

SAS/STAT[®] 14.2 User's Guide

The PSMATCH Procedure

This document is an individual chapter from *SAS/STAT® 14.2 User's Guide*.

The correct bibliographic citation for this manual is as follows: SAS Institute Inc. 2016. *SAS/STAT® 14.2 User's Guide*. Cary, NC: SAS Institute Inc.

SAS/STAT® 14.2 User's Guide

Copyright © 2016, SAS Institute Inc., Cary, NC, USA

All Rights Reserved. Produced in the United States of America.

For a hard-copy book: No part of this publication may be reproduced, stored in a retrieval system, or transmitted, in any form or by any means, electronic, mechanical, photocopying, or otherwise, without the prior written permission of the publisher, SAS Institute Inc.

For a web download or e-book: Your use of this publication shall be governed by the terms established by the vendor at the time you acquire this publication.

The scanning, uploading, and distribution of this book via the Internet or any other means without the permission of the publisher is illegal and punishable by law. Please purchase only authorized electronic editions and do not participate in or encourage electronic piracy of copyrighted materials. Your support of others' rights is appreciated.

U.S. Government License Rights; Restricted Rights: The Software and its documentation is commercial computer software developed at private expense and is provided with RESTRICTED RIGHTS to the United States Government. Use, duplication, or disclosure of the Software by the United States Government is subject to the license terms of this Agreement pursuant to, as applicable, FAR 12.212, DFAR 227.7202-1(a), DFAR 227.7202-3(a), and DFAR 227.7202-4, and, to the extent required under U.S. federal law, the minimum restricted rights as set out in FAR 52.227-19 (DEC 2007). If FAR 52.227-19 is applicable, this provision serves as notice under clause (c) thereof and no other notice is required to be affixed to the Software or documentation. The Government's rights in Software and documentation shall be only those set forth in this Agreement.

SAS Institute Inc., SAS Campus Drive, Cary, NC 27513-2414

November 2016

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.

SAS software may be provided with certain third-party software, including but not limited to open-source software, which is licensed under its applicable third-party software license agreement. For license information about third-party software distributed with SAS software, refer to <http://support.sas.com/thirdpartylicenses>.

Chapter 95

The PSMATCH Procedure

Contents

Overview: PSMATCH Procedure	7676
Process of Propensity Score Analysis	7677
Features of the PSMATCH Procedure	7679
Getting Started: PSMATCH Procedure	7680
Syntax: PSMATCH Procedure	7688
PROC PSMATCH Statement	7688
ASSESS Statement	7691
BY Statement	7694
CLASS Statement	7695
FREQ Statement	7695
MATCH Statement	7695
OUTPUT Statement	7701
PSDATA Statement	7702
PSMODEL Statement	7702
STRATA Statement	7703
Details: PSMATCH Procedure	7704
Observational Studies Contrasted with Randomized Trials	7704
Propensity Score Analysis	7705
Propensity Score Weighting	7707
Propensity Score Stratification	7709
Matching Process	7709
Matching Statistics	7711
Matching Methods	7712
Variable Balance Assessment	7713
Table Output	7716
ODS Table Names	7717
Graphics Output	7718
ODS Graphics	7719
Examples: PSMATCH Procedure	7720
Example 95.1: Propensity Score Weighting	7721
Example 95.2: Propensity Score Stratification	7728
Example 95.3: Optimal Variable Ratio Matching	7737
Example 95.4: Greedy Nearest Neighbor Matching	7742
Example 95.5: Matching with Replacement	7751
Example 95.6: Mahalanobis Distance Matching	7756
Example 95.7: Matching with Existing Propensity Scores in the Input Data Set	7761
References	7766

Overview: PSMATCH Procedure

In a randomized study, such as a randomized controlled trial, the subjects are randomly assigned to a treated (exposure) group or a control (non-exposure) group. Random assignment ensures that the distribution of the covariates is the same in both groups, and the treatment effect can be estimated by directly comparing the outcomes for the subjects in the two groups.

In contrast, the subjects in an observational study, such as a retrospective cohort study or a nonrandomized clinical trial, are not randomly assigned to the treated and control groups. Confounding can occur if some covariates are related to both the treatment assignment and the outcome. Consequently, there can be systematic differences between the treated subjects and the control subjects. In the presence of confounding, statistical approaches are required that remove the effects of confounding when estimating the effect of treatment.

One such approach is regression adjustment, which estimates the treatment effect after adjusting for differences in the baseline covariates. However, this approach has practical limitations, as discussed by Austin (2011a). Propensity score analysis is an alternative approach that circumvents many of these limitations.

The propensity score was defined by Rosenbaum and Rubin (1983, p. 47) as the probability of assignment to treatment conditional on a set of observed baseline covariates. Propensity score analysis minimizes the effects of confounding and provides some of the advantages of a randomized study. The basis for propensity score methods is the causal effect model introduced by Rubin (1974).

The PSMATCH procedure provides a variety of tools for propensity score analysis. The procedure either computes propensity scores or reads previously-computed propensity scores, and it provides the following methods for using the scores to allow for valid estimation of treatment effect in a subsequent outcome analysis:

- **Inverse probability of treatment weighting and ATT weighting (weighting by odds):** The procedure computes weights from the propensity scores. These weights can then be incorporated into a subsequent analysis that estimates the effect of treatment.
- **Stratification:** The procedure creates strata of observations that have similar propensity scores. In a subsequent analysis, the treatment effect can be estimated within each stratum, and the estimates can be combined across strata.
- **Matching:** The procedure matches each treated unit with one or more control units that have a similar value of the propensity score. In a subsequent analysis, the treatment effect can be estimated by comparing outcomes between treated and control subjects in the matched sample. If the outcome values for a study are not available prior to matching, only the matched units are needed for follow-up. Thus, the cost of the trial is reduced (Stuart 2010, p. 2).

The PSMATCH procedure also provides methods for assessing the balance of baseline covariates and other variables in the treated and control groups after matching, weighting, or stratification. The procedure itself does not carry out the outcome analysis, nor does it make use of the outcome variable.

After adequate variable balance has been achieved (as described in the section “[Process of Propensity Score Analysis](#)” on page 7677), and assuming that no other confounding variables are associated with both the treatment assignment and the outcome, the output data set that is created by the PSMATCH procedure serves as input for an appropriate statistical procedure for the outcome analysis.

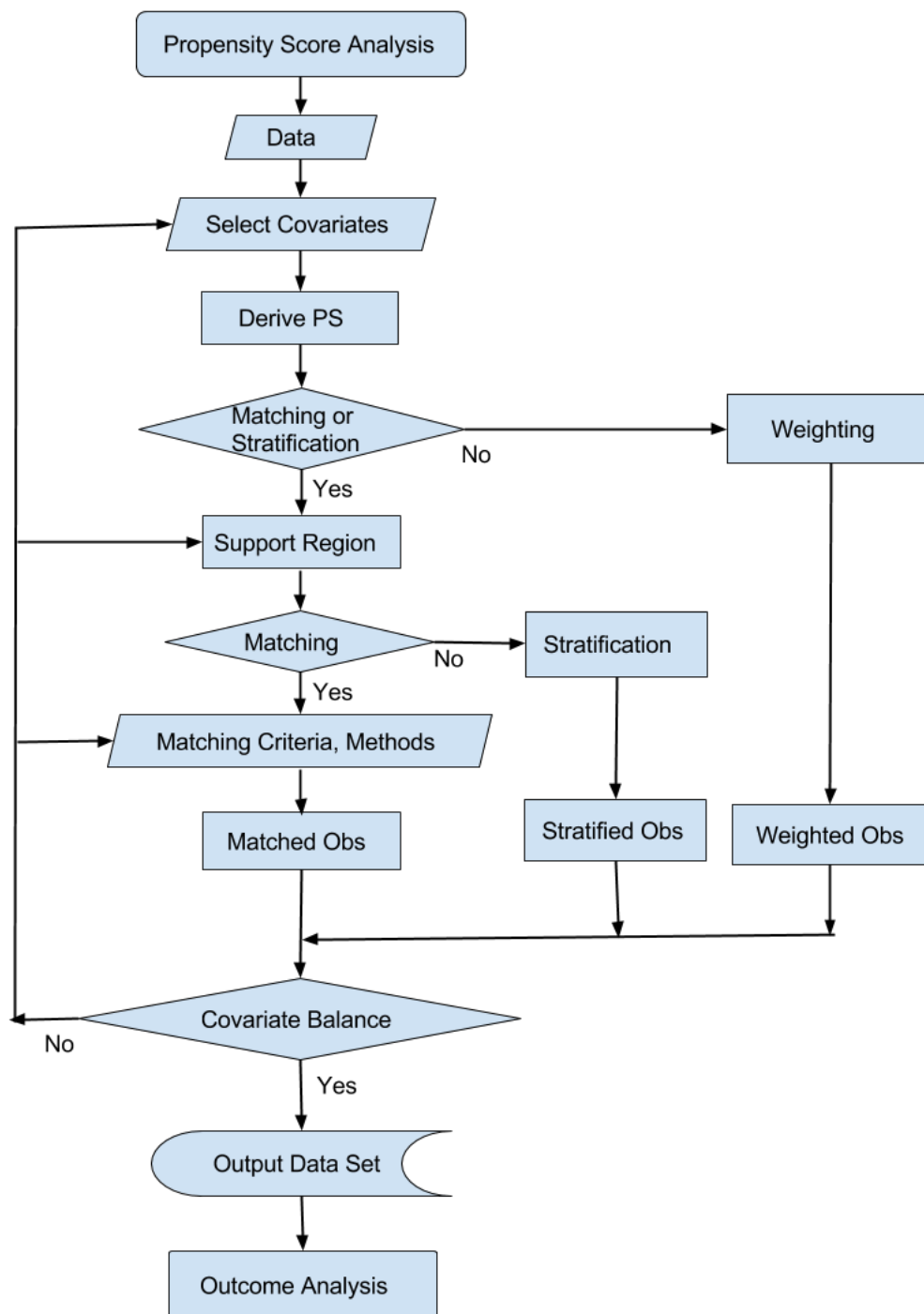
Process of Propensity Score Analysis

A propensity score analysis usually involves the following steps (Guo and Fraser 2015, p. 131):

1. You specify a set of confounding variables that might be related to both the treatment assignment and the outcome.
2. You use this set of variables to fit a logistic regression model and compute propensity scores. The response is the probability of assignment to the treatment group.
3. If you are using weighting, you compute observation weights for estimating the treatment effect in a weighted outcome analysis.
4. If you are using stratification or matching, you specify the support region of observations. Observations outside this region are not included in the stratification or matching.
5. If you are using stratification, you specify the number of strata and create strata of observations that have similar propensity scores.
6. If you are using matching, you specify criteria such as the matching statistic (the distance metric for comparing the similarity of subjects) and the method for creating matched sets of observations. You can also compute weights for matched observations.
7. You assess the balance of variables by comparing the distributions between the treated and control groups.
8. To improve the balance, you can repeat the process with a different set of variables for the logistic regression model, a different region of support for stratification and matching, a different set of matching criteria, or a different matching method.
9. When you are satisfied with the variable balance, you save the output data set for subsequent outcome analysis.

Note that the outcome variable is intentionally not used in this process, and the variable selection is not related to the observed outcomes (Rubin 2001; Stuart 2010, p. 5). Any variables that might have been affected by the treatment should not be included in the process (Rosenbaum and Rubin 1984; Stuart 2010, p. 5).

The flowchart in [Figure 95.1](#) summarizes these steps.

Figure 95.1 Propensity Score Analysis

After balance is achieved, you can add the response variable to the output data set that PROC PSMATCH created and perform an outcome analysis that mimics the analysis you would perform with data from a randomized study. For example, if you used matching with the PSMATCH procedure, a simple univariate test or analysis might be sufficient to estimate treatment effect.

Features of the PSMATCH Procedure

You can use the PSMATCH procedure to create propensity scores (PS) for observations from treated and control groups by fitting a binary logistic regression model. Alternatively, you can input propensity scores that have already been created by using a different model or even a different approach such as a tree-based method. For example, you can input propensity scores that have been computed by the LOGISTIC procedure using a binary probit model or by the HPSPLIT procedure using a classification tree.

By default, the PSMATCH procedure uses the propensity scores to compute weights for the observations. Various types of weights are available, depending on whether the outcome analysis will use the weights to estimate the average treatment effect at the population level (ATE) or the average treatment effect for subjects who receive treatment (ATT). For more information about propensity score weighting, see the section “[Propensity Score Weighting](#)” on page 7707.

The PSMATCH procedure optionally creates strata of observations that have similar propensity scores. For more information, see the section “[Propensity Score Stratification](#)” on page 7709.

The PSMATCH procedure optionally matches observations in the treated and control groups. The procedure provides three strategies for propensity score matching.

- Greedy nearest neighbor matching selects the control unit nearest to each treated unit. Greedy nearest neighbor matching is done sequentially for treated units and without replacement.
- Optimal matching selects all control units that match each treated unit by minimizing the total absolute difference in propensity score across all matches. Optimal matching selects all matches simultaneously and without replacement. Three methods for optimal matching are available: fixed ratio matching, variable ratio matching, and full matching.
- Matching with replacement selects the control unit that best matches each treated unit. Each control unit can be matched to more than one treated unit, but it can only be matched to the same treated unit once.

For all three matching methods, you can specify a caliper width which imposes a restriction on the quality of the matches. The difference in propensity score between the treated unit and its matching control unit must be less than or equal to the caliper width. For more information about these methods, see the section “[Matching Methods](#)” on page 7712.

Matching can be based on the difference in the logit of the propensity score (LPS), as well as the difference in the propensity score (PS). Furthermore, matching can be based on Mahalanobis distance computed from a set of continuous covariates (possibly including LPS and PS).

The PSMATCH procedure provides various ways to assess how well the distributions of variables are balanced between the treated and control groups. These variables include the propensity score, the logit of the propensity score, variables used in the logistic regression model, and other variables in the data set. The assessments include the following:

- differences in the distributions of the variables between the treated and control groups after weighting, stratification, and matching
- standardized differences in the variables between the treated and control groups after weighting, stratification, and matching

- percentage reductions of absolute differences after weighting, stratification, and matching.

When you use stratification, the differences are also computed within each stratum. For more information about these statistics, see the section “[Variable Balance Assessment](#)” on page 7713.

The PSMATCH procedure also provides various plots for assessing balance. These plots include the following:

- cloud plots, which are scatter plots in which the points are jittered to prevent overplotting
- box plots for continuous variables
- bar charts for classification variables
- a standardized differences plot that summarizes differences between the treated and control groups

When you use stratification, the plots are also produced by stratum.

The PSMATCH procedures saves propensity scores and weights in an output data set that contains a sample that has been adjusted either by weighting, stratification, or matching. If the sample is stratified, you can save the strata identification in the output data set. If the sample is matched, you can save the matching identification in the output data set.

Provided that the distributions of the variables in the adjusted sample are well balanced between the treated and control groups, the output data set serves as input for subsequent outcome analysis that incorporates weights or strata or that is based on matched observations. Although the PSMATCH procedure itself does not provide this analysis, many other SAS/STAT procedures can be used for this purpose.

Getting Started: PSMATCH Procedure

This example illustrates the use of the PSMATCH procedure to match observations for individuals in a treatment group with observations for individuals in a control group that have similar propensity scores. The matched observations are saved in an output data set which, with the addition of the outcome variable, can be used to provide an unbiased estimate of the treatment effect.

A pharmaceutical company is conducting a nonrandomized clinical trial to demonstrate the efficacy of a new treatment (Drug_X) by comparing it to an existing treatment (Drug_A). Patients in the trial can choose the treatment that they prefer; otherwise, physicians assign each patient to a treatment. The possibility of treatment selection bias is a concern because it can lead to systematic differences in the distributions of the baseline variables in the two groups, resulting in a biased estimate of treatment effect.

The data set `Drugs` contains baseline variable measurements for individuals from both treated and control groups. `PatientID` is the patient identification number, `Drug` is the treatment group indicator, `Gender` provides the gender, `Age` provides the age, and `Bmi` provides the body mass index (a measure of body fat based on height and weight). Typically, more variables are used in a propensity score analysis. In this example, only a few variables are used for a simple illustration of the use of the PSMATCH procedure.

[Figure 95.2](#) lists the first 10 observations.

Figure 95.2 Input Drug Data Set**First 10 Observations of the Input Drug Data Set**

Obs	PatientID	Drug	Gender	Age	Bmi
1	284	Drug_X	Male	29	22.02
2	201	Drug_A	Male	45	26.68
3	147	Drug_A	Male	42	21.84
4	307	Drug_X	Male	38	22.71
5	433	Drug_A	Male	31	22.76
6	435	Drug_A	Male	43	26.86
7	159	Drug_A	Female	45	25.47
8	368	Drug_A	Female	49	24.28
9	286	Drug_A	Male	31	23.31
10	163	Drug_X	Female	39	25.34

Note that the Drugs data set does not contain a response variable, because the response variable is not used in propensity score method. Instead, the response variable is added to the output data set of matched observations after matching for the outcome analysis.

The following statements invoke the PSMATCH procedure and request optimal matching to match observations for patients in the treatment group with observations for patients in the control group:

```
ods graphics on;
proc psmatch data=drugs region=cs;
  class Drug Gender;
  psmodel Drug(Treated='Drug_X')= Gender Age Bmi;
  match method=optimal(k=1) exact=Gender stat=lps caliper=0.25;
  assess lps var=(Gender Age Bmi) / weight=none plots=(boxplot barchart);
  output out(obs=match)=Outgs lps=_Lps matchid=_MatchID;
run;
```

The CLASS statement specifies the classification variables. The PSMODEL statement specifies the logistic regression model that creates the propensity score for each observation, which is the probability that the patient receives Drug_X. The Drug variable is the binary treatment indicator variable and TREATED='Drug_X' identifies Drug_X as the treated group. The Gender, Age, and Bmi variables are included in the model because they are believed to be related to the assignment.

The REGION= option specifies an interval region of propensity scores (or equivalently, logits of propensity scores) such that only observations that have propensity scores in the region are used in stratification and matching. Because the MATCH statement is also specified, the REGION=CS option requests that only observations that have propensity scores in the common support region be used for matching. By default, the region is extended by 0.25 times a pooled estimate of the common standard deviation of the logit of the propensity score. For details, see the description of the [EXTEND=](#) option on page 7689.

The MATCH statement specifies the criteria for matching. The STAT=LPS option (which is the default) requests that the logit of the propensity score be used in computing differences between pairs of observations. The METHOD=OPTIMAL(K=1) option (which is the default) requests optimal matching of one control unit to each unit in the treated group in order to minimize the total within-pair difference. The EXACT=GENDER option forces the treated unit and its matched control unit to have the same value of the Gender variable.

The CALIPER=0.25 option specifies the caliper requirement for matching. This means that for a match to be made, the difference in the logits of the propensity scores for pairs of individuals from the two groups must

be less than or equal to 0.25 times the pooled estimate of the common standard deviation of the logits of the propensity scores.

The “Data Information” table in [Figure 95.3](#) displays information about the input and output data sets, the numbers of observations in the treated and control groups, the lower and upper limits for the propensity score support region, and the numbers of observations in the treated and control groups that fall within the support region. Of the 373 observations in the control group, only 351 fall within the support region.

Figure 95.3 Data Information
The PSMATCH Procedure

Data Information	
Data Set	WORK.DRUGS
Output Data Set	WORK.OUTGS
Treatment Variable	Drug
Treatment Group	Drug_X
All Obs (Treatment)	113
All Obs (Control)	373
Support Region	Extended Common Support
Lower PS Support	0.050244
Upper PS Support	0.683999
Support Region Obs (Treatment)	113
Support Region Obs (Control)	351

The “Propensity Score Information” table in [Figure 95.4](#) displays summary statistics by treatment group for all observations, for support region observations, and for matched observations.

Figure 95.4 Propensity Score Information

Propensity Score Information										
	Treated (Drug = Drug_X)					Control (Drug = Drug_A)				
Observations	N	Mean	Std Dev	Minimum	Maximum	N	Mean	Std Dev	Minimum	Maximum
All	113	0.310773	0.132467	0.060231	0.641148	373	0.208801	0.131969	0.020157	0.685757
Region	113	0.310773	0.132467	0.060231	0.641148	351	0.217557	0.126747	0.050951	0.682374
Matched	113	0.310773	0.132467	0.060231	0.641148	113	0.308246	0.130999	0.061866	0.682374

The “Matching Information” table in [Figure 95.5](#) displays the matching criteria, the number of matched sets, the numbers of matched observations in the treated and control groups, and the total absolute difference in the logit of the propensity score for all matches.

Figure 95.5 Matching Information

Matching Information	
Difference Statistic	Logit of Propensity Score
Method	Optimal Fixed Ratio Matching
Control/Treated Ratio	1
Caliper (Logit PS)	0.191862
Matched Sets	113
Matched Obs (Treated)	113
Matched Obs (Control)	113
Total Absolute Difference	2.941869

The ASSESS statement produces the tables and plots which summarize differences in the specified variables between treated and control groups for all observations, for the support region observations, and for the matched observations. As requested by the LPS and VAR= options, the variables listed in the table are the logit of propensity score and the variables Gender, Age, and Bmi. The WEIGHT=NONE option suppresses display of differences for the weighted matched observations. Note that for a matching of one control unit to each treated unit, the weights are all 1 for matched treated and control units, and the results are identical for the weighted matched observations and the matched observations.

The “Standardized Variable Differences” table, as shown in [Figure 95.6](#), displays standardized differences between the treated and control groups for all observations, the support region observations, and the matched observations. For a binary classification variable (Gender), the difference is in the proportion of the first ordered level (Female).

Figure 95.6 Standardized Differences**The PSMATCH Procedure**

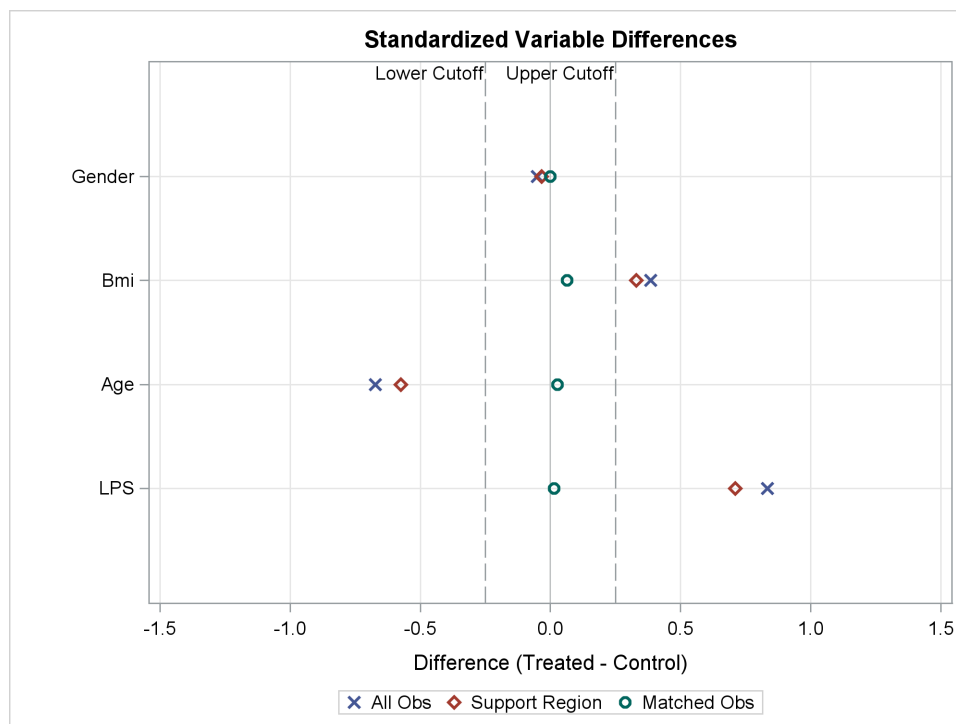
Standardized Variable Differences (Treated - Control)												
Standardized Mean Difference												
Variable	Mean Difference				Mean Difference				Percent Reduction		Variance Ratio	
	All Obs	Region Obs	Matched Obs	Divisor	All Obs	Region Obs	Matched Obs	Region Obs	Matched Obs	All Obs	Region Obs	Matched Obs
LPS	0.639971	0.545459	0.010556	0.767448	0.833894	0.710744	0.013755	14.77	98.35	0.6517	0.8314	1.0155
Age	-4.095091	-3.493684	0.168142	6.079104	-0.673634	-0.574704	0.027659	14.69	95.89	0.7076	0.8000	1.1262
Bmi	0.739296	0.632566	0.124248	1.923178	0.384414	0.328917	0.064605	14.44	83.19	0.8854	0.9288	1.1967
Gender	-0.024817	-0.016514	0	0.496925	-0.049941	-0.033233	0	33.46	100.00	0.9892	0.9922	1.0000

The divisor is computed from all observations, and it is used as the denominator to compute standardized differences for all observations, for support region observations, and for matched observations. The standardized mean differences are significantly reduced in the matched observations. The largest of these differences in absolute value is 0.0646, which is less than the upper limit of 0.25 recommended by Rubin (2001, p. 174) and Stuart (2010, p. 11). However, many authors use an upper limit of 0.10 (Normand et al. 2001; Mamdani et al. 2005; Austin 2009).

The variance ratios between the two groups are between 1 and 1.1967 for all variables in the matched observations, which is within the recommended range of 0.5 to 2. Because both EXACT=GENDER and METHOD=OPTIMAL are specified in the MATCH statement, the standardized difference for Gender is 0 in the matched observations.

When ODS Graphics is enabled, the PSMATCH procedure displays a standardized variable differences plot for the variables that are specified in the ASSESS statement, as shown in Figure 95.7.

Figure 95.7 Standardized Differences Plot



The “Standardized Variable Differences Plot” displays the standardized differences in the “Variable Differences” table in Figure 95.6. All differences for the matched observations are within the recommended limits of -0.25 and 0.25 , which are indicated by reference lines. Again, note that many authors use limits of -0.10 and 0.10 . (Normand et al. 2001; Mamdani et al. 2005; Austin 2009).

The PLOTS=BOXPLOT option requests a box plot for the logit of propensity score (LPS) and for each continuous variable specified in the ASSESS statement, as shown in Figure 95.8, Figure 95.9, and Figure 95.10. The box plots show good variable balance for the matched observations.

