



PhUSE US Connect 2019

Paper AB08

A SAS Macro for Variable Ratio Treatment vs. Placebo in Propensity Score Matching

Punjita Baranwal, Cytel, Hyderabad, India
Prashant Kulkarni, Cytel, Pune, India

ABSTRACT

In non-randomized trials there are important differences between treatment groups which give rise to unbalanced comparisons. Therefore, Propensity score matching (PSM) is a statistical matching technique to identify the pairs of subjects from two treatment groups which are similar in their profiles.

These matched groups of subjects can be used to compare two treatments. This will be similar to a randomized trial. There are several matching methods described in literature. Among these Nearest Available Mahalanobis Metric Matching within Caliper is considered optimal (Rosenbaum and Rubin, 1985).

A SAS[®] macro by Feng, Yu and Rong (2006) is available to implement this algorithm but has one limitation. It works for treatment: placebo ratio 1:1. Sometimes one group may have fewer subjects compared to other. In such circumstances one can have one-to-many matching (1:K where $K > 1$) to increase the number of matched pairs. Present paper describes a modification which allows for one-to-many matching.

KEY WORDS: Propensity Score, Caliper, Mahalanobis distance.

INTRODUCTION

Many authors since Fisher (1925), have pointed out that randomization is a powerful tool for taking care of many systematic sources of bias. However, randomization is not always possible, e.g. in observational studies or comparing two treatments from different studies. Drawing causal inference in such situations is challenging due to presence of confounders in unbalanced groups, and this has motivated statisticians to explore new analytic methods.

Commonly used methods for matching based on Propensity Score are:

1. Nearest available matching
2. Mahalanobis metric matching
3. Nearest available Mahalanobis metric matching within caliper

Rosenbaum and Rubin (1985) found the method Nearest Available Mahalanobis Metric matching within Caliper to be the best among the three, as it produced the best balance between the covariates in the treated and control groups, as well as the best balance of the covariates squares and cross-products between the two groups.

These methods can be implemented by using SAS[®] algorithms available in literature. For the best method the macro available is by Feng, Yu and Rong (2006). Sometimes, one of the groups may have fewer subjects compared to the others. In such circumstances, one can have one-to-many matching (1:K where $K > 1$). This helps in reducing bias (Ming and Rosenbaum, 2000). However, the available macro does not allow variable ratio of treated to controls for Matching. Here we describe a modification which overcomes this limitation.



PROPENSITY SCORE MATCHING (PSM)

Propensity score matching (PSM) is a statistical technique to identify the pairs of subjects from two treatment groups, not identical but as similar as possible. In case of non-randomization, it is virtually impossible to be convinced that the estimates of the treatment effects are unbiased. Drawing causal inference in such situations is challenging, and this has motivated statisticians to explore new analytic methods. PSM is often used for this purpose.

If there are many such pairs then the comparison between two groups may be fair. It will be close to a randomized trial. These matched groups of subjects can be used to compare the treatment vs Placebo. PSM attempts to reduce the bias due to confounding variables that could be found in an estimate of the treatment effect obtained from simply comparing outcomes among units that received the treatment versus those did not.

