

Implementing the treatment policy estimation for continuous variable with multiple imputation

Haolin Sun, Shanghai Xianxiang Medical Technology Co., Ltd;

ABBSTRACT

Many of the early approaches to treatment policy estimation continued the existing practices for ITT analysis. Commonly in longitudinal data, a standard Mixed Model Repeated Measures (MMRM) analysis was performed using all available data, following the common misconception that a treatment policy strategy meant that IEs are ignored. This approach may be reasonable with a negligible amount of missing data, but as missing data increases, the alignment between the target estimand and the estimation approach diverges for two reasons: Firstly, because standard MMRM cannot distinguish between data before and after an IE, and secondly because missing data are typically highly correlated with the occurrence of IEs. Combined, these can create a strong violation of the missing at random (MAR) assumption required for the validity of inference from the standard MMRM model.

Reference-based imputation (RBI) approaches were commonly advocated, particularly Jump-to-Reference (J2R) methods for treatment discontinuation in placebo-controlled trials with symptomatic treatments and Copy-Increment-from-Reference (CIR) methods in trials with disease-modifying treatments. These account for the occurrence of IEs when imputing missing data and respectively assume either an instant, complete loss of treatment effect, or a retention of treatment effects accrued until the time of the IE. Another set of approaches are the Retrieved Dropout (RD) class. In these, missing data are multiply imputed conditional upon whether they occur pre- or post-IE and/or pattern of the IE.

We described the difference between standard, retrieved dropout and reference-based multiple imputation of continuous longitudinal outcomes for statistical methodology. Software (GSK five macro, %mistep and MI procedure) and example sas code was showed how to do the multiple imputation. Also a simulation was done to evaluate the performance for 6 scenarios.

INTRODUCTION

A substantial percentage of the measurements of the outcome or outcomes of interest is often missing in the RCT trial. This “missingness” reduces the benefit provided by the randomization and introduces potential biases in the comparison of the treatment groups. Multiple imputation is a statistical technique for analyzing incomplete data sets. MI replaces each missing value with a set of plausible values that represent the uncertainty about the true value (valid under MAR). The ICH E9(R1) addendum distinguishes between IEs and missing data (ICH E9 working group (2019)). The addendum proposes that estimation methods to address the problem presented by missing data should be selected to align with the estimand. Missing data may occur in subjects without an IE or prior to the occurrence of an IE. As such missing outcomes are not associated with an IE, it is often plausible to impute them under a missing-at-random (MAR) assumption using a standard MMRM imputation model of the longitudinal outcomes. Informally, MAR occurs if the missing data can be fully accounted for by the baseline variables included in the model and the observed longitudinal outcomes, and if the model is correctly specified.

However, in some situations, it may be more reasonable to assume that missingness is “informative” and indicates a systematically better or worse outcome than in observed subjects. In such situations, MNAR imputation with a δ -adjustment could be explored as a sensitivity analysis. δ -adjustments add a fixed or random quantity to the imputations in order to make the imputed outcomes systematically worse or better than those observed.

The conventional MI methods based on Bayesian (or approximate Bayesian) posterior draws of model parameters combined with Rubin’s rules to make inferences as described in Carpenter, Roger, and Kenward (2013) and Cro et al. (2020).

Reference-based imputation methods formalize the idea to impute missing data in the intervention group based on data from a control or reference group. For a general description and review of reference-based imputation methods, we refer to Carpenter, Roger, and Kenward (2013), Cro et al. (2020), I. White, Royes, and Best (2020) and Wolbers et al. (2022).

STATISTICAL METHODOLOGY

ESTIMAND FOR ONE CASE STUDY

To illustrate different estimate approaches, the following example is used. For different approaches, the estimand we target is: “The difference in mean change from baseline HAMDTL17 depression score at week 26 for depressive disorder patients assigned to either experimental treatment or placebo including any subsequent effects of discontinuation and initiation of rescue medication upon discontinuation.” The attributes of the ESTIMAND are as follows:

Table 1. Attributes of the ESTIMAND

Five attributes	Description
Target population	depressive disorder patients met the inclusion and exclusion criteria
Treatment conditions	Assignment to treatment or control(placebo), including any subsequent discontinuation of assigned treatment and use of rescue medication.
Variable/endpoint	The mean change from baseline HAMDTL17 depression score at week 26
Population-level summary	The between-group comparison is performed using the difference in the mean change.
Other Intercurrent events(Handling Strategy)	Discontinuation from assigned treatment condition which may also include the use of rescue medication (Treatment Policy).

Notation

We define Y_{kij} as HAMDTL17 depression score for subject i at visit j assigned to group k , and $\Delta Y_{kij} = Y_{kij} - Y_{ki0}$ as the corresponding change from baseline. We also define two group $k \in \{C, T\}$ where C and T correspond to being assigned control and treatment group, respectively.

We consider a single IE and define D_{kij} to be an indicator variable of IE status, where $D_{kij} = 1$ indicating subject i from group k has experienced an IE prior to visit j and $D_{kij} = 0$ otherwise. α, β, δ for regression coefficients and ϵ for error for the regression model.

For the single IE considered, every patient belongs to one of $J + 1$ IE occurrence patterns: Either they do not experience an IE, or their first visit impacted by the IE is visit 1, $\dots$, J , respectively. We define P_{kij} as a pattern indicator, taking a value of 1 if visit j is the subject's first visit impacted by the IE, and a value of zero otherwise.

ESTIMATION FOR THE STUDY

We consider three general approaches to estimate the estimand defined above in presence of partial post-IE missing data. 1. Conventional MI which is based on “MAR”. The “MAR” assumption effectively says that the distribution of values after withdrawal for those who withdraw, conditional upon their observed values and on their baseline covariates, is the same as for those who do not leave and who have complete data. This is done computationally using Restricted Maximum-Likelihood (REML) or equivalently by using multiple imputation under the MAR assumption. In this context it is often called an MMRM model (Mixed Model for Repeated Measures) and fitted using mixed-models software such as proc MIXED. 2. Reference-based MI where a prespecified reference distribution is used for imputation of post-IE missing data (in our case the placebo group will provide the reference for both groups). J2R(Jump-to-Reference), CIR(Copy-Increments-to-Reference) and CR(Copy-Reference) was illustrated. 3. Retrieved Dropout MI which extended MAR assumptions, they assume that missing outcome data is similar to observed data

from subjects in the same treatment group with the same observed outcome history, and either the same IE status or the pattern.

General consideration for the MI

In general, the MI is conducted through 3 steps. 1. The missing data are filled in m times to generate m complete data sets. 2. The m complete data sets are analyzed by using standard procedures. 3. The results from the m complete data sets are combined for the inference.

