

STATISTICAL ANALYSIS PLAN

Efficacy and Safety of Cerebrolysin in the Treatment of Aphasia after Acute Ischemic Stroke (ESCAS)

Coordinating Institution

Foundation for the Study of Nanoneurosciences and Neuroregeneration (FSNN)

Principal Investigator

Prof. Dr. Dafin F Muresanu

Statistician

Vlad Florin Chelaru

Data Manager

Diana Chira

Background

The Efficacy and Safety of Cerebrolysin in the Treatment of Aphasia after Acute Ischemic Stroke (study code: ESCAS; protocol number: FSNN20200207) study is an exploratory, prospective, randomized-controlled, double-blinded Phase 4 clinical trial. The study aims to assess the efficacy and safety of Cerebrolysin in combination with speech therapy compared to a placebo (saline solution) in combination with speech therapy for the treatment of aphasia following acute ischemic stroke.

Purpose of this SAP

The purpose of this Statistical Analysis Plan (SAP) is to provide a detailed and comprehensive plan for the statistical analyses that will be conducted on the data collected during the ESCAS study. The SAP ensures the transparency, consistency, and reproducibility of the study's data analysis, and helps avoid potential biases and data-driven decisions by prespecifying all planned analyses before data collection. This document will outline the statistical methods for analyzing the primary and secondary endpoints, as well as any exploratory endpoints, and will describe the approach for handling missing data, data transformations, and subgroup analyses.

By following this SAP, the study team can ensure that the results of the ESCAS study are analyzed and reported in a scientifically sound and rigorous manner, in accordance with best practices and regulatory guidelines. This will ultimately contribute to the assessment of the safety and efficacy of Cerebrolysin in the treatment of aphasia after acute ischemic stroke, and inform clinical decision-making and future research in this area.

Objective

Primary objectives

The primary objective of the ESCAS study is to assess the efficacy of Cerebrolysin and speech therapy versus placebo (saline solution) and speech therapy at 30, 60, and 90 days after baseline. Efficacy will be evaluated using the Western Aphasia Battery (Romanian translated version) scores at each time point. The primary outcome measure is the change in Western Aphasia Battery score from baseline to each follow-up time point (30, 60, and 90 days).

Secondary objectives

The secondary objectives of the ESCAS study include the following:

- To assess the efficacy of Cerebrolysin and speech therapy versus placebo (saline solution) and speech therapy at 30, 60, and 90 days after baseline using measures of motor, neurological, and global functional outcome. These outcomes will be evaluated

using the National Institutes of Health Stroke Scale (NIHSS), Barthel Index (BI), and modified Rankin Scale (mRS) scores at each time point.

- To evaluate the safety of Cerebrolysin and speech therapy versus placebo (saline solution) and speech therapy at 30, 60, and 90 days after baseline. Safety will be assessed by comparing the incidence of adverse events (AEs) and severe adverse events (SAEs) between the two treatment groups.
 - Additionally, we will measure specifically cardiovascular (such as, but not limited to stroke, myocardial infarction, atherosclerosis, vascular stenosis, as well as their recurrence), hematological (including anemia and vitamin B9 or B12 deficiency), renal system (including hyperuremia, hyperuricemia and urinary tract infections) and metabolic (including dyslipidemia, diabetes mellitus and atherosclerosis) related adverse events.

The secondary outcome measures include changes in NIHSS, BI, and mRS scores from baseline to each follow-up time point (30, 60, and 90 days), as well as the incidence of AEs and SAEs.

By addressing both the primary and secondary objectives, the ESCAS study will provide a comprehensive evaluation of the efficacy and safety of Cerebrolysin in the treatment of aphasia after acute ischemic stroke, in comparison with a placebo control group.

Definitions and abbreviations

AE = adverse event

SAE = serious adverse event

CSP = clinical study protocol

SAP = statistical analysis plan

CSR = clinical study results

RDMNUM = randomization number of a patient, used as unique identifier of each trial participant across databases

NIHSS = National Institute of Health Stroke Scale / Score

WAB = Western Aphasia Battery

- WABSS = Western Aphasia Battery Spontaneous Speech
- WABC = Western Aphasia Battery Comprehension
- WABR = Western Aphasia Battery Repetition
- WABN = Western Aphasia Battery Naming
- WABAQ = Western Aphasia Battery Aphasia Quotient

mRS = Modified Rankin Scale

BI = Barthel Index

Abbreviations ending in T = in databases, signifies the total value of a score (used for NIHSS, WAB and its subscales, and BI)

Abbreviations starting with d = in databases, signifies a differential value of a score at a given visit, compared to visit 1 (used for NIHSS, WAB and its subscales)

Differential = usually a value computed as the difference between the values of a given score at two different visits

Statistician = person which is (part of the team) tasked with preparing the plan of the statistical analysis, preparing the data (including conversion to necessary formats), executing the statistical analysis and reporting the results

Software utilized

We will use Microsoft Excel 2019, part of Microsoft Office 2019 suite (Microsoft Corporation, Redmond, WA), for data preparation and cleanup.

We will use R v. 4.3.1 (R Core Team, Vienna, Austria) and RStudio (Posit Software, PBC, Boston, MA) for data analysis. In the R workspace, we will load the following libraries:

stringr	Wickham H. stringr: Simple, Consistent Wrappers for Common String Operations [Internet]. 2022. Available from: https://CRAN.R-project.org/package=stringr
stringr	Wickham H. stringr: Simple, Consistent Wrappers for Common String Operations [Internet]. 2022. Available from: https://CRAN.R-project.org/package=stringr
ggplot2	Wickham H. ggplot2: Elegant Graphics for Data Analysis [Internet]. Springer-Verlag New York; 2016. Available from: https://ggplot2.tidyverse.org
readxl	Wickham H, Bryan J. readxl: Read Excel Files [Internet]. 2023. Available from: https://CRAN.R-project.org/package=readxl
xlsx	Dragulescu A, Arendt C. xlsx: Read, Write, Format Excel 2007 and Excel 97/2000/XP/2003 Files [Internet]. 2020. Available from: https://CRAN.R-project.org/package=xlsx

Should additional libraries be required in processing of the data analysis, they will be mentioned and their use will be motivated in the CSR.

Coding systems utilized

The original database uses the following variables which are of interest for the data analysis. In the following list, unless otherwise noted, x represents the moment of measurement (1=baseline, 4=last visit at 90 days after treatment). Note that for Barthel Index and modified

Rankin Scale, only scores at visits 2, 3, and 4 (30, 60 and 90 days after baseline, respectively) are available, with no value at baseline.

- RDMNUM = randomization number (unique to each participant)
-
- Western Aphasia Battery
 - WABAQ_x = Aphasia Quotient (0-100%)
- NIH Stroke Scale (Score)
 - NIHSST_x = total score (0-42)
- Barthel Index
 - BIT_x = total score (0-100)
- Modified Rankin Scale
 - MRS_x = total score (0-6 where 0=no symptom and 6=dead)

AEs and SAEs were stored in a separate database, from which we will retrieve the following information:

- Randomization number
- Whether the AE represents a SAE or not
- Based on the nature of the AE, some will be further classified by a clinician as being of the following nature:
 - Neurological
 - Psychiatric
 - Cardiovascular
 - Renal
 - Hematological
 - Gastrointestinal
 - Metabolic
 - Respiratory
 - Immune-related
 - ENT
 - Ophthalmic

Osteoarticular

For ease of data analysis, the statistician will convert the scores from the original database to a “long” format, with the following columns:

- RDMNUM
- Visit (1/2/3/4)
- is_crb
- is_PP
- WABAQ
- dWABAQ *
- NIHSST
- dNIHSST *
- BIT *
- MRS *

Columns starting with a *d* represent changes in scores at visits 2/3/4 compared to visit 1 (baseline). Columns marked with * have values only for visit 2/3/4.

The column *is_crb* will have a true/false value, depending on whether the patient was assigned to the treatment or placebo group. The column *is_PP* will have a true/false value, depending on whether the specified data point will be analyzed as part of the *Per Protocol* analysis set. This is a mechanism through which the statistician or any other authorized person can exclude specific patients from each analysis set, depending on criteria such as protocol violations or patient withdrawal, and then prepare the subset for the statistical analysis routine. Assignment of each patient to the correct analysis sets is described later in this document.

AE-related data will be coded and analyzed in a separate database, with the following structure:

- RDMNUM
- SAE (is the AE a SAE?)
- AE_total (the number of total AEs accused by the patient)
- SAE_total (the number of total SAEs accused by the patient)
- Number of AEs accused by the patient, by category
 - AE_cardiovascular
 - AE_renal
 - ...

