods:

$$\frac{dP_d}{dP_\pi} = \prod_{j=1}^K \frac{\mathbb{I}\{a_j = d_j(h_j)\}}{\pi_j(a_j|h_j)}$$

Regime Mean or Value Function: The Key Estimand

- Embedded regimes are, by design, feasible
- For the embedded regime $d_1 = \left(BMOD, BMOD^R (BMOD + MEDS)^{1-R} \right)$, the value function is

$$\begin{aligned}\mu_{d_1} &= E_{d_1}(Y) \\ &= E \left[\frac{\mathbb{I}\{A_1 = BMOD, A_2 = BMOD^R (BMOD + MEDS)^{1-R}\}}{\pi_1(A_1)\pi_2(A_2|A_1, R)} Y \right] \\ &= E(W^{d_1} Y), \text{ say}\end{aligned}$$

- π_1 and π_2 are the randomization probabilities
 - Inverse probability weighting (IPW) approach
 - Same underlying idea as in the classical **Horvitz-Thompson estimator**
- Similarly, one can write expressions for other regimes d_2, d_3, d_4

Estimation of Regime Mean

- The regime mean can be estimated by the IPW estimator (*Robins et al., 2000*):

$$\bar{Y}_{d_1} = \sum_{i=1}^N W_i^{d_1} Y_i / \sum_{i=1}^N W_i^{d_1}$$

- Also, for large samples (*Murphy, 2005*):

$$\widehat{Var}(\bar{Y}_{d_1}) = \frac{1}{N^2} \sum_{i=1}^N (W_i^{d_1})^2 (Y_i - \bar{Y}_{d_1})^2$$

Primary Analysis

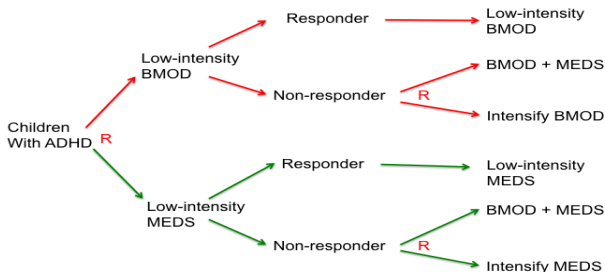
- Depending on the research question, it could be a comparison of two or more groups corresponding to two or more DTRs embedded in the SMART, or components thereof
- Standard methods of analysis, involving the above means and variances
- Standard software like SAS PROC GENMOD can be employed²

²Nahum-Shani I, Qian M, Almira D, Pelham W, Gnagy B, Fabiano G, et al. Experimental Design and Primary Data Analysis Methods for Comparing Adaptive Interventions. *Psychological Methods*. 2012;17:457-477

Primary Hypothesis and Sample Size: Scenario 1

Hypothesize that averaging over the secondary treatments, the initial treatment **BMOD** is as good as the initial treatment **MEDS**

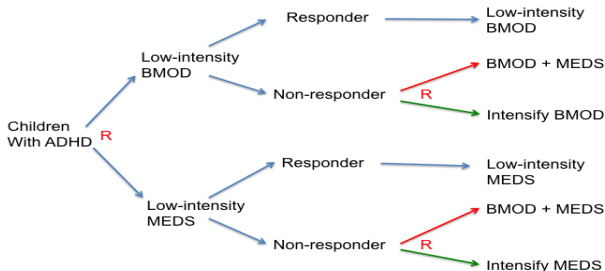
- *Sample size formula is same as that for a two-group comparison*



Primary Hypothesis and Sample Size: Scenario 2

Hypothesize that among non-responders an **augmentation (BMOD+MEDS)** is as good as an **intensification of treatment**

