

A Comprehensive Guide to Model Selection Using PROC GLMSELECT in SAS

Authored by
Mohammed looti

November 14, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *A Comprehensive Guide to Model Selection Using PROC GLMSELECT in SAS*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=1309>

In the realm of statistical modeling, identifying the most effective set of [predictor variables](#) for a [regression model](#) is a fundamental challenge. The **PROC GLMSELECT** statement in [SAS](#) provides a powerful and efficient mechanism for automated [model selection](#), helping researchers and analysts to navigate complex datasets and arrive at parsimonious, yet robust, models.

This procedure is particularly valuable when faced with a large number of potential independent variables, where manual selection could be time-consuming and prone to subjective bias. By systematically evaluating various subsets of predictors, **PROC GLMSELECT** aims to find the optimal balance between model fit and complexity, thereby enhancing the interpretability and predictive accuracy of the resulting model.

Throughout this guide, we will delve into the practical application of **PROC GLMSELECT**, illustrating its capabilities through a concrete example. We will cover the steps from data preparation to interpreting the comprehensive output, ensuring a clear understanding of how to leverage this powerful tool for effective model building in SAS.

Setting Up Your Data for Analysis in SAS

To demonstrate the utility of **PROC GLMSELECT**, let's consider a scenario where we aim to predict student final exam scores. Our objective is to construct a [multiple linear regression](#) model using several hypothesized predictor variables: the number of hours students spent studying, the number of preparatory exams they took, and their gender.

Our first step involves creating a dataset in SAS that contains this information for a sample of 20 students. This dataset, which we will name `exam_data`, will serve as the foundation for our model selection process. The data includes both continuous variables (hours, prep_exams, score) and a [categorical variable](#) (gender), which requires special handling in SAS procedures.

The following SAS code snippet illustrates how to construct this dataset using the `DATALINES` statement. This method is ideal for entering small datasets directly within your SAS program, making the example self-contained and reproducible for educational purposes.

```
/*create dataset*/  
data exam_data;  
input hours prep_exams gender $ score;  
datalines;  
1 1 0 76  
2 3 1 78  
2 3 0 85  
4 5 0 88  
2 2 0 72
```

```
1 2 1 69
5 1 1 94
4 1 0 94
2 0 1 88
4 3 0 92
4 4 1 90
3 3 1 75
6 2 1 96
5 4 0 90
3 4 0 82
4 4 1 85
6 5 1 99
2 1 0 83
1 0 1 62
2 1 0 76
;
run;
```

```
/*view dataset*/
```

```
proc print data=exam_data;
```

Obs	hours	prep_exams	gender	score
1	1	1	0	76
2	2	3	1	78
3	2	3	0	85
4	4	5	0	88
5	2	2	0	72
6	1	2	1	69
7	5	1	1	94
8	4	1	0	94
9	2	0	1	88
10	4	3	0	92
11	4	4	1	90
12	3	3	1	75
13	6	2	1	96
14	5	4	0	90
15	3	4	0	82
16	4	4	1	85
17	6	5	1	99
18	2	1	0	83
19	1	0	1	62
20	2	1	0	76

Upon execution of the data step and PROC PRINT, SAS will display the newly created exam_data table, allowing for a quick visual inspection to ensure the data has been loaded correctly. This preliminary check is crucial before proceeding with any statistical analysis, as it helps confirm data integrity.

Implementing PROC GLMSELECT for Automated Selection

With our dataset prepared, the next critical step is to invoke the [PROC GLMSELECT](#) procedure to initiate the automated [model selection](#) process. This procedure will systematically evaluate various combinations of our specified predictor variables to identify the subset that yields the most statistically sound [regression model](#) based on predefined criteria.

The core syntax for **PROC GLMSELECT** involves specifying the input dataset, declaring any [categorical variables](#) using the CLASS statement, and listing all potential predictor variables in the

`MODEL` statement. The ``CLASS`` statement is particularly important for variables like 'gender', ensuring SAS correctly treats it as a nominal variable rather than a continuous one.

