

Learning the F1 Score: Calculation and Implementation in R

Authored by
Mohammed loot

November 2, 2025

RECOMMENDED CITATION

Mohammed loot (2025). *Learning the F1 Score: Calculation and Implementation in R*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=8684>

The Crucial Role of F1 Score in Model Evaluation

The field of [machine learning](#) relies fundamentally on robust evaluation metrics to assess the true efficacy of predictive models. While simple accuracy is often the starting point, it frequently masks critical deficiencies, particularly when dealing with datasets exhibiting significant class imbalance. In such challenging classification environments, the **F1 Score** emerges as a superior and highly utilized metric.

The F1 Score is designed to provide a single, representative measure that harmonizes two essential, yet often conflicting, dimensions of performance: [Precision](#) and [Recall](#). By calculating the harmonic mean of these two metrics, the F1 Score ensures that a model cannot achieve a high score by maximizing one metric at the expense of the other. This balanced approach is vital in scenarios--such as fraud detection or medical diagnosis--where both minimizing false positives and minimizing false negatives carry substantial, distinct costs.

For any data science practitioner aiming to build and deploy reliable binary or multi-class classifiers, mastering the calculation and interpretation of the F1 Score is non-negotiable. This comprehensive guide details the mathematical foundation of this metric and provides precise instructions for its implementation within the powerful statistical programming environment, **R**.

Understanding the Foundations: Precision, Recall, and the Confusion Matrix

To fully appreciate the utility of the **F1 Score**, one must first grasp its constituent components, which are derived directly from the outcomes summarized in the [Confusion Matrix](#). This matrix is a pivotal tool that visually maps the relationship between the actual classes (ground truth) and the classes predicted by the model, classifying results into four categories: True Positives (TP), True Negatives (TN), False Positives (FP), and False Negatives (FN).

The mathematical definition of the F1 Score is rooted in the harmonic mean, expressed as follows:

$$\text{F1 Score} = 2 * (\text{Precision} * \text{Recall}) / (\text{Precision} + \text{Recall})$$

The contributing metrics are defined by the counts within the [Confusion Matrix](#):

Precision: This metric quantifies the quality or accuracy of the positive predictions made by the model. It is calculated as the ratio of correctly identified positive cases (True Positives) to the total number of cases the model labeled as positive (True Positives + False Positives). A high precision score signifies that when the model predicts a positive outcome, it is highly likely to be correct, thus minimizing "false alarms."

Recall (or Sensitivity): This metric, conversely, measures the model's ability to find and correctly identify all actual positive cases present in the dataset. It is calculated as the ratio of correctly

identified positive cases (True Positives) to the total number of actual positive cases (True Positives + False Negatives). High recall indicates that the model is highly effective at capturing positive instances, minimizing the number of missed opportunities.

The brilliance of the F1 Score lies in its use of the harmonic mean. Unlike the arithmetic mean, the harmonic mean heavily penalizes scenarios where there is a large disparity between Precision and Recall. If a model achieves near-perfect recall but terrible precision (or vice-versa), the [F1 Score](#) will remain low, ensuring that only truly balanced models are rated highly. The ideal score is 1.0, achieved only when both underlying metrics are perfect.

Practical Example: Step-by-Step Manual F1 Score Calculation

To solidify the theoretical understanding, let us walk through a practical example of calculating the **F1 Score** using a classification scenario. Imagine we have developed a logistic regression model tasked with predicting the outcome (drafted, labeled '1') for 400 college basketball players. Our goal is to determine the model's performance in a real-world context.

Upon evaluating the model against the test data, the results are summarized in the following confusion matrix, which clearly dictates the counts for the four possible outcomes:

		Predicted	
		Drafted = Yes	Drafted = No
Actual	Drafted = Yes	120 (True Positive)	40 (False Negative)
	Drafted = No	70 (False positive)	170 (True Negative)

From this matrix, we identify the counts: True Positives (TP) = 120, False Positives (FP) = 70, True Negatives (TN) = 170, and False Negatives (FN) = 40. We can now proceed with the necessary calculations to derive the [F1 Score](#) manually:

Calculate Precision: This step addresses the question: Of all the players predicted to be drafted, how many actually were?

Precision = True Positive / (True Positive + False Positive) = $120 / (120 + 70) = 120 / 190 = 0.63157$

Calculate Recall: This step addresses the question: Of all the players who were actually drafted, how many did the model correctly identify?

Recall = True Positive / (True Positive + False Negative) = 120 / (120 + 40) = 120 / 160 = **0.75**

Calculate F1 Score: Finally, we combine Precision and Recall using the harmonic mean formula to find the overall balanced score.

