

Multiple Imputation: Better than Single Imputation in Pain Studies?

Ashik Chowdhury, Cytel Statistical Software & Services Pvt. Ltd., Pune, India

ABSTRACT

Missing data is one of the genuine challenges we face during analysis of clinical trials. The main impact of missing data is that it can infuse bias in results which reduces the chance of getting the appropriate interpretation. Hence proper knowledge of techniques for handling missing data is crucial.

A common method of handling this problem is by imputing missing values. There can be single or multiple imputations. While single imputation methods like LOCF or BOCF have been widely used, multiple imputation is of a comparatively recent origin but is gaining popularity. The main reason behind its popularity is it overcomes some of the drawbacks of single imputation. The current paper presents comparative effectiveness of each with the help of real life cases involving pain studies.

INTRODUCTION

The most serious concern of missing data is that it can introduce bias into the results of a clinical study, particularly if the data are missing not at random. Missing data causes the loss of information and in turn loss of statistical power. So the major challenge for these pain studies becomes interpreting the results. It is often preferred to put maximum effort at the design stage to achieve complete capture of all data to minimize the bias and maximize the power (e.g. continue collecting pain outcomes even after subject discontinuation). However, even with the best designed study missing data is still a potential issue. As such, various imputation methods are often employed when handling data of this nature.

There are a few standard procedures to impute missing data. Single imputation usually identifies a particular record for a subject, e.g. baseline or just the previous non-missing value and repeats it for the missing data points. Multiple imputation uses a predicted value for a given subject and time point using statistical modelling of available data.

Osteoarthritis (OA) is a painful and debilitating musculoskeletal disease that is characterized by intra-articular (IA) inflammation. In these studies pain is typically measured on a daily basis following study drug injection for a prespecified length of time. Pain intensity is generally measured by self-assessment with validated methods such as the visual analogue scale (VAS) or numeric al rating scale (NRS, 0-10). Depending on the length of follow-up of these studies and the fact that pain assessments are collected daily, one potential problem is the frequent occurrence of missing data.

This paper will illustrate several different imputation methods that are possible for handling missing data in pain studies. It will also explore the challenges faced in both of the imputation techniques, single and multiple, using a comparative study.

REGULATORY GUIDANCE

There is no universally applicable method of handling missing data recommended by any regulatory body. Existing regulatory guidance does not have specific instruction on how to address the problem of missing data. Realizing this gap the U.S. Food and Drug Administration (FDA) created a panel on Handling of Missing Data in Clinical Trials. This panel prepared a report which summarized how to reduce the amount of missing data and appropriate statistical methods to address missing data for analysis. Table 1 displays some selected recommendations from that report (The Prevention and Treatment of Missing Data in Clinical Trials, Panel on Handling Missing Data in Clinical Trials; National Research Council)

PhUSE 2014

Table 1: Panel on Handling Missing Data in Clinical Trials Recommendations

Parameters	Recommendations
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
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, and this information should be recorded and used in the analysis
Technique to limit the amount of missing data	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
Data analysis	Statistical methods for handling missing data should be specified by clinical trial sponsors in study protocols, and their associated assumptions stated in a way that can be understood by clinicians
	Single imputation methods like last observation carried forward and baseline observation carried forward should not be used as the primary approach to the treatment of missing data unless the assumptions that underlie them are scientifically justified
	Sensitivity analyses should be part of the primary reporting of findings from clinical trials. Examining sensitivity to the assumptions about the missing data mechanism should be a mandatory component of reporting

There are several other documents which describe the techniques of handling missing data in clinical trials. Below are the three recent documents which are currently being used.

- ❖ Draft Guidance on Important Considerations for When Participation of Human Subjects in Research Is Discontinued, from the Office for Human Research Protections in the U.S. Department of Health and Human Services (2008)
- ❖ Guidance for Sponsors, Clinical Investigators, and IRBs: Data Retention When Subjects Withdraw from FDA-Regulated Clinical Trials, from the U.S. Food and Drug Administration (2008)
- ❖ Statistical Principles for Clinical Trials; Step 5: Note for Guidance on Statistical Principles for Clinical Trials, from the European Medicines Evaluation Agency (EMA) International Conference on Harmonisation (ICH) (2009) E9 (Statistical Principles for Clinical Trials)

MISSING DATA TYPES

Missing data can be classified in three main categories i) Missing Completely at Random (MCAR) ii) Missing at Random (MAR) and iii) Missing not at Random (MNAR).

MISSING COMPLETELY AT RANDOM

Missing data will be called missing completely at random (MCAR) if the probability of missingness does not depend on observed or unobserved measurements.

$$P(M_i | Y_i^{obs}, Y_i^{mis}) = P(M_i)$$

Where M_i is the i^{th} observation being missing, Y_i^{obs} and Y_i^{mis} are the i^{th} observed and unobserved record respectively.

The nice property of MCAR is that the analyses performed on the data are unbiased. The power may be lost, but the bias in parameter estimates due to missing data gets reduced. But in real life situation MCAR is a rare case.