Study objectives

Primary objectives

The primary objective of the ESCAS study is to assess the efficacy of Cerebrolysin and speech therapy versus placebo (saline solution) and speech therapy at 30, 60, and 90 days after baseline. Efficacy will be evaluated using the Western Aphasia Battery (Romanian translated version) scores at each time point. The primary outcome measure is the change in Western Aphasia Battery score from baseline to each follow-up time point (30, 60, and 90 days).

Secondary objectives

The secondary objectives of the ESCAS study include the following:

- To assess the efficacy of Cerebrolysin and speech therapy versus placebo (saline solution) and speech therapy at 30, 60, and 90 days after baseline using measures of motor, neurological, and global functional outcome. These outcomes will be evaluated using the National Institutes of Health Stroke Scale (NIHSS), Barthel Index (BI), and modified Rankin Scale (mRS) scores at each time point.
- To evaluate the safety of Cerebrolysin and speech therapy versus placebo (saline solution) and speech therapy at 30, 60, and 90 days after baseline. Safety will be assessed by comparing the incidence of adverse events (AEs) and severe adverse events (SAEs) between the two treatment groups.
 - Additionally, we will measure specifically neurological, psychiatric, cardiovascular, renal, hematological, gastrointestinal, metabolic, respiratory, immune-related, ENT, ophthalmic and osteoarticular related adverse events.

The secondary outcome measures include changes in NIHSS, BI, and mRS scores from baseline to each follow-up time point (30, 60, and 90 days), as well as the incidence of AEs and SAEs.

Study design

Overview

The ESCAS study is an exploratory, prospective, randomized-controlled, double-blinded Phase 4 clinical trial designed to assess the efficacy and safety of Cerebrolysin in combination with speech therapy compared to a placebo (saline solution) in combination with speech therapy for the treatment of aphasia following acute ischemic stroke. The study includes two treatment groups: one group receiving Cerebrolysin (N=60) and the other receiving a saline solution as placebo (N=60).

Sample size

The sample size for the ESCAS study was determined through a power analysis, considering several factors such as the type I error (alpha), type II error (beta), and the expected effect size (ES). The objective of the power analysis was to ensure that the study would have an adequate sample size to detect a statistically significant difference between the Cerebrolysin and placebo groups while minimizing the risks of false positive (type I error) and false negative (type II error) results.

Assumptions considered in sample size calculations

The sample size calculation was based on the following assumptions:

- Type I error (alpha): The probability of finding a statistically significant treatment effect when there is none. In this study, an alpha level of $\alpha=0.05$ was chosen, which is a commonly used threshold in clinical research.
- Type II error (beta): The probability of failing to detect a true treatment effect when one exists. In this study, a beta level of $\beta=0.2$ was chosen, which corresponds to a statistical power of $(1 - \beta)=0.8$. This means that there is a $(1 - \beta)$ probability of detecting a true effect, if present.
- Effect size (ES): The expected magnitude of the difference in treatment outcomes between the Cerebrolysin group and the placebo group. In this study, a medium effect size of ES was assumed, based on prior research and expert opinion.

Based on these assumptions, a power analysis was conducted to determine the necessary sample size for each treatment group. The analysis indicated that a sample size of N participants per group would be required to achieve the desired level of statistical power $(1 - \beta)$ while maintaining an acceptable type I error rate (α).

Therefore, the ESCAS study will enroll a total of at least 2N participants, with at least N participants allocated to the Cerebrolysin group and at least N participants allocated to the

placebo group. This sample size will ensure that the study has adequate power to detect a treatment effect of the expected magnitude while minimizing the risks of false positive and false negative results.

Assumed parameters and resulted sample size

Using G*Power 3.1.9.7 and assuming a medium effect size $d=0.5$, an alpha error of $\alpha=0.05$ and $\beta=0.2$, with an allocation ratio of 1 for Wilcoxon-Mann-Whitney test between two groups, we obtained a sample size of 53 patients per group. In order to provide a margin of error for situations such as patient withdrawal, SAEs or protocol violations, we have decided to recruit 120 patients in total for our study, 60 for each arm.

Randomisation

Randomisation was discussed at length in the CSP. Briefly, patients meeting inclusion and exclusion criteria will obtain a random number corresponding to the random list generated in advance by a biometrician selected by the Coordinator. Three opaque randomization envelopes will be distributed: to the nurse in charge of preparing the infusion solution, to the study center in case of suspicion of harm to the patient, as well as to the study coordinator.

The study and database will be unblinded after closure of the database and determination of the analysis populations.

Study schedule

Visit schedule

- Visit 1 - Day 0 - Baseline, 3-5 days post-stroke
 - Demographics
 - Medical history
 - Western Aphasia Battery
 - NIHSS
- Visit 2 - Day 30 $\pm$ 3
 - Western Aphasia Battery
 - NIHSS
 - Barthel Index
 - modified Rankin Scale
 - Adverse Events (AE)
 - Severe Adverse Events (SAE)
- Visit 3 - Day 60 $\pm$ 3
 - Western Aphasia Battery
 - NIHSS
 - Barthel Index
 - modified Rankin Scale

- Adverse Events (AE)
 - Severe Adverse Events (SAE)
- Visit 4 - Day 90 ± 3
 - Western Aphasia Battery
 - NIHSS
 - Barthel Index
 - modified Rankin Scale
 - Adverse Events (AE)
 - Severe Adverse Events (SAE)

Treatment schedule

- Treatment Cycle 1 – 30 ml Cerebrolysin/saline and ST for 2 x 5 days (2 weeks)
 - Study days 1 – 14
 - 30ml Cerebrolysin/saline i.v.
 - 1h Speech Therapy
- Treatment Cycle 2 – 30 ml Cerebrolysin/saline and ST for 2 x 5 days (2 weeks)
 - Study days 29-42
 - 30ml Cerebrolysin/saline i.v.
 - 1h Speech Therapy
- Treatment Cycle 3 – 30 ml Cerebrolysin/saline and ST for 2 x 5 days (2 weeks)
 - Study days 57-70
 - 30ml Cerebrolysin/saline i.v.
 - 1h Speech Therapy

Study endpoints

The statistical analysis methods selected for the ESCAS study are designed to address the study objectives and provide meaningful insights into the efficacy and safety of Cerebrolysin in the treatment of aphasia after acute ischemic stroke. The chosen methods are tailored to the types of variables included in the study and their distributions. A combination of nonparametric tests, contingency table analyses, and graphical representations will be used to analyze the data.

The study endpoints for the ESCAS study are designed to capture the efficacy and safety of Cerebrolysin in combination with speech therapy compared to placebo (saline solution) in combination with speech therapy for the treatment of aphasia following acute ischemic stroke. The endpoints are divided into primary and secondary endpoints, reflecting the prioritization of the study objectives.

Primary endpoint

The primary endpoint of the ESCAS study is the comparison of the Western Aphasia Battery (Romanian translated version) scores between the Cerebrolysin and placebo groups at each time

point (30, 60, and 90 days after study inclusion) compared to baseline. The Western Aphasia Battery is a standardized assessment tool used to evaluate language function in patients with aphasia. The study will use paired-measurement statistical tests to determine the efficacy of Cerebrolysin in improving language-related outcomes in the study population.

Additionally, unpaired-measurement statistical tests will be used to determine the differences between efficacy of the treatments between groups, at each time point.

Secondary endpoints

By evaluating these primary and secondary endpoints using paired-measurement statistical tests, the ESCAS study will provide a comprehensive assessment of the efficacy and safety of Cerebrolysin in the treatment of aphasia after acute ischemic stroke. This will help to inform clinical decision-making and guide future research in this area

Efficacy

- Comparison of the National Institutes of Health Stroke Scale (NIHSS) scores between the Cerebrolysin and placebo groups at each time point (30, 60, and 90 days after study inclusion) compared to baseline. The NIHSS is a widely used tool for assessing stroke severity and neurological deficits. Paired-measurement statistical tests will be used to analyze the NIHSS scores for each time point compared to baseline, in each group. Unpaired-measurement statistical tests will be used to determine the differences between efficacy of the treatment between groups, at each time point.
- Comparison of the modified Rankin Scale (mRS) scores between the Cerebrolysin and placebo groups at each follow-up time point (30, 60, and 90 days after study inclusion), as well as comparisons between the scores at 3rd and 4th visits compared to 2nd visit. The mRS is a commonly used measure of global disability and functional outcome in stroke patients.
- Comparison of the Barthel Index (BI) scores between the Cerebrolysin and placebo groups at each follow-up time point (30, 60, and 90 days after study inclusion), as well as comparisons between the scores at 3rd and 4th visits compared to 2nd visit. The BI is a measure of functional independence in activities of daily living.