F1 Score = 2 * (Precision * Recall) / (Precision + Recall)

F1 Score = 2 * (0.63157 * 0.75) / (0.63157 + 0.75) = 2 * (0.4736775) / (1.38157) = 0.947355 / 1.38157 = **0.6857**

The resulting **F1 Score** of 0.6857 confirms the model's effectiveness in achieving a reasonable balance. Although the model excels slightly more at identifying actual positive cases (Recall of 0.75), its performance is slightly dragged down by a higher rate of false alarms (Precision of 0.63157), demonstrating the penalizing nature of the harmonic mean calculation.

Automating Metric Calculation with the R caret Package

While manual calculations are invaluable for pedagogical purposes, real-world data science projects demand speed, accuracy, and efficiency. Within the [R](#) ecosystem, the [caret](#) (Classification and Regression Training) package is the standard toolkit for streamline model evaluation and metric generation. The `confusionMatrix()` function, in particular, provides a comprehensive, single-function solution for calculating the **F1 Score** and all related performance indicators directly from prediction vectors.

The following R code snippet demonstrates how to replicate the basketball drafting example within the R environment. We first define the actual (ground truth) and predicted vectors corresponding to the counts in our manual calculation example. We then call the core function, specifying the positive class (in this case, '1') and the mode to ensure all available metrics are reported:

```
library(caret)
```

```
#define vectors of actual values and predicted values
```

```
actual <- factor(rep(c(1, 0), times=c(160, 240)))
```

```
pred <- factor(rep(c(1, 0, 1, 0), times=c(120, 40, 70, 170)))
```

```
#create confusion matrix and calculate metrics related to confusion matrix
```

```
confusionMatrix(pred, actual, mode = "everything", positive="1")
```

```
Reference
```

```
Prediction 0 1
```

```
0 170 40
```

```
1 70 120
```

Accuracy : 0.725

95% CI : (0.6784, 0.7682)

No Information Rate : 0.6

P-Value : 1.176e-07

Kappa : 0.4444

Mcnemar's Test P-Value : 0.005692

Sensitivity : 0.7500

Specificity : 0.7083

Pos Pred Value : 0.6316

Neg Pred Value : 0.8095

Precision : 0.6316

Recall : 0.7500

F1 : 0.6857

Prevalence : 0.4000

Detection Rate : 0.3000

Detection Prevalence : 0.4750

Balanced Accuracy : 0.7292

'Positive' Class : 1

As the output confirms, the `confusionMatrix()` function efficiently calculates all derived statistics automatically, providing the full **Confusion Matrix** structure alongside metrics like Sensitivity (Recall), Specificity, and Precision. Critically, we see that the resulting **F1 Score** value is precisely **0.6857**, validating our precise manual steps and confirming the utility of the [caret](#) package for robust metric reporting in R.

Interpreting and Applying the F1 Score for Model Selection

The primary value of the **F1 Score** is its capacity to serve as an objective benchmark during the crucial model selection phase. When multiple candidate models (e.g., Logistic Regression, Support Vector Machines, Random Forests) are built to tackle the same classification problem, comparing their F1 Scores allows practitioners to quickly determine which model provides the most balanced and overall effective performance.

A higher F1 Score consistently indicates superior performance relative to the combined requirements of Precision and [Recall](#). For example, if a second, refined model applied to the basketball data were to achieve an F1 Score of 0.85, that model would be deemed unequivocally superior to our initial model (F1 = 0.6857). This high F1 value implies that the improved model

manages to increase both the accuracy of its positive predictions and its ability to capture nearly all true positive instances.

However, it is essential to contextualize the metric. The **F1 Score** is the ideal choice specifically when the costs associated with False Positives and False Negatives are deemed approximately equal--that is, when Precision and Recall must be weighted equally. If, in a certain application (like identifying spam), we are willing to tolerate more missed spam (lower Recall) to ensure absolutely no legitimate emails are flagged (high Precision), then prioritizing the Precision metric alone, or using a specialized metric like the F-beta score, might be more appropriate. The F1 Score, by definition, mandates equilibrium between the two performance factors.

Conclusion and Next Steps in R Modeling

The **F1 Score** stands as a cornerstone metric in classification tasks, offering a refined measure of model quality that moves beyond the simplicity of standard accuracy, particularly when class imbalances exist. By functioning as the harmonic mean of Precision and Recall, it drives model developers toward solutions that are both accurate in their positive assertions and comprehensive in their coverage of actual positive cases.

Implementing this and other sophisticated metrics in **R** is straightforward, thanks to powerful libraries such as [caret](#). Understanding how the underlying mathematics translates into functional code allows data scientists to generate reliable, reproducible performance reports necessary for making informed decisions regarding model deployment and optimization.

To further enhance your expertise in model evaluation and statistical programming, we encourage exploring the extensive documentation provided by the [R](#) community and specializing in advanced techniques available within the [machine learning](#) ecosystem.