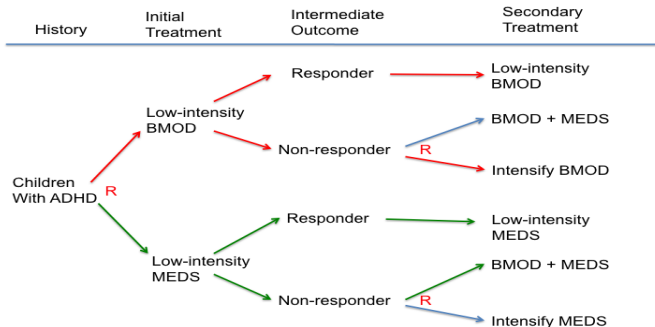
- *Sample size formula is same as that for a two-group comparison of non-responders* (overall sample size depends on the presumed **non-response rate**)



Primary Hypothesis and Sample Size: Scenario 3

Hypothesize that the “red” DTR (d_2) is as good as the “green” DTR (d_3)

- *Sample size formula involves a two-group comparison of “weighted” means (overall sample size depends on the presumed non-response rate)*



Sample Size Requirements

Assume **continuous** outcome, e.g., Teacher-rated Impairment Rating Scale in case of ADHD

Key Design Parameters:

Effect Size = $\frac{\Delta\mu}{\sigma}$ (Cohen's d)

Type I Error Rate = $\alpha = 0.05$

Desired Power = $1 - \beta = 0.8$

Initial Response Rate = $\gamma = 0.5$

Trial Size:

Effect Size	Scenario 1	Scenario 2	Scenario 3
0.3	$N_1 = 350$	$N_2 = \frac{N_1}{(1-\gamma)} = 700$	$N_3 = N_1 \times (2 - \gamma) = 525$
0.5	$N_1 = 128$	$N_2 = \frac{N_1}{(1-\gamma)} = 256$	$N_3 = N_1 \times (2 - \gamma) = 192$
0.8	$N_1 = 52$	$N_2 = \frac{N_1}{(1-\gamma)} = 104$	$N_3 = N_1 \times (2 - \gamma) = 78$

SMART Design: Other Comparisons

Comparing embedded DTR-1 (d_1) with embedded DTR-2 (d_2)

- d_1 and d_2 share the initial treatment of BMOD
- Patients who respond to BMOD are consistent with both d_1 and d_2 , and contribute information on the overall response to both regimens
- Specialized methods are needed to account for “re-using” their information

SMART Design: Other Comparisons

- Alternatively, we may wish to:
compare d_1 vs. d_2 vs. d_3 vs. d_4 to see which of the four regimes leads to the best expected outcome
- The above question requires more specialized methods to account for both **multiple comparisons** as well as the **re-use of data** for patients who are consistent with more than one regime

How about Comparing DTRs with Time-to-event Outcomes?

- Same principles of **inverse probability weighting** apply (e.g., weighted Kaplan-Meier estimator, weighted log-rank test, etc.)
- Under the **proportional hazards** assumption, Li and Murphy³ provided a conservative sample size formula using **weighted log-rank** test for comparing d_1 and d_3 :

$$n \leq \left\{ \frac{1}{\pi_1 \pi_2} + \frac{1}{(1 - \pi_1) \pi_2} \right\} \frac{(z_{1-\frac{\alpha}{2}} + z_{1-\beta})^2}{\xi^2 \cdot P_{d_1}(\text{event})}$$

where ξ is the **log hazard ratio** between the times-to-event under the two DTRs, and $P_{d_1}(\text{event})$ denotes the probability of observing an event before the end of the study among subjects following d_1 (**event rate**)

³Li Z and Murphy SA (2011). Sample size formulae for two-stage randomized trials with survival outcomes. *Biometrika*, 98(3): 503-518.

How about Comparing DTRs with Time-to-event Outcomes?