PhUSE 2014

Example:

- The responder is not able to answer on a questionnaire due to an accident or other randomly occurring event
- Breaking of laboratory instruments because of an accident

Table 2: Missing Completely at Random (MCAR)

Observation	Gender	Pain Score	Severity	Treatment	Missing
1	F	X	Mild	A	N
2	M	X	Worst	A	N
3	M	.	Moderate	A	Y
4	F	.	Moderate	A	Y
5	F	X	Mild	A	N
6	F	X	Mild	A	N
7	F	X	Moderate	A	N
8	M	.	Mild	A	Y
9	F	.	Moderate	A	Y
10	M	.	Worst	A	Y

Missingness neither depends on gender nor severity

MISSING AT RANDOM

Data are said to be missing at random (MAR) if, given the observed data, the probability of missingness depends on the observed data and does not depend on the unobserved data.

$$P(M_i | Y_i^{obs}, Y_i^{mis}) = P(M_i | Y_i^{obs})$$

This assumption of missing data implies that the characteristics of the missing data can be predicted from the observed variables and thus the unbiased response can be estimated. For example, suppose a participant drops out from the study due to lack of efficacy. This poor efficacy is reflected by a series of data before the participant drops out. Now if we have to choose a most appropriate value for the subsequent efficacy endpoint, it can be easily calculated using the observed data.

Many standard analyses (e.g., likelihood method) of missing data operate under MAR assumption.

Example:

- A participant refuses to answer questions concerning any of his/her personal information which is confidential (like income) and there is no correlation between this missingness and the subject's gender or the responses to the personal information being requested.
- Young participants are more likely to refuse to fill out the pain survey, but it does not depend on the level of their pain intensity.

Table 3: Missing at Random (MAR)

Observation	Gender	Pain Score	Severity	Treatment	Missing
1	F	X	Mild	A	N
2	M	.	Worst	A	Y
3	M	.	Moderate	A	Y
4	F	X	Moderate	A	N
5	F	X	Mild	A	N
6	F	X	Mild	A	N
7	F	X	Moderate	A	N
8	M	.	Mild	A	Y
9	F	X	Moderate	A	N
10	M	.	Worst	A	Y

Pain score is missing only when Gender is Male

MISSING NOT AT RANDOM

If missingness or dropout depends on the value of missing variables after conditioning on the observed variables, MCAR and MAR does not hold, then this missingness mechanism is said missing not at random (MNAR).

Example:

- In a hypertension study the blood pressure is measured less frequently for patients with lower blood pressure.
- In a question related to income, higher level income group people are less likely to report their income.
- Participants with severe pain, or having side effects from the experimental medication, are more likely to be missing at the end.

Table 4: Missing not at Random (MNAR)

Observation	Gender	Pain Score	Severity	Treatment	Missing
1	F	X	Mild	A	N
2	M	.	Worst	A	Y
3	M	X	Moderate	A	N
4	F	X	Moderate	A	N
5	F	X	Mild	A	N
6	F	X	Mild	A	N
7	F	X	Moderate	A	N
8	M	X	Mild	A	N
9	F	X	Moderate	A	N
10	M	.	Worst	A	Y

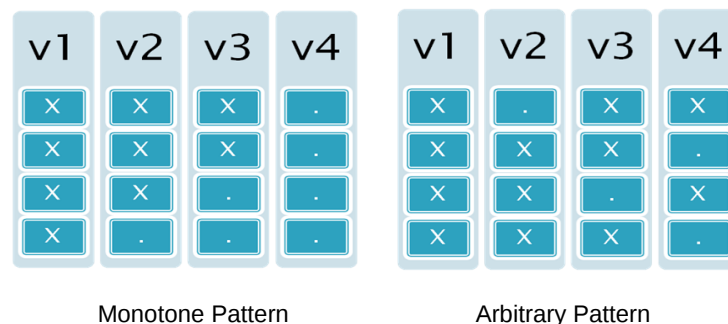
Pain score is missing when Gender is Male as well as Severity is worst

Missing data type MCAR and MAR are sometimes called as **ignorable** missing data whereas MNAR is called as **Non-Ignorable**. Life becomes more complicated when the dataset is MNAR. The reason is that the only way to get an unbiased estimate in this case is to model that missingness which leads to use of a complex statistical model.

MISSING DATA PATTERNS

The missing data pattern is one of the most important considerations in planning the imputation for missing data. The pattern of missingness can be classified as **monotone** or **arbitrary**. A dataset with variables $Y_1, Y_2, \dots, Y_n$ is said to have a **monotone** pattern when a variable Y_j is missing for a particular individual implies that the subsequent variables $Y_{j+1}, Y_{j+2}, \dots, Y_n$ are also missing. Alternatively if Y_j is observed then all previous variables $Y_{j-1}, Y_{j-2}, \dots, Y_1$ are also observed for that particular individual. If missingness occurs in a random fashion (i.e., for any variable and any participant) then the data set is said to have an **arbitrary** missing pattern. Figure 1 describes these common missing patterns.

