

Pooling Clinical Data: Key points and Pitfalls

Florence Buchheit, Novartis, Basel, Switzerland

ABSTRACT

Pooling of clinical data is used by all Pharmaceutical companies. Submission to Health Authorities of a new compound is one of the main reasons to pool data but there are multiple of others needs for pooling data such as ongoing safety review and publications.

The concept of pooling data is neither new nor defined by a pre-determined set of rules or methodologies. Inevitably pooling requirements are project/compound specific and are driven by the objective of the pooled data.

There is no obvious right or wrong way to pool data but the intention of this paper is to provide guidance on the creation of pooled datasets with general information, typical pitfalls and key points in order to obtain a robust pool. It will cover topics such as planning, data standards, programming strategies, validation and documentation.

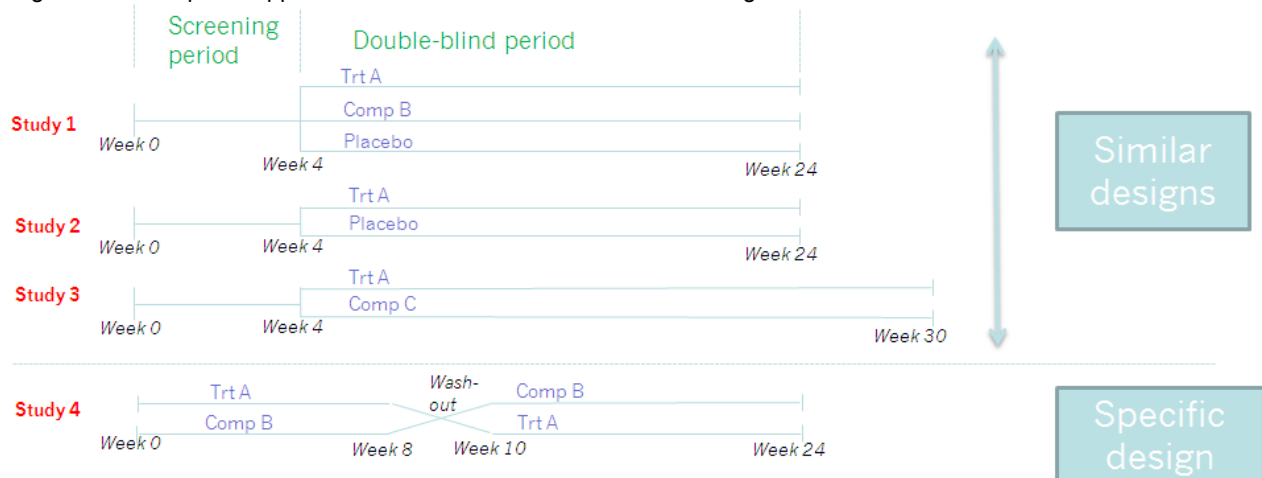
INTRODUCTION

The pooling of data can have a number of benefits like:

- showing safety profile across whole program
- improving the precision of incidence estimate for main efficacy analysis
- identifying rare safety signals
- exploration of possible drug-demographic, drug-disease or drug-drug interactions in subgroups of populations.

The quality of the pool is however dependent on the individual trials used as pooling may obscure real potentially meaningful differences between studies.

Figure 1 – Example of approach for studies with different/similar designs



In the figure 1 above it makes sense to pool the 3 first studies as designs are similar even if the lengths of studies are not identical. Team may decide in this pool to select data up to week 24 for the study 3 if this is more appropriate. Study 4 is a really specific trial and it would be really difficult to interpret results if data from this trial is pooled with the 3 others.

PhUSE 2012

There are a number of HA/ICH guidances that discuss integration of data across trials but there are rarely any fixed rules. Therefore, for each integrated analysis it is important to ask the question first of all - “do we need to pool at all?” As guidances change over time, it is important that teams keep up to date with the latest industry information.

PLANNING

Planning should not just focus on timelines and/ or milestones but also on a coherent pooling strategy where rationale is well defined and understood by all the relevant stake holders. To aid this, good documentation of any strategies or decisions made is key.

POOLING STRATEGY

Planning should start early to understand the future pooling needs. It is likely that a pooling plan will evolve during the life of a clinical development plan. At the earliest stage of a phase III program, developing a detailed plan or strategy with a submission in mind is unlikely to yield a permanent strategy that will meet the needs of the project later down the line. Thinking ahead of frequency of pooling, future studies to add, cut-off date is recommended for a good long-term strategy.

