

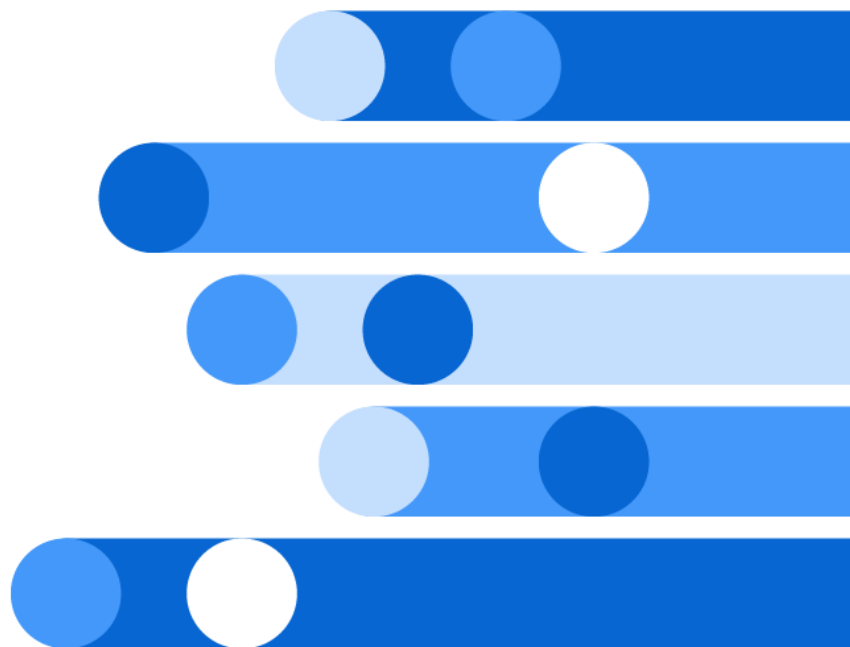


SAS/STAT 15.3[®]

User's Guide

The NPAR1WAY

Procedure



This document is an individual chapter from *SAS/STAT® 15.3 User's Guide*.

The correct bibliographic citation for this manual is as follows: SAS Institute Inc. 2023. *SAS/STAT® 15.3 User's Guide*. Cary, NC: SAS Institute Inc.

SAS/STAT® 15.3 User's Guide

Copyright © 2023, 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.

February 2025

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 [Third-Party Software Reference | SAS Support](#).

Chapter 90

The NPAR1WAY Procedure

Contents

Overview: NPAR1WAY Procedure	7194
Getting Started: NPAR1WAY Procedure	7194
Syntax: NPAR1WAY Procedure	7203
PROC NPAR1WAY Statement	7203
BY Statement	7212
CLASS Statement	7212
EXACT Statement	7213
FREQ Statement	7217
OUTPUT Statement	7218
STRATA Statement	7221
VAR Statement	7223
Details: NPAR1WAY Procedure	7223
Missing Values	7223
Tied Values	7224
Statistical Computations	7224
Simple Linear Rank Tests for Two-Sample Data	7224
One-Way ANOVA Tests	7226
Scores for Linear Rank and One-Way ANOVA Tests	7227
Stratified Analysis	7229
Hodges-Lehmann Estimation of Location Shift	7230
Fligner-Policello Test	7232
Ratio Mean Deviations (RMD) Exact Test	7234
Multiple Comparisons Based on Pairwise Rankings	7235
Empirical Distribution Function Tests	7235
Exact Tests	7237
Contents of the Output Data Set	7241
Displayed Output	7247
ODS Table Names	7255
ODS Graphics	7257
Examples: NPAR1WAY Procedure	7258
Example 90.1: Two-Sample Location Tests and Plots	7258
Example 90.2: EDF Statistics and EDF Plot	7262
Example 90.3: Exact Wilcoxon Two-Sample Test	7263
Example 90.4: Hodges-Lehmann Estimation	7265
Example 90.5: Exact Savage Multisample Test	7266
References	7267

Overview: NPAR1WAY Procedure

The NPAR1WAY procedure performs nonparametric tests for location and scale differences across a one-way classification. PROC NPAR1WAY also provides a standard analysis of variance on the raw data, empirical distribution function statistics, pairwise multiple comparison analysis, and stratified analysis.

PROC NPAR1WAY performs tests for location and scale differences based on the following rank-based scores of a response variable: Wilcoxon, median, Van der Waerden (normal), Savage, Siegel-Tukey, Ansari-Bradley, Klotz, Mood, and Conover. In addition, PROC NPAR1WAY provides tests that use the raw input data as scores. When the data are classified into two samples, tests are based on simple linear rank statistics. When the data are classified into more than two samples, tests are based on one-way analysis of variance (ANOVA) statistics. Both asymptotic and exact p -values are available for these tests. PROC NPAR1WAY also provides Hodges-Lehmann estimation of the location shift (with exact confidence limits), the Fligner-Policello test, and the RMD exact scale test for two-sample data.

PROC NPAR1WAY provides stratified analysis for two-sample data based on the following scores: Wilcoxon, median, Van der Waerden (normal), Savage, and raw data scores. Rank-based scores can be computed by using within-stratum ranks or overall ranks; strata can be weighted by stratum size or by equal weights. PROC NPAR1WAY also provides alignment by strata.

PROC NPAR1WAY computes empirical distribution function (EDF) statistics, which test whether the distribution of a variable is the same across different groups. These statistics include the Kolmogorov-Smirnov test, the Cramér-von Mises test, and the Kuiper test. Exact p -values are available for the two-sample Kolmogorov-Smirnov test.

PROC NPAR1WAY uses ODS Graphics to create graphs as part of its output. For general information about ODS Graphics, see Chapter 24, “[Statistical Graphics Using ODS](#).” For more information about the statistical graphics that PROC NPAR1WAY produces, see the `PLOTS=` option in the PROC NPAR1WAY statement and the section “[ODS Graphics](#)” on page 7257.

Getting Started: NPAR1WAY Procedure

This example illustrates how you can use PROC NPAR1WAY to perform a one-way nonparametric analysis. The data from Halverson and Sherwood (1930) consist of weight gain measurements for five different levels of gossypol additive in animal feed. Gossypol is a substance contained in cottonseed shells, and these data were collected to study the effect of gossypol on animal nutrition.

The following DATA step statements create the SAS data set Gossypol:

```
data Gossypol;
  input Dose n;
  do i=1 to n;
    input Gain @@;
    output;
  end;
  datalines;
0 16
228 229 218 216 224 208 235 229 233 219 224 220 232 200 208 232
```

```
.04 11
186 229 220 208 228 198 222 273 216 198 213
.07 12
179 193 183 180 143 204 114 188 178 134 208 196
.10 17
130 87 135 116 118 165 151 59 126 64 78 94 150 160 122 110 178
.13 11
154 130 130 118 118 104 112 134 98 100 104
;
```

The data set Gossypol contains the variable Dose, which represents the amount of gossypol additive, and the variable Gain, which represents the weight gain.

Researchers are interested in whether there is a difference in weight gain among animals receiving the different dose levels of gossypol. The following statements invoke the NPAR1WAY procedure to perform a nonparametric analysis of this problem:

```
proc npar1way data=Gossypol;
  class Dose;
  var Gain;
run;
```

The variable Dose is the CLASS variable, and the VAR statement specifies the variable Gain is the response variable. The CLASS statement is required, and you must name only one CLASS variable. You can name one or more analysis variables in the VAR statement. If you omit the VAR statement, PROC NPAR1WAY analyzes all numeric variables in the data set except the CLASS variable, the FREQ variable, and the BY variables.

When no analysis options are specified in the PROC NPAR1WAY statement, the ANOVA, WILCOXON, MEDIAN, VW, SAVAGE, and EDF options are invoked by default. The tables in the following figures show the results of these analyses.

The tables in [Figure 90.1](#) are produced by the ANOVA option. For each level of the CLASS variable Dose, PROC NPAR1WAY displays the number of observations and the mean of the analysis variable Gain. PROC NPAR1WAY displays a standard analysis of variance on the raw data. This gives the same results as the GLM and ANOVA procedures. The p -value for the F test is <0.0001 , which indicates that Dose accounts for a significant portion of the variability of the dependent variable Gain.

Figure 90.1 Analysis of Variance

The NPAR1WAY Procedure

Analysis of Variance for Variable Gain Classified by Variable Dose		
Dose	N	Mean
0	16	222.187500
0.04	11	217.363636
0.07	12	175.000000
0.1	17	120.176471
0.13	11	118.363636

Figure 90.1 continued

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Among	4	140082.986077	35020.74652	55.8143	<.0001
Within	62	38901.998997	627.45160		
Average scores were used for ties.					

The WILCOXON option produces the output in Figure 90.2. PROC NPAR1WAY first provides a summary of the Wilcoxon scores for the analysis variable Gain by class level. For each level of the CLASS variable Dose, PROC NPAR1WAY displays the following information: number of observations, sum of the Wilcoxon scores, expected sum under the null hypothesis of no difference among class levels, standard deviation under the null hypothesis, and mean score.

Next PROC NPAR1WAY displays the one-way ANOVA statistic, which for Wilcoxon scores is known as the Kruskal-Wallis test. The statistic is 52.6656, with 4 degrees of freedom, which is the number of class levels minus 1. The p -value (probability of a larger statistic under the null hypothesis) is <0.0001. This leads to rejection of the null hypothesis that there is no difference in location for Gain among the levels of Dose. This p -value is asymptotic, computed from the asymptotic chi-square distribution of the test statistic. For certain data sets it might also be useful to compute the exact p -value—for example, for small data sets or for data sets that are sparse, skewed, or heavily tied. You can use the EXACT statement to request exact p -values for any of the location or scale tests available in PROC NPAR1WAY.

Figure 90.2 Wilcoxon Score Analysis

Wilcoxon Scores (Rank Sums) for Variable Gain Classified by Variable Dose					
Dose	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
0	16	890.50	544.0	67.978966	55.656250
0.04	11	555.00	374.0	59.063588	50.454545
0.07	12	395.50	408.0	61.136622	32.958333
0.1	17	275.50	578.0	69.380741	16.205882
0.13	11	161.50	374.0	59.063588	14.681818
Average scores were used for ties.					

Kruskal-Wallis Test			
Chi-Square	DF	Pr >	ChiSq
52.6656	4	<.0001	

Figure 90.3 through Figure 90.5 display the analyses produced by the MEDIAN, VW, and SAVAGE options. For each score type, PROC NPAR1WAY provides a summary of scores and the one-way ANOVA statistic, as previously described for Wilcoxon scores. Other score types available in PROC NPAR1WAY are Siegel-Tukey, Ansari-Bradley, Klotz, and Mood, which can be used to test for scale differences. Conover scores can be used to test for differences in both location and scale. Additionally, you can specify the SCORES=DATA option, which uses the input data as scores. This option gives you the flexibility to construct any scores for your data with the DATA step and then analyze these scores with PROC NPAR1WAY.

Figure 90.3 Median Score Analysis

Median Scores (Number of Points Above Median) for Variable Gain Classified by Variable Dose					
Dose	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
0	16	16.0	7.880597	1.757902	1.00
0.04	11	11.0	5.417910	1.527355	1.00
0.07	12	6.0	5.910448	1.580963	0.50
0.1	17	0.0	8.373134	1.794152	0.00
0.13	11	0.0	5.417910	1.527355	0.00
Average scores were used for ties.					

Median One-Way Analysis		
Chi-Square	DF	Pr > ChiSq
54.1765	4	<.0001

Figure 90.4 Van der Waerden (Normal) Score Analysis

Van der Waerden Scores (Normal) for Variable Gain Classified by Variable Dose					
Dose	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
0	16	16.116474	0.0	3.325957	1.007280
0.04	11	8.340899	0.0	2.889761	0.758264
0.07	12	-0.576674	0.0	2.991186	-0.048056
0.1	17	-14.688921	0.0	3.394540	-0.864054
0.13	11	-9.191777	0.0	2.889761	-0.835616
Average scores were used for ties.					

Van der Waerden One-Way Analysis		
Chi-Square	DF	Pr > ChiSq
47.2972	4	<.0001

Figure 90.5 Savage Score Analysis

Savage Scores (Exponential) for Variable Gain Classified by Variable Dose					
Dose	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
0	16	16.074391	0.0	3.385275	1.004649
0.04	11	7.693099	0.0	2.941300	0.699373
0.07	12	-3.584958	0.0	3.044534	-0.298746
0.1	17	-11.979488	0.0	3.455082	-0.704676
0.13	11	-8.203044	0.0	2.941300	-0.745731
Average scores were used for ties.					

Savage One-Way Analysis		
Chi-Square	DF	Pr > ChiSq
39.4908	4	<.0001

The tables in Figure 90.6 display the empirical distribution function statistics, comparing the distribution of Gain for the different levels of Dose. These tables are produced by the EDF option, and they include Kolmogorov-Smirnov statistics and Cramér–von Mises statistics.

Figure 90.6 Empirical Distribution Function Analysis

Kolmogorov-Smirnov Test for Variable Gain Classified by Variable Dose			
Dose	N	EDF at Maximum	Deviation from Mean at Maximum
0	16	0.000000	-1.910448
0.04	11	0.000000	-1.584060
0.07	12	0.333333	-0.499796
0.1	17	1.000000	2.153861
0.13	11	1.000000	1.732565
Total	67	0.477612	
Maximum Deviation Occurred at Observation 36 Value of Gain at Maximum = 178.0			
Kolmogorov-Smirnov Statistics (Asymptotic)			
KS	0.457928	KSa	3.748300
Cramer-von Mises Test for Variable Gain Classified by Variable Dose			
Dose	N	Summed Deviation from Mean	
0	16	2.165210	
0.04	11	0.918280	
0.07	12	0.348227	
0.1	17	1.497542	
0.13	11	1.335745	
Cramer-von Mises Statistics (Asymptotic)			
CM	0.093508	CMA	6.265003

PROC NPAR1WAY uses ODS Graphics to create graphs as part of its output. The following statements produce a box plot of Wilcoxon scores for Gain classified by Dose. ODS Graphics must be enabled before producing graphs.

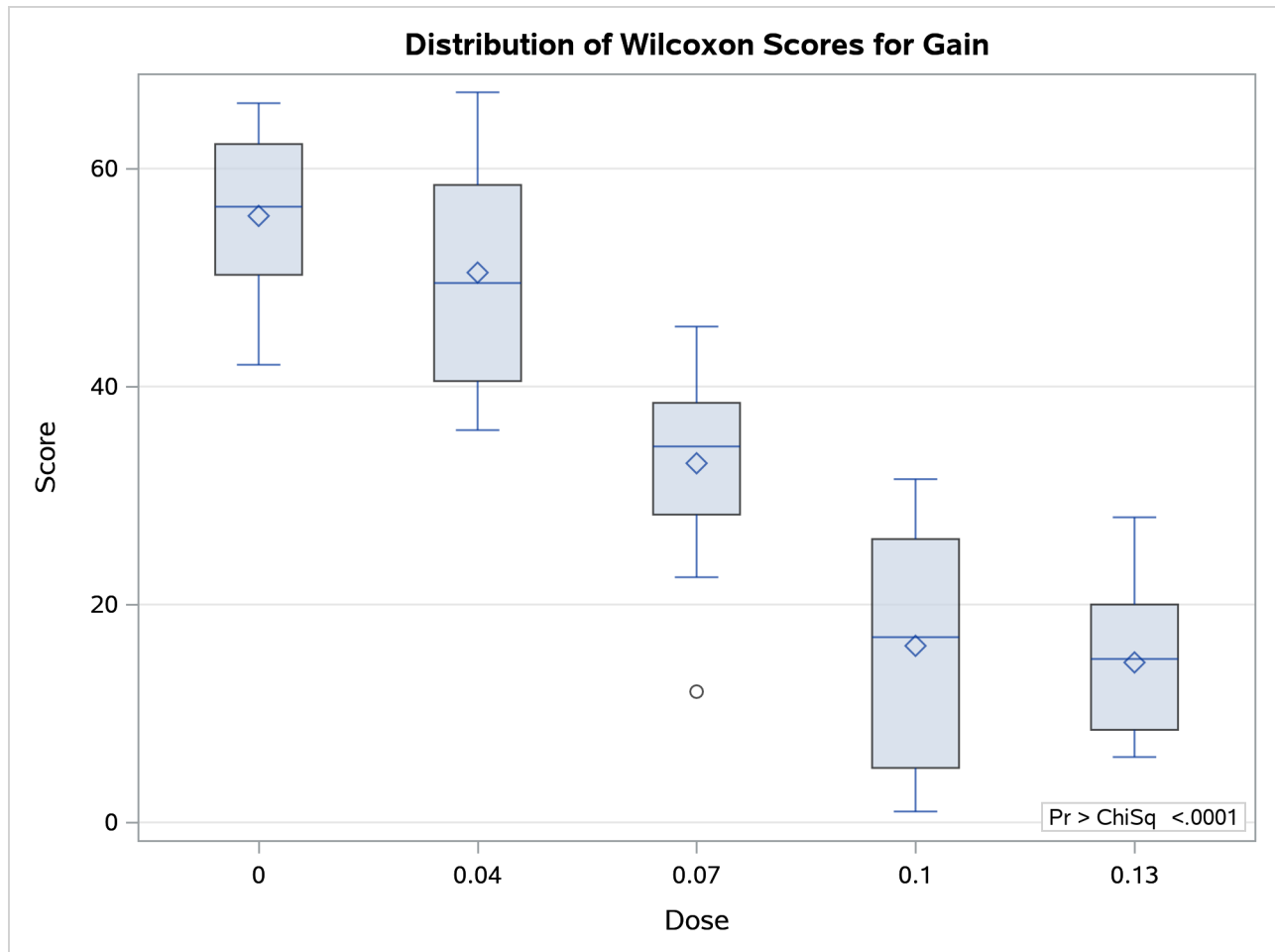
```
ods graphics on;
proc npar1way data=Gossypol plots(only)=wilcoxonboxplot;
  class Dose;
  var Gain;
run;
ods graphics off;
```

Figure 90.7 displays the box plot of Wilcoxon scores. This graph corresponds to the Wilcoxon scores analysis shown in Figure 90.2. To remove the p -value from the box plot display, you can specify the NOSTATS plot option in parentheses after the WILCOXONBOXPLOT option.

Box plots are available for all PROC NPAR1WAY score types except median scores, which are displayed

in a stacked bar chart. If ODS Graphics is enabled but you do not specify the PLOTS= option, PROC NPAR1WAY produces all plots that are associated with the analyses that you request.

Figure 90.7 Box Plot of Wilcoxon Scores



In the preceding example, the CLASS variable Dose has five levels, and the analyses examine possible differences among these five levels (samples). The following statements invoke the NPAR1WAY procedure to perform a nonparametric analysis of the two lowest levels of Dose:

```
proc npar1way data=Gossypol;
  where Dose <= .04;
  class Dose;
  var Gain;
run;
```

The tables in the following figures show the results of this two-sample analysis. The tables in [Figure 90.8](#) are produced by the ANOVA option.

Figure 90.8 Analysis of Variance for Two-Sample Data

The NPAR1WAY Procedure					
Analysis of Variance for Variable Gain Classified by Variable Dose					
Dose		N		Mean	
0		16		222.187500	
0.04		11		217.363636	

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Among	1	151.683712	151.683712	0.5587	0.4617
Within	25	6786.982955	271.479318		

Average scores were used for ties.

Figure 90.9 displays the output produced by the WILCOXON option. PROC NPAR1WAY provides a summary of the Wilcoxon scores for the analysis variable Gain for each of the two class levels. Because there are two levels, PROC NPAR1WAY displays the two-sample test, which is based on the simple linear rank statistic with Wilcoxon scores. The normal approximation includes a continuity correction. To remove the continuity correction, you can specify the CORRECT=NO option. PROC NPAR1WAY also gives a t approximation for the Wilcoxon two-sample test. Like the multisample analysis, PROC NPAR1WAY computes a one-way ANOVA statistic, which for Wilcoxon scores is known as the Kruskal-Wallis test. All these p -values show no difference in Gain for the two Dose levels at the 0.05 level of significance.

Figure 90.10 through Figure 90.12 display the two-sample analyses produced by the MEDIAN, VW, and SAVAGE options.

Figure 90.9 Wilcoxon Two-Sample Analysis

Wilcoxon Scores (Rank Sums) for Variable Gain Classified by Variable Dose					
Dose	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
0	16	253.50	224.0	20.221565	15.843750
0.04	11	124.50	154.0	20.221565	11.318182

Average scores were used for ties.

Wilcoxon Two-Sample Test					
Statistic	Z	Pr < Z	Pr > Z	t Approximation	
				Pr < Z	Pr > Z
124.5000	-1.4341	0.0758	0.1515	0.0817	0.1635

Z includes a continuity correction of 0.5.

Kruskal-Wallis Test			
Chi-Square	DF	Pr > ChiSq	
2.1282	1	0.1446	

Figure 90.10 Median Two-Sample Analysis

Median Scores (Number of Points Above Median) for Variable Gain Classified by Variable Dose					
Dose	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
0	16	9.0	7.703704	1.299995	0.562500
0.04	11	4.0	5.296296	1.299995	0.363636
Average scores were used for ties.					

Median Two-Sample Test			
Statistic	Z	Pr < Z	Pr > Z
4.0000	-0.9972	0.1593	0.3187

Median One-Way Analysis			
Chi-Square	DF	Pr > ChiSq	
0.9943	1	0.3187	

Figure 90.11 Van der Waerden (Normal) Two-Sample Analysis

Van der Waerden Scores (Normal) for Variable Gain Classified by Variable Dose					
Dose	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
0	16	3.346520	0.0	2.320336	0.209157
0.04	11	-3.346520	0.0	2.320336	-0.304229
Average scores were used for ties.					

Van der Waerden Two-Sample Test			
Statistic	Z	Pr < Z	Pr > Z
-3.3465	-1.4423	0.0746	0.1492

Van der Waerden One-Way Analysis			
Chi-Square	DF	Pr > ChiSq	
2.0801	1	0.1492	

Figure 90.12 Savage Two-Sample Analysis

Savage Scores (Exponential) for Variable Gain Classified by Variable Dose					
Dose	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
0	16	1.834554	0.0	2.401839	0.114660
0.04	11	-1.834554	0.0	2.401839	-0.166778
Average scores were used for ties.					

Savage Two-Sample Test			
Statistic	Z	Pr < Z	Pr > Z
-1.8346	-0.7638	0.2225	0.4450

Figure 90.12 *continued*

Savage One-Way Analysis		
Chi-Square	DF	Pr > ChiSq
0.5834	1	0.4450

The tables in Figure 90.13 display the empirical distribution function statistics, comparing the distribution of Gain for the two levels of Dose. The p -value for the Kolmogorov-Smirnov two-sample test is 0.6199, which indicates no rejection of the null hypothesis that the Gain distributions are identical for the two levels of Dose.