Figure 1: Pattern of Missingness



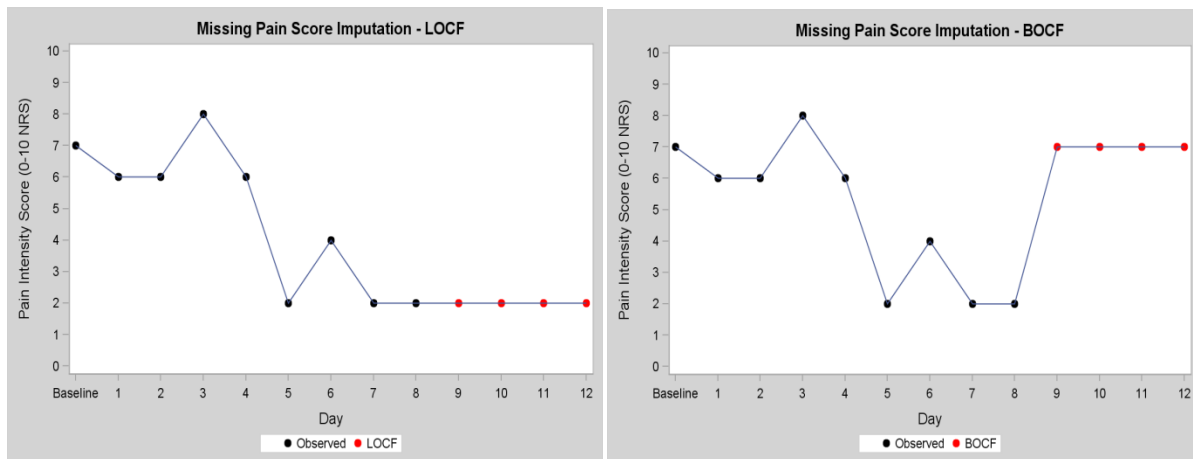
COMMON IMPUTATION STRATEGIES

Single-value imputation is widely used as an imputation method for pain studies. In this approach a missing value is replaced with a definite value based on some standard methods. For example, each missing value can be imputed using the mean of the observed data of that variable or it can be imputed by carrying forward the last observation (LOCF) or baseline observation (BOCF) for that particular individual. For pain studies, LOCF and BOCF methods are often considered the most conservative methods of handling missing data, particularly in cases where subject discontinuation is due to lack of efficacy. However, one potential drawback of these methods is that the uncertainty associated with the estimate of missing data cannot be assessed and the process might provide biased estimates with either underestimation or overestimation of variance. Multiple imputation, introduced by Rubin (1987) addresses the concerns of this uncertainty as the method rely on development of multiple estimates for each of the missing value. In the section below we will discuss last-observation-carry-forward (LOCF), baseline-observation-carry-forward (BOCF) and Multiple Imputation technique in detail.

LOCF & BOCF

It has been a common practice in pain studies to use LOCF or BOCF to impute missing values. In LOCF, a missing value is replaced by the last non-missing observation for that individual. This method is based on a strong assumption that the outcome of participant does not change after drop out from the study which is questionable in many settings. Another frequently used method to fill the pain score is BOCF, where baseline observation is used in place of missing score for each individual. This method assumes that a participant's pain score remain same as it was measured before study drug administration. This assumption might not always be true since most patients in pain studies improve substantially from baseline over time. One of the advantages of these methods is that they are very easy to implement. There is a misconception about LOCF and BOCF methods that they result in estimates of treatment effect that trend to favour control or study drug which is not always true. Figure 2, below describes missing data imputation using LOCF and BOCF.

Figure 2: Imputing Missing Pain Intensity Score (0-10 NRS) using LOCF and BOCF



FDA has strong recommendation (The Prevention and Treatment of Missing Data in Clinical Trials, National Research Council, 2010) not to use LOCF or BOCF as primary approach for handling missing data. *“Single imputation methods like last observation carried forward and baseline observation carried forward should not be used as the primary approach to the treatment of missing data unless the assumptions that underlie them are scientifically justified.”*

MULTIPLE IMPUTATION

The multiple imputation method was introduced by Donald Rubin (1987). In this technique, instead of filling the missing value with a single value, missing value is replaced with a set of possible values. Let us consider two variables –

X = Observed in every unit and Y = Observed on some units, $Y = (Y_M, Y_O)$, where Y_M and Y_O represents missing and observed records for variable Y .

Study the relationship between X and Y - use the data from units where Y and X are observed together with rest of the X 's. If $\tilde{X}$ represents the vector of X values from individuals with missing Y 's, we use this relationship to complete the

data set by drawing the missing observations randomly from $Y_M | \tilde{X}$. We repeat this m times to get m number of complete data sets. Then these m complete data sets are analyzed using the standard statistical methods and the results are combined. Now the question is how to combine these results? Here is the answer.