Figure 95.8 LPS Box Plot

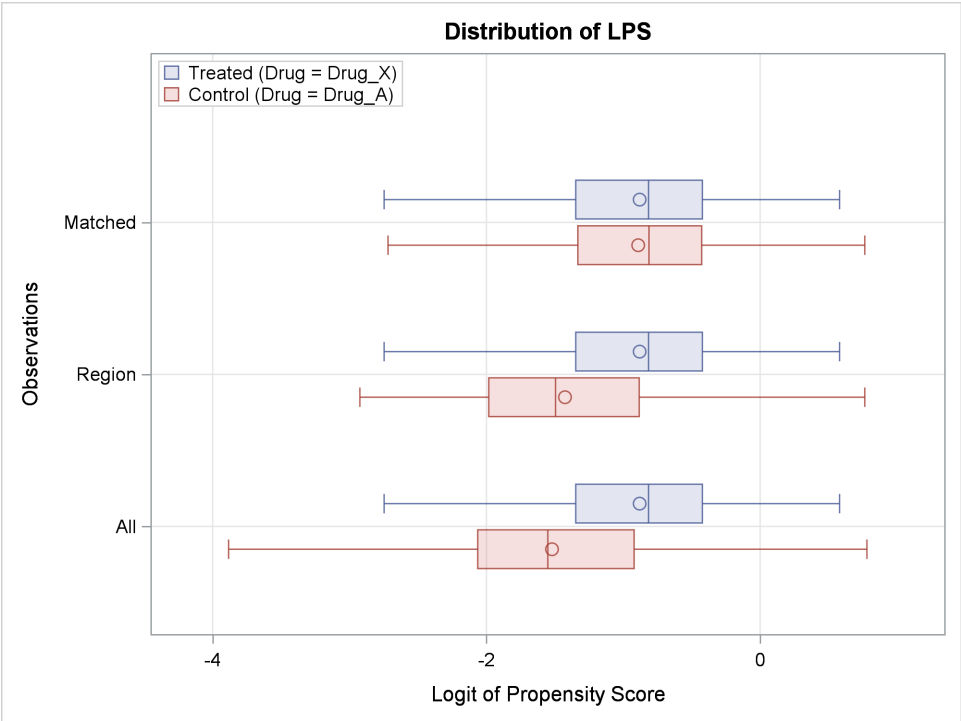


Figure 95.9 Age Box Plot

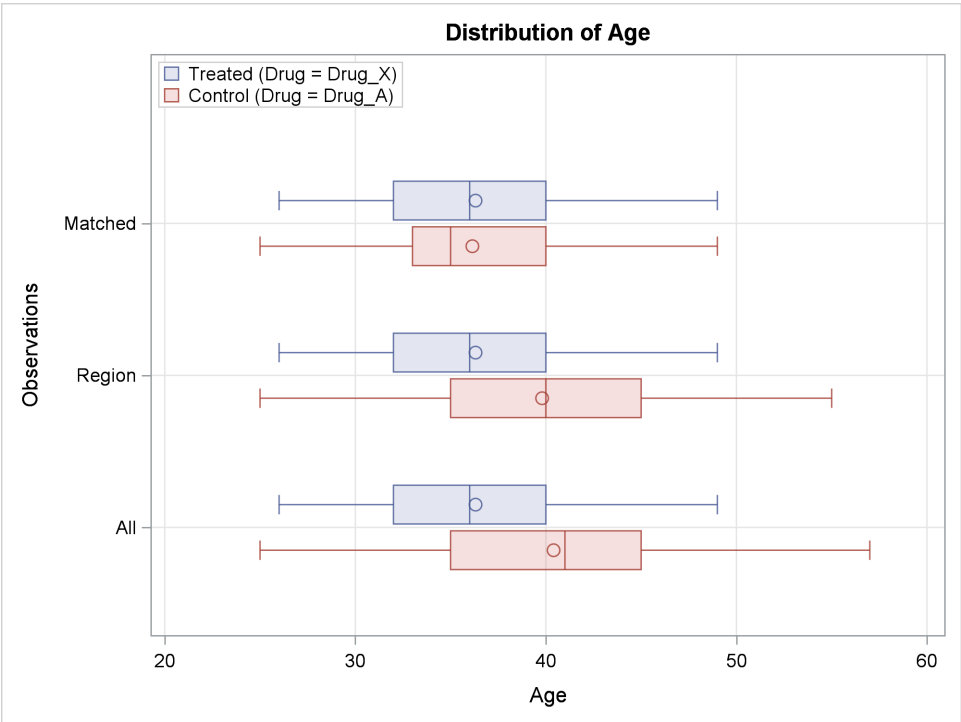
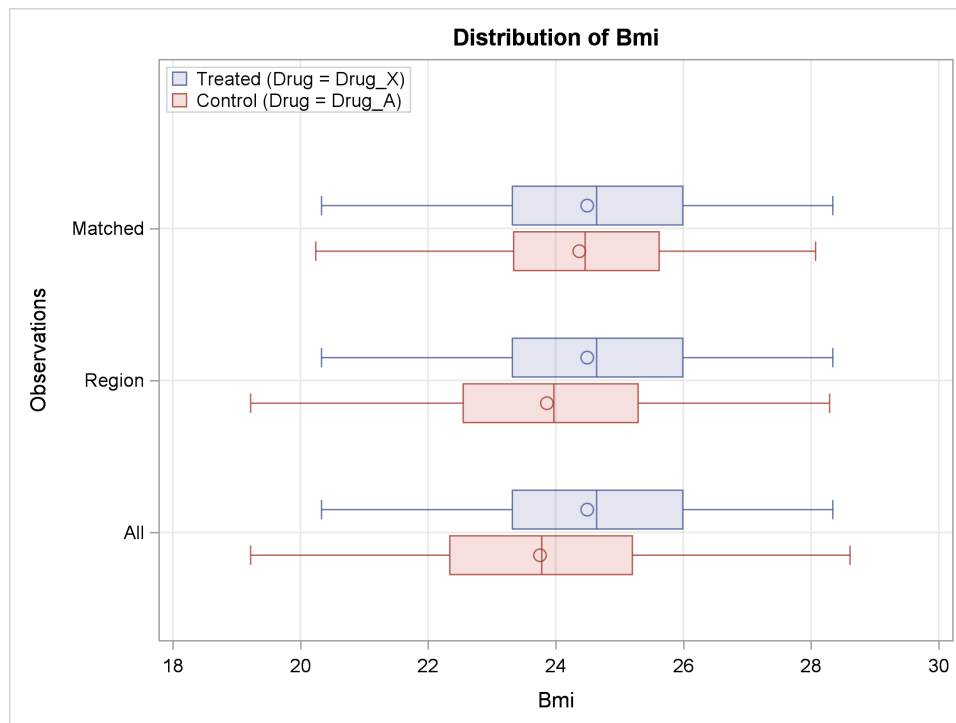
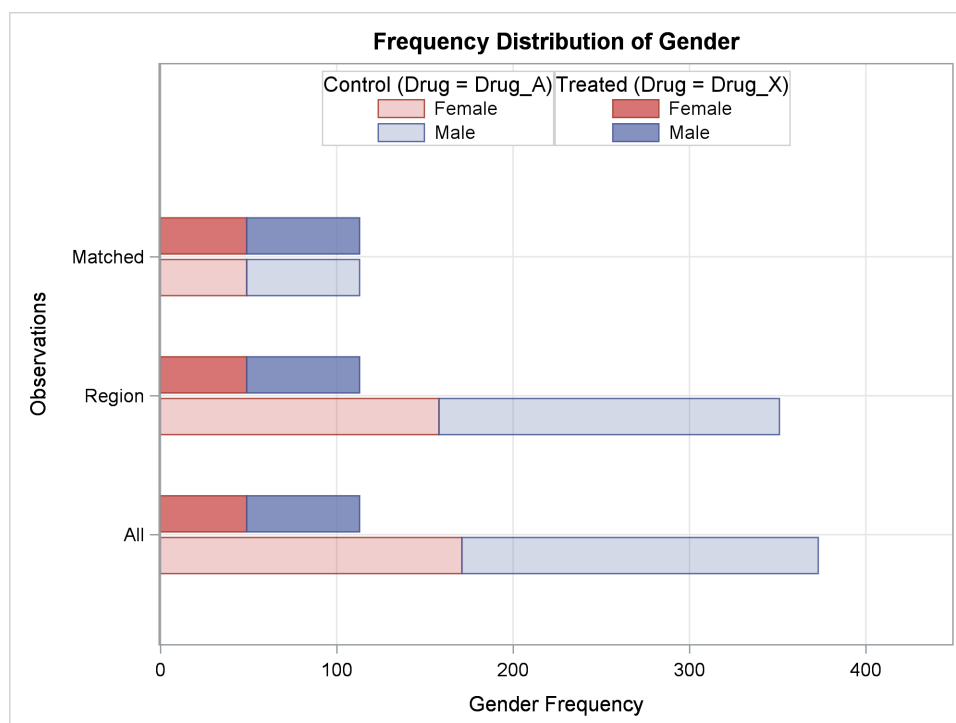


Figure 95.10 Bmi Box Plot

The PLOTS=BARCHART option displays a bar chart for each classification variable that is specified in the ASSESS statement, as shown in Figure 95.11. The bar chart displays identical distribution for matched observations for Gender because EXACT=GENDER is specified.

Figure 95.11 Gender Bar Chart

Because there is good balance in the distributions of the logit propensity score and the variables Gender, Age, and Bmi, you can output the matched observations for subsequent outcome analysis.

If you are not satisfied with the variable balance, you can do one or more of the following until you are satisfied: you can select another set of variables to fit the propensity score model, you can modify the specification of the propensity score model (for instance, by using nonlinear terms for the continuous variables or by adding interactions), you can modify the matching criteria, or you can choose another matching method.

The OUT(OBS=MATCH)=Outgs option in the OUTPUT statement creates an output data set Outgs that contains the matched observations. The following statements list the 10 observations in Outgs that have lowest propensity scores, as shown in [Figure 95.12](#).

```
proc sort data=outgs out=outgs1;
    by _ps_;
run;

proc print data=outgs1(obs=10);
    var PatientID Drug Gender Age Bmi _PS_ _LPS _MatchWgt_ _MatchID;
run;
```

Figure 95.12 Output Data Set with Matching Numbers

Obs	PatientID	Drug	Gender	Age	Bmi	_PS_	_Lps	_MATCHWGT_	_MatchID
1	89	Drug_X	Female	44	20.75	0.06023	-2.74744	1	1
2	213	Drug_A	Female	49	23.24	0.06187	-2.71892	1	1
3	141	Drug_A	Female	43	20.55	0.06401	-2.68256	1	2
4	323	Drug_X	Female	46	22.22	0.06763	-2.62375	1	2
5	217	Drug_X	Male	49	23.96	0.08772	-2.34184	1	3
6	290	Drug_A	Female	40	20.57	0.08778	-2.34104	1	4
7	420	Drug_A	Male	45	22.08	0.08801	-2.33813	1	3
8	234	Drug_X	Female	41	21.11	0.08904	-2.32538	1	4
9	320	Drug_X	Female	46	24.17	0.10323	-2.16183	1	5
10	473	Drug_A	Female	45	23.76	0.10464	-2.14669	1	5

By default, the output data set includes the variable `_PS_` (which provides the propensity score) and the variable `_MATCHWGT_` (which provides matched observation weights). The weight for each treated unit is 1. Because $K=1$ is specified in the `METHOD=OPTIMAL` option in the `MATCH` statement, one control unit is matched to each treated unit; so the weight for each matched control unit is also 1. The `LPS=_LPS` option creates a variable named `_LPS` which provides the logit of propensity score, and the `MATCHID=_MatchID` option creates a variable named `_MatchID` that identifies the matched sets of observations.

If you assume that no other confounding variables are associated with both the response variable and the treatment group indicator Drug, then after the responses for the trial are observed and added to the matched data set Outgs, you can use the same outcome analysis on this matched data set as you would have used on the original data set Drugs (augmented with responses) to estimate the treatment effect (Ho et al. 2007, p. 233).

Syntax: PSMATCH Procedure

The following statements are available in the PSMATCH procedure:

```

PROC PSMATCH < options > ;
    ASSESS < LPS > < PS > < VAR=(var-list) > < / assess-options > ;
    BY variables ;
    CLASS variables ;
    FREQ variable ;
    MATCH < options > ;
    OUTPUT OUT < (OBS=obs-value ) > = SAS-data-set < keyword=name < keyword=name ... > > ;
    PSDATA TREATVAR=treatvar < (TREATED='level' ) > < option > ;
    PSMODEL treatvar < (TREATED='level' ) > = < effects > < / WEIGHT=weight > ;
    STRATA < options > ;

```

The PROC PSMATCH statement invokes the PSMATCH procedure. The CLASS statement and either a PSMODEL or PSDATA statement are required. If a PSMODEL statement is specified, the CLASS statement must precede the PSMODEL statement. The STRATA statement is not used if a MATCH statement is also specified.

The following sections describe PROC PSMATCH statement and then describe the other statements in alphabetical order.

PROC PSMATCH Statement

The PROC PSMATCH statement invokes the PSMATCH procedure. [Table 95.1](#) summarizes the options available in the PROC PSMATCH statement.

Table 95.1 Summary of PROC PSMATCH Options

Option	Description
DATA=	Specifies the input data set
REGION=	Specifies the support region of observations for stratification and matching

DATA=SAS-data-set

names the input SAS data set. If the propensity scores are to be derived from this data set, you must also include a PSMODEL statement to specify the binary logistic model. Otherwise, a PSDATA statement is required to identify the variable that contains either the propensity scores or the logits of propensity scores.

REGION=region < (region-options) >

specifies an interval region of propensity scores (or equivalently, logits of propensity scores) such that only observations that have propensity scores in the region are used in stratification and matching. This option also specifies the observations to be included in the output data set if the OUT(OBS=REGION) option is specified in the OUTPUT statement (even when the STRATA and MATCH statements

are not specified). By default, `REGION=TREATED` if the `MATCH` statement is specified, and `REGION=ALLOBS` if the `MATCH` statement is not specified.

You can specify the following *regions* along with their *region-options*:

REGION=ALLOBS *<(region-options)>*

selects all available observations. You can specify the following *region-options* to select observations whose propensity scores lie in a specified range:

PSMIN=*pmin*

specifies the minimum propensity score in the support region, where $pmin \geq 0$. Observations that have propensity scores that are less than *pmin* are excluded from the support region. By default, `PSMIN=0`, so that observations with small propensity scores are not excluded.

PSMAX=*pmax*

specifies the maximum propensity score in the support region, where $pmax \leq 1$. Observations that have propensity scores that are greater than *pmax* are excluded from the support region. By default, `PSMAX=1`, so that observations with large propensity scores are not excluded.

You can also use the `PSMIN=` and `PSMAX=` options to exclude observations with extreme propensity scores from the output data set.

REGION=CS *<(region-option)>*

REGION=TREATED *<(region-option)>*

selects the region of common support for the propensity scores for observations in the treated and control groups or the region of propensity scores for observations in the treated group only:

CS selects observations whose propensity scores lie in the region of common support for the propensity scores for observations in the treated and control groups. This region is the widest interval such that both the treated and the control groups have subjects whose propensity scores lie within this interval. The lower endpoint of the region is the larger of the minimum propensity score for the treated group and the minimum propensity score for the control group. The upper endpoint is the smaller of the maximum propensity score for the treated group and the maximum propensity score for the control group.

TREATED selects observations whose propensity scores lie in the region of propensity scores for observations in the treated group.

You can use the following *region-option* to extend the specified support region:

EXTEND *<(ext-options)>* = *p* **<(LOWER=*p_l* UPPER=*p_u*)>**

specifies extension to the lower and upper ends of the common support region (`REGION=CS`) or the range of treated observations (`REGION=TREATED`) for the support region, $p \geq 0$. By default, `EXTEND=0.25`.

You can use the following *ext-options* to prescribe the extension requirement:

MULT=ONE | STDDEV

specifies the multiplier for the extension p to extend the support region:

ONE extends the region by p .

STDDEV extends the region by $p \times$ the pooled estimate of the common standard deviation of the specified statistic, where this estimate is computed as the square root of the average of the variance of the PS (LPS) in the treated group and the variance of the PS (LPS) in the control group.

By default, MULT=STDDEV.

STAT=LPS | PS

specifies the type of the statistic that is used to extend the support region:

LPS extends the region by using the logit of the propensity score.

PS extends the region by using the propensity score.

By default, STAT=LPS.

The MULT= and STAT= *ext-options* prescribe the extension requirement as follows:

- EXTEND(STAT=PS MULT=ONE)= p extends the specified support region by p in propensity score. That is, if (R_l, R_u) denotes the propensity score interval region that is computed from the specified *region*, then the range of the extended support region is given by $(R_l - p, R_u + p)$.
- EXTEND(STAT=PS MULT=STDDEV)= p extends the specified support region by $p \times \hat{\sigma}_{ps}$, the square root of the average variance of the propensity score in the treated and control groups. That is, if (R_l, R_u) denotes the propensity score interval region that is computed from the specified *region*, then the range of the extended support region is given by $(R_l - p \hat{\sigma}_{ps}, R_u + p \hat{\sigma}_{ps})$.
- EXTEND(STAT=LPS MULT=ONE)= p extends the specified support region by p in the logit of propensity score.
- EXTEND(STAT=LPS MULT=STDDEV)= p extends the specified support region by $p \times \hat{\sigma}_{lps}$, the square root of the average variance of the logit of propensity score in the treated and control groups.

You can specify one of the following two options to use an extension other than p :

LOWER= p_l extends the lower end of the specified region by p_l , where $p_l \geq 0$.

UPPER= p_u extends the upper end of the specified region by p_u , where $p_u \geq 0$.

ASSESS Statement

ASSESS <LPS> <PS> <VAR=(var-list)> </ assess-options> ;

The ASSESS statement assesses variable differences between the treated and control groups for all observations and for observations in the specified support region. It also assesses variable differences for matched observations if a MATCH statement is specified and assesses variable differences for observations by stratum if a STRATA statement is specified. In addition, the ASSESS statement assesses variable differences for weighted observations provided that the WEIGHT=NONE suboption is not specified.

You can specify variables for assessment by using the following options:

LPS

requests an assessment of differences in the logit of the propensity score.

PS

requests an assessment of differences in the propensity score.

VAR=(var-list)

requests an assessment of differences in the specified list of variables. These variables must be binary classification variables or continuous variables in the input data set.

If none of these options are specified, an assessment of differences in the propensity score is produced by default.

In addition, you can specify various *assess-options* after a slash (/). [Table 95.2](#) summarizes these options:

Table 95.2 ASSESS Statement Options

Option	Description
PLOTS=	Requests variable plots
STDDIFFDIV=	Specifies the divisor for the standardized difference
VARINFO	Displays variable information for the treated and control groups
WEIGHT=	Specifies the weight for the variable distribution

PLOTS <(global-option)> <= plot-request>

PLOTS <(global-option)> = (plot-request <... plot-request>)

specifies options that control the plots.

You can specify the following *global-options*:

ONLY

suppresses the default plots and displays only plots that are specifically requested.

ORIENT=HORIZONTAL | VERTICAL

controls the orientation of the plots:

HORIZONTAL places the lines and boxes horizontally for variable distribution plots, places the bar lengths horizontally for bar charts, places the variable values horizon-

tally for cloud plots, and places the standardized differences on the horizontal axis for the standardized differences plot.

VERTICAL places the lines and boxes vertically for variable distribution plots, places the bar lengths vertically for bar charts, places the variable values vertically for cloud plots, and places the standardized differences on the vertically axis for the standardized differences plot.

By default, ORIENT=HORIZONTAL.

You can specify the following *plot-requests*:

ALL

requests all applicable plots for all variables that are specified in the ASSESS statement. These plots include bar charts for binary classification variables, box plots for continuous variables, cloud plots for all variables, and a combined standardized differences plot for all variables. If you specify a STRATA statement, then PROC PSMATCH also produces the plots by stratum.

BAR <(DISPLAY=ALL | (*bar-list*))>

BARCHART <(DISPLAY=ALL | (*bar-list*))>

requests comparative bar charts for binary classification variables that are specified in the VAR= option. You can specify either of the following options:

DISPLAY=ALL

requests bar charts for all binary classification variables that are specified in the VAR= option.

DISPLAY=(*bar-list*)

specifies a subset of the binary classification variables for which bar charts are to be displayed.

By default, DISPLAY=ALL.

If you specify a STRATA statement, then the bar charts by stratum are also displayed.

BOX <(DISPLAY=ALL | (*box-list*))>

BOXPLOT <(DISPLAY=ALL | (*box-list*))>

requests box plots for LPS, PS, and all continuous variables that are specified in the VAR= option. You can specify either of the following options:

DISPLAY=ALL

requests box plots for LPS, PS, and all continuous variables that are specified in the VAR= option.

DISPLAY=(*box-list*)

specifies a subset of the continuous variables for which box plots are to be displayed.

By default, DISPLAY=ALL.

If you specify a STRATA statement, then the box plots by stratum are also displayed.

CLOUD <(DISPLAY=ALL | (*cloud-list*))>

CLOUDPLOT <(DISPLAY=ALL | (*cloud-list*))>

requests cloud plots for LPS, PS, and all variables that are specified in the VAR= option. The term cloud plot is used here to refer to scatter plots in which the points have been jittered by adding random noise to prevent overplotting, which typically occurs when a continuous variable (such as age) is rounded to some convenient unit (such as years).

You can specify either of the following options:

DISPLAY=ALL

requests cloud plots for LPS, PS, and all variables that are specified in the VAR= option.

DISPLAY=(*cloud-list*)

specifies a subset of the continuous variables for which box plots are to be displayed.

By default, DISPLAY=ALL.

If you specify a STRATA statement, then the cloud plots by stratum are also displayed.

NONE

suppresses all plots.

STDDIFF

STDDIFFPLOT

requests a standardized differences plot for PS, LPS, and all variables that are specified in the VAR= option. If you specify a STRATA statement, then a standardized differences by stratum plot is also displayed.

By default, PLOTS=STDDIFF.

STDDIFFDIV=POOLED | TREATED

specifies the divisor for the standardized difference:

POOLED uses the pooled standard deviation, which is computed as the square root of the average of the sample variance for the treated group and the sample variance for the control group.

TREATED uses the standard deviation of the variable values in the treated group only.

By default, STDDIFFDIV=POOLED.

VARINFO

requests a variable information table for the treated and control groups.

WEIGHT=ATEWGT | ATTWGT | MATCHWGT | NONE

requests (except when the WEIGHT=NONE is specified) additional variable assessment for weighted matched observations if a MATCH statement is specified, for weighted observations in each stratum if a STRATA statement is specified, and for weighted observations in the support region if neither the MATCH nor the STRATA statement is specified:

ATEWGT uses inverse probability of treatment weighting (IPTW) to weight the treatment and control groups up to the combined group. These weights are appropriate for estimation of the ATE. This option applies only if the MATCH statement is not specified.

ATTWGT	uses ATT weighting (also referred to as weighting by odds) to weight the control group up to the treatment group. These weights are appropriate for estimation of the ATT. This option applies only if the MATCH statement is not specified. For more information about ATT weighting, see the section “ ATT Weighting ” on page 7708.
MATCHWGT	uses match weighting to weight the control group up to the treatment group. That is, in each matched set, each treated unit has a weight of 1 and each control unit has a weight that equals the number of treated units divided by the number of control units in the matched set. For example, with one-to-one pair matching, each treated unit has a weight of 1 and each control unit has a weight of 1. With this weighting, the total weight of control units is the same as the total number of treated units in each matched set, and the total weight of matched control units is the same as the total number of treated units. This weighting is useful when multiple control units are matched to each treated unit, and is appropriate for estimating the ATT. This option applies only if a MATCH statement is specified.
NONE	does not add weighted variable assessment.

By default, WEIGHT=MATCHWGT if a MATCH statement is specified, WEIGHT=NONE if a STRATA statement is specified, and WEIGHT=ATTWGT if neither the MATCH nor the STRATA statements is specified.

For more information about these propensity score weights, See the section “[Propensity Score Weighting](#)” on page 7707.

BY Statement

BY variables ;

You can specify a BY statement with PROC PSMATCH to obtain separate analyses of observations in groups that are defined by the BY variables. When a BY statement appears, the procedure expects the input data set to be sorted in order of the BY variables. If you specify more than one BY statement, only the last one specified is used.

If your input data set is not sorted in ascending order, use one of the following alternatives:

- Sort the data by using the SORT procedure with a similar BY statement.
- Specify the NOTSORTED or DESCENDING option in the BY statement for the PSMATCH procedure. The NOTSORTED option does not mean that the data are unsorted but rather that the data are arranged in groups (according to values of the BY variables) and that these groups are not necessarily in alphabetical or increasing numeric order.
- Create an index on the BY variables by using the DATASETS procedure (in Base SAS software).

CLASS Statement

CLASS *variables* ;

The required CLASS statement specifies the following input variables that are used as classification variables:

- the variable to use as the treatment indicator in the PSDATA and PSMODEL statements
- the classification covariates in the logistic model in the PSMODEL statement
- the classification variables that are specified in the VAR= option in the ASSESS statement

If a PSMODEL statement is specified, the CLASS statement must precede the PSMODEL statement. Classification variables can be either character or numeric.

FREQ Statement

FREQ *variable* ;

The FREQ statement identifies a *variable* that contains the frequency of occurrence of each observation. PROC PSMATCH treats each observation as if it appears n times, where n is the value of the FREQ variable for the observation. The FREQ statement is not allowed if a MATCH statement is specified.

MATCH Statement

MATCH *< options >* ;

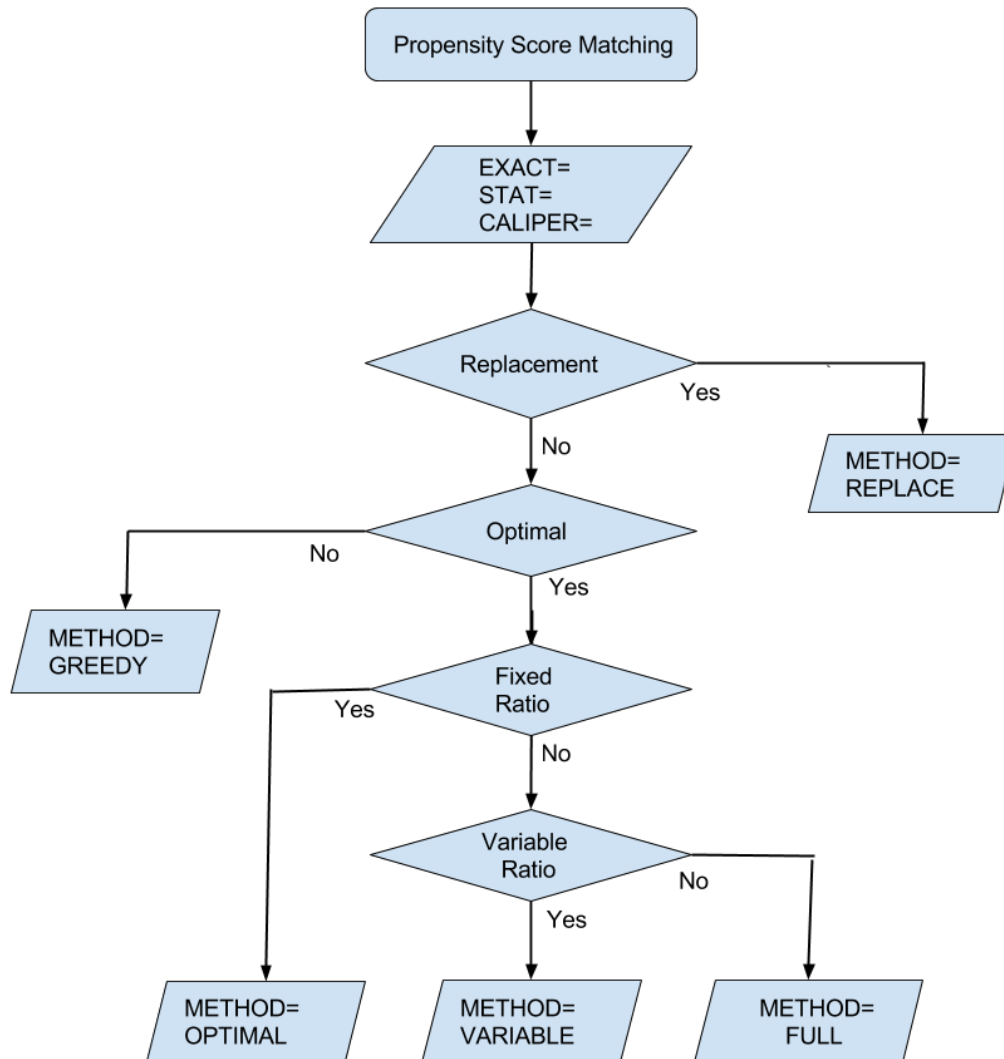
The MATCH statement matches observations in the control group to observations in the treatment group. The MATCH statement is not allowed if a FREQ statement is specified, and the STRATA statement does not apply if a MATCH statement is specified.

Table 95.3 summarizes the *options* in the MATCH statement.

Table 95.3 MATCH Statement Options

Option	Description
CALIPER=	Specifies the caliper width requirement for matching
EXACT=	Requests exact matching for specified classification variables
METHOD=	Specifies the method for matching
STAT=	Specifies the statistic for comparing treated units and control units

The flowchart in Figure 95.13 displays the steps in the propensity score matching process.

Figure 95.13 Propensity Score Matching Options

You can specify the following *options* in the MATCH statement:

CALIPER <(caliper-options)> = *r*

specifies the caliper width requirement for matching, where *r* is either missing or greater than 0. The difference between the treated unit and its matching control unit must be less than or equal to *r*. If you specify CALIPER=., then the caliper requirement is ignored. By default, CALIPER=0.25 (Rosenbaum and Rubin 1985, p. 37). Austin (2011a) has shown that CALIPER=0.20 is optimal in many settings.

You can use the following two *caliper-options* to prescribe the caliper requirement:

MULT=ONE | STDDEV

specifies the multiplier for the specified caliper width *r*:

- ONE** uses r for the caliper width.
- STDDEV** uses r times the pooled estimate of the common standard deviation of the STAT= statistic, where this estimate is computed as the square root of the average of the variances in the treated and control groups.

By default, MULT=STDDEV.

STAT=LPS | PS

specifies the statistic in the caliper width specification that is used for determining the distance between two units. This suboption applies only if you specify the STAT=MAH option in the MATCH statement.

LPS uses the logit of propensity score.

PS uses the propensity score scale.

By default, STAT=LPS.

If you do not specify the STAT=MAH option in the MATCH statement, the STAT= option in the MATCH statement is used to select the statistic used in the caliper width specification.

EXACT=variable | (variables)

specifies classification variables that are to be matched exactly. That is, observations in each matched set must have the same values for these variables. The *variables* must be specified in the CLASS statement.

METHOD=method <(method-options)>

specifies the method for the matching. You can specify the following *method* and *method-options*. By default, METHOD=OPTIMAL.

METHOD=FULL (KMAX=kmax <full-options>)

requests optimal full matching. Each treated unit is matched with one or more control units, and each control unit (if matched) is matched with one or more treated units. If the specified total number of control units to be matched is less than the number of available control units, then constrained full matching is performed—that is, not all observations are matched.

You must specify the following suboption:

KMAX=kmax

specifies the maximum number of control units to be matched with each treated unit, where $kmax \geq 1$.

You can also specify the following *full-options*:

KMAXTREATED=kmaxtrt

KMAXTRT=kmaxtrt

specifies the maximum number of treated units for each control, where $kmaxtrt \geq 1$. By default, KMAXTREATED=2.

KMEAN=*kmean*

specifies the average number of control units for each treated unit in the matched data set. If the resulting number of control units is greater than the number of control units in the support region, the number in the support region is used.

NCONTROL=*m*

specifies the total number of control units in the matched data set. If the number *m* is greater than the number of control units in the support region, the number in the support region is used.

PCTCONTROL=*p*

specifies the percentage of the total number of control units in the matched data set. If the resulting number of control units is greater than the total number of control units in the support region, the number in the support region is used.

If you do not specify any of the KMEAN=, NCONTROL=, and PCTCONTROL= options, KMEAN= ($kmax + 1 / kmaxtrt$) / 2 is used. If the resulting number of units computed from the KMEAN=, NCONTROL, or PCTCONTROL= option is less than the total number of control units, then not all control units are matched.

METHOD=GREEDY <(K=*k* ORDER=*order-option*)>

requests greedy nearest neighbor matching in which each treated unit is sequentially matched with the *k* nearest control units. Matching depends on the ordering of the treated units, which you can specify in the ORDER= suboption.