Monotone vs Non-Monotone Missing Patterns

The following code to check the missing patterns for the dataset.

```
proc mi data = tp_wide nimpute = 0;
  var basval change_1 - change_4;
run;
```

Group	basval	change_1	change_2	change_3	change_4	Freq	Percent
1	X	X	X	X	X	151	87.79
2	X	X	X	X	.	10	5.81
3	X	X	X	.	.	5	2.91
4	X	X	.	.	.	6	3.49

Group	basval	change_1	change_2	change_3	change_4	Freq	Percent
1	X	X	X	X	X	150	87.21
2	X	X	X	X	.	10	5.81
3	X	X	X	.	.	5	2.91
4	X	X	.	.	.	6	3.49
5	X	.	X	X	X	1	0.58

Figure 1. Missing Data Pattern.

We can see the missing patterns for the dataset, the left is monotone patterns while the right is Non-Monotone missing. For the intermediate the missing data are interspersed among full data values, while subjects who have only terminal missing values and no intermediate missing values are often referred to as having monotone missing values. A monotone missing pattern is easier to work with and more flexible, but this pattern often suggests a MNAR mechanism. Generally, first simplification is to assume that all intermediate missing values will be imputed assuming MAR. So the MNAR part of the model is restricted to patterns that are monotone. This simplifies the specification of the pattern-mixture models, and will often be a sensible assumption. The MAR assumption effectively says that the distribution of values after withdrawal for those who withdraw, conditional upon their observed values and on their baseline covariates, is the same as for those who do not leave and who have complete data. 2 methods can be used to impute the intermediate missing data into the monotone missing data for continuous variable.

- MCMC method

The same prior distributions as in the SAS implementation of the “five macros” are used, i.e. an improper flat priors for the regression coefficients and a weakly informative inverse Wishart prior for the covariance matrix (or matrices). The Part1B in the “five macros” is to fit parameter estimation model using the MCMC procedure and draws a pseudo-independent sample from the joint posterior distribution for the linear predictor parameters and the covariance parameters.

Alternatively, the MCMC statement in the MI procedure will create the similar results. Suppose we want to create 500 MI repetitions, the following sas code can be used.

```
proc mi data = tp_wide nimpute = 500 out = mi_monotone;
  var trt basval;
  MCMC chain = single nbiter = 200 niter = 100 impute = monotone;
run;
```

- FCS regression

For each pattern of missingness, the parameter was firstly estimated through the observed value. And then based on the multivariate normal distribution assumption with both the observed value and the covariates, the missing value was imputed from the distribution. Finally, the observed value and imputed value was combined to create the complete value for the given visit. The next round imputation start off based on the previous complete value, current visit observed value and covariates to estimate the parameters, and then sampling from the multivariate normal distribution and finally, the observed value

and imputed value was combined to create the complete value for the given visit. The imputation procedure was illustrated as in figure 2.

$$\begin{aligned}
 \theta_1^{(0)} &\sim P(\theta_1 | Y_{1(obs)}) \\
 Y_{1(*)}^{(0)} &\sim P(Y_1 | \theta_1^{(0)}) \\
 Y_1^{(0)} &= (Y_{1(obs)}, Y_{1(*)}^{(0)}) \\
 &\dots \\
 &\dots \\
 \theta_p^{(0)} &\sim P(\theta_p | Y_1^{(0)}, \dots, Y_{p-1}^{(0)}, Y_{p(obs)}) \\
 Y_{p(*)}^{(0)} &\sim P(Y_p | \theta_p^{(0)}) \\
 Y_p^{(0)} &= (Y_{p(obs)}, Y_{p(*)}^{(0)})
 \end{aligned}$$

Figure 2. Imputation process with FCS regression

The example sas code for the FCS regression is as below.

```

proc mi data = tp_wide out = midi_mi seed = 1999 nimpute = 5
  class sex trt;
  fcs reg(midi7 midi14 midi28 midi42 midi98);
  var sex trt age base midi7 midi14 midi28 midi42 midi98;
run;

```

Generally, the FCS method requires fewer iterations than the MCMC method. Often, as few as five or 10 iterations are enough to produce satisfactory results;

MCAR, MAR or MNAR

Rubin proposed a framework for handling missing data and described three different mechanisms by which data can be missing: Missing Completely at Random (MCAR), Missing at Random (MAR), and Missing Not at Random (MNAR), which was illustrated in table 1.

Table 2. Missing Data Mechanisms

Missing Mechanisms	Illustration
MCAR	✓ Missingness independent of unobserved and observed data ✓ Strong assumption ✓ Eg: planned missing ✓ Removing cases with missing data does not bias inference ✓ Valid methods: complete case, available case
MAR	✓ Conditional on observed data, missingness independent of unobserved data ✓ Probability of missingness depends only on obs. data ✓ Eg. men, more likely to answer negatively, also more likely to decline to answer. Male participants is more likely to drop out compared to female subjects. ✓ Assuming that data missingness was random after accounting for related variables, analyzing the available data may extend our knowledge about observed data to the missing data. ✓ Valid methods: MMRM, MI
MNAR	✓ Non-ignorable missing data or informative missing data ✓ Missingness depends on unobserved data ✓ Value of unobserved variable itself predicts missingness ✓ Eg. income/weight surveys ✓ the missing value itself influences the probability of missingness, or some unmeasured factor is linked to both the value of the missing variable and the probability of missingness. ✓ Valid methods: Pattern-mixture models, Tipping point

Simple Models

The simple approaches make no distinction between outcomes with respect to the IE and have no covariate terms relating to the IE in any models. They rely on a missing-at-random (MAR) assumption, where missing outcomes depend only upon observed and modelled covariates as well as previous outcomes. By not modelling the IE, they assume that any missing outcomes would have been similar to the observed data from similar subjects in the same treatment group.

$$\begin{aligned}\theta_2^{(*)} &\sim P(\theta_2 | Y_{1(obs)}, Y_{2(obs)}) \\ Y_2^{(*)} &\sim P(Y_2 | \theta_2^{(*)}) \\ &\dots \\ &\dots \\ \theta_p^{(*)} &\sim P(\theta_p | Y_{1(obs)}, \dots, Y_{p(obs)}) \\ Y_p^{(*)} &\sim P(Y_p | \theta_p^{(*)})\end{aligned}$$

Figure 3. Monotone Methods for Data Sets with Monotone Missing Patterns

Formular

This approach uses sequential imputation of missing outcomes at visit $j(j=1, \dots, J)$, respectively, and is based on a Bayesian linear regression model depending on the treatment group and past outcomes:

$$Y_{kij} = \alpha_{kj} + \beta_{j0}Y_{ki0} + \beta_{j1}Y_{ki1} + \dots + \beta_{j(j-1)}Y_{ki(j-1)} + \varepsilon_{kij}, \text{ where } \varepsilon_{kij} | Y_{ki0}, \dots, Y_{ki(j-1)} \sim N(0, \sigma_j^2)$$

Code

```
proc mi data = tp_wide nimpute = 500 seed = 1 out = miout1 noprint;
  class trt;
  var basval trt change_1 change_2 change_3 change_4;
  monotone regression (change_1 = trt basval);
  monotone regression (change_2 = change_1 trt basval);
  monotone regression (change_3 = change_1 change_2 trt basval);
  monotone regression (change_4 = change_1 change_2 change_3 trt basval);
run;
```

Alternatively, the %MISTEP by James Roger can be achieved the similar results as

```
%mistep(In=tp_wide, Out=miout1, Response=change_1, Class=trt, Model= trt basval,
Nimpute=500, seed=12345);
%mistep(In=miout1, Out=miout2, Response=change_2, Class=trt, Model= trt basval
change_1, seed=12321);
%mistep(In=miout2, Out=miout3, Response=change_3, Class=trt, Model= trt basval
change_1 change_2, Nimpute=500, seed=12322);
%mistep(In=miout3, Out=miout4, Response=change_4, Class=trt, Model= trt basval
change_1 change_2 change_3, seed=12323);
```

DATA

Figure 4 and Figure 5 presents the dataset before the MI procedure and after the MI procedure.