- ▶ The point estimate is the average of m repeated complete-data estimates

$$\hat{\theta}_{MI} = \sum_{k=1}^m \hat{\theta}_k / m$$

- ▶ Within imputation variance is the average variance within the imputed data sets

$$\hat{\sigma}_w^2 = \sum_{k=1}^m \hat{\sigma}_k^2 / m$$

- ▶ The between imputation variance is the variance across the m imputed data sets

$$\hat{\sigma}_b^2 = \sum_{k=1}^m (\hat{\theta}_k - \hat{\theta}_{MI})^2 / (m-1)$$

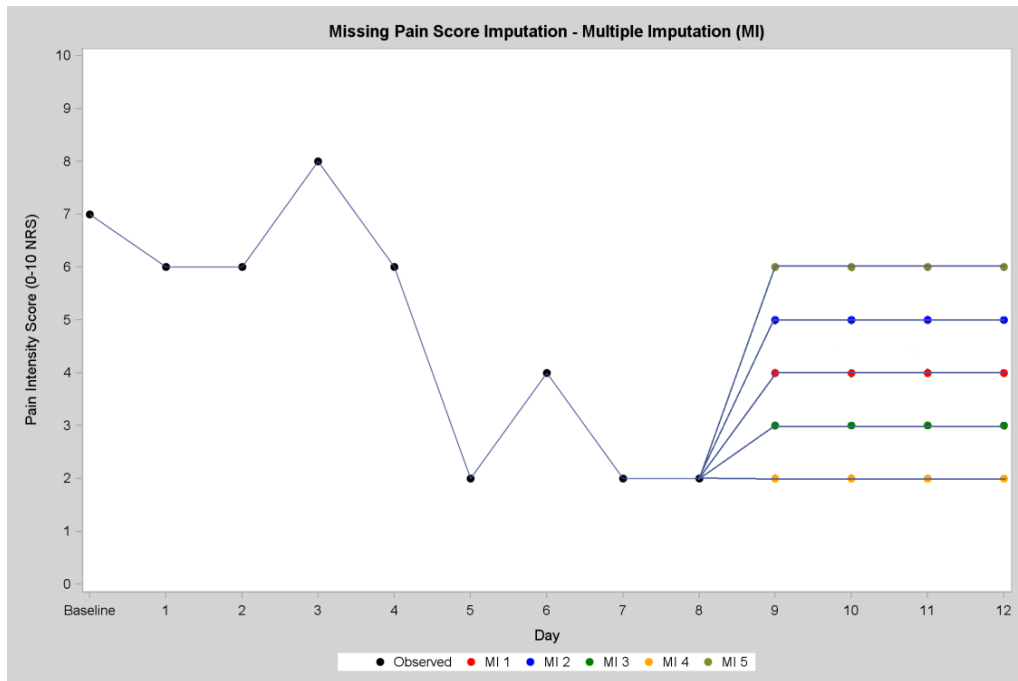
- ▶ And, total variance is the sum of the within and between variances

$$\hat{\sigma}_{MI}^2 = \hat{\sigma}_w^2 + (1 + m^{-1}) \hat{\sigma}_b^2$$

So the estimated standard error of $\hat{\theta}_{MI}$ is $\hat{\sigma}_{MI}$

In the illustration of multiple imputation method in figure 3, below it is assumed that from week 9 onwards the missing values are randomly replaced by a single number for each of the weeks. For example, for the first multiple imputation dataset (MI 1) all of the missing pain scores from weeks 9 to 12 are imputed by 4 and similarly for other imputation number as well. In reality it may take different value for different weeks.

Figure 3: Imputing Missing Pain Intensity Score (0-10 NRS) using Multiple Imputation



When performing multiple imputation, thought should be given to the number of multiple imputations, m , to be implemented. Theoretically m can take any value which is greater than 1 ($m > 1$). Historically, m typically takes the values $3 \leq m \leq 10$. We can determine this m using **Relative Efficiency** (RE) described by Rubin (1987) as

$$RE = \left(1 + \frac{\gamma}{m} \right)^{-1}$$

Where γ is the rate of missing information for the quantity being estimated. In an ideal case the imputation number should be chosen such that RE is close to 1. The table below shows how m changes with the change in the missing percentage. It is very clear from the table that m is directly proportional to γ .

Table 5: Relative Efficiency (RE)