Safety

- Incidence of AEs, SAEs and specific types of AEs. These safety endpoints will be used to evaluate the tolerability and safety profile of Cerebrolysin in the study population.

Statistical analysis

For comparing numeric values of paired samples, differential values for each pair of values will be computed. Then, the distribution of those differentials will be compared with the normal distribution, using the Shapiro-Wilk test. If the differentials are normally distributed, we will use a one-sample Student T test (equivalent with the Student T test for paired samples) with null

hypothesis $\mu=0$. If the differentials are not normally distributed, we will use a Wilcoxon signed-rank test with null hypothesis location=0.

For comparing numeric values of unpaired samples the normality of the values in each sample will be tested using Shapiro-Wilk test. If values from both samples are normally distributed, then variance equality between the samples will be assessed using the Bartlett test. If the variances are equal, then the differences between the two samples will be assessed using a Student T test for unpaired samples and equal variances. If the variances are not equal, then a Student T test for unpaired samples and unequal variances will be used. In both cases, the null hypothesis will be that mean difference is equal to 0. If values from at least one of the samples are not normally distributed, then a Wilcoxon rank sum with null hypothesis of location difference equal to 0 will be performed.

For comparing ordinal values of paired samples, a Wilcoxon signed-rank test will be used, and for unpaired samples, a Wilcoxon rank-sum test will be used.

For comparing the difference in prevalences of variants of one dichotomous or nominal variable, among the groups of another dichotomous or nominal variable (i.e. testing the association between two dichotomous/nominal variables), the Chi² test will be used or, where its assumptions would be violated (mainly due to small number of patients in any group), Fisher exact test will be used.

Where applicable, the two-tail p-value is reported. The type 1 error is assumed to be $\alpha=0.05$, and as such, results were considered statistically significant for $p<0.05$.

Analysis sets

The statistical analysis will be done separately and will be reported for each of the following study populations. Assessment of inclusion and exclusion will be done independently for each patient and each analysis set. Each patient will be included in all the relevant analysis sets.

Of mention is the fact that the same statistical procedures will be applied to the PP and ITT populations, the difference being in the selection of patients based on subsetting the database.

Efficacy analysis in Per Protocol population

The Per Protocol population (PP) includes all trial participants who have adhered to the study protocol, received the medication corresponding to the group they were assigned to after randomization, had at most minimal protocol deviations, and no missing values for total scores of Western Aphasia Battery and NIHSS at visits 2, 3 or 4, or missing values for modified Rankin Scale and Barthel Index at visits 3 or 4. This analysis set is used to assess the efficacy of the treatment in ideal conditions. From a missing data handling point of view, the PP population is subject to listwise deletion, meaning that a patient missing one data point is not taken into consideration in any analysis.

Efficacy analysis in Intention To Treat population

The Intention To Treat (ITT) population includes all trial participants who were registered in the trial and were randomized, irrespective of their subsequent adherence to the protocol or premature discontinuation. As such, compared to PP, ITT contains all the PP subjects, as well as subjects with major protocol deviations. This analysis set is used to assess the efficacy of the studied treatment, taking into consideration patients not adhering to the protocol (non-compliance, dropouts, SAEs, unforeseen events), thus better representing the expected results of the treatment in clinical practice. From a missing data handling point of view, the ITT population is subject to pairwise deletion, meaning that a patient missing one data point is still taken into consideration in any analysis that does not require the missing data point.

Safety analysis in safety population

The Safety Population (SP) includes all trial participants who were registered in the trial and received at least one dose of treatment, irrespective of their subsequent adherence to the protocol or premature discontinuation. Unlike ITT, SP analysis focuses on AEs and SAEs caused by the treatment, thus evaluating the safety-related parameters of the products. Patients are included in this population irrespective of missing data or premature discontinuation.

Data Review

Data management and transfer

In order to execute the statistical analysis, the statistician shall have access to the electronic database summarizing the CRFs of all patients included in the study. The data transfer shall be done through secure electronic means, such as e-mail or specific file sharing utilities. At no point in time should data be accessible on public servers, unsecured by credentials.

The statistician is required to keep the minimal number of copies of the study data on their devices, in order to facilitate the statistical analysis as well as easy restoration of the data in case of data loss.

After the finalization and delivery of the CSR and other relevant materials, the statistician is entitled, but not required to keep two copies of the database, statistical analysis routines, computer code, results or other materials resulting from their activity, for personal archiving reasons. The study coordinator is entitled to ask for a copy of this data, or request the statistician to permanently erase this data from their personal devices, up to one year after the delivery of the materials.

Data eligibility

Before unblinding, the statistician, having access to this SAP, will inspect the unblinded data to ensure correct format. Any abnormalities will be discussed and solved preferably before

unblinding. After unblinding, a second round of data inspection will be performed, and if required, corrections to the SAP will be issued.

Missing data handling

In our study, missing data is considered missing completely at random (MCAR).

Statistical procedures were chosen such that missing data points have minimal effect over data usability. As such, we avoid data imputation, choosing between two methods of record deletion: listwise and pairwise.

Listwise deletion means the analysis of only patients that have complete data for all variables; if any data point from any variable is missing, the patient is excluded from all analyses. This means that, for example, data from a patient missing one WAB value will not be taken into consideration during the analysis of mRS values. Listwise deletion has the advantage of being straightforward and easy to understand, at the cost of a reduced sample size. In our study, listwise deletion is used in analysis of the PP population.

Pairwise deletion means the analysis of all patients that have the required data points for individual statistical tests, irrespective of missing values from variables not included in the tests. This means that, for example, data from a patient missing one WAB value will be taken into consideration during the analysis of mRS values, but the same patient will not be included in analyses involving that specific missing value. Pairwise deletion has the advantage of retaining more data into the statistical analysis, with the cost of increased complexity (each statistical test having different sample sizes) and the assumption of independence of missing at random values. In our study, pairwise deletion is used in analysis of ITT population.

Statistical methodology

Data handling

Part of the data handling procedure was previously described in the Coding systems utilized. Briefly, the database containing the raw patient information from the CRFs will be converted to a long format database with visit number and patient randomization number to identify each measurement/score for each scale. This specific format is chosen in order to better facilitate the comparisons, as well as graphics generation with ggplot2.

Descriptive statistics

The following demographic data will be presented as part of the descriptive statistics: age at stroke onset, gender, educational level, alcohol and other substance abuse, as well as baseline WAB and NIHSS scores. Quantitative variables will be described as mean±standard deviation, while qualitative variables will be described through absolute and relative frequencies.

Average scores for each scale at each time point and for each group will be presented in table format, as well as in boxplot graphics.

Confirmatory statistics

No confirmatory analysis planned.

Comparability between study groups

Compatibility between groups will be tested for the following variables: age (at stroke onset), gender, education level, alcohol and other substance abuse, WAB scores (and their specific subscales) and NIHSS scores.

- Dichotomous variables will be tested using Chi² test or, where its assumptions are not met due to small group size, Fisher's exact test.
- Qualitative ordinal variables, such as education level or alcohol/substance abuse, will be tested using Wilcoxon rank-sum tests.
- For quantitative variables, initially, normality distribution will be assessed for each group individually using the Shapiro-Wilk test. If the values for at least one group are not normally distributed, Wilcoxon rank-sum tests will be used. If values from both groups are normally distributed, equality of variances will be assessed using Bartlett test and then a Student T test for unpaired samples and equal or unequal variances will be used.

Exploratory statistics

For WAB scores (and their specific subscales) as well as for NIHSS scores, we will calculate the differences between the scores at visits 2/3/4, and their corresponding baseline scores. These differences will be denoted with a lowercase letter "d." Throughout this section, we will refer to these values as *differentials*.

We will assess the normality of the differentials distributions for each group and visit using the Shapiro-Wilk test.

- If the distributions of the differentials show no statistically significant differences compared to the normal distribution, across all groups and visits, then those differentials will be tested against null hypotheses of $\mu=0$ using one sample Student T test. This would be the equivalent of testing the values of the scores for each group, at visit 2/3/4, against the baseline measurements of each respective group using paired-sample Student T test. Moreover, for each visit, we will compare the differentials between groups using Student T test for unpaired samples and equal or unequal variances, depending on the result of a Bartlett test for variance equality between groups.
- In cases where significant differences exist in the distribution of the differential values for certain groups and visits, compared to the normal distribution, we will test all differentials against the null hypothesis of location=0 using a one-sample Wilcoxon signed-rank test. This approach is equivalent to comparing the values of the scores for each group at visits 2, 3, and 4 against their corresponding baseline measurements using a Wilcoxon signed-rank test. Moreover, for each visit, we will compare the differentials between groups using a Wilcoxon rank-sum test.