You can specify the following suboptions:

K=*k*

specifies the number of matching control units, where $k > 0$, for each treated unit. PROC PSMATCH performs *k* separate loops of matching for treated units. In each loop, the nearest control unit is sequentially matched to each treated unit. By default, K=1 (one control unit for each treated unit).

ORDER=ASCENDING | DESCENDING | RANDOM <(SEED=*number*)>

specifies the ordering of treated units that are used to find the matching control units. You can specify one of the following values:

ASCENDING

orders the treated units in ascending order of the propensity score.

DESCENDING

orders the treated units in descending order of the propensity score.

RANDOM <(SEED=*number*)>

orders the treated units in random order of the propensity score. The SEED= suboption specifies a positive integer to start the pseudorandom number generator. If the SEED= option is not specified, the value is generated from reading the time of day from the computer's clock.

By default, ORDER=DESCENDING.

METHOD=OPTIMAL <(K=k)>

requests optimal fixed ratio matching. The $K=k$ suboption specifies the number of matching control units, where $k > 0$, for each treated unit. By default, $K=1$ (one control unit is matched with each treated unit).

METHOD=REPLACE <(K=k)>

requests a fixed number k of unique matching control units for each treated unit, where the matched control units are selected with replacement. This means that each control unit can be matched to more than one treated unit, but it can only be matched once to the same treated unit. The $K=k$ suboption specifies the number of matching control units, where $k > 0$, for each treated unit. By default, $K=1$ (one control unit is matched with each treated unit).

METHOD=VARRATIO (KMAX=kmax <varratio-options>)

requests optimal variable ratio matching. Each treated unit is matched with one or more control units.

You must specify the following suboption:

KMAX=kmax

specifies the maximum number of control units matched with each treated unit, where $kmax \geq 1$.

You can also specify the following *varratio-options*:

KMEAN=kmean

specifies the average number of control units for each treated unit in the matched data set. If the resulting number of control units is greater than the total number of control units in the support region, the number in the support region is used.

KMIN=kmin

specifies the minimum number of control units to be matched with each treated unit. By default, $KMIN=1$.

NCONTROL=m

specifies the total number of control units in the matched data set. If the number m is greater than the total number of control units in the support region, the number in the support region is used.

PCTCONTROL=p

specifies the percentage of total control units in the matched data set. If the resulting number of control units is greater than the total number of control units in the support region, the number in the support region is used.

If you do not specify any of the $KMEAN=$, $NCONTROL=$, and $PCTCONTROL=$ options, then $KMEAN = (kmin + kmax) / 2$ is used.

STAT=statistic

specifies the statistic to be compared when treated units are matched to control units. You can specify the following *statistics*:

LPS

specifies matching that minimizes the difference between the logits of the propensity scores for the two units.

PS

specifies matching that minimizes the difference between the propensity scores for the two units.

MAH (*var-options* < / *mah-options* >)

specifies matching that minimizes the Mahalanobis distance between the two units.

You use the following *var-options* to select at least one variable for computing the Mahalanobis distance:

LPS

includes the logit of the propensity score.

PS

includes the propensity score.

VAR=(*var-list*)

includes variables in the specified *var-list*. These variables must be continuous variables in the input data set.

You can also specify the following *mah-options*:

COV=CONTROL | IDENTITY | POOLED

specifies the type of covariance matrix in the Mahalanobis distance:

CONTROL uses the covariance matrix that is computed from observations in the control group.

IDENTITY uses the identity matrix, and the resulting distance is the Euclidean distance.

POOLED uses the pooled covariance matrix that is computed from observations in the treated group and observations in the control group.

By default, COV=CONTROL.

SQRT=YES | NO

specifies whether to apply the square root to the Mahalanobis distance in the difference computation. This *mah-option* does not affect matching results for greedy nearest neighbor matching and matching with replacement. It affects only results for optimal matching that minimizes the total absolute difference.

YES uses the square root of the Mahalanobis distance as the difference between treated and control units.

NO uses the Mahalanobis distance as the difference between treated and control units.

By default, SQRT=YES.

OUTPUT Statement

OUTPUT OUT < (**OBS**=*obs-value*) >=*SAS-data-set* < *keyword=name* < *keyword=name* ... > > ;

The OUTPUT statement specifies the output data set and variables. You must specify the following option:

OUT < (**OBS**=*obs-value*) >=*SAS-data-set*

names the output data set. The data set also includes the results of matching if you provide the MATCH statement. You can specify one of the following values for *obs-value*:

- ALL** the output data set contains all observations.
- REGION** the output data set contains only observations in the specified support region.
- MATCH** the output data set contains only the matched treated units and control units. This option applies only if you specify the MATCH statement.

By default, OBS=ALL.

You can also specify the one or more of the following *keywords* to create and name the output variables:

ATEWGT=*name*

creates and names the weight variable that provides inverse probability of treatment weighting. This weighting is appropriate for estimating the ATE.

ATTWGT=*name*

creates and names the weight variable for ATT weighting. This weighting is appropriate for estimating the ATT. If ATTWGT= is not specified and neither the MATCH nor the STRATA statement is specified, then this variable is automatically created with *name*=_ATTWGT_.

LPS=*name*

creates and names the variable that provides the logit of propensity score.

MATCHID=*name* | (*names*)

creates and names the variable that provides identification numbers for the matched treated and control units. This suboption applies only if you also specify the MATCH statement. If METHOD=REPLACE(K=*k*) is specified with *k* > 1, then you can use MATCHID=(*names*) to name the *k* matching groups for each treated unit.

MATCHWGT=*name*

creates and names the weight variable for the matching. This suboption applies only if you also specify the MATCH statement. In each matched set, each treated unit has a weight of 1 and each control unit has a weight that equals the number of treated units divided by the number of control units in the matched set. With this weighting, the total weight of control units is the same as the total number of treated units in each matched set, and the total weight of matched control units is the same as the total number of treated units. This weighting is appropriate for estimating the ATT.

If MATCHWGT= is not specified but the MATCH statement is specified, then this variable is automatically created with *name*=_MATCHWGT_.

PS=*name*

creates and names the variable that provides the propensity score.

If PS= is not specified and the PS= option in the PSDATA statement is also not specified, then this variable is automatically created with *name*=_PS_.

STRATA=*name*

creates and names the variable that numbers the strata. The suboption applies only if the STRATA statement is specified.

If STRATA= is not specified but the STRATA statement is specified, then this variable is automatically created with *name*=_STRATA_.

PSDATA Statement

PSDATA *TREATVAR*=*treatvar* <(TREATED='level')> <*option*> ;

The PSDATA statement specifies the treatment indicator variable and a variable for either the propensity score or the logit of propensity score for variables that are in the DATA= data set. Either the PSMODEL statement or the PSDATA statement is required.

You must specify the following option:

TREATVAR=*treatvar* <(TREATED='level')>

names the treatment indicator variable, *treatvar*, which must be a binary classification variable that is specified in the CLASS statement. The TREATED='level' suboption indicates the level that corresponds to treatment. If the TREATED='level' suboption is not specified, the first ordered level based on the formatted values is used to derive the propensity scores.

You must also specify one (and only one) of the following *options*:

PS=*name*

names the variable that contains propensity scores, where the variable *name* must be a variable in the DATA= data set.

LPS=*name*

names the variable that contains logits of propensity scores, where the variable *name* must be a variable in the DATA= data set.

PSMODEL Statement

PSMODEL *treatvar* <(TREATED='level')> = <*effects*> </ WEIGHT= *weight*> ;

The PSMODEL statement specifies the logistic regression model for computing the propensity score. Either the PSMODEL statement or the PSDATA statement is required to obtain the propensity scores.

The treatment indicator variable *treatvar* must be a binary classification variable that is listed in the CLASS statement. You can specify the following options:

TREATED=*'level'*

models the probability of the specified treated *level*. If this option is not specified, PROC PSMATCH models the probability of the first ordered level based on the formatted values.

effects

are the explanatory effects, which can include variables, main effects, interactions, and nested effects for the logistic regression model.

WEIGHT=*weight*

specifies a variable that contains the weight of each observation that is used in fitting the logistic regression model to derive the propensity scores. These weights should not be confused with weights derived from the propensity scores by the PSMATCH procedure.

STRATA Statement

STRATA < *options* > ;

The STRATA statement divides observations in the support region into strata based on propensity scores, where the support region is specified in the REGION= option in the PROC PSMATCH statement.

The STRATA statement does not apply when you specify the MATCH statement. You can specify the following *options*:

NSTRATA=*n*

specifies the number of strata, where $n \geq 2$. Only observations in the support region are stratified. By default, NSTRATA=5.

KEY=NONE | TREATED

specifies the type of observations that are used to construct the strata:

- | | |
|----------------|--|
| NONE | requests that each stratum contain approximately the same number of observations, which can be in either the treatment group or the control group. |
| TREATED | requests that each stratum contain approximately the same number of observations in the treatment group. |

By default, KEY=TREATED.

For more information, see the section “[Propensity Score Stratification](#)” on page 7709.

Details: PSMATCH Procedure

Observational Studies Contrasted with Randomized Trials

In a randomized study, such as a randomized controlled trial, the subjects are randomly assigned to a treated (exposure) group or a control (nonexposure) group. Random assignment ensures that the distribution of the covariates is the same in both groups, and the treatment effect can be estimated from a direct comparison of the outcomes for the subjects in the two groups.

In contrast, the subjects in an observational study are not randomly assigned to the treated and control groups. Confounding can occur if some covariates are related to both the treatment assignment and the outcome. Consequently, there can be systematic differences between the treated subjects and the control subjects. In the presence of confounding, statistical approaches are required that remove the effects of confounding when estimating the effect of treatment.

Observational studies are carried out when it is impractical or unethical to perform a randomized experiment. One example of an observational study is a retrospective cohort study that examines the relationship between a specific disease and a risk factor that occurred in the past; another example is a nonrandomized clinical trial that uses existing data such as control units that are extracted from a registry database.

The approach that the PSMATCH procedure uses and the following terminology are based on the framework for causal inference that was introduced by Rubin (1974) and Rosenbaum and Rubin (1983).

Under the potential outcomes framework, in an observational study whose goal is to estimate the effect of a treatment, each individual typically has two potential outcomes:

- $Y(1)$, the outcome that would be observed if the individual receives the treatment.
- $Y(0)$, the outcome that would be observed if the individual does not receive the treatment under identical circumstances to those under which the subject would have received the treatment.

However, only one outcome can be observed.

The treatment effect is defined as $Y(1) - Y(0)$, and the average treatment effect is defined as:

$$ATE = E(Y(1) - Y(0))$$

The average treatment effect for the treated (individuals who actually receive treatment) is defined as:

$$ATT = E(Y(1) - Y(0) \mid T = 1)$$

where T denotes the treatment assignment.

In a randomized trial, the potential outcomes $(Y(0), Y(1))$ and the treatment assignment T are independent:

$$(Y(0), Y(1)) \perp\!\!\!\perp T$$

Thus, the average treatment effect (ATE) is identical to the average treatment effect for the treated (ATT), which can be expressed as follows and can be estimated from the observed data:

$$E(Y(1) \mid T = 1) - E(Y(0) \mid T = 0)$$

In an observational study, the potential outcomes $(Y(0), Y(1))$ and the treatment assignment T might not be independent. In this case, the ATE and ATT are not the same. Furthermore, outcomes cannot be compared directly to estimate the treatment effect. In particular,

$$\begin{aligned} \text{ATT} &= E(Y(1) - Y(0) \mid T = 1) \\ &= E(Y(1) \mid T = 1) - E(Y(0) \mid T = 0) + E(Y(0) \mid T = 0) - E(Y(0) \mid T = 1) \end{aligned}$$

The following term can be estimated from the observed data:

$$E(Y(1) \mid T = 1) - E(Y(0) \mid T = 0)$$

However, the selection bias cannot be estimated from the observed data:

$$E(Y(0) \mid T = 0) - E(Y(0) \mid T = 1)$$

The selection bias is the average difference in the response that would be observed between individuals in the control group who do not receive treatment and individuals in the treatment group who do not receive treatment. Thus, the usual observed difference between the treated and control groups cannot be used to estimate the treatment effect. For subjects who are not randomly assigned to the treated and control groups, the baseline variables could be related to both the treatment assignment and the outcome, and consequently standard statistical methods of outcome analysis could result in biased estimates.

One strategy for correctly estimating the treatment effect is based on the propensity score, which is the conditional probability of the treatment assignment given the observed variables. You use propensity scores to account for confounding by weighting observations, by creating strata of subjects that have similar propensity scores, or by matching control subjects to treated subjects. This is done prior to the outcome analysis and without knowledge of the outcome variable (Rosenbaum and Rubin 1984; Stuart 2010, p. 5). The following section describes the propensity score approach.

Propensity Score Analysis

In a randomized study, the potential outcomes within treatment and control groups are unrelated to treatment assignment because individuals are randomly assigned to the groups. Consequently the treatment assignment given the variables X are strongly ignorable.

Rosenbaum and Rubin (1983) defined treatment assignment to be strongly ignorable when two conditions are met. The first condition (unconfoundedness) states that the potential outcomes $(Y(0), Y(1))$ and the treatment assignment T are conditionally independent given the observed baseline variables:

$$(Y(0), Y(1)) \perp\!\!\!\perp T \mid X = x$$

This condition is called the “no unmeasured confounders” assumption because it assumes that all the variables that affect both the outcome and the treatment assignment have been measured. The second condition (probabilistic assignment) states that there is a positive probability that a subject receives each treatment:

$$0 < \Pr(T = 1 \mid X = x) < 1$$

When the treatment assignment in an observational study is assumed to be strongly ignorable, Rosenbaum and Rubin (1983, p. 43) showed that unbiased estimates of average treatment effects can be obtained by conditioning on the propensity score $e(x)$, the probability of the treatment assignment conditional on a set of observed variables X :

$$e(x) = \Pr(T = 1 \mid X = x)$$

At any value of the propensity score $e(x)$, the difference between the treatment and control means is an unbiased estimate of the average treatment effect at $e(x)$. Consequently, matching on the propensity score and propensity score stratification also produce unbiased estimates of treatment effects (Rosenbaum and Rubin 1983, p. 44).

Furthermore, the propensity score is a balancing score. At each value of the propensity score, the distributions of the variables X are the same in the treated and control groups (Rosenbaum and Rubin 1983, p. 44; Stuart 2010, p. 6). Thus, the treatment assignment T and observed variables x are conditionally independent given the propensity score Rosenbaum (2010, p. 72):

$$x \perp\!\!\!\perp T \mid e(x)$$

Propensity score analysis attempts to replicate the properties of a randomized trial with respect to the observed variables X . The steps involved in this analysis are described in the section “[Process of Propensity Score Analysis](#)” on page 7677.

The following subsections describe the support region and the propensity score methods that are available in the PSMATCH procedure.

Support Region

For stratification and matching, the PSMATCH procedure selects observations whose propensity scores lie in a support region that can be defined in several ways:

- Selecting all available observations. You can request this definition by specifying `REGION=ALLOBS` in the `PROC PSMATCH` statement.
- Selecting observations whose propensity scores lie in a specified range. You can request this definition by specifying `REGION=ALLOBS` and then by additionally specifying range options.
- Selecting observations whose propensity scores lie in the region of common support for the propensity scores for observations in the treated and control groups. You can request this definition by specifying `REGION=CS`. This region can be extended by specifying the `EXTEND` suboption.
- Selecting observations whose propensity scores lie in the region of propensity scores for observations in the treated group. You can request this definition by specifying `REGION=TREATED`. This region can be extended by specifying the `EXTEND` suboption.

In combination with the `REGION=` option, you can specify the `OUT(OBS=REGION)` option in the `OUTPUT` statement to request that only those observations in the support region are to be included in the output data set. You can specify this combination even without the use of stratification or matching. For example, you can use the `REGION=ALLOBS(PMSIN=0.1 PSMAX=0.9)` option to include only observations with propensity scores greater than or equal to 0.1 and less than or equal to 0.9 in the output data set.

Propensity Score Methods

You can use the propensity score methods in the PSMATCH procedure to create an output data set that contains a sample that has been adjusted (either by matching, stratification, or weighting) so that the distributions of the variables are balanced between the treated and control groups. The two groups differ only randomly in their observed or measured variables, as in a randomized study. You can then use the output data set in an outcome analysis to estimate the effect of the treatment.

The following propensity score methods are available in the PSMATCH procedure:

- weighting, which creates weights that are appropriate for estimating the ATE and ATT
- stratification, which creates strata based on propensity scores
- matching, which matches treated units with control units

Note that the outcome variable is not involved in these methods. For more information about these methods, see the sections “[Propensity Score Weighting](#)” on page 7707, “[Propensity Score Stratification](#)” on page 7709, and “[Matching Process](#)” on page 7709.

Propensity Score Weighting

The PSMATCH procedure provides the following methods for weighting observations:

- inverse probability of treatment weighting, which is used to estimate the ATE
- ATT weighting (also referred to as weighting by odds), which is used to estimate the ATT
- weighting after matching, which is used to estimate the ATT

Inverse Probability of Treatment Weighting

Inverse probability of treatment weighting (IPTW) computes the weight for the j th observation with propensity score p_j as

$$w_j = \begin{cases} \frac{1}{p_j} & \text{for observations in the treated group} \\ \frac{1}{1-p_j} & \text{for observations in the control group} \end{cases}$$

These weights can be used in an outcome analysis to estimate the average treatment effect,

$$\text{ATE} = E(Y(1) - Y(0))$$

by weighting the two groups up to the full population. For example, for a treated unit with $p_j = 0.25$, the weight is 4, which represents four units in the full population.

You can specify the WEIGHT=ATEWGT option in the ASSESS statement to request weighted variable assessment that uses these weights, and you can use the ATEWGT= option in the OUTPUT statement to create a variable that contains these weights.

ATT Weighting

ATT weighting (also referred to as weighting by odds) computes the weight for the j th observation with propensity score p_j as

$$w_j = \begin{cases} 1 & \text{for observations in the treated group} \\ \frac{p_j}{1-p_j} & \text{for observations in the control group} \end{cases}$$

These weights can be used in an outcome analysis to estimate the average treatment effect for the treated units (individuals who actually receive treatment),

$$ATT = E(Y(1) - Y(0) \mid T = 1)$$

by weighting the control group up to the treated group. For example, for a control unit with $ps_j = 0.75$, the weight is 3, which represents three units in the treated population.

You can specify the `WEIGHT=ATTWGT` option in the `ASSESS` statement to request weighted variable assessment that uses these weights, and you can use the `ATTWGT=` option in the `OUTPUT` statement to create a variable that contains these weights.

Weighting after Matching

Except for matching with replacement, weights for use after matching are computed as

$$w_{gj} = \begin{cases} 1 & \text{for treated units in the } g\text{th matched set} \\ \frac{N_{gt}}{N_{gc}} & \text{for control units in the } g\text{th matched set} \end{cases}$$

where N_{gt} is the number of treated units and N_{gc} is the number of control units in the g th matched set.

The PSMATCH procedure computes these weights when you specify a `MATCH` statement, and they can be used to estimate the ATT because the total weight for the controls is equal to the total number of treated units in each matched group. For one-to-one greedy or optimal matching, the weight is 1 for both the treated and control units. Under a different matching algorithm, if the g th matched set contains $N_{gt}=1$ treated unit and $N_{gc}=3$ control units, then the weight for each treated unit is 1 and the weight for each control unit is $1/3$.

You can specify the `WEIGHT=MATCHWGT` option in the `ASSESS` statement to request weighted variable assessment that uses these weights, and you can use the `MATCHWGT=` option in the `OUTPUT` statement to create a variable that contains these weights.

For a k -to-one matching with replacement, the weight for each treated unit is 1 and the weight for each control unit is the number of its matched treated units divided by k . That is, if a control unit has three matched treated units in a one-to-one matching, then the weight for the control unit is 3. If a control unit has three matched treated units in a two-to-one matching, then the weight is $3/2$.

Propensity Score Stratification

Propensity stratification divides the observations into strata that have similar propensity scores, with the objective of balancing the observed variables between treated and control units within each stratum. The treatment effect can then be estimated by combining stratum-specific estimates of treatment effect. Rosenbaum and Rubin (1984, p. 521) show that an adjusted estimate of this type that is based on five strata can remove approximately 90% of the bias in the crude or unadjusted estimate.

The PSMATCH procedure performs stratification when you specify the STRATA statement and divide the observations contained in the support region (as specified in the REGION= option in the PROC PSMATCH statement) into the strata.

You can specify the KEY=TREATED option in the STRATA statement to allocate approximately the same number of treated units to each stratum. You can specify the KEY=NONE option to allocate approximately the same number of total units to each stratum.

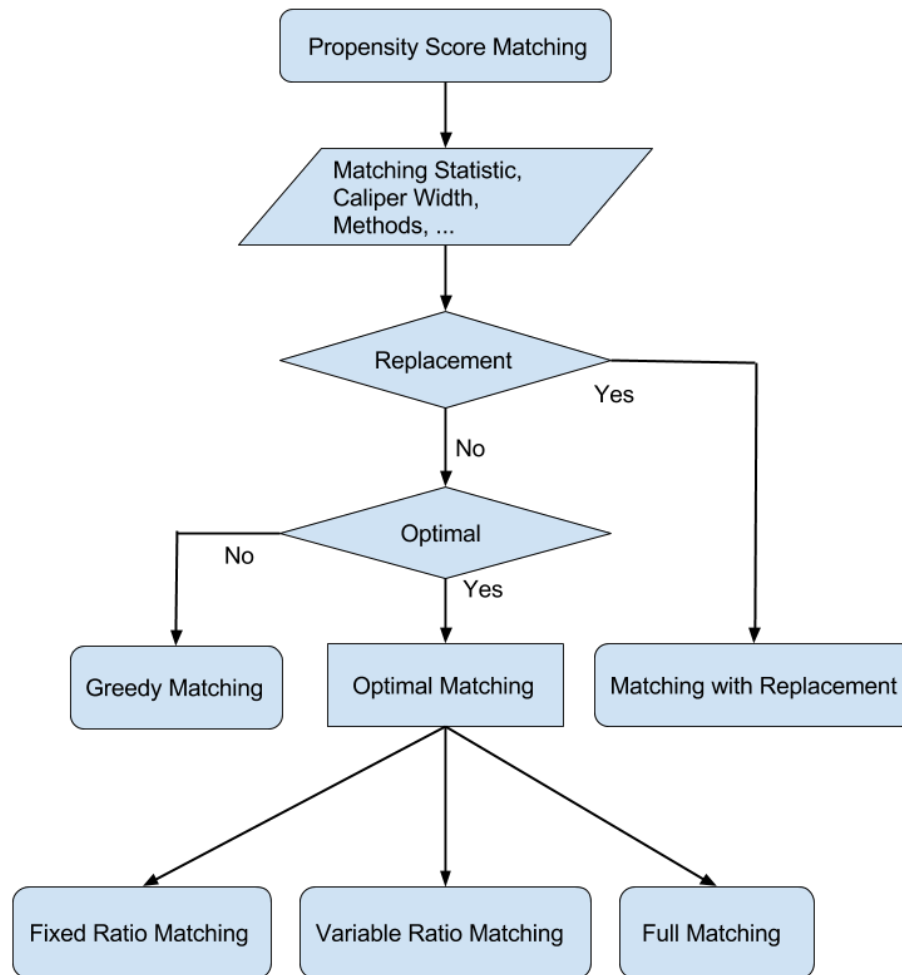
In the outcome analysis, you can use the weighted average of the stratum-specific treatment estimates to estimate the treatment effect. You can estimate the ATT if you weight by the stratum-specific number of treated units, and you can estimate the ATE if you weight by the sum of stratum-specific numbers of treated and control units (Stuart 2010, p. 13; Guo and Fraser 2015, pp. 76–77).

Matching Process

Except for matching with replacement in which multiple control units are matched to each treated unit, propensity score matching creates mutually exclusive sets of observations that have similar propensity scores. Each set has at least one treated unit and at least one control unit. The distribution of observed variables will be similar between treated units and control units in the matched sample.

For propensity score matching, Stuart (2010) reviews matching methods and provides guidance on their use. Austin (2014) provides a detailed comparison of algorithms for matching.

The flowchart in [Figure 95.14](#) summarizes the steps in propensity score matching.

Figure 95.14 Steps in Propensity Score Matching

The PSMATCH procedure provides the following strategies for matching observations in the control group to observations in the treatment group:

- greedy nearest neighbor matching, which sequentially and without replacement selects the control unit whose propensity score is closest to that of the given treated unit
- optimal matching, which selects all matches simultaneously and without replacement to minimize the total absolute difference in propensity score across all matches (this approach includes fixed ratio matching, variable ratio matching, and full matching)
- matching with replacement, which selects with replacement the control unit whose propensity score is closest to that of each treated unit

In addition to the propensity score, you can also use the logit of the propensity score and Mahalanobis distance as the matching statistic that is used to compare the closeness of two units. For more information about these matching methods, see the section “[Matching Methods](#)” on page 7712.

You can use the CALIPER= option in the MATCH statement to request that the difference in the propensity scores for a matched pair be less than or equal to a specified caliper width.

You can request exact matches of the levels of classification variables for treated and control units by specifying the EXACT= option in the MATCH statement.

Matching Statistics

The PSMATCH procedure provides the following types of statistic on which to match observations in the treated group with observations in the control group (you specify the statistic with the STAT= option in the MATCH statement):

- the difference in the logit of the propensity score (STAT=LPS; this is the default)
- the Mahalanobis distance between sets of continuous variables (STAT=MAH)
- the difference in the propensity score (STAT=PS)

Let p_{ti} and p_{cj} be the propensity scores of the i th treated unit and the j th control unit, respectively. When you specify STAT=PS, matching is based on the absolute difference,

$$|p_{ti} - p_{cj}|$$

When you specify STAT=LPS, matching is based on the absolute difference,

$$|l_{ti} - l_{cj}|$$

where $l_{ti} = \text{logit}(p_{ti})$ and $l_{cj} = \text{logit}(p_{cj})$ are the logits of the propensity scores.

When you specify STAT=MAH, two different distances can be used as the Mahalanobis distance in matching (as specified in the SQRT= suboption of the STAT=MAH option):

- $\sqrt{(\mathbf{X}_{ti} - \mathbf{X}_{cj})' \mathbf{V}^{-1} (\mathbf{X}_{ti} - \mathbf{X}_{cj})}$ (SQRT=YES; this is the default)
- $(\mathbf{X}_{ti} - \mathbf{X}_{cj})' \mathbf{V}^{-1} (\mathbf{X}_{ti} - \mathbf{X}_{cj})$ (SQRT=NO)

where $\mathbf{X}_{ti}$ is the set of variables of the i th treated unit, $\mathbf{X}_{cj}$ is the set of variables of the j th control unit, and $\mathbf{V}$ is the covariance matrix of $\mathbf{X}$.

Note that the SQRT= option does not affect the results for greedy nearest neighbor matching and matching with replacement; it affects only the results for optimal matching.

Three different covariance matrices can be used to compute the Mahalanobis distance (as specified in the COV= suboption of the STAT=MAH option):

- the covariance matrix that is based on observations in the control group (COV=CONTROL; this is the default)

- the pooled covariance matrix that is based on observations in the treated and control groups (COV=POOLED)
- the identity matrix (COV=IDENTITY) that yields the Euclidean distance.

Note that you can also include the propensity score and the logit of propensity in the Mahalanobis distance. For example, when you specify `STAT=MAH(PS VAR=(X1 X2 X3) / COV=POOLED)`, the PSMATCH procedure computes the Mahalanobis distance between observations in the treated and control groups by using the propensity score and variables X1, X2, and X3. The covariance matrix is the pooled covariance matrix of the treated and control groups.

Matching Methods

When you specify the MATCH statement, the PSMATCH procedure matches observations in the control group to observations in the treatment group by using one of the methods that are described in the following subsections. You can request the method with the METHOD= option.

Greedy Nearest Neighbor Matching

Greedy nearest neighbor matching, requested by the METHOD=GREEDY option, selects the control unit whose propensity score best matches the propensity score of each treated unit. Greedy nearest neighbor matching is done sequentially and without replacement.

The following criteria are available for greedy nearest neighbor matching:

- the number of control units matched to each treated unit (you can specify this number in the K= suboption)
- the order of propensity scores of treated units, which can be ascending, descending, or random (you can specify the order in the ORDER= suboption)

Replacement Matching

Replacement matching, requested by the METHOD=REPLACEMENT option, selects with replacement the control unit whose propensity score is closest to the propensity score for each treated unit. You can specify the number of control units to be matched to each treated unit in the K= suboption.

Optimal Matching

Optimal matching selects all matches simultaneously and without replacement to minimize the total absolute difference in propensity score across all matches. You can request the following optimal matching methods:

- fixed ratio matching, requested by the METHOD=OPTIMAL option, matches a fixed number of control units to each treated unit.
- variable ratio matching, requested by the METHOD=VARRATIO option, matches one or more control units to each treated unit.

- full matching, requested by the METHOD=FULL option, matches each treated unit to one or more control units and matches each control unit to one or more treated units. By additionally specifying the KMEAN=, NCONTROL=, or PCTCONTROL= suboptions, you can request constrained full matching in which the number of matched control units is less than the total number of available controls.

As alternatives to matching on the propensity score, you can match on the logit of the propensity score or use the Mahalanobis distance to match on a set of variables (possibly including the PS or the LPS). All three of these methods minimize the total absolute difference across all matches in the matching statistic, which is the total difference in the logit of propensity score by default.