Figure 90.13 EDF Two-Sample Analysis

Kolmogorov-Smirnov Test for Variable Gain Classified by Variable Dose			
Dose	N	EDF at Maximum	Deviation from Mean at Maximum
0	16	0.250000	-0.481481
0.04	11	0.545455	0.580689
Total	27	0.370370	
Maximum Deviation Occurred at Observation 4 Value of Gain at Maximum = 216.0			
Kolmogorov-Smirnov Two-Sample Test (Asymptotic)			
KS	0.145172	D	0.295455
KSa	0.754337	Pr > KSa	0.6199
Cramer-von Mises Test for Variable Gain Classified by Variable Dose			
Dose	N	Summed Deviation from Mean	
0	16	0.098638	
0.04	11	0.143474	
Cramer-von Mises Statistics (Asymptotic)			
CM	0.008967	CMa	0.242112
Kuiper Test for Variable Gain Classified by Variable Dose			
Dose	N	Deviation from Mean	
0	16	0.090909	
0.04	11	0.295455	
Kuiper Two-Sample Test (Asymptotic)			
K	0.386364	Ka	0.986440
		Pr > Ka	0.8383

Syntax: NPAR1WAY Procedure

The following statements are available in the NPAR1WAY procedure:

```
PROC NPAR1WAY < options > ;
  BY variables ;
  CLASS variable ;
  EXACT statistic-options < / computation-options > ;
  FREQ variable ;
  OUTPUT < OUT=SAS-data-set > < output-options > ;
  STRATA variables < / options > ;
  VAR variables ;
```

The **PROC NPAR1WAY** statement invokes the NPAR1WAY procedure. The **PROC NPAR1WAY** and **CLASS** statements are required. Table 90.1 summarizes the statements available in the NPAR1WAY procedure.

The rest of this section provides detailed syntax information for each of these statements, beginning with the **PROC NPAR1WAY** statement. The remaining statements are described in alphabetical order.

Table 90.1 Summary of PROC NPAR1WAY Statements

Statement	Description
BY	Provides separate analyses for each BY group
CLASS	Identifies the classification variable
EXACT	Requests exact tests
FREQ	Identifies a frequency variable
OUTPUT	Requests an output data set
STRATA	Identifies strata variables
VAR	Identifies analysis variables

PROC NPAR1WAY Statement

```
PROC NPAR1WAY < options > ;
```

The **PROC NPAR1WAY** statement invokes the NPAR1WAY procedure. Optionally, it identifies the input data set and requests analyses. By default, the procedure uses the most recently created SAS data set. If you do not specify any analysis options or a **STRATA** statement, **PROC NPAR1WAY** performs an analysis of variance (**ANOVA** option), tests for location differences (**WILCOXON**, **MEDIAN**, **SAVAGE**, and **VW** options), and empirical distribution function tests (**EDF** option).

Table 90.2 lists the *options* available in the **PROC NPAR1WAY** statement. Descriptions of the *options* follow in alphabetical order.

Table 90.2 PROC NPAR1WAY Statement Options

Option	Description
Input Data Set	
DATA=	Names the input SAS data set
Missing Values	
MISSING	Treats missing values as a valid level
Analyses	
AB	Requests analysis of Ansari-Bradley scores
ANOVA	Requests standard analysis of variance
CONOVER	Requests analysis of Conover scores
D	Requests one-sided Kolmogorov-Smirnov statistics
DSCF	Requests pairwise multiple comparison analysis
EDF KS	Requests empirical distribution function statistics
FP	Requests the Fligner-Policello test
HL	Requests Hodges-Lehmann estimation for two-sample data
KLOTZ	Requests analysis of Klotz scores
MEDIAN	Requests analysis of median scores
MOOD	Requests analysis of Mood scores
RMD	Requests the ratio mean deviations (RMD) test
SAVAGE	Requests analysis of Savage scores
SCORES=DATA	Requests analysis of input data scores
ST	Requests analysis of Siegel-Tukey scores
VW NORMAL	Requests analysis of Van der Waerden (normal) scores
WILCOXON	Requests analysis of Wilcoxon scores
Analysis Details	
ADJUST	Requests median adjustment for scale analyses
ALIGN=STRATA	Requests alignment by strata
ALPHA=	Specifies the confidence level for Hodges-Lehmann estimation
CORRECT=NO	Suppresses the continuity correction for Wilcoxon and Siegel-Tukey two-sample tests
Displayed Output	
NOPRINT	Suppresses the display of all output
TABLES=RESTORE	Displays score test tables in factoid form
VARINFO	Displays variable information
Plots	
PLOTS=	Requests plots from ODS Graphics

You can specify the following *options*:

AB <(ADJUST)>

requests an analysis of Ansari-Bradley scores. For more information, see the section “[Ansari-Bradley Scores](#)” on page 7228. The ADJUST suboption subtracts class medians from the input observations before performing the analysis.

ADJUST

adjusts for location differences among classes before performing tests for scale differences. This option applies to the following analyses: Ansari-Bradley (**AB**), Klotz (**KLOTZ**), Mood (**MOOD**), and Siegel-Tukey (**ST**). If you specify the **ADJUST** option, PROC NPAR1WAY subtracts the corresponding class median from each observation before scoring the data and performing the scale tests that you request.

You can also request adjustment for an individual scale test by specifying **ADJUST** in parentheses after the scale test option; for example, you can specify **ST(ADJUST)** to adjust by class medians before performing the Siegel-Tukey test.

ALIGN=STRATA <(option)>

aligns the response variable values by strata (Hodges and Lehmann 1962; Mehrotra, Lu, and Li 2010) before performing the analyses that you request. When you specify this option, you must also specify a **STRATA** statement to define the strata. PROC NPAR1WAY subtracts the stratum location measure from each response value in the stratum. You can specify a location measure *option* in parentheses after the **ALIGN=STRATA** option; by default, PROC NPAR1WAY uses the stratum median (**MEDIAN**) as the stratum location measure.

You can specify one of the following *options*:

HL

specifies the Hodges-Lehmann shift as the stratum location measure. The stratum Hodges-Lehmann shift is the median of all pairwise (Walsh) averages of response values in the stratum.

MEAN

specifies the mean (average) of the response values in the stratum as the stratum location measure.

MEDIAN

specifies the median of the response values in the stratum as the stratum location measure. This is the default location measure for **ALIGN=STRATA**.

ALPHA= α

specifies the level of the confidence limits for the Hodges-Lehmann location shift, which you can request by specifying the **HL** option. The value of α must be between 0 and 1; a confidence level of α produces $100(1 - \alpha)\%$ confidence limits. By default, **ALPHA=0.05**, which produces 95% confidence limits for the location shift.

ANOVA

requests a standard analysis of variance on the raw data.

CONOVER

requests an analysis of Conover scores. For more information, see the section “[Conover Scores](#)” on page 7229.

CORRECT=NO

suppresses the continuity correction for the Wilcoxon two-sample test and the Siegel-Tukey two-sample test. For more information, see the section “[Continuity Correction](#)” on page 7226.

D

requests the one-sided Kolmogorov-Smirnov $D+$ and $D-$ statistics and their asymptotic p -values, in addition to the two-sided D statistic that the [EDF](#) option produces for two-sample data. The **D** option invokes the [EDF](#) option. If you request exact Kolmogorov-Smirnov tests by specifying the **KS** option in the **EXACT** statement for two-sample data, PROC NPAR1WAY provides $D+$ and $D-$ by default. For more information about Kolmogorov-Smirnov statistics, see the section “[Empirical Distribution Function Tests](#)” on page 7235.

DATA=SAS-data-set

names the *SAS-data-set* to be analyzed by PROC NPAR1WAY. If you omit the **DATA=** option, the procedure uses the most recently created SAS data set.

DSCF

requests the Dwass, Steel, Critchlow-Fligner multiple comparison procedure, which is based on pairwise two-sample rankings. This procedure is available for multisample data, where the number of **CLASS** variable levels is greater than 2. For more information, see the section “[Multiple Comparisons Based on Pairwise Rankings](#)” on page 7235.

EDF**KS**

requests statistics based on the empirical distribution function. These include the Kolmogorov-Smirnov and Cramér-von Mises tests and, if there are only two classification levels, the Kuiper test. For more information, see the section “[Empirical Distribution Function Tests](#)” on page 7235.

The **EDF** option produces the Kolmogorov-Smirnov D statistic for two-sample data. You can request the one-sided $D+$ and $D-$ statistics for two-sample data by specifying the **D** option.

FP <(REFCLASS=class-number | 'class-value')>

requests the Fligner-Policello test for two-sample data. For more information, see the section “[Fligner-Policello Test](#)” on page 7232.

The **REFCLASS=** suboption specifies which of the two **CLASS** variable levels (samples) to use as the reference class X in the difference between class placement sums, $P_y - P_x$. **REFCLASS=1** specifies the first class that is listed in the “Fligner-Policello Placements” table, and **REFCLASS=2** specifies the second class. (The table displays class levels in the order in which they appear in the input data set.) **REFCLASS='class-value'** identifies the reference class by the formatted value of the **CLASS** variable.

By default, PROC NPAR1WAY uses the larger of the two classes as the reference class X . If both classes have the same number of observations, PROC NPAR1WAY uses the class that appears second in the “Fligner-Policello Placements” table as the reference class.

HL <(REFCLASS=class-number | 'class-value')>

requests Hodges-Lehmann estimation of the location shift for two-sample data. This option also provides asymptotic confidence limits for the location shift (which are sometimes known as Moses confidence limits). For more information, see the section “[Hodges-Lehmann Estimation of Location Shift](#)” on page 7230. You can specify the confidence level in the **ALPHA=** option. By default, **ALPHA=0.05**, which produces 95% confidence limits for the location shift.

The **REFCLASS=** suboption specifies which of the two **CLASS** variable levels (samples) to use as the reference class X in the location shift, $Y - X$. **REFCLASS=1** specifies the first class that is listed in the “Wilcoxon Scores” table, and **REFCLASS=2** specifies the second class. (The table displays class

levels in the order in which they appear in the input data set.) REFCLASS='class-value' identifies the reference class by the formatted value of the **CLASS** variable.

By default, PROC NPAR1WAY uses the larger of the two classes as the reference class *X*. If both classes have the same number of observations, PROC NPAR1WAY uses the class that appears second in the “Wilcoxon Scores” table as the reference class.

KLOTZ <(ADJUST)>

requests an analysis of Klotz scores. For more information, see the section “[Klotz Scores](#)” on page 7228. If you specify the ADJUST suboption, PROC NPAR1WAY subtracts the corresponding class median from each observation before performing the analysis.

MEDIAN

requests an analysis of median scores. When there are two classification levels, this option produces the two-sample median test. When there are more than two samples, this option produces the multisample median test, which is also known as the Brown-Mood test. For more information, see the section “[Median Scores](#)” on page 7227.

MISSING

treats missing values as a valid nonmissing level for **CLASS** and **STRATA** variables. By default, PROC NPAR1WAY excludes an observations from the analysis if the value of the CLASS variable is missing. By default, PROC NPAR1WAY excludes an observations from the stratified analysis if the value of the STRATA variable is missing. For more information, see the section “[Missing Values](#)” on page 7223.

MOOD <(ADJUST)>

requests an analysis of Mood scores. For more information, see the section “[Mood Scores](#)” on page 7229. If you specify the ADJUST suboption, PROC NPAR1WAY subtracts the corresponding class median from each observation before performing the analysis.

NOPRINT

suppresses the display of all output. You can use the NOPRINT option when you only want to create an output data set. This option temporarily disables the Output Delivery System (ODS). For more information, see Chapter 23, “[Using the Output Delivery System](#).”

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

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

controls the plots that are produced through ODS Graphics. *Plot-requests* specify the plots to produce, and *plot-options* control the appearance and content of the plots. You can specify *plot-options* in parentheses after a *plot-request*. A *global-plot-option* applies to all plots for which it is available, unless it is altered by a specific *plot-option*. You can specify *global-plot-options* in parentheses after the PLOTS option.

When you specify only one *plot-request*, you can omit the parentheses around the request. For example:

```
plots=all
plots=wilcoxonboxplot
plots=(wilcoxonboxplot edfplot)
plots(only)=(medianplot normalboxplot)
```

ODS Graphics must be enabled before plots can be requested. For example:

```
ods graphics on;
proc npar1way plots=wilcoxonboxplot;
    variable response;
    class treatment;
run;
ods graphics off;
```

For more information about enabling and disabling ODS Graphics, see the section “[Enabling and Disabling ODS Graphics](#)” on page 663 in Chapter 24, “[Statistical Graphics Using ODS](#).”

If ODS Graphics is enabled but you do not specify the PLOTS= option, PROC NPAR1WAY produces all plots that are associated with the analyses that you request. If you request a plot by specifying the PLOTS= option but do not request the corresponding analysis, PROC NPAR1WAY automatically invokes that analysis. For example, if you specify PLOTS=CONOVERBOXPLOT but do not also specify the CONOVER option in the PROC NPAR1WAY statement, PROC NPAR1WAY produces the Conover scores analysis in addition to the box plot.

You can suppress default plots and request specific plots by using the **PLOTS(ONLY)=** option; **PLOTS(ONLY)=(*plot-requests*)** produces only the plots that are specified as *plot-requests*. You can suppress all plots by specifying the **PLOTS=NONE** option. The PLOTS= option has no effect when you specify the **NOPRINT** option.

PROC NPAR1WAY provides box plots of scored data, median plots, and empirical distribution plots. See [Figure 90.7](#), [Output 90.1.2](#), [Output 90.1.4](#), and [Output 90.2.2](#) for examples of plots that PROC NPAR1WAY produces. For general information about ODS Graphics, see Chapter 24, “[Statistical Graphics Using ODS](#).”

Global Plot Options

A *global-plot-option* applies to all plots for which the option is available unless it is altered by a specific *plot-option*. You can specify the following *global-plot-options* in parentheses after the PLOTS option. You cannot specify both the STATS and the NOSTATS *global-plot-options* in the same statement.

NOSTATS

suppresses the *p*-values that are displayed by default in the plots.

ONLY

suppresses the default plots and requests only the plots that are specified as *plot-requests*.

STATS

displays *p*-values in the plots. This is the default.

Plot Requests

You can specify the following *plot-requests*:

ABBOXPLOT**AB**

requests a box plot of Ansari-Bradley scores. This plot is associated with the Ansari-Bradley analysis, which you request by specifying the **AB** option.

ALL

requests all plots that are associated with the specified analyses. This is the default if you do not specify the **ONLY** *global-plot-option*.

ANOVABOXPLOT**ANOVA**

requests a box plot of the raw data. This plot is associated with the analysis of variance based on the raw data, which you request by specifying the **ANOVA** option.

CONOVERBOXPLOT**CONOVER**

requests a box plot of Conover scores. This plot is associated with the Conover analysis, which you request by specifying the **CONOVER** option.

DATASCORESBXPLOT**DATASCORES**

requests a box plot of raw data scores. This plot is associated with the analysis that uses input data as scores, which you request by specifying the **SCORES=DATA** option.

EDFPLOT**EDF**

requests an empirical distribution plot. This plot is associated with the analyses based on the empirical distribution function, which you request by specifying the **EDF** option.

FPBOXPLOT**FP**

requests a box plot of Fligner-Policello placements. This plot is associated with the Fligner-Policello analysis, which you request by specifying the **FP** option.

KLOTZBOXPLOT**KLOTZ**

requests a box plot of Klotz scores. This plot is associated with the Klotz analysis, which you request by specifying the **KLOTZ** option.

MEDIANPLOT**MEDIAN**

requests a stacked bar chart showing the frequencies above and below the overall median. This plot is associated with the median score analysis, which you request by specifying the **MEDIAN** option.

MOODBOXPLOT**MOOD**

requests a box plot of Mood scores. This plot is associated with the Mood analysis, which you request by specifying the **MOOD** option.

NONE

suppresses all plots.

SAVAGEBOXPLOT**SAVAGE**

requests a box plot of Savage scores. This plot is associated with the Savage analysis, which you request by specifying the **SAVAGE** option.

STBOXPLOT**ST**

requests a box plot of Siegel-Tukey scores. This plot is associated with the Siegel-Tukey analysis, which you request by specifying the **ST** option.

VWBOXPLOT**VW****NORMALBOXPLOT****NORMAL**

requests a box plot of Van der Waerden (normal) scores. This plot is associated with the Van der Waerden analysis, which you request by specifying the **VW** (**NORMAL**) option.

WILCOXONBOXPLOT**WILCOXON**

requests a box plot of Wilcoxon scores. This plot is associated with the Wilcoxon analysis, which you request by specifying the **WILCOXON** option.

Plot Options

The following *plot-options* are available for any *plot-request*. You cannot specify both the **STATS** and the **NOSTATS** *plot-options* for the same plot. If you specify both the **NOSTATS** *global-plot-option* and the **STATS** *plot-option* for an individual *plot-request*, the **STATS** *plot-option* overrides the *global-plot-option* and displays statistics in the individual plot.

NOSTATS

suppresses the *p*-values that are displayed in the plot by default.

STATS

displays *p*-values in the plot. This is the default.

RMD

requests the ratio mean deviations (RMD) exact test of scale for two-sample data. For more information, see the section “**Ratio Mean Deviations (RMD) Exact Test**” on page 7234.

You can also request the RMD test by specifying the **RMD** option in the **EXACT** statement, which enables you to specify *computation-options* and request the following additional statistics: the exact mid *p*-value, the exact point probability, and Monte Carlo estimates of the exact *p*-values.

SAVAGE

requests an analysis of Savage scores. For more information, see the section “[Savage Scores](#)” on page 7228.

SCORES=DATA <(ADJUST)>**PERM <(ADJUST)>**

requests an analysis that uses input data as scores. This option gives you the flexibility to construct any scores for your data with the DATA step and then analyze these scores with PROC NPAR1WAY. For more information, see the section “[Scores for Linear Rank and One-Way ANOVA Tests](#)” on page 7227.

To produce the two-sample permutation test that is known as Pitman’s test, provide raw (unscored) data in the input data set and specify the SCORES=DATA option in the EXACT statement. For more information, see the section “[Exact Tests](#)” on page 7237.

If you specify the ADJUST suboption, PROC NPAR1WAY subtracts the corresponding class median from each observation before performing the analysis.

ST <(ADJUST)>

requests an analysis of Siegel-Tukey scores. For more information, see the section “[Siegel-Tukey Scores](#)” on page 7228. If you specify the ADJUST suboption, PROC NPAR1WAY subtracts the corresponding class median from each observation before performing the analysis.

TABLES=RESTORE

displays score test tables in factoid (label-value) format, which is the format of these tables in releases before SAS/STAT 15.1. By default (beginning in SAS/STAT 15.1), PROC NPAR1WAY displays score test tables in tabular format.

Score test tables include the two-sample test, one-way analysis, and exact test tables that PROC NPAR1WAY produces when you specify any of the following options in the PROC NPAR1WAY or EXACT statement: AB, CONOVER, FP, KLOTZ, MEDIAN, MOOD, SAVAGE, SCORES=DATA, ST, VW | NORMAL, or WILCOXON. For more information, see the sections “[Displayed Output](#)” on page 7247 and “[ODS Table Names](#)” on page 7255.

VARINFO <(TIES)>

displays a “Variable Information” table for each analysis variable that you specify in the VAR statement. This table provides the following data summary information: number of observations and number of classes used in the analysis, number of missing observations and number of empty classes, sum of frequencies used (if you specify a FREQ statement), and number of strata (if you specify the ALIGN=STRATA option).

If you specify the TIES suboption, the “Variable Information” table also includes the number of variable levels and the number of tied observations.

VW**NORMAL**

requests an analysis of Van der Waerden (normal) scores. For more information, see the section “[Van der Waerden \(Normal\) Scores](#)” on page 7227.

WILCOXON

requests an analysis of Wilcoxon scores. When there are two classification levels (samples), this option produces the Wilcoxon rank-sum test. For any number of classification levels, this option produces the Kruskal-Wallis test. For more information, see the section “[Wilcoxon Scores](#)” on page 7227.

BY Statement

BY *variables* ;

You can specify a BY statement in PROC NPAR1WAY 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 in the NPAR1WAY 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).

For more information about BY-group processing, see the discussion in *[SAS Language Reference: Concepts](#)*. For more information about the DATASETS procedure, see the discussion in the *[Base SAS Procedures Guide](#)*.

CLASS Statement

CLASS *variable* ;

The CLASS statement, which is required, names one and only one classification variable. The variable can be character or numeric. The CLASS variable identifies groups (samples) in the data, and PROC NPAR1WAY provides analyses to examine differences among these groups. There can be two or more groups in the data.

EXACT Statement

EXACT *statistic-options* < / *computation-options* > ;

The EXACT statement requests exact tests and confidence limits for selected statistics. The *statistic-options* identify which statistics to compute, and the *computation-options* specify options for computing exact statistics. For more information, see the section “[Exact Tests](#)” on page 7237.

NOTE: PROC NPAR1WAY computes exact tests by using fast and efficient algorithms that are superior to direct enumeration. Exact tests are appropriate when a data set is small, sparse, skewed, or heavily tied. For some large problems, computation of exact tests might require a large amount of time and memory. Consider using asymptotic tests for such problems. Alternatively, when asymptotic methods might not be sufficient for such large problems, consider using Monte Carlo estimation of exact *p*-values. You can request Monte Carlo estimation by specifying the **MC** *computation-option* in the EXACT statement. For more information, see the section “[Computational Resources](#)” on page 7239.

Statistic Options

The *statistic-options* specify which exact tests to compute. [Table 90.3](#) lists the available *statistic-options* and the exact statistics that are computed. Descriptions of the *statistic-options* follow [Table 90.3](#) in alphabetical order.

PROC NPAR1WAY provides exact *p*-values for the following nonparametric location and scale tests: Wilcoxon, median, Van der Waerden (normal), Savage, Siegel-Tukey, Ansari-Bradley, Klotz, Mood, and Conover. The procedure also provides exact *p*-values for tests that use the raw input data as scores. These exact tests are available when the data are classified into two levels (two-sample tests) and when the data are classified into more than two levels (multisample tests). The two-sample exact tests are based on simple linear rank statistics, and the multisample exact tests are based on one-way ANOVA statistics.

For two-sample data, PROC NPAR1WAY also provides the exact Kolmogorov-Smirnov test, the exact ratio mean deviations (RMD) test, and exact Hodges-Lehmann confidence limits for the location shift.

If you list no *statistic-options* in the EXACT statement, PROC NPAR1WAY computes all available exact *p*-values for those tests that you request in the PROC NPAR1WAY statement.