The following code snippet demonstrates how to apply **PROC GLMSELECT** to our `exam_data`. Here, `score` is our dependent variable, while `hours`, `prep_exams`, and `gender` are the candidate independent variables from which the procedure will select the best combination.

```
/*perform model selection*/  
proc glmselect data=exam_data;  
class gender;  
model score = hours prep_exams gender;  
run;
```

It is vital to include `gender` in the `CLASS` statement because it is a binary [categorical variable](#) (represented by 0 and 1 in our dataset). Failing to declare it in the `CLASS` statement would lead SAS to treat it as a continuous numeric variable, which would yield incorrect statistical inferences and potentially a misleading model.

Interpreting the GLMSELECT Output: Selection Process Overview

Upon successful execution of the **PROC GLMSELECT** statement, SAS generates a series of output tables. The initial tables provide a high-level overview of the model selection process, detailing the criteria used and the general settings of the procedure. Understanding this initial summary is crucial for contextualizing the subsequent detailed results.

A key piece of information presented here is the criterion chosen by the procedure to guide the variable selection. In our example, the output clearly indicates that the criterion used to stop adding or removing variables from the model was **SBC**. This acronym stands for [Schwarz Information Criterion](#), which is also widely known as the [Bayesian Information Criterion](#).

The **SBC** is a widely recognized metric for [model selection](#) that balances model fit with model complexity. Its primary goal is to select the model that best explains the data with the fewest possible predictor variables, thereby guarding against overfitting. A model with a lower **SBC** value is generally preferred, as it signifies a better trade-off between explanatory power and parsimony.

Essentially, the **PROC GLMSELECT** statement iteratively evaluates different models, adding or removing variables one at a time, until it identifies the model that achieves the minimum **SBC** value. This model is then considered the "best" model according to this criterion, as it represents the most efficient explanation of the dependent variable.

The GLMSELECT Procedure

Data Set	WORK.EXAM_DATA
Dependent Variable	score
Selection Method	Stepwise
Select Criterion	SBC
Stop Criterion	SBC
Effect Hierarchy Enforced	None

Number of Observations Read	20
Number of Observations Used	20

Class Level Information		
Class	Levels	Values
gender	2	0 1

Dimensions	
Number of Effects	4
Number of Parameters	5

Understanding the Stepwise Selection Results

Following the initial overview, **PROC GLMSELECT** provides a detailed breakdown of the stepwise selection process. This section of the output is particularly insightful as it chronicles each step taken by the algorithm, illustrating how variables were considered for inclusion or exclusion and how the selection criterion (**SBC** in our case) changed at each stage.

The table typically displays the model effects currently in the model, the effect considered for entry or removal, and the corresponding **SBC** value for the model at that step. This allows analysts to trace the path of the selection algorithm and understand why certain variables were ultimately included or excluded.

In our specific example, the output reveals that the process began with a model containing only an [intercept term](#), which yielded an **SBC** value of **93.4337**. The procedure then evaluated potential improvements by considering the addition of other variables.

Subsequently, the addition of 'gender' as a predictor variable was considered. However, this action resulted in an increase in the **SBC** value to **71.7383**, which is counterintuitive as a lower **SBC** is

preferred. This indicates that including 'gender' did not improve the model sufficiently to justify the added complexity, according to the **SBC** criterion.

Therefore, based on the principle of minimizing **SBC**, the selection process concluded that the most parsimonious and effective model for predicting exam scores includes only the [intercept term](#) and the variable 'hours studied'. This outcome highlights the ability of **PROC GLMSELECT** to identify the simplest model that still provides a good fit to the data.

The GLMSELECT Procedure

Stepwise Selection Summary

Step	Effect Entered	Effect Removed	Number Effects In	Number Parm's In	SBC
0	Intercept		1	1	93.4337
1	hours		2	2	70.4452*
* Optimal Value of Criterion					

Selection stopped at a local minimum of the SBC criterion.