Table 95.4 lists the suboptions available for optimal matching. The symbol "X" indicates that the option is applicable for the specified method.

Table 95.4 Applicable Options for Optimal Matching

METHOD=	K=	KMIN=	KMAX=	KMAXTRT=	KMEAN= NCONTROL= PCTCONTROL=
OPTIMAL	X				
VARRATIO		X	X		X
FULL			X	X	X

- K= specifies the number of control units that are matched to each treated unit.
- KMIN= specifies the minimum number of control units that are matched to each treated unit.
- KMAX= specifies the maximum number of control units that are matched to each treated unit.
- KMAXTRT= specifies the maximum number of treated units that are matched to each matched control unit.
- KMEAN= specifies the average number of control units that are matched to each treated unit.
- NCONTROL= specifies the total number of control units that are matched.
- PCTCONTROL= specifies the percentage of control units that are matched.

You can specify only one of the KMEAN=, NCONTROL=, and PCTCONTROL= options.

Variable Balance Assessment

Propensity score analysis assumes that the true propensity scores are known. When the propensity scores are estimated—as is usually the case in practice—you need to assess how well the distributions of the propensity scores (or the logit propensity scores) and the adjusted variables are balanced between the treatment group and the control group.

The ASSESS statement in the PSMATCH procedure provides a variety of statistical measures and graphical displays for comparing these distributions. You can make these assessments for all the observations in the data set, the observations in the support region, or the matched observations (if you specify a MATCH statement).

Two statistical measures for variable balance assessment are the standardized difference between the treatment and control groups and the variance ratio. For good variable balance, the absolute standardized difference should be less than or equal to 0.25, and the variance ratio should be between 0.5 and 2 (Rubin 2001, p. 174; Stuart 2010, p. 11).

Note that in addition to the threshold of 0.25 for the standardized difference, a smaller threshold of 0.1 has also been used to indicate meaningful imbalance in the variables (Normand et al. 2001; Mamdani et al. 2005; Austin 2009).

The standardized difference is computed by dividing the mean difference by an estimate of its standard deviation. Two estimates of the standard deviation are available:

- the square root of the average of the variances in the treatment and control groups (Rosenbaum and Rubin 1985, p. 37),
- the standard deviation of observations in the treatment group only (Stuart 2010, p. 11)

For binary classification variables, the mean is taken to be the average proportion p of the first classification level, and the variance is computed as $p(1 - p)$ (Austin, Grootendorst, and Anderson 2007, p. 737).

If you specify a STRATA statement, then stratum-specific standardized mean differences are computed for observations in the support region.

The PSMATCH procedure displays the standardized differences in plots. You can also request box plots for continuous variables, bar charts for binary classification variables, and cloud plots for both continuous and binary classification variables. These plots are also produced by stratum if you specify a STRATA statement.

The next three subsections describe how standardized mean differences and variance ratios are computed for all observations, observations in the support region, and matched observations.

Standardized Mean Differences for All Observations

For all observations in the data set, let $\bar{x}_{t(\text{all})}$ be the mean of the observations in the treatment group and let $\bar{x}_{c(\text{all})}$ be the mean of the observations in the control group, with corresponding sample variances $V(x_{t(\text{all})})$ and $V(x_{c(\text{all})})$. Then the standardized mean difference is

$$d_{(\text{all})} = \frac{\bar{x}_{t(\text{all})} - \bar{x}_{c(\text{all})}}{s_{(\text{all})}}$$

By default (or if you specify the STDDIFFDIV=POOLED option), the divisor is the pooled standard deviation

$$s_{(\text{all})} = \sqrt{\frac{V(x_{t(\text{all})}) + V(x_{c(\text{all})})}{2}}$$

Alternatively, if you specify the STDDIFFDIV=TREATED option, the divisor is the standard deviation for observations in the treatment group only,

$$s_{(\text{all})} = \sqrt{V(x_{t(\text{all})})}$$

The variance ratio is

$$\frac{V(x_{t(\text{all})})}{V(x_{c(\text{all})})}$$

Standardized Mean Differences for Observations in the Support Region

For observations in the support region, let $\bar{x}_{t(\text{region})}$ be the mean of observations in the treatment group and $\bar{x}_{c(\text{region})}$ be the mean of observations in the control group, with corresponding sample variances $V(x_{t(\text{region})})$ and $V(x_{c(\text{region})})$. Then the standardized difference is

$$d_{(\text{region})} = \frac{\bar{x}_{t(\text{region})} - \bar{x}_{c(\text{region})}}{s_{(\text{all})}}$$

where the divisor is derived from all the observations, and the variance ratio is

$$\frac{V(x_{t(\text{region})})}{V(x_{c(\text{region})})}$$

The reduction percentage for the standardized mean difference is computed as

$$100 \times \frac{\max(|d_{(\text{all})}| - |d_{(\text{region})}|, 0)}{|d_{(\text{all})}|}$$

If you specify a STRATA statement, the stratum-specific standardized difference is

$$d_{(g)} = \frac{\bar{x}_{t(g)} - \bar{x}_{c(g)}}{s_{(\text{all})}}$$

where g is the stratum index, and $\bar{x}_{t(g)}$ and $\bar{x}_{c(g)}$ are the means of the observations in the treatment and control groups, respectively, in the g th stratum of the support region.

Standardized Mean Differences for Matched Observations

Let $\bar{x}_{t(\text{matched})}$ be the mean of matched observations in the treatment group, and let $\bar{x}_{c(\text{matched})}$ be the mean of matched observations in the control group, with corresponding sample variances $V(x_{t(\text{matched})})$ and $V(x_{c(\text{matched})})$. Then the standardized difference is

$$\frac{\bar{x}_{t(\text{matched})} - \bar{x}_{c(\text{matched})}}{d_{(\text{all})}}$$

where the divisor is derived by using all observations, and the variance ratio is

$$\frac{V(x_{t(\text{matched})})}{V(x_{c(\text{matched})})}$$

The reduction percentage for the standardized mean difference is computed as

$$100 \times \frac{\max(|d_{(\text{all})}| - |d_{(\text{matched})}|, 0)}{|d_{(\text{all})}|}$$

Table Output

By default, the PSMATCH procedure displays the “Data Information” and “Propensity Score Information” tables. If you specify a MATCH statement, the procedure also displays the “Matching Information” table. If you specify a STRATA statement, the procedure also displays the “Strata Information” table.

If you specify the ASSESS statement, the “Standardized Variable Differences” table is displayed. In addition, if you specify a STRATA statement, the “Strata Standardized Variable Differences” table is also displayed.

If you specify the VARINFO option in the ASSESS statement, the “Variable Information” table is displayed. In addition, if you specify a STRATA statement, the “Strata Variable Information” table is also displayed.

Data Information

The “Data Information” table displays the names of the input and output data sets, the number of observations in the treated group and the control group, and the number of observations in the support region that are in the treated group and the control group. The minimum and maximum propensity scores for observations in the support region are also displayed.

Matching Information

The “Matching Information” table displays the matching statistic, the matching method, and the caliper width. The table also displays the number of matched sets of observations, the numbers of matched observations in the treated and control groups, and the total absolute difference across all matches.

Propensity Score Information

The “Propensity Score Information” table displays descriptive statistics (the number of observations, mean, standard deviation, minimum, and maximum) for the propensity scores of observations in the treated group and the control group. These statistics are computed using all observations, observations in the support region, and matched observations (if you specify a MATCH statement).

Standardized Variable Differences

The “Standardized Variable Differences” table displays statistics that summarize the differences in the variables and the logit propensity score (LPS) between the treated and control groups. These statistics are computed using all observations, observations in the support region, and matched observations (if you specify a MATCH statement).

The statistics include the following:

- the mean difference between observations in the treated and control groups
- the divisor that is used to compute the standardized mean difference, $\sqrt{(V_t + V_c)/2}$, where V_t and V_c are the sample variances of all observations in the treated and control groups
- the standardized mean difference, which is the mean difference divided by the divisor
- the reduction percentage of the standardized mean difference for observations in the support region, compared with the standardized mean difference of all observations (this statistic is also computed for matched observations if you specify a MATCH statement)

- the ratio of variances for observations in the treated and control groups

Strata Information

The “Strata Information” table displays descriptive statistics that include the propensity score range, the number of observations in the treated group and the control group, and the total number of observations in each stratum.

Strata Standardized Variable Differences

The “Strata Standardized Variable Differences” table displays the variable difference statistics between the treated and control groups in each stratum.

For each variable, the statistics include the following:

- the mean difference between observations in the treated and control groups
- the standardized mean difference, which is the mean difference divided by the divisor (that is displayed in the “Standardized Variable Differences” table)
- the reduction percentage of the standardized mean difference for observations in the stratum, compared with the standardized mean difference in absolute value of all observations
- the ratio of variances between observations in the treated and control groups in each stratum

Strata Variable Information

The “Strata Variable Information” table displays descriptive statistics that include the number of observations, variable mean, and standard deviation of the observations in each of the treatment and control groups in each stratum. For continuous variables, the statistics also include the minimum and maximum.

Variable Information

For variables that are specified in the ASSESS statement, the “Variable Information” table displays descriptive statistics that are computed using all observations and observations in each of the treatment and control groups in the support region.

These statistics include the sample size, mean, and standard deviation. For continuous variables, the statistics also include the minimum and maximum. If you specify a MATCH statement, the table also displays descriptive statistics for the matched observations in the treatment and control groups.

ODS Table Names

PROC PSMATCH assigns a name to each table it creates. You must use these names to refer to tables when you use the Output Delivery System (ODS). These names are listed in [Table 95.5](#). For more information about ODS, see Chapter 20, “[Using the Output Delivery System](#).”

Table 95.5 ODS Tables Produced by PROC PSMATCH

ODS Table Name	Description	Statement	Option
DataInfo	Data information		
MatchInfo	Matching information	MATCH	
PSInfo	Propensity score information		
StdVarDiff	Standardized differences between the treated group and the control group	ASSESS	
StrataInfo	Strata information	STRATA	
StrataStdVarDiff	Strata standardized differences between the treated group and the control group	ASSESS STRATA	
StrataVarInfo	Strata variable information	ASSESS STRATA	VARINFO
VarInfo	Variable information	ASSESS	VARINFO

Graphics Output

This section describes the use of ODS for creating graphics with the PSMATCH procedure. To request these graphs, ODS Graphics must be enabled and you must specify the ASSESS option. In addition, except for the standardized differences plot (which is the default) you must use the PLOTS= option in the ASSESS statement to specify the plots. For more information about ODS Graphics, see Chapter 21, “[Statistical Graphics Using ODS](#).”

Variable Bar Chart

The PLOTS=BARCHART option displays bar charts for binary classification variables in the treated and control groups for all observations and for observations in the support region. If you specify the MATCH statement, bar charts are also created for matched observations.

Variable Box Plot

The PLOTS=BOXPLOT option displays box plots for continuous variables in the treated and control groups for all observations and for observations in the support region. If you specify the MATCH statement, box plots are also created for matched observations.

Variable Cloud Plot

The PLOTS=CLOUDPLOT option displays cloud plots for continuous and binary classification variables in the treated and control groups for all observations and for observations in the support region. If you specify the MATCH statement, cloud plots are also created for matched observations. Here the term cloud plot refers to a scatter plot in which the points are jittered to prevent overplotting by adding random noise to data in the plot. For example, with a continuous variable and the default ORIENT=HORIZONTAL option, the variable values are displayed horizontally and the treated and control groups are displayed vertically. While the exact variable values are displayed along the horizontal axis, the points are jittered in the vertical direction.

Standardized Variable Differences Plot

The `PLOTS=STDDIFFPLOT` option displays a plot of the standardized differences for continuous and binary classification variables for all observations and for observations in the support region. If you specify the `MATCH` statement, plots are also created for matched observations.

Strata Variable Bar Chart

If you specify a `STRATA` statement, the `PLOTS=BARCHART` option displays bar charts for binary classification variables in the treated and control groups for the observations in each stratum.

Strata Variable Box Plot

If you specify a `STRATA` statement, the `PLOTS=BOXPLOT` option displays box plots for continuous variables in the treated and control groups for the observations in each stratum.

Strata Variable Cloud Plot

If you specify a `STRATA` statement, the `PLOTS=CLOUDPLOT` option displays cloud plots for continuous and binary classification variables in the treated and control groups for the observations in each stratum. The cloud plot is also referred as a jittered scatter plot and is used to prevent overplotting by adding random noise to data in the plot.

Strata Standardized Variable Differences Plot

If you specify a `STRATA` statement, the `PLOTS=STDDIFFPLOT` option displays standardized differences plots for continuous and binary classification variables for the observations in each stratum.

ODS Graphics

Statistical procedures use ODS Graphics to create graphs as part of their output. ODS Graphics is described in detail in Chapter 21, [“Statistical Graphics Using ODS.”](#)

Before you create graphs, ODS Graphics must be enabled (for example, by specifying the `ODS GRAPHICS ON` statement). For more information about enabling and disabling ODS Graphics, see the section [“Enabling and Disabling ODS Graphics”](#) on page 607 in Chapter 21, [“Statistical Graphics Using ODS.”](#)

The overall appearance of graphs is controlled by ODS styles. Styles and other aspects of using ODS Graphics are discussed in the section [“A Primer on ODS Statistical Graphics”](#) on page 606 in Chapter 21, [“Statistical Graphics Using ODS.”](#)

`PROC PSMATCH` assigns a name to each graph it creates. You can use these names to refer to the graphs when you use ODS. To request the graph, ODS Graphics must be enabled and you must specify the `ASSESS` option. In addition, except for the standardized differences plot (which is the default), you must use the `PLOTS=` option in the `ASSESS` statement to specify the plots, as indicated in [Table 95.6](#).

Table 95.6 Graphs Produced by PROC PSMATCH

ODS Graph Name	Plot Description	Statement	PLOTS=
VarBarChart	Binary variable bar chart	ASSESS	BARChart
VarBoxPlot	Continuous variable box plot	ASSESS	BOXPLOT
VarCloudPlot	Variable cloud plot	ASSESS	CLOUDPLOT
StdVarDiffPlot	Standardized differences plot	ASSESS	STDDIFFPLOT
StrataVarBarChart	Strata binary variable bar chart	ASSESS, STRATA	BARChart
StrataVarBoxPlot	Strata continuous variable box plot	ASSESS, STRATA	BOXPLOT
StrataVarCloudPlot	Strata variable cloud plot	ASSESS, STRATA	CLOUDPLOT
StrataStdVarDiffPlot	Strata standardized differences plot	ASSESS, STRATA	STDDIFFPLOT

Examples: PSMATCH Procedure

In practice, the outcome data for an observational study might or might not be available at the time that a propensity score analysis is done. You can handle these situations as follows:

- When the outcome data are not yet available, you might not need to retain the covariate data for all individuals in the study in order to carry out the outcome analysis. For example, if you use the matching method for propensity score analysis, only the matched units are needed for follow-up. Retaining only the matched units reduces the cost of the study (Stuart 2010, p. 2). The clinical trial described in “[Getting Started: PSMATCH Procedure](#)” on page 7680 is an example in which outcome data are not yet available.
- When the outcome data are available at the time of the propensity score analysis, they should not be used in the analysis (Stuart 2010, p. 2). In [Example 95.4](#), the question is whether taking a music class improves grade point averages, and the grades together with other measures are available when the propensity score analysis is done at the completion of the school year.

The examples in this section illustrate the main methods for propensity score analysis that are available in the PSMATCH procedure. For simplicity, the examples use only a few variables. In practice, propensity score analysis often involves many more variables.

Example 95.1: Propensity Score Weighting

This example creates observation weights for all patients in the trial of a propensity score analysis. The Drugs data set contains the patient information and is described in the section “[Getting Started: PSMATCH Procedure](#)” on page 7680.

The following statements invoke the PSMATCH procedure and create observation weights that are appropriate to estimate the average treatment effect for the treated (ATT):

```
ods graphics on;
proc psmatch data=drugs region=allobs;
  class Drug Gender;
  psmodel Drug (Treated='Drug_X')= Gender Age Bmi;
  assess lps var=(Gender Age Bmi)
    / varinfo plots=(boxplot barchart)
    weight=attwgt;
  output out (obs=all)=OutEx1;
run;
```

The CLASS statement specifies the classification variables. The PSMODEL statement specifies the logistic regression model that creates the propensity score for each observation, which is the probability that the patient receives Drug_X. The Drug variable is the binary treatment indicator variable, and TREATED='Drug_X' identifies Drug_X as the treated group. The Gender, Age, and Bmi variables are included in the model because they are believed to be related to the assignment.

The REGION= option specifies an interval region of propensity scores (or equivalently, logits of propensity scores) such that only observations that have propensity scores in the region are used in stratification and matching. Even without stratification and matching, you can still use the REGION= option to select observations in the region to compare variable differences between observations in the treatment and control groups. The REGION=ALLOBS option selects all available observations.

The “Data Information” table in [Output 95.1.1](#) displays information about the input and output data sets, the numbers of observations in the treated and control groups, the lower and upper limits for the propensity score support region, and the numbers of observations in the treated and control groups that fall within the support region. Because REGION=ALLOBS is specified, all 373 observations in the control group fall within the support region.

Output 95.1.1 Data Information
The PSMATCH Procedure

Data Information	
Data Set	WORK.DRUGS
Output Data Set	WORK.OUTEX1
Treatment Variable	Drug
Treatment Group	Drug_X
All Obs (Treatment)	113
All Obs (Control)	373
Support Region	All Obs
Lower PS Support	0.020157
Upper PS Support	0.685757
Support Region Obs (Treatment)	113
Support Region Obs (Control)	373

The “Propensity Score Information” table in [Output 95.1.2](#) displays summary statistics by treatment group for all observations and for the support region observations. Because `REGION=ALLOBS` is specified, all observations are in the support region.

Output 95.1.2 Propensity Score Information

Propensity Score Information										
		Treated (Drug = Drug_X)				Control (Drug = Drug_A)				
Observations	N	Mean	Std Dev	Minimum	Maximum	N	Mean	Std Dev	Minimum	Maximum
All	113	0.310773	0.132467	0.060231	0.641148	373	0.208801	0.131969	0.020157	0.685757
Region	113	0.310773	0.132467	0.060231	0.641148	373	0.208801	0.131969	0.020157	0.685757

The `ASSESS` statement produces the tables and plots that summarize differences in the specified variables between treated and control groups for all observations and for the support region observations. As requested by the `LPS` and `VAR=` options, the variables listed in the table are the logit of propensity score and the variables `Gender`, `Age`, and `Bmi`. The `WEIGHT=ATTWGT` option also summarizes differences in the specified `VAR=` variables between treated and control group of differences for the weighted observations.

The “Variable Information” table in [Output 95.1.3](#) displays variable differences between the treated and control groups for all observations and for the support region observations. With the `WEIGHT=ATTWGT` option, the differences for the weighted region observations are also displayed. For a binary classification variable (`Gender`), the difference is in the proportion of the first ordered level (`Female`).

Output 95.1.3 Variable Information**The PSMATCH Procedure**

Variable Information						
Treated (Drug = Drug_X)						
Variable	Observations	N	Weight	Mean	Std Dev	Minimum Maximum
LPS	All	113		-0.880615	0.681761	-2.747444 0.580348
	Region	113		-0.880615	0.681761	-2.747444 0.580348
Age	All	113		36.309735	5.534114	26.000000 49.000000
	Region	113		36.309735	5.534114	26.000000 49.000000
	ATT Weighted Region	113	113.00	36.309735	5.534114	26.000000 49.000000
Bmi	All	113		24.492566	1.863797	20.330000 28.340000
	Region	113		24.492566	1.863797	20.330000 28.340000
	ATT Weighted Region	113	113.00	24.492566	1.863797	20.330000 28.340000
Gender	All	113		0.433628	0.495575	
	Region	113		0.433628	0.495575	
	ATT Weighted Region	113	113.00	0.433628	0.495575	

Variable Information						
Control (Drug = Drug_A)						
Variable	Observations	N	Weight	Mean	Std Dev	Minimum Maximum
LPS	All	373		-1.520586	0.844486	-3.883858 0.780357
	Region	373		-1.520586	0.844486	-3.883858 0.780357
Age	All	373		40.404826	6.579103	25.000000 57.000000
	Region	373		40.404826	6.579103	25.000000 57.000000
	ATT Weighted Region	373	116.59	35.877442	6.226667	25.000000 57.000000
Bmi	All	373		23.753271	1.980778	19.220000 28.610000
	Region	373		23.753271	1.980778	19.220000 28.610000
	ATT Weighted Region	373	116.59	24.600042	1.948818	19.220000 28.610000
Gender	All	373		0.458445	0.498270	
	Region	373		0.458445	0.498270	
	ATT Weighted Region	373	116.59	0.443106	0.496753	

With REGION=ALLOBS, the statistics are identical between all observations and the support region observations. In addition, the statistics are also identical to weighted support region observations in the treated group, because each treated unit receives a weight of 1 when WEIGHT=ATTWGT. The total weight of the control units is 116.59, which is close to 113, the total weight of treated units. Also, the weighted variable means for support region control units are closer to the corresponding variable means for support region treated units. The statistics for the logit of propensity score are not displayed because the ATT weights for observations are computed from their propensity scores.

The “Standardized Variable Differences” table, as shown in [Output 95.1.4](#), displays standardized differences between the treated and control groups for all observations, the support region observations, and the weighted support region observations.

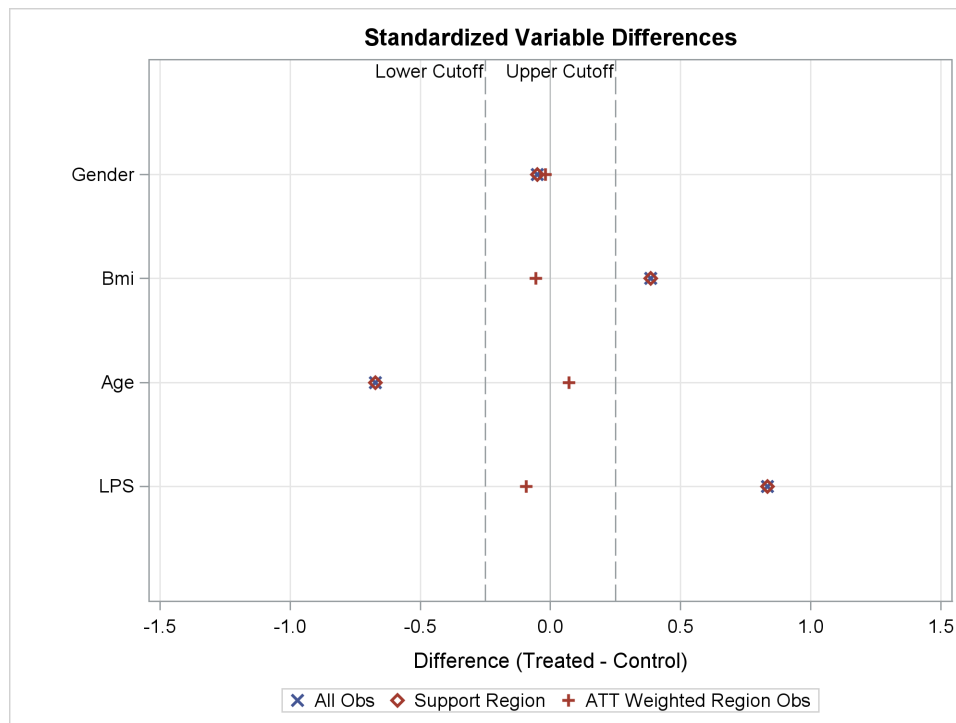
Output 95.1.4 Standardized Differences

Standardized Variable Differences (Treated - Control)									
Standardized Mean Difference									
Mean Difference					Mean Difference			Percent Reduction	
Variable	All Obs	Region Obs	ATT	Divisor	All Obs	Region Obs	ATT	Region Obs	ATT
			Weighted Region Obs				Weighted Region Obs		Weighted Region Obs
LPS	0.639971	0.639971		0.767448	0.833894	0.833894		0.00	
Age	-4.095091	-4.095091	0.432293	6.079104	-0.673634	-0.673634	0.071111	0.00	89.44
Bmi	0.739296	0.739296	-0.107476	1.923178	0.384414	0.384414	-0.055884	0.00	85.46
Gender	-0.024817	-0.024817	-0.009478	0.496925	-0.049941	-0.049941	-0.019073	0.00	61.81

Standardized Variable Differences (Treated - Control)				
Variance Ratio				
Variable	All Obs	Region Obs	ATT	
			Weighted Region Obs	
LPS	0.6517	0.6517		
Age	0.7076	0.7076	0.7899	
Bmi	0.8854	0.8854	0.9147	
Gender	0.9892	0.9892	0.9953	

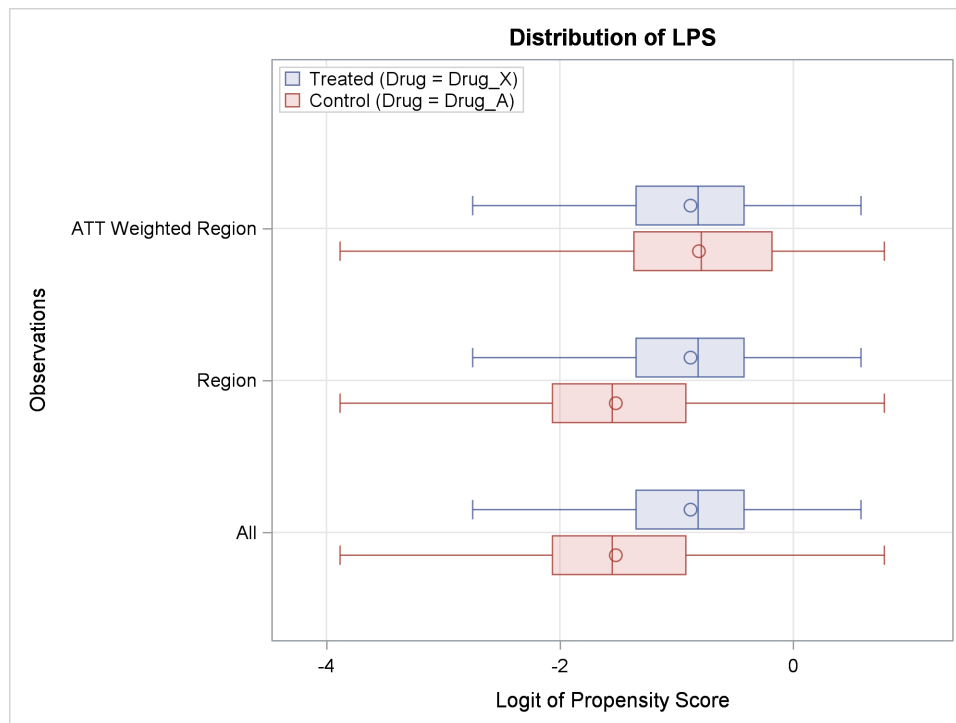
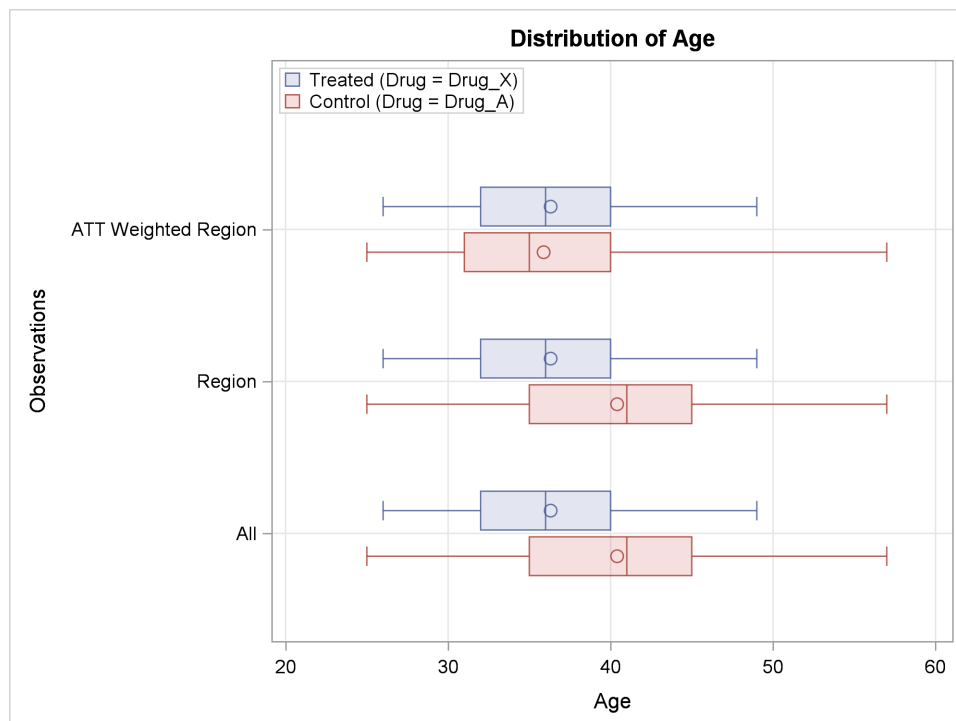
The standardized mean differences are significantly reduced in the weighted region observations; the largest of these differences is 0.0711 in absolute value, which is less than the recommended upper limit of 0.25 (Rubin 2001, p. 174; Stuart 2010, p. 11). The variance ratios between the two groups are within the recommended range of 0.5 to 2. With REGION=ALLOBS, the percentage of reduction in variable mean difference is 0 for the support region observations.

