

Handling Missing Data in Clinical Trials: Techniques and Methods

Pennidhi Karlakunta, Merck & Co., Inc., North Wales, PA, USA

Naveen Kommuru, Merck & Co., Inc., North Wales, PA, USA

ABSTRACT

The reliability and interpretability of results from randomized clinical trials are greatly influenced by the quality of data. Missing data in clinical trials can be a serious problem, especially in randomized trials where missing data could produce biased estimates and compromise inferences. The preferred and most satisfactory approach to address missing data is to identify ways to prevent it. However, missing data are ubiquitous as some of the trials could span over months or years. Since each trial has its own set of design and measurement characteristics, currently no universal method for handling missing data in clinical trials.

This paper will present programming techniques and methods to handle missing data with a major focus on multiple imputation.

INTRODUCTION

Missing data are a potential source of bias and seriously compromise inferences when analyzing clinical trials. Interpretation of the results of a trial is always problematic when the proportion of missing values is substantial. Data entry errors or missing critical datapoints can affect conclusions in trials of treatments for many diseases. Even though there are many possible reasons for missing data (e.g. patient refusal to continue in the study, patient withdrawals due to treatment failure, treatment success or adverse events, data entry errors, patients moving, etc.), only some of which are related to study treatment.

The concept of missing data can be defined as both the existence of missing data and the mechanism that explains the reason for the data being missing. The missing values lead to biased conclusions due to many factors such as the relationship between missingness, treatment assignment and outcome. The type of measure applied to quantify the treatment effect and the expected changes over time for the variables being measured. The way missing data are handled can have, depending upon the amount and type of missing data, a crucial influence on the results of a clinical trial and on the certainty with which conclusions can be drawn.

It should be noted that the strategy applied to handle missing values might constitute a source of bias since there is no universal best approach for all situations. The acceptability of an approach will depend upon the assumptions made and whether it is reasonable to make those assumptions in the particular case of interest. It is also very important when designing a study that the likely pattern of missing data is considered when specifying the primary analysis and the predefined sensitivity analyses.

THE EFFECT OF MISSING VALUES ON DATA ANALYSIS

POWER, VARIABILITY AND BIAS

The sample size and variability of the outcomes affect the power of a clinical trial. The power of a trial will increase if the sample size is increased or if the variability of the outcomes is reduced.

If missing values are handled by simply excluding any patients with missing values from the analysis, it will result in a reduction of datapoints, which leads to reduction of the statistical power. If we have higher number of missing values, the power of the trial will be reduced.

Conversely, non-completers might be more likely to have extreme values such as treatment failure leading to dropout, extremely good response leading to loss of follow-up. Therefore, the loss of these non-completers could lead to an underestimation of variability and hence artificially narrow the confidence interval for the treatment effect. If the methods used to handle missing data do not adequately take this into account, the resulting confidence interval cannot be considered a valid summary of the uncertainty of the treatment effect.

While the reduction of the statistical power is mainly related to the number of missing values, the risk of bias in the estimation of the treatment effect from the observed data depends upon the relationship between missingness and treatment/outcome. If the unmeasured observation is related to the real value of the outcome (e.g. the unobserved measurements have a higher proportion of poor outcomes), it may lead to a biased estimate of the treatment effect even if the missing values are not related to treatment.

PATTERNS OF MISSING DATA

The most common missing data pattern is termed generalized or arbitrary, where there is no particular pattern in the missing data structure. Missing observations are distributed across cases and variables in a nonsystematic fashion.

Some data collection processes produce a more structured or systematic pattern of missingness in the data, which is called monotonic missing data pattern.

A third pattern of missing data arises in studies that incorporate randomization procedures to allow item missing data on selected variables for subsets of study observations. The technique is termed matrix sampling or “missing by design” sampling.

MISSING DATA MECHANISMS

Missing data for a single variable is classified into one of three categories: missing completely at random (MCAR), missing at random (MAR), and missing not at random (MNAR).

MCAR is that the likelihood of missing data is unrelated to any observed or unobserved variables. That is, the chance of missing data is the same for individuals in different treatment groups and those who have differential disease severity or treatment response. For example, a dropped test tube in a lab or an equipment failure may lead to missing data. As this is equally likely to occur in any subject in the study, it represents a completely random process.

When the likelihood of missing data is related to observed variables but not to unobserved variables, the missing data mechanism is referred to as missing at random (MAR). For example, missing data caused by features of the study design such as participants being removed from the trial if their conditions are not controlled sufficiently well according to protocol criteria, or dropout based on recorded side-effects.