	PATIENT	trt	basval	lastvis	lastobs	_NAME_	change_1	change_2	change_3	change_4
1	1503	1	32	4	4	change	-11	-12	-13	-15
2	1507	0	14	4	4	change	-3	0	-5	-9
3	1509	1	21	4	4	change	-1	-3	-5	-8
4	1511	0	21	4	4	change	-5	-3	-3	-9
5	1513	1	19	1	4	change	5	3.157	-0.147	0.127
6	1514	0	21	1	1	change	2	.	.	.
7	1516	0	23	4	4	change	-8	-2	-3	-8
8	1517	1	19	1	4	change	5	3.465	-0.017	0.327

Figure 4. Dataset before the MI procedure

	Imputation	PATIENT	trt	basval	lastvis	lastobs	NAME	change_1	change_2	change_3	change_4
1	1	1503	1	32	4	4	change	-11	-12	-13	-15
2	1	1507	0	14	4	4	change	-3	0	-5	-9
3	1	1509	1	21	4	4	change	-1	-3	-5	-8
4	1	1511	0	21	4	4	change	-5	-3	-3	-9
5	1	1513	1	19	1	4	change	5	3.157	-0.147	0.127
6	1	1514	0	21	1	1	change	2	-2.636578557	-5.661778227	-3.807301327
7	1	1516	0	23	4	4	change	-8	-2	-3	-8
8	1	1517	1	19	1	4	change	5	3.465	-0.017	0.327

Figure 5. Dataset after the MI procedure

Generally, the simulation shows If the “MAR” is hold, the MMRM and MI approaches create almost exactly match results.

```
proc mixed data = dt;
  class visit trt subjid;
  model change = visit*trt visit*basval/noint ddfm = kr;
  repeated visit/patient=subjid type=un;
  lsmeans visit*trt;
run;
```

Retrieved dropout approach

For simplicity, we create figure 6 to illustrate the retrieved dropout approach.

1. For Pattern 1 the subject completes the study treatments and the primary endpoint assessments (on-treatment completers).
2. For Pattern 2 the subject completes the study treatments but miss the primary endpoint assessments (on-treatment in-completers).
3. For Pattern 3 the subject experience intercurrent events, resulting in discontinuing or changing study treatments, but the subjects stay in the study and complete their primary endpoint assessments (retrieved dropouts (RDs)).
4. For Pattern 4 the subject experience intercurrent events, resulting in discontinuing or changing study treatment, and the subjects drop out from study and miss their primary endpoint assessments (non-retrieved dropouts (non-RDs)).

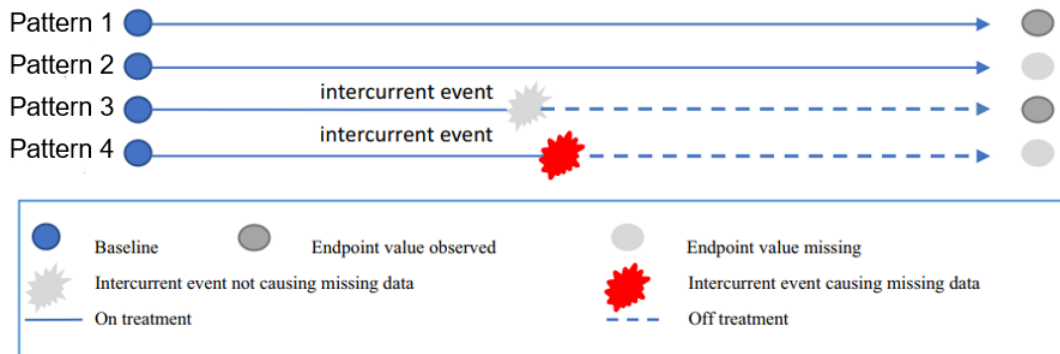


Figure 6. Patterns of observed and missing data

There should be very few subjects in Pattern 2. For such subjects, it may be reasonable to assume data are missing at random (MAR) and observed data from Pattern 1 can be used to impute missing data. The special interest is for Pattern 4, that is the missing data occurring after and likely caused by intercurrent events. These data are very likely missing not at random. With enough RDs, missing data in non-RDs could be addressed, to the best of our knowledge, as how the measurement would have been, had it been measured. Missing data in an endpoint for non-RDs (Patterns 4, Figure 6) can be imputed based on the observed data from the RDs in the same treatment arm (i.e., Patterns 3 in the same treatment group).

Formular

This model allows a two-level distinction between patient outcomes as pre- or post-IE at each visit. This is achieved by including indicators D_{kij} for IE status at each visit per treatment group, resulting in the following model. In addition, rather than regressing on the observed earlier outcomes as done in simple model, the sequential imputation model regresses on the residuals calculated using predicted values from earlier imputation models. Using the residuals in this way is similar to the residual conditioning in the MMRM approaches and has been shown to reduce bias and variance inflation.

$$Y_{kij} = \alpha_{kj} + \delta_{kj} D_{kij} + \beta_{j0} Y_{ki0} + \beta_{j1} (Y_{ki1} - \hat{Y}_{ki1}) + \dots + \beta_{j(j-1)} (Y_{ki(j-1)} - \hat{Y}_{ki(j-1)}) + \varepsilon_{kij}, \text{ where} \\ \varepsilon_{kij} | D_{ki1}, \dots, D_{kij}, Y_{ki0}, \dots, Y_{ki(j-1)} \sim N(0, \sigma_j^2).$$

The term $(Y_{kil} - \hat{Y}_{kil})$ define the residuals using the predicted value $\hat{Y}_{kil}$ from the previous $\ell - 1$ visits.

Code

The code below shows how to do the multiple imputation for visit 2 based on the previous residuals. ROC MI and PROC MIXED sequentially to obtain the residual after each imputation.