- Li and Murphy's formula builds on sample size formula based on log rank test for classical survival analysis⁴
- Li and Murphy also developed an alternative formula based on weighted Kaplan-Meier estimator, but that requires more inputs at the design stage
- The log rank test based formula has been implemented in the online sample size calculator (web app):

<http://methodologymedia.psu.edu/logranktest/samplesize>

⁴Schoenfeld DA (1981). The asymptotic properties of nonparametric tests for comparing survival distributions. *Biometrika*, 68: 316-319.

Outline

- 1 Chronic Diseases, Personalized Medicine and Dynamic Treatment Regimes
- 2 SMART Design
 - Primary Analysis and Sample Size
- 3 SMART in Mobile Health (mHealth)**
 - Non-Inferiority Testing
 - Adaptive SMART for Mobile Health
- 4 Estimation of Optimal DTRs
 - Q-learning
 - Analysis of Smoking Cessation Data: A Simple Case Study
- 5 Discussion

Life is Getting Digital...

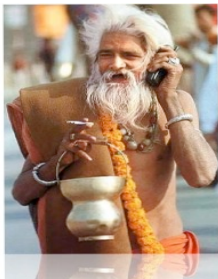


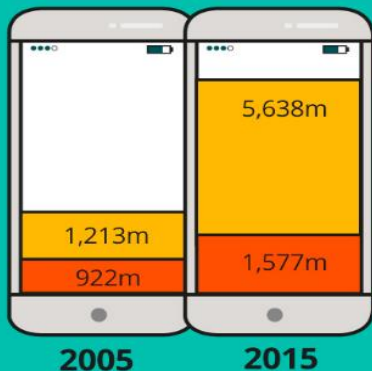
Image courtesy of Gary Bennett

Source: Duke Global Health Institute

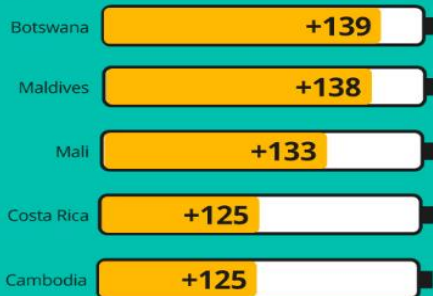
Mobile Phone Subscriptions, Globally

Mobile phone subscriptions

● Developed ● Developing



Highest increases in subscriptions per 100 inhabitants (2005–2015)



LONDON
SCHOOL of
HYGIENE
& TROPICAL
MEDICINE



Source: International Telecommunications Union (2016)