When the likelihood of missing data depends on the unobserved data, the missing data is termed missing not at random (MNAR). For example, in substance abuse trials with abstinence as an outcome, it is usual that dropout is higher for those who have relapsed.

PREVENTION OF MISSING DATA

TRIAL OUTCOMES AND ESTIMANDS

The trial protocol should explicitly define (a) the objective(s) of the trial; (b) the associated primary outcome or outcomes; (c) how, when, and on whom the outcome or outcomes will be measured; and (d) the measures of intervention effects. These measures should be meaningful for all study participants, and able to derive the estimates with minimal assumptions. Concerning the latter, the protocol should address the potential impact and treatment of missing data.

MINIMIZING DROPOUTS IN TRIAL DESIGN

Investigators, sponsors, and regulators should design clinical trials consistent with the goal of maximizing the number of participants who are maintained on the protocol-specified intervention until the outcome data are collected.

CONTINUING DATA COLLECTION FOR DROPOUTS

Trial sponsors should continue to collect information on key outcomes on participants who discontinue their protocol specified intervention in the course of the study, except in those cases for which a compelling cost-benefit analysis argues otherwise. This information should be recorded and used in the analysis.

The trial design team should consider whether or not participants who discontinue the protocol intervention should have access to and be encouraged to use specific alternative treatments. Such treatments should be specified in the study protocol.

Data collection and information about all relevant treatments and key covariates should be recorded for all initial study participants, whether or not participants received the intervention specified in the protocol.

ACTIONS FOR INVESTIGATORS AND SITE PERSONNEL

Study sponsors should explicitly anticipate potential problems of missing data. In particular, the trial protocol should contain a section that addresses missing data issues, including the anticipated amount of missing data, and steps taken in trial design and trial conduct to monitor and limit the impact of missing data.

Informed consent documents should emphasize the importance of collecting outcome data from individuals who choose to discontinue treatment during the study, and they should encourage participants to provide this information whether or not they complete the anticipated course of study treatment.

TARGETS FOR ACCEPTABLE RATES OF MISSING DATA

All trial protocols should recognize the importance of minimizing the amount of missing data, and, in particular, they should set a minimum rate of completeness for the primary outcome(s), based on what has been achievable in similar past trials.

METHODS OF HANDLING MISSING DATA

- Complete case analysis
- Single Imputation
 - LOCF
 - BOCF
 - WOCF
 - Mean Substitution
- Inverse probability weighting
- Likelihood-based analysis
- Event time analysis
- Non-responder Imputation
- Multiple Imputation (MI)

Complete Case Analysis approach may be considered in exploratory studies, especially in the initial phases of drug development. However, it cannot be recommended as the primary analysis in a confirmatory trial. It can be used in confirmatory trials as a secondary supportive analysis (sensitivity analysis) to illustrate the robustness of conclusions.

Single imputation methods replace a missing datapoint by a single value and analyses are conducted as if all data were observed. One widely used single imputation method is Last Observation Carried Forward (LOCF). This analysis imputes last measured value of the endpoint to all subsequent, scheduled, but missing, evaluations. Only under certain restrictive assumptions LOCF produce an unbiased estimate of the treatment effect. In some situations, LOCF does not produce conservative estimates. Baseline Observation Carried Forward (BOCF), Worst Observation Carried Forward (WOCF) and Mean Substitution are other single imputation methods.

Estimates obtained from this complete case analysis may be biased if the excluded individuals are systematically different from those included. Inverse probability weighting (IPW) can reduce this bias. In this method, complete cases are weighted by the inverse of their probability of being a complete case. Another use of IPW is to correct for unequal sampling fractions.

Maximum likelihood method does not impute any data, but rather uses each case available data to compute maximum likelihood estimates. The maximum likelihood estimate of a parameter is the value of the parameter that is most likely to have resulted in the observed data. Like multiple imputation, this method gives unbiased parameter estimates and standard errors. One advantage is that it does not require the careful selection of variables used to impute values that multiple imputation requires; however, it is limited to linear models.

In an event time analysis, missingness can take at least two important forms. First, with repeated event times, it is possible that some event times in a sequence are unobserved. Second, it is common for event times to be unobserved because the event has not occurred when follow-up terminates. This situation is referred to as right censoring in the time-to-event.