Table 90.3 EXACT Statement Statistic Options

Statistic Option	Exact Test
AB	Ansari-Bradley test
CONOVER	Conover test
HL	Hodges-Lehmann confidence limits
KLOTZ	Klotz test
KS EDF	Two-sample Kolmogorov-Smirnov test
MEDIAN	Median test
MOOD	Mood test
RMD	Ratio mean deviations (RMD) test
SAVAGE	Savage test
SCORES=DATA	Test with input data as scores
ST	Siegel-Tukey test
VW NORMAL	Van der Waerden (normal scores) test
WILCOXON	Wilcoxon test for two-sample data or Kruskal-Wallis test for multisample data

You can specify the following *statistic-options*:

AB

requests an exact Ansari-Bradley test. For more information, see the sections “[Ansari-Bradley Scores](#)” on page 7228 and “[Exact Tests](#)” on page 7237. The **AB** option in the PROC NPAR1WAY statement provides Ansari-Bradley score analysis and asymptotic tests.

CONOVER

requests an exact Conover test. For more information, see the sections “[Conover Scores](#)” on page 7229 and “[Exact Tests](#)” on page 7237. The **CONOVER** option in the PROC NPAR1WAY statement provides Conover score analysis and asymptotic tests.

HL

requests exact Hodges-Lehmann confidence limits for the location shift for two-sample data. For more information, see the section “[Hodges-Lehmann Estimation of Location Shift](#)” on page 7230. The **HL** option in the PROC NPAR1WAY statement provides asymptotic Hodges-Lehmann confidence limits.

You can specify the level of the confidence limits in the **ALPHA=** option in the PROC NPAR1WAY statement. By default, ALPHA=0.05, which produces 95% confidence limits for the location shift.

KLOTZ

requests an exact Klotz test. For more information, see the sections “[Klotz Scores](#)” on page 7228 and “[Exact Tests](#)” on page 7237. The **KLOTZ** option in the PROC NPAR1WAY statement provides Klotz score analysis and asymptotic tests.

KS**EDF**

requests an exact Kolmogorov-Smirnov two-sample test. For more information, see the section “[Empirical Distribution Function Tests](#)” on page 7235. The **EDF** option in the PROC NPAR1WAY statement provides the asymptotic Kolmogorov-Smirnov test and other statistics that are based on

the empirical distribution function. The **D** option in the PROC NPAR1WAY statement provides the asymptotic one-sided Kolmogorov-Smirnov tests for two-sample data.

MEDIAN

requests an exact median test. For more information, see the sections “[Median Scores](#)” on page 7227 and “[Exact Tests](#)” on page 7237. The **MEDIAN** option in the PROC NPAR1WAY statement provides median score analysis and asymptotic tests.

MOOD

requests an exact Mood test. For more information, see the sections “[Mood Scores](#)” on page 7229 and “[Exact Tests](#)” on page 7237. The **MOOD** option in the PROC NPAR1WAY statement provides Mood score analysis and asymptotic tests.

RMD

requests the ratio mean deviations (RMD) exact test for two-sample data. For more information, see the sections “[Ratio Mean Deviations \(RMD\) Exact Test](#)” on page 7234 and “[Exact Tests](#)” on page 7237.

SAVAGE

requests an exact Savage test. For more information, see the sections “[Savage Scores](#)” on page 7228 and “[Exact Tests](#)” on page 7237. The **SAVAGE** option in the PROC NPAR1WAY statement provides Savage score analysis and asymptotic tests.

SCORES=DATA

PERM

requests an exact test that uses the input data as scores. For two-sample data, the test is based on the rank-sum statistic. For multisample data, the test is based on the one-way ANOVA statistic. For more information, see the sections “[Scores for Linear Rank and One-Way ANOVA Tests](#)” on page 7227 and “[Exact Tests](#)” on page 7237. The **SCORES=DATA** option in the PROC NPAR1WAY statement provides analysis of the data scores and the corresponding asymptotic test.

ST

requests an exact Siegel-Tukey test. For more information, see the sections “[Siegel-Tukey Scores](#)” on page 7228 and “[Exact Tests](#)” on page 7237. The **ST** option in the PROC NPAR1WAY statement provides analysis of Siegel-Tukey scores and asymptotic tests.

VW

NORMAL

requests an exact Van der Waerden (normal scores) test. For more information, see the sections “[Van der Waerden \(Normal\) Scores](#)” on page 7227 and “[Exact Tests](#)” on page 7237. The **VW (NORMAL)** option in the PROC NPAR1WAY statement provides analysis of Van der Waerden (normal) scores and asymptotic tests.

WILCOXON

requests an exact Wilcoxon test. When the data consist of two classification levels (samples), the exact test is based on the Wilcoxon rank-sum statistic. When the data consist of more than two levels (multisample data), the exact test is based on the one-way ANOVA statistic for Wilcoxon scores, which is the Kruskal-Wallis statistic. For more information, see the sections “[Wilcoxon Scores](#)” on page 7227 and “[Exact Tests](#)” on page 7237. The **WILCOXON** option in the PROC NPAR1WAY statement provides analysis of Wilcoxon scores and the asymptotic Wilcoxon and Kruskal-Wallis tests.

Computation Options

Computation-options specify options for computing exact statistics. You can specify the following *computation-options*:

ALPHA= α

specifies the level of the confidence limits for Monte Carlo p -value estimates. The value of α must be between 0 and 1; a confidence level of α produces $100(1 - \alpha)\%$ confidence limits. By default, ALPHA=0.01, which produces 99% confidence limits for the exact p -values.

This option invokes the MC option.

MAXTIME=*value*

specifies the maximum clock time (in seconds) that PROC NPAR1WAY can use to compute an exact p -value. If the procedure does not complete the computation within the specified time, the computation terminates. The maximum time value must be a positive number. This option is available for exact p -value computation and for Monte Carlo estimation of exact p -values. For more information, see the section “[Computational Resources](#)” on page 7239.

MC

requests Monte Carlo estimation of exact p -values instead of direct exact p -value computation. Monte Carlo estimation can be useful for large problems where exact computations require a substantial amount of time and memory but asymptotic approximations might not be sufficient. For more information, see the section “[Monte Carlo Estimation](#)” on page 7240.

This option is available for all EXACT *statistic-options* except the HL option, which produces exact Hodges-Lehmann confidence limits. The ALPHA=, N=, and SEED= options invoke the MC option.

MIDP

requests exact mid p -values for the exact tests. The exact mid p -value is defined as the exact p -value minus half the exact point probability. For two-sample data, PROC NPAR1WAY provides exact mid p -values for the one-sided tests of the linear rank statistics. For multisample data, PROC NPAR1WAY provides exact mid p -values for the one-way ANOVA tests. For more information, see the section “[Definition of \$p\$ -Values](#)” on page 7238.

The MIDP option is available for all EXACT statement *statistic-options* except the HL option, which produces exact Hodges-Lehmann confidence limits. You cannot specify both the MIDP option and the MC option.

N= n

specifies the number of samples for Monte Carlo estimation. The value of n must be a positive integer. Larger values of n produce more precise estimates of exact p -values. Because larger values of n generate more samples, the computation time increases. By default, N=10,000.

This option invokes the MC option.

PFORMAT=*format-name* | EXACT

specifies the display format for exact p -values. PROC NPAR1WAY applies this format to one- and two-sided exact p -values, exact point probabilities, and exact mid p -values. By default, PROC NPAR1WAY displays exact p -values in the PVALUE6.4 format.

You can provide a *format-name* or you can specify PFORMAT=EXACT to control the format of exact p -values. The value of *format-name* can be any standard SAS numeric format or a user-defined format.

The format length must not exceed 24. For information about formats, see the FORMAT procedure in the *Base SAS Procedures Guide* and the FORMAT statement and SAS format in *SAS Formats and Informats: Reference*.

If you specify PFORMAT=EXACT, PROC NPAR1WAY uses the 6.4 format to display exact p -values that are greater than or equal to 0.001; the procedure uses the E10.3 format to display values that are between 0.000 and 0.001.

POINT

requests exact point probabilities for the exact tests. The exact point probability is the exact probability that the test statistic equals the observed value. For two-sample data, PROC NPAR1WAY provides exact point probabilities for the one-sided tests of the linear rank statistics. For multisample data, PROC NPAR1WAY provides exact point probabilities for the one-way ANOVA tests. For more information, see the section “[Definition of \$p\$ -Values](#)” on page 7238.

The POINT option is available for all EXACT statement *statistic-options* except the [HL](#) option, which produces exact Hodges-Lehmann confidence limits. You cannot specify both the POINT option and the [MC](#) option.

SEED=*number*

specifies the initial seed for random number generation for Monte Carlo estimation. The value of the SEED= option must be an integer. If you do not specify the SEED= option or if the SEED= value is negative or 0, PROC NPAR1WAY uses the time of day from the computer’s clock to obtain the initial seed.

This option invokes the [MC](#) option.

FREQ Statement

FREQ *variable* ;

The FREQ statement names a numeric variable that represents a frequency of occurrence for each observation in the input data set. PROC NPAR1WAY treats the data set as if each observation appears n times, where n is the value of the FREQ variable for the observation. The sum of the FREQ variable values represents the total number of observations, and the analysis is based on this expanded number of observations.

Noninteger values of the FREQ variable are truncated to the largest integer less than the FREQ value. An observation is used in the analysis only if the value of the FREQ variable is greater than or equal to 1.

OUTPUT Statement

OUTPUT <OUT=SAS-data-set> <output-options> ;

The OUTPUT statement creates a SAS data set that contains statistics that PROC NPAR1WAY computes. Table 90.4 lists the statistics that can be stored in the output data set. You identify which statistics to include by specifying *output-options*.

The output data set contains one observation for each analysis variable that you name in the VAR statement. If you use a BY statement, the output data set contains an observation or set of observations for each BY group. For more information, see the section “Contents of the Output Data Set” on page 7241.

As an alternative to the OUTPUT statement, you can use the Output Delivery System (ODS) to store statistics that PROC NPAR1WAY computes. ODS can create a SAS data set from any table that PROC NPAR1WAY produces. For more information, see the section “ODS Table Names” on page 7255 and Chapter 23, “Using the Output Delivery System.”

The OUTPUT statement data set does not include Monte Carlo estimates of exact *p*-values, multiple comparison statistics, or stratified analysis statistics. You can use ODS to store these statistics in output data sets.

You can specify the following *options*:

OUT=SAS-data-set

names the output data set. When you use an OUTPUT statement but do not specify the OUT= option, PROC NPAR1WAY creates a data set and names it by using the DATA n convention.

output-options

specifies the statistics to include in the output data set. Table 90.4 lists the *output-options* that are available in the OUTPUT statement. Descriptions of the *output-options* follow the table in alphabetical order.

When you specify an *output-option*, the output data set includes the test statistic and associated values from the analysis that you specify. The associated values might include standardized statistics, one- and two-sided *p*-values, exact *p*-values, degrees of freedom, and confidence limits. See the section “Contents of the Output Data Set” on page 7241 for a list of output data set variables that each *output-option* produces.

If you do not specify any *output-options* in the OUTPUT statement, the output data set includes all available statistics from the analyses that you request in the PROC NPAR1WAY statement. If you request a statistic in the OUTPUT statement but do not request the corresponding analysis in the PROC NPAR1WAY statement, PROC NPAR1WAY performs the corresponding analysis.

Table 90.4 OUTPUT Statement Output Options

Output Option	Output Data Set Statistics
AB	Ansari-Bradley test
ANOVA	Analysis of variance
CONOVER	Conover test
EDF KS	Kolmogorov-Smirnov test, Cramér–von Mises test, and Kuiper test for two-sample data
FP	Fligner-Policello test
HL	Hodges-Lehmann estimates
KLOTZ	Klotz test
MEDIAN	Median test
MOOD	Mood test
RMD	Ratio mean deviations (RMD) test
SAVAGE	Savage test
SCORES=DATA	Test with input data as scores
ST	Siegel-Tukey test
VW NORMAL	Van der Waerden (normal scores) test
WILCOXON	Wilcoxon test for two-sample data and Kruskal-Wallis test

You can specify the following *output-options*:

AB

includes statistics from the Ansari-Bradley analysis in the output data set. The **AB** option in the PROC NPAR1WAY statement requests analysis of Ansari-Bradley scores. For more information, see the section “[Ansari-Bradley Scores](#)” on page 7228.

ANOVA

includes analysis of variance statistics in the output data set. The **ANOVA** option in the PROC NPAR1WAY statement requests a standard analysis of variance for the raw data.

CONOVER

includes statistics from the Conover analysis in the output data set. The **CONOVER** option in the PROC NPAR1WAY statement requests analysis of Conover scores. For more information, see the section “[Conover Scores](#)” on page 7229.

EDF**KS**

includes empirical distribution function statistics in the output data set. The **EDF** option in the PROC NPAR1WAY statement requests computation of these statistics, which include the Kolmogorov-Smirnov test, the Cramér–von Mises test, and the Kuiper test for two sample-data. For more information, see the section “[Empirical Distribution Function Tests](#)” on page 7235.

FP

includes the Fligner-Policello test in the output data set. The **FP** option in the PROC NPAR1WAY statement requests the Fligner-Policello test, which is available for two-sample data. For more information, see the section “[Fligner-Policello Test](#)” on page 7232.

HL

includes the Hodges-Lehmann estimate and confidence limits in the output data set. The **HL** option in the PROC NPAR1WAY statement requests Hodges-Lehmann estimation, which is available for two-sample data. For more information, see the section “[Hodges-Lehmann Estimation of Location Shift](#)” on page 7230.

KLOTZ

includes statistics from the Klotz analysis in the output data set. The **KLOTZ** option in the PROC NPAR1WAY statement requests analysis of Klotz scores. For more information, see the section “[Klotz Scores](#)” on page 7228.

MEDIAN

includes statistics from the median analysis in the output data set. The **MEDIAN** option in the PROC NPAR1WAY statement requests analysis of median scores. For more information, see the section “[Median Scores](#)” on page 7227.

MOOD

includes statistics from the Mood analysis in the output data set. The **MOOD** option in the PROC NPAR1WAY statement requests analysis of Mood scores. For more information, see the section “[Mood Scores](#)” on page 7229.

RMD

includes statistics from the ratio mean deviations (RMD) test in the output data set. The **RMD** option in the PROC NPAR1WAY statement and the **RMD** option in the EXACT statement request the RMD test. For more information, see the section “[Ratio Mean Deviations \(RMD\) Exact Test](#)” on page 7234.

SAVAGE

includes statistics from the Savage analysis in the output data set. The **SAVAGE** option in the PROC NPAR1WAY statement requests analysis of SAVAGE scores. For more information, see the section “[Savage Scores](#)” on page 7228.

SCORES=DATA**PERM**

includes statistics from the analysis of data scores in the output data set. The **SCORES=DATA** option in the PROC NPAR1WAY statement requests an analysis that uses the raw input data as scores. For more information, see the section “[Scores for Linear Rank and One-Way ANOVA Tests](#)” on page 7227.

ST

includes statistics from the Siegel-Tukey analysis in the output data set. The **ST** option in the PROC NPAR1WAY statement requests analysis of Siegel-Tukey scores. For more information, see the section “[Siegel-Tukey Scores](#)” on page 7228.

VW**NORMAL**

includes statistics from the Van der Waerden (normal scores) analysis in the output data set. The **VW (NORMAL)** option in the PROC NPAR1WAY statement requests analysis of Van der Waerden scores. For more information, see the section “[Van der Waerden \(Normal\) Scores](#)” on page 7227.

WILCOXON

includes statistics from the Wilcoxon analysis in the output data set. The **WILCOXON** option in the PROC NPAR1WAY statement requests analysis of Wilcoxon scores. For two-sample data, this analysis includes the Wilcoxon rank-sum test. For all data (two-sample and multisample), the analysis includes the Kruskal-Wallis test. For more information, see the section “[Wilcoxon Scores](#)” on page 7227.

STRATA Statement

STRATA *variables* < / *options* > ;

The STRATA statement names one or more *variables* that define the strata. The levels of the STRATA variables combine to form the strata, which are nonoverlapping subgroups.

The STRATA variables can be either character or numeric. PROC NPAR1WAY uses the formatted values of the STRATA variables to determine the STRATA variable levels. Thus, you can use formats to group values into levels. For more information, see the discussion of the FORMAT procedure in the [Base SAS Procedures Guide](#) and the discussions of the FORMAT statement and SAS formats in [SAS Formats and Informats: Reference](#).

PROC NPAR1WAY excludes an observation from the stratified analysis if it has a missing value for any STRATA variable unless you specify the **MISSING** option in the PROC NPAR1WAY statement. When you specify the **MISSING** option, PROC NPAR1WAY treats missing values of the STRATA variables as valid, nonmissing levels. For more information, see the section “[Missing Values](#)” on page 7223.

You can request stratified analyses by specifying one or more score type *options* (**WILCOXON**, **MEDIAN**, **VW | NORMAL**, **SAVAGE**, and **SCORES=DATA**) in the STRATA statement. If you do not specify a score type *option* in the STRATA statement (or the **ALIGN=STRATA** option in the PROC NPAR1WAY statement), PROC NPAR1WAY performs stratified analysis of Wilcoxon scores by default. The statistic options in the PROC NPAR1WAY statement request unstratified analyses; the statistic options in the STRATA statement request stratified analyses.

Stratified analysis is available for two-sample data (when the data are classified into two levels by the **CLASS** variable). For more information, see the section “[Stratified Analysis](#)” on page 7229.

If you specify the **ALIGN=STRATA** option in the PROC NPAR1WAY statement along with a STRATA statement, PROC NPAR1WAY aligns the analysis variable values by strata before performing the unstratified analyses that you request (by specifying options in the PROC NPAR1WAY statement). When you specify the **ALIGN=STRATA** option to compute aligned analysis, you cannot also specify STRATA statement *options* to perform stratified analysis.

You can specify the following *options*:

MEDIAN

requests a stratified analysis of median scores. For more information, see the sections “[Median Scores](#)” on page 7227 and “[Stratified Analysis](#)” on page 7229.

RANKS=STRATUM | OVERALL

specifies the type of ranking to be used for stratified analysis. RANKS=STRATUM ranks (and scores) the observations separately within each stratum. RANKS=OVERALL ranks (and scores) the observations overall (without regard to stratum classification) before stratified analysis is performed. By default, RANKS=STRATUM. For more information, see the section “[Stratified Analysis](#)” on page 7229.

SAVAGE

requests a stratified analysis of Savage scores. For more information, see the sections “[Savage Scores](#)” on page 7228 and “[Stratified Analysis](#)” on page 7229.

SCORES=DATA**PERM**

requests a stratified analysis of the raw (unscored) values of the analysis variable. This option enables you to construct any scores for your data (for example, by using the DATA step) and analyze the scores by using PROC NPAR1WAY. For more information, see the sections “[Scores for Linear Rank and One-Way ANOVA Tests](#)” on page 7227 and “[Stratified Analysis](#)” on page 7229.

VW**NORMAL**

requests a stratified analysis of Van der Waerden (normal) scores. For more information, see the sections “[Van der Waerden \(Normal\) Scores](#)” on page 7227 and “[Stratified Analysis](#)” on page 7229.

WEIGHTS=STRATUM | EQUAL

specifies the stratum weights for stratified analysis. WEIGHTS=STRATUM weights the stratum statistics by stratum size; the weight of stratum i is $1/(n_i + 1)$, where n_i is the number of observations (frequency) in stratum i . WEIGHTS=EQUAL weights the stratum statistics equally; the weight of each stratum is 1. For more information, see the section “[Stratified Analysis](#)” on page 7229.

When you specify **RANKS=OVERALL**, WEIGHTS=EQUAL by default; otherwise, WEIGHTS=STRATUM by default.

WILCOXON

requests a stratified analysis of Wilcoxon scores. For more information, see the sections “[Wilcoxon Scores](#)” on page 7227 and “[Stratified Analysis](#)” on page 7229. By default (when **WEIGHTS=STRATUM** and **RANKS=STRATUM**), the stratified Wilcoxon test is the van Elteren test (Van Elteren 1960).

VAR Statement

VAR *variables* ;

The VAR statement names the response (dependent) variables to be included in the analysis. These variables must be numeric. If you omit the VAR statement, PROC NPAR1WAY includes all numeric variables in the data set except the **BY**, **CLASS**, **FREQ**, and **STRATA** variables.

Details: NPAR1WAY Procedure

Missing Values

If an observation has a missing value for a response (**VAR**) variable, PROC NPAR1WAY excludes the observation from the analysis of that variable. If an observation has a missing or nonpositive value for the **FREQ** variable, PROC NPAR1WAY excludes that observation from the analysis.

To display the number of missing response values, you can specify the **VARINFO** option in the PROC NPAR1WAY statement. This provides a “Variable Information” table for each response variable. The table includes the number of observations that are used in the analysis and the number of missing values.

If an observation has a missing value for the **CLASS** variable, PROC NPAR1WAY excludes that observation from the analysis unless you specify the **MISSING** option in the PROC NPAR1WAY statement. If you specify the **MISSING** option, PROC NPAR1WAY treats missing values of the **CLASS** variable as a valid level.

If an observation has a missing value for a **STRATA** variable, PROC NPAR1WAY excludes that observation from the stratified analyses unless you specify the **MISSING** option in the PROC NPAR1WAY statement. If you specify the **MISSING** option, PROC NPAR1WAY treats missing values of the **STRATA** variables as valid levels.

When you specify a **STRATA** statement, PROC NPAR1WAY displays a “Stratum Information” table, which lists the stratum variable levels and the number of observations (frequency) in each stratum. This table summarizes the valid observations in the input data set (or **BY** group). An observation is valid if it has nonmissing values of the **CLASS** and **STRATA** variables (unless you specify the **MISSING** option) and a positive value of the **FREQ** variable. PROC NPAR1WAY then analyzes each response (**VAR**) variable separately; each analysis includes valid observations that have nonmissing values for the individual response variable.

PROC NPAR1WAY treats missing **BY** variable values like any other **BY** variable value. The missing values form a separate, valid **BY** group.

Tied Values

Tied values occur when two or more observations are equal, whether the observations occur in the same sample or in different samples. In theory, nonparametric tests were developed for continuous distributions where the probability of a tie is zero. In practice, however, ties often occur. PROC NPAR1WAY uses the same method to handle ties for all score types. The procedure computes the scores as if there were no ties, averages the scores for tied observations, and assigns this average score to each observation that has the same value.

To display the number of tied response values, you can specify the [VARINFO\(TIES\)](#) option in the PROC NPAR1WAY statement. This provides a “Variable Information” table for each response variable. The table includes the number of variable levels and the number of tied observations.

When there are tied values, PROC NPAR1WAY first sorts the observations in ascending order and assigns ranks as if there were no ties. Then the procedure computes the scores based on these ranks by using the formula for the specified score type. The procedure averages the scores for tied observations and assigns this average score to each of the tied observations. Thus, all equal data values have the same score value. PROC NPAR1WAY then computes the test statistic from these scores.