Countries with Shortage of Healthcare Providers

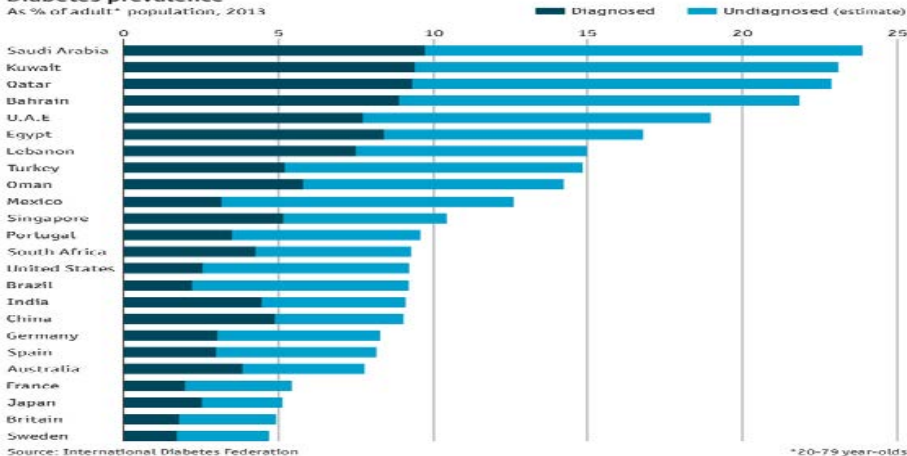


Source: Duke Global Health Institute

The Global Epidemic of Diabetes

Diabetes prevalence

As % of adult* population, 2013



Source: International Diabetes Federation

Source: Duke Global Health Institute

mHealth Interventions for Diabetes Management

- Good glycaemic control is key to managing diabetes and its many complications
- Treatments for Type 2 Diabetes Mellitus (T2DM) include diet, exercise, oral medications and insulin injections
- Initiation of insulin in insulin-naïve patients traditionally involves multiple visits to a physician to adjust insulin dose which can be costly and time-consuming
 - Delay in scheduling can result delay in achieving glycaemic control
- **Self-titration for insulin dose calculation** is also an option but many patients hesitate and express uncertainty

The “Diabetes Pal” App

- A smartphone app, developed in Singapore (Duke-NUS and SGH), with an in-built algorithm to compute daily insulin dosage based on lowest of the previous 3 fasting blood glucose (FBG) readings (patients are required to measure their FBG daily and input into the app)

DIABETES PAL

Phone Application



The “Diabetes Pal” App

The **feasibility** of the app to deliver the insulin titration algorithm in insulin-naïve patients has been validated in a pilot RCT ($n = 66$):

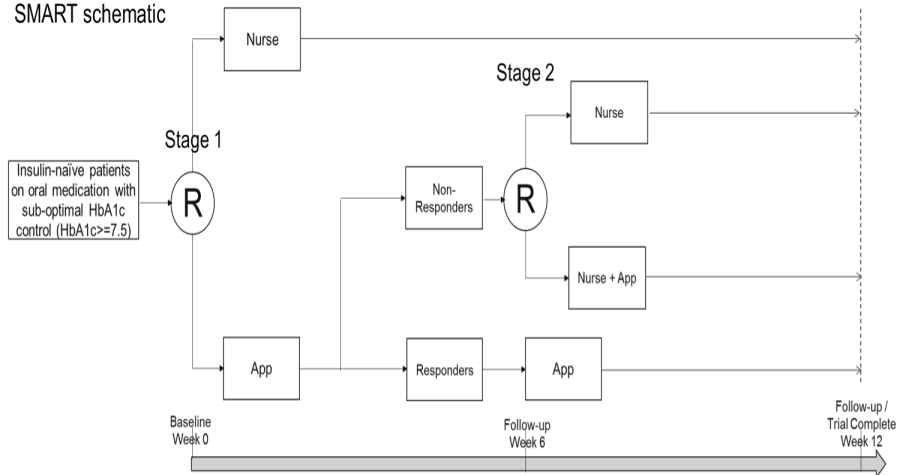
Bee YM, et al. (2016). A smartphone application to deliver a treat-to-target insulin titration algorithm in insulin-naïve patients with type 2 diabetes: A pilot randomized controlled trial. Diabetes Care, 39(10): e174-e176.

The “Diabetes Pal” App: What’s Next?

- Not everyone in the population is good with technology (e.g., app); some may need human intervention (e.g., nurse)
 - But human intervention **adds more cost** to the healthcare system
- Perhaps a **stepped-care approach** involving the app may be cost-effective
 - Start with a low-cost, low-intensity intervention (app) first; and then **step up to a high-cost, high-intensity intervention (e.g., nurse) only for patients who show early signs of non-response** to low-cost, low-intensity intervention
- It would be interesting to study if a “stepped-care regime” is almost as good as (**non-inferior** to) a “more aggressive or resource-intensive regime” – if yes, then the stepped-care regime should be selected
- This sort of research question can be investigated using a **SMART design**

A SMART Design involving The Diabetes Pal App

SMART schematic



Non-Inferiority Testing Framework

- Embedded regimes are:

$$d_1 = (App, App^R Nurse^{1-R}), d_2 = (App, App^R (App + Nurse)^{1-R}) \text{ and} \\ d_3 = (Nurse, Nurse^R Nurse^{1-R})$$