For the binary variables, the Non-responder Imputation (NRI) approach will be used for imputing missing data which is highly conservative method that assumes non-response for all missing data.

Multiple imputation is designed to fill in missing data under one or more models and to properly reflect uncertainty associated with the “filling in” process. Instead of imputing a single value for each missing observation, a set of values for each missing observation is generated from its predictive distribution.

MULTIPLE IMPUTATION (MI)

Multiple imputation (MI) is not simply a technique for imputing missing data. It is also a method for obtaining estimates and correct inferences for statistics ranging from simple descriptive statistics to the parameters of complex multivariate models.

This approach consists of a three-step process:

- 1) The analyst defines a multivariate “imputation model” for the data, and under this model PROC MI is used to independently impute missing values in the original dataset M times, generating M complete “repetition” versions of the analysis dataset.
- 2) Standard SAS procedures or SURVEY procedures for the analysis of each of the M completed datasets and output of the results.
- 3) PROC MIANALYZE inputs the results of the M separate analyses and applies multiple imputation formulae to generate estimates, standard errors, confidence intervals, and test statistics for the descriptive statistics or model parameters of interest.

WHY THE MULTIPLE IMPUTATION METHOD?

No imputation method or statistical modeling technique is optimal for all forms of missing data problems. The strengths of the multiple imputation approach to missing data rest on the following attributes.

- Model-based
- Multivariate
- Multiple independent repetitions
- Robust
- Very usable

METHODS FOR MULTIPLE IMPUTATION

Depending on the pattern of missing data and variable types, PROC MI provides the following primary classes of methods for generating the multiple imputations.

Missing Data Pattern	Variable Type	Method	PROC MI Statement
Monotone	Continuous	Linear regression predictive mean matching propensity score	MONOTONE REG MONOTONE REGPMM MONOTONE PROPENSITY
	Binary	Logistic regression	MONOTONE LOGISTIC
	Nominal	Discriminant function	MONOTONE DISCRIM
Arbitrary	Continuous	With continuous covariates: MCMC monotone method MCMC full-data imputation	MCMC IMPUTE=MONOTONE MCMC IMPUTE=FULL
	Continuous	With mixed covariates: FCS regression FCS predictive mean matching	FCS REG FCS REGPMM
	Binary	FCS logistic regression	FCS LOGISTIC
	Nominal	FCS discriminant function	FCS DISCRIM

Table 1.1 – Multiple Imputation methods based on the missing data pattern

HOW MANY MULTIPLE IMPUTATION REPETITIONS ARE NEEDED?

If the rates of missing data and therefore fraction of missing information are modest (< 20%), MI analyses based on as few as M=5 or M=10 repetitions will achieve > 96% of the maximum statistical efficiency. If the fraction of missing information is high (30% to 50%), analysts are advised to specify M=20 or M=30 to maintain a minimum relative efficiency of 95% or greater. Historically, the rule of thumb for most practical applications of MI has been to use M=5, and this is the current default in PROC MI.

CASE STUDY: TESTING RESULTS USING MULTIPLE IMPUTATION

Below table was produced using ePRO analysis dataset (ADPRO) with complete case analysis (without any imputations). Dummy data was used for this analysis. We will use the same ePRO dataset and reproduce the same table using multiple imputation method as explained below.

Analysis of Change from Baseline in EORTC QLQ-C30 Global Health Status/QoL at Week 24
(FAS Population)

Treatment	Baseline		Week 24		Change from Baseline at Week 24	
	N	Mean (SD)	N	Mean (SD)	N	LS Mean (95% CI) [†]
Treatment A	202	66.87 (18.982)	125	74.53 (16.943)	222	4.65 (-1.55, 7.75)
Treatment B	69	69.57 (19.742)	46	70.83 (18.151)	74	0.07 (-4.73, 4.88)
Pairwise Comparison					Difference in LS Means (95% CI)	p-Value
Treatment A vs. Treatment B					4.57 (-0.79, 9.94)	0.0943

[†] Based on cLDA model with the PRO scores as the response variable, and treatment by study visit interaction, stratification factors (prior auto-SCT (yes, no) and disease status after frontline therapy (primary refractory, relapsed less than 12 months, relapsed 12 months or more)) as covariates.
For baseline and Week 24, N is the number of subjects in each treatment group with non-missing assessments at the specific time point; for change from baseline, N is the number of subjects in the analysis population in each treatment group.
Database Cutoff Date: 16JAN2020

Source: [adam-adsl; adpro]

Report 1 – Global Health Status QOL (based on Non-imputed data).