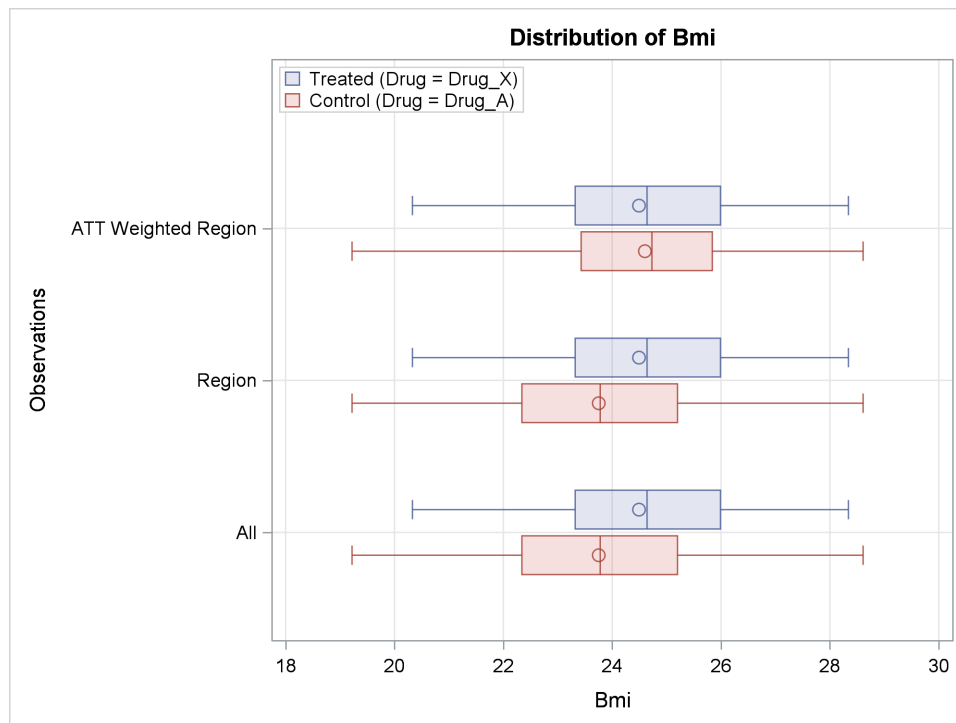
When ODS Graphics is enabled, the PSMATCH procedure displays a standardized variable differences plot for the variables that are specified in the ASSESS statement, as shown in [Output 95.1.5](#).

Output 95.1.5 Standardized Differences Plot

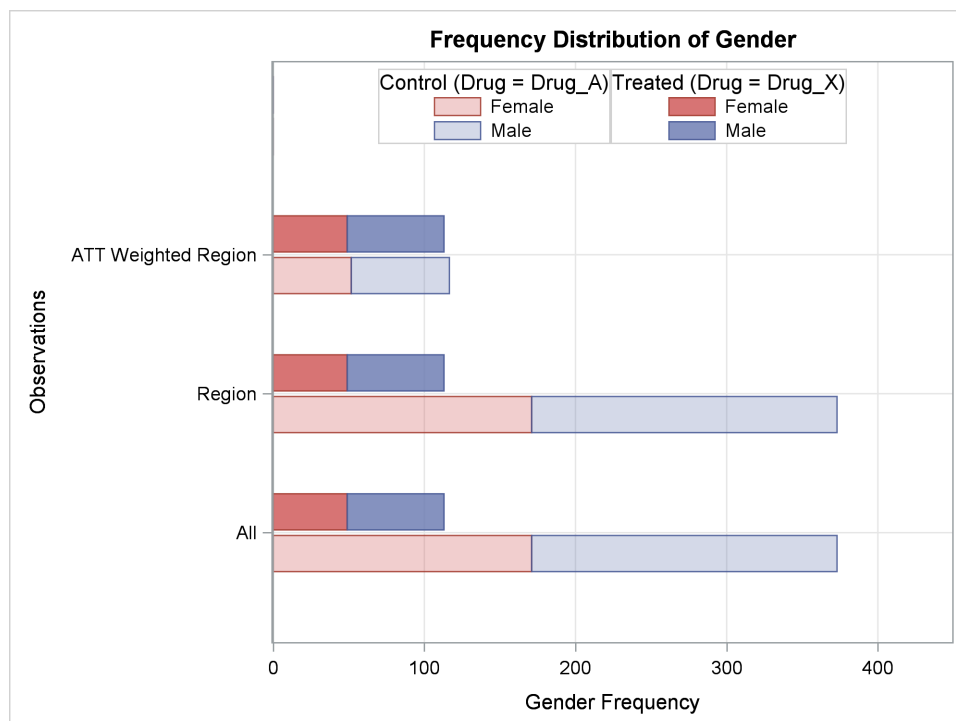
The “Standardized Variable Differences Plot” displays the standardized differences in the “Variable Differences” table in [Output 95.1.4](#). All differences for the matched observations are within the recommended limits of -0.25 and 0.25 , which are indicated by reference lines.

The `PLOTS=BOXPLOT` option requests a box plot for the logit of propensity score (LPS) and for each continuous variable specified in the `ASSESS` statement, as shown in [Output 95.1.6](#), [Output 95.1.7](#), and [Output 95.1.8](#). The box plots show good variable balance for the weighted support region observations.

Output 95.1.6 LPS Box Plot**Output 95.1.7** Age Box Plot

Output 95.1.8 Bmi Box Plot

The PLOTS=BARCHART option displays a bar chart for each classification variable that is specified in the ASSESS statement, as shown in [Output 95.1.9](#). The bar chart displays similar distributions of Gender between males and females for the weighted support region observations.

Output 95.1.9 Gender Bar Chart

Because there is good balance in the weighted distributions of the variables Gender, Age, and Bmi, you can output all observations (including added observation weights) so that they can be used for subsequent weighted outcome analysis.

If you are not satisfied with the variable balance, you can do one or more of the following until you are satisfied: you can select another set of variables to fit the propensity score model, you can modify the specification of the propensity score model by using nonlinear terms for the continuous variables or by adding interactions (Rosenbaum and Rubin 1984), or you can choose another propensity score method (such as matching).

The OUT(OBS=ALL)=OutEx1 option in the OUTPUT statement creates an output data set, OutEx1, that contains all available observations. The following statements list the first 10 observations in OutEx1, as shown in [Output 95.1.10](#).

```
proc print data=OutEx1(obs=10);
  var PatientID Drug Gender Age Bmi _ps_ _AttWgt_;
run;
```

Output 95.1.10 Output Data Set with PS Weights

Obs	PatientID	Drug	Gender	Age	Bmi	_PS_	_ATTWGT_
1	284	Drug_X	Male	29	22.02	0.36444	1.00000
2	201	Drug_A	Male	45	26.68	0.22296	0.28694
3	147	Drug_A	Male	42	21.84	0.11323	0.12768
4	307	Drug_X	Male	38	22.71	0.19733	1.00000
5	433	Drug_A	Male	31	22.76	0.35311	0.54586
6	435	Drug_A	Male	43	26.86	0.27263	0.37482
7	159	Drug_A	Female	45	25.47	0.14911	0.17523
8	368	Drug_A	Female	49	24.28	0.07780	0.08437
9	286	Drug_A	Male	31	23.31	0.38341	0.62182
10	163	Drug_X	Female	39	25.34	0.24995	1.00000

By default, the output data set includes the variable `_PS_` (which provides the propensity score) and the variable `_ATTWGT_` (which provides matched observation weights). The weight for each treated unit is 1, and the weight for each control unit is computed as $p / (1 - p)$, where p is the propensity score.

If you assume that no other confounding variables are associated with both the response variable and the treatment group indicator Drug, then after the responses for the trial are observed and added to the data set OutEx1, you can use the same outcome analysis with weights on this output data set as you would have used on the original data set Drugs (augmented with responses) to estimate the treatment effect.

Example 95.2: Propensity Score Stratification

This example creates strata of observations that are based on propensity scores for patients in the trial in a propensity score analysis. The Drugs data set contains the patient information and is described in the section “Getting Started: PSMATCH Procedure” on page 7680.

The following statements invoke the PSMATCH procedure and create five strata that are based on propensity scores:

```
ods graphics on;
proc psmatch data=drugs region=allobs;
  class Drug Gender;
  psmodel Drug(Treated='Drug_X')= Gender Age Bmi;
  strata nstrata=5;
  assess ps var=(Gender Age Bmi)
    / varinfo plots=(boxplot barchart);
  output out(obs=all)=OutEx2;
run;
```

The CLASS statement specifies the classification variables. The PSMODEL statement specifies the logistic regression model that creates the propensity score for each observation, which is the probability that the patient receives Drug_X. The Drug variable is the binary treatment indicator variable, and TREATED='Drug_X' identifies Drug_X as the treated group. The Gender, Age, and Bmi variables are included in the model because they are believed to be related to the assignment.

The REGION= option specifies an interval region of propensity scores such that only observations that have propensity scores in the region are used in stratification and matching. You can also use the REGION= option to select observations in the region to compare variable differences between observations in the treatment and control groups. The REGION=ALLOBS option selects all available observations.

The STRATA statement creates strata of observations based on propensity scores. The NSTRATA=5 option (which is the default) stratifies observations in the support region into five strata.

The “Data Information” table in [Output 95.2.1](#) displays information about the input and output data sets, the numbers of observations in the treated and control groups, the lower and upper limits for the propensity score support region, and the numbers of observations in the treated and control groups that fall within the support region. Note that because REGION=ALLOBS, all 373 observations in the control group fall within the support region.

Output 95.2.1 Data Information

The PSMATCH Procedure

Data Information	
Data Set	WORK.DRUGS
Output Data Set	WORK.OUTEX2
Treatment Variable	Drug
Treatment Group	Drug_X
All Obs (Treatment)	113
All Obs (Control)	373
Support Region	All Obs
Lower PS Support	0.020157
Upper PS Support	0.685757
Support Region Obs (Treatment)	113
Support Region Obs (Control)	373
Number of Strata	5

The “Propensity Score Information” table in [Output 95.2.2](#) displays summary statistics by treatment group for all observations and for the support region observations.

Output 95.2.2 Propensity Score Information

Propensity Score Information										
Treated (Drug = Drug_X)						Control (Drug = Drug_A)				
Observations	N	Mean	Std Dev	Minimum	Maximum	N	Mean	Std Dev	Minimum	Maximum
All	113	0.310773	0.132467	0.060231	0.641148	373	0.208801	0.131969	0.020157	0.685757
Region	113	0.310773	0.132467	0.060231	0.641148	373	0.208801	0.131969	0.020157	0.685757

When you specify a STRATA statement, the “Strata Information” table in [Output 95.2.3](#) displays the minimum and maximum propensity scores, the total number of treated observations, and the total number of controls in each stratum.

Output 95.2.3 Strata Information

Strata Information						
			Frequencies			
Stratum Index	Propensity Score Range		Treated	Control	Total	
1	0.020157	0.194358	22	209	231	
2	0.196742	0.261300	23	59	82	
3	0.261861	0.322300	23	38	61	
4	0.325937	0.434208	23	41	64	
5	0.437927	0.685757	22	26	48	

The ASSESS statement produces the tables and plots that summarize differences in the specified variables between treated and control groups for all observations, for the support region observations, and for the matched observations. As requested by the PS and VAR= options, the variables listed in the table are the propensity score and the variables Gender, Age, and Bmi. When you specify a STRATA statement, WEIGHT=NONE by default, suppressing display of differences for the weighted observations.

The VARINFO option displays the “Variable Information” table, which contains variable differences between the treated and control groups for all observations and for the support region observations, as shown in [Output 95.2.4](#). For a binary classification variable (Gender), the difference is in the proportion of the first ordered level (Female).

Output 95.2.4 Variable Information**The PSMATCH Procedure**

Variable Information											
		Treated (Drug = Drug_X)						Control (Drug = Drug_A)			
Variable	Observations	N	Mean	Std Dev	Minimum	Maximum	N	Mean	Std Dev	Minimum	Maximum
PS	All	113	0.310773	0.132467	0.060231	0.641148	373	0.208801	0.131969	0.020157	0.685757
	Region	113	0.310773	0.132467	0.060231	0.641148	373	0.208801	0.131969	0.020157	0.685757
Age	All	113	36.309735	5.534114	26.000000	49.000000	373	40.404826	6.579103	25.000000	57.000000
	Region	113	36.309735	5.534114	26.000000	49.000000	373	40.404826	6.579103	25.000000	57.000000
Bmi	All	113	24.492566	1.863797	20.330000	28.340000	373	23.753271	1.980778	19.220000	28.610000
	Region	113	24.492566	1.863797	20.330000	28.340000	373	23.753271	1.980778	19.220000	28.610000
Gender	All	113	0.433628	0.495575			373	0.458445	0.498270		
	Region	113	0.433628	0.495575			373	0.458445	0.498270		

When REGION=ALLOBS, the statistics are identical between all observations and the support region observations.

The “Standardized Variable Differences” table in [Output 95.2.5](#) displays the standardized differences between the treated and control groups for all observations and for the support region observations,

Output 95.2.5 Standardized Differences

Standardized Variable Differences (Treated - Control)								
Standardized Mean Difference								
Mean Difference			Mean Difference			Percent	Variance Ratio	
Variable	All Obs	Region Obs	Divisor	All Obs	Region Obs	Reduction	Region	
						Region Obs	All Obs	Region Obs
PS	0.101972	0.101972	0.132218	0.771242	0.771242	0.00	1.0076	1.0076
Age	-4.095091	-4.095091	6.079104	-0.673634	-0.673634	0.00	0.7076	0.7076
Bmi	0.739296	0.739296	1.923178	0.384414	0.384414	0.00	0.8854	0.8854
Gender	-0.024817	-0.024817	0.496925	-0.049941	-0.049941	0.00	0.9892	0.9892

When ODS Graphics is enabled, the PLOTS option displays plots for the specified variables. The plots without strata information are not shown here because the REGION=ALLOBS option results in all observations being in the support region.

When you specify a STRATA statement, the ASSESS statement also produces tables and plots that summarize differences in the specified variables between treated and control groups by stratum.

The VARINFO option in the ASSESS statement displays the variable information of the treated and control groups for observations in each stratum, as shown in [Output 95.2.6](#).

Output 95.2.6 Strata Variable Information

The PSMATCH Procedure

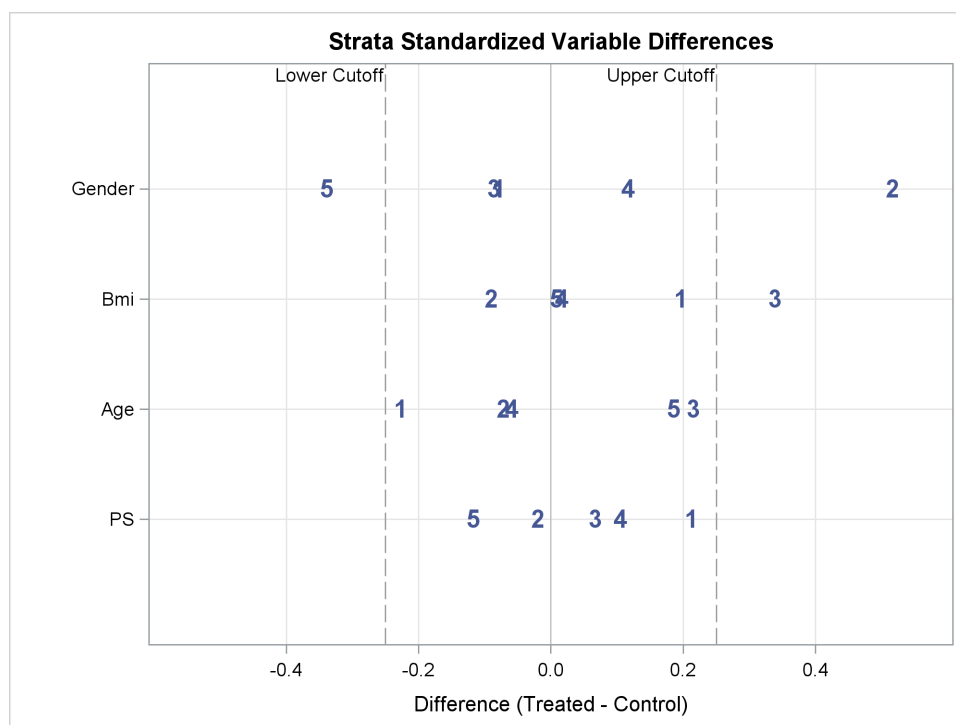
Strata Variable Information											
Region Observations											
Treated (Drug = Drug_X)							Control (Drug = Drug_A)				
Variable	Stratum Index	N	Mean	Std Dev	Minimum	Maximum	N	Mean	Std Dev	Minimum	Maximum
PS	1	22	0.140401	0.041360	0.060231	0.194358	209	0.115891	0.043859	0.020157	0.194132
	2	23	0.221993	0.019418	0.196742	0.259357	59	0.228211	0.018395	0.197342	0.261300
	3	23	0.299861	0.018811	0.263497	0.322300	38	0.294568	0.017541	0.261861	0.321559
	4	23	0.380869	0.026077	0.326680	0.434184	41	0.370552	0.030646	0.325937	0.434208
	5	22	0.512086	0.058200	0.437927	0.641148	26	0.531182	0.071893	0.441204	0.685757
Age	1	22	42.909091	4.150836	35.000000	49.000000	209	44.449761	4.811706	32.000000	57.000000
	2	23	37.652174	3.938259	29.000000	44.000000	59	38.254237	4.305543	29.000000	47.000000
	3	23	36.565217	3.870942	29.000000	43.000000	38	35.421053	3.636389	27.000000	43.000000
	4	23	33.695652	4.247296	26.000000	41.000000	41	34.219512	3.086359	29.000000	41.000000
	5	22	30.772727	2.827279	27.000000	35.000000	26	29.807692	2.939649	25.000000	37.000000
Bmi	1	22	23.500909	1.751203	20.330000	26.110000	209	23.175407	1.917237	19.240000	27.850000
	2	23	23.653043	1.794401	20.430000	26.660000	59	23.878475	1.951062	19.220000	27.680000
	3	23	24.707826	1.764444	20.850000	27.560000	38	24.108158	1.698325	20.240000	27.600000
	4	23	24.915217	1.950177	20.980000	28.340000	41	24.935854	1.484916	22.370000	28.290000
	5	22	25.695000	1.130338	23.320000	28.060000	26	25.730769	1.337953	23.410000	28.610000
Gender	1	22	0.454545	0.497930			209	0.507177	0.499948		
	2	23	0.565217	0.495728			59	0.322034	0.467256		
	3	23	0.391304	0.488042			38	0.447368	0.497222		
	4	23	0.434783	0.495728			41	0.390244	0.487805		
	5	22	0.318182	0.465770			26	0.500000	0.500000		

The “Strata Standardized Variable Differences” table in [Output 95.2.7](#) displays the variable differences, standardized differences, reduction percentages, and ratios of variances for observations in each stratum. The standardized difference is the variable difference divided by the divisor (which is displayed in the “Standardized Variable Differences” table in [Output 95.2.5](#)), and the reduction percentage compares the standardized difference with the standardized difference of all observations.

Output 95.2.7 Strata Standardized Differences

Strata Standardized Variable Differences (Treated - Control) Region Observations					
Variable	Stratum Index	Mean Difference	Standardized Difference	Percent Reduction	Variance Ratio
PS	1	0.024509	0.185370	75.96	0.889269
	2	-0.006219	-0.047033	93.90	1.114307
	3	0.005293	0.040035	94.81	1.149989
	4	0.010317	0.078030	89.88	0.724072
	5	-0.019096	-0.144430	81.27	0.655343
Age	1	-1.540670	-0.253437	62.38	0.744171
	2	-0.602063	-0.099038	85.30	0.836667
	3	1.144165	0.188213	72.06	1.133163
	4	-0.523860	-0.086174	87.21	1.893792
	5	0.965035	0.158746	76.43	0.925010
Bmi	1	0.325502	0.169252	55.97	0.834299
	2	-0.225431	-0.117218	69.51	0.845857
	3	0.599668	0.311811	18.89	1.079380
	4	-0.020636	-0.010730	97.21	1.724822
	5	-0.035769	-0.018599	95.16	0.713731
Gender	1	-0.052632	-0.105915	0.00	0.991940
	2	0.243183	0.489377	0.00	1.125585
	3	-0.056064	-0.112822	0.00	0.963416
	4	0.044539	0.089629	0.00	1.032750
	5	-0.181818	-0.365887	0.00	0.867769

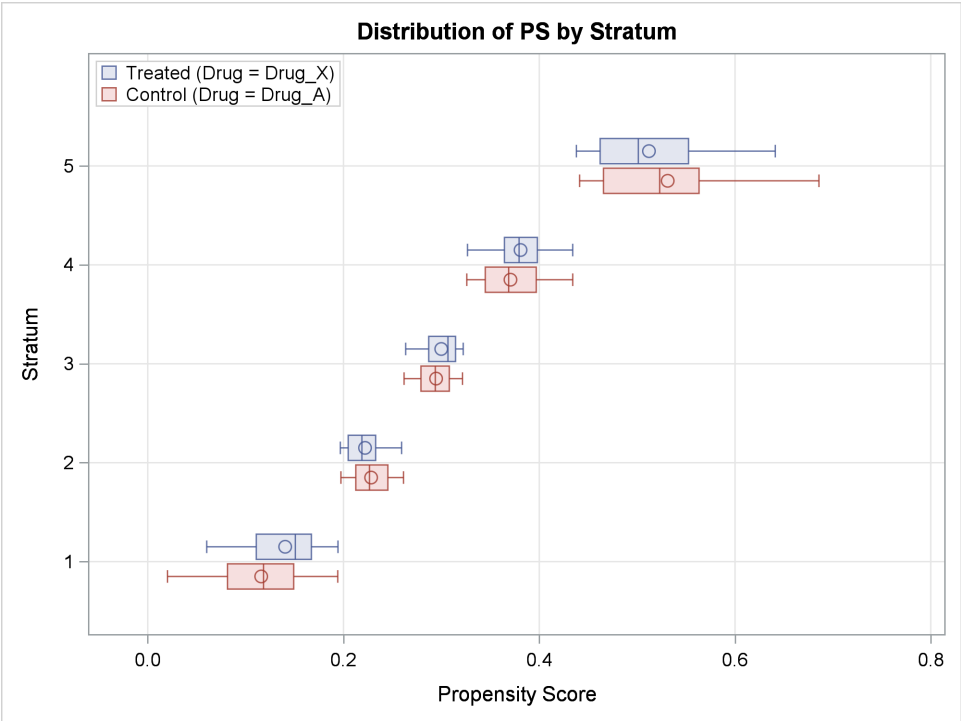
When ODS Graphics is enabled, the strata standardized differences for the specified variables are displayed, as shown in [Output 95.2.8](#).

Output 95.2.8 Strata Standardized Differences Plot

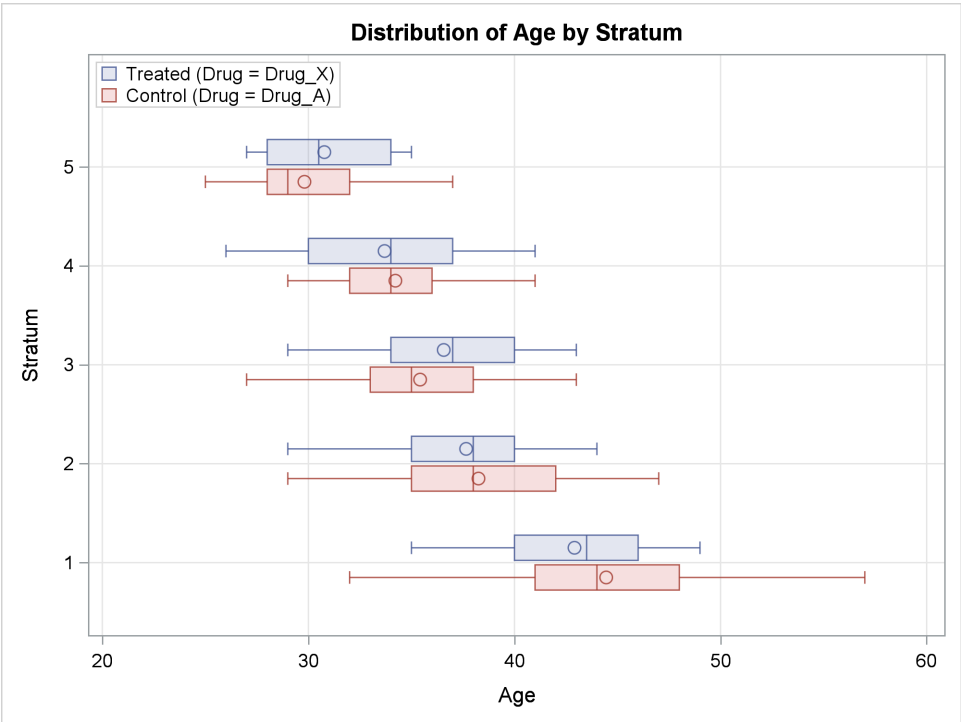
The “Strata Standardized Variable Differences Plot” displays the standardized differences in the “Strata Standardized Variable Differences” table in [Output 95.2.7](#). The plot shows larger differences in Stratum 2 and Stratum 5 for Gender.

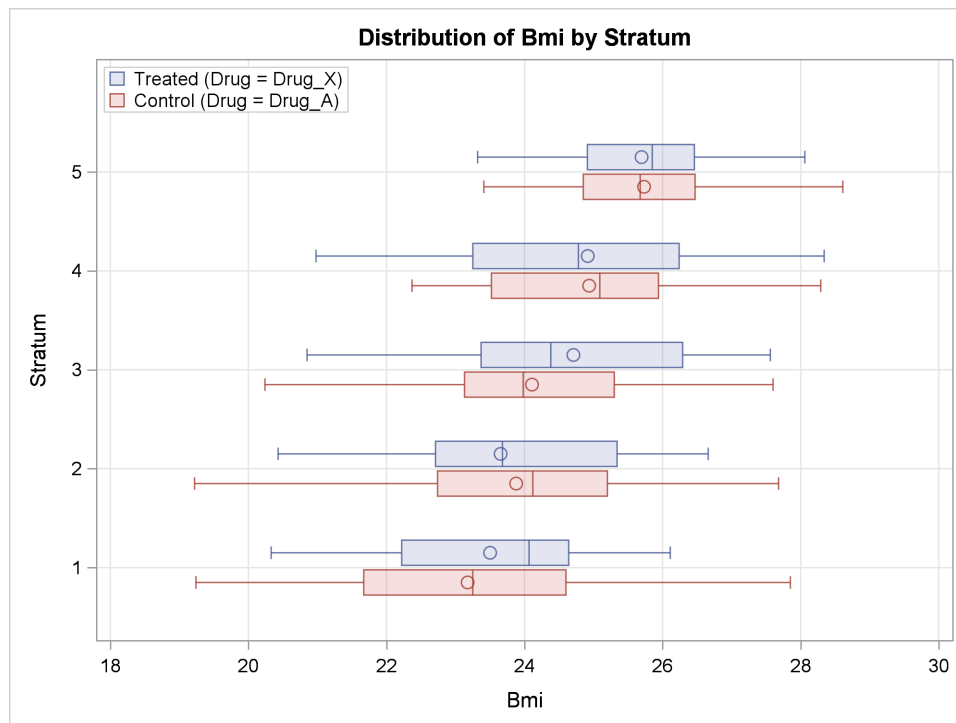
The PLOTS=BOXPLOT option displays a strata variable box plot for each continuous variable, as shown in [Output 95.2.9](#) for PS, [Output 95.2.10](#) for Age, and [Output 95.2.11](#) for Bmi. The box plots show reasonably good variable balance in each stratum.

Output 95.2.9 PS Strata Box Plot

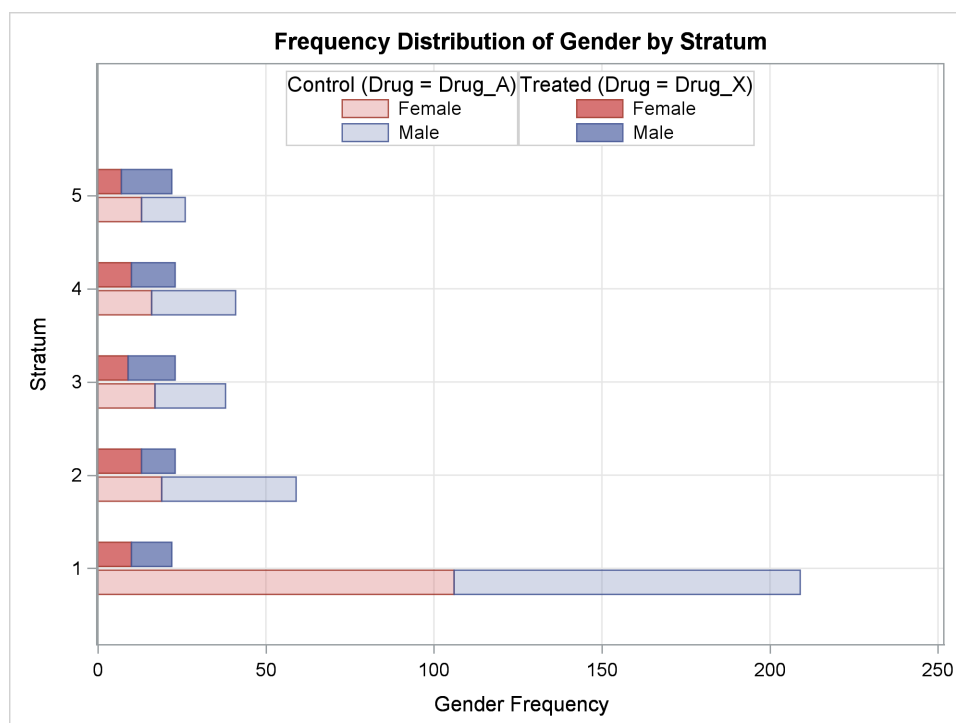


Output 95.2.10 Age Strata Box Plot



Output 95.2.11 Bmi Strata Box Plot

PLOTS=BARCHART displays a strata variable bar chart for each classification variable, as shown in [Output 95.2.12](#) for Gender. The bar chart shows differences in the female and male relative frequencies between the treated and control groups in Stratum 2 and Stratum 5.