- Here d_1 and d_2 are stepped-care regimes, whereas d_3 is the active control regime
- In non-inferiority testing⁵, generally the goal is to show that the **efficacy** of the experimental treatment (regime), when compared to an active control treatment (regime), is not **below** the pre-specified **non-inferiority margin**
- In practice, a **non-inferiority margin** is the maximum **clinically acceptable difference** from the active control on average that researchers agree to accept in exchange of the secondary benefits (e.g. cost, burden, side-effects) of the new treatment

⁵D'Agostino RB, Massaro JM, Sullivan LM (2003). Non-inferiority trials: design concepts and issues-the encounters of academic consultants in statistics. *Statistics in Medicine*, 22(2): 169-86

Non-Inferiority Testing in SMARTs

- For two regimes d_1 and d_3 with corresponding regime means μ_{d_1} and μ_{d_3} respectively, the hypothesis for the non-inferiority test is

$$H_0 : \mu_{d_1} \leq \mu_{d_3} - \theta \text{ vs. } H_1 : \mu_{d_1} > \mu_{d_3} - \theta,$$

where θ is a pre-specified non-inferiority margin ($\theta > 0$)

- This hypothesis tests that the **average efficacy** of regime d_1 is **not inferior** to that of the regime d_3 , with the non-inferiority margin θ
- The **choice** of θ depends on both **statistical** reasoning and **clinical** judgment
- The unscaled test statistic is $\bar{Y}_{d_1} - \bar{Y}_{d_3}$, with mean $\mu_{d_1} - \mu_{d_3}$ and variance $\nu_{d_1 d_3}$ (obtained from Murphy's calculations)

Test Statistic and Sample Size Calculation

- Under H_0 , the **large-sample distribution** of the test statistic

$$Z_{d_1 d_3} = \frac{(\bar{Y}_{d_1} - \bar{Y}_{d_3}) - (\mu_{d_1} - \mu_{d_3})}{\sqrt{\text{var}(\bar{Y}_{d_1} - \bar{Y}_{d_3})}} = \frac{\bar{Y}_{d_1} - \bar{Y}_{d_3} + \theta}{\sqrt{\nu_{d_1 d_3}}} \rightarrow N(0, 1)$$

- Reject H_0 at level α and **conclude non-inferiority** if $Z_{d_1 d_3} > z_\alpha$ (one-sided)
- Deriving sample size formula requires large sample approximations and some assumptions
- Given a postulated effect size $\delta = \mu_{d_3} - \mu_{d_1}$, the required n depends on the **difference** between θ and δ , rather than their **individual values**

“Realistic” Synthetic Data Analysis

- Data simulated from the pilot study sample, with Y taken as $-\text{HbA1c}$
- Response rate to app = 51% and response rate to nurse = 69% (according to Deloitte Global Mobile Consumer Survey UK Edition)
- $\bar{Y}_{d_3} = -8.107$ with s.d. = 0.086; $\bar{Y}_{d_1} = -8.269$ with s.d. = 0.491
- For $\theta = 0.5$, test statistic = $0.6781 < z_\alpha = 1.645$, so non-inferiority can't be concluded
- For $\theta = 1$, test statistic = $1.6811 > z_\alpha = 1.645$, so non-inferiority can be concluded
- Choice of θ is tricky!

From SMART to SMART-AR

- Technologies evolve very quickly, so interventions using mobile devices **must be evaluated and disseminated very quickly**; otherwise they will lose relevance
- This is a very different setting from traditional RCTs (or even traditional SMARTs)
 - SMARTs can be made more useful by incorporating **adaptive randomization**
 - In modern contexts like **mHealth**, SMART with **Adaptive Randomization (SMART-AR)**⁶ has been developed
 - SMART-AR is also **ethically appealing** even in a non-mHealth context
 - In general, how best to do this is not known yet

⁶Cheung YK, Chakraborty B, and Davidson K (2015). Sequential multiple assignment randomized trial (SMART) with adaptive randomization for quality improvement in depression treatment program. *Biometrics*, 71(2): 450-459.