1. The first step is to do the MI for visit 2 based on the imputation results from visit 1, variable `ontreat_2=0,1` indicate an IE not occurred and occurred at the visit.

```
proc mi data=change1 nimpute=500 seed=1 out=change2 noprint;
  class trt ontreat_2;
  var trt ontreat_2 basval resid1 change_2; /*resid1 is the resid1 from
  visit 1*/
  monotone regression (change_2=trt ontreat_2 trt*ontreat_2 basval
  resid1/details);
run;
```

2. The second step is to do the prediction based on the imputation from visit 1

```
proc mixed data= change1;
  class trt ontreat_2;
  model change_2=trt ontreat_2 trt*ontreat_2 basval resid1
  outp=pred_change2(rename=(pred=predicted2) keep=patient pred) s
  ddfm=kr;
run;
```

3. To calculate the residual per the imputed and the predicted from the mixed model to facilitate the imputation in visit 3.

```
proc sql;
  create table resid2 as
  select * from change2, pred_change2
  where change2.patient=pred_change2.patient;
quit;
data resid2;
  set resid2;
  resid2=change_2-predicted2;
run;
```

DATA

We can see from Figure 7, the IE occurs at visit 2 for patient 1513 but the primary endpoint assessments were available at the following visit, so this is retrieved dropouts. For patient 1514 the primary endpoint assessments as missing for visit 3 and visit 4 after the IE (IE also occurs at visit 2), so this is non-retrieved dropouts for this patient. Suppose both patients come from the same treatment group, the primary endpoint assessments from patient 1513 will be used to impute patient 1514.

	_Imputati	PATIENT	lastvis	change_1	change_2	change_3	change_4	ontreat_1	ontreat_2	ontreat_3	ontreat_4	predicted	resid1	predicted2	resid2
1	1	1503	4	-11	-12	-13	-15	1	1	1	1	-6	-5	-13	1.2
2	1	1507	4	-3	0	-5	-9	1	1	1	1	-1	-2	-4	3.6
3	1	1509	4	-1	-3	-5	-8	1	1	1	1	-2	1	-4	.96
4	1	1511	4	-5	-3	-3	-9	1	1	1	1	-3	-2	-6	2.8
5	1	1513	1	5	3.157	-0.147	0.127	1	0	0	0	-2	7	3.5	-.3
6	1	1514	1	2	3.4933	.	.	1	0	0	0	-3	5	.89	2.6
7	1	1516	4	-8	-2	-3	-8	1	1	1	1	-3	-5	-9	6.5
8	1	1517	1	5	3.465	-0.017	0.327	1	0	0	0	-2	7	3.5	.01

Figure 7. Data structure before start imputation for visit 3

Additionally, the %MISTEP by James Roger can be achieved the similar results as

```
%mistep(data=tp_wide,out=tp1a,nimpute=500,response=change_1,Class=trt
ontreat_1,Model=trt ontreat_1 trt*ontreat_1
basval,suffix=1, seed=145);

%mistep(data=tp1a,out= tp2a,response=change_2,Class=trt ontreat_2,
Model=residual1 trt ontreat_2 trt*ontreat_2 basval,seed=121);

%mistep(data=tp2a,out= tp3a,response=change_3,Class=trt ontreat_3,
Model=residual1 residual2 trt ontreat_3 trt*ontreat_3 basval,seed=126);

%mistep(data=tp3a,out=imputed,response=change_4,Class=trt ontreat_4,
Model=residual1 residual2 residual3 trt ontreat_4 trt* ontreat_4
basval,seed=111);
```

Extension of the RD approach

Furthermore, except add IE indicator into the model, we can add time dependent covariate that indicates the current status of patient with respect to their IE pattern. This is achieved by including a set of indicators $P_{ki1}, \dots, P_{kij}$ at visit j . As we consider the case where IE occurrence is monotone, this corresponds to the pattern of IE occurrence and results in the following model:

$$Y_{kij} = \alpha_{kj} + \delta_{kj} P_{kij} + \beta_{j0} Y_{ki0} + \beta_{j1} (Y_{ki1} - \hat{Y}_{ki1}) + \dots + \beta_{j(j-1)} (Y_{ki(j-1)} - \hat{Y}_{ki(j-1)}) + \varepsilon_{kij}, \text{ where} \\ \varepsilon_{kij} | P_{ki1}, \dots, P_{kij}, Y_{ki0} \dots Y_{ki(j-1)} \sim N(0, \sigma_j^2)$$

The IE occurs at visit 1 for patient 105 and visit 3 for patient 114 respectively. The missing data occurs at visit 2 and afterward for patient 105 and occurs at visit 5 for patient 114. The detailed data can be found in figure 8.

	sim_run	subjid	visitrn	group	response	groupn	change	disctime	rdrate
11521	1	105	1	trt	7.879393	1	-0.53253619	1	6
11522	1	105	2	trt	.	1	.	1	6
11523	1	105	3	trt	.	1	.	1	6
11524	1	105	4	trt	.	1	.	1	6
11525	1	105	5	trt	.	1	.	1	6
11566	1	114	1	trt	7.785229	1	-0.348443823	3	6
11567	1	114	2	trt	7.340857	1	-0.792815823	3	6
11568	1	114	3	trt	8.318454	1	0.1847811773	3	6
11569	1	114	4	trt	8.605899	1	0.4722261773	3	6
11570	1	114	5	trt	.	1	.	3	6

Figure 8. Example data structure before start imputation for IE indicator and IE pattern

For IE indicator imputation, the non-missing visit is coded as 6 while all the rest is coded as 12345. There are totally only 2 kinds of results to indicate whether the IE occurs till the current visit. The variables for IE indicators are disc1 to disc5 as in figure 9. The coding for the IE pattern and IE indicator is a little bit different. Take visit 4 as an example. The pattern variable pat4 is coded as 123 for patient 105 and coded as 4 even both of them are after the IE occurrence. That's because both of them occurs after the IE but the IE occurred at different visit before them. the impact to the underlying true value at visit 4 is different.

VIEWTABLE: Data.Dar_drop_data WHERE(pat4 in (12, 45, 123))

	sim_run	subjid	disctime	diso1	diso2	diso3	diso4	diso5	pat1	pat2	pat3	pat4	pat5
1305	1	105	1	6	12345	12345	12345	12345	6	123	123	123	123
1314	1	114	3	6	6	6	12345	12345	6	6	6	4	4

Figure 9. Coding difference for between the IE indicator and IE pattern

1. The first step is to do the MI for visit 4 based on the imputation results from visit 3, variable pat3 indicate an IE not occurred and occurred at the visit.

```
proc mi data=change3 nimpute=500 seed=1 out=change4 noprint;
  class trt pat3;
  var trt pat3 basval resid1 resid2 change_3;
  monotone regression (change_3=trt pat3 trt*pat3 basval resid1 resid2
  /details);
run;
```

2. The second step is to do the prediction based on the imputation from visit 1

```
proc mixed data= change3;
  class trt pat3;
  model change_3=trt pat3 trt*pat3 basval resid1 resid2
  outp=pred_change3(rename=(pred=predicted3) keep=patient pred) s
  ddfm=kr;
run;
```

3. To calculate the residual per the imputed and the predicted from the mixed model to facilitate the imputation in visit 3.

```
proc sql;
  create table resid3 as
  select * from change3, pred_change3
  where change3.patient=pred_change3.patient;
quit;

data resid3;
  set resid3;
  resid3=change_3-predicted3;
run;
```

Reference based approach

In brief, reference-based approach is implemented in four steps. First, a MMRM imputation model is fitted to the observed outcome data. Second, conditional mean imputation of missing data is performed based on the parameter estimates from the MMRM imputation model and the chosen imputation method. Third, the completed data is analyzed using a simple ANCOVA model. Finally, these analysis results is combined according to Rubin's rules to make inferences. Difference between 3 kind of reference-based imputation can be found in figure 10.

It generally Includes model estimation phase and imputation phase. The purpose of the model estimation is to estimate the mean trajectories and covariance matrices for each treatment group in the absence of IEs. That means records occur after the IE will be taken as missing from both groups. In any case, all observed post-IE data will be included as is in the subsequent imputation and analysis steps.

Imputation phase is also known as conditional mean imputation step. In order to impute missing data for a subject, the marginal imputation distribution is first defined. For MAR imputation, the marginal imputation distribution for subject i is a multivariate normal distribution with mean $\tilde{\mu}_i = X_i\hat{\beta}$ and covariance $\tilde{\Sigma}_i = \hat{\Sigma}$. For subjects randomized to the control group, the same marginal imputation distribution as for MAR imputation is also used for reference-based imputation methods. For a subject i randomized to the intervention group for whom a reference-based imputation method is used after visit $\tilde{t}_i < J$, denote the predicted mean based on the assigned treatment as $\mu_i = X_i\hat{\beta}$. Denote the corresponding design matrix

assuming that (counter to fact) the subject had been randomized to the control group (but using the subject's observed baseline characteristics and time-varying covariates) by $X_{i,ref}$ and denote the corresponding predicted mean from the imputation model by $\mu_{ref,i} = X_{i,ref}\hat{\beta}$. Then, the mean of the marginal imputation distribution depends on the chosen reference-based imputation is as below.

1. Copy reference (CR): $\tilde{\mu}_i = \mu_{ref,i}$
2. Jump to reference (J2R): $\tilde{\mu}_i = (\mu_i[1], \dots, \mu_i[\tilde{t}_i], \mu_{ref,i}[\tilde{t}_i + 1], \dots, \mu_{ref,i}[J])$
3. Copy increments in reference (CIR): $\tilde{\mu}_i = (\mu_i[1], \dots, \mu_i[\tilde{t}_i], \mu_i[\tilde{t}_i] + (\mu_{ref,i}[\tilde{t}_i + 1] - \mu_{ref,i}[\tilde{t}_i]), \dots, \mu_i[\tilde{t}_i] + (\mu_{ref,i}[J] - \mu_{ref,i}[\tilde{t}_i]))$

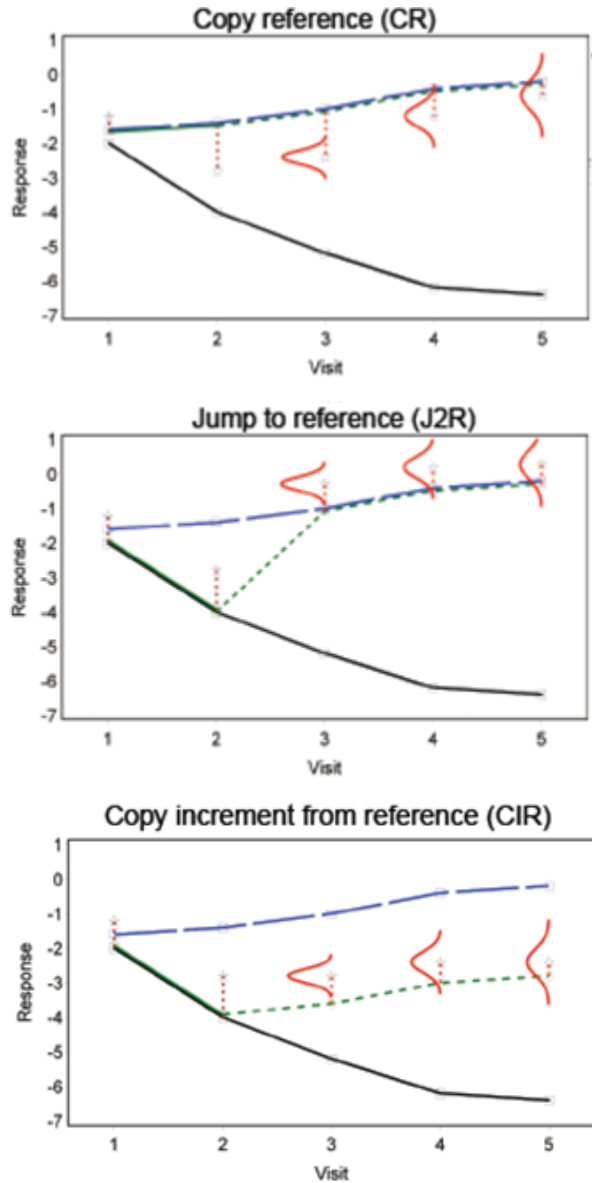


Figure 10. Diagram of reference-based MI

Formular

1. **J2R**: The J2R method assumes subjects in the experimental arm resemble subjects in the control arm immediately after occurrence of IE. For a subject $g_j \in G_j$, the missing endpoint values can be imputed based on the following normal distribution:

$y_{g_j, (j+1) \dots J}^{mis} \sim N \left(\mu_{PBO, (j+1) \dots J}^{(obs)} + \Sigma_{21} \Sigma_{11}^{-1} (y_{g_j, 1 \dots j}^{(obs)} - \mu_{G_j, 1 \dots j}^{(obs)}), \Sigma_{22} - \Sigma_{21} \Sigma_{11}^{-1} \Sigma_{12} \right)$ for the experimental arm.

- $\mu_{PBO, (j+1) \dots J}^{(obs)}$ is the mean vector of the observed values from the placebo completers at timepoint $j+1, \dots, J$
- $y_{g_j, 1 \dots j}^{(obs)}$ is the observed intermediate measurements at timepoints $1, 2, \dots, j$ for subject $g_j \in G_j$.
- $\mu_{G_j, 1 \dots j}^{(obs)}$ is the mean vector at timepoints $1, 2, \dots, j$ for all subjects in group G_j
- $\Sigma = \begin{bmatrix} \Sigma_{11} & \Sigma_{12} \\ \Sigma_{21} & \Sigma_{22} \end{bmatrix}$ is the $J \times J$ sample covariance matrix based on data from placebo completers at time points $1, 2, \dots, J$. Σ is partitioned based on j , the timepoint at which the last measurement before IE is taken. Specifically, $\Sigma_{11} = \text{cov}(y_{\{PBO\}, 1 \dots j}^{(obs)})$ is the $j \times j$ sample covariance matrix based on placebo completers at timepoint $1, 2, \dots, j$, and $\Sigma_{22} = \text{cov}(y_{\{PBO\}, (j+1) \dots J}^{(obs)})$ is the $(J - j) \times (J - j)$ sample covariance matrix based on placebo completers at timepoints $j + 1, \dots, J$. Σ_{12} and Σ_{21} is the corresponding submatrices.