Output 95.2.12 Gender Strata Bar Chart

If you are not satisfied with the variable balance, you can do one or more of the following until you are satisfied: you can select another set of variables to fit the propensity score model, you can modify the specification of the propensity score model (for instance, by using nonlinear terms for the continuous variables or by adding interactions), you can increase the number of strata, or you can choose another propensity score method (such as matching).

The `OUT(OBS=ALL)=OutEx2` option in the `OUTPUT` statement creates an output data set, `OutEx2`, that contains all observations. The following statements list the first 10 observations in `OutEx2`, as shown in [Output 95.2.13](#):

```
proc print data=OutEx2(obs=10);
  var PatientID Drug Gender Age Bmi _ps_ _Strata_;
run;
```

Output 95.2.13 Output Data Set with Strata

Obs	PatientID	Drug	Gender	Age	Bmi	_PS_	_STRATA_
1	284	Drug_X	Male	29	22.02	0.36444	4
2	201	Drug_A	Male	45	26.68	0.22296	2
3	147	Drug_A	Male	42	21.84	0.11323	1
4	307	Drug_X	Male	38	22.71	0.19733	2
5	433	Drug_A	Male	31	22.76	0.35311	4
6	435	Drug_A	Male	43	26.86	0.27263	3
7	159	Drug_A	Female	45	25.47	0.14911	1
8	368	Drug_A	Female	49	24.28	0.07780	1
9	286	Drug_A	Male	31	23.31	0.38341	4
10	163	Drug_X	Female	39	25.34	0.24995	2

By default, the output data set includes the variable `_PS_` (which provides the propensity score) and the variable `_STRATA_` (which provides stratum indices).

If you assume that no other confounding variables are associated with both the response variable and the treatment group indicator `Drug`, then after the responses for the trial are observed and added to the data set `OutEx2`, you can estimate the treatment effect within each stratum and combine these estimates across strata to estimate the treatment effect (Stuart 2010, pp. 13–14).

Example 95.3: Optimal Variable Ratio Matching

This example performs optimal matching of variable numbers of patients in the control group with each patient in the treatment group in a propensity score analysis. The `Drugs` data set contains the patient information and is described in the section “[Getting Started: PSMATCH Procedure](#)” on page 7680.

The following statements invoke the `PSMATCH` procedure and request optimal variable ratio matching to match each observation for patients in the treatment group with a variable number of observations for patients in the control group:

```
ods graphics on;
proc psmatch data=drugs region=treated(extend(stat=ps mult=one)=0.025);
  class Drug Gender;
  psmodel Drug(Treated='Drug_X')= Gender Age Bmi;
```

```

match stat=ps method=varratio(kmin=1 kmax=3) exact=(Gender) caliper=.;
assess ps var=(Gender Age Bmi)
      / plots(orient=vertical);
output out (obs=match)=OutEx3 matchid=_MatchID;
run;

```

The CLASS statement specifies the classification variables. The PSMODEL statement specifies the logistic regression model that creates the propensity score for each observation, which is the probability that the patient receives Drug_X. The Drug variable is the binary treatment indicator variable, and TREATED='Drug_X' identifies Drug_X as the treated group. The Gender, Age, and Bmi variables are included in the model because they are believed to be related to the assignment.

The REGION= option specifies an interval region of propensity scores (or equivalently, logits of propensity scores) such that only observations that have propensity scores in the region are used in stratification and matching. Because the MATCH statement is also specified, the REGION=TREATED(EXTEND(STAT=PS MULT=ONE)=0.025) option requests that only observations that have propensity scores in the region defined by the treated observations be used for matching. The EXTEND(STAT=PS MULT=ONE)=0.025 option requests that the region be extended by the specified 0.025 in propensity score.

The MATCH statement specifies the criteria for matching. The STAT=PS option requests that the propensity score be used in computing differences between pairs of observations. The METHOD=VARRATIO(KMIN=1 KMAX=3) option requests optimal variable ratio matching of one to three control units to each unit in the treated group in order to minimize the total absolute difference in propensity score across all matches.

The default average number of control units to each treated unit is computed as the mean of the KMIN= and KMAX= values, so an average of two control units are matched to each treated unit. The EXACT=GENDER option forces the treated unit and its matched control unit to have the same value of Gender. The CALIPER=. option ignores the caliper requirement for matching.

The “Data Information” table in [Output 95.3.1](#) displays information about the input and output data sets, the numbers of observations in the treated and control groups, the lower and upper limits for the propensity score support region, and the numbers of observations in the treated and control groups that fall within the support region. Of the 373 observations in the control group, 366 fall within the support region.

Output 95.3.1 Data Information
The PSMATCH Procedure

Data Information	
Data Set	WORK.DRUGS
Output Data Set	WORK.OUTEX3
Treatment Variable	Drug
Treatment Group	Drug_X
All Obs (Treatment)	113
All Obs (Control)	373
Support Region	Extended Treatment Group
Lower PS Support	0.035231
Upper PS Support	0.666148
Support Region Obs (Treatment)	113
Support Region Obs (Control)	366

The “Propensity Score Information” table in [Output 95.3.2](#) displays summary statistics by treatment group for all observations, for the support region observations, and for the matched observations.

Output 95.3.2 Propensity Score Information

Propensity Score Information										
Observations	N	Treated (Drug = Drug_X)				N	Control (Drug = Drug_A)			
		Mean	Std Dev	Minimum	Maximum		Mean	Std Dev	Minimum	Maximum
All	113	0.310773	0.132467	0.060231	0.641148	373	0.208801	0.131969	0.020157	0.685757
Region	113	0.310773	0.132467	0.060231	0.641148	366	0.208677	0.126739	0.037141	0.635131
Matched	113	0.310773	0.132467	0.060231	0.641148	226	0.266700	0.121379	0.056086	0.635131

The “Matching Information” table in [Output 95.3.3](#) displays the matching criteria, the number of matched sets, the numbers of matched observations in the treated and control groups, and the total absolute difference in the propensity score for all matches. Note that with an average of two control units to each treated unit, 226 control units are matched.

Output 95.3.3 Matching Information

Matching Information	
Difference Statistic	Propensity Score
Method	Optimal Variable Ratio Matching
Min Control/Treated Ratio	1
Max Control/Treated Ratio	3
Matched Sets	113
Matched Obs (Treated)	113
Matched Obs (Control)	226
Total Absolute Difference	1.436979

The ASSESS statement produces the tables and plots that summarize differences in the specified variables between treated and control groups for all observations, for the support region observations, and for the matched observations. As requested by the PS and VAR= options, the variables listed in the table are the logit of propensity score and the variables Gender, Age, and Bmi. By default (or if WEIGHT=MATCHWGT is specified), each treated unit receives a weight of 1 and each control unit receives a weight that is computed as the number of treated units divided by the number of control units in the matched set. That is, if three control units are matched to a treated unit in a matched set, then each control unit receives a weight of 1/3.

The “Standardized Variable Differences” table, as shown in [Output 95.3.4](#), displays standardized differences between the treated and control groups for all observations, for the support region observations, and for the matched observations. For a binary classification variable (Gender), the difference is in the proportion of the first ordered level (Female).

Output 95.3.4 Standardized Differences**The PSMATCH Procedure****Standardized Variable Differences (Treated - Control)**

Variable	Mean Difference			Weighted Matched Obs
	All Obs	Region Obs	Matched Obs	
PS	0.101972	0.102096	0.044073	0.005141
Age	-4.095091	-3.982615	-1.504425	-0.131268
Bmi	0.739296	0.728714	0.235531	0.029572
Gender	-0.024817	-0.022656	-0.004425	0

Standardized Variable Differences (Treated - Control)**Standardized Mean Difference**

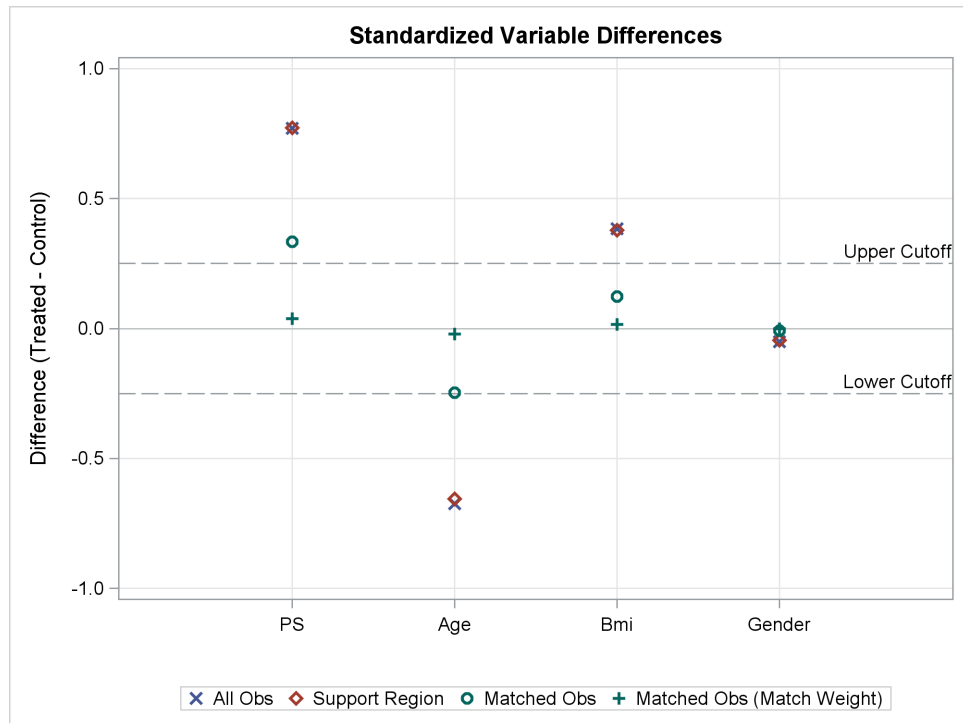
Variable	Divisor	Mean Difference			Percent Reduction			
		All Obs	Region Obs	Matched Obs	Weighted Matched Obs	Region Obs	Matched Obs	Weighted Matched Obs
PS	0.132218	0.771242	0.772181	0.333333	0.038886	0.00	56.78	94.96
Age	6.079104	-0.673634	-0.655132	-0.247475	-0.021593	2.75	63.26	96.79
Bmi	1.923178	0.384414	0.378911	0.122470	0.015377	1.43	68.14	96.00
Gender	0.496925	-0.049941	-0.045592	-0.008904	0	8.71	82.17	100.00

**Standardized Variable Differences
(Treated - Control)****Variance Ratio**

Variable	All Obs	Region		Weighted Matched Obs
		Obs	Matched Obs	
PS	1.0076	1.0924	1.1910	1.0773
Age	0.7076	0.7393	0.8812	0.9654
Bmi	0.8854	0.9227	1.0249	1.1025
Gender	0.9892	0.9899	0.9977	1.0000

The standardized mean differences are significantly reduced in the matched observations, the standardized differences are less than the recommended upper limit of 0.25, and the variance ratios between the two groups are within the recommended range of 0.5 to 2.

When ODS Graphics is enabled, the PSMATCH procedure displays a standardized differences plot for the variables that are specified in the ASSESS statement, as shown in [Output 95.3.5](#).

Output 95.3.5 Standardized Differences Plot

When you specify the `ORIENT=VERTICAL` option, the standardized differences are placed on the vertical axis. The “Standardized Variable Differences Plot” displays the standardized differences in the “Variable Differences” table in [Output 95.3.4](#). All differences for the matched observations are within the recommended limits of -0.25 and 0.25 , which are indicated by reference lines.

If you are not satisfied with the variable balance, you can do one or more of the following until you are satisfied: you can select another set of variables to fit the propensity score model, you can modify the matching criteria, or you can choose another matching method.

The `OUT(OBS=MATCH)=OutEx3` option in the `OUTPUT` statement creates an output data set, `OutEx3`, that contains the matched observations. The following statements list the 10 observations that have lowest propensity scores, as shown in [Output 95.3.6](#):

```
proc sort data=OutEx3 out=OutEx3a;
  by _PS_;
run;

proc print data=OutEx3a(obs=10);
  var PatientID Drug Gender Age Bmi _PS_ _MATCHWGT_ _MatchID;
run;
```

Output 95.3.6 Output Data Set with Optimal Variable Ratio Matches

Obs	PatientID	Drug	Gender	Age	Bmi	_PS_	_MATCHWGT_	_MatchID
1	311	Drug_A	Female	49	22.80	0.056086	0.33333	1
2	89	Drug_X	Female	44	20.75	0.060231	1.00000	1
3	213	Drug_A	Female	49	23.24	0.061866	0.33333	1
4	141	Drug_A	Female	43	20.55	0.064010	0.33333	1
5	323	Drug_X	Female	46	22.22	0.067625	1.00000	2
6	245	Drug_A	Female	52	25.32	0.071559	0.33333	2
7	137	Drug_A	Female	45	22.04	0.072150	0.33333	2
8	40	Drug_A	Female	42	20.65	0.072655	0.33333	2
9	341	Drug_A	Male	55	26.76	0.086895	0.33333	3
10	269	Drug_A	Female	48	24.35	0.087566	0.33333	4

By default, the output data set includes the variable `_PS_` (which provides the propensity score) and the variable `_MATCHWGT_` (which provides matched observation weights). The weight for each treated unit is 1. Because `METHOD=VARRATIO(KMIN=1 KMAX=3)` is specified in the `MATCH` statement, one, two, or three control units are matched to each treated unit; so the weight for each matched control unit is 1, 1/2, or 1/3. The `MATCHID=_MatchID` option creates a variable named `_MatchID` that identifies the matched sets of observations.

If you assume that no other confounding variables are associated with both the response variable and the treatment group indicator `Drug`, then after the responses for the trial are observed and added to the data set `OutEx3`, you can use the same outcome analysis on this output data set as you would have used on the original data set `Drugs` (augmented with responses) to estimate the treatment effect.

Example 95.4: Greedy Nearest Neighbor Matching

This example performs greedy matching in a propensity score analysis.

At the completion of a school year, a school administrator asks whether taking a music class causes an improvement in the grade point averages (GPAs) of students. The reasoning behind this question is that learning to read and perform music might improve general reading ability, concentration, and memory.

The data set `School` contains information about students that is available at the end of the school year. `StudentID` is the student identification number, `Music` indicates whether the student took a music class, `Gender` provides the gender of the student, and `Absence` is the percentage of absences. For simplicity, this example uses only three covariates (`Music`, `Gender`, and `Absence`), but in practice a propensity score analysis often involves many more covariates.

[Output 95.4.1](#) lists the first 10 observations.

Output 95.4.1 Input School Data Set**First 10 Observations of the Input Music Data Set**

Obs	StudentID	Music	Gender	Absence
1	18	No	Female	3.71200
2	61	No	Male	2.07552
3	95	No	Female	2.53865
4	41	No	Male	3.00637
5	19	Yes	Female	0.08081
6	51	No	Female	1.20229
7	110	No	Male	2.20710
8	87	No	Female	2.30150
9	103	No	Female	3.08102
10	175	No	Female	1.12169

In this example, the outcome data (GPAs) for the students happen to be available at the time of the propensity score analysis, but the recommended practice is not to use the outcome values in the propensity score analysis (Stuart 2010, p. 2). Instead, the response variable is added to the output data set created by the PSMATCH procedure; that output data set consists of matched observations that are subsequently used in an outcome analysis.

The following statements invoke the PSMATCH procedure and request greedy nearest neighbor matching to sequentially match each observation for students in the treatment group (those who took music) with one observation for students in the control group (those who did not take music):

```
ods graphics on;
proc psmatch data=School region=treated;
  class Music Gender;
  psmodel Music(Treated='Yes')= Gender Absence;
  match method=greedy(k=1) exact=Gender caliper=0.5;
  assess lps var=(Gender Absence) / plots=all weight=none;
  output out(obs=match)=OutEx4 matchid=_MatchID;
run;
```

The CLASS statement specifies the classification variables. The PSMODEL statement specifies the logistic regression model that creates the propensity score for each student, which is the probability that the student enrolled in the music class. The Music variable is the binary treatment indicator variable and TREATED='Yes' identifies Yes as the treated group. The Gender and Absence variables are included in the model because they are believed to be related to enrolling in the music class.

The REGION= option specifies an interval region of propensity scores such that only observations that have propensity scores in the region are used in stratification and matching. Because the MATCH statement is also specified, the REGION=TREATED option requests that only observations whose propensity scores lie in the range that corresponds to observations in the treated group be used for matching. By default, the region is extended by 0.25 times the pooled estimate of the common standard deviation of the logit of the propensity score statistic.

The MATCH statement requests matching and specifies the criteria for matching. The STAT=LPS option (which is the default) requests that the logit of the propensity score be used in computing differences between pairs of observations. The METHOD=GREEDY(K=1) option requests greedy nearest neighbor matching in which one control unit is matched with each unit in the treated group such that the matching produces

the smallest within-pair difference among all available pairs with this treated unit. The EXACT=GENDER option forces the treated unit and its matched control unit to have the same value of the Gender variable. The CALIPER=0.5 option specifies the caliper requirement for matching: units are matched only if the difference in the logits of the propensity score for pairs of units from the two groups is less than or equal to 0.5 times the pooled estimate of the common standard deviation of the logits of the propensity scores.

The “Data Information” table, which is produced by the PSMATCH procedure and shown in [Output 95.4.2](#), displays information about the input and output data sets, the numbers of observations in the treated and control groups, the lower and upper limits for the propensity score support region, and the numbers of observations in the treated and control groups that fall within the support region. Of the 140 observations in the control group, 132 fall within the support region.

Output 95.4.2 Data Information
The PSMATCH Procedure

Data Information	
Data Set	WORK.SCHOOL
Output Data Set	WORK.OUTEX4
Treatment Variable	Music
Treatment Group	Yes
All Obs (Treatment)	60
All Obs (Control)	140
Support Region	Extended Treatment Group
Lower PS Support	0.079975
Upper PS Support	0.530932
Support Region Obs (Treatment)	60
Support Region Obs (Control)	132

The “Propensity Score Information” table in [Output 95.4.3](#) displays summary statistics for the treatment and control groups, which are computed for all observations, support region observations, and matched observations.

Output 95.4.3 Propensity Score Information

Propensity Score Information										
Treated (Music = Yes)						Control (Music = No)				
Observations	N	Mean	Std Dev	Minimum	Maximum	N	Mean	Std Dev	Minimum	Maximum
All	60	0.347143	0.096184	0.092831	0.490191	140	0.279796	0.124997	0.026465	0.488875
Region	60	0.347143	0.096184	0.092831	0.490191	132	0.294024	0.113997	0.083346	0.488875
Matched	60	0.347143	0.096184	0.092831	0.490191	60	0.340188	0.098482	0.092963	0.488875

The “Matching Information” table in [Output 95.4.4](#) displays the matching criteria, the number of matched sets, the numbers of matched observations in the treated and control groups, and the total absolute difference in the logit of the propensity score for all matches.

Output 95.4.4 Matching Information

Matching Information	
Difference Statistic	Logit of Propensity Score
Method	Greedy Matching
Control/Treated Ratio	1
Order	Descending
Caliper (Logit PS)	0.326257
Matched Sets	60
Matched Obs (Treated)	60
Matched Obs (Control)	60
Total Absolute Difference	2.943871

The ASSESS statement produces tables and plots that summarize differences in the specified variables between treated and control groups for all observations, for the support region observations, and for the matched observations. You can use these results to assess how well matching achieves a balance in the distributions of these variables. As requested by the LPS and VAR= options, the variables are the logit of propensity score and the covariates Gender and Absence. The WEIGHT=NONE option suppresses the display of differences for weighted matched observations. For a matching of one control unit to each treated unit, the weights are all 1 for matched treated and control units, and the results are identical for the weighted matched observations and the matched observations.

The “Standardized Variable Differences” table in [Output 95.4.5](#) displays standardized differences between the treated and control groups, which are computed for all observations, support region observations, and matched observations. For the binary classification variable (Gender), the computed difference is in the proportion of the first ordered level (Female).

Output 95.4.5 Standardized Differences**The PSMATCH Procedure**

Standardized Variable Differences (Treated - Control)									
				Standardized Mean Difference				Percent Reduction	
		Mean Difference		Divisor			Mean Difference		
Variable	All Obs	Region Obs	Matched Obs		All Obs	Region Obs	Matched Obs	Region Obs	Matched Obs
LPS	0.406809	0.284163	0.032679	0.652514	0.623449	0.435489	0.050081	30.15	91.97
Absence	-0.697568	-0.485721	-0.057759	1.136945	-0.613546	-0.427216	-0.050802	30.37	91.72
Gender	-0.045238	-0.034848	0	0.496344	-0.091143	-0.070210	0	22.97	100.00

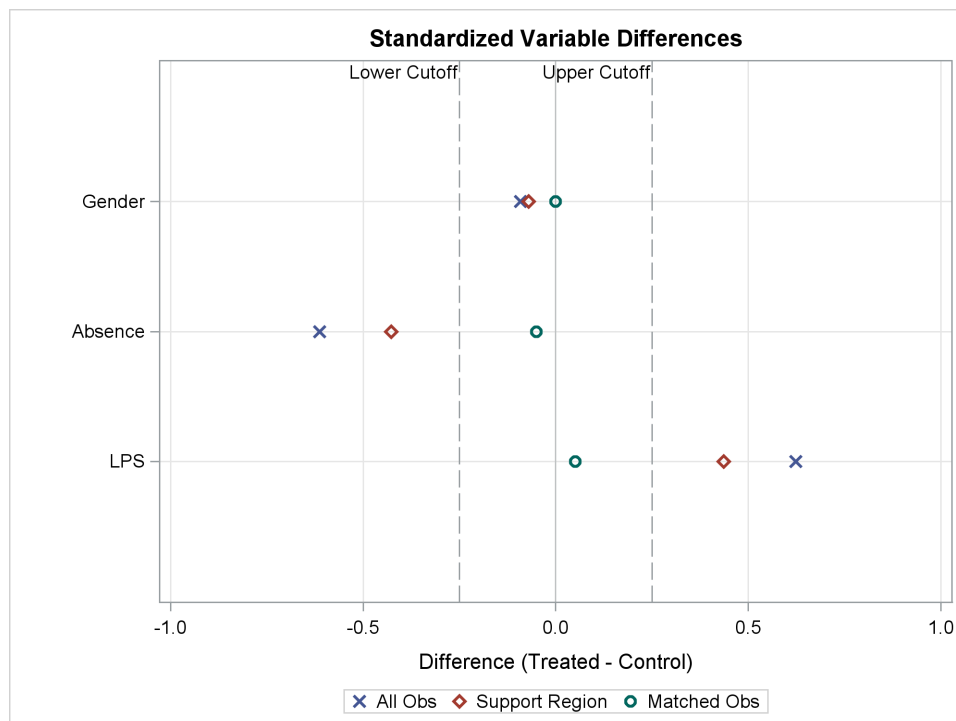
Standardized Variable Differences
(Treated - Control)

Variance Ratio			
		Region Obs	Matched Obs
Variable	All Obs	Obs	Obs
LPS	0.3810	0.6135	0.9696
Absence	0.3550	0.5560	0.9375
Gender	1.0208	1.0144	1.0000

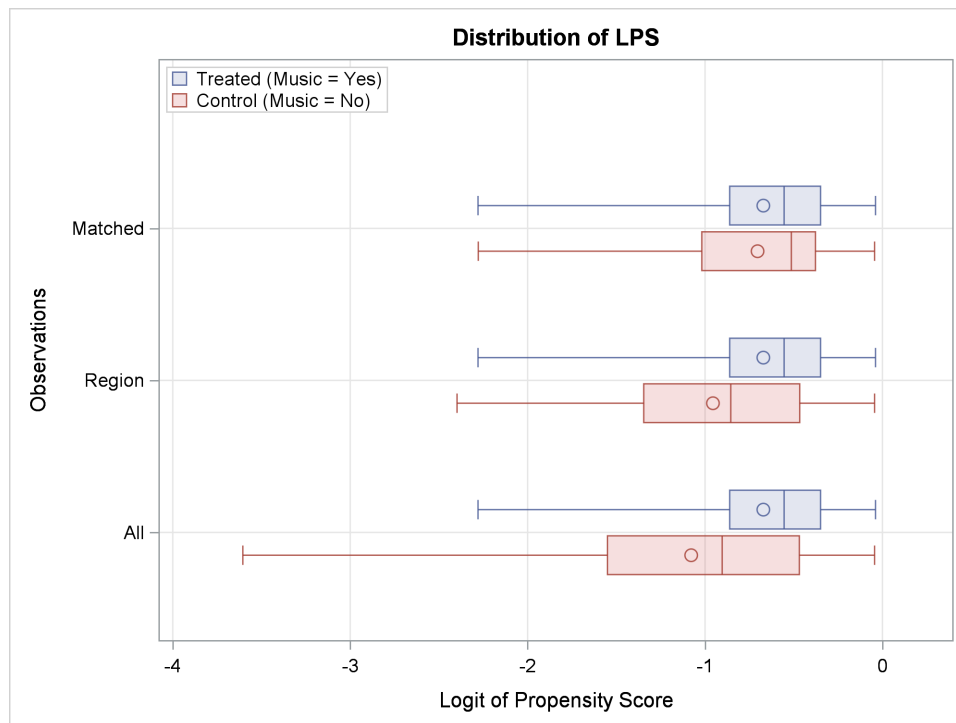
The standardized mean differences are significantly reduced in the matched observations, and the largest of these differences is 0.0508 in absolute value, which is less than the recommended upper limit of 0.25 (Rubin 2001, p. 174; Stuart 2010, p. 11). The variance ratios between the two groups are between 0.9375 and 1 for all variables in the matched observations, which is within the recommended range of 0.5 to 2. Because both EXACT=GENDER and METHOD=GREEDY are specified in the MATCH statement, the standardized difference for Gender is 0 in the matched observations.

When ODS Graphics is enabled and you specify PLOTS=ALL, the PSMATCH procedure uses ODS Graphics to create all applicable plots. [Output 95.4.6](#) displays a plot of the standardized differences in Gender, Absence, and the logit propensity score for all observations, observations in the support region, and matched observations.

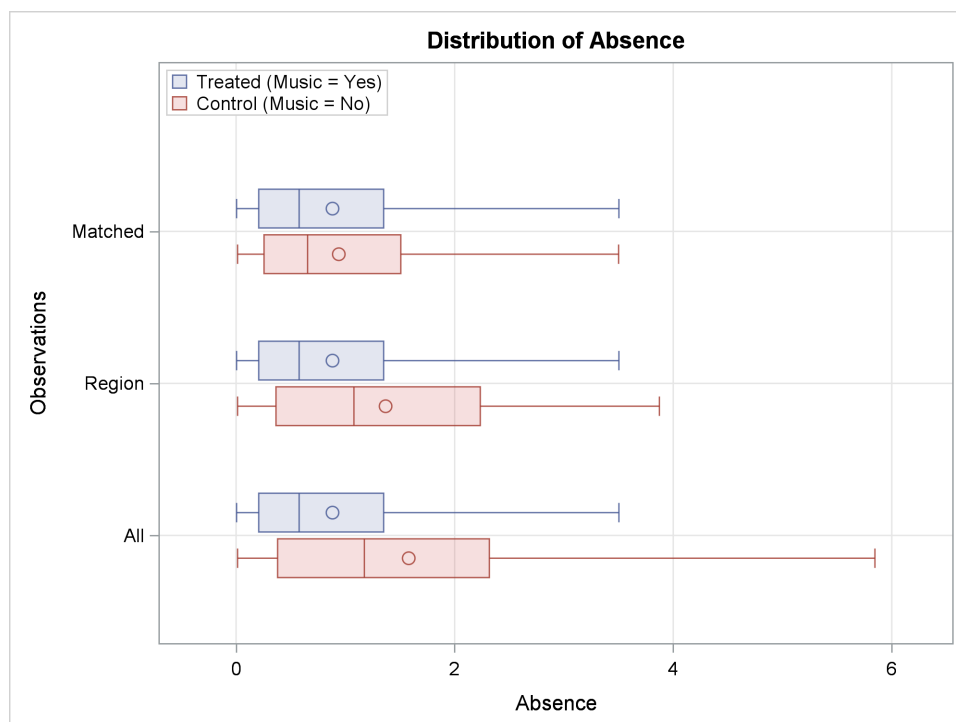
Output 95.4.6 Standardized Differences Plot



[Output 95.4.7](#) displays box plots that compare the distributions of the logit propensity score for units in the treated and control groups, based on all observations, observations in the support region, and matched observations. Note that the two distributions are well-balanced for matched observations.