ADPRO ANALYSIS DATASET

Below is the ADaM dataset structure (BDS) of the ePRO analysis dataset used for the analysis.

	USUBJID	TRT01P	STRATA1N	PARAM	PARAMCD	BASE	AVAL	CHG	CHGCAT1	AVISIT
1	XXXX-XXX_XX00A	Treatment A	6	Global health status/QoL	QL2	.	83.333	.	.	Week 6
2	XXXX-XXX_XX00A	Treatment A	6	Global health status/QoL	QL2	.	83.333	.	.	Week 12
3	XXXX-XXX_XX00A	Treatment A	6	Global health status/QoL	QL2	.	83.333	.	.	Week 18
4	XXXX-XXX_XX00B	Treatment A	4	Global health status/QoL	QL2	91.667	91.667	.	.	Baseline
5	XXXX-XXX_XX00B	Treatment A	4	Global health status/QoL	QL2	91.667	91.667	0	Stable	Week 18
6	XXXX-XXX_XX00B	Treatment A	4	Global health status/QoL	QL2	91.667	83.333	-8.334	Stable	Week 24
7	XXXX-XXX_XX00C	Treatment A	6	Global health status/QoL	QL2	50	50.000	.	.	Baseline
8	XXXX-XXX_XX00C	Treatment A	6	Global health status/QoL	QL2	50	66.667	16.667	Improved	Week 6
9	XXXX-XXX_XX00C	Treatment A	6	Global health status/QoL	QL2	50	66.667	16.667	Improved	Week 12
10	XXXX-XXX_XX00C	Treatment A	6	Global health status/QoL	QL2	50	66.667	16.667	Improved	Week 18

RESTRUCTURING ANALYSIS DATASET

ePRO dataset was re-structured using below code to have visit level columns (Y1 – Y5 for Week 6 to Week 24) so that the model applied in multiple imputation uses all the given visit variable values as covariates.

```
proc transpose data=adpro out=adpro_t(drop=_name_ _label_) prefix=y;
  by usubjid trt01pn trt01p STRATA1 paramcd param;
  id avisitn;
  var aval;
run;
```

The re-structured data sample of the ADPRO dataset using the above sample code:

USUBJID	TRT01P	PARAMCD	PARAM	Y2	Y3	Y4	Y1	Y5
XXXX-XXX_XX00A	Treatment A	QL2	Global health status/QoL	83.333	83.333	83.333	.	.
XXXX-XXX_XX00B	Treatment A	QL2	Global health status/QoL	.	.	91.667	91.667	83.333
XXXX-XXX_XX00C	Treatment A	QL2	Global health status/QoL	66.667	66.667	66.667	50.000	.

CHECKING THE MISSING DATA PATTERN

Before applying the multiple imputation method, we need to find the pattern and percentage of missing data using the below sample code. Based on the pattern the correct imputation model to be used, will be decided.

```
proc mi data=adplda_t nimpute=0 simple;
  class paramcd trt01pn;
  fcs;
  var y1-y5 paramcd trt01pn;
run;
```

Missing Data Patterns														
Group	Y1	Y2	Y3	Y4	Y5	PARAMCD	TRT01PN	Freq	Percent	Group Means				
										Y1	Y2	Y3	Y4	Y5
1	X	X	X	X	X	X	X	131	44.26	67.875328	73.345977	74.109397	73.345962	73.218756
2	X	X	X	X	.	X	X	42	14.19	67.857190	65.674667	66.865071	64.880952	.
3	X	X	X	.	X	X	X	15	5.07	68.333267	72.222133	71.666667	.	71.666533
4	X	X	X	.	.	X	X	41	13.85	64.024415	67.479707	60.365878	.	.
5	X	X	.	X	X	X	X	4	1.35	81.249750	60.416750	.	72.916750	72.916750
6	X	X	.	X	.	X	X	4	1.35	70.833500	72.916750	.	77.083250	.
7	X	X	.	.	.	X	X	18	6.08	67.592500	60.648111	.	.	.

From the above output, missing data pattern is arbitrary, and the cumulative percent of missing data is more than 30%. Based on the Table 1.1, we decided to use FCS REG method for multiple imputation of the missing values, since the targeted (to be imputed) variable type is 'continuous' and includes the mixed covariates in the imputation model. In addition, we decided to use 50 repetitions to achieve the maximum statistical efficiency.