2. **CR**: Unlike the J2R approach in which experimental arm dropouts start to resemble subjects in the placebo group only after the dropout, the CR approach expects the experimental arm dropouts to have similar profiles as the placebo group from the beginning of the study. For subject $g_j \in G_j$, CR imputes the missing values using the following distribution:

- $y_{g_j, (j+1) \dots J}^{mis} \sim N \left(\mu_{PBO, (j+1) \dots J}^{(obs)} + \Sigma_{21} \Sigma_{11}^{-1} (y_{g_j, 1 \dots j}^{(obs)} - \mu_{PBO, 1 \dots j}^{(obs)}), \Sigma_{22} - \Sigma_{21} \Sigma_{11}^{-1} \Sigma_{12} \right)$ for the experimental arm.
- $\mu_{PBO, 1 \dots j}^{(obs)}$ consists of means of the observed values from the placebo completers at timepoints $1, 2, \dots, j$.

3. **CIR**: CIR assumes that the benefit accrued up to IE is retained but that there is no additional residual benefit after the IE. But from IE they progress in the same way as the subjects in the reference arm (parallel to the profile for the reference arm). For subject $g_j \in G_j$, CIR imputes the missing values using the following distribution:

- $y_{g_j, (j+1) \dots J}^{mis} \sim N \left(\mu_{PBO, (j+1) \dots J}^{(obs)} - \mu_{PBO, j}^{(obs)} + \Sigma_{21} \Sigma_{11}^{-1} (y_{g_j, 1 \dots j}^{(obs)} - \mu_{G_j, 1 \dots j}^{(obs)}), \Sigma_{22} - \Sigma_{21} \Sigma_{11}^{-1} \Sigma_{12} \right)$

Code

If the original data have intermediate missing data which we can treat as MAR then we can impute these intermediate missing data before using the code. Often this original data set will be created by using the MI procedure with the MCMC statement using the MONOTONE option to impute intermediate missing data and generate a data set with only monotone missing data.

1. **CR**: The code below shows how to do the Copy Reference Multiple imputation sequentially. The example below creates 500 times for the missing records.

```
proc mi data= test nimpute = 500 seed = 123
  out = test1 noprint;
  var basval change_1 change_2 change_3 change_4;
  mnar model(change_1 change_2 change_3 change_4/modelobs = (trt = "1"))
run;
```

Alternatively, this can be achieved by %mistep by James Roger

```

data test1;
  Set test;
  if change_4<.z then trt4="PLACEBO"; else trt4=trt;
  if change_5<.z then trt5="PLACEBO"; else trt5=trt;
  if change_6<.z then trt6="PLACEBO"; else trt6=trt;
  if change_7<.z then trt7="PLACEBO"; else trt7=trt;
run;

%mistep (Data=test1,Out=test2,Response=change_4,Class=trt4
,Model=trt4 basval,seed=12345);
%mistep (Data=test2,Out=test3,Response=change_5,Class=trt5
,Model=Imputed4 trt basval,seed=12321);
%mistep (Data=test3,Out=test4,Response=change_6,Class=trt6
,Model=Imputed4 Imputed5 trt basval,seed=1216);
%mistep (Data=test4,Out=test5,Response=change_7,Class=trt7
,Model=Imputed4 Imputed5 Imputed6 trt basval,seed=1221);

```

2. J2R: The J2R can't be done directly by MI procedure but the code below with %mistep can do the Jump to Reference Multiple imputation sequentially for each visit.

```

%mistep (Data=test1,Out=test2,Response=change_4,Class=trt4
,Model=trt4 basval,seed=1881);
%mistep (Data=test2,Out=test3,Response=change_5 ,Class=trt5
,Model=Residual4 trt5 basval, seed=1228);
%mistep (Data=test3,Out=test4,Response=change_6 ,Class=trt6
,Model=Residual4 Residual5 trt6 basval, seed=8221);
%mistep (Data=test4,Out=test5,Response=change_7,Class=trt7
,Model=Residual4 Residual5 Residual6 trt7 basval,seed=1281);

```

3. CIR: The CIR can't be done directly by MI procedure but the code below can do the copy increments in reference multiple imputation.

```

%mistep (Data=test1,Out=test2,Response=change_4,Class=trt4,Model=trt4
basval,Suffix=4,Predict= %str(trt4="PLACEBO"),carry=%str((Mu4 - Predict4)*
(1 - missing(change4)*change4)),seed=1881);

%mistep (Data=test2,Out=test3,Response=change_5,Class=trt5,Model=Residual4
trt5 basval,Predict=%str(trt5="PLACEBO"),carry=%str((Mu5 - Predict5)*(1 -
missing(change5)*change5)),delta= mistep_carry,seed=1228);

%mistep (Data=test3,Out=test4,Response=change_6,Class=trt6,Model=Residual4
Residual5 trt6 basval,Predict= %str(trt6="PLACEBO"),carry=%str((Mu6 -
Predict6)*(1-missing(change6)*change6)),delta=mistep_carry,seed=8221);

%mistep (Data=test4,Out=test5,Response=change_7,Class=trt7,Model=Residual4
Residual5 Residual6 trt7 basval,Delta=mistep_Carry,seed=1281);

```

%mistep macro generates IMPUTED1 – IMPUTED4 and also generate MU1 – MU4 and RESIIDUAL1-RESIDUAL4 variable automatically.

DATA

VIEWTABLE: Work.P1							
	PATIENT	trt	basval	change_4	change_5	change_6	change_7
46	3436	DRUG	30	-5	-11	-18	-19
47	3450	DRUG	15	8	6	4	11
48	3453	DRUG	31	-2	-4	-7	-7
49	3606	DRUG	7	5	0	1	4
50	3618	DRUG	8	7	4.918801026	6	2
51	3704	DRUG	17	-5	-9	-7	-7
52	3714	DRUG	26	-7	-15	.	.
53	3716	DRUG	18	-7	-10	-15	-16

Figure 11. Diagram of data structure before the MI

VIEWTABLE: Work.P																							
	PATIENT	trt	basval	change_4	chang	change	change_7	trt4	trt5	trt	trt7	MU4	Imputed4	residual4	MU5	Impute	resid	MU6	Imputed6	resi	MU7	Imputed7	residual7
46	3436	DRUG	30	-5	-11	-18	-19	DRUG	DRUG	DRU	DRUG	-5.7	-5	.7	-7.97	-11	*****	-12.86	-18	****	-16	-19	-2.6242
47	3450	DRUG	15	8	6	4	11	DRUG	DRUG	DRU	DRUG	-0.67	8	9	4.063	6	2	3.0078	4	1	2.48	11	8.51821
48	3453	DRUG	31	-2	-4	-7	-7	DRUG	DRUG	DRU	DRUG	-6.04	-2	4	-5.95	-4	2	-7.913	-7	1	-7.6	-7	0.63987
49	3606	DRUG	7	5	0	1	4	DRUG	DRUG	DRU	DRUG	2.019	5	3	2.203	0	*****	-0.518	1	2	-33	4	4.32775
50	3618	DRUG	8	7	5	6	2	DRUG	DRUG	DRU	DRUG	1.683	7	5	3.79	5	1	2.8951	6	3	3.77	2	-1.7724
51	3704	DRUG	17	-5	-9	-7	-7	DRUG	DRUG	DRU	DRUG	-1.34	-5	-4	-6.93	-9	*****	-10.14	-7	3	-8.2	-7	1.21953
52	3714	DRUG	26	-7	-15	.	.	DRUG	DRUG	PLA	PLACE	-4.36	-7	-3	-9.32	-15	*****	-15.36	*****	****	-18	*****	0.0176
53	3716	DRUG	18	-7	-10	-15	-16	DRUG	DRUG	DRU	DRUG	-1.67	-7	-5	-6.68	-10	*****	-11.48	-15	****	-15	-16	-1.4772

Figure 12. Diagram of data structure after the MISTEP

SIMULATION

BACKGROUND

the simulation code comes from James Bell estimation methods comparisons. The patient-level data for a two-arm randomized clinical trial with 200 subjects per arm. The multivariate normal distribution for the HAMDTL17 depression score across the visits (weeks 0,4,8,14,20,26) with mean and covariance matrix differing between arms. The IE (discontinuation of treatment) for 2 models were simulated as well, of which one is the “discontinuation at random” model and the other is “discontinuation not at random”. The conditional probability for a patient to have an IE at visit j without having an IE at visit $j - 1$ is defined through a logistic regression model. If a patient had an IE, the HAMDTL17 depression score simulated is shifted by $\delta(\cdot)$. We also assume that missing observations can only occur for those patients who already had the IE and that missing is not intermittent, that is once a patient has a missing observation, the observations from all future visits will also be missing.

The conditional probability for a patient, who had IE and has not yet ad any missing observations, to have a missing at visit j is varied between $\text{logit}^{-1}(0.1)$ to $\text{logit}^{-1}(0.6)$ by 0.1 on the logit-scale. Here, $\text{logit}^{-1}(x) = 1/(1 + \exp(-x))$. This defines 6 missing scenarios which are characterized by an increasing proportion of missing values. Table 3 shows the average percentage of patient status at final visit. 5000 simulations were conducted for each scenario.

Table 3. The average percentage of patient status at Visit 5.

Group	Status(%)	Missingness Scenario					
		1	2	3	4	5	6
Control	Receiving assigned treatment (all observed)	74.5	74.5	74.5	74.5	74.5	74.5
	Discontinued assigned treatment and observed	19.3	14.4	10.6	7.6	5.3	3.5
	Discontinued assigned treatment and missing	6.2	11.1	14.9	17.9	20.2	22.0
Treatment	Receiving assigned treatment (all observed)	84.8	84.8	84.8	84.8	84.8	84.8

Discontinued assigned treatment and observed	10.8	7.5	5.1	3.4	2.2	1.3
Discontinued assigned treatment and missing	4.4	7.7	10.1	11.8	13.0	13.9

The operating characteristics of interest in the simulation study are the bias, estimated standard error for the treatment policy estimand, that is an estimator for the difference between treatment and control in mean change of HAMDTL17 depression score from baseline to week 26 including any effects of the intercurrent event. With $m = 1, \dots, M$ the index of the simulated data sets, δ_m the point estimate, and δ the true estimand value, the operating characteristics of bias are estimated as follows, A positive (negative) bias implies that the imputation method underestimates (overestimates) the underlying real treatment effect.

$$\widehat{Bias} = \frac{1}{M} \sum_{m=1}^M \delta_m - \delta$$

standard error S of the treatment effect point estimate are derived by applying Rubin's Rules to combine the model estimates from the 5000 imputed datasets.

SCENARIO ASSUMPTION

Missing scenario 1 is a setting where substantial amounts of post-IE data are collected with approximately 70-75% of patients post-IE continuing to provide data (and approximately 20-25% missing), closely reflecting the PIONEER 1 trial data. Missing scenario 2 and missing scenario 3 show a more moderate amount of post-IE data collected with approximately 35-55% of patients post-IE continuing to provide data (and approximately 45-65% withdrawing). Finally, missingness scenarios 4, 5 and 6 are settings where low amounts of postIE data are collected with approximately 10-30% of patients continuing to provide data (and approximately 70-90% withdrawing).

Table 4. Mean treatment effect bias at different scenario

Assumption	Methods	Missingness Scenarios					
		1	2	3	4	5	6
MAR	FULL	0	0	0	0	0	0
	MI1	-0.017	-0.023	-0.04	-0.052	-0.065	-0.077
	MI2	0	0	0	0	0	-0.077
	MI3	0	0	0	0	0	-0.014
	MI4	-0.009	-0.013	-0.022	-0.03	-0.037	-0.045
	MI5	0	-0.003	-0.007	-0.009	-0.013	-0.019
	MI6	0.002	0.003	0.003	0.003	0.002	0
MNAR	FULL	0	0	0	0	0	0
	MI1	-0.01	-0.02	-0.028	-0.038	-0.046	-0.053
	MI2	0	0	-0.001	-0.002	-0.005	-0.013
	MI3	0	0	0	-0.001	-0.002	-0.009
	MI4	-0.003	-0.005	-0.008	-0.012	-0.015	-0.02
	MI5	0.005	0.007	0.009	0.01	0.009	0.009
	MI6	0.007	0.013	0.017	0.02	0.022	0.023

FULL refer to there is no missing data which serve as the base case; MI1 refer to the conventional MI assume MAR. MI2 refer to IE indicator RD approach. MI3 refer to IE pattern RD approach. MI4, MI5 and MI6 refer to CIR, CR and J2R reference based approach. The number in bold over-estimate the treatment effect.

Table 5. Mean treatment effect standard error at different scenario

Assumption	Methods	Missingness Scenarios					
		1	2	3	4	5	6
MAR	FULL	0.11	0.11	0.11	0.11	0.11	0.11
	MI1	0.112	0.114	0.116	0.118	0.1198	0.121
	MI2	0.112	0.117	0.123	0.132	0.145	0.158
	MI3	0.115	0.123	0.134	0.147	0.161	0.171
	MI4	0.113	0.116	0.117	0.119	0.121	0.122
	MI5	0.113	0.116	0.117	0.119	0.121	0.122
	MI6	0.113	0.116	0.117	0.12	0.122	0.123
MNAR	FULL	0.11	0.11	0.11	0.11	0.11	0.11
	MI1	0.11	0.112	0.114	0.116	0.117	0.118
	MI2	0.112	0.115	0.121	0.13	0.141	0.156
	MI3	0.113	0.12	0.132	0.145	0.158	0.168
	MI4	0.111	0.114	0.116	0.117	0.119	0.121
	MI5	0.111	0.114	0.116	0.117	0.119	0.121
	MI6	0.111	0.114	0.116	0.117	0.119	0.122

FULL refer to there is no missing data, which serves as the base case; MI1 refer to the simple method. MI2 refer to IE indicator RD approach. MI3 refer to IE pattern RD approach. MI4, MI5 and MI6 refer to CIR, CR and J2R reference based approach. The number in bold are variability inflation cases(>0.13).