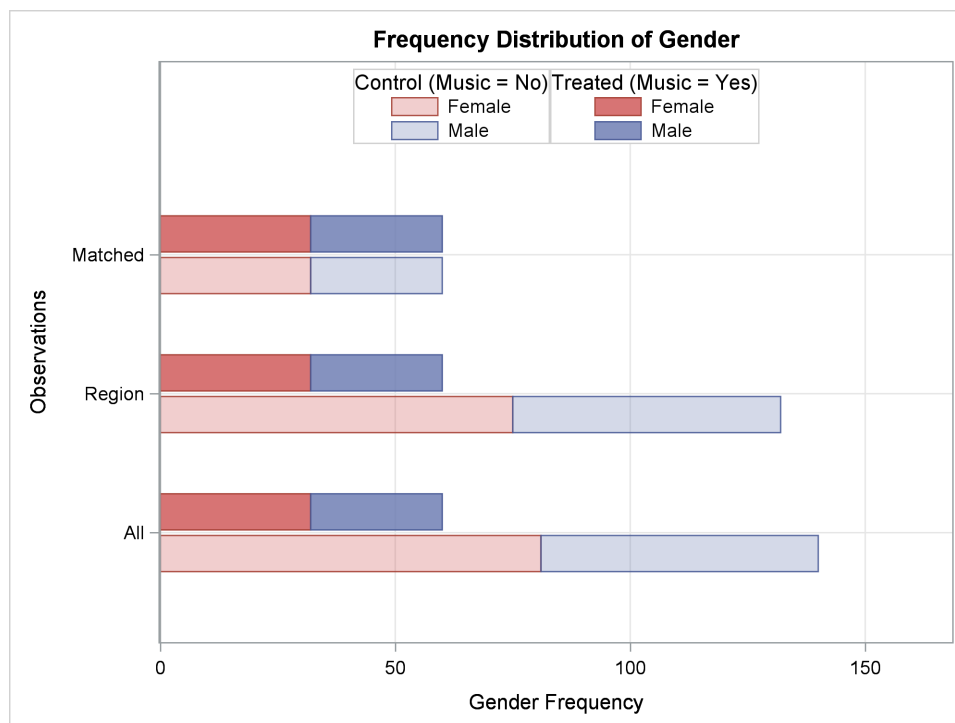
Output 95.4.7 LPS Box Plot

Output 95.4.8 displays box plots that compare the distributions of Absence for units in the treated and control groups, based on all observations, observations in the support region, and matched observations. Again, note that the two distributions are well-balanced for matched observations.

Output 95.4.8 Absence Box Plot

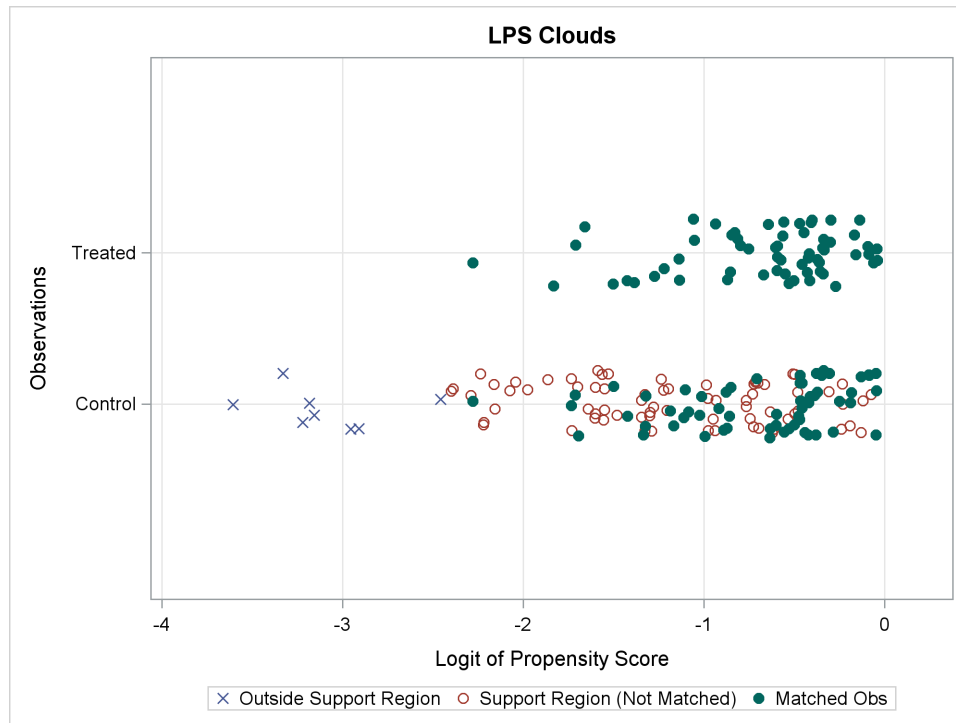
Output 95.4.9 displays bar charts that compare the distributions of Gender for units in the treated and control groups, based on all observations, observations in the support region, and matched observations. Again, note that the two distributions are well-balanced for matched observations.

Output 95.4.9 Gender Bar Chart



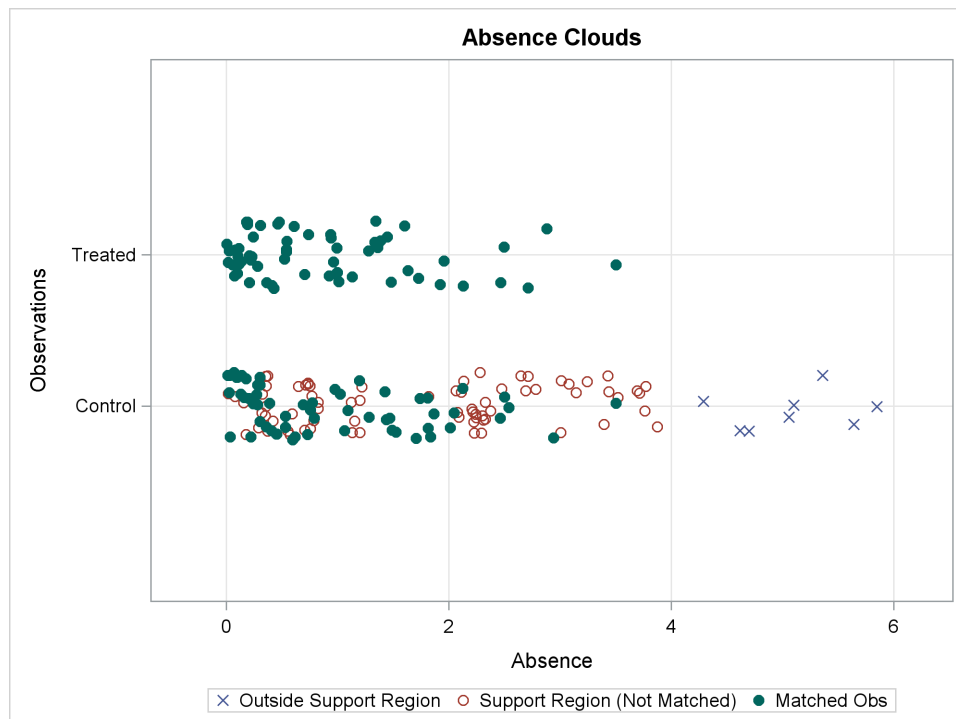
Output 95.4.10 displays a cloud plot that compares the values of the logit propensity score LPS for observations in the treated and control groups, based on all observations, observations in the support region, and matched observations. The points are jittered in the vertical direction to avoid overlap.

Output 95.4.10 LPS Cloud Plot



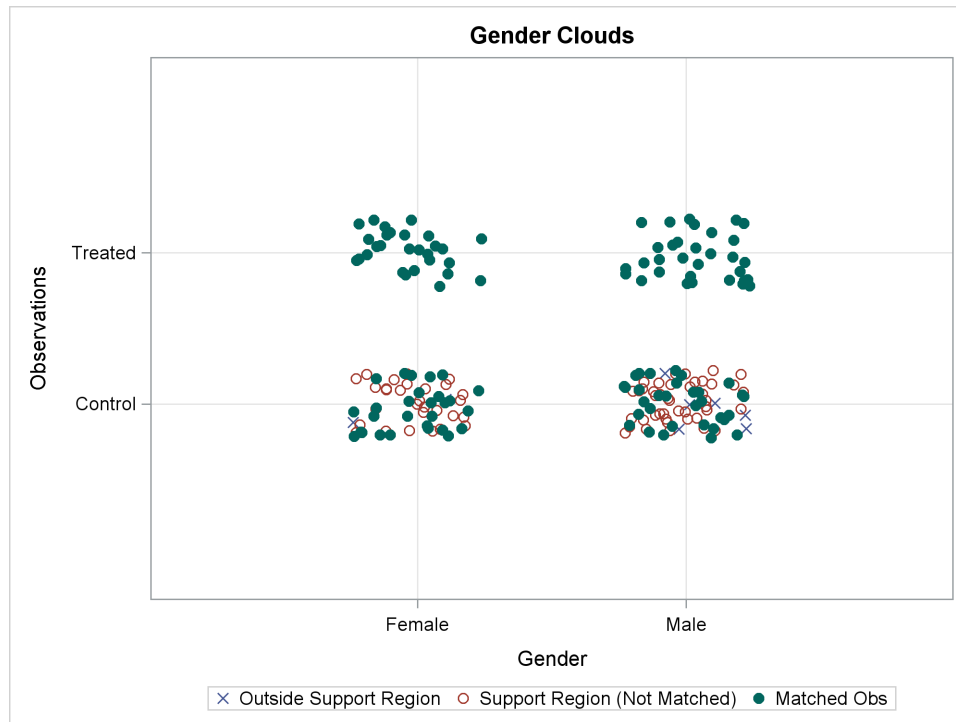
Output 95.4.11 displays a cloud plot that compares the values of Absence for observations in the treated and control groups, based on all observations, observations in the support region, and matched observations.

Output 95.4.11 Absence Cloud Plot



Output 95.4.12 displays a cloud plot that compares the values of Gender for observations in the treated and control groups, based on all observations, observations in the support region, and matched observations.

Output 95.4.12 Gender Cloud Plot



If you are not satisfied with the variable balance, you can do one or more of the following until you are satisfied: you can select another set of variables to fit the propensity score model, you can modify the matching criteria, or you can choose another matching method.

The OUT(OBS=MATCH)=OutEx4 option in the OUTPUT statement creates an output data set, OutEx4, that contains the matched observations. The following statements list the 10 observations in OutEx4 that have lowest propensity scores, as shown in Output 95.4.13:

```
proc sort data=OutEx4 out=OutEx4a;
  by _PS_;
run;

proc print data=OutEx4a(obs=10);
  var StudentID Music Gender Absence _PS_ _MATCHWGT_ _MatchID;
run;
```

Output 95.4.13 Output Data Set with Multiple Matches

Obs	StudentID	Music	Gender	Absence	_PS_	_MATCHWGT_	_MatchID
1	33	Yes	Female	3.50313	0.09283	1	60
2	82	No	Female	3.50036	0.09296	1	60
3	67	Yes	Female	2.71352	0.13790	1	59
4	95	No	Female	2.53865	0.15009	1	59
5	47	No	Female	2.49866	0.15300	1	58
6	4	Yes	Female	2.49425	0.15333	1	58
7	37	No	Male	2.93955	0.15549	1	57
8	152	Yes	Male	2.88102	0.15988	1	57
9	20	Yes	Female	2.12751	0.18224	1	56
10	121	No	Female	2.12239	0.18267	1	56

By default, the output data set includes the variable `_PS_` (which provides the propensity score) and the variable `_MATCHWGT_` (which provides matched observation weights). The `MATCHID=_MatchID` option creates a variable named `_MatchID` that identifies the matched sets of observations.

If you assume that no other confounding variables are associated with both the GPA and the music class indicator `Music`, you can add the GPAs for the students to the data set `OutEx4` and perform an outcome analysis of GPA on this data set to estimate the music class effect.

Example 95.5: Matching with Replacement

This example performs matching with replacement in the propensity score analysis. The data set `School` contains the student information and is described in [Example 95.4](#).

The following statements invoke the `PSMATCH` procedure and request matching with replacement to match observations for students in the treatment group with observations for students in the control group:

```
ods graphics on;
proc psmatch data=School region=allobs(psmin=0.05);
  class Music Gender;
  psmodel Music(Treated='Yes')= Gender Absence;
  match method=replace(k=1) stat=ps exact=Gender caliper=.;
  assess ps var=(Gender Absence);
  output out(obs=match)=outex5 matchid=_MatchID;
run;
```

The `CLASS` statement specifies the classification variables. The `PSMODEL` statement specifies the logistic regression model that creates the propensity score for each student, which is the probability that the student enrolled in the music class. The `Music` variable is the binary treatment indicator variable, and `TREATED='Yes'` identifies Yes as the treated group. The `Gender` and `Absence` variables are included in the model because they are believed to be related to enrolling in the music class.

The `REGION=` option specifies an interval region of propensity scores such that only observations that have propensity scores in the region are used in stratification and matching. Because the `MATCH` statement is also specified, the `REGION=ALLOBS(PSMIN=0.05)` option requests that all available observations whose propensity scores are greater than or equal to 0.05 be used for matching.

The MATCH statement requests matching and specifies the criteria for matching. The STAT=PS option requests that the propensity score be used in computing differences between pairs of observations. The METHOD=REPLACE(K=1) option requests matching with replacement in which each treated unit is matched to the closest control unit.

The EXACT=GENDER option forces the treated unit and its matched control unit to have the same value of the Gender variable. The CALIPER=. option ignores the caliper requirement for matching.

The “Data Information” table in [Output 95.5.1](#) displays information about the input and output data sets, the numbers of observations in the treated and control groups, the lower and upper limits for the propensity score support region, and the numbers of observations in the treated and control groups that fall within the support region. Of the 140 observations in the control group, 134 fall within the support region.

Output 95.5.1 Data Information
The PSMATCH Procedure

Data Information	
Data Set	WORK.SCHOOL
Output Data Set	WORK.OUTEX5
Treatment Variable	Music
Treatment Group	Yes
All Obs (Treatment)	60
All Obs (Control)	140
Support Region	PS Bounded Obs
Lower PS Support	0.05
Upper PS Support	0.490191
Support Region Obs (Treatment)	60
Support Region Obs (Control)	134

The “Propensity Score Information” table in [Output 95.5.2](#) displays summary statistics by treatment group for all observations, for support region observations, and for matched observations.

Output 95.5.2 Propensity Score Information

Propensity Score Information										
Treated (Music = Yes)						Control (Music = No)				
Observations	N	Mean	Std Dev	Minimum	Maximum	N	Mean	Std Dev	Minimum	Maximum
All	60	0.347143	0.096184	0.092831	0.490191	140	0.279796	0.124997	0.026465	0.488875
Region	60	0.347143	0.096184	0.092831	0.490191	134	0.290611	0.116521	0.051746	0.488875
Matched	60	0.347143	0.096184	0.092831	0.490191	41	0.335146	0.103487	0.092963	0.488875

Note that the number of matched control units (41) is less than the number of matched treated units (60) because for matching with replacement, a control unit can be matched with more than one treated unit.

The “Matching Information” table in [Output 95.5.3](#) displays the matching criteria, the number of matched sets, the numbers of matched observations in the treated and control groups, and the total absolute difference in the propensity score for all matches. For matching with replacement, one control unit might be matched to more than one treated unit. In this example, 41 control units are matched to 60 treated units.

Output 95.5.3 Matching Information

Matching Information	
Difference Statistic	Propensity Score
Method	Replacement Matching
Control/Treated Ratio	1
Matched Sets	41
Matched Obs (Treated)	60
Matched Obs (Control)	41
Total Absolute Difference	0.264451

The ASSESS statement produces tables and plots that summarize differences in the specified variables between treated and control groups for all observations, for the support region observations, and for the matched observations. You can use these results to assess how well the matching achieves a balance in the distributions of these variables. As requested by the PS and VAR= options, the variables are the propensity score and the covariates Gender and Absence.

The “Standardized Variable Differences” table displays standardized differences between the treated and control groups for all observations, the support region observations, and the matched observations, as shown in [Output 95.5.4](#).

Output 95.5.4 Standardized Differences**The PSMATCH Procedure****Standardized Variable Differences (Treated - Control)**

Variable	Mean Difference			Weighted Matched Obs
	All Obs	Region Obs	Matched Obs	
PS	0.067347	0.056532	0.011997	0.001176
Absence	-0.697568	-0.531749	-0.087541	-0.008541
Gender	-0.045238	-0.033831	-0.052033	0

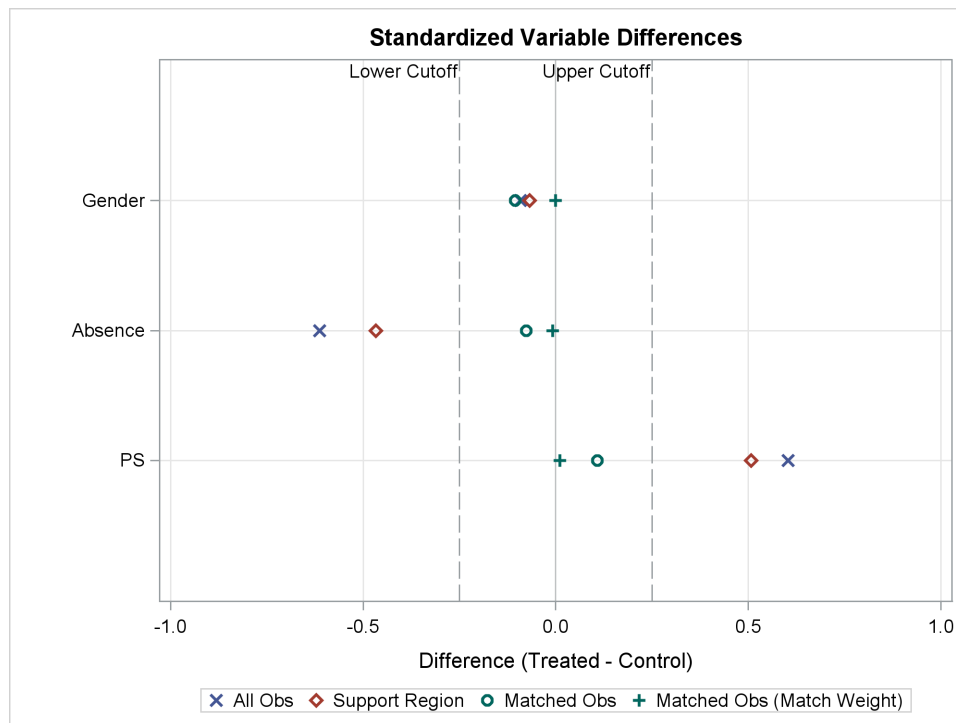
Standardized Variable Differences (Treated - Control)**Standardized Mean Difference**

Variable	Divisor	Mean Difference			Percent Reduction			
		All Obs	Region Obs	Matched Obs	Weighted Matched Obs	Region Obs	Matched Obs	Weighted Matched Obs
PS	0.111525	0.603876	0.506900	0.107574	0.010546	16.06	82.19	98.25
Absence	1.136945	-0.613546	-0.467700	-0.076997	-0.007513	23.77	87.45	98.78
Gender	0.496344	-0.091143	-0.068160	-0.104831	0	25.22	0.00	100.00

**Standardized Variable Differences
(Treated - Control)****Variance Ratio**

Variable	All Obs	Region Matched		Weighted Matched Obs
		Obs	Obs	
PS	0.5921	0.6814	0.8638	1.0062
Absence	0.3550	0.5050	0.8192	1.0029
Gender	1.0208	1.0138	1.0254	1.0000

When ODS Graphics is enabled, the PSMATCH procedure displays a standardized variable differences plot for the variables that are specified in the ASSESS statement, as shown in [Output 95.5.5](#).

Output 95.5.5 Standardized Differences Plot

The “Standardized Variable Differences Plot” displays the standardized differences in the “Variable Differences” table in [Output 95.5.4](#). All differences for the matched observations are within the recommended limits of -0.25 and 0.25 , which are indicated by reference lines.

If you are not satisfied with the variable balance, you can do one or more of the following until you are satisfied: you can select another set of variables to fit the propensity score model, you can modify the matching criteria, or you can choose another matching method.

The `OUT(OBS=MATCH)=OutEx5` option in the `OUTPUT` statement creates an output data set, `OutEx5`, that contains the matched observations. The following statements list the 10 observations in `OutEx5` that have the largest propensity scores, as shown in [Output 95.5.6](#).

```
proc sort data=OutEx5 out=OutEx5a;
  by descending _PS_;
run;

proc print data=OutEx5a(obs=10);
  var StudentID Music Gender Absence _PS_ _MATCHWGT_ _MatchID;
run;
```

Output 95.5.6 Output Data Set of Matched Observations with Replacement

Obs	StudentID	Music	Gender	Absence	_PS_	_MATCHWGT_	_MatchID
1	156	Yes	Male	0.01792	0.49019	1	41
2	142	Yes	Male	0.02451	0.48926	1	41
3	105	No	Male	0.02723	0.48888	2	41
4	64	No	Male	0.03317	0.48804	1	40
5	98	Yes	Male	0.05443	0.48503	1	40
6	30	No	Male	0.10043	0.47853	2	39
7	182	Yes	Male	0.10352	0.47810	1	39
8	115	Yes	Male	0.11002	0.47718	1	39
9	130	No	Male	0.17651	0.46780	1	38
10	104	Yes	Male	0.18986	0.46592	1	38

By default, the output data set includes the variable `_PS_` (which provides the propensity score) and the variable `_MATCHWGT_` (which provides matched observation weights). The `MATCHID=_MatchID` option creates a variable named `_MatchID` that identifies the matched sets of observations.

If you assume that no other confounding variables are associated with both the GPA and the music class indicator `Music`, you can add the GPAs for the students to the data set `OutEx5` and perform an outcome analysis of GPA on this data set to estimate the music class effect.

Example 95.6: Mahalanobis Distance Matching

This example performs Mahalanobis distance matching, where the distances between patients in the treatment group and patients in the control group are computed from a set of variables. The data set `Drugs` contains the patient information and is described in the section “[Getting Started: PSMATCH Procedure](#)” on page 7680.

The following statements invoke the `PSMATCH` procedure and request an optimal matching to match patients in the treatment group to patients in the control group, based on Mahalanobis distances:

```
ods graphics on;
proc psmatch data=drugs region=cs;
  class Drug Gender;
  psmodel Drug(Treated='Drug_X')= Gender Age Bmi;
  match method=optimal(k=1) exact=Gender
        stat=mah(lps var=(Age Bmi)) caliper=.25;
  assess lps var=(Gender Age Bmi) / weight=none;
  output out(obs=match)=OutEx6 matchid=_MatchID;
run;
```

The `CLASS` statement specifies the classification variables. The `PSMODEL` statement specifies the logistic regression model that creates the propensity score for each observation, which is the probability that the patient receives `Drug_X`. The `Drug` variable is the binary treatment indicator variable, and `TREATED='Drug_X'` identifies `Drug_X` as the treated group. The `Gender`, `Age`, and `Bmi` variables are included in the model because they are believed to be related to the assignment.

The `REGION=` option specifies an interval region of propensity scores (or equivalently, logits of propensity scores) such that only observations that have propensity scores in the region are used in stratification and matching. Because the `MATCH` statement is also specified, the `REGION=CS` option requests that only

observations that have propensity scores in the common support region be used for matching. By default, the region is extended by 0.25 times the pooled estimate of the common standard deviation of the logit of the propensity score, where this estimate is the square root of the average of the variances in the treated and control groups.

The MATCH statement specifies the criteria for matching. The STAT=MAH(LPS VAR=(AGE Bmi)) option requests that the Mahalanobis distance, computed from the logit of propensity score and the Age and Bmi variables, be used in computing differences between pairs of observations. The METHOD=OPTIMAL(K=1) option (which is the default) requests optimal matching of one control unit to each unit in the treated group in order to minimize the total within-pair difference. The EXACT=GENDER option forces the treated unit and its matched control unit to have the same value of the Gender variable. The CALIPER=0.25 option specifies the caliper requirement for matching: for a match to be made, the difference in the logits of the propensity score for pairs of individuals from the two groups must be less than or equal to 0.25 times the pooled estimate of the common standard deviation of the logits of the propensity scores.

The “Data Information” table in [Output 95.6.1](#) displays information about the input and output data sets, the numbers of observations in the treated and control groups, the lower and upper limits for the propensity score support region, and the numbers of observations in the treated and control groups that fall within the support region. Of the 373 observations in the control group, 351 fall within the support region.

Output 95.6.1 Data Information

The PSMATCH Procedure

Data Information	
Data Set	WORK.DRUGS
Output Data Set	WORK.OUTEX6
Treatment Variable	Drug
Treatment Group	Drug_X
All Obs (Treatment)	113
All Obs (Control)	373
Support Region	Extended Common Support
Lower PS Support	0.050244
Upper PS Support	0.683999
Support Region Obs (Treatment)	113
Support Region Obs (Control)	351

The “Propensity Score Information” table in [Output 95.6.2](#) displays summary statistics by treatment group for all observations, for support region observations, and for matched observations.

Output 95.6.2 Propensity Score Information

Propensity Score Information										
	Treated (Drug = Drug_X)					Control (Drug = Drug_A)				
Observations	N	Mean	Std Dev	Minimum	Maximum	N	Mean	Std Dev	Minimum	Maximum
All	113	0.310773	0.132467	0.060231	0.641148	373	0.208801	0.131969	0.020157	0.685757
Region	113	0.310773	0.132467	0.060231	0.641148	351	0.217557	0.126747	0.050951	0.682374
Matched	113	0.310773	0.132467	0.060231	0.641148	113	0.305309	0.133560	0.064010	0.682374

The “Matching Information” table in [Output 95.6.3](#) displays the matching criteria, the number of matched sets, the numbers of matched observations in the treated and control groups, and the total Mahalanobis difference for all matches.

Output 95.6.3 Matching Information

Matching Information	
Difference Statistic	Mahalanobis Distance
Mahalanobis Covariance	Control Group
Method	Optimal Fixed Ratio Matching
Control/Treated Ratio	1
Caliper (Logit PS)	0.191862
Matched Sets	113
Matched Obs (Treated)	113
Matched Obs (Control)	113
Total Absolute Difference	24.46784

The ASSESS statement produces the tables and plots that summarize differences in the specified variables between treated and control groups for all observations, for the support region observations, and for the matched observations. As requested by the LPS and VAR= options, the variables listed in the table are the logit of propensity score and the variables Gender, Age, and Bmi. The WEIGHT=NONE option suppresses display of differences for the weighted matched observations. Note that for a matching of one control unit to each treated unit, the weights are all 1 for matched treated and control units, and the results are identical for the weighted matched observations and the matched observations.

The “Standardized Variable Differences” table, as shown in [Output 95.6.4](#), displays standardized differences between the treated and control groups for all observations, the support region observations, and the matched observations. For a binary classification variable (Gender), the difference is in the proportion of the first ordered level (Female).

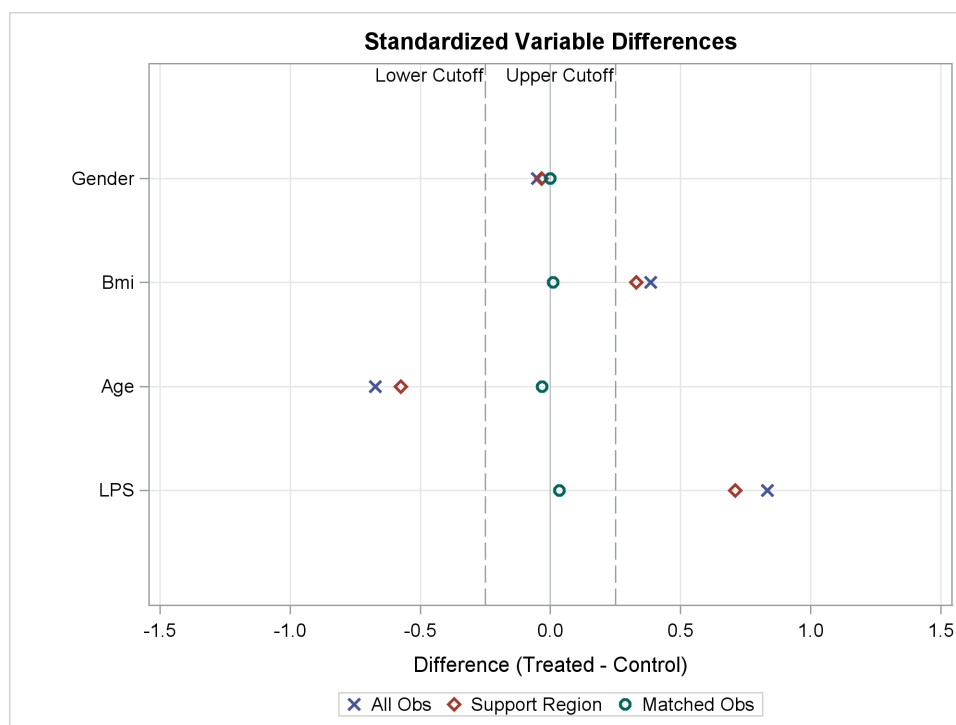
Output 95.6.4 Standardized Differences**The PSMATCH Procedure**

Standardized Variable Differences (Treated - Control)									
					Standardized Mean Difference				
Mean Difference					Mean Difference			Percent Reduction	
Variable	All Obs	Region Obs	Matched Obs	Divisor	All Obs	Region Obs	Matched Obs	Region Obs	Matched Obs
LPS	0.639971	0.545459	0.026321	0.767448	0.833894	0.710744	0.034297	14.77	95.89
Age	-4.095091	-3.493684	-0.194690	6.079104	-0.673634	-0.574704	-0.032026	14.69	95.25
Bmi	0.739296	0.632566	0.018850	1.923178	0.384414	0.328917	0.009801	14.44	97.45
Gender	-0.024817	-0.016514	0	0.496925	-0.049941	-0.033233	0	33.46	100.00

Standardized Variable Differences (Treated - Control)				
Variance Ratio				
Variable	All Obs	Region Obs	Matched Obs	
LPS	0.6517	0.8314	0.9899	
Age	0.7076	0.8000	1.0214	
Bmi	0.8854	0.9288	1.0744	
Gender	0.9892	0.9922	1.0000	

The standardized mean differences are significantly reduced in the matched observations, and the largest of these differences is 0.0343 in absolute value, which is less than the recommended upper limit of 0.25. The variance ratios between the two groups are also in the recommended range of 0.5 to 2. Because both EXACT=GENDER and METHOD=OPTIMAL are specified in the MATCH statement, the standardized difference for Gender is 0 in the matched observations.