Outline

- 1 Chronic Diseases, Personalized Medicine and Dynamic Treatment Regimes
- 2 SMART Design
 - Primary Analysis and Sample Size
- 3 SMART in Mobile Health (mHealth)
 - Non-Inferiority Testing
 - Adaptive SMART for Mobile Health
- 4 **Estimation of Optimal DTRs**
 - Q-learning
 - Analysis of Smoking Cessation Data: A Simple Case Study
- 5 Discussion

Secondary Analysis of SMART Data

- Goal: To find **optimal** treatment sequence for each individual patient by **deeply tailoring** on their time-varying covariates and intermediate outcomes
 - This is to develop the **evidence-based decision support system** for clinicians (they can apply these regimens while treating future patients)
- Methodologically challenging – custom-made analytic tools necessary
 - One popular approach is **Q-learning**, a stage-wise regression-based method
 - We developed an **R package** called **qLearn** (*Xin et al., 2012*) that conducts Q-learning (Freely available at CRAN):

<http://cran.r-project.org/web/packages/qLearn/>

Q-learning: Introduction

- Q-learning (*Watkins, 1989*)
 - A popular method from Reinforcement (Machine) Learning
 - A generalization of least squares regression to multistage decision problems (*Murphy, 2005*)
 - Adapted to and implemented in the DTR context with several variations (*Zhao et al., 2009; Chakraborty et al., 2010; Schulte et al., 2012; Song et al., 2014*)
 - Relatively easy to understand and implement, and forms a good basis for more complex approaches
- The intuition comes from dynamic programming (*Bellman, 1957*) in case the multivariate distribution of the data is known
 - Q-learning is an approximate dynamic programming approach

Notation and Data Structure (Recap)

K stages (or decision points) on a single patient:

$$O_1, A_1, \dots, O_K, A_K, O_{K+1}$$

O_j : Observation (pre-treatment) at the j -th stage

A_j : Treatment (action) at the j -th stage, $A_j \in \{-1, 1\}$

H_j : History at the j -th stage, $H_j = \{O_1, A_1, \dots, O_{j-1}, A_{j-1}, O_j\}$

Y : Primary Outcome (assume larger is better, without loss of generality)

A DTR is a sequence of decision rules:

$$d \equiv (d_1, \dots, d_K) \text{ with } d_j(h_j) \in \mathcal{A}_j$$

Dynamic Programming (DP): The Background

Two Stages

- Move backward in time to take care of the **delayed effects**
- Define the “Quality of treatment”, **Q-functions**:

$$\begin{aligned}Q_2(h_2, a_2) &= \mathbb{E}\left[Y \mid H_2 = h_2, A_2 = a_2\right] \\Q_1(h_1, a_1) &= \mathbb{E}\left[\underbrace{\max_{a_2} Q_2(H_2, a_2)}_{\text{delayed effect}} \mid H_1 = h_1, A_1 = a_1\right]\end{aligned}$$

- Optimal DTR:

$$d_j(h_j) = \arg \max_{a_j} Q_j(h_j, a_j), \quad j = 1, 2$$

What if the true Q-functions are unknown?

From DP to Q-learning

- DP requires modeling the joint multivariate distribution (likelihood) of time-varying covariates and outcome
- This is hard!
 - The knowledge needed to model the joint distribution is often unavailable
 - Model mis-specification can lead to incorrect conclusions
- Turn to semi-parametric methods
- In particular, Q-learning estimates the true Q-functions from data using regression models

Q-learning: Typical Implementation ($K = 2$)

- Linear regression models for Q-functions:

$$Q_j(H_j, A_j; \beta_j, \psi_j) = \beta_j^T H_j + (\psi_j^T H_j) A_j, j = 1, 2,$$