MATCHING ALGORITHM

- (1) Propensity scores are computed for every subject using a logistic regression model with all possible independent variables that may have affected choosing treatment for study participants included.
- (2) The subjects in Control group (the group with less study participants) are randomly ordered and the first subject from this group is selected. Subjects in Treated group whose propensity score is within the caliper (one quarter of standard deviation of the logit of the propensity score as suggested by Rosenbaum and Rubin⁵) are identified as initial matched candidates.
- (3) Three situations may occur in this step:
 - I. No candidate is located within the caliper, then this round of matching will stop here and the next round will start.
 - II. Only one candidate is identified within the caliper and this candidate is considered as final.
 - III. Two or More than two candidates are identified, and then the search process will move to Step 4.
- (4) Mahalanobis distances based on a smaller numbers of key variables and propensity score are calculated between the subject in Control group and those initially selected subjects in Treated group. The treated subject with the smallest distance to the control subject is selected as a final matched candidate. The matched pair is then removed from the pool, and the process will repeat for the next subject in control group. All remaining subjects in Treatment group are available for the remaining matching rounds. The matching process will repeat until the control group is exhausted.
- (5) In the 1:K (where $K > 1$) matching macro, all control are initially matched to their "best" treatment group in the first iteration. An individual treated subject can be picked at most one time. The set of matched controls then matched 1:1 to the set of un-matched treated in the remaining $K-1$ additional iterations.
- (6) If a Control fails to receive a matched treated in the successive iteration, then this un-matched Control are removed completely from matching process from each of the successive iteration. However, the corresponding matched treated subjects are added back to the pool set of unmatched treated, and allowed to match to another control for the next iteration.
- (7) Matched dataset(s) from iteration(s) before the K^{th} iteration are then combined with matched set at K^{th} iteration (this is the final iteration), by keeping only common matched controls from these datasets and pooling all possible matched treated from all iteration.
- (8) The dataset with all these matched pairs is created. This gives the 1:K (where $K > 1$) matching by Mahalanobis optimal matching method.

METHOD

Example Data

Hypothetical data is presented here with some knowledge about covariates. The cases (N = 2629) received the treatment. The controls (N = 600) received Placebo.

Table 1 describes the baseline characteristics of the original population. All of these baseline characteristics were independent variables in the multivariate logistic regression model to create the propensity score. Difference between groups was evaluated by simple t-test for continuous data and the Chi-square test for categorical data.

The propensity score was calculated by using PROC LOGISTIC for multivariate logistic regression model. One quarter of standard deviation of logit of propensity score (about 0.18) was used as a caliper for initial matching of subjects.

Table 1: Original Population (Un-matched Data)			
	Placebo	Treatment	p-value
Total Patients, n	600	2629	
Age, mean (SD), years	34.7 (9.17)	36.2 (9.04)	0.0002
Sex - Female, n (%)	392 (65.33)	1806 (68.70)	0.111
Country - USA, n (%)	177 (29.50)	472 (17.95)	<.0001
Ethnicity - Not Hispanic, n (%)	357 (59.50)	723 (27.55)	<.0001
Race - White, n (%)	363 (60.50)	1432 (54.51)	<.0001
Height, mean (SD), cm	169.6 (9.05)	169.1 (9.19)	0.1899
Weight, mean (SD), kg	75.5 (19.15)	72.3 (16.91)	0.0002
Baseline EDSS score, mean (SD)	2.2 (1.08)	2.5 (1.27)	<.0001
Time since first MS symptom, mean (SD), years	3.8 (4.26)	6.5 (6.22)	<.0001
Number of relapses past 1 year, mean (SD)	1.6 (0.83)	1.5 (0.76)	0.0065
Number of relapses past 2 years, mean (SD)	2.4 (0.96)	2.3 (1.12)	0.1893
Number of prior DMTs taken, mean (SD)	0.5 (0.71)	0.4 (0.71)	0.3036
Propensity Score	1.3 (0.53)	1.7 (0.76)	<.0001

1: K MATCH (WHERE K=1)

Figure 1 displays what the matching macro does in the initial 1:1 match iteration. The numbers of records at each step are from the 1:1 match of the example data.

Step (A): The macro creates separate files for the treated and controls.

Step (B): Next, it performs matching using Mahalanobis optimal matching technique.

Step (C): It then outputs a file of the matched pairs. This file contains the field to link the matched pairs (MATCH1).

Step (D): Finally, it outputs a file of un-matched cases. This will be the pool from which the matches are made in the next iteration.