The asymptotic tests might be less accurate when the distribution of the data is heavily tied. For such data, it might be appropriate to use the exact tests provided by PROC NPAR1WAY as described in the section [“Exact Tests”](#) on page 7237.

When computing empirical distribution function statistics for data with ties, PROC NPAR1WAY uses the formulas given in the section [“Empirical Distribution Function Tests”](#) on page 7235. No special handling of ties is necessary.

PROC NPAR1WAY bases its computations on the internal numeric values of the analysis variables; the procedure does not format or round these values before analysis. When values differ in their internal representation, even slightly, PROC NPAR1WAY does not treat them as tied values. If this is a concern for your data, round the analysis variables by an appropriate amount before invoking PROC NPAR1WAY. For information about the ROUND function, see the discussion in [SAS Functions and CALL Routines: Reference](#).

Statistical Computations

Simple Linear Rank Tests for Two-Sample Data

Statistics of the form

$$S = \sum_{j=1}^n c_j a(R_j)$$

are called *simple linear rank statistics*, where

- R_j is the rank of observation j
- $a(R_j)$ is the score based on the rank of observation j
- c_j indicates the class of observation j
- n is the total number of observations

For two-sample data (where the observations are classified into two levels), PROC NPAR1WAY calculates simple linear rank statistics for the scores that you specify. The section “[Scores for Linear Rank and One-Way ANOVA Tests](#)” on page 7227 describes the available scores, which you can use to test for differences in location and differences in scale.

To compute the linear rank statistic S , PROC NPAR1WAY sums the scores of the observations in the smaller of the two samples. If both samples have the same number of observations, PROC NPAR1WAY sums those scores for the sample that appears first in the input data set.

For each score that you specify, PROC NPAR1WAY computes an asymptotic test of the null hypothesis of no difference between the two classification levels. Exact tests are also available for these two-sample linear rank statistics. PROC NPAR1WAY computes exact tests for each score type that you specify in the **EXACT** statement. For more information, see the section “[Exact Tests](#)” on page 7237.

To compute an asymptotic test for a linear rank sum statistic, PROC NPAR1WAY uses a standardized test statistic z , which has an asymptotic standard normal distribution under the null hypothesis. The standardized test statistic is computed as

$$z = (S - E_0(S)) / \sqrt{\text{Var}_0(S)}$$

where $E_0(S)$ is the expected value of S under the null hypothesis, and $\text{Var}_0(S)$ is the variance under the null hypothesis. As shown in Randles and Wolfe (1979),

$$E_0(S) = \frac{n_1}{n} \sum_{j=1}^n a(R_j)$$

where n_1 is the number of observations in the first (smaller) class level (sample), n_2 is the number of observations in the other class level, and

$$\text{Var}_0(S) = \frac{n_1 n_2}{n(n-1)} \sum_{j=1}^n (a(R_j) - \bar{a})^2$$

where $\bar{a}$ is the average score,

$$\bar{a} = \frac{1}{n} \sum_{j=1}^n a(R_j)$$

Definition of p -Values

PROC NPAR1WAY computes one-sided and two-sided asymptotic p -values for each two-sample linear rank test. When the test statistic z is greater than its null hypothesis expected value of 0, PROC NPAR1WAY computes the right-sided p -value, which is the probability of a larger value of the statistic occurring under the null hypothesis. When the test statistic is less than or equal to 0, PROC NPAR1WAY computes the left-sided p -value, which is the probability of a smaller value of the statistic occurring under the null hypothesis. The one-sided p -value $P_1(z)$ can be expressed as

$$P_1(z) = \begin{cases} \text{Prob}(Z > z) & \text{if } z > 0 \\ \text{Prob}(Z < z) & \text{if } z \leq 0 \end{cases}$$

where Z has a standard normal distribution. The two-sided p -value $P_2(z)$ is computed as

$$P_2(z) = \text{Prob}(|Z| > |z|)$$

Continuity Correction

PROC NPAR1WAY uses a continuity correction for the asymptotic two-sample Wilcoxon and Siegel-Tukey tests by default. You can remove the continuity correction by specifying the **CORRECT=NO** option. PROC NPAR1WAY incorporates the continuity correction when computing the standardized test statistic z by subtracting 0.5 from the numerator ($S - E_0(S)$) if it is greater than 0. If the numerator is less than 0, PROC NPAR1WAY adds 0.5. Some sources recommend a continuity correction for nonparametric tests that use a continuous distribution to approximate a discrete distribution. For more information, see Sheskin (1997).

If you specify **CORRECT=NO**, PROC NPAR1WAY does not use a continuity correction for any test.

One-Way ANOVA Tests

PROC NPAR1WAY computes a one-way ANOVA test for each score type that you specify. Under the null hypothesis of no difference among class levels (samples), this test statistic has an asymptotic chi-square distribution with $r - 1$ degrees of freedom, where r is the number of class levels. For Wilcoxon scores, this test is known as the Kruskal-Wallis test.

Exact one-way ANOVA tests are also available for multisample data (where the data are classified into more than two levels). For two-sample data, exact simple linear rank tests are available. PROC NPAR1WAY computes exact tests for each score type that you specify in the **EXACT** statement. For more information, see the section “Exact Tests” on page 7237.

PROC NPAR1WAY computes the one-way ANOVA test statistic as

$$C = \left(\sum_{i=1}^r (T_i - E_0(T_i))^2 / n_i \right) / S^2$$

where T_i is the total of scores for class level i , $E_0(T_i)$ is the expected total for level i under the null hypothesis of no difference among levels, n_i is the number of observations in level i , and S^2 is the sample variance of the scores. The total of scores for class level i is given by

$$T_i = \sum_{j=1}^n c_{ij} a(R_j)$$

where $a(R_j)$ is the score for observation j , and c_{ij} indicates whether observation j is in level i . The expected total of scores for class level i under the null hypothesis is equal to

$$E_0(T_i) = \frac{n_i}{n} \sum_{j=1}^n a(R_j)$$

The sample variance of the scores is computed as

$$S^2 = \frac{1}{(n-1)} \sum_{j=1}^n (a(R_j) - \bar{a})^2$$

where $\bar{a}$ is the average score,

$$\bar{a} = \frac{1}{n} \sum_{j=1}^n a(R_j)$$

Scores for Linear Rank and One-Way ANOVA Tests

For each score type that you specify, PROC NPARIWAY computes a one-way ANOVA statistic and also a linear rank statistic for two-sample data. The following score types are used primarily to test for differences in location: Wilcoxon, median, Van der Waerden (normal), and Savage. The following scores types are used to test for scale differences: Siegel-Tukey, Ansari-Bradley, Klotz, and Mood. Conover scores can be used to test for differences in both location and scale. This section gives formulas for the score types available in PROC NPARIWAY. For further information about the formulas and the applicability of each score, see Randles and Wolfe (1979), Gibbons and Chakraborti (2010), Conover (1999), and Hollander and Wolfe (1999).

In addition to the score types described in this section, you can specify the **SCORES=DATA** option to use the input data observations as scores. This enables you to produce a wide variety of tests. You can construct any scores by using the DATA step, and then you can use PROC NPARIWAY to compute the corresponding linear rank and one-way ANOVA tests for these scores. You can also analyze raw (unscored) data by using the **SCORES=DATA** option; for two-sample data, the corresponding exact test is a permutation test that is known as Pitman's test.

Wilcoxon Scores

Wilcoxon scores are the ranks of the observations, which can be written as

$$a(R_j) = R_j$$

where R_j is the rank of observation j , and $a(R_j)$ is the score of observation j .

Using Wilcoxon scores in the linear rank statistic for two-sample data produces the rank sum statistic of the Mann-Whitney-Wilcoxon test. Using Wilcoxon scores in the one-way ANOVA statistic produces the Kruskal-Wallis test. Wilcoxon scores are locally most powerful for location shifts of a logistic distribution.

When computing the asymptotic Wilcoxon two-sample test, PROC NPARIWAY uses a continuity correction by default, as described in the section “**Continuity Correction**” on page 7226. If you specify the **CORRECT=NO** option in the PROC NPARIWAY statement, the procedure does not use a continuity correction.

Median Scores

Median scores equal 1 for observations greater than the median, and 0 otherwise. In terms of the observation ranks, median scores are defined as

$$a(R_j) = \begin{cases} 1 & \text{if } R_j > (n + 1)/2 \\ 0 & \text{if } R_j \leq (n + 1)/2 \end{cases}$$

Using median scores in the linear rank statistic for two-sample data produces the two-sample median test. Using median scores in the one-way ANOVA statistic for multisample data produces the Brown-Mood test. Median scores are particularly powerful for distributions that are symmetric and heavy-tailed.

Van der Waerden (Normal) Scores

Van der Waerden scores are the quantiles of a standard normal distribution and are also known as *quantile normal scores*. Van der Waerden scores are computed as

$$a(R_j) = \Phi^{-1} \left(\frac{R_j}{n + 1} \right)$$

where Φ is the cumulative distribution function of a standard normal distribution. These scores are powerful for normal distributions.

Savage Scores

Savage scores are expected values of order statistics from the exponential distribution, with 1 subtracted to center the scores around 0. Savage scores are computed as

$$a(R_j) = \sum_{i=1}^{R_j} \left(\frac{1}{n-i+1} \right) - 1$$

Savage scores are powerful for comparing scale differences in exponential distributions or location shifts in extreme value distributions (Hajek 1969, p. 83).

Siegel-Tukey Scores

Siegel-Tukey scores are defined as

$$\begin{aligned} a(1) &= 1, & a(n) &= 2, & a(n-1) &= 3, & a(2) &= 4, \\ a(3) &= 5, & a(n-2) &= 6, & a(n-3) &= 7, & a(4) &= 8, \dots \end{aligned}$$

where the score values continue to increase in this pattern toward the middle ranks until all observations have been assigned a score.

When computing the asymptotic Siegel-Tukey two-sample test, PROC NPAR1WAY uses a continuity correction by default, as described in the section “[Continuity Correction](#)” on page 7226. If you specify the `CORRECT=NO` option in the PROC NPAR1WAY statement, the procedure does not use a continuity correction.

Ansari-Bradley Scores

Ansari-Bradley scores are similar to Siegel-Tukey scores, but Ansari-Bradley scoring assigns the same score value to corresponding extreme ranks. (Siegel-Tukey scores are a permutation of the ranks $1, 2, \dots, n$.) Ansari-Bradley scores are defined as

$$\begin{aligned} a(1) &= 1, & a(n) &= 1, \\ a(2) &= 2, & a(n-1) &= 2, \dots \end{aligned}$$

Equivalently, Ansari-Bradley scores are equal to

$$a(R_j) = \frac{n+1}{2} - \left| R_j - \frac{n+1}{2} \right|$$

Klotz Scores

Klotz scores are the squares of the Van der Waerden (normal) scores. Klotz scores are computed as

$$a(R_j) = \left(\Phi^{-1} \left(\frac{R_j}{n+1} \right) \right)^2$$

where Φ is the cumulative distribution function of a standard normal distribution.

Mood Scores

Mood scores are computed as the square of the difference between the observation rank and the average rank. Mood scores can be written as

$$a(R_j) = \left(R_j - \frac{n+1}{2} \right)^2$$

Conover Scores

Conover scores are based on the squared ranks of the absolute deviations from the sample means. For observation j the absolute deviation from the mean is computed as

$$U_j = |X_{j(i)} - \bar{X}_i|$$

where $X_{j(i)}$ is the value of observation j , observation j belongs to sample i , and $\bar{X}_i$ is the mean of sample i . The values of U_j are ranked, and the Conover score for observation j is computed as

$$a(U_j) = (\text{Rank}(U_j))^2$$

Following Conover (1999), when there are ties among the values of U_j , PROC NPAR1WAY assigns the average rank to each of the tied observations. To compute the average rank, PROC NPAR1WAY first ranks the U_j as if there were no ties and then averages the ranks of the tied observations.

The Conover score test is also known as the squared ranks test for variances. For more information, see Conover (1999).

Stratified Analysis

PROC NPAR1WAY provides stratified analysis of two-sample data for the following score types: Wilcoxon, median, Van der Waerden (normal), Savage, and data scores. It computes the stratified test statistic by combining the stratum-level linear rank statistics for the score type that you specify. The section “[Scores for Linear Rank and One-Way ANOVA Tests](#)” on page 7227 describes the computation of the rank-based scores. For stratified analysis, you can compute the scores by using within-stratum ranks or overall ranks. By default, or if you specify the **RANKS=STRATUM** option, PROC NPAR1WAY ranks and scores the response values separately within each stratum. If you specify the **RANKS=OVERALL** option, PROC NPAR1WAY ranks and scores the response values overall (without regard to stratum classification). For more information, see Mehrotra, Lu, and Li (2010), Lehmann and D’Abrera (2006), and Van Elteren (1960).

The stratified rank sum statistic T is computed as

$$T = \sum_{k=1}^K S_k / w_k$$

where S_k is the sum of the rank-based scores for stratum k (as described in the section “[Simple Linear Rank Tests for Two-Sample Data](#)” on page 7224), w_k is the weight of stratum k , and K is the total number of strata. By default, or if you specify the **WEIGHTS=STRATUM** option, PROC NPAR1WAY computes the stratum weights as $1/(n_k + 1)$, where n_k is the number of observations in stratum k . If you specify the **WEIGHTS=EQUAL** option, the stratum weight (w_k) is 1 for all strata.

The expected value and variance of T are similarly computed by combining stratum-level values as

$$E_0(T) = \sum_{k=1}^K E_0(S_k)/w_k$$

$$\text{Var}_0(T) = \sum_{k=1}^K \text{Var}_0(S_k)/w_k^2$$

where $E_0(S_k)$ and $\text{Var}_0(S_k)$ are the expected value and variance of the rank sum statistic for stratum k . For more information, see the section “Simple Linear Rank Tests for Two-Sample Data” on page 7224.

The stratified test statistic is computed as

$$z = (T - E_0(T)) / \sqrt{\text{Var}_0(T)}$$

Under the null hypothesis of no difference between the two classification levels, the test statistic has a standard normal distribution. PROC NPAR1WAY provides one-sided and two-sided asymptotic p -values for the test as described in the section “Definition of p -Values” on page 7225.

PROC NPAR1WAY defines the reference class group (sample) for the stratified two-sample analysis by considering the number of observations (total frequency) across all strata. The reference group is defined as the smaller of the two samples. If both samples have the same number of observations, the reference group is defined as the level that appears first in the input data set. Each stratum rank sum statistic is then computed as the sum of the scores for the observations in the reference group. PROC NPAR1WAY lists the reference group first in the “Class Information” table.

Hodges-Lehmann Estimation of Location Shift

If you specify the **HL** option, PROC NPAR1WAY computes the Hodges-Lehmann estimate of location shift for two-sample data. This option also provides asymptotic confidence limits for the location shift (which are sometimes known as Moses confidence limits). You can specify the confidence level in the **ALPHA=** option in the PROC NPAR1WAY statement. By default, ALPHA=0.05, which produces 95% confidence limits. Additionally, you can request exact confidence limits for the location shift by specifying the **HL** option in the EXACT statement.

The Hodges-Lehmann estimator of location shift is associated with the Wilcoxon linear rank statistic. For more information, see Hollander and Wolfe (1999) and Hodges and Lehmann (1983).

PROC NPAR1WAY computes the Hodges-Lehmann estimate $\hat{\Delta}$ as the median of all paired differences between observations in the two samples (classes), which can be written as

$$\hat{\Delta} = \text{median} \left((Y_j - X_i) \quad \text{where } j = 1, 2, \dots, n_1; i = 1, 2, \dots, n_2 \right)$$

The Y_j are observations in class 1, the X_i are observations in class 2, and n_1 and n_2 denote the number of observations in class 1 and class 2, respectively.

By default, PROC NPAR1WAY uses the larger of the two classes as the reference class X (class 2 in the Hodges-Lehmann difference). If both classes have the same number of observations, PROC NPAR1WAY uses the class that appears second in the “Wilcoxon Scores” table as the reference class. (The table displays class levels in the order in which they appear in the input data set.) You can specify the reference class by using the **HL(REFCLASS=)** option. REFCLASS=1 specifies the first class that is listed in the “Wilcoxon Scores”

table, and REFCLASS=2 specifies the second class. REFCLASS='class-value' identifies the reference class by the formatted value of the **CLASS** variable.

Let m denote the total number of differences ($n_1 \times n_2$), and let $U^{(k)}$ denote the k th value of $(Y_j - X_i)$ among the ordered differences. When m is an odd number, the median difference is the value that has rank $(m + 1)/2$,

$$\hat{\Delta} = U^{(k)} \quad \text{where } k = (m + 1)/2$$

When m is an even number, the median difference is the average of the values that have ranks $(m/2)$ and $((m/2) + 1)$,

$$\hat{\Delta} = (U^{(k)} + U^{(k+1)})/2 \quad \text{where } k = m/2$$

Following Hollander and Wolfe (1999), the asymptotic lower and upper confidence limits for the location shift are

$$(\Delta_L = U^{(C_\alpha)}, \Delta_U = U^{(m+1-C_\alpha)})$$

where C_α is the largest integer less than or equal to C_α^* , which is computed as

$$C_\alpha^* = E_0(S) - z_{\alpha/2} \sqrt{\text{Var}_0(S)}$$

where $E_0(S)$ and $\text{Var}_0(S)$ are the expected value and variance, respectively, of the Wilcoxon statistic S under the null hypothesis (as described in the section “Simple Linear Rank Tests for Two-Sample Data” on page 7224), and $z_{\alpha/2}$ is the $100(1 - \alpha/2)$ th percentile of the standard normal distribution. For Wilcoxon rank scores,

$$E_0(S) = n_1 n_2 / 2$$

When there are no tied values, $\text{Var}_0(S)$ for Wilcoxon scores equals

$$\text{Var}_0(S) = n_1 n_2 (n_1 + n_2 + 1) / 12$$

PROC NPAR1WAY displays the midpoint of the confidence interval (Δ_L, Δ_U) , which can also be used as an estimate of location shift. For more information, see Lehmann (1963). Additionally, PROC NPAR1WAY provides an estimate of the asymptotic standard error of $\hat{\Delta}$ based on the length of the confidence interval, which is computed as

$$\text{se}(\hat{\Delta}) = (\Delta_U - \Delta_L) / 2 z_{\alpha/2}$$

Exact Confidence Limits

If you specify the **HL** option in the EXACT statement, PROC NPAR1WAY computes exact confidence limits for the location shift between the two samples. You can specify the level of the confidence limits in the **ALPHA=** option in the PROC NPAR1WAY statement. By default, ALPHA=0.05, which produces 95% confidence limits.

PROC NPAR1WAY computes exact confidence limits for the location shift as described in Randles and Wolfe (1979, p. 180). PROC NPAR1WAY first generates the exact conditional distribution of the Mann-Whitney

U statistic, which is the number of pairwise differences $(Y_j - X_i)$ that are positive plus half the number of pairwise differences that are 0. The Mann-Whitney statistic is defined as

$$M = \sum_{j=1}^{n_1} \sum_{i=1}^{n_2} \phi(Y_j, X_i)$$

where

$$\phi(Y_j, X_i) = \begin{cases} 1 & \text{if } Y_j > X_i \\ 1/2 & \text{if } Y_j = X_i \\ 0 & \text{otherwise} \end{cases}$$

From the exact conditional distribution of the Mann-Whitney statistic M , PROC NPAR1WAY chooses $C_{L,\alpha}^*$ as the smallest value such that $\text{Prob}(M \geq C_{L,\alpha}^*) \leq \alpha/2$. Rounding $C_{L,\alpha}^*$ up to the nearest integer $C_{L,\alpha}$, the lower confidence limit is the difference $(Y_i - X_j)$ that has a rank of $(n_1 n_2 - C_{L,\alpha} + 1)$.

To find the upper confidence limit, PROC NPAR1WAY chooses $C_{U,\alpha}^*$ as the largest Mann-Whitney value such that $\text{Prob}(M \leq C_{U,\alpha}^*) \leq \alpha/2$. Rounding $C_{U,\alpha}^*$ down to the nearest integer $C_{U,\alpha}$, the upper confidence limit is the difference $(Y_i - X_j)$ that has a rank of $(n_1 n_2 - C_{U,\alpha})$.

Because this is a discrete problem, the confidence coefficient is not exactly $(1 - \alpha)$ but is at least $(1 - \alpha)$; thus, these confidence limits are conservative.

Fligner-Policello Test

If you specify the **FP** option, PROC NPAR1WAY computes the Fligner-Policello location test for two-sample data (Fligner and Policello 1981). The null hypothesis for the test is $H_0: \theta_x = \theta_y$, where θ_x and θ_y are the population medians of the two classes. The Fligner-Policello test assumes that the distribution in each class is symmetric around the class median, but it does not require that the two class distributions have the same form or that the class variances be equal. For more information, see Hollander and Wolfe (1999) and Juneau (2007).

The Fligner-Policello test is based on placement scores (Orban and Wolfe 1979). The placement of an observation X_i from class X , $P(X_i)$, is defined as the number of observations in class Y that are less than X_i . If there are ties, the placement of X_i is adjusted by adding half the number of observations in class Y that are equal to X_i . The placement of an observation Y_j from class Y , $P(Y_j)$, is defined in the same way. The placements can be expressed as

$$P(X_i) = \sum_{j=1}^{n_y} (I(Y_j < X_i) + 0.5 I(Y_j = X_i))$$

$$P(Y_j) = \sum_{i=1}^{n_x} (I(X_i < Y_j) + 0.5 I(X_i = Y_j))$$

where $I(\cdot)$ is an indicator function and n_x and n_y denote the number of observations in class X and class Y , respectively.

The average placements for class X and class Y are computed as

$$\bar{P}_x = \left(\sum_{i=1}^{n_x} P(X_i) \right) / n_x$$

$$\bar{P}_y = \left(\sum_{j=1}^{n_y} P(Y_j) \right) / n_y$$

The Fligner-Policello test statistic is computed as

$$z = \left(\sum_{j=1}^{n_y} P(Y_j) - \sum_{i=1}^{n_x} P(X_i) \right) / \left(2\sqrt{V_x + V_y + \bar{P}_x \bar{P}_y} \right)$$

where

$$V_x = \sum_{i=1}^{n_x} (P(X_i) - \bar{P}_x)^2$$

$$V_y = \sum_{j=1}^{n_y} (P(Y_j) - \bar{P}_y)^2$$

and the standard deviation of the placements is $\sqrt{V_x/(n_x - 1)}$ for class X and $\sqrt{V_y/(n_y - 1)}$ for class Y .

