



Geospatial Problem Solving

Spatial Interpolation

Image from NASA:
https://commons.wikimedia.org/wiki/File:Earth_Western_Hemisp_here_transparent_background.png#filelinks



AmericaView™
Empowering Earth Observation Education
AmericaView.org - Est. 2003

This module explores density estimation and spatial interpolation methods. I also provide a brief conceptual introduction to kriging.



Module Topics



1. Density estimation
2. Inverse distance weighting interpolation
3. Spline interpolation
4. Kriging interpolation
5. Types of kriging

These are the topics covered in this module.

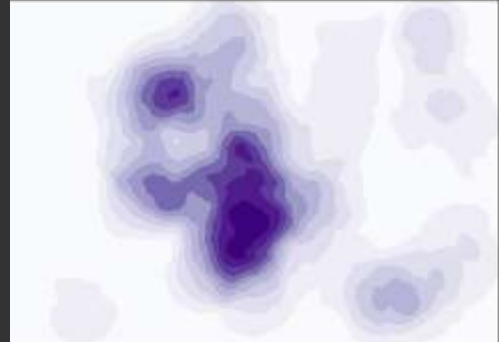
Review of Density Estimation



Point Density/Line Density



- ❖ Calculates the magnitude per unit area for point or line features that fall within a neighborhood around each cell
- ❖ Line length/area
- ❖ Number of points/area
- ❖ Can incorporate weights
- ❖ Must specify neighborhood shape and size



4

You may be interested in obtaining a raster-based measure of density for point or line features. This can be accomplished using the Point Density or Line Density methods.

These tools use neighborhood analysis or moving windows. The user must define the size and shape of the moving window. The tool will then calculate the number of points or length of lines within the window.

For the center cell of each window, the tool will return a measure of density. For points this will be the number of points per unit area. For lines, this will be the length of lines per unit area.

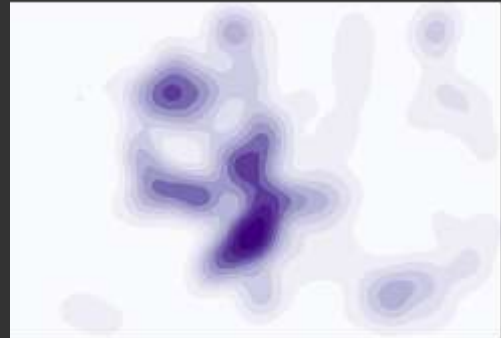
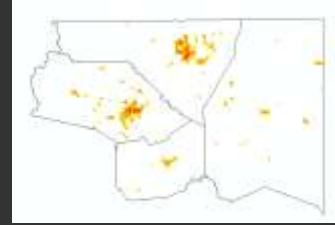
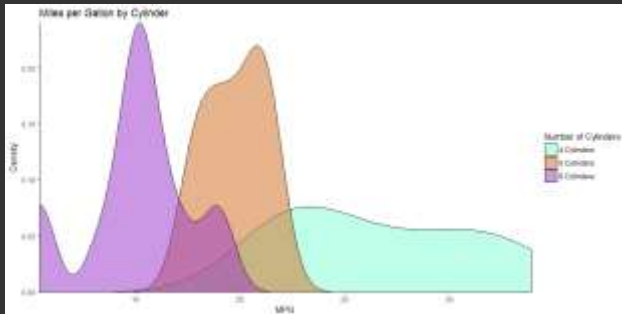
The size and shape of the window can have a large impact on the output. More generalized patterns will be produced with larger windows whereas more local patterns will be observed if a smaller window is used. There is not necessarily a correct window size, as this depends on the goals of the analysis.



Kernel Density



- ❖ Fits a kernel function to estimate density



5

Kernel density is similar to point or line density, except that a different method is used to calculate density. Specifically, this method makes use of a kernel function. The next slide provides an explanation of this method. Please take the time to read through this explanation.