Findings from trials, changes in procedural guidance from health authorities or governing bodies, and other external or internal influences may mean the strategy needs to be revised or refined to answer questions or support arguments not considered at the start.

Figure 2 – Example of different strategies of pooling



OPERATIONAL CONSIDERATIONS

When planning the details of how and when to execute an integration/pooling exercise there are a number of logistic and operational factors that should be considered.

The timelines for the key milestones are often defined by the wider project team and depend on the deliverable.

These high level milestones should be used to create a more comprehensive project plan considering the specific timelines and milestones needed to ensure operational delivery. Below are a suggested number of milestones that could be considered in any plan.

- Develop, review and finalize pooling strategy – This may require less work for certain pooling activities such as DSUR, RMP updates, and more work if the integration is for a submission or filing
- Develop, review and finalize pooling plan
- Develop, review and revise statistical sections of integrated data reporting and analysis plan
- Develop, review and finalize programming specifications
- Creation and validation of pooled/integrated analysis datasets – this milestone may need to be split into more granular deliverables depending on approach. If an ongoing module approach is adopted – adding each study to a pool once it is locked-, timelines for preparing/pooling individual studies may be appropriate.
- Creation and validation of pooled data outputs
- Dry runs
- Final delivery

POOLING CONTENT

Studies may have to be pooled by indication, development program or across projects or therapeutic areas. They may be pooled by treatment comparator. These and which data (domains) are required to be pooled are primarily dependent on the deliverable and objectives of any analysis. Below is a list of data which are commonly pooled, however, some deviations may be required or reasonable depending on the focus of the analysis. The first step is to decide what studies the integrated database should contain as outlined above. The second step is to decide the relevant information to pool, this will commonly include:

- **SAFETY**
- Duration of treatment exposure
- Adverse events, SAEs (including deaths) and ADRs
- Specific defined risks of interest which are often defined based on adverse event SMQs but could also contain data other than AEs. These are often especially relevant for the RMP.
- Lab values, often with a focus on parameters that represent a specific safety marker significant to the indication or a drug class. There is often special interest in lab abnormalities in relation to clinically significant abnormality ranges

Further areas of potential interest, depending on indication, compound characteristics and data captured may include:

- Vital signs
- Specifically solicited events captured outside of standard panels (e.g. immunogenicity problems)
- ECG/EKG, x-ray, other relevant tests/assessments (e.g. MRI or neurological examinations)
- Adjudicated endpoints such as deaths, hospitalizations, Cardiovascular events, liver abnormalities or other endpoints that may be specific to an indication or drug class
- Concomitant medications: all or only selected subset, for example indication-related concomitant medications.
- Patient disposition
- Medical history – in particular may be relevant to identify defined risks of interest (Risk History)

- **EFFICACY**

The selection of which efficacy data to be pooled is very much dependent on the project. Selection should primarily focus on key endpoints that support the overall objectives of a program or submission or that may be used to make label or commercial claims or add clarity or additional weight to results seen across individual studies of a development program.

Frequently one of the objectives of pooling is to facilitate the identification of subgroups of patients who may get the most benefit. It often requires pooling of this data to increase the patient numbers in order to make any subgroup analysis feasible.

EXECUTION

A key point in the execution phase is the mapping of study level data to pool level. This process involves the comparison of the data structures (including metadata, code lists, endpoint and algorithms) from the studies with the final pooled data specification. The effort is likely to be more significant if a data pool includes older studies or studies that were developed outside of the project standard (e.g. some outsourced or locally run studies) as special attention is needed to ensure data are mapped correctly. Developing reports to do this comparison may be a more efficient and less error prone approach than performing manual review.

PROGRAMMING APPROACH: RAW VERSUS DERIVED

- **POOLING FROM DERIVED/ANALYSIS DATA**

Utilizing derived (analysis) datasets generated for CSR reporting can reduce work in re-deriving common endpoints or data points and help ensure consistency in derivations with the CSR.

However care should be taken to ensure that all derivations and algorithms are consistent across individual studies and with those needed for any pooled analysis. This assessment should not only focus on algorithms defining the derivation rules of given endpoints, for example AUC (area under the curve), but also on windowing or imputation rules for endpoints that handle missing data, such as LOCF.

- **POOLING FROM RAW DATA**

Pooling from raw data may be a suitable option in cases where there is significant variability in the data structures, standards or endpoint derivations across studies.