Under the null hypothesis, the Fligner-Policello statistic has an asymptotic standard normal distribution. PROC NPAR1WAY provides one- and two-sided asymptotic p -values for the Fligner-Policello test. For the one-sided test, PROC NPAR1WAY displays the right-sided p -value when the test statistic z is greater than its null hypothesis expected value of 0. PROC NPAR1WAY displays the left-sided p -value when the test statistic z is less than or equal to 0. The one-sided p -value $P_1(z)$ can be expressed as

$$P_1(z) = \begin{cases} \text{Prob}(Z > z) & \text{if } z > 0 \\ \text{Prob}(Z < z) & \text{if } z \leq 0 \end{cases}$$

where Z has a standard normal distribution. The two-sided p -value $P_2(z)$ is computed as $\text{Prob}(|Z| > |z|)$.

When you specify the FP option, PROC NPAR1WAY displays a “Fligner-Policello Placements” table and a “Fligner-Policello Test” table. The “Fligner-Policello Placements” table contains the following information for each of the two classes: number of observations, sum of the placements, average placement, and standard deviation of the placements. The “Fligner-Policello Test” table contains the test statistic z and the corresponding one- and two-sided p -values. This table also displays the difference between the class placement sums, which is the numerator of the test statistic. When ODS Graphics is enabled and you specify the FP or PLOTS=FPBOXPLOT option, PROC NPAR1WAY provides a box plot of the Fligner-Policello placements.

PROC NPAR1WAY computes the Fligner-Policello difference (the numerator of the test statistic) as the placement sum for class Y minus the placement sum for class X (the reference class). By default, PROC NPAR1WAY uses the larger of the two classes as the reference class X . If both classes have the same number of observations, PROC NPAR1WAY uses the class that appears second in the “Fligner-Policello Placements”

table as the reference class. (The table displays class levels in the order in which they appear in the input data set.)

You can specify the reference class by using the **FP(REFCLASS=)** option. **REFCLASS=1** specifies the first class that is listed in the “Fligner-Policello Placements” table, and **REFCLASS=2** specifies the second class. **REFCLASS='class-value'** identifies the reference class by the formatted value of the **CLASS** variable.

Ratio Mean Deviations (RMD) Exact Test

If you specify the **RMD** option in the **PROC NPAR1WAY** statement or the **EXACT** statement, **PROC NPAR1WAY** computes the ratio mean deviations (RMD) exact test of scale differences for two-sample data (Higgins 2004). This test is based on the ratio of absolute deviations of the observations (response variable values) from their class medians. For more information, see Butar and Park (2008) and Richter and McCann (2017).

The RMD test statistic is computed as

$$\text{RMD} = \left(\sum_{j=1}^{n_1} d_{1j} / n_1 \right) / \left(\sum_{j=1}^{n_2} d_{2j} / n_2 \right)$$

where the absolute deviations for class i are

$$d_{ij} = |X_{ij} - m_i|$$

and where m_i is the sample median of class i , X_{ij} are the observations in class i ($j = 1, \dots, n_i$), and n_i is the number of observations in class i ($i = 1, 2$).

PROC NPAR1WAY uses the first class that is displayed in the “Deviations” table as class 1 and the second class as class 2. The “Deviations” table displays class levels in the order in which they appear in the input data set.

The RMD p -value is computed by using an exact permutation test that is conditional on the observed deviations (Higgins 2004). **PROC NPAR1WAY** displays the exact one-sided p -value as the left-sided p -value ($\text{Pr} \leq \text{RMD}$) when the value of the RMD statistic is less than or equal to 1; otherwise, **PROC NPAR1WAY** displays the right-sided p -value ($\text{Pr} \geq \text{RMD}$).

If you specify the **MIDP computation-option** in the **EXACT** statement, **PROC NPAR1WAY** also provides the exact mid p -value for the RMD test. If you specify the **POINT computation-option** in the **EXACT** statement, **PROC NPAR1WAY** provides the exact point probability for the RMD test.

PROC NPAR1WAY computes the modified RMD two-sided statistic as

$$\text{RMD2} = \max(\text{RMD}, 1/\text{RMD})$$

The corresponding exact two-sided p -value is computed by using a permutation test for the RMD2 statistic.

To obtain Monte Carlo estimates of exact p -values for the RMD test, you can specify the **MC computation-option** together with the **RMD statistic-option** in the **EXACT** statement. For more information, see the section “Exact Tests” on page 7237.

Multiple Comparisons Based on Pairwise Rankings

If you specify the **DSCF** option, PROC NPAR1WAY computes the Dwass, Steel, Critchlow-Fligner (DSCF) multiple comparison analysis, which is based on pairwise two-sample Wilcoxon comparisons (Dwass 1960; Steel 1960; Critchlow and Fligner 1991). The DSCF analysis is available when the number of **CLASS** variable levels (samples) is greater than 2. There are $r(r-1)/2$ pairs of samples, where r is the total number of samples.

For each pair of samples, PROC NPAR1WAY ranks the observations and computes the standardized Wilcoxon test statistic, which is described in the sections “Simple Linear Rank Tests for Two-Sample Data” on page 7224 and “Wilcoxon Scores” on page 7227. The procedure does not include a continuity correction in the Wilcoxon test statistics that it computes for the DSCF analysis. The DSCF statistic for a pair of samples is computed as $\sqrt{2} z$, where z is the two-sample standardized Wilcoxon statistic. Under the null hypothesis of no location differences among the r samples, the distribution of the DSCF statistics can be approximated by the studentized range distribution for r independent standard normal variables. The p -value for a two-sample DSCF comparison is the percentile of the studentized range distribution that corresponds to the value of the DSCF statistic. For more information, see Hollander and Wolfe (1999), Juneau (2007), and Juneau (2004).

When you specify the DSCF option, PROC NPAR1WAY displays the “Pairwise Two-Sided Multiple Comparison Analysis” table, which contains the following information for each two-sample comparison: comparison identification (two CLASS variable levels), two-sample Wilcoxon test statistic, DSCF statistic, and two-sided DSCF p -value. You can store these statistics in a SAS data set by using the Output Delivery System (ODS). The ODS name for this table is DSCF. For more information, see the section “ODS Table Names” on page 7255 and Chapter 23, “Using the Output Delivery System.”

Empirical Distribution Function Tests

If you specify the **EDF** option, PROC NPAR1WAY computes tests based on the empirical distribution function. These include the Kolmogorov-Smirnov and Cramér-von Mises tests, and also the Kuiper test for two-sample data. This section gives formulas for these test statistics. For further information about the formulas and the interpretation of EDF statistics, see Hollander and Wolfe (1999) and Gibbons and Chakraborti (2010). For information about the k -sample analogs of the Kolmogorov-Smirnov and Cramér-von Mises statistics, see Kiefer (1959).

The *empirical distribution function* (EDF) of a sample $\{x_j\}$, $j = 1, 2, \dots, n$, is defined as

$$F(x) = \frac{1}{n}(\text{number of } x_j \leq x) = \frac{1}{n} \sum_{j=1}^n I(x_j \leq x)$$

where $I(\cdot)$ is an indicator function. PROC NPAR1WAY uses the subsample of values within the i th class level to generate an EDF for the class, F_i . The EDF for the overall sample, pooled over classes, can also be expressed as

$$F(x) = \frac{1}{n} \sum_i n_i F_i(x)$$

where n_i is the number of observations in the i th class level, and n is the total number of observations.

Kolmogorov-Smirnov Test

The Kolmogorov-Smirnov statistic measures the maximum deviation of the EDF within the classes from the pooled EDF. PROC NPAR1WAY computes the Kolmogorov-Smirnov statistic as

$$KS = \max_j \sqrt{\frac{1}{n} \sum_i n_i (F_i(x_j) - F(x_j))^2} \quad \text{where } j = 1, 2, \dots, n$$

The asymptotic Kolmogorov-Smirnov statistic is computed as

$$KS_a = KS \times \sqrt{n}$$

For each class level i and overall, PROC NPAR1WAY displays the value of F_i at the maximum deviation from F and the value $\sqrt{n_i} (F_i - F)$ at the maximum deviation from F . PROC NPAR1WAY also gives the observation where the maximum deviation occurs.

If there are only two class levels, PROC NPAR1WAY computes the two-sample Kolmogorov-Smirnov test statistic D as

$$D = \max_j |F_1(x_j) - F_2(x_j)| \quad \text{where } j = 1, 2, \dots, n$$

The p -value for this test is the probability that D is greater than the observed value d under the null hypothesis of no difference between class levels (samples). PROC NPAR1WAY computes the asymptotic p -value for D by using the approximation

$$\text{Prob}(D > d) = 2 \sum_{i=1}^{\infty} (-1)^{(i-1)} e^{(-2i^2 z^2)}$$

where

$$z = d \sqrt{n_1 n_2 / n}$$

For more information, see Hodges (1957).

If you specify the **D** option, or if you request exact Kolmogorov-Smirnov p -values by specifying the **KS** option in the EXACT statement, PROC NPAR1WAY also computes the one-sided Kolmogorov-Smirnov statistics $D+$ and $D-$ for two-sample data as

$$D+ = \max_j (F_1(x_j) - F_2(x_j)) \quad \text{where } j = 1, 2, \dots, n$$

$$D- = \max_j (F_2(x_j) - F_1(x_j)) \quad \text{where } j = 1, 2, \dots, n$$

The asymptotic probability that $D+$ is greater than the observed value d^+ , under the null hypothesis of no difference between the two class levels, is computed as

$$\text{Prob}(D+ > d^+) = e^{-2z^2} \quad \text{where } z = d^+ \sqrt{n_1 n_2 / n}$$

Similarly, the asymptotic probability that $D-$ is greater than the observed value d^- is computed as

$$\text{Prob}(D- > d^-) = e^{-2z^2} \quad \text{where } z = d^- \sqrt{n_1 n_2 / n}$$

To request exact p -values for the Kolmogorov-Smirnov statistics, you can specify the **KS** option in the EXACT statement. For more information, see the section “Exact Tests” on page 7237.

Cramér–von Mises Test

The Cramér–von Mises statistic is defined as

$$\text{CM} = \frac{1}{n^2} \sum_i \left(n_i \sum_{j=1}^p t_j (F_i(x_j) - F(x_j))^2 \right)$$

where t_j is the number of ties at the j th distinct value and p is the number of distinct values. The asymptotic value is computed as

$$\text{CM}_a = \text{CM} \times n$$

PROC NPAR1WAY displays the contribution of each class level to the sum CM_a .

Kuiper Test

For data with two class levels, PROC NPAR1WAY computes the Kuiper statistic, its scaled value for the asymptotic distribution, and the asymptotic p -value. The Kuiper statistic is computed as

$$K = \max_j (F_1(x_j) - F_2(x_j)) - \min_j (F_1(x_j) - F_2(x_j)) \quad \text{where } j = 1, 2, \dots, n$$

The asymptotic value is

$$K_a = K \sqrt{n_1 n_2 / n}$$

PROC NPAR1WAY displays the value of $(\max_j |F_1(x_j) - F_2(x_j)|)$ for each class level.

The p -value for the Kuiper test is the probability of observing a larger value of K_a under the null hypothesis of no difference between the two classes. PROC NPAR1WAY computes this p -value according to Owen (1962, p. 441).

Exact Tests

PROC NPAR1WAY provides exact p -values for the following nonparametric location and scale tests: Wilcoxon, median, van der Waerden (normal), Savage, Siegel-Tukey, Ansari-Bradley, Klotz, Mood, and Conover. The procedure also provides exact p -values for tests that use the raw input data as scores. These exact tests are available for two-sample and multisample data. The two-sample exact tests are based on simple linear rank statistics, and the multisample exact tests are based on one-way ANOVA statistics.

For two-sample data, PROC NPAR1WAY also provides the exact Kolmogorov-Smirnov test, the exact ratio mean deviations (RMD) test, and exact Hodges-Lehmann confidence limits for the location shift.

Exact tests can be useful in situations where the asymptotic assumptions are not met and therefore the asymptotic p -values might not be close approximations for the true p -values. Standard asymptotic methods involve the assumption that the test statistic follows a particular distribution when the sample size is sufficiently large. When the sample size is not large, asymptotic results might not be valid. Asymptotic results might also be unreliable when the distribution of the data is sparse, skewed, or heavily tied. For more information, see Agresti (2007) and Bishop, Fienberg, and Holland (1975). Exact computations are based on the statistical theory of exact conditional inference for contingency tables, which is reviewed by Agresti (1992).

In addition to the computation of exact p -values, PROC NPAR1WAY provides the option to estimate exact p -values by Monte Carlo simulation. This can be useful for large problems where exact computations require a substantial amount of time and memory but asymptotic approximations might not be sufficient.

The following sections summarize the exact computational algorithms, define the exact p -values that PROC NPAR1WAY computes, discuss the computational resource requirements, and describe the Monte Carlo estimation option.

Computational Algorithms

PROC NPAR1WAY computes exact p -values for nonparametric location and scale tests by using the network algorithm, which was developed by Mehta and Patel (1983). This algorithm provides a substantial advantage over direct enumeration, which can be very time-consuming and feasible only for small problems. See Agresti (1992) for a review of algorithms for computation of exact p -values, and see Mehta, Patel, and Tsiatis (1984) and Mehta, Patel, and Senchaudhuri (1991) for information about the performance of the network algorithm.

To implement the network algorithm, PROC NPAR1WAY constructs a contingency table from the input data; the contingency table rows correspond to the levels of the classification variable and the table columns correspond to the levels of the response variable. The reference set for the observed contingency table includes all tables that have the same marginal row and column sums as the observed table. Corresponding to this reference set, the network algorithm forms a directed acyclic network that consists of nodes in a number of stages. A path through the network corresponds to a distinct table in the reference set. The distances between nodes are defined so that the total distance of a path through the network is the corresponding value of the test statistic. At each node, the algorithm computes the shortest and longest path distances for all paths that pass through that node. For the two-sample linear rank statistics, which can be expressed as linear combinations of cell frequencies multiplied by increasing row and column scores, PROC NPAR1WAY computes shortest and longest path distances by using the algorithm of Agresti, Mehta, and Patel (1990). For the multisample one-way test statistics, PROC NPAR1WAY computes an upper bound for the longest path and a lower bound for the shortest path by following the approach of Valz and Thompson (1994).

The longest and shortest path distances (bounds) for a node are compared to the value of the observed test statistic to determine whether all paths through the node contribute to the p -value, no paths through the node contribute to the p -value, or neither of these situations occurs. If all paths through the node contribute, the p -value is incremented accordingly, and these paths are eliminated from further analysis. If no paths contribute, these paths are eliminated from further analysis. Otherwise, the algorithm continues to process this node and the associated paths. The algorithm finishes when all nodes have been accounted for.

PROC NPAR1WAY performs the network algorithm by using full numerical precision to represent all statistics, row and column scores, and other quantities in the computations. Although it is possible to use rounding to improve the speed and memory requirements of the algorithm, PROC NPAR1WAY does not use rounding because it might reduce the accuracy of the results.

Definition of p -Values

For two-sample linear rank tests, PROC NPAR1WAY computes exact one-sided and two-sided p -values for each test that you specify in the EXACT statement. For one-sided tests, PROC NPAR1WAY displays the right-sided p -value when the observed value of the test statistic is greater than its expected value. The right-sided p -value is the sum of probabilities for those tables with a test statistic that is greater than or equal to the observed test statistic. Otherwise, when the observed test statistic is less than or equal to its expected value, PROC NPAR1WAY displays the left-sided p -value. The left-sided p -value is the sum of probabilities

for those tables with a test statistic that is less than or equal to the observed value. The one-sided p -value P_1 can be expressed as

$$P_1(t) = \begin{cases} \text{Prob}(\text{Test Statistic} \geq t) & \text{if } t > E_0(T) \\ \text{Prob}(\text{Test Statistic} \leq t) & \text{if } t \leq E_0(T) \end{cases}$$

where t is the observed value of the test statistic and $E_0(T)$ is the expected value of the test statistic under the null hypothesis. PROC NPAR1WAY computes the two-sided p -value as the sum of the one-sided p -value and the corresponding area in the opposite tail of the distribution of the statistic, equidistant from the expected value. The two-sided p -value P_2 can be expressed as

$$P_2(t) = \text{Prob}(|\text{Test Statistic} - E_0(T)| \geq |t - E_0(T)|)$$

Tests for multisample data are based on one-way ANOVA statistics. For a test of this form, large values of the test statistic indicate a departure from the null hypothesis; the test is inherently two-sided. The exact p -value is the sum of probabilities for those tables having a test statistic greater than or equal to the value of the observed test statistic.

If you specify the **POINT** option in the EXACT statement, PROC NPAR1WAY provides point probabilities for the exact tests. The point probability is the exact probability that the test statistic equals the observed value. For two-sample data, PROC NPAR1WAY provides point probabilities for the one-sided tests of the linear rank statistics. For multisample data, PROC NPAR1WAY provides point probabilities for the one-way ANOVA tests.

If you specify the **MIDP** option in the EXACT statement, PROC NPAR1WAY provides exact mid p -values. The exact mid p -value is defined as the exact p -value minus half the exact point probability, which equals the average of $\text{Prob}(\text{Test Statistic} \geq t)$ and $\text{Prob}(\text{Test Statistic} > t)$ for a right-sided test. The exact mid p -value is smaller and less conservative than the nonadjusted exact p -value. For more information, see Agresti (2013, section 1.1.4) and Hirji (2006, sections 2.5 and 2.11.1). For two-sample data, PROC NPAR1WAY provides mid p -values for the one-sided tests of the linear rank statistics. For multisample data, PROC NPAR1WAY provides mid p -values for the one-way ANOVA tests.

Computational Resources

PROC NPAR1WAY uses relatively fast and efficient algorithms for exact computations. These algorithms, together with improvements in computing power, make it feasible to perform exact computations for data where previously only asymptotic methods could be applied. Nevertheless, depending on your available computing resources, exact computations for some very large problems might require a prohibitive amount of time and memory. For such large problems, consider whether exact methods are really needed or whether asymptotic methods might give results that are very close to the exact results while requiring much less computing time and memory. When asymptotic methods might not be sufficient for such large problems, consider using Monte Carlo estimation of exact p -values, which is described in the section “[Monte Carlo Estimation](#)” on page 7240.

There is no formula that can predict in advance how much time and memory are needed to compute an exact p -value for a specific data set and test. The time and memory requirements depend on several factors, which include the following: the total number of observations, the number of rows (class levels) and columns (response levels), the particular arrangement of the observations, and the test to be performed. Generally, larger problems (in terms of total sample size, number of rows, and number of columns) tend to require more time and memory. For a fixed total sample size, time and memory requirements tend to increase as

the number of rows and columns increase because this corresponds to an increase in the number of tables in the reference set. For a fixed total sample size, time and memory requirements also tend to increase as the marginal row and column sums become more homogeneous. For more information, see Agresti, Mehta, and Patel (1990) and Gail and Mantel (1977).

While PROC NPAR1WAY is computing an exact p -value, you can terminate the computation by pressing the system interrupt key sequence and choosing to stop computations. For more information, see the *SAS Companion* for your system. After you terminate an exact computation, PROC NPAR1WAY completes all other remaining tasks. The procedure reports missing values for any exact p -values that were not computed before termination.

To limit the amount of time that PROC NPAR1WAY uses for exact computations, you can specify the **MAXTIME=** option in the EXACT statement. This option sets the maximum amount of clock time (in seconds) that PROC NPAR1WAY can use to compute an exact p -value. If PROC NPAR1WAY does not finish an exact computation in the time that you specify, the procedure terminates the computation and completes the remaining tasks.

Monte Carlo Estimation

When you specify the **MC** option in the EXACT statement, PROC NPAR1WAY computes Monte Carlo estimates of exact p -values. Monte Carlo estimation can be useful for large problems where exact computations require a substantial amount of time and memory but asymptotic approximations might not be sufficient.

To describe the precision of a Monte Carlo estimate, PROC NPAR1WAY provides the asymptotic standard error and $100(1 - \alpha)\%$ confidence limits. You can specify the confidence level α in the **ALPHA=** option in the EXACT statement; by default, ALPHA=0.01, which produces 99% confidence limits.

You can specify the number of Monte Carlo samples by using the **N=n** option in the EXACT statement. By default, PROC NPAR1WAY uses 10,000 samples to compute a Monte Carlo estimate. To improve the precision of the Monte Carlo estimates, you can specify a larger value of n ; this increases the computation time because more samples are generated. To reduce the computation time, you can specify a smaller value of n .

PROC NPAR1WAY computes a Monte Carlo estimate of an exact p -value by generating a random sample of tables from the reference set for the exact test. The reference set is based on the contingency table of the input data; the contingency table rows correspond to the levels of the classification variable and the table columns correspond to the levels of the response variable. The reference set includes all tables that have the same marginal row and column sums as the observed table.

PROC NPAR1WAY generates a random sample of tables from the reference set by using the algorithm of Agresti, Wackerly, and Boyett (1979), which generates tables in proportion to their hypergeometric probabilities conditional on the marginal frequencies. For each sample table, PROC NPAR1WAY computes the value of the test statistic and compares it to the value of the observed test statistic. To estimate a right-sided p -value, PROC NPAR1WAY counts all sample tables for which the test statistic is greater than or equal to the observed test statistic. The estimate of the p -value is the number of these tables divided by the total number of sample tables, which can be expressed as

$$\begin{aligned}\hat{p}_{\text{mc}} &= m / n \\ m &= \text{number of samples for which (test statistic} \geq t_o) \\ n &= \text{total number of samples} \\ t_o &= \text{observed test statistic}\end{aligned}$$

PROC NPARIWAY computes estimates of left-sided and two-sided exact p -values similarly. For left-sided exact p -values, PROC NPARIWAY evaluates whether the sample test statistics are less than or equal to the observed test statistic. For two-sided exact p -values, PROC NPARIWAY compares sample test statistics to the observed test statistic by using the definition of the two-sided p -value (P_2) for the test. For more information, see the section “Definition of p -Values” on page 7238 and descriptions of the individual tests.

The variable m has a binomial distribution with n trials and success probability p . The asymptotic standard error of the Monte Carlo estimate is

$$\text{se}(\hat{p}_{\text{mc}}) = \sqrt{\hat{p}_{\text{mc}}(1 - \hat{p}_{\text{mc}}) / (n - 1)}$$

PROC NPARIWAY constructs asymptotic confidence limits for the exact p -value as

$$\hat{p}_{\text{mc}} \pm (z_{\alpha/2} \times \text{se}(\hat{p}_{\text{mc}}))$$

where $z_{\alpha/2}$ is the $100(1 - \alpha/2)$ th percentile of the standard normal distribution, and the confidence level α is determined by the ALPHA= option in the EXACT statement.

When the Monte Carlo estimate $\hat{p}_{\text{mc}}$ is 0, PROC NPARIWAY computes confidence limits for the p -value as

$$(0, 1 - \alpha^{(1/n)})$$

When the Monte Carlo estimate $\hat{p}_{\text{mc}}$ is 1, PROC NPARIWAY computes confidence limits for the p -value as

$$(\alpha^{(1/n)}, 1)$$

Contents of the Output Data Set

The **OUTPUT** statement creates a SAS data set that contains statistics that PROC NPARIWAY computes. You identify which statistics to store in the output data set by specifying *output-options* in the OUTPUT statement. For more information, see the description of the **OUTPUT** statement.

The output data set contains one observation for each analysis variable for each **BY** group. (You can specify the analysis variables in the **VAR** statement.) The OUTPUT data set includes the following variables:

- BY variables, if you use a **BY** statement
- **_VAR_**, which identifies the analysis variable
- variables that contain the statistics

When you specify an *output-option* in the OUTPUT statement, the output data set contains the test statistic and associated values from the analysis that you specify. The associated values might include standardized statistics, one- and two-sided p -values, exact p -values, degrees of freedom, and confidence limits. If you request an exact test for the specified analysis by using the **EXACT** statement, the output data set includes the exact p -values. The statistics that are included also depend on the classification of the data. Some statistics are available only for two-sample data (where the **CLASS** variable groups the data into two classes); other statistics are available only for multisample data.

Table 90.5 lists variable names and descriptions or the statistics that are available in the output data set.

Monte Carlo estimates of exact p -values, multiple comparison statistics, and stratified analysis statistics are not available in this output data set. You can store any table that PROC NPAR1WAY produces in a SAS data set by using the Output Delivery System (ODS). For more information, see the section “ODS Table Names” on page 7255 and Chapter 23, “Using the Output Delivery System.”

Table 90.5 Output Data Set Variable Names and Descriptions

Output Option	Output Variables	Variable Descriptions
AB	_AB_	* Two-sample Ansari-Bradley statistic
	Z_AB	* Ansari-Bradley statistic, standardized
	PL_AB	* p -value (left-sided), Ansari-Bradley test
	PR_AB	* p -value (right-sided), Ansari-Bradley test
	P2_AB	* p -value (two-sided), Ansari-Bradley test
	XPL_AB	* Exact p -value (left-sided), Ansari-Bradley test
	XPR_AB	* Exact p -value (right-sided), Ansari-Bradley test
	XPT_AB	* Exact point probability, Ansari-Bradley test
	XMP_AB	* Exact mid p -value, Ansari-Bradley test
	XP2_AB	* Exact p -value (two-sided), Ansari-Bradley test
	CHAB	Ansari-Bradley chi-square
	DF_CHAB	Degrees of freedom, Ansari-Bradley chi-square
	P_CHAB	p -value, Ansari-Bradley chi-square test
	XP_CHAB	** Exact p -value, Ansari-Bradley chi-square test
	XPT_CHAB	** Exact point probability, Ansari-Bradley chi-square test
	XMP_CHAB	** Exact mid p -value, Ansari-Bradley chi-square test
ANOVA	_MSA_	ANOVA effect mean square, among MS
	MSE	ANOVA error mean square, within MS
	F	F statistic for ANOVA
	P_F	p -value, F statistic for ANOVA
CONOVER	_CON_	* Two-sample Conover statistic
	Z_CON	* Conover statistic, standardized
	PL_CON	* p -value (left-sided), Conover test
	PR_CON	* p -value (right-sided), Conover test
	P2_CON	* p -value (two-sided), Conover test
	XPL_CON	* Exact p -value (left-sided), Conover test
	XPR_CON	* Exact p -value (right-sided), Conover test
	XPT_CON	* Exact point probability, Conover test
	XMP_CON	* Exact mid p -value, Conover test
	XP2_CON	* Exact p -value (two-sided), Conover test
	CHCON	Conover chi-square
	DF_CHCON	Degrees of freedom, Conover chi-square
	P_CHCON	p -value, Conover chi-square test
	XP_CHCON	** Exact p -value, Conover chi-square test
	XPT_CHCO	** Exact point probability, Conover chi-square test
	XMP_CHCON	** Exact mid p -value, Conover chi-square test

Table 90.5 *continued*

Output Option	Output Variables	Variable Descriptions
EDF	_KS_	Kolmogorov-Smirnov statistic
	KSA	Kolmogorov-Smirnov statistic (asymptotic)
	Dp	* Two-sample Kolmogorov-Smirnov $D+$
	P_Dp	* p -value, $D+$
	Dm	* Two-sample Kolmogorov-Smirnov $D-$
	P_Dm	* p -value, $D-$
	D	* Two-sample Kolmogorov-Smirnov statistic D
	P_KSA	* p -value, D
	XP_Dp	* Exact p -value, $D+$
	XPT_Dp	* Exact point probability, $D+$
	XMP_Dp	* Exact mid p -value, $D+$
	XP_Dm	* Exact p -value, $D-$
	XPT_Dm	* Exact point probability, $D-$
	XMP_Dm	* Exact mid p -value, $D-$
	XP_D	* Exact p -value, D
	XPT_D	* Exact point probability, D
	XMP_D	* Exact mid p -value, D
	CM	Cramér-von Mises statistic
	CMA	Cramér-von Mises statistic (asymptotic)
	K	* Kuiper two-sample statistic
	KA	* Kuiper two-sample statistic (asymptotic)
	P_KA	* p -value, two-sample Kuiper test
FP	_FP_	* Fligner-Policello statistic
	PL_FP	* p -value (left-sided), Fligner-Policello test
	PR_FP	* p -value (right-sided), Fligner-Policello test
	P2_FP	* p -value (two-sided), Fligner-Policello test
HL	_HL_	* Hodges-Lehmann estimate, location shift
	L_HL	* Lower confidence limit, Hodges-Lehmann
	U_HL	* Upper confidence limit, Hodges-Lehmann
	M_HL	* Confidence limit midpoint, Hodges-Lehmann
	E_HL	* ASE of Hodges-Lehmann estimate
	XL_HL	* Exact lower confidence limit, Hodges-Lehmann
	XU_HL	* Exact upper confidence limit, Hodges-Lehmann
	XM_HL	* Exact confidence limit midpoint
KLOTZ	_KLOTZ_	* Two-sample Klotz statistic
	Z_K	* Klotz statistic, standardized
	PL_K	* p -value (left-sided), Klotz test
	PR_K	* p -value (right-sided), Klotz test
	P2_K	* p -value (two-sided), Klotz test
	XPL_K	* Exact p -value (left-sided), Klotz test
	XPR_K	* Exact p -value (right-sided), Klotz test

Table 90.5 continued

Output Option	Output Variables	Variable Descriptions
	XPT_K	* Exact point probability, Klotz test
	XMP_K	* Exact mid p -value, Klotz test
	XP2_K	* Exact p -value (two-sided), Klotz test
	CHK	Klotz chi-square
	DF_CHK	Degrees of freedom, Klotz chi-square
	P_CHK	p -value, Klotz chi-square test
	XP_CHK	** Exact p -value, Klotz chi-square test
	XPT_CHK	** Exact point probability, Klotz chi-square test
	XMP_CHK	** Exact mid p -value, Klotz chi-square test
MEDIAN	_MED_	* Two-sample median statistic
	Z_MED	* Median statistic, standardized
	PL_MED	* p -value (left-sided), median test
	PR_MED	* p -value (right-sided), median test
	P2_MED	* p -value (two-sided), median test
	XPL_MED	* Exact p -value (left-sided), median test
	XPR_MED	* Exact p -value (right-sided), median test
	XPT_MED	* Exact point probability, median test
	XMP_MED	* Exact mid p -value, median test
	XP2_MED	* Exact p -value (two-sided), median test
	CHMED	Median chi-square (Brown-Mood test)
	DF_CHMED	Degrees of freedom, median chi-square
	P_CHMED	p -value, median chi-square test
	XP_CHMED	** Exact p -value, median chi-square test
	XPT_CHME	** Exact point probability, median chi-square test
	XMP_CHMED	** Exact mid p -value, median chi-square test
MOOD	_MOOD_	* Two-sample Mood statistic
	Z_MOOD	* Mood statistic, standardized
	PL_MOOD	* p -value (left-sided), Mood test
	PR_MOOD	* p -value (right-sided), Mood test
	P2_MOOD	* p -value (two-sided), Mood test
	XPL_MOOD	* Exact p -value (left-sided), Mood test
	XPR_MOOD	* Exact p -value (right-sided), Mood test
	XPT_MOOD	* Exact point probability, Mood test
	XMP_MOOD	* Exact mid p -value, Mood test
	XP2_MOOD	* Exact p -value (two-sided), Mood test
	CHMOOD	Mood chi-square
	DF_CHMOO	Degrees of freedom, Mood chi-square
	P_CHMOOD	p -value, Mood chi-square test
	XP_CHMOO	** Exact p -value, Mood chi-square test
	XPT_CHMO	** Exact point probability, Mood chi-square test
	XMP_CHMOOD	** Exact mid p -value, Mood chi-square test

Table 90.5 *continued*

Output Option	Output Variables	Variable Descriptions
RMD	_RMD_	* Ratio mean deviations (RMD) statistic
	XPL_RMD	* Exact p -value (left-sided), RMD test
	XPR_RMD	* Exact p -value (right-sided), RMD test
	XPT_RMD	* Exact point probability, RMD test
	XMP_RMD	* Exact mid p -value, RMD test
	RMD2	* RMD2 two-sided statistic
	XP2_RMD	* Exact p -value (two-sided), RMD test
SAVAGE	_SAV_	* Two-sample Savage statistic
	Z_SAV	* Savage statistic, standardized
	PL_SAV	* p -value (left-sided), Savage test
	PR_SAV	* p -value (right-sided), Savage test
	P2_SAV	* p -value (two-sided), Savage test
	XPL_SAV	* Exact p -value (left-sided), Savage test
	XPR_SAV	* Exact p -value (right-sided), Savage test
	XPT_SAV	* Exact point probability, Savage test
	XMP_SAV	* Exact mid p -value, Savage test
	XP2_SAV	* Exact p -value (two-sided), Savage test
	CHSAV	Savage chi-square
	DF_CHSAV	Degrees of freedom, Savage chi-square
	P_CHSAV	p -value, Savage chi-square test
	XP_CHSAV	** Exact p -value, Savage chi-square test
	XPT_CHSA	** Exact point probability, Savage chi-square test
	XMP_CHSAV	** Exact mid p -value, Savage chi-square test
SCORES=DATA	_DATA_	* Two-sample data scores statistic
	Z_DATA	* Data scores statistic, standardized
	PL_DATA	* p -value (left-sided), data scores test
	PR_DATA	* p -value (right-sided), data scores test
	P2_DATA	* p -value (two-sided), data scores test
	XPL_DATA	* Exact p -value (left-sided), data scores test
	XPR_DATA	* Exact p -value (right-sided), data scores test
	XPT_DATA	* Exact point probability, data scores test
	XMP_DATA	* Exact mid p -value, data scores test
	XP2_DATA	* Exact p -value (two-sided), data scores test
	CHDATA	Data scores chi-square
	DF_CHDAT	Degrees of freedom, data scores chi-square
	P_CHDATA	p -value, data scores chi-square test
	XP_CHDAT	** Exact p -value, data scores chi-square test
	XPT_CHDA	** Exact point probability, data scores chi-square test
	XMP_CHDATA	** Exact mid p -value, data scores chi-square test
ST	_ST_	* Two-sample Siegel-Tukey statistic
	Z_ST	* Siegel-Tukey statistic, standardized
	PL_ST	* p -value (left-sided), Siegel-Tukey test

Table 90.5 continued

Output Option	Output Variables	Variable Descriptions
	PR_ST	* p -value (right-sided), Siegel-Tukey test
	P2_ST	* p -value (two-sided), Siegel-Tukey test
	XPL_ST	* Exact p -value (left-sided), Siegel-Tukey test
	XPR_ST	* Exact p -value (right-sided), Siegel-Tukey test
	XPT_ST	* Exact point probability, Siegel-Tukey test
	XMP_ST	* Exact mid p -value, Siegel-Tukey test
	XP2_ST	* Exact p -value (two-sided), Siegel-Tukey test
	CHST	Siegel-Tukey chi-square
	DF_CHST	Degrees of freedom, Siegel-Tukey chi-square
	P_CHST	p -value, Siegel-Tukey chi-square test
	XP_CHST	** Exact p -value, Siegel-Tukey chi-square test
	XPT_CHST	** Exact point probability, Siegel-Tukey chi-square test
	XMP_CHST	** Exact mid p -value, Siegel-Tukey chi-square test
VW NORMAL	_VW_	* Two-sample Van der Waerden statistic
	Z_VW	* Van der Waerden statistic, standardized
	PL_VW	* p -value (left-sided), Van der Waerden test
	PR_VW	* p -value (right-sided), Van der Waerden test
	P2_VW	* p -value (two-sided), Van der Waerden test
	XPL_VW	* Exact p -value (left-sided), Van der Waerden test
	XPR_VW	* Exact p -value (right-sided), Van der Waerden test
	XPT_VW	* Exact point probability, Van der Waerden test
	XMP_VW	* Exact mid p -value, Van der Waerden test
	XP2_VW	* Exact p -value (two-sided), Van der Waerden test
	CHVW	Van der Waerden chi-square
	DF_CHVW	Degrees of freedom, Van der Waerden chi-square
	P_CHVW	p -value, Van der Waerden chi-square test
	XP_CHVW	** Exact p -value, Van der Waerden chi-square test
	XPT_CHVW	** Exact point probability, Van der Waerden chi-square test
	XMP_CHVW	** Exact mid p -value, Van der Waerden chi-square test
WILCOXON	_WIL_	* Two-sample Wilcoxon statistic
	Z_WIL	* Wilcoxon statistic, standardized
	PL_WIL	* p -value (left-sided), Wilcoxon test
	PR_WIL	* p -value (right-sided), Wilcoxon test
	P2_WIL	* p -value (two-sided), Wilcoxon test
	PTL_WIL	* p -value (left-sided), Wilcoxon t approximation
	PTR_WIL	* p -value (right-sided), Wilcoxon t approximation
	PT2_WIL	* p -value (two-sided), Wilcoxon t approximation
	XPL_WIL	* Exact p -value (left-sided), Wilcoxon test
	XPR_WIL	* Exact p -value (right-sided), Wilcoxon test
	XPT_WIL	* Exact point probability, Wilcoxon test
	XMP_WIL	* Exact mid p -value, Wilcoxon test
	XP2_WIL	* Exact p -value (two-sided), Wilcoxon test

Table 90.5 *continued*

Output Option	Output Variables	Variable Descriptions
	KW	Kruskal-Wallis statistic
	DF_KW	Degrees of freedom, Kruskal-Wallis test
	P_KW	<i>p</i> -value, Kruskal-Wallis test
	XP_KW	** Exact <i>p</i> -value, Kruskal-Wallis test
	XPT_KW	** Exact point probability, Kruskal-Wallis test
	XMP_KW	** Exact mid <i>p</i> -value, Kruskal-Wallis test

*Statistic included only for two-sample data.

**Statistic included only for multisample data.

Displayed Output

If you specify the **VARINFO** option, PROC NPAR1WAY displays a “Variable Information” table for each analysis variable that you specify in the **VAR** statement. This table includes the following information:

- Number of Observations Used
- Number of Missing Observations
- Sum of Frequencies Used, if you specify a **FREQ** statement
- Number of Class Levels Used
- Number of Empty Class Levels Excluded
- Number of Strata, if you specify the **ALIGN=STRATA** option
- Number of Variable Levels, if you specify the **VARINFO(TIES)** option
- Number of Tied Observations, if you specify the **VARINFO(TIES)** option

If you specify the **ANOVA** option, PROC NPAR1WAY displays a “Class Means” table and an “Analysis of Variance” table for each response variable. The “Class Means” table includes the following information for each **CLASS** variable value (level):

- N, which is the number of observations
- Mean of the response variable

The “Analysis of Variance” table includes the following information for each Source of variation (Among classes and Within classes):

- DF, which is the degrees of freedom associated with the source

- Sum of Squares
- Mean Square, which is the sum of squares divided by the degrees of freedom

The “Analysis of Variance” table also includes the following:

- F Value for testing the hypothesis that the class means are equal, which is computed by dividing the Mean Square (Among) by the Mean Square (Within)
- $\text{Pr} > F$, which is the significance probability corresponding to the F Value

For each score type that you specify, PROC NPAR1WAY displays a “Class Scores” table. The available score types include Wilcoxon, median, Van der Waerden (normal), Savage, Siegel-Tukey, Ansari-Bradley, Klotz, Mood, Conover, and raw data scores. PROC NPAR1WAY computes the scores for the response variable values and classifies the scored observations according to the **CLASS** variable values. The “Class Scores” table includes the following information for each CLASS variable level:

- N, which is the number of observations
- Sum of Scores
- Expected Under H_0 , which is the expected sum of scores under the null hypothesis of no difference among classes
- Std Dev Under H_0 , which is the standard deviation under the null hypothesis
- Mean Score

When there are two levels of the **CLASS** variable, PROC NPAR1WAY displays a “Two-Sample Test” table for each analysis of scores. The “Two-Sample Test” table includes the following information:

- Statistic, which is the sum of scores for the class with the smaller sample size
- Z, which is the standardized test statistic and has an asymptotic standard normal distribution under the null hypothesis
- One-Sided $\text{Pr} < Z$ or One-Sided $\text{Pr} > Z$, which is the asymptotic one-sided p -value. This is displayed as $\text{Pr} < Z$ or $\text{Pr} > Z$, depending on whether Z is ≤ 0 or > 0 .
- Two-Sided $\text{Pr} > |Z|$, which is the asymptotic two-sided p -value

For Wilcoxon scores, the “Two-Sample Test” table also includes a t Approximation for the Wilcoxon two-sample test.

If you request an exact test by specifying the score type in the **EXACT** statement, the “Two-Sample Test” table also includes the following exact p -values:

- One-Sided $\text{Pr} \leq S$ or One-Sided $\text{Pr} \geq S$, which is the exact one-sided p -value. This is displayed as $\text{Pr} \leq S$ or $\text{Pr} \geq S$, depending on whether $S \leq \text{Mean}$ or $S > \text{Mean}$, where S is the test statistic and Mean is its expected value under the null hypothesis.

- Point $Pr = S$, which is the exact point probability. This is displayed if you specify the **POINT** option in the EXACT statement.
- Mid p-Value, which is displayed if you specify the **MIDP** option in the EXACT statement
- Two-Sided $Pr \geq |S - \text{Mean}|$, which is the exact two-sided p -value

If you request Monte Carlo estimates for a two-sample exact test by specifying the **MC** option in the EXACT statement, PROC NPAR1WAY displays the “Monte Carlo Estimates for the Exact Test” table, which includes the following information:

- Estimate of One-Sided $Pr \leq S$ or One-Sided $Pr \geq S$, which is the exact one-sided p -value, together with the Lower and Upper Confidence Limits
- Estimate of Two-Sided $Pr \geq |S - \text{Mean}|$, which is the exact two-sided p -value, together with the Lower and Upper Confidence Limits
- Number of Samples used to compute the Monte Carlo estimates
- Initial Seed used to compute the Monte Carlo estimates

For both two-sample and multisample data, PROC NPAR1WAY displays a “One-Way Analysis” table, which includes the following information:

- Chi-Square, which is the one-way ANOVA statistic for testing the null hypothesis of no difference among classes
- DF, which is the degrees of freedom
- $Pr > \text{Chi-Square}$, which is the asymptotic p -value

For multisample data, if you request an exact test by specifying the score type in the EXACT statement, the “One-Way Analysis” table also displays the exact p -value as follows:

- Exact $Pr \geq \text{Chi-Square}$
- Exact $Pr = \text{Chi-Square}$, which is the point probability. This is displayed if you specify the **POINT** option in the EXACT statement.
- Exact Mid p-Value, which is displayed if you specify the **MIDP** option in the EXACT statement

For multisample data, if you specify the **MC** option in the EXACT statement, PROC NPAR1WAY displays the following information in the “Monte Carlo Estimate for the Exact Test” table:

- Estimate of Exact $Pr \geq \text{Chi-Square}$, together with the Lower and Upper Confidence Limits
- Number of Samples used to compute the Monte Carlo estimate
- Initial Seed used to compute the Monte Carlo estimate

If you specify the **HL** option for two-sample data, PROC NPAR1WAY produces a “Hodges-Lehmann Estimation” table, which includes the following information:

- Location Shift estimate
- Confidence Limits for the Location Shift
- Confidence Interval Midpoint
- Asymptotic Standard Error estimate, which is based on the confidence interval

If you request exact Hodges-Lehmann confidence limits by specifying the **HL** option in the EXACT statement, the “Hodges-Lehmann Estimation” table also includes Exact Confidence Limits and the exact Interval Midpoint.

If you specify the **FP** option for two-sample data, PROC NPAR1WAY produces the “Fligner-Policello Placements” table and the “Fligner-Policello Test” table. The “Fligner-Policello Placements” table includes the following information for each **CLASS** variable level:

- N, which is the number of observations
- Sum, which is the sum of the placements
- Mean, which is the average placement
- Std Dev, which is the standard deviation of the placements

The “Fligner-Policello Test” table includes the following information:

- Difference, which is the difference in placement sums between the two **CLASS** levels (samples)
- Statistic (Z), which is the standardized test statistic and has an asymptotic standard normal distribution under the null hypothesis
- One-Sided $\Pr < Z$ or One-Sided $\Pr > Z$, which is the one-sided p -value. This is displayed as $\Pr < Z$ or $\Pr > Z$, depending on whether Z is ≤ 0 or > 0 .
- Two-Sided $\Pr > |Z|$, which is the two-sided p -value

If you specify the **RMD** option for two-sample data, PROC NPAR1WAY produces the “Deviations” table and the “Ratio Mean Deviations (RMD) Exact Test” table. The “Deviations” table includes the following information for each **CLASS** variable level:

- N, which is the number of observations
- Median, which is the sample median
- Sum of Deviations, which is the sum of the absolute deviations from the median
- Mean Deviation, which is the average of the absolute deviations

The “Ratio Mean Deviations (RMD) Exact Test” table includes the following information:

- RMD, which is the ratio mean deviations statistic
- $\Pr \leq \text{RMD}$ or $\Pr \geq \text{RMD}$, which is the exact one-sided p -value. PROC NPAR1WAY displays $\Pr \leq \text{RMD}$ when RMD is ≤ 1 and $\Pr \geq \text{RMD}$ otherwise.
- RMD2, which is the modified RMD two-sided statistic
- $\Pr \geq \text{RMD2}$, which is the exact two-sided p -value

If you specify the [DSCF](#) option for multisample data, PROC NPAR1WAY produces the “Pairwise Two-Sided Multiple Comparison Analysis” table, which includes the following information for each two-sample comparison:

- Comparison, which identifies the two CLASS levels that are compared
- Wilcoxon Z, which is the standardized two-sample Wilcoxon statistic
- DSCF Value, which is the Dwass, Steel, Critchlow-Fligner statistic
- $\Pr > \text{DSCF}$, which is the two-sided p -value

If you specify the [EDF](#) option, PROC NPAR1WAY produces tables for the Kolmogorov-Smirnov test, the Cramér-von Mises test, and for two-sample data only, the Kuiper test.