For BI and mRS scores, differences from visit 2 to visit 3 and 4 respectively will be assessed in the same way, with paired-sample tests, and differences between groups will be assessed using unpaired-samples tests. .

The number of patients with at least an active AE, SAE and specific AE class, will be reported. Incidence of AEs, SAEs and AE specific classes will be compared between groups using Chi² test or where its assumptions are violated, Fisher's exact test. Moreover, the number of AEs and SAEs for each patient will be compared between groups using Wilcoxon rank-sum tests.

Interim analysis

None planned.

Analysis of subgroups

Subgroup analysis will be planned and executed after unblinding, if their distribution of demographic groups across trial arms allows for it.

Source code

Supplementary file 1 is provided at the end of this SAP, representing the source code for statistical analysis using R. Ideally, this code should not be changed after unblinding. Still, due to discrepancies between databases before and after unblinding, minimal adjustments to the statistical routines (especially for data loading) might be required. All subsequent source code changes should be documented, and the final source code version used in statistical analysis should be made available together with the CSR.

Planned tables and graphs

Planned tables

- table1a_demographics_of_all_patients_at_visit_1.xlsx: demographics and baseline characteristics of each group; this table will also contain the results of the statistical tests done for comparability between groups
- table1b_demographics_of_all_patients_ITT_at_visit_2.xlsx: demographics and baseline characteristics of each group, computed only for patients that were also part of the ITT analysis at visit 2
- table1c_demographics_of_all_patients_PP.xlsx: demographics and baseline characteristics of each group, computed only for patients that were part of the PP analysis
- itt/3dWABAQ.xlsx: WAB Aphasia Quotient differentials at visits 2/3/4, comparisons with baseline and comparisons of differentials between groups, at each visit
- itt/4dNIHSST.xlsx: NIHSS differentials at visits 2/3/4, comparisons with baseline and comparisons of differentials between groups, at each visit

- itt/5BIT.xlsx: Barthel Index scores at visits 2/3/4, comparisons of differentials from visits 3/4 against visit 2, and comparisons of those differentials between groups, at visits 3/4
- itt/6MRS.xlsx: modified Rankin Score at visits 2/3/4, comparisons of differentials from visits 3/4 against visit 2, and comparisons of those differentials between groups, at visits 3/4
- pp/3dWABAQ.xlsx: WAB Aphasia Quotient differentials at visits 2/3/4, comparisons with baseline and comparisons of differentials between groups, at each visit
- pp/4dNIHSST.xlsx: NIHSS differentials at visits 2/3/4, comparisons with baseline and comparisons of differentials between groups, at each visit
- pp/5BIT.xlsx: Barthel Index scores at visits 2/3/4, comparisons of differentials from visits 3/4 against visit 2, and comparisons of those differentials between groups, at visits 3/4
- pp/6MRS.xlsx: modified Rankin Score at visits 2/3/4, comparisons of differentials from visits 3/4 against visit 2, and comparisons of those differentials between groups, at visits 3/4
- AE_table.xlsx: counting of total AEs, SAEs and specific AE categories, as well as statistical tests between groups.

Planned graphs

- itt/1WABAQ.png: Boxplot showing the WAB aphasia quotient score (Oy) for each group (fill) at each visit (1/2/3/4) (Ox)
- itt/2NIHSST.png: Boxplot showing the NIHSS score (Oy) for each group (fill) at each visit (1/2/3/4) (Ox)
- itt/3dWABAQ.png: Boxplot showing the WAB aphasia quotient differential (Oy) for each group (fill) at each visit (2/3/4 compared to 1) (Ox)
- 4dNIHSST.png: Boxplot showing the NIHSS score differential (Oy) for each group (fill) at each visit (2/3/4 compared to 1) (Ox)
- itt/5BIT.png and itt/6MRS.png: Boxplot showing the respective score (Oy) for each group (fill) at each visit (2/3/4)
- pp/1* to 6*: the same graphs, redone for the Per Protocol analysis set
- AE_number_of_AEs.png: barplot showing the total number (Oy) of AEs, SAEs and specific categories (Ox) by group (fill)
- AE_number_of_patients.png: barplot showing the total number of patients (Oy) accusing AEs, SAEs or specific categories (Ox) by group (fill)
- AE_percent_of_patients.png: barplot showing the percentage of patients (Oy) accusing AEs, SAEs and specific categories (Ox) by group (fill)

Supplementary file: R source code

```
# ESCAS statistical analysis code
# Author: Vlad-Florin Chelaru, vlad.chelaru@brainscience.ro
library(readxl)
library(xlsx)
```

```
library(stringr)
library(ggplot2)
#library(MANOVA.RM)
# Some general functions ----
beautiful.p.value<-function(x){ #beautiful p value, if strictly under 0.001 print as <0.001,
otherwise print as =0.***
  x<-as.numeric(x)
  if(is.na(x))
    return("NA")
  if(is.nan(x))
    return("NaN")
  if(x<0.001)
    return("<0.001")
  return(paste("",round(x,3),sep=""))
}
test.numeric.unpaired<-function(t1,t2){
  ret<-list()
  ret$sum_nas<-sum(is.na(t1))+sum(is.na(t2))
  t1<-t1[!is.na(t1)]
  t2<-t2[!is.na(t2)]
  if(length(unique(t1))==1)
    ret$shapiro.p.t1<-1
  else
    ret$shapiro.p.t1<-shapiro.test(t1)$p.value
  if(length(unique(t2))==1)
    ret$shapiro.p.t2<-1
  else
    ret$shapiro.p.t2<-shapiro.test(t2)$p.value
  if(ret$shapiro.p.t1<0.05|ret$shapiro.p.t2<0.05){ # if any of them is not normally distributed
    ret$final.test<-"Wilcoxon rank sum"
    ret$p.value<-wilcox.test(t1,t2,paired = F)$p.value
  }else{
    ret$bartlett.p<-bartlett.test(list(t1,t2))$p.value
    if(ret$bartlett.p<0.05){ # unequal distrib
      ret$final.test<-"Student T test for unpaired samples and unequal variances"
      ret$p.value<-t.test(t1,t2,paired = F,var.equal = F)$p.value
    }else{
      ret$final.test<-"Student T test for unpaired samples and equal variances"
      ret$p.value<-t.test(t1,t2,paired = F,var.equal = T)$p.value
    }
  }
  return(ret)
}
test.numeric.paired.2<-function(t1,t2){
```

```

ret<-list()
if(length(t1)!=length(t2))
  stop("Different lengths!")
ret$na_vals<-is.na(t1)|is.na(t2)
t1<-t1[!ret$na_vals]
t2<-t2[!ret$na_vals]
ret$shapiro.p.t1<-shapiro.test(t1)$p.value
ret$shapiro.p.t2<-shapiro.test(t2)$p.value
if((ret$shapiro.p.t1<0.05|ret$shapiro.p.t2<0.05)){ # if any of them is not normally distributed
  ret$final.test<-"Wilcoxon signed rank"
  ret$p.value<-wilcox.test(t1,t2,paired = T)$p.value
}else{
  ret$final.test<-"Student T test for paired samples"
  ret$p.value<-t.test(t1,t2,paired = T)$p.value
}
return(ret)
}
test.numeric.paired.1<-function(tx){ #the same idea, but for when I supply only one group of
values
  ret<-list()
  ret$na_vals<-is.na(tx)
  tx<-tx[!ret$na_vals]
  ret$shapiro.p<-shapiro.test(tx)$p.value
  if(ret$shapiro.p<0.05){
    ret$final.test<-"Wilcoxon signed rank"
    ret$p.value<-wilcox.test(tx)$p.value
  }else{
    ret$final.test<-"Student T test for one sample"
    ret$p.value<-t.test(tx)$p.value
  }
  return(ret)
}
contingency.table.test<-function(t){
  if(nrow(t)==2 & ncol(t)==2){
    f<-fisher.test(t)
    bigtotal<-sum(t)
    can_use_chisq<-
      (sum(t[1,])*sum(t[,1])/bigtotal>=5) &
      (sum(t[1,])*sum(t[,2])/bigtotal>=5) &
      (sum(t[2,])*sum(t[,1])/bigtotal>=5) &
      (sum(t[2,])*sum(t[,2])/bigtotal>=5)
    ret<-c(
      "p_chi2_cor"=chisq.test(t)$p.value,
      "p_chi2"=chisq.test(t,correct = F)$p.value,