- At stage 2, regress Y on $(H_2, H_2 A_2)$ to obtain $(\hat{\beta}_2, \hat{\psi}_2)$
- Construct stage-1 Pseudo-outcome:

$$\tilde{Y}_{1i} = \max_{a_2} Q_2(H_{2i}, a_2; \hat{\beta}_2, \hat{\psi}_2), i = 1, \dots, n$$

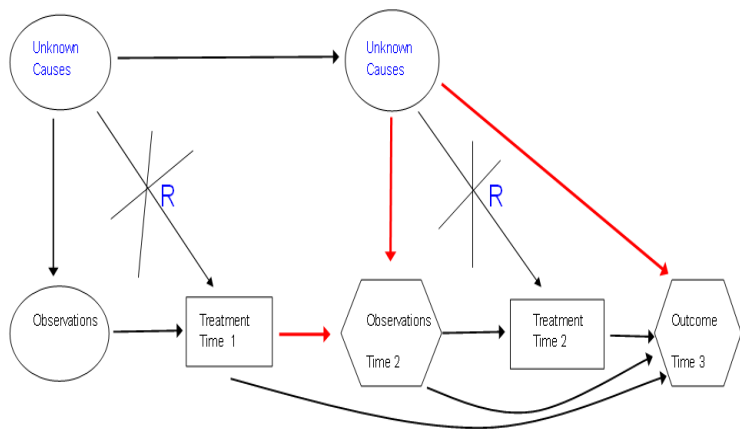
- At stage 1, regress $\tilde{Y}_1$ on $(H_1, H_1 A_1)$ to obtain $(\hat{\beta}_1, \hat{\psi}_1)$
- Estimated Optimal DTR:

$$\hat{d}_j(h_j) = \arg \max_{a_j} Q_j(h_j, a_j; \hat{\beta}_j, \hat{\psi}_j) = \text{sign}(\hat{\psi}_j^T h_j)$$

Q-learning: Remarks

- Q-learning is appealing because it is easy to implement in standard software, and easy to explain to clinical collaborators who are familiar with regression
- The approach has several limitations, e.g.:
 - Not robust to model mis-specification
 - Only limited results are available for discrete outcomes
- More sophisticated approaches (e.g., A-learning or outcome-weighted learning) exist, at least for some treatment/outcome types

Why move through stages as in Q-learning? Why not run an “all-at-once” multivariable regression?



Berkson's Paradox or Collider-stratification Bias: There may be **non-causal** association(s) even with randomized data, leading to biased stage-1 effects (*Berkson, 1946; Greenland, 2003; Murphy, 2005*)

Project Quit: A Smoking Cessation Trial (Simplified)

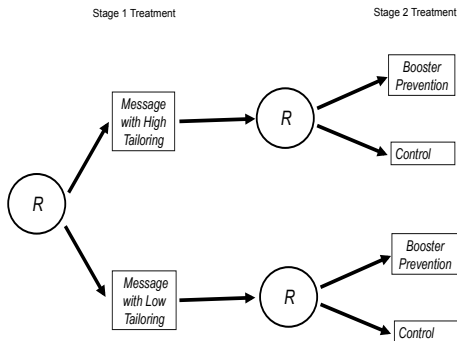
Two-stage Web-based (*eHealth*) Behavioral Intervention Trial for Smoking Cessation conducted at the University of Michigan⁷

- Stage-1 Covariate : education ($\leq$ high school vs. $>$ high school)
- Stage-1 Intervention : **tailoring of success story, low vs. high**
(in addition to free nicotine patch)
- Stage-2 Covariates : quit status at 6 months (1 = quit, 0 = not quit),
months non-smoking over 6 months
- Stage-2 Intervention : **booster prevention vs. control**
- Primary Outcome : months non-smoking over 12 months

Center for Health Communications Research, University of Michigan

⁷Strecher et al. (2008). Web-based smoking cessation components and tailoring depth: Results of a randomized trial. *American Journal of Preventive Medicine*, 34(5): 373-381.