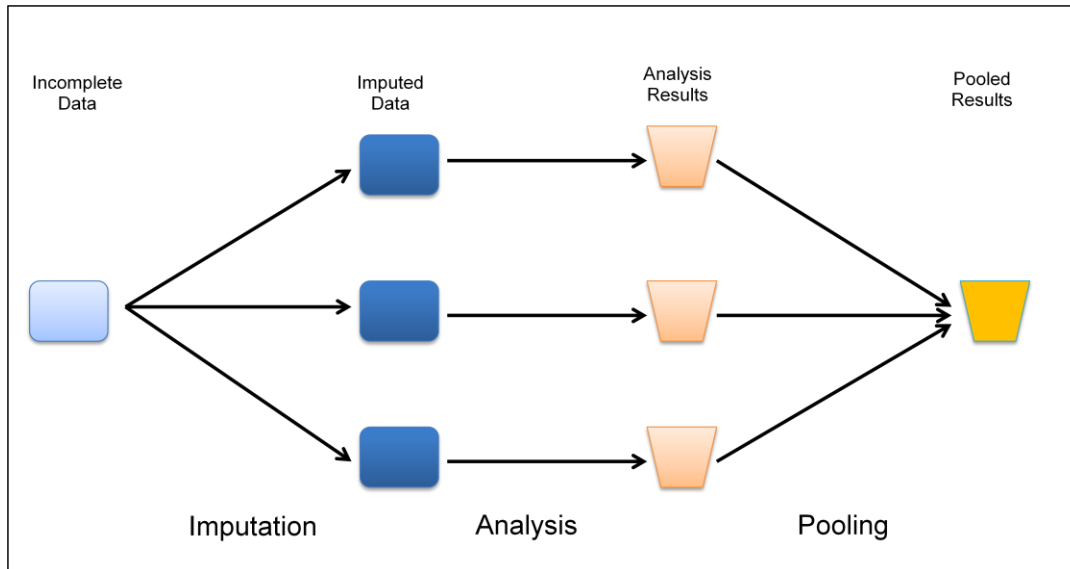
	γ				
m	10%	20%	30%	50%	70%
3	0.9677	0.9375	0.9091	0.8571	0.8108
5	0.9804	0.9615	0.9434	0.9091	0.8772
10	0.9901	0.9804	0.9709	0.9524	0.9346
20	0.9950	0.9901	0.9852	0.9756	0.9662

[Source: <http://sites.stat.psu.edu/~jls/mifag.html#minf>]

SAS® PROCS FOR MULTIPLE IMPUTATION

The whole multiple imputation procedure can be divided in to three parts as shown in the following figure. For each of these three steps there are SAS® PROCS which are also described in this section.

Figure 4: Schematic Diagram of Multiple Imputation Technique



Imputation: Fill in missing values m times from the distribution of observed data. Below is a very basic sample SAS[®] code for imputation.

There are many other options available in PROC MI.
(visit <http://support.sas.com/documentation/onlinedoc/stat/132/mi.pdf>)

```
PROC MI
  DATA = /* Input Data */
  OUT = /* Output Data */
  SEED = /* specifies seed to begin random number generator */
  NIMPUTE = /* Specifies number of Imputations */;
RUN;
```

Several methods are available in PROC MI. The choice of these methods depend on pattern of missing data and type of variables to be imputed. Table 6 shows some of the available methods in PROC MI.

Table 5: Imputation Methods - PROC MI

Missing Data Pattern	Type of Variables to be Imputed	Imputation Methods Recommended	SAS [®] Options in PROC MI
Monotone	Continuous	Regression method	MONOTONE REG
		Predicted mean matching	MONOTONE REGPMM
		Propensity score	MONOTONE PROPENSITY
Monotone	Categorical (Ordinal)	Logistic regression	MONOTONE LOGISTIC
	Categorical (Nominal)	Discriminant function method	MONOTONE DISCRIM
Arbitrary	Continuous	MCMC full data imputation	MCMC IMPUTE = FULL
		MCMC monotone data imputation	MCMC IMPUTE = MONOTONE

Analysis: Analyze multiple (m) complete datasets by imputation using standard statistical methods (PROC REG, PROC GLM, PROC MIXED, PROC GENMOD, etc.).

Pooling: Combine the m estimates into a single result using the SAS[®] MIANALYZE procedure. Below is sample SAS code –

```
PROC MIANALYZE;
  BY; /* Specifies group variables */
  CLASS; /* Lists the classification variables */
  MODELEFFECTS /* Lists the effects to be analyzed */;
  STDERR /* lists the standard errors associated with the effects in the
          MODELEFFECTS statement */;
RUN;
```

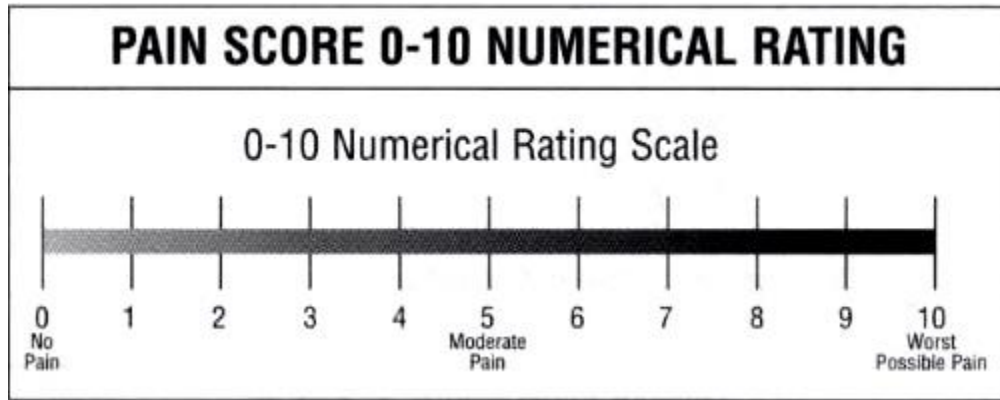
SIMULATING DATA

To evaluate the performances of the imputation techniques so far discussed in this paper data are simulated for a pain intensity score (0-10 NRS) in a Osteoarthritis (OA) trial. A brief description of OA and the NRS is given below.