The “Kolmogorov-Smirnov Test” table includes the following information for each [CLASS](#) variable level:

- N, which is the number of observations
- EDF at Maximum, which is the value of the class EDF (empirical distribution function) at its maximum deviation from the pooled EDF
- Deviation from Mean at Maximum, which is the value of $\sqrt{n_i} (F_i - F)$ at its maximum, where n_i is the class sample size, F_i is the class EDF, and F is the pooled EDF

The “Kolmogorov-Smirnov Test” table displays the following statistics:

- KS, which is the Kolmogorov-Smirnov statistic
- KSa, which is the asymptotic Kolmogorov-Smirnov statistic, $\text{KSa} = \sqrt{n} \text{KS}$

For two-sample data, the “Kolmogorov-Smirnov Test” table also displays the following statistics:

- $\Pr > \text{KSa}$, which is the asymptotic p -value for KSa and equals $\Pr > D$
- D, which is the two-sample Kolmogorov-Smirnov statistic, $\max_j |F_1(x_j) - F_2(x_j)|$

If you specify the **D** option for two-sample data, PROC NPAR1WAY displays the following one-sided Kolmogorov-Smirnov statistics and their asymptotic p -values in the “Kolmogorov-Smirnov Two-Sample Test” table:

- $D+$, which is $\max_j (F_1(x_j) - F_2(x_j))$
- $\Pr > D+$
- $D-$, which is $\max_j (F_2(x_j) - F_1(x_j))$
- $\Pr > D-$

For two-sample data, if you request exact Kolmogorov-Smirnov tests by specifying the **KS** option in the EXACT statement, PROC NPAR1WAY displays the following exact p -values in the “Kolmogorov-Smirnov Two-Sample Test” table:

- Exact $\Pr \geq D$
- Exact $\Pr \geq D+$
- Exact $\Pr \geq D-$
- Exact Point $\Pr = D$, Exact Point $\Pr = D+$, and Exact Point $\Pr = D-$, if you specify the **POINT** option in the EXACT statement
- Exact Mid p -Values for D , $D+$, and $D-$, if you specify the **MIDP** option in the EXACT statement

If you request Monte Carlo estimates for the two-sample exact Kolmogorov-Smirnov test, PROC NPAR1WAY displays the following information in the “Kolmogorov-Smirnov Two-Sample Test” table:

- Estimate of Exact $\Pr \geq D$, together with the Lower and Upper Confidence Limits
- Estimate of Exact $\Pr \geq D+$, together with the Lower and Upper Confidence Limits
- Estimate of Exact $\Pr \geq D-$, together with the Lower and Upper Confidence Limits
- Number of Samples used to compute the Monte Carlo estimates
- Initial Seed used to compute the Monte Carlo estimates

The “Cramér–von Mises Test” table includes the following information for each CLASS variable level:

- N , which is the number of observations
- Summed Deviation from Mean, which is $(n_i/n) \sum_{j=1}^p t_j (F_i(x_j) - F(x_j))^2$

The “Cramér–von Mises Statistics” table displays the following statistics:

- CM , which is the Cramér–von Mises statistic

- C_{Ma}, which is the asymptotic Cramér–von Mises statistic, $C_{Ma} = n \text{ CM}$

For two-sample data, PROC NPAR1WAY displays the “Kuiper Test” table, which includes the following information for each CLASS variable level:

- N, which is the number of observations
- Deviation from Mean, which is $\max_j |F_1(x_j) - F_2(x_j)|$

The “Kuiper Two-Sample Statistics” table displays the following statistics:

- K, which is the Kuiper two-sample test statistic
- K_a, which is the asymptotic Kuiper two-sample test statistic, $K_a = K \sqrt{n_1 n_2 / n}$
- Pr > K_a

If you specify a STRATA statement along with the ALIGN=STRATA option in the PROC NPAR1WAY statement, PROC NPAR1WAY displays a “Stratum Alignment” table for each analysis variable. This table includes the following information for each stratum:

- Stratum index, which is a sequential stratum identification number
- STRATA variables, which list the levels of STRATA variables
- N Obs, which is the number of observations in the stratum
- Median, Mean, or HL Shift

If you specify a STRATA statement and do not specify the ALIGN=STRATA option, PROC NPAR1WAY displays the “Stratum Information,” “Stratified Analysis Method,” and “Class Information” tables.

The “Stratum Information” table displays the following information for each stratum:

- Stratum index, which is a sequential stratum identification number
- STRATA variables, which list the levels of STRATA variables
- N Obs, which is the number of observations in the stratum

The “Stratified Analysis Method” table displays the following information:

- Stratum Weights (Stratum Size or Equal)
- Ranks (Within Strata or Overall)

The “Class Information” tables displays the following information for each class:

- Class index, which is a sequential class identification number

- CLASS variable level
- N Obs, which is the number of observations in the class

If you request stratified analysis by using the **STRATA** statement, PROC NPAR1WAY displays the “Scores by Strata” and “Stratified Test” tables for each analysis variable and score type that you specify.

The “Scores by Strata” table includes the following information for each stratum:

- Stratum index, which is a sequential stratum identification number
- N Obs, which is the number of observations in the stratum
- Class N Obs, which is the number of observations in class 1
- Statistic, which is the sum of scores in class 1
- Expected, which is the expected value under H_0
- Std Dev, which is the standard deviation under H_0
- Mean score

The “Stratified Test” table includes the following information:

- N Obs, which is the total number of observations
- N Strata, which is the number of strata
- Statistic, which is the stratified sum of scores
- Expected, which is the expected value under H_0
- Std Dev, which is the standard deviation under H_0
- Z, which is the standardized test statistic
- $\Pr < Z$ or $\Pr > Z$, which is the one-sided p -value
- $\Pr > |Z|$, which is the two-sided p -value

ODS Table Names

PROC NPAR1WAY assigns a name to each table that it creates. You can use these names to refer to tables when you use the Output Delivery System (ODS) to select tables and create output data sets. For more information about ODS, see Chapter 23, “Using the Output Delivery System.”

Table 90.6 lists the ODS table names together with their descriptions and the options required to produce the tables. If you do not specify any analysis options in the PROC NPAR1WAY statement and do not specify a STRATA statement, PROC NPAR1WAY provides the ANOVA, WILCOXON, MEDIAN, VW (NORMAL), SAVAGE, and EDF analyses by default.

Table 90.6 ODS Tables Produced by PROC NPAR1WAY

ODS Table Name	Description	Statement	Option
ABAnalysis	Ansari-Bradley one-way analysis	PROC	AB
ABMC	Monte Carlo estimates for the Ansari-Bradley exact test	EXACT	AB / MC
ABScores	Ansari-Bradley scores	PROC	AB
ABTest*	Ansari-Bradley two-sample test	PROC	AB
ANOVA	Analysis of variance	PROC	ANOVA
ClassInfo*	Class information	STRATA	
ClassMeans	Class means	PROC	ANOVA
ConoverAnalysis	Conover one-way analysis	PROC	CONOVER
ConoverMC	Monte Carlo estimates for the Conover exact test	EXACT	CONOVER / MC
ConoverScores	Conover scores	PROC	CONOVER
ConoverTest*	Conover two-sample test	PROC	CONOVER
CVMStats	Cramér–von Mises statistics	PROC	EDF
CVMTest	Cramér–von Mises test	PROC	EDF
DataScores	Data scores	PROC	SCORES=DATA
DataScoresAnalysis	Data scores one-way analysis	PROC	SCORES=DATA
DataScoresMC	Monte Carlo estimates for the data scores exact test	EXACT	SCORES=DATA / MC
DataScoresStrata*	Data scores by strata	STRATA	SCORES=DATA
DataScoresStrataTest*	Stratified data scores test	STRATA	SCORES=DATA
DataScoresTest*	Data scores two-sample test	PROC	SCORES=DATA
DSCF**	DSCF multiple comparison analysis	PROC	DSCF
FPPlacements*	Fligner-Policello placements	PROC	FP
FPTTest*	Fligner-Policello test	PROC	FP
HodgesLehmann*	Hodges-Lehmann estimation	PROC	HL
KlotzAnalysis	Klotz one-way analysis	PROC	KLOTZ
KlotzMC	Monte Carlo estimates for the Klotz exact test	EXACT	KLOTZ / MC
KlotzScores	Klotz scores	PROC	KLOTZ
KlotzTest*	Klotz two-sample test	PROC	KLOTZ
KruskalWallisMC**	Monte Carlo estimates for the Kruskal-Wallis exact test	EXACT	WILCOXON / MC

Table 90.6 continued

ODS Table Name	Description	Statement	Option
KruskalWallisTest	Kruskal-Wallis test	PROC	WILCOXON
KS2Stats*	Kolmogorov-Smirnov two-sample statistics	PROC	EDF
KSExactTest*	Kolmogorov-Smirnov exact test	EXACT	KS EDF
KSMC*	Monte Carlo estimates for the Kolmogorov-Smirnov exact test	EXACT	KS EDF / MC
KSStats**	Kolmogorov-Smirnov statistics	PROC	EDF
KSTest	Kolmogorov-Smirnov test	PROC	EDF
KuiperStats*	Kuiper two-sample statistics	PROC	EDF
KuiperTest*	Kuiper test	PROC	EDF
MedianAnalysis	Median one-way analysis	PROC	MEDIAN
MedianMC	Monte Carlo estimates for the median exact test	EXACT	MEDIAN / MC
MedianScores	Median scores	PROC	MEDIAN
MedianStrata*	Median scores by strata	STRATA	MEDIAN
MedianStrataTest*	Stratified median test	STRATA	MEDIAN
MedianTest*	Median two-sample test	PROC	MEDIAN
MoodAnalysis	Mood one-way analysis	PROC	MOOD
MoodMC	Monte Carlo estimates for the Mood exact test	EXACT	MOOD / MC
MoodScores	Mood scores	PROC	MOOD
MoodTest*	Mood two-sample test	PROC	MOOD
RMDDeviations*	RMD deviations	PROC or EXACT	RMD RMD
RMDMC*	Monte Carlo estimates for the RMD exact test	EXACT	RMD / MC
RMDTest*	RMD two-sample test	PROC or EXACT	RMD RMD
SavageAnalysis	Savage one-way analysis	PROC	SAVAGE
SavageMC	Monte Carlo estimates for the Savage exact test	EXACT	SAVAGE / MC
SavageScores	Savage scores	PROC	SAVAGE
SavageStrata*	Savage scores by strata	STRATA	SAVAGE
SavageStrataTest*	Stratified Savage test	STRATA	SAVAGE
SavageTest*	Savage two-sample test	PROC	SAVAGE
STAnalysis	Siegel-Tukey one-way analysis	PROC	ST
STMC	Monte Carlo estimates for the Siegel-Tukey exact test	EXACT	ST / MC
StrataAlign	Stratum alignment	PROC	ALIGN=STRATA
StrataInfo	Stratum information	STRATA	
StrataMethod*	Stratified analysis method	STRATA	
STScores	Siegel-Tukey scores	PROC	ST
STTest*	Siegel-Tukey two-sample test	PROC	ST
VarInfo	Variable information	PROC	VARINFO

Table 90.6 *continued*

ODS Table Name	Description	Statement	Option
VWAnalysis	Van der Waerden one-way analysis	PROC	VW NORMAL
VWMC	Monte Carlo estimates for the Van der Waerden exact test	EXACT	VW NORMAL / MC
VWScores	Van der Waerden scores	PROC	VW NORMAL
VWStrata*	Van der Waerden scores by strata	STRATA	VW NORMAL
VWStrataTest*	Stratified Van der Waerden test	STRATA	VW NORMAL
VWTest*	Van der Waerden two-sample test	PROC	VW NORMAL
WilcoxonMC*	Monte Carlo estimates for the Wilcoxon two-sample exact test	EXACT	WILCOXON / MC
WilcoxonScores	Wilcoxon scores	PROC	WILCOXON
WilcoxonStrata*	Wilcoxon scores by strata	STRATA	WILCOXON
WilcoxonStrataTest*	Stratified Wilcoxon test	STRATA	WILCOXON
WilcoxonTest*	Wilcoxon two-sample test	PROC	WILCOXON

*Table produced only for two-sample data.

**Table produced only for multisample data.

ODS Graphics

Statistical procedures use ODS Graphics to create graphs as part of their output. ODS Graphics is described in detail in Chapter 24, “[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 663 in Chapter 24, “[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 662 in Chapter 24, “[Statistical Graphics Using ODS](#).”

When ODS Graphics is enabled, you can request specific plots by specifying the **PLOTS=** option in the PROC NPARIWAY statement. If you do not specify the PLOTS= option but have enabled ODS Graphics, PROC NPARIWAY produces all plots that are associated with the analyses that you request.

PROC NPARIWAY assigns a name to each graph that it creates by using ODS Graphics. You can use these names to refer to the graphs. [Table 90.7](#) lists the names of the graphs that PROC NPARIWAY generates together with their descriptions and the options that are required to produce the graphs.

Table 90.7 Graphs Produced by PROC NPAR1WAY

ODS Graph Name	Description	Option
ABBoxPlot	Box plot of Ansari-Bradley scores	AB
ANOVABoxPlot	Box plot of raw data	ANOVA
ConoverBoxPlot	Box plot of Conover scores	CONOVER
DataScoresBoxPlot	Box plot of data scores	SCORES=DATA
EDFPlot	Empirical distribution function plot	EDF
FPBoxPlot	Box plot of Fligner-Policello placements	FP
KlotzBoxPlot	Box plot of Klotz scores	KLOTZ
MedianPlot	Median plot	MEDIAN
MoodBoxPlot	Box plot of Mood scores	MOOD
SavageBoxPlot	Box plot of Savage scores	SAVAGE
STBoxPlot	Box plot of Siegel-Tukey scores	ST
VWBoxPlot	Box plot of Van der Waerden scores	VW NORMAL
WilcoxonBoxPlot	Box plot of Wilcoxon scores	WILCOXON

Examples: NPAR1WAY Procedure

Example 90.1: Two-Sample Location Tests and Plots

Fifty-nine female patients with rheumatoid arthritis who participated in a clinical trial were assigned to two groups, active and placebo. The response status (excellent=5, good=4, moderate=3, fair=2, poor=1) of each patient was recorded.

The following SAS statements create the data set *Arthritis*, which contains the observed status values for all the patients. The variable *Treatment* denotes the treatment received by a patient, and the variable *Response* contains the response status of the patient. The variable *Freq* contains the frequency of the observation, which is the number of patients with the *Treatment* and *Response* combination.

```
data Arthritis;
  input Treatment $ Response Freq @@;
  datalines;
Active 5 5 Active 4 11 Active 3 5 Active 2 1 Active 1 5
Placebo 5 2 Placebo 4 4 Placebo 3 7 Placebo 2 7 Placebo 1 12
;
```

The following PROC NPAR1WAY statements test the null hypothesis that there is no difference in the patient response status against the alternative hypothesis that the patient response status differs in the two treatment groups. The WILCOXON option requests the Wilcoxon test for difference in location, and the MEDIAN option requests the median test for difference in location. The variable *Treatment* is the CLASS variable, and the VAR statement specifies that the variable *Response* is the analysis variable.

The PLOTS= option requests a box plot of the Wilcoxon scores and a median plot for *Response* classified by *Treatment*. ODS Graphics must be enabled before producing plots.

```
ods graphics on;
proc npar1way data=Arthritis wilcoxon median
      plots=(wilcoxonboxplot medianplot);
  class Treatment;
  var Response;
  freq Freq;
run;
ods graphics off;
```

Output 90.1.1 shows the results of the Wilcoxon analysis. The Wilcoxon two-sample test statistic equals 999.0, which is the sum of the Wilcoxon scores for the smaller sample (Active). This sum is greater than 810.0, which is the expected value under the null hypothesis of no difference between the two samples, Active and Placebo. The one-sided p -value is 0.0016, which indicates that the patient response for the Active treatment is significantly more than for the Placebo group.

Output 90.1.1 Wilcoxon Two-Sample Test

The NPAR1WAY Procedure

Wilcoxon Scores (Rank Sums) for Variable Response Classified by Variable Treatment

Treatment	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
Active	27	999.0	810.0	63.972744	37.000000
Placebo	32	771.0	960.0	63.972744	24.093750

Average scores were used for ties.

Wilcoxon Two-Sample Test

Statistic	t Approximation				
	Z	Pr > Z	Pr > Z	Pr > Z	Pr > Z
999.0000	2.9466	0.0016	0.0032	0.0023	0.0046

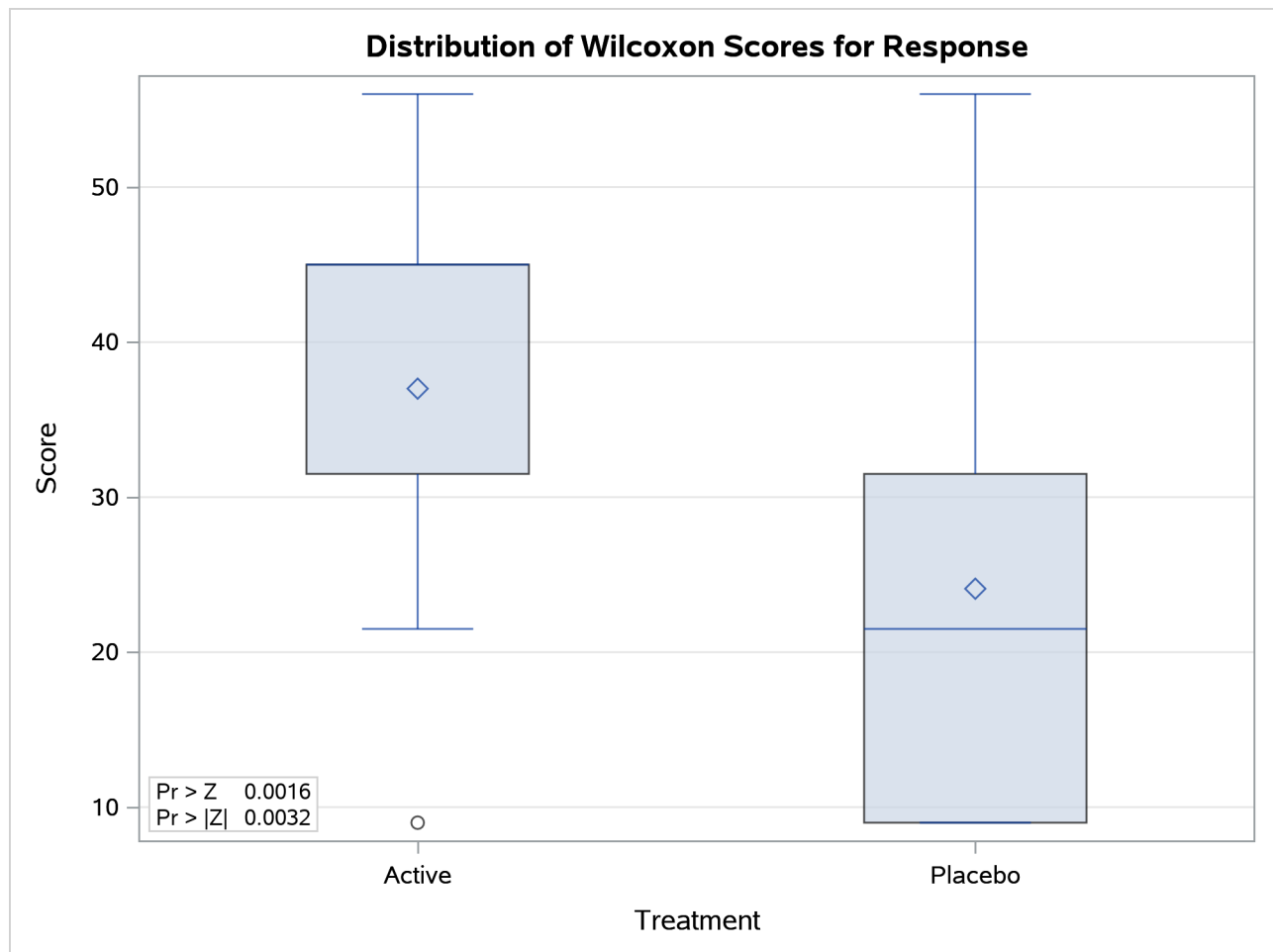
Z includes a continuity correction of 0.5.

Kruskal-Wallis Test

Chi-Square	DF	Pr > ChiSq
8.7284	1	0.0031

Output 90.1.2 displays the box plot of Wilcoxon scores classified by Treatment, which corresponds to the Wilcoxon analysis in Output 90.1.1. To remove the p -values from the box plot display, you can specify the NOSTATS plot option in parentheses after the WILCOXONBOXPLOT option.

Output 90.1.2 Box Plot of Wilcoxon Scores



Output 90.1.3 shows the results of the median two-sample test. The value of the test statistic is 18.9167, and its standardized Z value is 3.2667. The one-sided p -value $\text{Pr} > Z$ is 0.0005. This supports the alternative hypothesis that the effect of the Active treatment is greater than that of the Placebo.

Output 90.1.4 displays the median plot for the analysis of Response classified by Treatment. The median plot is a stacked bar chart showing the frequencies above and below the overall median. This plot corresponds to the median scores analysis in Output 90.1.3.