PhUSE 2012

The use of raw data for a pooled database will inevitably require the reprogramming of endpoints, variables and derivations already programmed at a study level, however, this may be a necessity if endpoints in the final pooling differ from those defined for one or multiple studies.

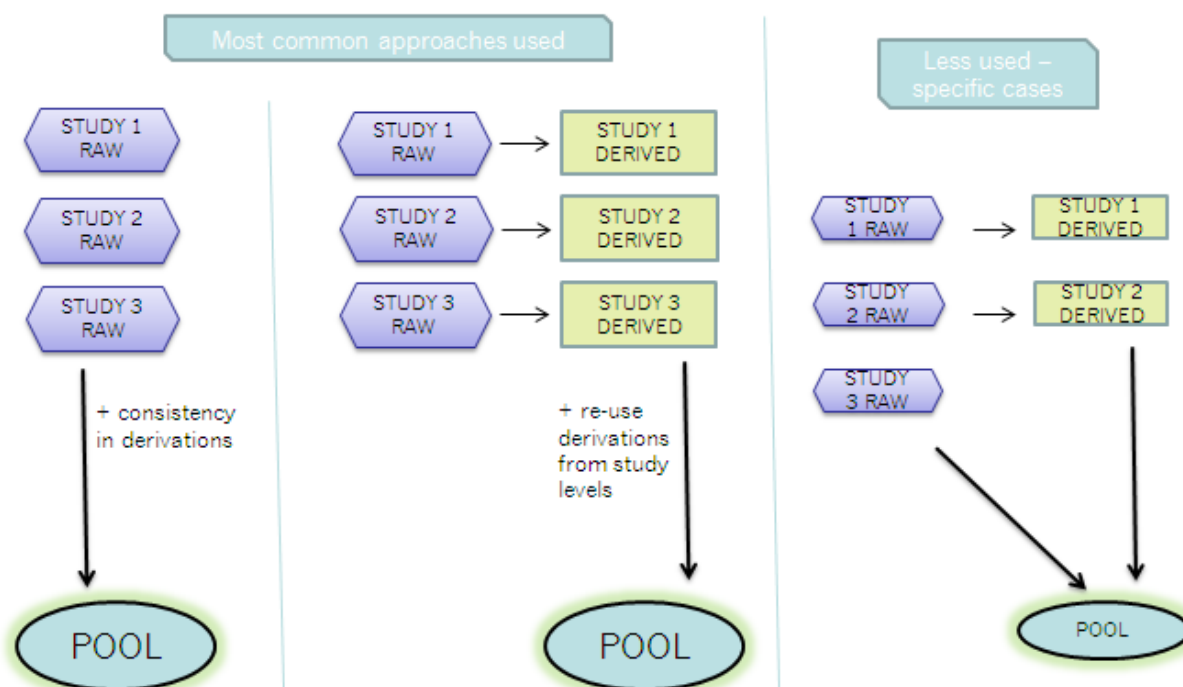
- POOLING FROM RAW AND DERIVED DATA

In special situations it might happen that raw and derived datasets together need to be pooled.

Such situations could arise when:

- some 'standard' studies have the required endpoints available in derived datasets but some 'non-standard' studies do not: we could use the derived datasets for 'standard' studies and use the raw datasets for 'non-standard' studies
- for an ongoing pool raw data might be used for the most recent study as the corresponding validated derived datasets are not yet available.
- in derived datasets not all patients are included, but for specific analyses we need all patients; in this case data for those missing patients could be taken from raw dataset.

Figure 3 – Summary of raw/derived approaches



EXECUTION STRATEGY

When considering the execution strategy as the technical implementation plan of pooling, it is advised that teams should look at how the pool will or could be used in the future.

A given data pool may be produced initially to support a specific deliverable (e.g. submission) but could then grow as more studies are additional to support a data warehouse / future exploratory analysis. This could guide whether a modular or "one hit" approach should be adopted.

- THE "ONE HIT" APPROACH

Often programs are developed at a data type level i.e. one program per analysis dataset that standardizes data for each study to a common structure with common derivations/algorithms and includes global derivations relevant to the data of all study. For example one program would be to create the pooled AE analysis dataset containing the code to pool all studies. Often this approach results in a program that contains large number of macros loops executing different code to process data for different studies as needed.

Benefits of this approach include:

- Reduced coding effort i.e. fewer programs (one per pooled dataset)
- Common code to process common data across studies via common derivations
- Changes can be applied to one line of code that would handle a number derivations across studies
- The approach is most appropriate to support a specific pooling event rather than an ongoing data pool.
- All the data handling is described in the programming code, and easy to confirm the handling.

PhUSE 2012

Possible risks or issues include:

- The code can become significantly complex especially if there is variability in the source data across studies
- This may be a sub-optimal approach if only data from certain studies are available at a given time as the program may need to be finalized and re-finalized as additional studies are added. This approach could result in programming the dataset for 4 studies for example, then adding another and needing to revalidate all of the code to ensure that code related to the new study does not impact previously validated results unexpectedly.
- This process also limits the ability to upscale resources if needed as only one person can work on a given program. This person would be responsible for pooling all studies for a given domain and may need to finish their activity before work could start on dependent datasets.
- **- THE “MODULAR” APPROACH**

This approach looks to split the pooling process in to multiple steps. One example of a modular approach might be the following.

i) Firstly standardizing the data at a study and/or domain (dataset) level. This would involve the standardization of data to meet the pooled dataset structure/code list needs. This may mean creating new variables, treatment groupings/codings, visit and code list mapping etc. to give a standardized “pre-pool” base data set for each domain and/or each study e.g. AE study 1, AE study 2, AE study 3 etc.

ii) Secondly the creation of programs that combine all the “pre-pool” standardized datasets for each study into one dataset and perform additional derivations at a wider level e.g. creating variables for reporting sub-group or populations (e.g. study length etc.).

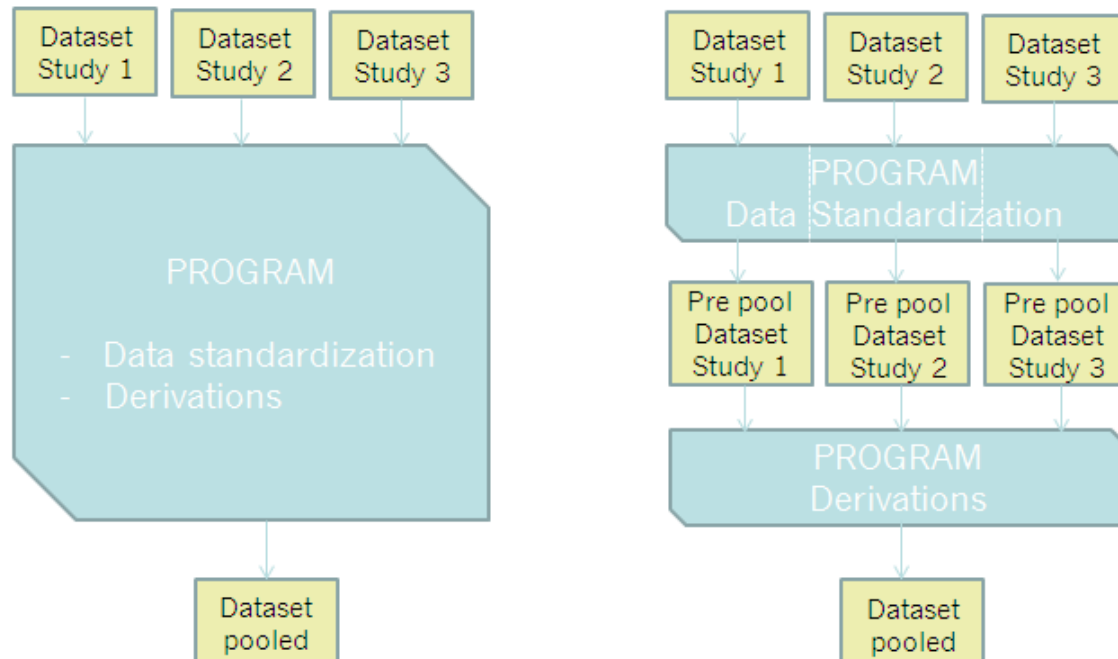
Benefits of this approach include:

- Allows modular working preparing studies as they become ready – “pre-pool” datasets can be finalized and fully validated once the clinical database for a new study to be added is locked, i.e. pooling preparation activities can proceed on an ongoing basis
- May support additional flexibility to resource – additional resource can be added if there is a critical need to work on the same data domains across studies
- A final complete pooled dataset specification is not essential as long as an interim specification is available defining core variables and data structures. Any reporting specific sub groups and derivations could be defined in a second stage plan and implemented in stage 2

Possible risks or issues include:

- There is potentially a larger code burden (more programs)
- Changes that affect multiple studies may require multiple programs to be updated (but should be minimized if macros are used)

Figure 4 - Summary of programming strategies



PhUSE 2012

A well thought out approach can help mitigate some of these risks/issues. The recommended approach is the development of utility macros to handle common derivations across studies, macros to handle processing of similar or identical study structures, the use of study level formats to map values from source to final variables. This approach may not be the most appropriate or efficient to handle pooling for a quick hit specific deliverable however standardizing of datasets from individual trials before pooling can help simplify the pooling process and reduce mapping and implementation errors greatly.

DOCUMENTATION AND VALIDATION CONSIDERATION

As with any deliverable, good documentation is essential. The tracking of the contents of any data pool can be beneficial later in the development cycle for a product. Whilst much of the detailed information regarding the contents of a data pool is defined in various reporting and analysis plans, tracking information regarding data pools centrally can provide an easier reference in cases where multiple requests come in that could be supported by an existing pool e.g. publications, HA questions, analyses to aid decision making.

Validation is a critical element in the execution of any pooling exercise and should be based on a risk assessment and according to a validation strategy. As such, teams should prepare for this well in advance and this aspect should be covered in any execution strategy. It is vital that validation efforts are made according to the pooling strategy and carried out thoroughly.

Whilst the basic concepts for validation of programming are no different to any other programming or reporting activity, the pooling of data from multiple studies can be complex and error prone. Verification in addition to standard programming validation is strongly advised.

- REPORTS TO SUPPORT REVIEW

As pooling often requires the derivation of common derived variables to enable the mapping of data of similar times (visits), time points etc to a common reportable way, validation may also require a formal review to ensure that data is mapped as expected and required.

A simple example is the mapping of study visits and assessment time points. This would be traditionally handled by creating an analysis visit variable. In some cases the programming definition of the derivation may contain a more general rule stating the final analysis visit values and stating that study level visit and time points should be mapped to this. This is especially likely to be the case if the pooling plan includes a large number of studies. Adding the exact mapping detail for each study in terms of variable values could make the plan huge, complex and cumbersome.

In an example like this it is advised that the source visit and time point variables are retained in the final dataset. This would support traceability as well as allow a simple report to be written to allow a review of the mapping from source to the pooled analysis variable across all studies. This mapping matrix type report is a useful way to review that mapping has been implemented as expected. This may also help flag where mappings represent an equivalence rather than an exact mapping, for example, where a 15 min pre-dose, 30 min pre-dose, non-time specific pre-dose time points collected at the study level are mapped to "pre-dose" for the purpose of pooled data reporting.

- CROSS CHECKS BETWEEN POOLED AND STUDY LEVEL DATA

Once data are pooled it is important to ensure that the pooled data represent the data from individual studies consistently where expected. Use of tools to track the population counts and treatment exposure from individual trials on an ongoing basis as trial databases lock is essential to support this cross check.

Inevitably pooling from raw data will result in reprogramming of the same data and endpoints produced for CSR analysis.

In these cases teams may consider performing suitable consistency checks to ensure that any re-derivations yield the same results as the study or that any discrepancies are due to known differences in algorithms.

Differences that may be identified could be documented in case questions arise regarding any inconsistencies between a pooled level report and the individual CSR summaries so this information is readily available.

One mechanism to support this additional level of QC is to develop programs that will produce pooled analysis in such a way that they can be easily subsetted to report the data from the pooled database for a single study only. These subsetted reports generated from the pooled database can then be compared to the equivalent trial level outputs.

CONSIDERATIONS AND POSSIBLE PITFALLS

MAPPING OF TREATMENTS

It is often the case in pooled data presentations that treatments are presented differently to those in individual studies. For example, the key treatment arms and doses are likely to be presented individually but treatment arms and regimens investigated for dose ranging, or exploratory studies may be grouped into an "other active treatment" type group.

Special care should be taken for studies with multiple periods of treatment where up-/down titrations as well as add-on therapies might have to be taken into account.

PhUSE 2012

It is important that the required mappings are well documented.

MAPPING OF VISITS AND TIME POINTS

During the planning of mapping it is essential that any plans define how time points and visits will be pooled. Some considerations may include:

- Are pre-dose assessments in all pooled studies consistent? Different protocols may define pre-dose measures at different time points (e.g. 15 minutes pre-dose, 30 minutes pre-dose, pre-dose – with no specific time point). It may be useful to consider if there is a rationale to present all pre-dose assessments under one single pre-dose time point or if there is something to consider in the trial design that may mean that a pre-dose value at a given visit is influenced by the previous dose that means these values may not be directly comparable
- Are key reporting visits consistent across trials? Would it make sense to pool a week 24 visit with a week 26 visit, if planned visit differs across studies but essentially represent the same time point in the study (e.g. 6 months)?
- Is there value in pooling visits or time points that occur in only 1 or 2 of 20 pooled studies? This may be a necessity if endpoints are derived that utilize all data e.g. minimum or maximum post baseline value or change, however, the inclusion of these values in a pooled dataset does not always mean that these infrequently occurring time points need to be reported, e.g. in by visit summaries

ALGORITHM DEFINITIONS AND MISSING DATA

Though endpoints may be consistent across studies, the algorithms used to derive these may be dependent on a number of factors such as the length of the trial.

For example, a given endpoint may be analyzed using LOCF. Part of the algorithm for selection of the LOCF value may include a rule to cut-off data based on a number of weeks or days prior to the time point of interest e.g. LOCF is derived as the last non-missing value within X weeks of the primary end point. As the primary endpoint may vary between studies depending on the length, objectives and visit schedule of the trial, X may be variable across trials/studies. Once the data are pooled, does the LOCF need to be revised so a consistent algorithm is implemented?

Other common algorithm issues may include baseline definitions and other key covariates. These should all be considered.

It may be of value to note key algorithmic differences across studies while working on the study reporting to facilitate later planning for pooling.

EXTENSIONS AND INTERIM ANALYSIS

If extension data need to be pooled special attention to the related data should be considered; for examples:

- Treatment: on which treatment should patients who switch treatment from core to extension be considered.
- Baseline: for patients who switch treatment from core to extension, should their core baseline or end of core study as baseline be used?
- handling of duplicate records (events starting in a core study and continuing to its extension) need careful consideration (like AEs, Concomitant medications...)

If interim analysis data is included in the pool, we should also think at the specific points like how to determine the cut-off point (include all data up to a specific date? a specific visit? include only patients who completed the study? etc.)

CODING

When pooling it is often the case that studies in the pool were coded using different versions of a coding dictionary or even different coding dictionaries. This is especially relevant to MedDRA and WHO DRL dictionaries.

In general, for pooled analyses data should be reported under a consistent coding dictionary with the latest version available at the time of reporting by remapping based on lowest level code. For AE data using latest SMQ terms is required for RMP update.

If the pool includes a significant legacy of older studies, some recoding of data may be required to ensure data can be reported consistently. For very old studies in the legacy this may require translation from a foreign language or recoding to a common dictionary however this should be done by a specialists in those areas, not by the statistical and programming team.

Retaining the study level coded values may help in easily verifying differences between the CSR and pool reported summaries, however, care should be taken to ensure that these variables are clearly differentiated from those to be used for reporting (i.e. clear variable names and labels)

CONCLUSION

There is no obvious right or wrong way to pool data but there are many points that need to be well thought in advance. The early planning and a clear planning strategy would definitely help to create a robust pool. In terms of programming the possibilities to pool from raw or derived should be considered depending on the pool objectives, update, data availability... Also the strategy for the programs set up should be decided with clear rationale. Some specific points such as treatment re-mapping, inclusion of interim analysis, re-coding... have to be clearly documented to avoid possible errors.

PhUSE 2012

LIST OF ABBREVIATIONS

ADR – Adverse Drug Reaction
AE – Adverse Event
CSR – Clinical Study Report
DSUR - Drug Safety Update Report
HA – Health Authority
ICH - International Conference on Harmonization
LOCF – Last observation carried forward
SMQ - Standardized MedDRA Query

ACKNOWLEDGMENTS

The author thanks: Martin Blogg, Francesca Callegari, Joanna Cheng, Andrea Fabel, David Lawrence, Deborah Matos, Rainer Mertes, Miwa Okada, Simon Walsh, ShinRu Wang, Abraham Yeh, Xiaohong Zhang from Novartis.

CONTACT INFORMATION

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

Florence Buchheit
Novartis Pharma AG
Postfach
CH-4002 Basel
SWITZERLAND
+41 61 3246474
Florence.buchheit@novartis.com
www.novartis.com

Brand and product names are trademarks of their respective companies.