```

```

    "p_fisher"=f$p.value,
    "p_chosen"=ifelse(can_use_chisq,chisq.test(t,correct = F)$p.value,f$p.value),
    "can_use_chi2"=can_use_chisq,
    "OR"=t[1,1]*t[2,2]/(t[1,2]*t[2,1]),
    "OR_inf"=exp(log(t[1,1]*t[2,2]/(t[1,2]*t[2,1]))-1.96*sqrt(1/t[1,1]+1/t[1,2]+1/t[2,1]+1/t[2,2])),
    "OR_sup"=exp(log(t[1,1]*t[2,2]/(t[1,2]*t[2,1]))+1.96*sqrt(1/t[1,1]+1/t[1,2]+1/t[2,1]+1/t[2,2])),
    "OR_Fisher"=unname(f$estimate),
    "OR_F_inf"=f$conf.int[1],
    "OR_F_sup"=f$conf.int[2],
    "shorthand"=paste(
      ifelse(can_use_chisq,"Chi2","Fisher")," ",
      beautiful.p.value(ifelse(can_use_chisq,chisq.test(t,correct = F)$p.value,f$p.value)),
      " OR=",
      round(t[1,1]*t[2,2]/(t[1,2]*t[2,1]),3),
      sep=""
    )
  )
}
ret
}
else
{
  f<-fisher.test(t)
  y<-0
  for(j in 1:nrow(t))
    for(k in 1:ncol(t))
      y<-y+(sum(t[j,])*sum(t[,k])/sum(t))>=5)
  can_use_chisq<-(y/(nrow(t)*ncol(t))>=0.8)
  ret<-c(
    "p_chi2_cor"=chisq.test(t)$p.value,
    "p_chi2"=chisq.test(t,correct = F)$p.value,
    "p_fisher"=fisher.test(t)$p.value,
    "p_chosen"=ifelse(can_use_chisq,chisq.test(t,correct = F)$p.value,fisher.test(t)$p.value),
    "can_use_chi2"=can_use_chisq,
    "shorthand"=paste(
      ifelse(can_use_chisq,"Chi2","Fisher")," ",
      beautiful.p.value(ifelse(can_use_chisq,chisq.test(t,correct = F)$p.value,f$p.value)),
      sep=""
    )
  )
}
ret
}
}
descr.numeric.meansd<-function(x){
  return(paste0(round(mean(x),3),"±",round(sd(x),3)))
}

```

```
}
descr.numeric.median<-function(x){
  x<-quantile(x)
  return(paste0(round(x[3],3)," (",round(x[2],3)," - ",round(x[4],3),")"))
}
# Load data - this might change ----
db_scores<-data.frame(read_excel("dbfin.xlsx","only_total_scores"))
db_bigdat<-data.frame(read_excel("dbfin.xlsx","escas_data"))
db_assign<-data.frame(read_excel("rndtest.xlsx"))
db_advers<-data.frame(read_excel("dbfin.xlsx","adverse_events"))
names(db_assign)<-c("RDMNUM","is_crb","is_pp")
# NB: rndtest is a database for test, which contains the group assignment data (RDMNUM,
is_crb and is_pp)
# this portion of code is subject to changes after unblinding to adapt the unblinding data format
to preexistent code
## Check uniqueness of the RDMNUM just for sanity ----
t1<-table(db_scores$RDMNUM)
t2<-table(db_bigdat$RDMNUM)
t3<-table(db_assign$RDMNUM)
sum(!t1==1)
sum(!t2==1)
sum(!t3==1)
sum(names(t1)!=names(t2))
sum(names(t1)!=names(t3))
## Introduce grouping and population info into the data frames ----
for(i in 1:nrow(db_scores)){
  db_scores[i,"is_crb"]<-db_assign[db_assign$RDMNUM==db_scores$RDMNUM[i],"is_crb"]
  db_bigdat[i,"is_crb"]<-db_assign[db_assign$RDMNUM==db_bigdat$RDMNUM[i],"is_crb"]
  db_scores[i,"is_pp"]<-db_assign[db_assign$RDMNUM==db_scores$RDMNUM[i],"is_pp"]
  db_bigdat[i,"is_pp"]<-db_assign[db_assign$RDMNUM==db_bigdat$RDMNUM[i],"is_pp"]
}
# TABLE 1 DEMOGRAPHICS ----
compute_demographics<-function(db_bigdat,name_out){
  beautiful_table<-data.frame(
    Characteristic=as.character(),
    Variant=as.character(),
    CRB=as.character(),
    PLC=as.character(),
    p.value=as.character(),
    test=as.character()
  )
  ## Number of patients in each group ----
  beautiful_table[1,]<-c(
    "N=",
```

```

    """,
    paste0(sum(db_bigdat$sis_crb),"
(",round(sum(db_bigdat$sis_crb)/nrow(db_bigdat)*100,3),"%")),
    paste0(sum(!db_bigdat$sis_crb),"
(",round(sum(!db_bigdat$sis_crb)/nrow(db_bigdat)*100,3),"%")),
    """, """)
)
## Age ----
db_bigdat$computed_age<-floor(as.numeric((db_bigdat$SODTC-
db_bigdat$BRTHTD)/365.25))
x<-test.numeric.unpaired(
  (db_bigdat$computed_age[db_bigdat$sis_crb]),
  (db_bigdat$computed_age[!db_bigdat$sis_crb])
)
beautiful_table[2,]<-c(
  "Age","mean+sd",
  descr.numeric.meansd(db_bigdat$computed_age[db_bigdat$sis_crb]),
  descr.numeric.meansd(db_bigdat$computed_age[!db_bigdat$sis_crb]),
  beautiful.p.value(x$p.value),x$final.test
)
beautiful_table[3,]<-c(
  """, "median",
  descr.numeric.median(db_bigdat$computed_age[db_bigdat$sis_crb]),
  descr.numeric.median(db_bigdat$computed_age[!db_bigdat$sis_crb]),
  """, """)
)
## Sex ----
x<-contingency.table.test(table(db_bigdat$SEX,db_bigdat$sis_crb))
beautiful_table[4,]<-c(
  "Sex","Female",
  paste0(sum(db_bigdat$SEX==2&db_bigdat$sis_crb),"
(",round(sum(db_bigdat$SEX==2&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*100,3),"%")),
  paste0(sum(db_bigdat$SEX==2&(!db_bigdat$sis_crb)),
  ("round(sum(db_bigdat$SEX==2&(!db_bigdat$sis_crb))/sum(!db_bigdat$sis_crb)*100,3),"%")),
  beautiful.p.value((x["p_chosen"])),ifelse(x["can_use_chi2"]=="TRUE","Chi^2 test","Fisher
test")
)
beautiful_table[5,]<-c(
  """, "Male",
  paste0(sum(db_bigdat$SEX==1&db_bigdat$sis_crb),"
(",round(sum(db_bigdat$SEX==1&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*100,3),"%")),
  paste0(sum(db_bigdat$SEX==1&(!db_bigdat$sis_crb)),
  ("round(sum(db_bigdat$SEX==1&(!db_bigdat$sis_crb))/sum(!db_bigdat$sis_crb)*100,3),"%")),
  """, """)

```

```

)

## Education level ----
x<-contingency.table.test(table(db_bigdat$EDU,db_bigdat$sis_crb))
beautiful_table[6,]<-c(
  "Education","No formal education",
  paste0(sum(db_bigdat$EDU==1&db_bigdat$sis_crb),"
(",round(sum(db_bigdat$EDU==1&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*100,3,"%"),
  paste0(sum(db_bigdat$EDU==1&(!db_bigdat$sis_crb)),
  ("round(sum(db_bigdat$EDU==1&(!db_bigdat$sis_crb))/sum(!db_bigdat$sis_crb)*100,3,"%"),
  beautiful.p.value((x["p_chosen"])),ifelse(x["can_use_chi2"]=="TRUE","Chi^2 test","Fisher
test")
)
beautiful_table[7,]<-c(
  "", "Primary school",
  paste0(sum(db_bigdat$EDU==2&db_bigdat$sis_crb),"
(",round(sum(db_bigdat$EDU==2&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*100,3,"%"),
  paste0(sum(db_bigdat$EDU==2&(!db_bigdat$sis_crb)),
  ("round(sum(db_bigdat$EDU==2&(!db_bigdat$sis_crb))/sum(!db_bigdat$sis_crb)*100,3,"%"),
  "", ""
)
beautiful_table[8,]<-c(
  "", "Secondary school",
  paste0(sum(db_bigdat$EDU==3&db_bigdat$sis_crb),"
(",round(sum(db_bigdat$EDU==3&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*100,3,"%"),
  paste0(sum(db_bigdat$EDU==3&(!db_bigdat$sis_crb)),
  ("round(sum(db_bigdat$EDU==3&(!db_bigdat$sis_crb))/sum(!db_bigdat$sis_crb)*100,3,"%"),
  "", ""
)
beautiful_table[9,]<-c(
  "", "High school",
  paste0(sum(db_bigdat$EDU==4&db_bigdat$sis_crb),"
(",round(sum(db_bigdat$EDU==4&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*100,3,"%"),
  paste0(sum(db_bigdat$EDU==4&(!db_bigdat$sis_crb)),
  ("round(sum(db_bigdat$EDU==4&(!db_bigdat$sis_crb))/sum(!db_bigdat$sis_crb)*100,3,"%"),
  "", ""
)
beautiful_table[10,]<-c(
  "", "University",
  paste0(sum(db_bigdat$EDU==5&db_bigdat$sis_crb),"
(",round(sum(db_bigdat$EDU==5&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*100,3,"%"),
  paste0(sum(db_bigdat$EDU==5&(!db_bigdat$sis_crb)),
  ("round(sum(db_bigdat$EDU==5&(!db_bigdat$sis_crb))/sum(!db_bigdat$sis_crb)*100,3,"%"),
  "", ""
)