When ODS Graphics is enabled, the PSMATCH procedure displays a standardized variable differences plot for the variables that are specified in the ASSESS statement, as shown in [Output 95.6.5](#).

Output 95.6.5 Standardized Differences Plot

The “Standardized Variable Differences Plot” displays the standardized differences in the “Variable Differences” table in [Output 95.6.4](#). All differences for the matched observations are within the recommended limits of -0.25 and 0.25 , which are indicated by reference lines.

If you are not satisfied with the variable balance, you can do one or more of the following until you are satisfied: you can select another set of variables to fit the propensity score model, you can modify the matching criteria, or you can choose another matching method.

The `OUT(OBS=MATCH)=OutEx6` option in the `OUTPUT` statement creates an output data set, `OutEx6`, that contains the matched observations. The following statements list the 10 observations in `OutEx6` that have lowest propensity scores, as shown in [Output 95.6.6](#):

```
proc sort data=OutEx6 out=OutEx6a;
  by _ps_;
run;

proc print data=OutEx6a(obs=10);
  var PatientID Drug Gender Age Bmi _ps_ _MatchWgt_ _MatchID;
run;
```

Output 95.6.6 Output Data Set with Mahalanobis Distance Matches

Obs	PatientID	Drug	Gender	Age	Bmi	_PS_	_MATCHWGT_	_MatchID
1	89	Drug_X	Female	44	20.75	0.06023	1	1
2	141	Drug_A	Female	43	20.55	0.06401	1	1
3	323	Drug_X	Female	46	22.22	0.06763	1	2
4	137	Drug_A	Female	45	22.04	0.07215	1	2
5	111	Drug_A	Female	41	21.01	0.08714	1	4
6	217	Drug_X	Male	49	23.96	0.08772	1	3
7	429	Drug_A	Male	49	24.00	0.08848	1	3
8	234	Drug_X	Female	41	21.11	0.08904	1	4
9	183	Drug_A	Female	45	23.62	0.10157	1	6
10	189	Drug_A	Female	46	24.10	0.10171	1	5

By default, the output data set includes the variable `_PS_` (which provides the propensity score) and the variable `_MATCHWGT_` (which provides matched observation weights). The weight for each treated unit is 1. Because `K=1` is specified in the `METHOD=` option in the `MATCH` statement, one control unit is matched to each treated unit; so the weight for each matched control unit is also 1. The `MATCHID=_MatchID` option creates a variable named `_MatchID` that identifies the matched sets of observations.

If you assume that no other confounding variables are associated with both the response variable and the treatment group indicator `Drug`, then after the responses for the trial are observed and added to the data set `OutEx6`, you can use the same outcome analysis on this output data set as you would have used on the original data set `Drugs` (augmented with responses) to estimate the treatment effect.

Example 95.7: Matching with Existing Propensity Scores in the Input Data Set

This example performs optimal matching with propensity scores already available in the input data set. The `Drugs` data set contains the patient information and is described in the section “[Getting Started: PSMATCH Procedure](#)” on page 7680.

The following statements use the `LOGISTIC` procedure to fit a binary complementary log-log model and to derive propensity scores:

```
ods select none;
proc logistic data=drugs;
  class Drug Gender;
  model Drug(Event='Drug_X')= Gender Age Bmi / link=cloglog;
  output out=drug1 p=pscore;
run;
ods select all;
```

The output data set `Drug1` is constructed from the data set `Drugs` and contains the `PScore` variable for propensity scores.

Output 95.7.1 lists the first 10 observations.

Output 95.7.1 Output Drug Data Set with Propensity Scores
First 10 Observations of the Input Drug Data Set

Obs	PatientID	Drug	Gender	Age	Bmi	pscore
1	284	Drug_X	Male	29	22.02	0.35498
2	201	Drug_A	Male	45	26.68	0.21794
3	147	Drug_A	Male	42	21.84	0.12261
4	307	Drug_X	Male	38	22.71	0.19821
5	433	Drug_A	Male	31	22.76	0.34298
6	435	Drug_A	Male	43	26.86	0.26261
7	159	Drug_A	Female	45	25.47	0.15077
8	368	Drug_A	Female	49	24.28	0.08713
9	286	Drug_A	Male	31	23.31	0.37211
10	163	Drug_X	Female	39	25.34	0.24005

The following statements invoke the PSMATCH procedure and request an optimal matching to match patients in the treatment group to patients in the control group:

```
ods graphics on;
proc psmatch data=Drug1 region=cs;
  class Drug Gender;
  psdata treatvar=Drug(Treated='Drug_X') ps=pscore;
  match method=optimal(k=1) exact=Gender stat=lps caliper=0.5;
  assess lps var=(Gender Age Bmi) / weight=none;
  output out(obs=match)=OutEx7 lps=_Lps matchid=_MatchID;
run;
```

The CLASS statement specifies the classification variables. The PSDATA statement specifies the binary treatment variable and the variable for propensity score information in the input DATA= data set. The TREATVAR=DRUG option specifies DRUG as the binary treatment indicator variable, and TREATED='Drug_X' identifies Drug_X as the treated group.

The REGION= option specifies an interval region of propensity scores (or equivalently, logits of propensity scores) such that only observations that have propensity scores in the region are used in stratification and matching. Because the MATCH statement is also specified, the REGION=CS option requests that only observations that have propensity scores in the common support region be used for matching. By default, the region is extended by 0.25 times the pooled estimate of the common standard deviation of the logit of the propensity score, where this estimate is the square root of the average of the variances in the treated and control groups.

The MATCH statement specifies the criteria for matching. The STAT=LPS option (which is the default) requests that the logit of the propensity score be used in computing differences between pairs of observations. The METHOD=OPTIMAL(K=1) option (which is the default) requests optimal matching of one control unit to each unit in the treated group in order to minimize the total within-pair difference. The EXACT=GENDER option forces the treated unit and its matched control unit to have the same value of the Gender variable. The CALIPER=0.5 option specifies the caliper requirement for matching: for a match to be made, the difference in the logits of the propensity score for pairs of individuals from the two groups must be less than or equal to 0.5 times the pooled estimate of the common standard deviation of the logits of propensity scores.

The “Data Information” table in [Output 95.7.2](#) displays information about the input and output data sets, the numbers of observations in the treated and control groups, the lower and upper limits for the propensity score support region, and the numbers of observations in the treated and control groups that fall within the support region. Of the 373 observations in the control group, 352 fall within the support region.

Output 95.7.2 Data Information

The PSMATCH Procedure

Data Information	
Data Set	WORK.DRUG1
Output Data Set	WORK.OUTEX7
Treatment Variable	Drug
Treatment Group	Drug_X
All Obs (Treatment)	113
All Obs (Control)	373
Support Region	Extended Common Support
Lower PS Support	0.060563
Upper PS Support	0.698199
Support Region Obs (Treatment)	113
Support Region Obs (Control)	352

The “Propensity Score Information” table in [Output 95.7.3](#) displays summary statistics by treatment group for all observations, for the support region observations, and for the matched observations.

Output 95.7.3 Propensity Score Information

Propensity Score Information										
		Treated (Drug = Drug_X)				Control (Drug = Drug_A)				
Observations	N	Mean	Std Dev	Minimum	Maximum	N	Mean	Std Dev	Minimum	Maximum
All	113	0.304022	0.128669	0.071521	0.659420	373	0.208864	0.125520	0.029452	0.713529
Region	113	0.304022	0.128669	0.071521	0.659420	352	0.214597	0.117684	0.060627	0.651891
Matched	113	0.304022	0.128669	0.071521	0.659420	113	0.298444	0.121541	0.072344	0.651891

The “Matching Information” table in [Output 95.7.4](#) displays the matching criteria, the number of matched sets, the numbers of matched observations in the treated and control groups, and the total absolute difference in the logit of the propensity score for all matches.

Output 95.7.4 Matching Information

Matching Information	
Difference Statistic	Logit of Propensity Score
Method	Optimal Fixed Ratio Matching
Control/Treated Ratio	1
Caliper (Logit PS)	0.356051
Matched Sets	113
Matched Obs (Treated)	113
Matched Obs (Control)	113
Total Absolute Difference	3.616259

The ASSESS statement produces the tables and plots that summarize differences in the specified variables between treated and control groups for all observations, for the support region observations, and for the matched observations. As requested by the LPS and VAR= options, the variables listed in the table are the logit of propensity score and the variables Gender, Age, and Bmi. The WEIGHT=NONE option suppresses display of differences for the weighted matched observations. Note that for a matching of one control unit to each treated unit, the weights are all 1 for matched treated and control units, and the results are identical for the weighted matched observations and the matched observations.

The “Standardized Variable Differences” table displays standardized differences between the treated and control groups for all observations, for the support region observations, and for the matched observations, as shown in [Output 95.7.5](#). For a binary classification variable (Gender), the difference is in the proportion of the first ordered level (Female).

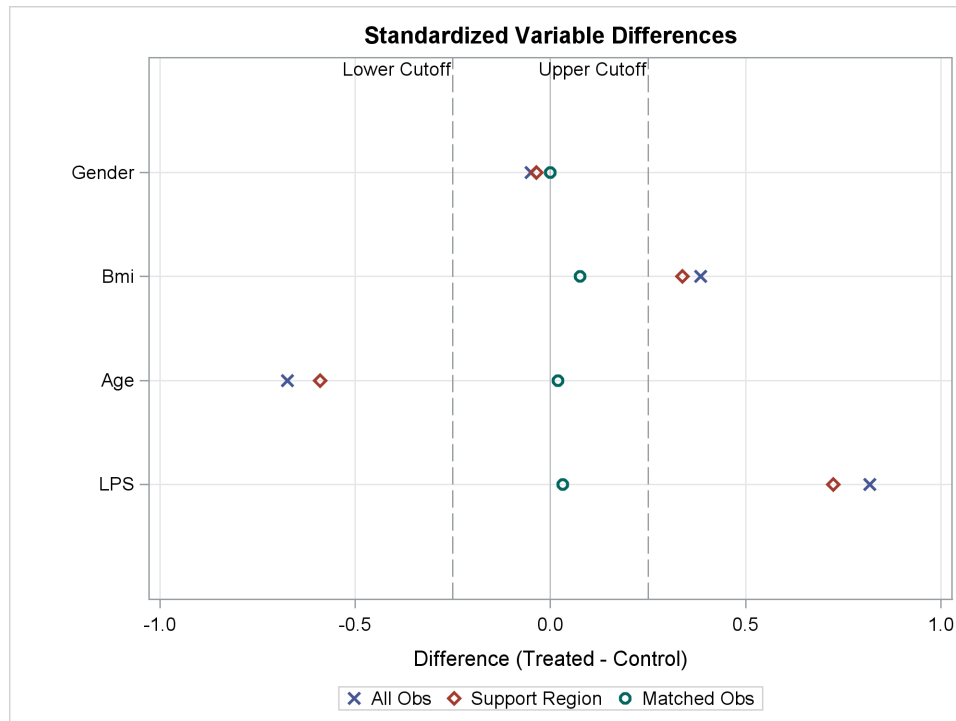
Output 95.7.5 Standardized Differences

The PSMATCH Procedure

Standardized Variable Differences (Treated - Control)													
Standardized Mean Difference													
Variable	Mean Difference				Mean Difference				Percent Reduction		Variance Ratio		
	All Obs	Region Obs	Matched Obs	Divisor	All Obs	Region Obs	Matched Obs		Region Obs	Matched Obs	All Obs	Region Obs	Matched Obs
LPS	0.582391	0.516130	0.022677	0.712102	0.817848	0.724797	0.031845	11.38	96.11	0.7177	0.9052	1.0929	
Age	-4.095091	-3.585152	0.115044	6.079104	-0.673634	-0.589750	0.018925	12.45	97.19	0.7076	0.7928	1.0143	
Bmi	0.739296	0.650890	0.146195	1.923178	0.384414	0.338445	0.076017	11.96	80.23	0.8854	0.9394	1.3509	
Gender	-0.024817	-0.018076	0	0.496925	-0.049941	-0.036376	0	27.16	100.00	0.9892	0.9916	1.0000	

The standardized mean differences are significantly reduced in the matched observations, and the largest of these differences is 0.076 in absolute value, which is less than the recommended upper limit of 0.25. The variance ratios between the two groups are between 1 and 1.1967 for all variables in the matched observations, which is within the recommended range of 0.5 to 2. Because both EXACT=GENDER and METHOD=OPTIMAL are specified in the MATCH statement, the standardized difference for Gender is 0 in the matched observations.

When ODS Graphics is enabled, the PSMATCH procedure displays a standardized variable differences plot for the variables that are specified in the ASSESS statement, as shown in [Output 95.7.6](#).

Output 95.7.6 Standardized Differences Plot

The “Standardized Variable Differences Plot” displays the standardized differences in the “Standardized Variable Differences” table in [Output 95.7.5](#). All differences for the matched observations are within the recommended limits of -0.25 and 0.25 , which are indicated by reference lines.

If you are not satisfied with the variable balance, you can do one or more of the following until you are satisfied: you can select another set of variables to fit the propensity score model, you can modify the matching criteria, or you can choose another matching method.

The `OUT(OBS=MATCH)=OutEx7` option in the `OUTPUT` statement creates an output data set, `OutEx7`, that contains the matched observations. The following statements list the 10 observations `OutEx7` that have lowest propensity scores, as shown in [Output 95.7.7](#):

```
proc sort data=OutEx7 out=OutEx7a;
  by pscore;
run;

proc print data=OutEx7a(obs=10);
  var PatientID Drug Gender Age Bmi pscore _LPS _MatchWgt_ _MatchID;
run;
```

Output 95.7.7 Output Data Set With Optimal Matches

Obs	PatientID	Drug	Gender	Age	Bmi	pscore	_Lps	_MATCHWGT_	_MatchID
1	89	Drug_X	Female	44	20.75	0.07152	-2.56356	1	1
2	213	Drug_A	Female	49	23.24	0.07234	-2.55123	1	1
3	323	Drug_X	Female	46	22.22	0.07822	-2.46677	1	2
4	245	Drug_A	Female	52	25.32	0.08090	-2.43015	1	2
5	217	Drug_X	Male	49	23.96	0.09796	-2.22013	1	3
6	429	Drug_A	Male	49	24.00	0.09865	-2.21228	1	3
7	234	Drug_X	Female	41	21.11	0.09887	-2.20987	1	4
8	66	Drug_A	Female	48	24.53	0.09927	-2.20531	1	4
9	183	Drug_A	Female	45	23.62	0.10931	-2.09786	1	5
10	320	Drug_X	Female	46	24.17	0.11056	-2.08507	1	5

By default, the output data set includes the variable `_PS_` (which provides the propensity score) and the variable `_MATCHWGT_` (which provides matched observation weights). The weight for each treated unit is 1. Because $K=1$ is specified in the `METHOD=` option in the `MATCH` statement, one control unit is matched to each treated unit; so the weight for each matched control unit is also 1. The `LPS=_LPS` option creates a variable named `_LPS`, which provides the logit of propensity score, and the `MATCHID=_MatchID` option creates a variable named `_MatchID`, which identifies the matched sets of observations.

If you assume that no other confounding variables are associated with both the response variable and the treatment group indicator `Drug`, then after the responses for the trial are observed and added to the data set `OutEx7`, you can use the same outcome analysis on this output data set as you would have used on the original data set `Drugs` (with added responses) to estimate the treatment effect.

References

- Austin, P. C. (2007). "The Performance of Different Propensity Score Methods for Estimating Marginal Odds Ratios." *Statistics in Medicine* 26:3078–3094.
- Austin, P. C. (2009). "Balance Diagnostics for Comparing the Distribution of Baseline Covariates between Treatment Groups in Propensity-Score Matched Samples." *Statistics in Medicine* 28:3083–3107.
- Austin, P. C. (2011a). "An Introduction to Propensity Score Methods for Reducing the Effects of Confounding in Observational Studies." *Multivariate Behavioral Research* 46:399–424.
- Austin, P. C. (2011b). "Optimal Caliper Widths for Propensity-Score Matching When Estimating Differences in Means and Differences in Proportions in Observational Studies." *Pharmaceutical Statistics* 10:150–161.
- Austin, P. C. (2014). "A Comparison of 12 Algorithms for Matching on the Propensity Score." *Statistics in Medicine* 33:1057–1069.
- Austin, P. C., Grootendorst, P., and Anderson, G. M. (2007). "A Comparison of the Ability of Different Propensity Score Models to Balance Measures Variables between Treated and Untreated Subjects: A Monte Carlo Study." *Statistics in Medicine* 26:734–753.

- Austin, P. C., and Stuart, E. A. (2015a). “Moving towards Best Practice When Using Inverse Probability of Treatment Weighting (IPTW) Using the Propensity Score to Estimate Causal Treatment Effects in Observational Studies.” *Statistics in Medicine* 34:3661–3679.
- Austin, P. C., and Stuart, E. A. (2015b). “The Performance of Inverse Probability of Treatment Weighting and Full Matching on the Propensity Score in the Presence of Model Misspecification When Estimating the Effect of Treatment on Survival Outcomes.” *Statistical Methods in Medical Research* <http://smm.sagepub.com/content/early/2015/04/30/0962280215584401.full.pdf+html>.
- Crump, R. K., Hotz, V. J., Imbens, G. W., and Mitnik, O. A. (2009). “Dealing with Limited Overlap in Estimation of Average Treatment Effects.” *Biometrika* 96:187–199.
- Faries, D. E., Leon, A. C., Haro, J. M., and Obenchain, R. L., eds. (2010). *Analysis of Observational Health Care Data Using SAS*. Cary, NC: SAS Institute Inc.
- Guo, S., and Fraser, M. W. (2015). *Propensity Score Analysis: Statistical Methods and Applications*. 2nd ed. Thousand Oaks, CA: Sage Publications.
- Hansen, B. B. (2004). “Full Matching in an Observational Study of Coaching for the SAT.” *Journal of the American Statistical Association* 99:609–618.
- Hill, J., and Reiter, J. P. (2006). “Interval Estimation for Treatment Effects Using Propensity Score Matching.” *Statistics in Medicine* 25:2230–2256.
- Ho, D., Imai, K., King, G., and Stuart, E. A. (2007). “Matching as Nonparametric Preprocessing for Reducing Model Dependence in Parametric Causal Inference.” *Political Analysis* 15:199–236.
- Imbens, G. W., and Rubin, D. B. (2015). *Causal Inference for Statistics, Social, and Biomedical Sciences: An Introduction*. New York: Cambridge University Press.
- Lunceford, J. K., and Davidian, M. (2004). “Stratification and Weighting via the Propensity Score in Estimation of Causal Treatment Effects: A Comparative Study.” *Statistics in Medicine* 23:2937–2960.
- Mamdani, M., Sykora, K., Li, P., Normand, S. L., Streiner, D. L., Austin, P. C., Rochon, P. A., and Anderson, G. M. (2005). “Reader’s Guide to Critical Appraisal of Cohort Studies: 2. Assessing Potential for Confounding.” *BMJ* 330:960–962.
- Normand, S.-L. T., Landrum, M. B., Guadagnoli, E., Ayanian, J. Z., Ryan, T. J., Cleary, P. D., and McNeil, B. J. (2001). “Validating Recommendations for Coronary Angiography Following Acute Myocardial Infarction in the Elderly: A Matched Analysis Using Propensity Scores.” *Journal of Clinical Epidemiology* 54:387–398.
- Rosenbaum, P. R. (2010). *Design of Observational Studies*. New York: Springer-Verlag.
- Rosenbaum, P. R., and Rubin, D. B. (1983). “The Central Role of the Propensity Score in Observational Studies for Causal Effects.” *Biometrika* 70:41–55.
- Rosenbaum, P. R., and Rubin, D. B. (1984). “Reducing Bias in Observational Studies Using Subclassification on the Propensity Score.” *Journal of the American Statistical Association* 79:516–524.
- Rosenbaum, P. R., and Rubin, D. B. (1985). “Constructing a Control Group Using Multivariate Matched Sampling Methods That Incorporate the Propensity Score.” *American Statistician* 39:33–38.

- Rubin, D. B. (1974). "Estimating Causal Effects of Treatments in Randomized and Nonrandomized Studies." *Journal of Educational Psychology* 66:688–701.
- Rubin, D. B. (1980a). "Bias Reduction Using Mahalanobis-Metric Matching." *Biometrics* 36:293–298.
- Rubin, D. B. (1980b). "Comment on D. Basu, 'Randomization Analysis of Experimental Data: The Fisher Randomization Test'." *Journal of the American Statistical Association* 75:591–593.
- Rubin, D. B. (1990). "Comment: Neyman (1923) and Causal Inference in Experiments and Observational Studies." *Statistical Science* 5:472–480.
- Rubin, D. B. (2001). "Using Propensity Scores to Help Design Observational Studies: Application to the Tobacco Litigation." *Health Services and Outcomes Research Methodology* 2:169–188.
- Rubin, D. B. (2005). "Causal Inference Using Potential Outcomes: Design, Modeling, Decisions." *Journal of the American Statistical Association* 100:322–331.
- Stuart, E. A. (2010). "Matching Methods for Causal Inference: A Review and a Look Forward." *Statistical Science* 25:1–21.
- Yue, L. Q., Lu, N., and Xu, Y. (2014). "Designing Premarket Observational Comparative Studies Using Existing Data as Controls: Challenges and Opportunities." *Journal of Biopharmaceutical Statistics* 24:994–1010.

Subject Index

- data information
 - PSMATCH procedure, [7716](#)
- matching information
 - PSMATCH procedure, [7716](#)
- propensity score analysis
 - process, [7677](#)
- propensity score information
 - PSMATCH procedure, [7716](#)
- PSMATCH procedure
 - data information, [7716](#)
 - features, [7679](#)
 - introductory example, [7680](#)
 - matching information, [7716](#)
 - ODS Graphics names, [7719](#)
 - ODS table names, [7717](#)
 - propensity score information, [7716](#)
 - standardized differences plot, [7719](#)
 - standardized variable differences, [7716](#)
 - strata information, [7717](#)
 - strata standardized variable differences, [7717](#)
 - strata variable information, [7717](#)
 - syntax, [7688](#)
 - table output, [7716](#)
 - variable bar chart, [7718](#), [7719](#)
 - variable box plot, [7718](#), [7719](#)
 - variable cloud plot, [7718](#), [7719](#)
 - variable information, [7717](#)
- standardized differences plot
 - PSMATCH procedure, [7719](#)
- standardized variable differences
 - PSMATCH procedure, [7716](#)
- strata information
 - PSMATCH procedure, [7717](#)
- strata standardized variable differences
 - PSMATCH procedure, [7717](#)
- strata variable information
 - PSMATCH procedure, [7717](#)
- variable bar chart
 - PSMATCH procedure, [7718](#), [7719](#)
- variable box plot
 - PSMATCH procedure, [7718](#), [7719](#)
- variable cloud plot
 - PSMATCH procedure, [7718](#), [7719](#)
- variable information
 - PSMATCH procedure, [7717](#)

Syntax Index

- ASSESS statement
 - PSMATCH procedure, [7691](#)
- ATEWGT= option
 - OUTPUT statement (PSMATCH), [7701](#)
- ATTWGT= option
 - OUTPUT statement (PSMATCH), [7701](#)
- BY statement
 - PSMATCH procedure, [7694](#)
- CALIPER option
 - MATCH statement (PSMATCH), [7696](#)
- CLASS statement
 - PSMATCH procedure, [7695](#)
- DATA= option
 - PROC PSMATCH statement, [7688](#)
- FREQ statement
 - PSMATCH procedure, [7695](#)
- K= option
 - MATCH statement (PSMATCH), [7698](#), [7699](#)
- KEY= option
 - PROC PSMATCH statement, [7703](#)
- KMAX= option
 - MATCH statement (PSMATCH), [7697](#), [7699](#)
- KMAXTREATED= option
 - MATCH statement (PSMATCH), [7697](#)
- KMEAN= option
 - MATCH statement (PSMATCH), [7698](#), [7699](#)
- KMIN= option
 - MATCH statement (PSMATCH), [7699](#)
- LPS option
 - ASSESS statement, [7691](#)
 - MATCH statement, [7700](#)
- LPS= option
 - OUTPUT statement (PSMATCH), [7701](#)
 - PSDATA statement (PSMATCH), [7702](#)
- MATCH statement
 - PSMATCH procedure, [7695](#)
- MATCHID= option
 - OUTPUT statement (PSMATCH), [7701](#)
- MATCHWGT= option
 - OUTPUT statement (PSMATCH), [7701](#)
- NCONTROL= option
 - MATCH statement (PSMATCH), [7698](#), [7699](#)
- NSTRATA= option
 - PROC PSMATCH statement, [7703](#)
- ORDER= option
 - MATCH statement (PSMATCH), [7698](#)
- OUT= option
 - OUTPUT statement (PSMATCH), [7701](#)
- OUTPUT statement
 - PSMATCH procedure, [7701](#)
- PCTCONTROL= option
 - MATCH statement (PSMATCH), [7698](#), [7699](#)
- PLOTS option
 - ASSESS statement, [7691](#)
- PS option
 - ASSESS statement, [7691](#)
 - MATCH statement, [7700](#)
- PS= option
 - OUTPUT statement (PSMATCH), [7702](#)
 - PSDATA statement (PSMATCH), [7702](#)
- PSMATCH procedure
 - EXACT= option, [7697](#)
 - METHOD= option, [7697](#)
 - STAT= option, [7699](#)
- PSMATCH procedure, ASSESS statement, [7691](#)
- LPS option, [7691](#)
- PLOTS option, [7691](#)
- PS option, [7691](#)
- STDDIFFDIV= option, [7693](#)
- VAR= option, [7691](#)
- VARINFO option, [7693](#)
- WEIGHT= option, [7693](#)
- PSMATCH procedure, BY statement, [7694](#)
- PSMATCH procedure, CLASS statement, [7695](#)
- PSMATCH procedure, FREQ statement, [7695](#)
- PSMATCH procedure, MATCH statement, [7695](#)
 - CALIPER option, [7696](#)
 - K= option, [7698](#), [7699](#)
 - KMAX= option, [7697](#), [7699](#)
 - KMAXTREATED= option, [7697](#)
 - KMEAN= option, [7698](#), [7699](#)
 - KMIN= option, [7699](#)
 - LPS option, [7700](#)
 - NCONTROL= option, [7698](#), [7699](#)
 - ORDER= option, [7698](#)
 - PCTCONTROL= option, [7698](#), [7699](#)
 - PS option, [7700](#)
 - SEED= option, [7698](#)
 - VAR= option, [7700](#)

- PSMATCH procedure, OUTPUT statement, [7701](#)
 - ATEWGT option, [7701](#)
 - ATTWGT option, [7701](#)
 - LPS= option, [7701](#)
 - MATCHID= option, [7701](#)
 - MATCHWGT= option, [7701](#)
 - OUT= option, [7701](#)
 - PS= option, [7702](#)
 - STRATA= option, [7702](#)
- PSMATCH procedure, PROC PSMATCH statement
 - DATA= option, [7688](#)
 - KEY= option, [7703](#)
 - NSTRATA= option, [7703](#)
 - REGION= option, [7689](#)
- PSMATCH procedure, PSDATA statement, [7702](#)
 - LPS option, [7702](#)
 - PS option, [7702](#)
 - TREATVAR= option, [7702](#)
- PSMATCH procedure, PSMODEL statement, [7702](#)
- PSMATCH procedure, STRATA statement, [7703](#)
- PSMDATA statement
 - PSMATCH procedure, [7702](#)
- PSMODEL statement
 - PSMATCH procedure, [7702](#)
- REGION= option
 - PROC PSMATCH statement, [7689](#)
- SEED= option
 - MATCH statement (PSMATCH), [7698](#)
- STDDIFFDIV= option
 - ASSESS statement (PSMATCH), [7693](#)
- STRATA statement
 - PSMATCH procedure, [7703](#)
- STRATA= option
 - OUTPUT statement (PSMATCH), [7702](#)
- TREATVAR= option
 - PSDATA statement (PSMATCH), [7702](#)
- VAR= option
 - ASSESS statement, [7691](#)
 - MATCH statement, [7700](#)
- VARINFO option
 - ASSESS statement, [7693](#)
- WEIGHT= option
 - ASSESS statement (PSMATCH), [7693](#)