SMART Design Schematic (Simplified)



Secondary Research Questions

- **Stage-1 question:** (In future) How should a web-based smoking cessation intervention be designed so as to maximize each individual's chance of quitting over the two stages? Should this intervention be **adapted** to the smoker's baseline education?
- **Stage-2 question:** Should the stage-2 intervention be **adapted** to either the stage-1 intervention the smoker has already received and/or the smoker's intermediate outcome (e.g., stage-1 quit status)?

Results from Q-learning Analysis

No significant stage-2 treatment effect ($n = 479$)

Stage-1 Analysis Summary ($n = 1848$)

Variable	Coefficient	95% Bootstrap CI
education	0.01	(-0.18, 0.20)
high vs. low tailoring	0.07	(-0.01, 0.11)
tailoring:education	-0.11	(-0.24, -0.00)*

- The “highly individually tailored” level of story appears more effective for subjects with less education ($\leq$ high school)
- This finding is consistent with that of *Strecher et al.* (2008) – a logistic regression analysis of the stage-1 quit status (binary) data

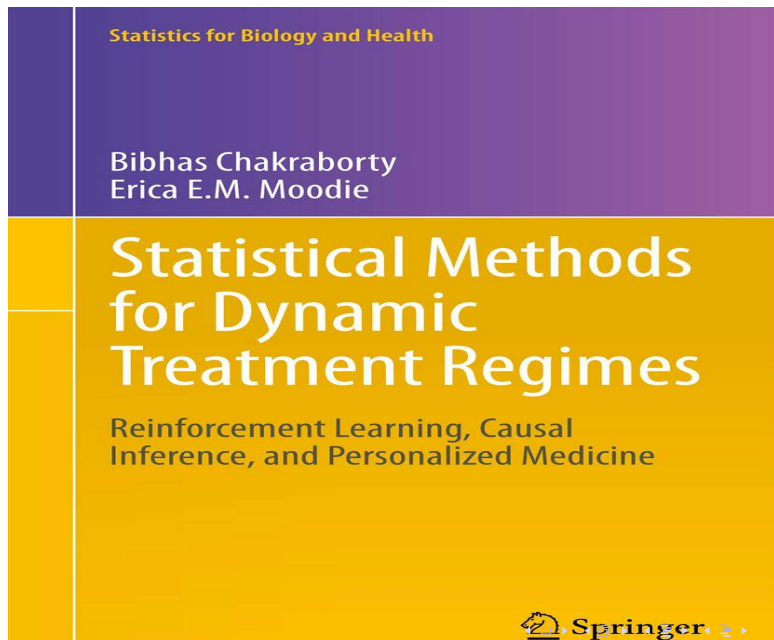
Outline

- 1 Chronic Diseases, Personalized Medicine and Dynamic Treatment Regimes
- 2 SMART Design
 - Primary Analysis and Sample Size
- 3 SMART in Mobile Health (mHealth)
 - Non-Inferiority Testing
 - Adaptive SMART for Mobile Health
- 4 Estimation of Optimal DTRs
 - Q-learning
 - Analysis of Smoking Cessation Data: A Simple Case Study
- 5 Discussion

Take-home Messages

- SMART is a cutting-edge trial design that formalizes the sequential decision making in clinical practice
 - Very relevant for chronic conditions
 - Does not necessarily require a lot of sample size (depends on the primary question)
- SMARTs are useful in the modern context of mHealth
- Non-inferiority testing in SMART is very appealing from a **cost-effectiveness** perspective
- Analysis methods exist for continuous, binary and time-to-event data coming from a SMART
- Relevant softwares are also available
- Lots of open problems, thus many opportunities for statisticians!

Shameless Self-Promotion!!



- Shoot your questions, comments, criticisms, request for slides to:
bibhas.chakraborty@duke-nus.edu.sg