Osteoarthritis is a painful and debilitating musculoskeletal disease that is characterized by intra-articular (IA) inflammation, deterioration of articular cartilage, and degenerative changes to peri-articular and subchondral bone (Creamer and Hochberg, 1997; Goldring and Goldring, 2006). Symptomatic OA is estimated to affect 9.6% of men and 18.0% of women worldwide, is generally more prevalent in Europe and the USA than other parts of the world, and is expected to become the fourth leading cause of disability by the year 2020 due to increases in life expectancy and aging populations (Woolf and Pfleger, 2003).

The pain numeric rating scale (NRS) is a generic, unidimensional questionnaire of pain intensity in adults. In the NRS, a respondent selects a whole number (0-10 integers) on a single 11-point scale that best reflects the intensity of her/his pain. In this scale, the 0 is anchored with the descriptor “no pain” and the 10 is anchored with the descriptor “worst pain imaginable” (Hawker GA et al, 2011).

Figure 5: Numeric Rating Scale (NRS, 0-10)



The following design was chosen for simulation of trials to be considered as case studies in this paper.

A randomized, placebo-controlled, parallel-group, single dose study in patients with Osteoarthritis of the knee. Total of 200 patients with knee OA are randomized (1:1) with either Active or Placebo. Each patient is evaluated for a total of 12 weeks following a single dose of study drug.

Altogether 100 trials each with 200 patients were simulated, where pain score are collected on a daily basis and then weekly averages are computed. After simulating the data for pain score (forming complete data), data were deleted according to a predefined random procedure to create trial data sets with missing data which will be henceforth called incomplete data in this paper. The data has been deleted in such a way that after the first deleted observation all subsequent observations were also deleted (e.g. for a patient if data was deleted for day i , then day $i+1$ onwards data was also deleted). This was done to create a monotone missing pattern (e.g. missing data is due to withdrawal of patients). The withdrawal rate in this simulation was assumed to be 35%.

ANALYSIS

The primary endpoint for this study is change from baseline to week i ($i = 1$ to 12) of weekly mean of the average daily (24-hr) pain intensity score. This endpoint is analyzed with a longitudinal mixed effects model with fixed effects for treatment study week, treatment-by-week interaction, and baseline as a covariate. Subject will be the random effect. Treatment differences from placebo will be presented, and estimated via least squares means from the analysis model along with 90% confidence intervals and associated 1-sided p-values.

The main focus for this paper was the primary endpoint for the case study considered. All the analyses were done for complete data, incomplete data and imputed (using LOCF, BOCF and MI) datasets.

RESULTS

The results of complete data analysis were considered as reference to compare with others. Table 6 displays the simulation results for primary endpoint for each treatment arm. The table reports the ranges of mean change from baseline and standard error (SE) over 12 weeks for 100 trials.

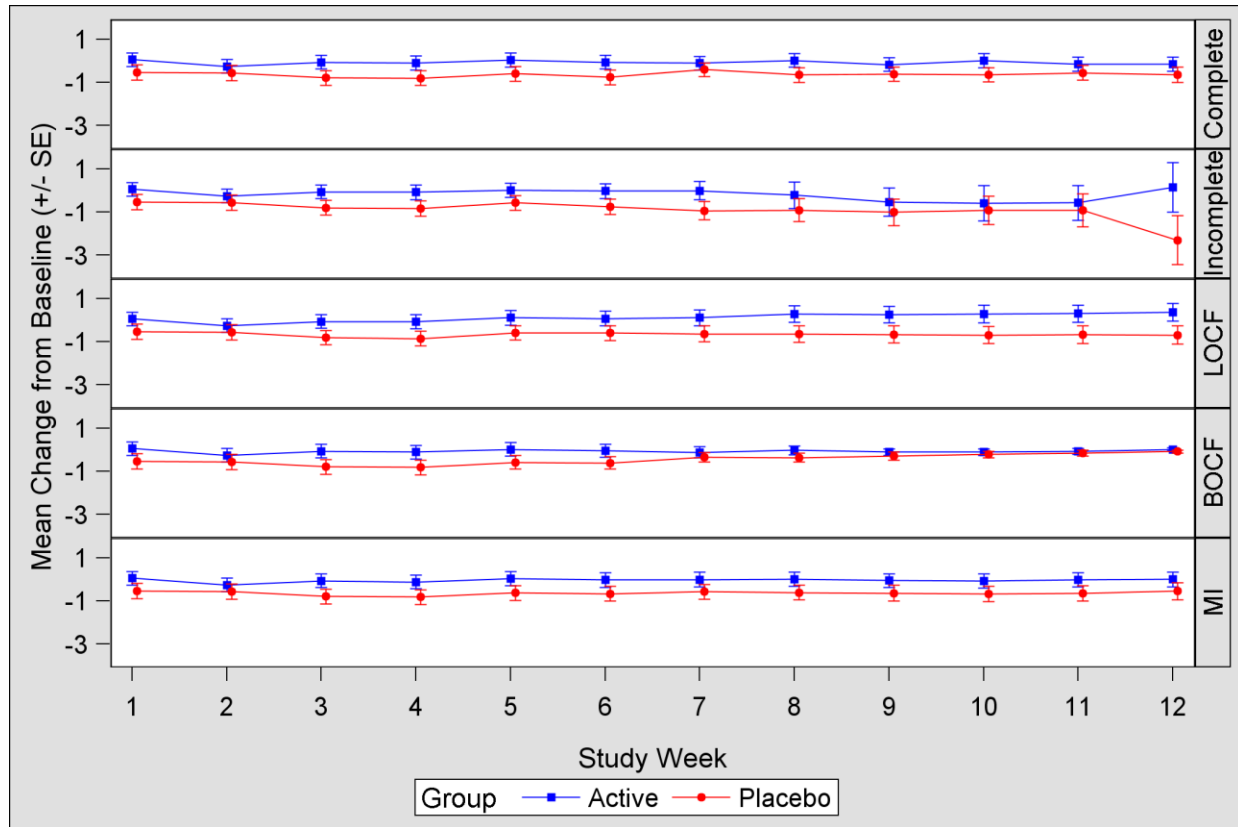
Table 6: Mean and SE – Change from Baseline in Weekly Average Daily Pain Intensity Score (0-10 NRS)