MODEL BASED MULTIPLE IMPUTATION (Step 1)

Multiple imputation applied with SAS procedure PROC MI based on FCG REG method resulting 50 repetitions of dataset which are indexed by column _IMPUTATION_. Each missing value is imputed based on statistical modeling, and this process is repeated 50 times. Below is the sample code and resulting imputed output dataset.

```
PROC MI DATA=adpro_t OUT=adpro_mi SEED=3475 NIMPUTE=50 noprint minmaxiter=1000
  minimum = . 0 0 0 0 0 maximum=. 100 100 100 100 100;
  by paramcd trt01pn;
  class STRATA1;
  FCS REG (y1-y5);
  VAR STRATA1 y1-y5;
RUN;
```

IMPUTATION	USUBJID	TRT01P	PARAMCD	PARAM	Y2	Y3	Y4	Y1	Y5
1	XXXX-XXX_XX00A	Treatment A	QL2	Global health status/QoL	83.333	83.333	83.333	83.884	92.218
1	XXXX-XXX_XX00B	Treatment A	QL2	Global health status/QoL	89.132	76.850	91.667	91.667	83.333
1	XXXX-XXX_XX00C	Treatment A	QL2	Global health status/QoL	66.667	66.667	66.667	50.000	71.330
2	XXXX-XXX_XX00A	Treatment A	QL2	Global health status/QoL	83.333	83.333	83.333	81.288	93.142
2	XXXX-XXX_XX00B	Treatment A	QL2	Global health status/QoL	73.708	91.064	91.667	91.667	83.333
2	XXXX-XXX_XX00C	Treatment A	QL2	Global health status/QoL	66.667	66.667	66.667	50.000	50.737
3	XXXX-XXX_XX00A	Treatment A	QL2	Global health status/QoL	83.333	83.333	83.333	87.961	59.878
3	XXXX-XXX_XX00B	Treatment A	QL2	Global health status/QoL	96.062	91.519	91.667	91.667	83.333
3	XXXX-XXX_XX00C	Treatment A	QL2	Global health status/QoL	66.667	66.667	66.667	50.000	72.608

BACK TO ADAM BDS STRUCTURE

In order to analyze the data after the imputation and to derive the final estimates, above multiple imputed dataset is transposed back to the original ADaM BDS structure along with re-derived required variables such as BASE, CHG, CHGCAT1 and AVISIT. Below is the sample BDS structure dataset with the imputations, which includes 50 repetitions of the data, indexed by the column _IMPUTATION_.

```
proc transpose data=adpro_mi out=adpromi;
  by _imputation_ trt01pn trt01p usubjid STRATA1 parcat2 paramcd paramn param;
run;
```

IMPUTATION	USUBJID	TRT01P	STRATA1N	PARAM	PARAMCD	BASE	AVAL	CHG	CHGCAT1	AVISIT
1	XXXX-XXX_XX00A	Treatment A	6	Global health status/QoL	QL2	83.54324	83.333	-0.2102396	Stable	Week 6
1	XXXX-XXX_XX00A	Treatment A	6	Global health status/QoL	QL2	83.54324	83.333	-0.2102396	Stable	Week 12
1	XXXX-XXX_XX00A	Treatment A	6	Global health status/QoL	QL2	83.54324	83.333	-0.2102396	Stable	Week 18
1	XXXX-XXX_XX00A	Treatment A	6	Global health status/QoL	QL2	83.54324	83.54324	0	Stable	Baseline
1	XXXX-XXX_XX00A	Treatment A	6	Global health status/QoL	QL2	83.54324	45.70502	-37.83822	Deteriorated	Week 24
1	XXXX-XXX_XX00C	Treatment A	6	Global health status/QoL	QL2	50	66.667	16.667	Improved	Week 6
1	XXXX-XXX_XX00C	Treatment A	6	Global health status/QoL	QL2	50	66.667	16.667	Improved	Week 12
1	XXXX-XXX_XX00C	Treatment A	6	Global health status/QoL	QL2	50	66.667	16.667	Improved	Week 18
1	XXXX-XXX_XX00C	Treatment A	6	Global health status/QoL	QL2	50	50	0	Stable	Baseline
1	XXXX-XXX_XX00C	Treatment A	6	Global health status/QoL	QL2	50	47.18979	-2.8102134	Stable	Week 24