```

```

)
x<-
test.numeric.unpaired(db_bigdat$EDU[db_bigdat$sis_crb],db_bigdat$EDU[!db_bigdat$sis_crb])
beautiful_table[11,]<-c(
  "", "treat as numbers",
  descr.numeric.median(db_bigdat$EDU[db_bigdat$sis_crb]),
  descr.numeric.median(db_bigdat$EDU[!db_bigdat$sis_crb]),
  beautiful.p.value(x$p.value),x$final.test
)
## Alcohol ----
x<-contingency.table.test(table(db_bigdat$Substances_alcohol,db_bigdat$sis_crb))
beautiful_table[12,]<-c(
  "Alcohol use", "None",
  paste0(sum(db_bigdat$Substances_alcohol==1&db_bigdat$sis_crb),"
(",round(sum(db_bigdat$Substances_alcohol==1&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*10
0,3),"%")),
  paste0(sum(db_bigdat$Substances_alcohol==1&!db_bigdat$sis_crb)),
  ("round(sum(db_bigdat$Substances_alcohol==1&!db_bigdat$sis_crb))/sum(!db_bigdat$sis_crb)*
100,3),"%")),
  beautiful.p.value((x["p_chosen"])),ifelse(x["can_use_chi2"]=="TRUE","Chi^2 test","Fisher
test")
)
beautiful_table[13,]<-c(
  "", "Use",
  paste0(sum(db_bigdat$Substances_alcohol==2&db_bigdat$sis_crb),"
(",round(sum(db_bigdat$Substances_alcohol==2&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*10
0,3),"%")),
  paste0(sum(db_bigdat$Substances_alcohol==2&!db_bigdat$sis_crb)),
  ("round(sum(db_bigdat$Substances_alcohol==2&!db_bigdat$sis_crb))/sum(!db_bigdat$sis_crb)*
100,3),"%")),
  "", ""
)
beautiful_table[14,]<-c(
  "", "Abuse",
  paste0(sum(db_bigdat$Substances_alcohol==3&db_bigdat$sis_crb),"
(",round(sum(db_bigdat$Substances_alcohol==3&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*10
0,3),"%")),
  paste0(sum(db_bigdat$Substances_alcohol==3&!db_bigdat$sis_crb)),
  ("round(sum(db_bigdat$Substances_alcohol==3&!db_bigdat$sis_crb))/sum(!db_bigdat$sis_crb)*
100,3),"%")),
  "", ""
)

```

```

x<-
test.numeric.unpaired(db_bigdat$Substances_alcohol[db_bigdat$sis_crb],db_bigdat$Substances
_alcohol[!db_bigdat$sis_crb])
beautiful_table[15,]<-c(
  "", "treat as numbers",
  descr.numeric.median(db_bigdat$Substances_alcohol[db_bigdat$sis_crb]),
  descr.numeric.median(db_bigdat$Substances_alcohol[!db_bigdat$sis_crb]),
  beautiful.p.value(x$p.value),x$final.test
)
## Other substances ----
x<-contingency.table.test(table(db_bigdat$Substances_other,db_bigdat$sis_crb))
beautiful_table[16,]<-c(
  "Other substances use", "None",
  paste0(sum(db_bigdat$Substances_other==1&db_bigdat$sis_crb),
  ("round(sum(db_bigdat$Substances_other==1&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*100,
  3),"%)" ),
  paste0(sum(db_bigdat$Substances_other==1&!db_bigdat$sis_crb),
  ("round(sum(db_bigdat$Substances_other==1&!db_bigdat$sis_crb)/sum(!db_bigdat$sis_crb)*1
  00,3),"%)" ),
  beautiful.p.value((x["p_chosen"])),ifelse(x["can_use_chi2"]=="TRUE","Chi^2 test","Fisher
  test")
)
beautiful_table[17,]<-c(
  "", "Use",
  paste0(sum(db_bigdat$Substances_other==2&db_bigdat$sis_crb),
  ("round(sum(db_bigdat$Substances_other==2&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*100,
  3),"%)" ),
  paste0(sum(db_bigdat$Substances_other==2&!db_bigdat$sis_crb),
  ("round(sum(db_bigdat$Substances_other==2&!db_bigdat$sis_crb)/sum(!db_bigdat$sis_crb)*1
  00,3),"%)" ),
  "", ""
)
beautiful_table[18,]<-c(
  "", "Abuse",
  paste0(sum(db_bigdat$Substances_other==3&db_bigdat$sis_crb),
  ("round(sum(db_bigdat$Substances_other==3&db_bigdat$sis_crb)/sum(db_bigdat$sis_crb)*100,
  3),"%)" ),
  paste0(sum(db_bigdat$Substances_other==3&!db_bigdat$sis_crb),
  ("round(sum(db_bigdat$Substances_other==3&!db_bigdat$sis_crb)/sum(!db_bigdat$sis_crb)*1
  00,3),"%)" ),
  "", ""
)

```

```

x<-
test.numeric.unpaired(db_bigdat$Substances_other[db_bigdat$is_crb],db_bigdat$Substances_
other[!db_bigdat$is_crb])
beautiful_table[19,]<-c(
  "", "treat as numbers",
  descr.numeric.median(db_bigdat$Substances_other[db_bigdat$is_crb]),
  descr.numeric.median(db_bigdat$Substances_other[!db_bigdat$is_crb]),
  beautiful.p.value(x$p.value),x$final.test
)
## WAB Aphasia quotient ----
x<-
test.numeric.unpaired(db_bigdat$WABAQ_1[db_bigdat$is_crb],db_bigdat$WABAQ_1[!db_bigd
at$is_crb])
beautiful_table[20,]<-c(
  "WAB Aphasia quotient", "mean+sd",
  descr.numeric.meansd(db_bigdat$WABAQ_1[db_bigdat$is_crb]),
  descr.numeric.meansd(db_bigdat$WABAQ_1[!db_bigdat$is_crb]),
  beautiful.p.value(x$p.value),x$final.test
)
beautiful_table[21,]<-c(
  "", "median",
  descr.numeric.median(db_bigdat$WABAQ_1[db_bigdat$is_crb]),
  descr.numeric.median(db_bigdat$WABAQ_1[!db_bigdat$is_crb]),
  "", ""
)
## NIHSS ----
x<-
test.numeric.unpaired(db_bigdat$NIHSST_1[db_bigdat$is_crb],db_bigdat$NIHSST_1[!db_bigda
t$is_crb])
beautiful_table[22,]<-c(
  "NIHSS", "mean+sd",
  descr.numeric.meansd(db_bigdat$NIHSST_1[db_bigdat$is_crb]),
  descr.numeric.meansd(db_bigdat$NIHSST_1[!db_bigdat$is_crb]),
  beautiful.p.value(x$p.value),x$final.test
)
beautiful_table[23,]<-c(
  "", "median",
  descr.numeric.median(db_bigdat$NIHSST_1[db_bigdat$is_crb]),
  descr.numeric.median(db_bigdat$NIHSST_1[!db_bigdat$is_crb]),
  "", ""
)
write.xlsx(beautiful_table,name_out)
}
compute_demographics(db_bigdat,"table1a_demographics_of_all_patients_at_visit_1.xlsx")