Analysis Performed	Active		Placebo	
	Mean	SE	Mean	SE
Complete Data	-1.12 to 0.83	0.27 to 0.39	-1.17 to 1.05	0.29 to 0.39
Incomplete Data	-2.72 to 1.93	0.29 to 1.61	-3.83 to 2.71	0.29 to 1.52
Imputed Data using LOCF	-1.12 to 1.00	0.30 to 0.51	-1.10 to 1.00	0.29 to 0.51
Imputed Data using BOCF	-1.04 to 0.79	0.03 to 0.39	-1.10 to 0.98	0.04 to 0.40
Imputed Data using MI	-1.11 to 0.87	0.29 to 0.41	-1.10 to 1.07	0.29 to 0.40

PhUSE 2014

The figure 6 below represents mean change from baseline ($\pm$ SE) in weekly average of daily pain intensity score (0-10 NRS) over time using complete data, incomplete data (missing completely at random), and applying LOCF, BOCF, and MI to the missing data. The SEs are quite similar for all the cases except incomplete data where it shows higher value for both the treatment arms. On average MI provided estimates very close to the complete data. The trend displayed in the below figure is consistent with rest of the trials.

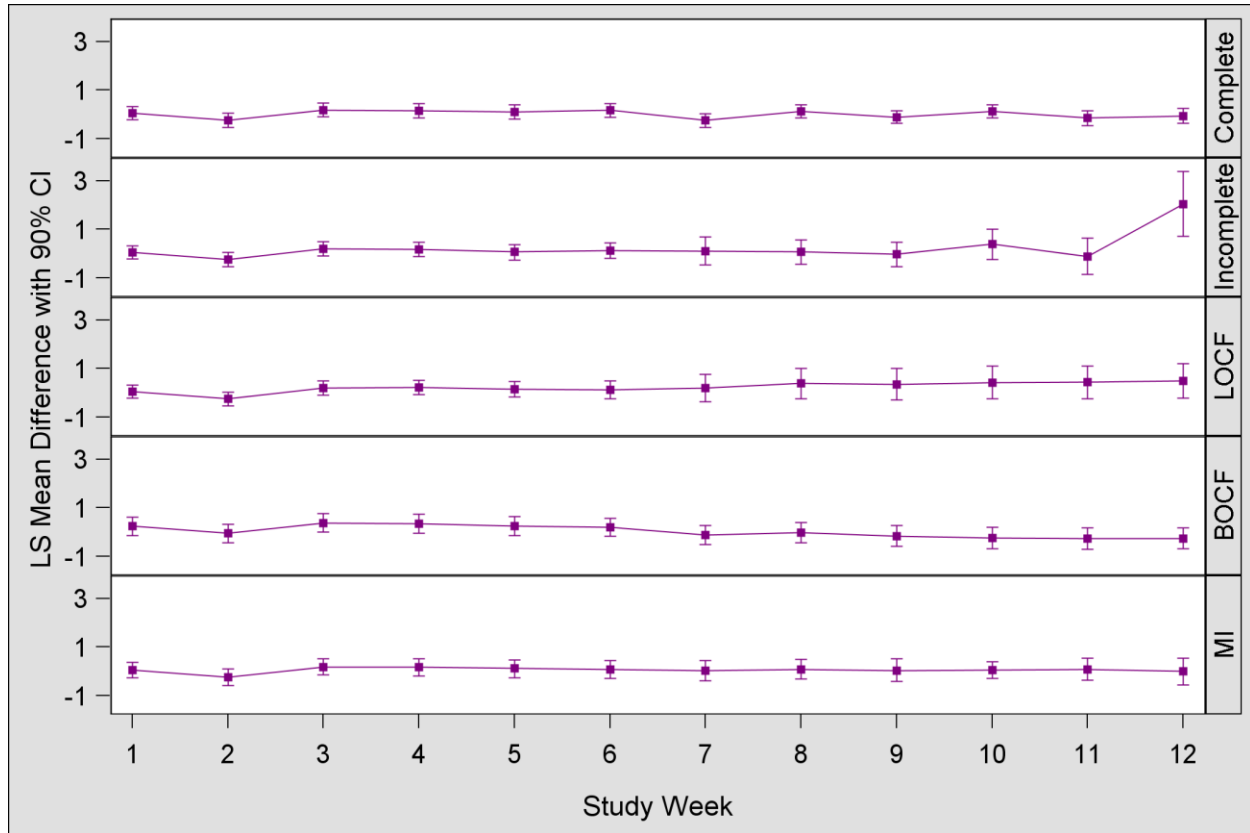
Figure 6: Mean Change from Baseline ($\pm$ SE) in Weekly Average of Daily Pain Intensity Score (0-10 NRS) Over Time (for a single simulation)



To test the treatment effect a linear mixed effect model analysis was carried out on the primary endpoint i.e., change from baseline of weekly average of daily pain intensity score. Figure 7 presents treatment differences from placebo which is estimated via least squares means (along with 90% confidence intervals) from the analysis model. In terms of treatment difference all the imputation methods (BOCF, LOCF and MI) are showing results close to complete data. The trend shown in the below figure are following consistent results in all other trials.

PhUSE 2014

Figure 7: LS Mean Difference from Placebo in Change from Baseline (with 90% CI) in Weekly Average Daily Pain Intensity Score (0-10 NRS) Over Time (for single simulation)



A test was included to comparing the imputation methods. For each of 100 simulations, the continuous response (LS Means change from baseline of weekly average of daily pain intensity score) is analyzed through an ANOVA with fixed factors week, treatment and methods. The Complete data analysis method is considered as reference to which the rest of the methods are compared. These data were analyzed with PROC GLIMMIX from SAS[®] 9.2, assuming equal variances across the methods.

Table 6 displays number of cases where significant differences between average LS Means (across the visits) for complete data versus imputed data are observed. All the imputation methods except MI show statistically significant differences in LS means relative to the complete data.

Table 6: Significant LS Mean Differences from 100 Trials

Group	Comparison	Number of Cases p-value < 0.1
Active	Complete vs Incomplete	27
	Complete vs LOCF	37
	Complete vs BOCF	33
Placebo	Complete vs Incomplete	28
	Complete vs LOCF	38
	Complete vs BOCF	35

DISCUSSION

Although this exercise only employed 100 simulated samples, results of this analysis show estimates of change in pain from multiple imputation are consistently close to results obtained from the complete data sets. However, this is not to say that MI is always the best imputation method to be used over LOCF or BOCF, but rather that MI should also be considered when determining imputation techniques for clinical trials where pain data are analyzed. When designing the analysis for pain studies, careful consideration should be given to the study design in order to determine which imputation method should be employed as sensitivity analysis to the primary results. One can also use this paper as a reference for how multiple imputation can be performed in pain studies. Since in this paper I have only focused on a particular type of pain study with a specific missing pattern, more research is warranted to further explore the use of multiple imputation in the setting of pain studies.

REFERENCE

- Creamer P, Hochberg MC. Osteoarthritis. *Lancet* 1997;350(906):503–508.
- Goldring SR, Goldring MB. Clinical aspects, pathology and pathophysiology of osteoarthritis. *J Musculoskeletal Neuronal Interact* 2006;6(4):376–378.
- Hawker GA, Mian S, Kendzerska T, French M. Measures of Adult Pain. *Arthr Care Res.* 2011;63(S11):S240–S252.
- European Medicines Agency. 2010. “Guideline on Missing Data in Confirmatory Clinical Trials. 2 July 2010. EMA/CPMP/EWP/1776/99 Rev. 1”. Available at http://www.ema.europa.eu/docs/en_GB/document_library/Scientific_guideline/2010/09/WC500096793.pdf
- 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 Academic Press.
- Rubin, D.B. (1987). Multiple imputation for nonresponse in surveys. New York: John Wiley & Sons, Inc.
- Woolf AD and Pfleger B. Burden of major musculoskeletal conditions. *Bull World Health Organ* 2003;81:646-656.

ACKNOWLEDGMENTS

The author would like to express thanks to Chris Schoonmaker and Aniruddha Deshmukh who reviewed the manuscript and provided valuable feedback.

RECOMMENDED READING

Clinical Trials with Missing Data: A Guide for Practitioners By Michael O'Kelly, Bohdana Ratitch

CONTACT INFORMATION

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

Ashik Chowdhury
Cytel Statistical Software & Services Pvt. Ltd
6th Floor, Lohia-Jain IT Park – A Wing, Survey #150, Paud Road, Kothrud,
Pune 411 038, India
Work Phone +91 (20) 6709-0255 Cell: +91 8806311426
Ashik.chowdhury@cytel.com
www.cytel.com

SAS and all other SAS Institute Inc. product or service names are registered trademarks or trademarks of SAS Institute Inc. in the USA and other countries. ® indicates USA registration.