Stop Details

Candidate For	Effect	Candidate SBC		Compare SBC
Entry	gender	71.7383	>	70.4452
Removal	hours	93.4337	>	70.4452

Deriving and Evaluating the Final Regression Model

Once the optimal set of predictor variables has been identified by **PROC GLMSELECT**, the procedure provides the detailed results for the final chosen [regression model](#). This final output section is crucial for understanding the model's structure, the significance of its components, and its overall performance.

The most important part of this section for constructing the model equation is the [Parameter Estimates](#) table. This table lists the estimated coefficients for each predictor variable included in the final model, along with the [intercept term](#). These coefficients are the foundational elements for writing out the predictive equation.

From the values presented in the **Parameter Estimates** table, we can directly formulate our fitted [regression model](#). Based on our analysis, the equation for predicting a student's final exam score is:

$$\text{Exam Score} = 67.161689 + 5.250257(\text{hours studied})$$

This equation indicates that for every additional hour a student studies, their predicted final exam score increases by approximately 5.25 points, assuming all other factors remain constant. The [intercept term](#) of 67.161689 represents the predicted exam score for a student who studies zero hours.

The GLMSELECT Procedure Selected Model

The selected model is the model at the last step (Step 1).

Effects: Intercept hours

Analysis of Variance				
Source	DF	Sum of Squares	Mean Square	F Value
Model	1	1338.29063	1338.29063	48.00
Error	18	501.90937	27.88385	
Corrected Total	19	1840.20000		

Root MSE	5.28052
Dependent Mean	83.70000
R-Square	0.7273
Adj R-Sq	0.7121
AIC	90.45375
AICC	91.95375
SBC	70.44521

Parameter Estimates				
Parameter	DF	Estimate	Standard Error	t Value
Intercept	1	67.161689	2.663270	25.22
hours	1	5.250257	0.757847	6.93

Beyond the equation, the output also provides several key metrics that assess how well the chosen model fits the observed data. These statistics are essential for evaluating the model's overall reliability and its practical utility.

One such critical metric is the [R-Square](#) value, which quantifies the proportion of the variance in the dependent variable (exam scores) that can be explained by the predictor variables included in

the model. A higher [R-Square](#) indicates a better fit.

In our model, an **R-Square** value of **72.73%** implies that approximately 72.73% of the variation in student exam scores can be accounted for by the number of hours studied. This suggests a reasonably strong explanatory power of the model.

Another important diagnostic is the [Root Mean Square Error \(Root MSE\)](#). This statistic represents the standard deviation of the residuals (the differences between observed and predicted values) and provides a measure of the average distance that the observed values fall from the [regression line](#). A lower **Root MSE** indicates greater precision in the model's predictions.

For this [regression model](#), the observed exam scores deviate by an average of **5.28052** units from the predicted scores. This value helps in understanding the typical magnitude of prediction errors made by our model.

Concluding Thoughts and Further Exploration

The [PROC GLMSELECT](#) statement in [SAS](#) stands as an invaluable tool for automated [model selection](#), particularly beneficial for complex datasets with numerous potential predictor variables. It streamlines the process of identifying a parsimonious and effective [regression model](#), thereby enhancing the reliability and interpretability of statistical findings.

While our example utilized the default stepwise selection method with the **SBC** criterion, **PROC GLMSELECT** offers a wide array of options for different selection methods (e.g., forward, backward, LASSO, LAR) and alternative criteria (e.g., AIC, AICC, ADJRSQ). Exploring these options can provide further flexibility to tailor the model selection process to specific research questions and data characteristics.

For a comprehensive understanding of all available arguments and advanced functionalities, it is highly recommended to refer to the [official SAS documentation for PROC GLMSELECT](#). This resource provides in-depth explanations and examples to help users maximize the potential of this powerful procedure.

Additional Resources

For those interested in expanding their SAS proficiency, the following tutorials offer guidance on performing other common statistical tasks and data manipulations within the SAS environment: