

Learning Curve Fitting Techniques with Python: A Practical Guide

Authored by
Mohammed looti

November 4, 2025

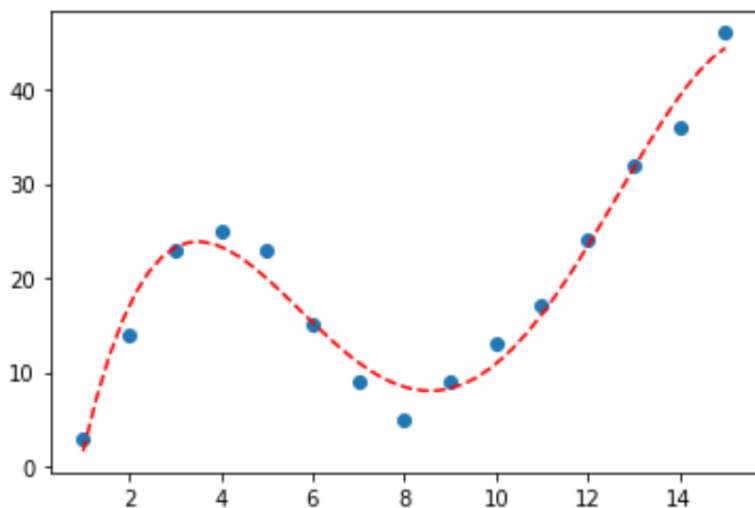
RECOMMENDED CITATION

Mohammed looti (2025). *Learning Curve Fitting Techniques with Python: A Practical Guide*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=10177>

In the realm of data science, predictive modeling, and advanced statistical analysis, the ability to accurately represent the relationship between variables is fundamentally important. Often, real-world data does not conform to simple straight lines; instead, datasets frequently exhibit complex, non-linear patterns. This necessity drives the application of [Curve Fitting](#)--a powerful technique used to select the optimal mathematical function that best captures the underlying trend between an independent variable (x) and a dependent variable (y).

The objective of curve fitting transcends merely drawing a line through scattered data points. It is about constructing a mathematically sound and **robust model** capable of performing reliable interpolation (predicting within the known range) and extrapolation (predicting outside the known range). This task is made highly efficient by leveraging the sophisticated data manipulation and processing capabilities inherent in the [Python](#) ecosystem, specifically utilizing core libraries such as [NumPy](#), [Pandas](#), and [Matplotlib](#).

This comprehensive, step-by-step tutorial provides a detailed walkthrough of fitting multiple polynomial curves to a sample dataset using Python. Crucially, we will move beyond visual inspection by employing rigorous statistical criteria, focusing specifically on the **Adjusted R-squared value**, to objectively identify the curve that offers the most accurate, yet parsimonious, representation of your data structure.



Understanding the Importance of Curve Fitting in Data Analysis

Curve fitting is not just a computational exercise; it is a fundamental pillar of applied science, finding critical uses across diverse fields including engineering, financial economics, clinical biology, and physics. This technique enables practitioners to distill complex, noisy empirical data into a simple, continuous mathematical function. When dealing with raw data points that are inherently scattered or contain measurement noise, fitting a curve provides an indispensable

smoothing approximation that clearly reveals the systemic behavior or underlying mechanisms being investigated.

The successful execution of curve fitting hinges entirely upon the correct selection of the curve type. While data exhibiting a simple, proportional relationship may be adequately handled by standard linear regression, the majority of real-world phenomena follow more complex trajectories. These may include patterns of exponential growth, asymptotic saturation, or cyclical oscillation, all of which necessitate the use of higher-order functions, such as [Polynomial Regression](#). By systematically testing and evaluating polynomials of varying degrees, we can effectively manage the critical trade-off between the simplicity of the model and its goodness of fit to the observed data.

A poor choice of curve type can lead to significant modeling errors. If the chosen model is too basic, it results in [underfitting](#)--the inability of the model to capture the essential patterns in the data. Conversely, if the model is excessively complex, it leads to [overfitting](#)--where the model captures random noise and idiosyncrasies rather than the true signal. Our robust methodology, which emphasizes a structured comparative evaluation using statistical metrics, is designed precisely to mitigate these risks, ensuring that the selected model possesses strong generalization capabilities when applied to new, previously unseen data points.

Step 1: Data Preparation and Initial Visualization in Python

The foundation of any successful statistical modeling endeavor is meticulous data preparation and a thorough visual assessment of the data's distribution. Visualization is perhaps the most critical initial step, as it provides immediate, intuitive clues regarding the relationship between the variables, guiding the subsequent selection of appropriate mathematical models.

To begin, we generate a synthetic dataset, ideal for demonstrating the technique, utilizing the [Pandas](#) library, which is expertly designed for handling and structuring tabular data. Following the dataset creation, we immediately leverage [Matplotlib](#) to produce a scatterplot. This visual representation allows us to immediately observe the initial pattern--in this particular case, a clear and pronounced non-linear trend--before we proceed to the more complex computational analysis required for curve fitting.

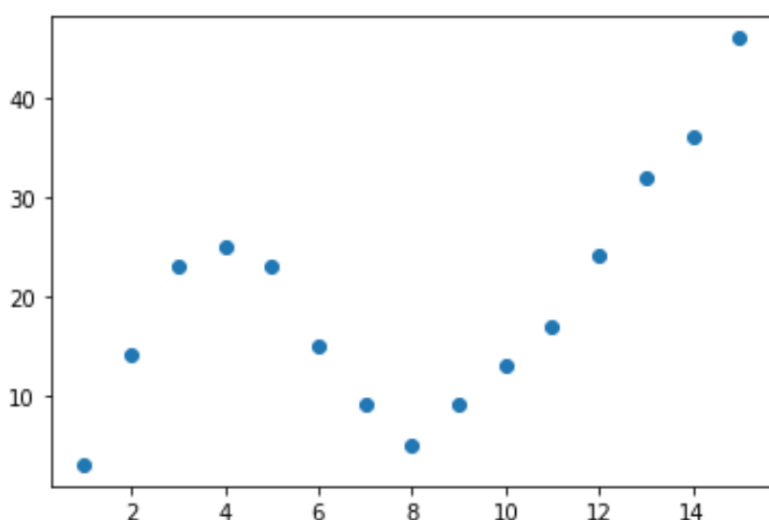
The following code snippet details the construction of our sample DataFrame and the subsequent generation of the preliminary scatterplot, serving as the baseline for all future modeling efforts:

```
import pandas as pd
import matplotlib.pyplot as plt
```

```
#create DataFrame
```

```
df = pd.DataFrame({'x': ,  
'y': })  
  
#create scatterplot of x vs. y  
plt.scatter(df.x, df.y)
```

As clearly depicted in the resulting visualization, the distribution of the data points emphatically demonstrates that they do not follow a straight line. The distinct pattern observed suggests a complex, potentially parabolic or higher-order relationship, thereby confirming the necessity of exploring various degrees of **polynomial models** to achieve an optimal fit.



Step 2: Implementing and Visualizing Polynomial Regression Models

Given the strongly non-linear characteristics identified in our dataset, our strategy involves fitting a sequence of [polynomial regression](#) models. We will systematically test degrees ranging from 1 (a simple linear model) up to degree 5 (a quintic model). This range allows us to directly observe how incrementally increasing the mathematical complexity of the model impacts its adherence to the empirical data points.

This implementation relies heavily on [NumPy](#)'s foundational functions for polynomial manipulation. Specifically, we use `polyfit()`, which is responsible for computing the coefficients that define the best-fit polynomial of the specified degree. These coefficients are then wrapped into a convenient, callable function object using `poly1d()`. By defining five distinct models (Model 1 through Model 5), we establish a necessary foundation for rigorous comparative statistical analysis.

Following the definition of these five mathematical models, we proceed to visualize all five fitted

curves simultaneously overlaid onto the original scatterplot. This direct visual comparison provides immediate insights: the linear model (Degree 1) is clearly seen to be severely underfitting the data, while the curves generated by the higher-degree models appear to conform much more accurately to the dense scatter of points, although visual assessment alone is insufficient for final model selection.

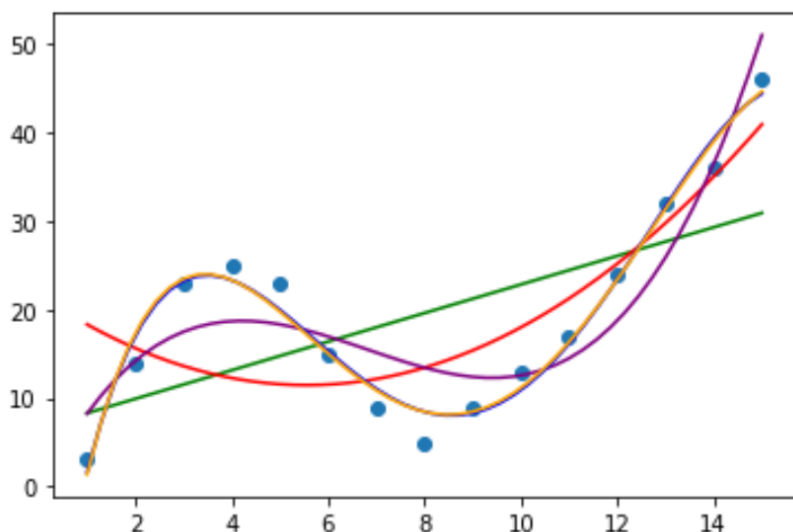
import numpy as np

```
#fit polynomial models up to degree 5
model1 = np.poly1d(np.polyfit(df.x, df.y, 1)) # Linear Model (Degree 1)
model2 = np.poly1d(np.polyfit(df.x, df.y, 2)) # Quadratic Model (Degree 2)
model3 = np.poly1d(np.polyfit(df.x, df.y, 3)) # Cubic Model (Degree 3)
model4 = np.poly1d(np.polyfit(df.x, df.y, 4)) # Quartic Model (Degree 4)
model5 = np.poly1d(np.polyfit(df.x, df.y, 5)) # Quintic Model (Degree 5)
```

```
#create scatterplot
polyline = np.linspace(1, 15, 50)
plt.scatter(df.x, df.y)
```

```
#add fitted polynomial lines to scatterplot
plt.plot(polyline, model1(polyline), color='green')
plt.plot(polyline, model2(polyline), color='red')
plt.plot(polyline, model3(polyline), color='purple')
plt.plot(polyline, model4(polyline), color='blue')
plt.plot(polyline, model5(polyline), color='orange')
plt.show()
```

The resulting plot clearly illustrates the dramatic visual differences in the quality of the fit. While the linear (green) and quadratic (red) models are unable to effectively capture the complex dips and peaks of the data, the higher-degree polynomials (purple, blue, and orange) demonstrate a much tighter adherence to the data points. This visualization confirms that a higher-order model is necessary, but it does not definitively tell us which specific degree is optimal without statistical validation.



Step 3: Evaluating Model Performance using Adjusted R-squared

Relying solely on visual assessment is inherently subjective and prone to favoring overfit models. To objectively and reliably determine which curve provides the best balance between explanatory power and model complexity, we must utilize a robust statistical metric: the [Adjusted R-squared](#) (R^2_{adj}) value.

The Adjusted R-squared is a far more reliable measure of **goodness of fit** than the standard R-squared (R^2) when comparing models that contain differing numbers of predictors (i.e., polynomial degrees). Standard R^2 is fundamentally flawed for model comparison because it will always increase, or at least remain constant, whenever a new variable or term is added to the model, even if that term offers no genuine explanatory improvement. The Adjusted R-squared elegantly corrects this bias by applying a penalty proportional to the number of predictors included in the model, thereby strongly favoring models that are both accurate and **parsimonious**.

A higher Adjusted R-squared value signifies a superior fit, indicating that the model explains a larger percentage of the variance in the response variable, after accounting for the inherent complexity introduced by the degree of the polynomial. We now define a custom function in Python to efficiently calculate this crucial metric for each of the five polynomial models we have defined.

#define function to calculate adjusted r-squared

```
def adjR(x, y, degree):
    results = {}
    coeffs = np.polyfit(x, y, degree)
    p = np.poly1d(coeffs)
    yhat = p(x)
```

```

ybar = np.sum(y)/len(y)
ssreg = np.sum((yhat-ybar)**2)
sstot = np.sum((y - ybar)**2)
results = 1- (((1-(ssreg/sstot))*(len(y)-1))/(len(y)-degree-1))

return results

#calculated adjusted R-squared of each model
adjR(df.x, df.y, 1)
adjR(df.x, df.y, 2)
adjR(df.x, df.y, 3)
adjR(df.x, df.y, 4)
adjR(df.x, df.y, 5)

{'r_squared': 0.3144819}
{'r_squared': 0.5186706}
{'r_squared': 0.7842864}
{'r_squared': 0.9590276}
{'r_squared': 0.9549709}

```

Upon meticulous review of the calculated output, we observe a clear and statistically significant trend: the Adjusted R-squared value increases substantially as the model degree is raised, peaking specifically at the fourth-degree polynomial. Beyond this point (at degree 5), the value slightly decreases. Critically, the fourth-degree polynomial yields the highest Adjusted R-squared value of **0.9590276**. This decisive result confirms that the 4th-degree model strikes the best possible balance between the variance explained and the intrinsic complexity of the model, establishing it as the most appropriate and optimal choice for fitting this particular dataset.

Step 4: Visualizing the Optimal Fit and Deriving the Predictive Equation

With the fourth-degree polynomial statistically validated as the optimal model, the final stages of the process involve generating a dedicated visualization and extracting the mathematical equation for practical predictive use. This visualization serves as the final confirmation of the high quality of the fit that was rigorously determined by the Adjusted R-squared calculation.

The following Python code block re-fits the optimal 4th-degree polynomial and then plots the resulting curve using a highly visible dashed red line. This plot distinctly highlights the curve's precise alignment and smooth conformity with the underlying scatter of observed data points:

```

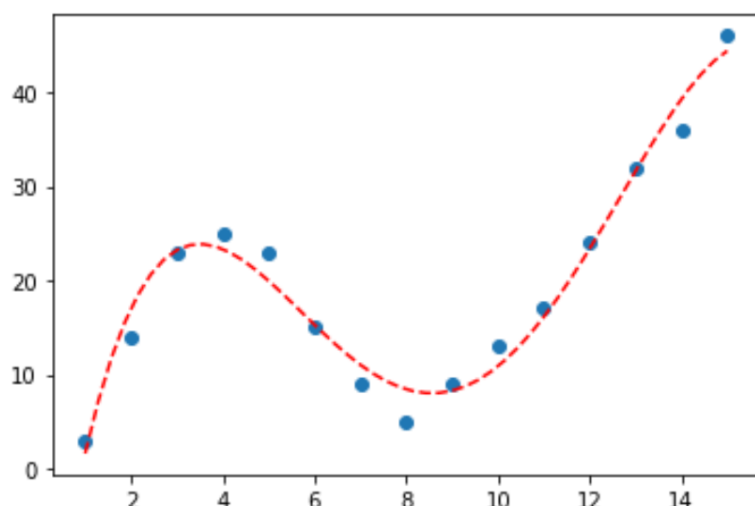
#fit fourth-degree polynomial
model4 = np.poly1d(np.polyfit(df.x, df.y, 4))

```

```
#define scatterplot
polyline = np.linspace(1, 15, 50)
plt.scatter(df.x, df.y)

#add fitted polynomial curve to scatterplot
plt.plot(polyline, model4(polyline), '--', color='red')
plt.show()
```

The resulting image visually confirms that the 4th-degree polynomial curve smoothly and accurately traverses the path defined by the data points, demonstrating an excellent, statistically justified fit. This curve is now ready for use in forecasting.



The core utility derived from curve fitting is the resulting mathematical formula, which allows for precise **predictive analysis**. We can easily retrieve the coefficients that define this optimal polynomial by using the simple `print()` function on the `model4` object:

```
print(model4)
```

```
4 3 2
-0.01924 x + 0.7081 x - 8.365 x + 35.82 x - 26.52
```

This output translates directly into the following explicit mathematical relationship that governs the dataset:

$$y = -0.01924x^4 + 0.7081x^3 - 8.365x^2 + 35.82x - 26.52$$

This precise equation can now be used to predict the value of the response variable (\$y\$) for any

new, hypothetical value of the predictor variable (x). For example, if we substitute $x = 4$ into our derived 4th-degree equation, we predict that the corresponding y value will be approximately **23.32**, a prediction based on the statistically optimal model:

$$y = -0.0192(4)^4 + 0.7081(4)^3 - 8.365(4)^2 + 35.82(4) - 26.52 = 23.32$$

Conclusion and Further Resources

Curve fitting in [Python](#), expertly facilitated by powerful and flexible libraries like [NumPy](#) and [Matplotlib](#), provides an efficient, systematic, and statistically rigorous methodology for modeling complex, non-linear data structures. By carefully comparing the performance of multiple polynomial degrees and grounding our final selection in the statistically sound metric of the [Adjusted R-squared](#), we were able to confidently select the fourth-degree model as the optimal fit for our sample data.

This comprehensive methodology ensures that the resulting predictive equation is not only visually accurate but also fundamentally justified from a statistical standpoint. By adhering to these principles, practitioners can successfully navigate the challenges of [overfitting](#) and [underfitting](#), leading directly to models with superior generalization capabilities and enhanced predictive power.

Additional Resources

To further deepen your expertise and technical understanding of these essential data modeling techniques, we highly recommend exploring the official documentation and related resources for the following libraries and concepts:

NumPy Documentation: For detailed information on core array manipulation, numerical processing, and polynomial manipulation functions, including `polyfit` and `poly1d`.

Matplotlib Gallery: For exploring advanced plotting techniques, customization options, and generating high-quality visualizations of fitted curves.

Statistical Modeling Texts: For a comprehensive theoretical background on regression analysis, model selection criteria, and the statistical foundations of the Adjusted R-squared metric.