ANALYSIS WITH MULTIPLE IMPUTED DATASET (Step 2)

Same SAS procedures PROC MEANS and PROC MIXED are used to analyze this imputed data, in keeping consistent with the analysis of non-imputed data to produce the **Report 1**. The below sample codes produce the summary statistics required for the final Report for each repetition of the imputed dataset using by variable _IMPUTATION_.

```
proc means data= adpromi MEAN STD median min max ql;
  by _imputation_ trt01pn trt01p;
  var aval base;
run;
```

	_TRTNAM	BIGN	_IMPUTATION_	NB	MEANB	STDB	NV	MEANV	STDV	LABEL	ESTIMATE	STDERR
1	Treatment A	222	1	222	67.639703048	18.543853223	222	72.521502239	17.472059936	P1: Week 24; TRT: 1	4.6801	1.3535
2	Treatment A	222	2	222	67.466466093	18.862392895	222	70.628639731	17.97807418	P1: Week 24; TRT: 1	2.9400	1.2920
3	Treatment A	222	3	222	66.415047728	18.999528841	222	70.444187458	18.601435802	P1: Week 24; TRT: 1	3.6335	1.3533
4	Treatment A	222	4	222	66.838788003	18.933755662	222	70.803907265	19.200848462	P1: Week 24; TRT: 1	3.6908	1.3811
5	Treatment A	222	5	222	66.524958789	19.025194811	222	71.336788971	17.523955587	P1: Week 24; TRT: 1	4.3633	1.3449
6	Treatment A	222	6	222	67.013206753	18.774523193	222	70.46192564	17.616884323	P1: Week 24; TRT: 1	3.2280	1.2996

Dataset: est_sum

```
proc mixed data= adpromi;
  by _imputation_;
  class avisitn usubjid STRATA1N;
  model aval=avisitn STRATA1N v2t1 v3t1 v4t1 v5t1 / ddfm=kr;
  repeated avisitn / subject=usubjid type=un R;
  estimate "P1: Week 24; TRT: 1" avisitn -1 0 0 0 1 v5t1 1 / divisor=1 cl alpha=0.05;
  estimate "P1: Week 24; TRT: 2" avisitn -1 0 0 0 1/divisor=1 cl alpha=0.05;
  estimate "P2: Week 24; TRT: 1 - 2" v5t1 1 / cl alpha=0.05;
run;
```

	IMPUTATION	LABEL	ESTIMATE	STDERR	DF	TVALUE	PROBT	ALPHA	LOWER	UPPER	_BY_TRT	_1STTRT	_2NDTRT
1	1	P2: Week 24; TRT: 1 - 2	8.0335	2.3208	289	3.46	0.0006	0.05	3.4658	12.6013	TRT1_2	1	2
2	2	P2: Week 24; TRT: 1 - 2	4.4622	2.2132	286	2.02	0.0447	0.05	0.1060	8.8184	TRT1_2	1	2
3	3	P2: Week 24; TRT: 1 - 2	5.6605	2.3819	292	2.38	0.0181	0.05	0.9726	10.3484	TRT1_2	1	2
4	4	P2: Week 24; TRT: 1 - 2	4.1921	2.4069	288	1.74	0.0826	0.05	-0.5453	8.9295	TRT1_2	1	2
5	5	P2: Week 24; TRT: 1 - 2	5.0368	2.2596	290	2.23	0.0266	0.05	0.5895	9.4842	TRT1_2	1	2
6	6	P2: Week 24; TRT: 1 - 2	1.6045	2.1621	286	0.74	0.4586	0.05	-2.6511	5.8601	TRT1_2	1	2

Dataset: est_comp

POOLING THE RESULTS (Step 3)

PROC MIANALYZE is being used for each of the above two datasets (est_sum, est_comp – which includes 50 repetitions data) to produce the pooled statistical inferences. Below is the sample code and the pooled statistical inferences for these two datasets.

```
proc mianalyze data=est_sum;
  by _trtnam;
  modeleffects estimate meanb meanv;
  stderr stderr stdb stdv;
  ods output ParameterEstimates=est0sum;
run;
```

```
proc mianalyze data=est_comp;
  modeleffects estimate;
  stderr stderr;
  ods output ParameterEstimates=est0comp;
run;
```

	_TRTNAM	ESTIMATE	MEANB	MEANV	STDB	STDV	LOWER	UPPER	BIGN	NB	NV
1	Treatment A	4.146304	66.873930	71.375857	18.816351	17.704110	0.81327	7.4793	222	222	222
2	Treatment B	-0.339998	69.082885	67.676018	19.748969	19.359143	-4.90083	4.2208	74	74	74

