

Implementing ADNCA in the PK Process: A Cross-Departmental Approach

Saad Frédéric, Johnson & Johnson Innovative Medicine, Beerse, Belgium
Vuurstaek Daphne, Johnson & Johnson Innovative Medicine, Beerse, Belgium

ABSTRACT

PK-related data formatted as ADaM datasets are being requested more frequently by health authorities. This requires study teams to prepare submission-ready ADaM PK-related data, on an ad hoc basis without prior planning. Over the years, the International Society of Pharmacometrics (ISoP) & CDISC have partnered closely to standardize Non-Compartmental Analysis (NCA) data formats establishing the new regulatory standards which have been adopted by J&J. The new ADaM BDS data structure, ADNCA, necessitated process improvements to optimize the delivery of the required datasets. To address challenges and minimize last-minute re-work moving forward, we initiated a collaborative effort across multiple departments to reform our NCA data flow. This poster presents insights into the modifications we implemented utilizing an ADNCA dataset as input for the NCA analysis. We will outline the steps taken to enhance the rollout of this revised process and discuss the challenges encountered during implementation of the initial studies using this streamlined approach. Through these changes, we aim to improve efficiency and compliance with the latest regulatory standards in our programming efforts.

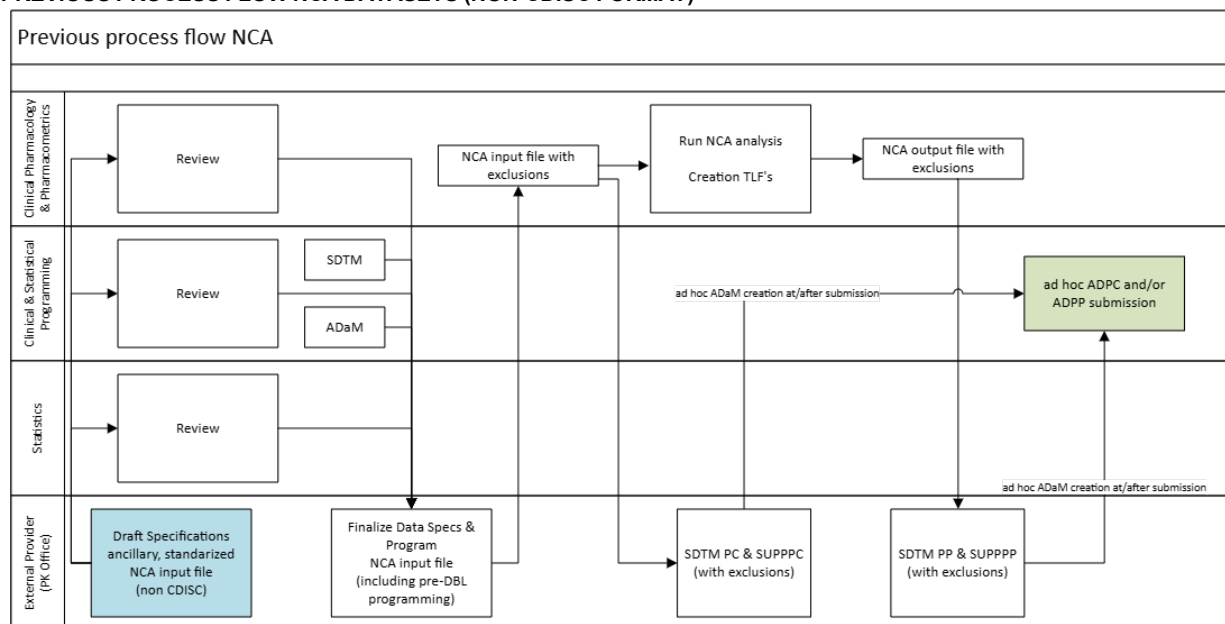
INTRODUCTION

Previous PK/PD data flows often required ad hoc ADaM dataset creation for FDA & PMDA submission purposes, few weeks prior to the actual ADaM dataset submission to include ADPC and ADPP datasets, as required, based on locked SDTM PP, PC, their supplemental data and ADSL. Introducing the ADNCA ADaM BDS dataset ensures a clear NCA data flow, in line with CDISC data standards. The ADNCA dataset specifications are reviewed early on by key stakeholders. The dataset is programmed by an external, unblinded PK Office, and analyzed by Clinical Pharmacology & Pharmacometrics within a Phoenix WinNonlin workflow immediately after unblinding. By working with an unblinded external party (PK Office), the ADNCA dataset can be prepared before unblinding, minimizing timelines after database release. Post-processing of the ADNCA dataset using R shiny allows efficient P21 check procedures by Clinical & Statistical Programming to retrieve the submission ready ADNCA dataset ahead of submission timelines, resulting in submission readiness for FDA and PMDA. J&J initiated production pilots last year and adjusted down the road to mitigate challenges.

OVERVIEW OF THE PREVIOUS NCA PROCESS FLOW

Non-Compartmental Analysis (NCA) is a method used in pharmacokinetics to evaluate drug concentration data to estimate pharmacokinetic parameters such as the area under the curve (AUC), clearance (CL), half-life ($t_{1/2}$), maximum concentration (C_{max}) and several other PK parameters. Over the past decade, J&J has internally standardized their workflow to obtain a standardized NCA input and output dataset with exclusions in a non-CDISC standard as source for the NCA analysis and NCA related tables and figures. In close collaboration with Clinical Data Standards & Data Management, these legacy NCA input and output datasets were converted by an external provider in SDTM SUPPPC, PP and SUPPPP. At times, FDA or PMDA requested the ADaM datasets to support NCA analysis, these were not available and were usually created ad hoc just weeks prior to the submission or after a submission based on an information request. The submitted datasets were also not the actual datasets used to perform NCA analysis or create the related outputs.

PREVIOUS PROCESS FLOW NCA DATASETS (NON-CDISC FORMAT)



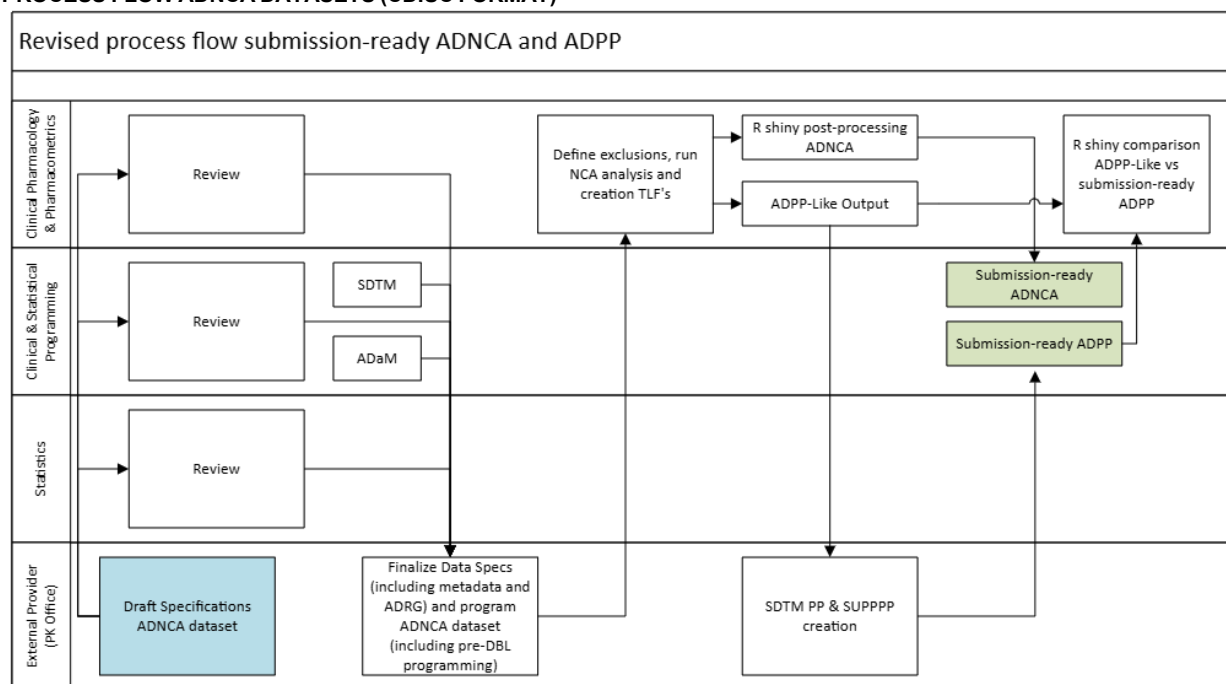
OVERVIEW OF THE IMPLEMENTED ADNCA PROCESS FLOW

Driven by the necessity to further standardize and clarify the roles and responsibilities within the NCA dataset workflow, a collaborative effort between CDISC and the International Society of Pharmacometrics resulted in the release of the ADNCA CDISC ADaM Basic Data Structure. J&J Clinical Pharmacology & Pharmacometrics revised their workflow together with Clinical & Statistical Programming to integrate the new CDISC standard within the NCA analysis workflow using Phoenix WinNonlin.

The input dataset created by the external provider for NCA analysis is now considered an ADaM dataset from the start. A template metadata file for ADNCA was created such that the ADNCA generated by the external provider follows requirements for the Phoenix WinNonlin software but also following ADaM requirements. This dataset is then updated by the PK Scientist to update exclusion flags and used for the NCA analysis and TLF creation. An R shiny application is used to check the ADNCA dataset date and time formatting pulled from Phoenix WinNonlin prior to providing the Statistical Programming ensuring smooth and consistent data processing to obtain submission ready ADNCA datasets. The package provided to Statistical Programming also includes the associated metadata and analysis drug reviewers guide clarification on the ADNCA dataset flow. After this, Statistical Programming will take the necessary steps, including P21 validation, to include the dataset in the ADaM eSub package.

In the revised flow, the ADPP-like dataset is used as input for any PK/PD parameter related outputs created within Phoenix WinNonlin. Per current CDISC IG, the PK/PD parameters from NCA should still be tabulated within the SDTM PP and SUPPPP dataset. As such, after SDTM finalization and P21 validation, an ADaM ADPP dataset is generated for ADaM submission purposes, using SDTM PP, SUPPPP and the ADaM ADSL dataset. The content and data structure of the ADPP-like and ADPP datasets match per revised process, a shiny R application compares the ADPP-like and ADPP to avoid additional reruns downstream. Both the ADNCA and ADPP are available within weeks after database lock and can be submitted readily for upcoming FDA or PMDA submissions.

PROCESS FLOW ADNCA DATASETS (CDISC FORMAT)



CHALLENGES THROUGHOUT THE PILOT PHASE

Throughout the initial pilots several learnings popped up, which required continuous adjustments of the initial planned approach, improving the quality of the specifications of the ADNCA dataset moving forward. Common challenges included date and time formatting issues when exporting dates and times from Phoenix WinNonlin, mitigated by R shiny applications. Dealing with Pinnacle21 errors for ADNCA & ADPP was leading us to make changes in the dataset and its attributes for required variables and expected labels and required some updates during dataset programming. We mitigated this for future trials by updating our dataset specification templates.

Since ADNCA is required to contain records for all scheduled samples, even if the record is not available in PC domain, P21 also has error AD0196 (Required Variable value is null) popping up. In the guidance documents for Statistical Programming, we provide advice on how to document this error in ADRG when popping up.

We chose to work together with an unblinded external provider (PK Office) to be able to pre-program ADNCA before DBL, decreasing timelines after database release. When they deliver the initial ADNCA dataset and metadata, they are also providing a study-specific text (based on a template text), which includes an overview of company owned SAS macros or R functions with a brief explanation on what those are used for, to be included in the ADRG, together with code used to generate ADNCA. This way, the statistical programmer has everything available to prepare the eSub ADaM package considering requirements from health authorities.

Another key challenge we found is related to implementing the ADNCA CDISC standard for a PK-PD combined ADNCA dataset, which wasn't obvious based on the CDISC guidance. Collaboration between Clinical Pharmacology & Pharmacometrics, Clinical & Statistical Programming, Clinical Data Standards and External Providers to combine key expertise has been essential throughout this process.

CONCLUSION

The roll-out and adoption of the CDISC ADaM BDS dataset structure of ADNCA has been implemented successfully at Johnson & Johnson Innovative Medicine through close collaboration between Clinical Pharmacology & Pharmacometrics, Clinical & Statistical Programming and Clinical Data Standards. The implementation ensures in time submission-ready non-compartmental analysis datasets including exclusions to descriptive concentration-time data, as well as data excluded from PK/PD parameter calculation, actively used in most recent ADaM submission packages, per regulatory requirements. At this point, we are still including the PK analysis results in PP/SUPPPP as part of the SDTM package. However, once this is no longer required based on CDISC guidelines, we are well-prepared to immediately have the ADPP-like dataset used for the creation of submission-ready ADPP.

REFERENCES

CDISC Analysis Data Model Implementation Guide for Non-compartmental Analysis Input Data v1.0

ACKNOWLEDGMENTS

We would like to thank:

- David Wexler & Kevin Shalayda as NCA analysis experts providing detailed analysis specific feedback
- Kirankumar Padhiar as clinical data standards expert, providing metadata guidance
- Payton Woodall & Stephen Mackie for the creation of the R Shiny app for ADNCA/ADPP
- All early adopters at J&J for making this new data standard implementation a success

CONTACT INFORMATION

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

Author Name: Frédéric Saad

Company: Johnson & Johnson Innovative Medicine

Address: Turnhoutseweg 30, 2340 Beerse, Belgium

Email: fsaad3@its.jnj.com

Website: <https://www.jnj.com>

Co-author Name: Daphne Vuurstaek

Company: Johnson & Johnson Innovative Medicine

Address: Turnhoutseweg 30, 2340 Beerse, Belgium

Email: DVuurst1@its.jnj.com

Website: <https://www.jnj.com>

Johnson&Johnson