```

```
compute_demographics(db_bigdat[!is.na(db_bigdat$WABAQ_2),],"table1b_demographics_of_all_patients_ITT_at_visit_2.xlsx")
compute_demographics(db_bigdat[db_bigdat$is_pp,],"table1c_demographics_of_all_patients_PP.xlsx")
# Creation of LONG database ----
db_long<-data.frame(
  RDMNUM=as.numeric(),
  is_crb=as.logical(),
  is_pp=as.logical(),
  visit=as.numeric()
)
db_names<-data.frame(
  short=c(
    "WABAQ","NIHSST",
    "dWABAQ","dNIHSST",
    "BIT","MRS",
    # AE stuff starts at 7
    "AE_total","SAE_total",
    "AE_neurological","AE_psychiatric","AE_cardiovascular",
    "AE_renal","AE_hematological","AE_gastrointestinal",
    "AE_metabolic","AE_respiratory","AE_immune.related.AE",
    "AE_ENT","AE_ophthalmic","AE_osteoarticular"
  ),
  long=c(
    "WAB Aphasia Quotient",
    "NIHSS",
    "WAB Aphasia Quotient differential",
    "NIHSS differential",
    "Barthel Index","Modified Rankin Scale",
    "Adverse events","Serious adverse events",
    "Neurological AEs","Psychiatric AEs","Cardiovascular AEs",
    "Renal AEs","Hematological AEs","Gastrointestinal AEs",
    "Metabolic AEs","Respiratory AEs","Immune-related AEs",
    "ENT AEs","Ophthalmic AEs","Osteoarticular AEs"
  )
)
for(score_to_diff in c("WABAQ","NIHSST"))
  for(i in 2:4)
    db_scores[,paste0("d",score_to_diff,"_",i)]<-db_scores[,paste0(score_to_diff,"_",i)]-
db_scores[,paste0(score_to_diff,"_1")]
for(pati in 1:nrow(db_scores))
  for(visit in 1:4){
    j=nrow(db_long)+1
    db_long[j,"RDMNUM"]<-db_scores[pati,"RDMNUM"]
```

```

db_long[j,"is_crb"]<-db_scores[pati,"is_crb"]
db_long[j,"is_pp"]<-db_scores[pati,"is_pp"]
db_long[j,"visit"]<-visi
for(i in 1:2) # i represents the score that I transform to long format, and WAB and NIHSS have
4 visits so they are treated separately
  db_long[j,db_names$short[i]]<-db_scores[pati,paste0(db_names$short[i],"_",visi)]
  if(visi>1) for(i in 3:6)
    db_long[j,db_names$short[i]]<-db_scores[pati,paste0(db_names$short[i],"_",visi)]
}
# A function that generates a boxplot graph ----
# DFL should have the following; they should be in this order, not necessarily this names
# RDMNUM as whatever
# is_crb as logical
# visit as numeric
# value as numeric
graph_them_all<-function(dfl,ylb,t1l){
  names(dfl)<-c("RDMNUM","is_crb","visit","value")
  dfl$visit<-factor(dfl$visit)
  n<-table(dfl$is_crb)/sum(table(dfl$RDMNUM,dfl$is_crb)[1,])
  dfl$is_crb<-factor(dfl$is_crb,levels=c(F,T),labels=c(
    paste0("Placebo\n(N=",n[1],")"),
    paste0("Cerebrolysin\n(N=",n[2],")")))
  g<-
    ggplot(dfl,aes(x=factor(visit,levels=1:4),y=value,fill=is_crb))+
    geom_boxplot(outlier.size = 3,lwd=1)+
    stat_summary(fun=mean,position = position_jitterdodge(jitter.width = 0,jitter.height =
0),size=1,shape=5,show.legend = F) +
    stat_summary(fun.data=mean_sdl,fun.args = list(mult=1),geom="errorbar",position =
position_jitterdodge(jitter.width = 0,jitter.height =
0),linewidth=0.5,linetype="dashed",width=0.5,show.legend = F)+
    scale_fill_manual(values=c("#6495ED","#ee82ee"))+
    labs(title=t1l,x="Visit",y=ylb,fill="Treatment")+
    theme_light()+
    theme(plot.title=element_text(hjust=0),plot.subtitle = element_text(hjust=0),text =
element_text(size = 25))
  return(g)
}
# A function that generates a table ----
# the following rows:
# N (number of valid data points)    1
# mean-sd of group                  2
# median of group                    3
# p value diff inside group (paired) 4
# statistical test                    5

```

```

# the following columns:
# visit (2/3/4)          1
# sub                    2
# crb                    3
# plc                    4
# difference between them (unpaired) 5
# already_diff should be true for WAB and NIHSS, because they are already differentials
#                        false for mRS and BI, because then I will manually compute diff 3-2 4-2
table_them_all<-function(dfl,already_diff){
  ret<-data.frame(
    visit=as.character(),
    sub=as.character(),
    crb=as.character(),
    plc=as.character(),
    diff_between_grps=as.character()
  )
  names(dfl)<-c("RDMNUM","is_crb","visit","value")
  ret[1:15,2]<-c("Number of valid data points N","Mean+SD of group","Median of group","p value
(differential)","Statistical test")
  for(visi in 2:4){
    ret[(visi-2)*5+1,1]<-visi
    vc<-dfl[dfl$is_crb&dfl$visit==visi,4]
    vp<-dfl[(!dfl$is_crb)&dfl$visit==visi,4]
    vc<-vc[!is.na(vc)]
    vp<-vp[!is.na(vp)]
    ret[(visi-2)*5+1,3]<-length(vc)
    ret[(visi-2)*5+1,4]<-length(vp)
    ret[(visi-2)*5+2,3]<-descr.numeric.meansd(vc)
    ret[(visi-2)*5+2,4]<-descr.numeric.meansd(vp)
    ret[(visi-2)*5+3,3]<-descr.numeric.median(vc)
    ret[(visi-2)*5+3,4]<-descr.numeric.median(vp)
    if(already_diff)
    {
      tv<-test.numeric.paired.1(vc)
      tvp<-test.numeric.paired.1(vp)
      ret[(visi-2)*5+4,3]<-beautiful.p.value(tv$p.value)
      ret[(visi-2)*5+4,4]<-beautiful.p.value(tvp$p.value)
      ret[(visi-2)*5+5,3]<-tv$final.test
      ret[(visi-2)*5+5,4]<-tvp$final.test
    }
  }else if(visi!=2){
    vc<-dfl[dfl$is_crb&dfl$visit==visi,4]-dfl[dfl$is_crb&dfl$visit==2,4]
    vp<-dfl[(!dfl$is_crb)&dfl$visit==visi,4]-dfl[(!dfl$is_crb)&dfl$visit==2,4]
    vc<-vc[!is.na(vc)]
  }
}

```

```

    vp<-vp[!is.na(vp)]
    tv<-test.numeric.paired.1(vc)
    tvp<-test.numeric.paired.1(vp)
    ret[(visi-2)*5+4,3]<-beautiful.p.value(tvc$p.value)
    ret[(visi-2)*5+4,4]<-beautiful.p.value(tvp$p.value)
    ret[(visi-2)*5+5,3]<-tvc$final.test
    ret[(visi-2)*5+5,4]<-tvp$final.test
  }
  tv<-test.numeric.unpaired(vc, vp)
  ret[(visi-2)*5+4,5]<-beautiful.p.value(tv$p.value)
  ret[(visi-2)*5+5,5]<-(tv$final.test)
}
for(i in 1:nrow(ret))
  for(j in 1:ncol(ret))
    if(is.na(ret[i,j]))ret[i,j]< "-"
return(ret)
}
# Iteration for tabling and graphing FOR ITT ----
if(dir.exists("itt"))
  unlink("itt",recursive = T,force=T)
dir.create("itt")
setwd("itt/")
for(i in 1:2){ # For NIHSS and WAB raw scores only graph them
  dfl<-db_long[,c("RDMNUM","is_crb","visit",db_names$short[i])]
  g<-graph_them_all(dfl,db_names$long[i],paste0(i," ",db_names$long[i]))
  png(paste0(i,db_names$short[i],".png"),width = 800,height = 600)
  print(g)
  dev.off()
}
for(i in 1:2+2){ # For NIHSS and WAB differentials graph them and do tables
  dfl<-db_long[db_long$visit!=1,c("RDMNUM","is_crb","visit",db_names$short[i])]
  g<-graph_them_all(dfl,db_names$long[i],paste0(i," ",db_names$long[i]))
  t<-table_them_all(dfl,already_diff = T)
  png(paste0(i,db_names$short[i],".png"),width = 800,height = 600)
  print(g)
  dev.off()
  write.xlsx(t,paste0(i,db_names$short[i],".xlsx"))
}
for(i in 5:6){ # For mRS and BI graph and make tables, but with already_diff=F
  dfl<-db_long[db_long$visit!=1,c("RDMNUM","is_crb","visit",db_names$short[i])]
  g<-graph_them_all(dfl,db_names$long[i],paste0(i," ",db_names$long[i]))
  t<-table_them_all(dfl,already_diff = F)
  png(paste0(i,db_names$short[i],".png"),width = 800,height = 600)
  print(g)
}

```

```
dev.off()
write.xlsx(t,paste0(i,db_names$short[i],".xlsx"))
}
setwd("..")
# Iteration for tabling and graphing FOR PP ----
db_long<-db_long[db_long$is_pp,]
if(dir.exists("pp"))
  unlink("pp",recursive = T,force=T)
dir.create("pp")
setwd("pp/")
for(i in 1:2){ # For NIHSS and WAB raw scores only graph them
  dfl<-db_long[,c("RDMNUM","is_crb","visit",db_names$short[i])]
  g<-graph_them_all(dfl,db_names$long[i],paste0(i," ",db_names$long[i]))
  png(paste0(i,db_names$short[i],".png"),width = 800,height = 600)
  print(g)
  dev.off()
}
for(i in 1:2+2){ # For NIHSS and WAB differentials graph them and do tables
  dfl<-db_long[db_long$visit!=1,c("RDMNUM","is_crb","visit",db_names$short[i])]
  g<-graph_them_all(dfl,db_names$long[i],paste0(i," ",db_names$long[i]))
  t<-table_them_all(dfl,already_diff = T)
  png(paste0(i,db_names$short[i],".png"),width = 800,height = 600)
  print(g)
  dev.off()
  write.xlsx(t,paste0(i,db_names$short[i],".xlsx"))
}
for(i in 5:6){ # For mRS and BI graph and make tables, but with already_diff=F
  dfl<-db_long[db_long$visit!=1,c("RDMNUM","is_crb","visit",db_names$short[i])]
  g<-graph_them_all(dfl,db_names$long[i],paste0(i," ",db_names$long[i]))
  t<-table_them_all(dfl,already_diff = F)
  png(paste0(i,db_names$short[i],".png"),width = 800,height = 600)
  print(g)
  dev.off()
  write.xlsx(t,paste0(i,db_names$short[i],".xlsx"))
}
setwd("..")
# AE preprocessing ----
# NB: if the AE classification is not complete, i.e. there are still NAs, fill them with the following
command (pay attention to which columns you execute this on)
for(i in 11:22) db_advers[is.na(db_advers[,i]),i]<-0
# NB: if the SAE classification is not complete, ditto
db_advers[is.na(db_advers[,3]),3]<-"no"
# AE information table ----
db_adver2<-db_assign
```

```

for(i in 1:nrow(db_adver2)){
  rdmnum<-db_adver2$RDMNUM[i]
  db_adver2[i,"AE_total"]<-sum(db_advers$RDMNUM==rdmnum)
  db_adver2[i,"SAE_total"]<-sum(db_advers$RDMNUM==rdmnum&db_advers$SAE=="yes")
  for(j in
c("neurological","psychiatric","cardiovascular","renal","hematological","gastrointestinal","metabol
ic","respiratory","immune.related.AE","ENT","ophthalmic","osteoarticular"))
    db_adver2[i,paste0("AE_",j)]<-
sum(db_advers$RDMNUM==rdmnum&db_advers[,paste0("is_",j)]==1)
}
tbl_out<-data.frame(
  AE_category=as.character(),
  sub=as.character(),
  Crb=as.character(),
  Plc=as.character(),
  pval=as.character(),
  test=as.character()
)
db_aelong<-data.frame(
  AE_category=as.character(),
  is_crb=as.logical(),
  no_patients=as.numeric(),
  no_aes=as.numeric(),
  percent_patients=as.numeric()
)
for(i in 7:20){
  rws<-nrow(tbl_out)
  tbl_out[rws+1,1]<-db_names[i,"short"]
  tbl_out[rws+1,5,2]<-c("Patients with AEs in group","Total number of AEs in group","Mean
number of AEs per patient","Median (Q1-Q3) number of AEs per patient","Max number of AEs
per patient")
  xc<-db_adver2[db_adver2$is_crb,db_names[i,"short"]]
  xp<-db_adver2[!db_adver2$is_crb,db_names[i,"short"]]
  tbl_out[rws+1,3]<-paste0(sum(xc!=0),"(",round(sum(xc!=0)/length(xc)*100,3),"%)")
  tbl_out[rws+1,4]<-paste0(sum(xp!=0),"(",round(sum(xp!=0)/length(xp)*100,3),"%)")
  x<-contingency.table.test(table(db_adver2$is_crb,db_adver2[,db_names[i,"short"]]!=0))
  tbl_out[rws+1,5]<-beautiful.p.value(x[["p_chosen"]])
  tbl_out[rws+1,6]<-x[["shorthand"]]
  tbl_out[rws+2,3]<-sum(xc)
  tbl_out[rws+2,4]<-sum(xp)
  tbl_out[rws+3,3]<-descr.numeric.meansd(xc)
  tbl_out[rws+3,4]<-descr.numeric.meansd(xp)
  tbl_out[rws+4,3]<-descr.numeric.median(xc)
  tbl_out[rws+4,4]<-descr.numeric.median(xp)

```

```

tbl_out[rws+5,3]<-max(xc)
tbl_out[rws+5,4]<-max(xp)
x<-test.numeric.unpaired(xc,xp)
tbl_out[rws+3,5]<-beautiful.p.value(x$p.value)
tbl_out[rws+3,6]<-x$final.test
for(j in 1:max(xc,xp)){
  tbl_out[rws+5+j,2]<-paste0("Patient with ",j," AEs")
  tbl_out[rws+5+j,3]<-paste0(sum(xc==j)," (",round(sum(xc==j)/length(xc)*100,3),"%)")
  tbl_out[rws+5+j,4]<-paste0(sum(xp==j)," (",round(sum(xp==j)/length(xp)*100,3),"%)")
}
rws<-(i-7)*2
db_aelong[rws+1:2,1]<-db_names[i,"short"]
db_aelong[rws+1:2,2]<-c(T,F)
db_aelong[rws+1,3]<-sum(xc!=0)
db_aelong[rws+2,3]<-sum(xp!=0)
db_aelong[rws+1,4]<-sum(xc)
db_aelong[rws+2,4]<-sum(xp)
db_aelong[rws+1,5]<-sum(xc!=0)/length(xc)*100
db_aelong[rws+2,5]<-sum(xp!=0)/length(xp)*100
}
for(i in 1:nrow(tbl_out))
  for(j in 1:ncol(tbl_out))
    if(is.na(tbl_out[i,j])) tbl_out[i,j]<-""
write.xlsx(tbl_out,"AE_table.xlsx")
db_aelong$AE_category2<-factor(rep(1:14,each=2),levels=1:14,labels=c(
  "AE\nall types",
  "SAE",
  "Neurol.", "Psych.", "CV", "Renal",
  "Hematol.", "GI", "Metab.", "Resp.",
  "Immune\nrelated", "ENT", "Ophthalmic", "Osteo-\narticular"
))
db_aelong$fill<-factor(db_aelong$is_crb,levels=c(F,T),labels=c(
  paste0("Placebo\n(N=",sum(!db_adver2$is_crb),")"),
  paste0("Cerebrolysin\n(N=",sum(db_adver2$is_crb),")")
))
g<-ggplot(db_aelong,aes(x=AE_category2,y=no_patients,fill=fill))+
  geom_col(position="dodge")+
  scale_fill_manual(values=c("#6495ED", "#ee82ee"))+
  labs(x="AE category",y="Number of patients",fill="Treatment")+
  theme_light()+
  theme(plot.title=element_text(hjust=0),plot.subtitle = element_text(hjust=0),text =
element_text(size = 25))
png("AE_number_of_patients.png",width=2000,height = 600)
print(g)

```

```
dev.off()
g<-ggplot(db_aelong,aes(x=AE_category2,y=no_aes,fill=fill))+
  geom_col(position="dodge")+
  scale_fill_manual(values=c("#6495ED","#ee82ee"))+
  labs(x="AE category",y="Number of AEs",fill="Treatment")+
  theme_light()+
  theme(plot.title=element_text(hjust=0),plot.subtitle = element_text(hjust=0),text =
element_text(size = 25))
png("AE_number_of_AEs.png",width=2000,height = 600)
print(g)
dev.off()
g<-ggplot(db_aelong,aes(x=AE_category2,y=percent_patients,fill=fill))+
  geom_col(position="dodge")+
  scale_fill_manual(values=c("#6495ED","#ee82ee"))+
  labs(x="AE category",y="Percent patients of total (%)",fill="Treatment")+
  theme_light()+
  theme(plot.title=element_text(hjust=0),plot.subtitle = element_text(hjust=0),text =
element_text(size = 25))
png("AE_percent_of_patients.png",width=2000,height = 600)
print(g)
dev.off()
```