Dataset: est0sum

	NIMPUTE	ESTIMATE	STDERR	LOWER	UPPER	DF	TVALUE	PROBT	ALPHA	LABEL	_BY_TRT
1	50	4.486302	2.661546	-0.74058	9.713181	611.5	1.69	0.0924	0.05	P2: Week 24; TRT: 1 - 2	TRT1_2

Dataset: est0comp

FINAL REPORT WITH MULTIPLE IMPUTATION

Using the pooled statistical inferences listed in the above two datasets (est0sum, est0comp), the final output (Report 2) is produced consistent with the same template as Report 1 (based on the Non-imputed data).

Analysis of Change from Baseline in EORTC QLQ-C30 Global Health Status/QoL with Multiple Imputation Based on MAR at Week 24 (FAS Population)

Treatment	Baseline		Week 24		Change from Baseline at Week 24	
	N	Mean (SD)	N	Mean (SD)	N	LS Mean (95% CI) [†]
Treatment A	222	66.87 (18.816)	222	71.38 (17.704)	222	4.15 (0.81, 7.48)
Treatment B	74	69.08 (19.749)	74	67.68 (19.359)	74	-0.34 (-4.90, 4.22)
Pairwise Comparison					Difference in LS Means (95% CI)	
Treatment A vs. Treatment B					4.49 (-0.74, 9.71)	
					p-Value	
					0.0924	

[†] Based on cLDA model with the PRO scores as the response variable, and treatment by study visit interaction, stratification factors (prior auto-SCT (yes, no) and disease status after frontline therapy (primary refractory, relapsed less than 12 months, relapsed 12 months or more)) as covariates.

For baseline and Week 24, N is the number of subjects in each treatment group with non-missing assessments at the specific time point; for change from baseline, N is the number of subjects in the analysis population in each treatment group.

Database Cutoff Date: 16JAN2020

Source: [adam-adsl; adplda]

Report 2 – Global Health Status QOL (based on multiple imputation).

CONCLUSION

We have taken the Report 1 as a reference table (based on the non-imputed data) to do our case study using the Multiple imputation (MI) method. The pattern of the missing data is identified as 'Arbitrary' using PROC MI and we decided to use the FCS REG imputation method since variable type is 'continuous' and has mixed covariates. We have also produced the report 2 using multiple imputation in the same template as Report 1. This Report 2 could be used as a supporting report as part of the sensitivity analysis, which can be requested by the submission agencies or the internal committees. As per our case study results, we observe that the statistical inferences in Report 2 are close to the statistical inferences from the main analysis Report 1.

REFERENCES

- [1] Guideline on Missing Data in Confirmatory Clinical Trials, 2 July 2010 EMA/CPMP/EWP/1776/99 Rev. 1 Committee for Medicinal Products for Human Use (CHMP).
- [2] Roderick J. Little, Ph.D., Ralph D'Agostino, Ph.D., et al., The Prevention and Treatment of Missing Data in Clinical Trials, October 4, 2012, N Engl J Med 2012; 367:1355-1360
- [3] James D. Dziura, Lori A. Post, Qing Zhao, Zhixuan Fu, and Peter Peduzzi, Strategies for Dealing with Missing Data in Clinical Trials: From Design to Analysis, Yale J Biol Med. 2013 Sep; 86(3): 343–358.

[4] National Research Council 2010. The Prevention and Treatment of Missing Data in Clinical Trials. Washington, DC: The National Academies Press.

[5] Berglund, Patricia and Heeringa, Steven, 2014. Multiple Imputation of Missing Data Using SAS®. Cary, NC: SAS Institute Inc.

ACKNOWLEDGEMENTS

The authors would like to thank the management team for their encouragement and review of this paper.

CONTACT INFORMATION

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

Pennidhi Karlakunta
Merck & Co., Inc. USA
351 N Sumneytown Pike
North Wales, PA 19454
Work Phone: +1 (267) 305 2899
Work email: pennidhi.karlakunta@merck.com

Naveen Kommuru
Merck & Co., Inc. USA
351 N Sumneytown Pike
North Wales, PA 19454
Work Phone: +1 (267) 305 3740
Work email: naveen.kommuru@merck.com

Brand and product names are trademarks of their respective companies.