Figure 1→ 1: K Match Where K=1

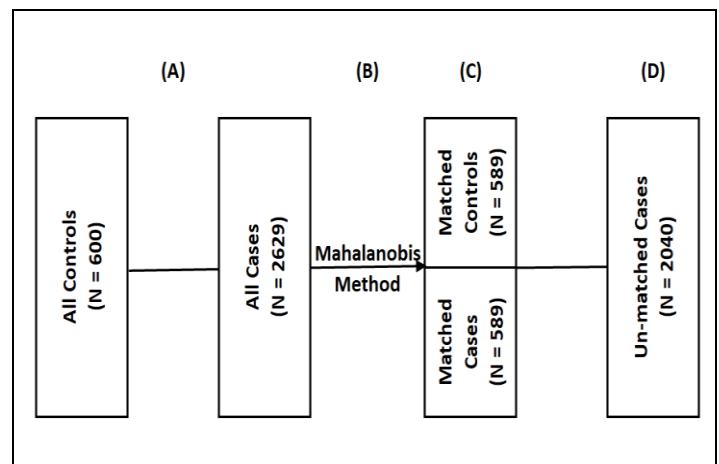


Table 2 describes the matched population based on the matching algorithm implemented in our macro. In summary, 589 patients from treatment group were successfully matched by 589 patients from Placebo group A. Difference between the matched pairs is evaluated using t-test for continuous variables and Chi-Square test for categorical ones. So, after matching, there is no significant difference between the two matched groups among all the baseline variables, which shows that the two matched groups are well balanced.

Table 2: 1:1 Matched Population			
	Placebo	Treatment	p-value
Total Patients, n	589	589	
Age, mean (SD), years	34.6 (9.17)	34.6 (8.70)	0.9974
Sex - Female, n (%)	385 (65.37)	391 (66.38)	0.7123
Country - USA, n (%)	177 (30.05)	131(22.24)	<.0001
Ethnicity - Not Hispanic, n (%)	352 (59.76)	210 (35.71)	<.0001
Race - White, n (%)	358 (60.78)	337 (57.31)	<.0001
Height, mean (SD), cm	169.6 (9.05)	169.6 (8.43)	0.9735
Weight, mean (SD), kg	75.5 (19.19)	75.0 (17.15)	0.6535
Baseline EDSS score, mean (SD)	2.2 (1.07)	2.2 (1.03)	0.9448
Time since first MS symptom, mean (SD), years	3.7 (4.23)	3.6 (4.05)	0.6085
Number of relapses past 1 year, mean (SD)	1.6 (0.83)	1.6 (0.74)	0.6572
Number of relapses past 2 years, mean (SD)	2.4 (0.96)	2.3 (1.05)	0.7724
Number of prior DMTs taken, mean (SD)	0.5 (0.71)	0.4 (0.71)	0.2882
Propensity Score	1.3 (0.53)	1.3 (0.52)	0.9872

1: K MATCH (WHERE K=2)

Figure 2 displays what the matching macro does for the remaining iterations. The numbers of records at each step are from the 1:2 match of the example data.

Step (A): The macro starts with the previously matched controls and pool of un-matched treated.

Step (B): It then performs matching using Mahalanobis optimal matching technique.

Step (C): Next, it outputs a file of new matched pairs. This file also contains the field to link this set of matched pairs (MATCHK, where K is the number of the iteration). Note that any cases from a previous iteration, that did not receive a match in the current iteration, are no longer in the current set of matched pairs.

Step (D): The macro then outputs the file of un-matched treated.

Step (E): Next it determines which treated matched to a control in a previous iteration. If that control still exists in the current set of matched pairs, the macro appends the treated to the file of current matched pairs.

Step (F): Finally, it determines which treated matched to a control in a previous iteration, but whose matched control did not receive a match in the current iteration. These treated are appended to the file of un-matched treated. Thus, these treated are added back to the pool of available treated to be potential matches in the next iteration.

Figure 2 → 1: K Where K = 2

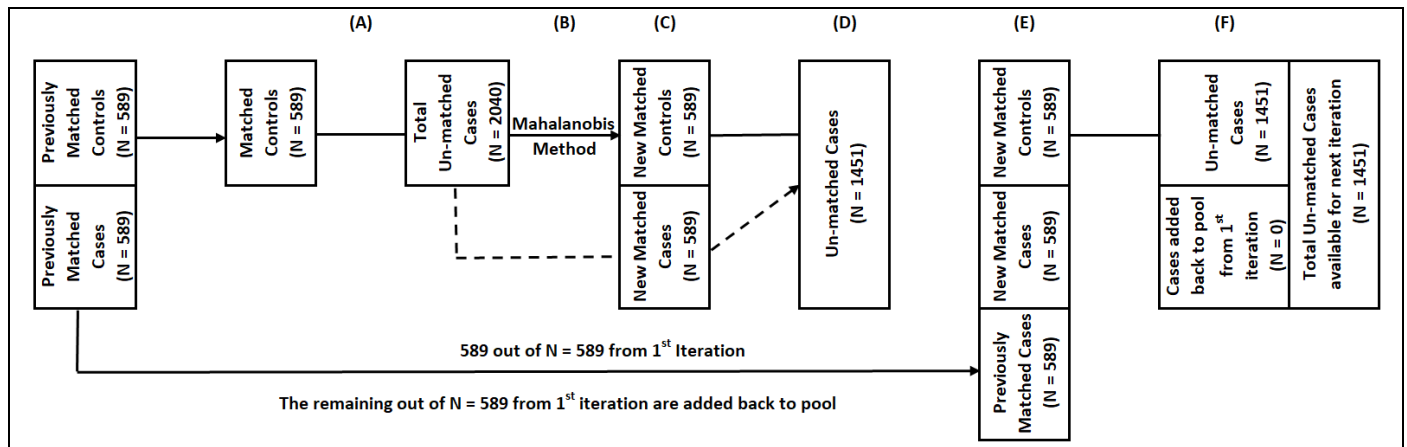


Table 3 describes the matched population based on the matching algorithm implemented in our macro. In summary, 1178 patients from treatment group were successfully matched by 589 patients from Placebo group A. Difference between the matched pairs is evaluated using t-test for continuous variables and Chi-Square test for categorical ones. So, after matching, there is no significant difference between the two matched groups among all the baseline variables, which shows that the two matched groups are well balanced.

Table 3: 1:2 Matched Population			
	Placebo	Treatment	p-value
Total Patients, n	589	1178	
Age, mean (SD), years	34.6 (9.17)	34.5 (8.53)	0.8263
Sex - Female, n (%)	385(65.37)	781 (66.30)	0.6961
Country - USA, n (%)	177 (30.05)	267 (22.67)	<.0001
Ethnicity - Not Hispanic, n (%)	352 (59.76)	431 (36.65)	<.0001
Race - White, n (%)	358 (60.78)	669 (56.84)	<.0001
Height, mean (SD), cm	169.6(9.05)	169.7 (8.67)	0.8625
Weight, mean (SD), kg	75.5 (19.19)	74.8 (16.94)	0.4413
Baseline EDSS score, mean (SD)	2.2 (1.07)	2.2 (1.04)	0.6891
Time since first MS symptom, mean (SD), years	3.7 (4.23)	3.7 (4.07)	0.7175
Number of relapses past 1 year, mean (SD)	1.6 (0.83)	1.6 (0.78)	0.5829
Number of relapses past 2 years, mean (SD)	2.4 (0.96)	2.3 (1.05)	0.8387
Number of prior DMTs taken, mean (SD)	0.5 (0.71)	0.4 (0.74)	0.419
Propensity Score	1.3 (0.53)	1.3 (0.52)	0.89

Figure 3 → 1: K Where K = 3 (Execution of next iteration after second iteration shown in Figure 2)

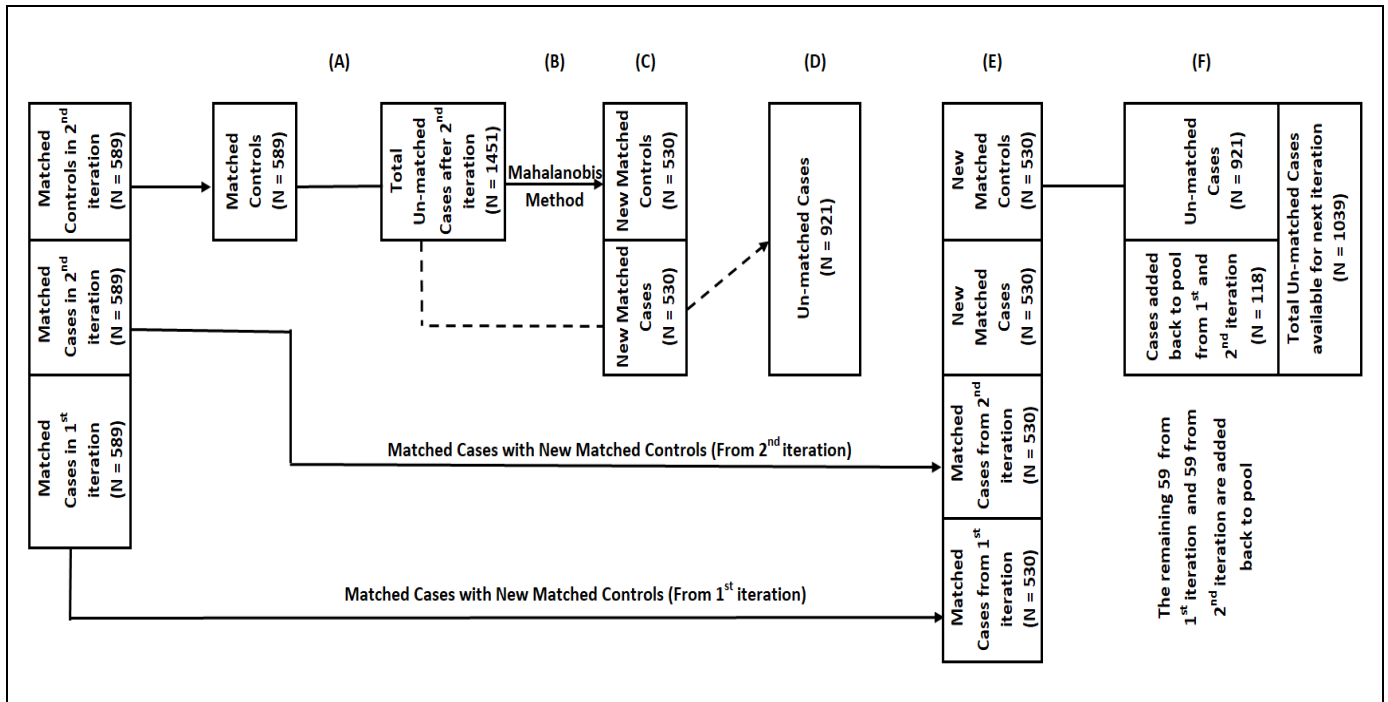


Table 4 describes the matched population based on the matching algorithm implemented in our macro. In summary, 1590 patients from treatment group were successfully matched by 530 patients from Placebo group A. Difference between the matched pairs is evaluated using t-test for continuous variables and Chi-Square test for categorical ones. So, after matching, there is no significant difference between the two matched groups among all the baseline variables, which shows that the two matched groups are well balanced.

Table 4: 1:3 Matched Population			
	Placebo	Treatment	p-value
Total Patients, n	530	1590	
Age, mean (SD), years	34.7 (9.24)	34.5 (8.83)	0.6016
Sex - Female, n (%)	356 (67.17)	1069 (67.23)	0.9787
Country - USA, n (%)	142 (26.79)	316 (19.87)	<.0001
Ethnicity - Not Hispanic, n (%)	315 (59.43)	565 (35.58)	<.0001
Race - White, n (%)	324 (61.13)	911 (57.33)	<.0001
Height, mean (SD), cm	169.4 (8.97)	169.2 (8.97)	0.7332
Weight, mean (SD), kg	73.5 (17.31)	72.8 (16.78)	0.4172
Baseline EDSS score, mean (SD)	2.3 (1.06)	2.3 (1.09)	0.9173
Time since first MS symptom, mean (SD), years	3.9 (4.14)	3.9 (4.00)	0.9869
Number of relapses past 1 year, mean (SD)	1.6 (0.84)	1.6 (0.78)	0.7161
Number of relapses past 2 years, mean (SD)	2.4 (0.98)	2.4 (1.09)	0.9406
Number of prior DMTs taken, mean (SD)	0.4 (0.68)	0.4 (0.70)	0.6916
Propensity Score	1.4 (0.50)	1.4 (0.49)	0.7213

MATCHED SAMPLE

The results of performing a 1:1, 1:2 and 1:3 case-control match on the example data are given in the above tables. For the matched analysis, differences between matched pairs were evaluated using the t-test for continuous data and Chi-square test for categorical data. With the exception of Country, Race, and Ethnicity there is no longer a significant difference between the characteristics of the treatment (cases) and Placebo (controls). **Table 2** describes the descriptive statistics in 1:1 matched population, **Table 3** describes the descriptive statistics in 1:2 matched population, and **Table 4** describes the descriptive statistics in 1:3 matched population.

INCOMPLETE MATCHES

If the data for some covariates are missing for some subjects, then that case will be completely removed from the analysis and propensity score will be not calculated for that case. In case of disjoint ranges of case and control propensity scores (if no candidate is located within the caliper) also results in incomplete matching.

CONCLUSIONS

The macro presented here will allow an analyst to perform a 1: K (where $K > 1$) case-control match on propensity score. The match will be a well-matched sample and contain well-matched pairs. By applying the macro presented in this paper, subjects in treatment group can be successfully matched to Placebo group using Mahalanobis metric matching within the calipers defined by propensity score method. Such matching process would in general result in two treatment groups having very similar baseline characteristics, thus treatment difference can be appropriately estimated. The source SAS code is attached in the end of this paper and SAS users could apply this macro in their own research work.

A limitation to this method is the multivariate model from which the propensity score was computed. A poor model will result in a poor predicted probability of the outcome (propensity score). This method can only reduce bias in measured characteristics.

Also, the analyst should be very careful for deciding the number of iteration (i.e. K where $K > 1$), as it may reduce the sample size resulting in decrease in power of the trial.



REFERENCES

1. Feng, Yu and Rong "A Method/Macro Based on Propensity Score and Mahalanobis Distance to Reduce Bias in Treatment Comparison in Observational Study", 2006.
2. Lori S. Parsons, "Performing a 1:N Case-Control Match on Propensity Score"
3. Rosenbaum, P. R. and Rubin, D. B. (1983). "The Central Role of the Propensity Score in Observational Studies for Causal Effects". *Biometrika* Vol. 70 (1), pp 41– 55. doi:[10.1093/biomet/70.1.41](https://doi.org/10.1093/biomet/70.1.41)
4. Dehejia R., H., and Wahba, S. (2002). "Propensity Score - Matching Methods For Non-experimental Causal Studies", *The Review of Economics and Statistics*, Vol. 84(1), pp 151–161. <http://www.uh.edu/~adkugler/Dehejia&Wahba.pdf>
5. Cochran, W. G. and Rubin, D. B. "Controlling bias in observational studies: A review". *Sankhya*, Ser. A. 35 (1973), 417-446



APPENDIX 1:

```
/**--Start of the Program--**/;

libname XXXX 'X:\XXXX'; /** Please specify the library name and location **/;

/**--Dataset used for PSM containing all the Case and Control subject information--
**/;
data test;
    set XXXX.test;
run;

/**--Start of Macro for 1:K Propensity Score Matching using Mahalanobis Distance--**/;

%macro match(data, class, yvar, xvar, var, id, K);

/**--Defining Variable--**/
    1) data- The dataset containing all the case and control subject information
    2) class-The categorical independent variables for calculating Propensity Score
        in PROC LOGISTIC
    3) yvar- Treatment variable i.e. Treatment=0 for Control and Treatment=1 for
        Case. Also, the dependent variable for calculating Propensity Score in PROC
        LOGISTIC
    4) xvar-The set of independent variables for calculating Propensity Score in
        PROC LOGISTIC
    5) var-The key variables (including propensity score as one key variable) for
        calculating the Mahalanobis distance
    6) id- The Patient ID i.e. USUBJID
    7) K- Variable Ratio/The number of iteration for propensity score matching **/;

/**--Perform logistic regression, select independent variables, and create the
propensity score--**/;
proc logistic descending data=&data noprint;
    class &class;
    model &yvar=&xvar;
    output out=propen(drop=_level_) prob=prob xbeta=logit;
run;

proc univariate data=propen noprint;
    var logit;
    output out=propen_sd std=sd;
run;

/**--Create datasets of case and control, which contain patients from case and control
groups--**/;
data ctrl case;
    set propen;
    if &yvar=0 and prob ne . then do;
        rannum=ranuni(1234567);
        label rannum='Uniform Randomization Score';
        output ctrl;
    end;
    else if &yvar=1 and prob ne . then do;
        output case;
    end;
run;

data ctrl;
    if _n_=1 then set propen_sd(keep=sd);
    set ctrl;
run;

/**--Start of 1:K matching--**/;

%do num=1 %to &K;
```




```
/**--Set up a dataset which will contain all control patients and matched case
patients--***/;
data case_ctrl_together; run;

/**--Initial Dataset for matching--***/;
data ctrl1; set ctrl; run;
data case1; set case; run;

/**--Sort ctrl by random number generated--***/;
proc sort data=ctrl&num; by rannum; run;

data _null_;
    set ctrl&num;
    by sd;
    if last.sd then call symput('tot',put(_n_,8.));
run;

data case_temp;
    set case&num;
/**--Select one patient once time from dataset ctrl, and search for matched-pair in
case group--***/;
    %let j=%eval(&tot);
    %do i=1 %to &j;
data ctrl_temp;
    set ctrl&num;
    if _n_=&i;
    low=logit - 0.25*sd;
    up =logit + 0.25*sd;
run;

data match_temp;
    if _n_=1 then set ctrl_temp(keep=low up);
    set case_temp(drop=rannum);
    if low<=logit<=up;
run;
/**--Calculation number of patients who were included in the match dataset **/
    1) if there is no match found, then select next patient from case group
    2) if there is one matched, select this patient and go to next round
    3) if there are more than 1 patients, calculate Mahalanobis distance and select
        the patient with smallest distance, then go to next round--***/;
proc sql;
    create table casen as select count(&id) as n from match_temp;
quit;

data _null_;
    set casen;
    if _n_=1 then call symput('casen',put(n,8.));
run;

%if %eval(&casen)=0 %then %do;
data match_temp1;
    set match_temp;
    newid=&i;
    m_casen=&casen;
run;

data case_ctrl_together;
    set case_ctrl_together match_temp1(keep=&id newid m_casen);
    if &id='' then delete;
run;
%end;

%else %if %eval(&casen)=1 %then %do;
data match_temp1;
    set ctrl_temp match_temp;
    newid=&i;
    m_casen=&casen;
```



```
run;

data case_ctrl_together;
    set case_ctrl_together match_temp1(keep=&id newid m_casen);
    if &id=' ' then delete;
run;

proc sql;
    create table case_temp
    as select a.*
    from case_temp a, match_temp b
    where a.&id^=b.&id;
quit;
%end;

%else %do;
    /**--Please specify the key variables (including propensity score as one key variable)
    which are used to calculate the mahalanobis distance--**/;
proc princomp data=match_temp out=out outstat=outstat noprint;
    var &var;
run;

proc score data=ctrl_temp score=outstat out=reference_point;
    var &var;
run;

proc append data=out base=reference_point force nowarn;
run;

proc fastclus data=reference_point maxc=1 replace=none maxiter=0 noprint
    out=mahalanobis_to_point(drop=cluster);
    var prin;;
run;

proc sql undo_policy=none;
    create table mahalanobis_to_point
    as select a.*
    from mahalanobis_to_point a, ctrl_temp b
    where a.&id^=b.&id;
quit;

proc sort data=mahalanobis_to_point;
    by distance;
run;

data mahalanobis_to_point;
    set mahalanobis_to_point;
    by distance;
    if _n_=1;
run;

data case_ctrl;
    set ctrl_temp(keep=&id) mahalanobis_to_point(keep=&id);
    newid=&i;
    m_casen=&casen;
run;

data case_ctrl_together;
    set case_ctrl_together case_ctrl;
    if &id=' ' then delete;
run;

/**--Exclude patients in case group which are selected in the Mahalanobis macro--**/;
proc sql undo_policy=none;
    create table case_temp
    as select a.*
    from case_temp a, mahalanobis_to_point b
```



```
        where a.&id^=b.&id;
quit;
%end;
%end;

/**--Create test dataset that contains patients from case_control_together, and all
variables and propensity scores---**/;
proc sql undo_policy=none;
    create table test as select a.*, b.newid, b.m_casen
        from propen a, case_ctrl_together b
        where a.&id=b.&id order by newid,&yvar;
quit;

proc transpose data=test out=test_tran&num (drop=newid _name_ _label_ ) ;
    by newid;
    id &yvar;
    var &id;
run;
Proc sort data=test_tran&num ; by _0; Run;

/**--Create Final matched dataset with all N matched cases---**/
/**--Create dataset for Cases and controls for next round of matching---**/;

proc sort data=ctrl&num ; by &id; run;
proc sort data=case&num ; by &id; run;

proc sql undo_policy=none;
%if &num=1 %then %do;
    create table match&num as select * from test order by newid,&yvar;

    create table test&num as select * from test order by &id;

    create table ctrl%eval(&num+1) as select a.*
        from ctrl&num a inner join test&num b on a.&id=b.&id;

    create table case%eval(&num+1) as select a.*
        from case a left join test&num b on a.&id=b.&id where b.&id is null;
%end;

%if &num ge 2 %then %do;

    %if &num=2 %then %do;
        create table mtch&num as select a.*, b._1 as _1&num, monotonic() as newid
            from test_tran%eval(&num-1) a left join test_tran&num b on
a._0=b._0 where (_1&num ne '');
        %end;
    %else %do;
        create table mtch&num as select a.*, b._1 as _1&num, monotonic() as newid
            from mtch%eval(&num-1) a left join test_tran&num b on a._0=b._0
where (_1&num ne '');
        %end;
quit;

proc transpose data=mtch&num out=mtch_tran&num (where=( _name_ ne 'newid')
rename=(coll=&id)) ;
    by newid;
    var _0--_1&num;
run;

proc sql undo_policy=none;
    create table test&num as select * from test order by &id;

    create table ctrl%eval(&num+1) as select a.*
        from ctrl&num a inner join test&num b on a.&id=b.&id;

    create table case%eval(&num+1) as select a.*
        from case a left join mtch_tran&num b on a.&id=b.&id where b.&id is null;
```



```
        create table match&num as select a.*, b.newid
        from propen a, mtch_tran&num b
        where a.&id=b.&id order by newid,&yvar;
quit;
%end;
%end;
%mend match;
/**--This is the end of macro--**/;
options mprint;

/**--Calling the macro by specifying the variable name--**/;

%match(data=test, class=sex ethnic country, yvar=trtn, xvar=age sex weight height
edblres mxttsy mxrlnum1 mxrlnum2 pdmt, var=age weight height edblres logit,
id=usubjid, K=3);

/**--End of the Program--**/;
```



ACKNOWLEDGEMENT

The authors would like to thank all the colleagues at Cytel who gave their valuable inputs as and when required.

CONTACT INFORMATION

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

Author Name: Punjita Baranwal
Company: Cytel Statistical Software & Services Pvt. Ltd.
Address: Waverock, Survey No.115, Nanakramguda, Hyderabad, India
City / Postcode: Hyderabad 500032
Work Phone: Desk +91.40.6635.0525
Email: Punjita.Baranwal@cytel.com

Author Name: Prashant Kulkarni
Company: Cytel Statistical Software & Services Pvt. Ltd.
Address: Lohia-Jain IT Park– A Wing, Survey No.150, Paud Road, Kothrud, India
City / Postcode: Pune 411038
Work Phone: Desk +91. 20.6709.0175
Email: Prashant.Kulkarni@cytel.com

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