RESULTS

1. Conventional MI based on the “MAR” assumption displayed anti-conservative bias for all scenarios while the mean treatment effect standard error is the smallest as the missingness increased comparing with Retrieved dropout approach.
2. The IE indicator retrieved dropout approach displayed unbiased estimation except the scenario 6. There is medium mean standard error of the treatment effect estimates. The model is not suitable to situations where patients take more than one visit to fully transition between clinical status (pre or post IE) as the model cannot account for the number of visits takes to transition.
3. The IE pattern retrieved dropout approach showed unbiased estimation but in all cases there’s highest standard error of the treatment effect estimates. The models also suffer from fitting problems. These fitting problems likely prevent the use of even more complex retrieved dropout models to handle different types of intercurrent events separately. In general, it is questionable if the situationally smaller bias than other methods is worth the fitting problems and substantial decrease in precision, unless there is a strong expectation of at least 60-70% post-IE data (e.g. scenario1)
4. J2R analysis was close to unbiased, showing a slight increase in conservative bias as the missingness scenarios increased. The CR is the next best, with anti-conservative bias. The CIR approach was the worst performing, with more extensive anti-conservative bias that worsened with increasing missingness. The J2R, CIR and CR approaches were almost identical and constant across all missingness scenarios. In all cases, they resulted in slightly smaller variation than the full data estimates. This is because

reference-based approach derives the uncertainty for the missing data from the whole control arm, whereas the RD approaches derive it from only the much smaller post-IE data within each arm

CONCLUSION

For regulatory decision making, the primary interest is generally to estimate how well a drug will work in clinical practice. To this end, the treatment policy strategy, by which subjects are followed, assessed, and analyzed regardless of their compliance to the planned treatment, is usually utilized. As noted in the National Academy of Sciences report titled "The Prevention and Treatment of Missing Data in Clinical Trials", the best approach of handling missing data is always to prevent the problem by well-planned study designs and conduct, and efficacy and safety data should be collected for all subjects including those who prematurely discontinue treatments during the trial.

Despite efforts to prevent their occurrence, missing data are practically inevitable in clinical trials. The goal of the missing data handling is to address what the measurement would have been if it was not missing, so that we can determine the most appropriate estimate of treatment effect with corresponding uncertainty.

We used multiple imputation method to handle the missing data to reduce the bias and increase the accuracy. Simple approaches to deal with missing data, that ignore distinctions between pre- and post-IE data, are biased and not appropriate unless only a trivial amount of missing data is expected. More complex RD models that do condition on IE occurrences appear to address the bias issue but at the expense of fitting problems or higher mean treatment effect standard error. Although many researches were done for the reference based multiple imputation, the J2R and CR methods are not considered top choices for missing data imputation for primary efficacy analyses in antihyperglycemic product development programs for a placebo-controlled trial. The reason is that the J2R is computationally much intensive. Additionally, J2R imputes missing endpoint values from the experimental arm conditional on the intermediate measurements in the experimental arm it's not always conserved. On the other hand, the CR approach assumes that subjects in the experimental arm have the same mean response as the subjects in the placebo group even before the dropout, which does not align with the actual observations. Finally, the analysis result based on the CR approach is not always conservative,

As statisticians, we would prefer not to have to make such unenviable choices. We believe it ought to be possible to deal with the variance inflation of RD methods by Bayesian methods incorporating into them a clinically plausible prior belief of, e.g., no treatment effect post-IE, so that with a lack of data the estimator becomes similar to a de facto RBI method. This would have the added advantage that the arbitrariness of the RBI variance is now directly relatable to the strength of the prior.

REFERENCES

- [1] Alessandro Noci, Craig Gower-Page, and Marcel Wolbers. "rbmi: Statistical Specifications". 2023-11-24. Available at https://cran.r-project.org/web/packages/rbmi/vignettes/stat_specs.html#scope-of-this-document
- [2] National Research Council. (2010). The Prevention and Treatment of Missing Data in Clinical Trials. Panel on Handling Missing Data in Clinical Trials. Committee on National Statistics, Division of Behavioral and Social Sciences and Education. Washington, DC: The National Academies Press.
- [3] Wang Y, Tu W, Kim Y, et al. Statistical methods for handling missing data to align with treatment policy strategy[J]. Pharmaceutical statistics, 2023, 22(4): 650-670.
- [4] Carpenter, James R., James H. Roger, and Michael G. Kenward. "Analysis of longitudinal trials with protocol deviation: a framework for relevant, accessible assumptions, and inference via multiple imputation." Journal of biopharmaceutical statistics 23.6 (2013): 1352-1371.
- [5] Bell, James, et al. "Estimation methods for estimands using the treatment policy strategy; a simulation study based on the PIONEER 1 Trial." arXiv preprint arXiv:2402.12850 (2024).
- [6] Wolbers, Marcel, et al. "Standard and reference - based conditional mean imputation." Pharmaceutical statistics 21.6 (2022): 1246-1257.

[7] White, Ian, Royes Joseph, and Nicky Best. "A causal modelling framework for reference-based imputation and tipping point analysis in clinical trials with quantitative outcome." *Journal of Biopharmaceutical Statistics* 30.2 (2020): 334-350.

ACKNOWLEDGEMENTS

I would like to thank Li gaoyang, Wu elli, Yao xiangming and DIA DHC group in China to give me a lot of support during the simulation and paper preparation. I would like to thank Thomas Drury to share the simulation code in the github, that's really helpful for me to produce the simulation results. We would also thank Sun Shuguang to guide me how to use the server SAS and server R to do the simulation.

RECOMMENDED READING

- *Base SAS® Procedures Guide*
- *SAS® For Dummies®*
- *Estimation methods for estimands using the treatment policy strategy; a simulation study based on the PIONEER 1 Trial*
- *Statistical Considerations and Methods for Handling Missing Outcome Data During the Era of COVID-19 PharmaSUG 2023 - Paper SA-269*
- *FDA. Insulin degludec and liraglutide injection, briefing document for endocrinologic and metabolic drugs advisory committee meeting (EMDAC). 2016. Accessed July 10, 2022. <https://fda.report/media/97814/FDA-Briefing-Information-for-the-May-24-2016-Meeting-of-the-Endocrinologic-and-Metabolic-Drugs-Advisory-Committee.pdf>*
- *FDA. Semaglutide, briefing document for endocrinologic and metabolic drugs advisory committee meeting (EMDAC). 2017. Accessed July 10, 2022. <https://fda.report/media/108291/FDA-Briefing-Information-for-the-October-18-2017-Meeting-of-the-Endocrinologic-and-Metabolic-Drugs-Advisory-Committee.pdf>*

CONTACT INFORMATION

Your comments and questions are valued and encouraged. Contact the author at:

Sun Haolin
Shanghai Xianxiang Medical Technology Co., Ltd
+8615000773163
sunlanlin@163.com

Any brand and product names are trademarks of their respective companies.