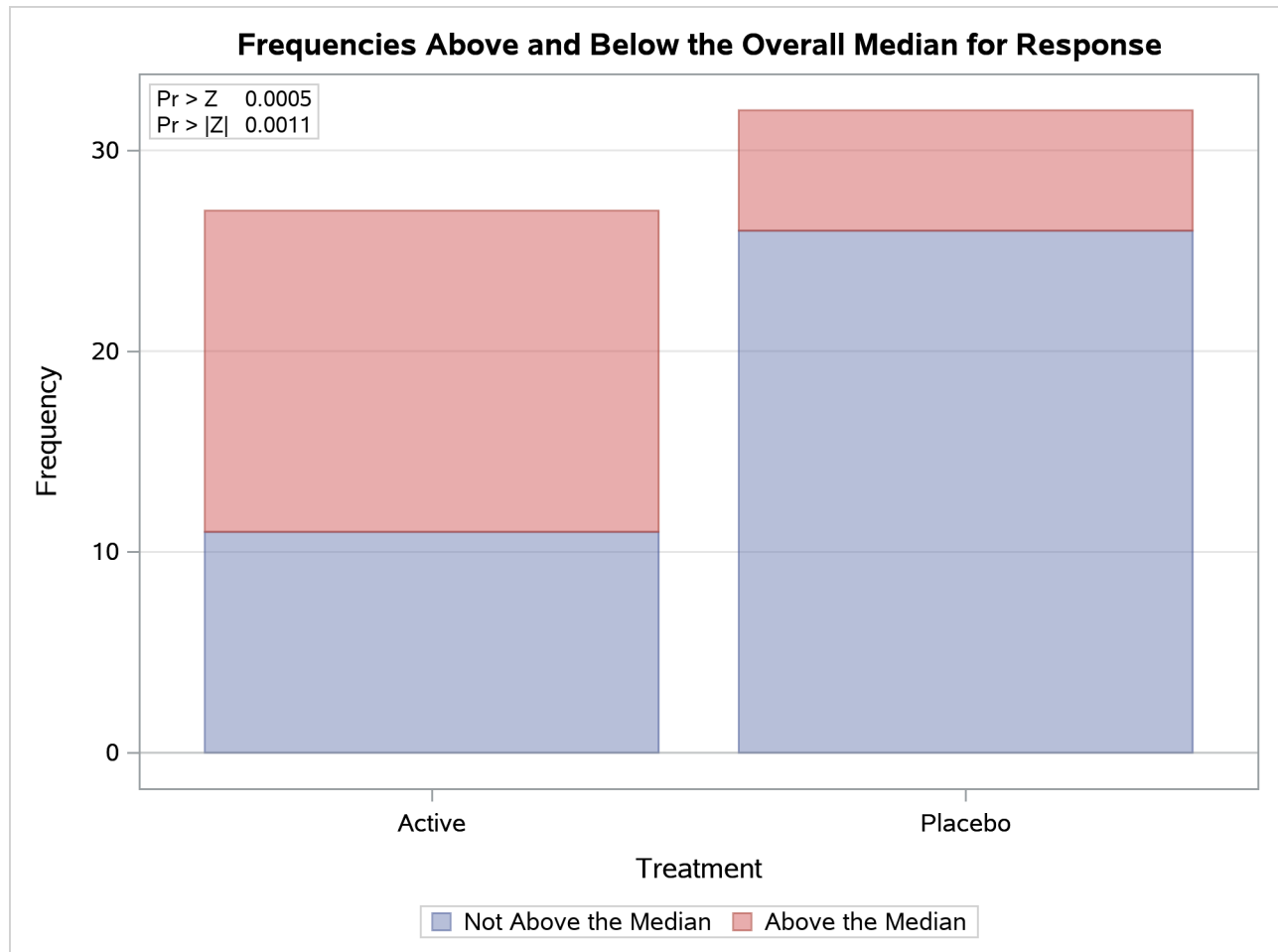
Output 90.1.3 Median Two-Sample Test

Median Scores (Number of Points Above Median) for Variable Response Classified by Variable Treatment					
Treatment	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
Active	27	18.916667	13.271186	1.728195	0.700617
Placebo	32	10.083333	15.728814	1.728195	0.315104
Average scores were used for ties.					

Output 90.1.3 *continued*

Median Two-Sample Test			
Statistic	Z	Pr > Z	Pr > Z
18.9167	3.2667	0.0005	0.0011

Median One-Way Analysis		
Chi-Square	DF	Pr > ChiSq
10.6713	1	0.0011

Output 90.1.4 Median Plot

Example 90.2: EDF Statistics and EDF Plot

This example uses the SAS data set *Arthritis* created in [Example 90.1](#). The data set contains the variable *Treatment*, which denotes the treatment received by a patient, and the variable *Response*, which contains the response status of the patient. The variable *Freq* contains the frequency of the observation, which is the number of patients with the *Treatment* and *Response* combination.

The following statements request empirical distribution function (EDF) statistics, which test whether the distribution of a variable is the same across different groups. The EDF option requests the EDF analysis. The variable *Treatment* is the CLASS variable, and the variable *Response* specified in the VAR statement is the analysis variable. The FREQ statement names *Freq* as the frequency variable.

The PLOTS= option requests an EDF plot for *Response* classified by *Treatment*. ODS Graphics must be enabled before producing plots.

```
ods graphics on;
proc npar1way edf plots=edfplot data=Arthritis;
  class Treatment;
  var Response;
  freq Freq;
run;
ods graphics off;
```

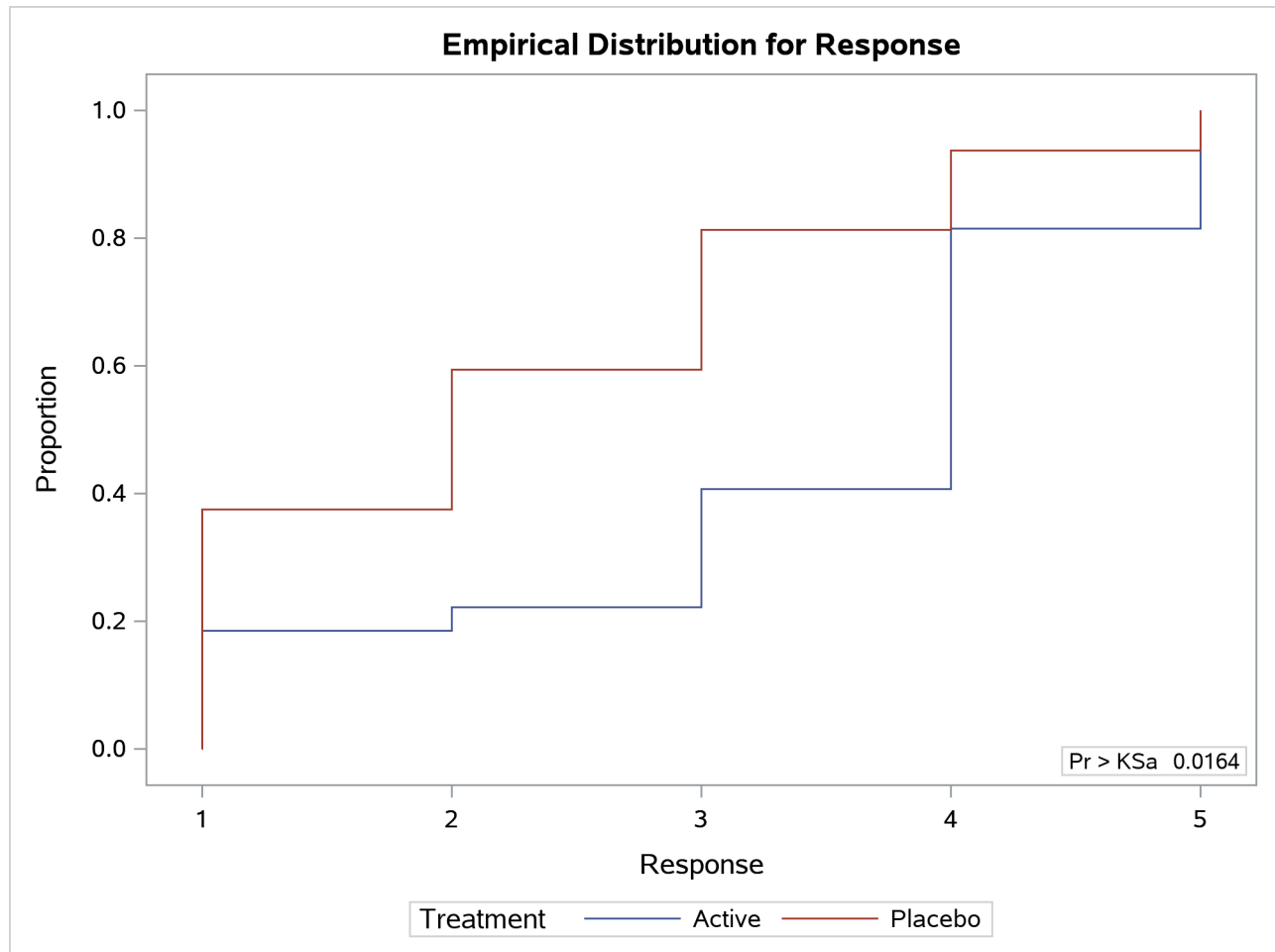
[Output 90.2.1](#) shows EDF statistics that compare the two levels of *Treatment*, *Active* and *Placebo*. The asymptotic *p*-value for the Kolmogorov-Smirnov test is 0.0164. This supports rejection of the null hypothesis that the distributions are the same for the two samples.

[Output 90.2.2](#) shows the EDF plot for *Response* classified by *Treatment*.

Output 90.2.1 Empirical Distribution Function Statistics

The NPAR1WAY Procedure

Kolmogorov-Smirnov Test for Variable Response Classified by Variable Treatment			
Treatment	N	EDF at Maximum	Deviation from Mean at Maximum
Active	27	0.407407	-1.141653
Placebo	32	0.812500	1.048675
Total	59	0.627119	
Maximum Deviation Occurred at Observation 3 Value of Response at Maximum = 3.0			
Kolmogorov-Smirnov Two-Sample Test (Asymptotic)			
KS	0.201818	D	0.405093
KSa	1.550191	Pr > KSa	0.0164

Output 90.2.2 Empirical Distribution Function Plot

Example 90.3: Exact Wilcoxon Two-Sample Test

Researchers conducted an experiment to compare the effects of two stimulants. Thirteen randomly selected subjects received the first stimulant, and six randomly selected subjects received the second stimulant. The reaction times (in minutes) were measured while the subjects were under the influence of the stimulants.

The following SAS statements create the data set `React`, which contains the observed reaction times for each stimulant. The variable `Stim` represents Stimulant 1 or 2. The variable `Time` contains the reaction times observed for subjects under the stimulant.

```
data React;
  input Stim Time @@;
  datalines;
1 1.94 1 1.94 1 2.92 1 2.92 1 2.92 1 2.92 1 3.27
1 3.27 1 3.27 1 3.27 1 3.70 1 3.70 1 3.74
2 3.27 2 3.27 2 3.27 2 3.70 2 3.70 2 3.74
;
```

The following statements request a Wilcoxon test of the null hypothesis that there is no difference between the effects of the two stimulants. Stim is the CLASS variable, and Time is the analysis variable. The WILCOXON option requests an analysis of Wilcoxon scores. The CORRECT=NO option removes the continuity correction from the computation of the standardized z test statistic. The WILCOXON option in the EXACT statement requests exact p -values for the Wilcoxon test. Because the sample size is small, the large-sample normal approximation might not be adequate, and it is appropriate to compute the exact test. These statements produce the results shown in [Output 90.3.1](#).

```
proc npar1way wilcoxon correct=no data=React;
  class Stim;
  var Time;
  exact wilcoxon;
run;
```

[Output 90.3.1](#) displays the results of the Wilcoxon two-sample test. The value of the Wilcoxon statistic is 79.50. Because this value is greater than 60.0 (the expected value under the null hypothesis), PROC NPAR1WAY displays the right-sided p -values. The normal approximation for the Wilcoxon two-sample test yields a one-sided p -value of 0.0382 and a two-sided p -value of 0.0764. For the exact Wilcoxon test, the one-sided p -value is 0.0527, and the two-sided p -value is 0.1054.

Output 90.3.1 Exact Wilcoxon Two-Sample Test

The NPAR1WAY Procedure

Wilcoxon Scores (Rank Sums) for Variable Time Classified by Variable Stim					
Stim	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
1	13	110.50	130.0	11.004784	8.500
2	6	79.50	60.0	11.004784	13.250
Average scores were used for ties.					

Wilcoxon Two-Sample Test							
Statistic (S)	Z	Pr > Z	Pr > Z	t		Pr >= S	Exact Pr >= S-Mean
				Approximation			
79.5000	1.7720	0.0382	0.0764	0.0467	0.0933	0.0527	0.1054

Kruskal-Wallis Test			
Chi-Square	DF	Pr > ChiSq	
3.1398	1	0.0764	

Example 90.4: Hodges-Lehmann Estimation

This example uses the SAS data set `React` created in [Example 90.3](#). The data set contains the variable `Stim`, which represents Stimulant 1 or 2, and the variable `Time`, which contains the reaction times observed for subjects under the stimulant.

The following statements request Hodges-Lehmann estimation of the location shift between the two groups. `Stim` is the CLASS variable, and `Time` is the analysis variable. The `HL` option requests Hodges-Lehmann estimation. The `ALPHA=` option sets the confidence level for the Hodges-Lehmann confidence limits. The `HL` option in the `EXACT` statement requests exact confidence limits for the estimate of location shift. The `ODS SELECT` statement selects which tables to display. [Output 90.4.1](#) shows the Hodges-Lehmann results.

```
proc npar1way hl alpha=.02 data=React;
  class Stim;
  var Time;
  exact hl;
  ods select WilcoxonScores HodgesLehmann;
run;
```

The `HL` option invokes the `WILCOXON` option, which produces a table of Wilcoxon scores ([Output 90.4.1](#)). The Hodges-Lehmann estimate of location shift is 0.35, and the asymptotic confidence limits are 0.00 and 0.82. The confidence interval midpoint is 0.41, which can also be used as an estimate of the location shift. The ASE estimate of 0.1762 is based on the length of the confidence interval. The exact confidence limits are 0.00 and 1.33.

Output 90.4.1 Hodges-Lehmann Estimate of Location Shift

The NPAR1WAY Procedure

Wilcoxon Scores (Rank Sums) for Variable Time Classified by Variable Stim					
Stim	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
1	13	110.50	130.0	11.004784	8.500
2	6	79.50	60.0	11.004784	13.250
Average scores were used for ties.					
Hodges-Lehmann Estimation					
Location Shift (2 - 1) 0.3500					
Type	98% Confidence Limits		Interval Midpoint	Asymptotic Standard Error	
Asymptotic (Moses)	0.0000	0.8200	0.4100	0.1762	
Exact	0.0000	1.3300	0.6650		

Example 90.5: Exact Savage Multisample Test

A researcher conducting a laboratory experiment randomly assigned 15 mice to receive one of three drugs. The survival time (in days) was then recorded.

The following SAS statements create the data set `Mice`, which contains the observed survival times for the mice. The variable `Treatment` denotes the treatment received. The variable `Days` contains the number of days the mouse survived.

```
data Mice;
  input Treatment $ Days @@;
  datalines;
1 1 1 1 1 3 1 3 1 4
2 3 2 4 2 4 2 4 2 15
3 4 3 4 3 10 3 10 3 26
;
```

The following statements request a Savage test of the null hypothesis that there is no difference in survival time among the three drugs. `Treatment` is the CLASS variable, and `Days` is the analysis variable. The SAVAGE option requests an analysis of Savage scores. The SAVAGE option in the EXACT statement requests exact p -values for the Savage test. Because the sample size is small, the large-sample normal approximation might not be adequate, and it is appropriate to compute the exact test.

PROC NPAR1WAY tests the null hypothesis that there is no difference in the survival times among the three drugs against the alternative hypothesis of difference among the drugs. The SAVAGE option specifies an analysis based on Savage scores. The variable `Treatment` is the CLASS variable, and the variable `Days` is the response variable. The EXACT statement requests the exact Savage test.

```
proc npar1way savage data=Mice;
  class Treatment;
  var Days;
  exact savage;
run;
```

Output 90.5.1 shows the results of the Savage test. The exact p -value is 0.0445, which supports a difference in survival times among the drugs at the 0.05 level. The asymptotic p -value based on the chi-square approximation is 0.0638.

Output 90.5.1 Exact Savage Multisample Test
The NPAR1WAY Procedure

Savage Scores (Exponential) for Variable Days Classified by Variable Treatment					
Treatment	N	Sum of Scores	Expected Under H0	Std Dev Under H0	Mean Score
1	5	-3.367980	0.0	1.634555	-0.673596
2	5	0.095618	0.0	1.634555	0.019124
3	5	3.272362	0.0	1.634555	0.654472
Average scores were used for ties.					

Output 90.5.1 continued

Savage One-Way Analysis				
Chi-Square	DF	Pr > ChiSq	Exact Pr >= ChiSq	
5.5047	2	0.0638	0.0445	

References

- Agresti, A. (1992). "A Survey of Exact Inference for Contingency Tables." *Statistical Science* 7:131–177.
- Agresti, A. (2002). *Categorical Data Analysis*. 2nd ed. New York: John Wiley & Sons.
- Agresti, A. (2007). *An Introduction to Categorical Data Analysis*. 2nd ed. New York: John Wiley & Sons.
- Agresti, A. (2013). *Categorical Data Analysis*. 3rd ed. Hoboken, NJ: John Wiley & Sons.
- Agresti, A., Mehta, C. R., and Patel, N. R. (1990). "Exact Inference for Contingency Tables with Ordered Categories." *Journal of the American Statistical Association* 85:453–458.
- Agresti, A., Wackerly, D., and Boyett, J. M. (1979). "Exact Conditional Tests for Cross-Classifications: Approximation of Attained Significance Levels." *Psychometrika* 44:75–83.
- Bishop, Y. M. M., Fienberg, S. E., and Holland, P. W. (1975). *Discrete Multivariate Analysis: Theory and Practice*. Cambridge, MA: MIT Press.
- Butar, F. B., and Park, J.-W. (2008). "Permutation Tests for Comparing Two Populations." *Journal of Mathematical Sciences and Mathematics Education* 3:19–30.
- Conover, W. J. (1999). *Practical Nonparametric Statistics*. 3rd ed. New York: John Wiley & Sons.
- Critchlow, D. E., and Fligner, M. A. (1991). "On Distribution-Free Multiple Comparisons in the One-Way Analysis of Variance." *Communications in Statistics—Theory and Methods* 20:127–139.
- Dwass, M. (1960). "Some k-Sample Rank-Order Tests." In *Contributions to Probability and Statistics: Essays in Honor of Harold Hotelling*, edited by I. Olkin, S. G. Ghurye, W. Hoeffding, W. G. Madow, and H. B. Mann, 198–202. Stanford, CA: Stanford University Press.
- Fligner, M. A., and Policello, G. E. (1981). "Robust Rank Procedures for the Behrens-Fisher Problem." *Journal of the American Statistical Association* 76:162–168.
- Gail, M. H., and Mantel, N. (1977). "Counting the Number of $r \times c$ Contingency Tables with Fixed Margins." *Journal of the American Statistical Association* 72:859–862.
- Gibbons, J. D., and Chakraborti, S. (2010). *Nonparametric Statistical Inference*. 5th ed. New York: Chapman & Hall.
- Hajek, J. (1969). *A Course in Nonparametric Statistics*. San Francisco: Holden-Day.
- Halverson, J. O., and Sherwood, F. W. (1930). "Investigations in the Feeding of Cottonseed Meal to Cattle." *North Carolina Agricultural Experiment Station Technical Bulletin* 39, 158 pages.

- Higgins, J. J. (2004). *An Introduction to Modern Nonparametric Statistics*. Belmont, CA: Brooks/Cole.
- Hirji, K. F. (2006). *Exact Analysis of Discrete Data*. Boca Raton, FL: Chapman & Hall/CRC.
- Hodges, J. L., Jr. (1957). "The Significance Probability of the Smirnov Two-Sample Test." *Arkiv för Matematik* 3:469–486.
- Hodges, J. L., Jr., and Lehmann, E. L. (1962). "Rank Methods for Combination of Independent Experiments in Analysis of Variance." *Annals of Mathematical Statistics* 33:482–497.
- Hodges, J. L., Jr., and Lehmann, E. L. (1983). "Hodges-Lehmann Estimators." In *Encyclopedia of Statistical Sciences*, vol. 3, edited by S. Kotz, N. L. Johnson, and C. B. Read. New York: John Wiley & Sons.
- Hollander, M., and Wolfe, D. A. (1999). *Nonparametric Statistical Methods*. 2nd ed. New York: John Wiley & Sons.
- Juneau, P. (2004). "Simultaneous Nonparametric Inference in a One-Way Layout Using the SAS System." In *Proceedings of PharmaSUG 2004 (Pharmaceutical Industry SAS Users Group)*. Paper no. 04. Cary, NC: SAS Institute Inc.
- Juneau, P. (2007). "Nonparametric Methods in Pharmaceutical Statistics." In *Pharmaceutical Statistics Using SAS: A Practical Guide*, edited by A. Dmitrienko, C. Chuang-Stein, and R. D'Agostino, 117–150. Cary, NC: SAS Institute Inc.
- Kiefer, J. (1959). "*K*-Sample Analogues of the Kolmogorov-Smirnov and Cramér–von Mises Tests." *Annals of Mathematical Statistics* 30:420–447.
- Lehmann, E. L. (1963). "Nonparametric Confidence Intervals for a Shift Parameter." *Annals of Mathematical Statistics* 34:1507–1512.
- Lehmann, E. L., and D'Abrera, H. J. M. (2006). *Nonparametrics: Statistical Methods Based on Ranks*. Rev. ed. New York: Springer.
- Mehrotra, D. V., Lu, X., and Li, X. (2010). "Rank-Based Analyses of Stratified Experiments: Alternatives to the van Elteren Test." *American Statistician* 64:121–130.
- Mehta, C. R., and Patel, N. R. (1983). "A Network Algorithm for Performing Fisher's Exact Test in $r \times c$ Contingency Tables." *Journal of the American Statistical Association* 78:427–434.
- Mehta, C. R., Patel, N. R., and Senchaudhuri, P. (1991). "Exact Stratified Linear Rank Tests for Binary Data." In *Computing Science and Statistics: Proceedings of the Twenty-Third Symposium on the Interface*, edited by E. M. Keramidas, 200–207. Fairfax Station, VA: Interface Foundation.
- Mehta, C. R., Patel, N. R., and Tsiatis, A. A. (1984). "Exact Significance Testing to Establish Treatment Equivalence with Ordered Categorical Data." *Biometrics* 40:819–825.
- Orban, J., and Wolfe, D. A. (1979). *A Class of Distribution-Free Two-Sample Tests Based on Placements*. Technical Report 178, Department of Statistics, Ohio State University.
- Owen, D. B. (1962). *Handbook of Statistical Tables*. Reading, MA: Addison-Wesley.
- Quade, D. (1966). "On Analysis of Variance for the *k*-Sample Problem." *Annals of Mathematical Statistics* 37:1747–1758.

- Randles, R. H., and Wolfe, D. A. (1979). *Introduction to the Theory of Nonparametric Statistics*. New York: John Wiley & Sons.
- Richter, S. J., and McCann, M. H. (2017). "Permutation Tests of Scale Using Deviances." *Communications in Statistics—Simulation and Computation* 46:5553–5565.
- Sheskin, D. J. (1997). *Handbook of Parametric and Nonparametric Statistical Procedures*. Boca Raton, FL: CRC Press.
- Steel, R. G. D. (1960). "A Rank Sum Test for Comparing All Pairs of Treatments." *Technometrics* 2:197–207.
- Valz, P. D., and Thompson, M. E. (1994). "Exact Inference for Kendall's S and Spearman's ρ with Extensions to Fisher's Exact Test in $r \times c$ Contingency Tables." *Journal of Computational and Graphical Statistics* 3:459–472.
- Van Elteren, P. H. (1960). "On the Combination of Independent Two-Sample Tests of Wilcoxon." *Bulletin of the International Statistical Institute* 37:351–361.