Point Density vs. Kernel Density



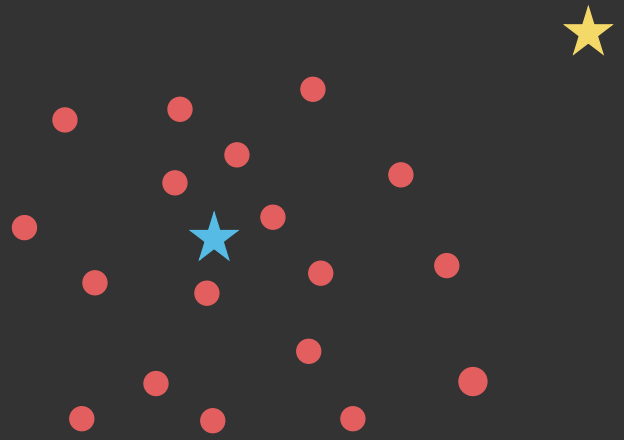
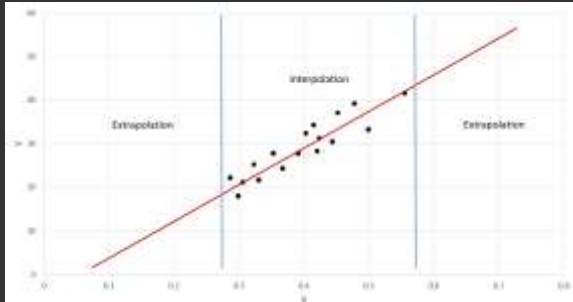
- ❖ In **point** and **line density**, a neighborhood is specified that calculates the density of the population around each output cell.
- ❖ **Kernel density** spreads the known quantity of the population for each point out from the point location. The resulting surfaces surrounding each point in kernel density are based on a quadratic formula with the highest value at the center of the surface (the point location) and tapering to zero at the search radius distance. For each output cell, the total number of the accumulated intersections of the individual spread surfaces is calculated.

This slide provides a comparison of the point density and kernel density methods.

Spatial Interpolation Methods



What is Spatial Interpolation?



8

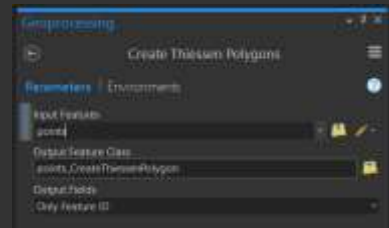
We can't measure everything everywhere. For example, it is common to collect measurements at point locations, such as collecting climate or weather data at weather stations. If you would like to estimate values at locations that were not sampled, you will need to apply interpolation or extrapolation as a means to obtain a continuous surface from your discrete measurements. Since these locations were not actually measured but estimated from nearby sample locations, there will be some error or uncertainty in the product.

Extrapolation is generally considered to be more uncertain than interpolation. As an example, if you are trying to predict a variable y using a variable x with a line or linear relationship, estimates within the bounds of your sample data will likely be less uncertain than estimates outside of the bounds of your data. Maybe outside of the bounds the relationship is no longer linear, begins to increase exponentially, or levels off. Without data points or samples, this would be hard to determine.

In a spatial context, extrapolation would be making an estimate outside of the geographic bounds of your data, as represented by the yellow star. In contrast, interpolation would be making an estimate within the geographic bounds of your data, as represented by the blue star.

Thiessen (Voronoi) Polygons

- ❖ Defines the area that is closer to one point than any other points in a dataset



One fairly simple method to perform interpolation is Thiessen (Voronoi) Polygons. Thiessen (Voronoi) polygons represent the area that is closer to one point than any other points in a dataset. The polygon can be attributed with the attributes of the input point layer.

In the examples provided, Thiessen polygons and point measurements in Ontario that represent weather stations are being used to estimate mean annual temperature for the entire province. In the West Virginia example, the entire area of the state has been divided into regions that are closer to one city than any other provided city.



Inverse Distance Weighting (IDW)



- ❖ Inverse distance weighted average of near sample points
- ❖ **Power** = controls the significance of surrounding points
 - ❖ Higher power = less influence by distant points
 - ❖ Lower power = more influence by distant points
- ❖ **Radius** = specifies which input point will be considered
 - ❖ **Variable** = does not specify a distance but a number of points
 - ❖ Can specify number of points to use
 - ❖ Can specify the maximum distance over which to search
 - ❖ **Fixed** = specify a distance over which to search
 - ❖ Can specify a minimum number of points that is used to increase the search radius

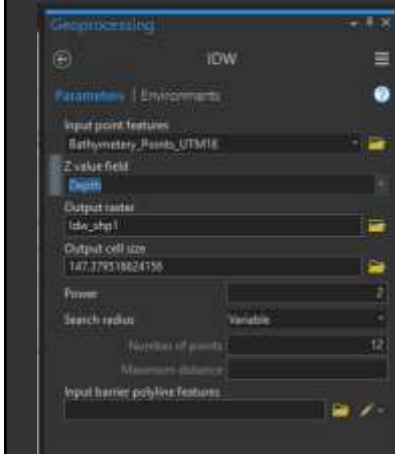
10

Another spatial interpolation method is inverse distance weighting or IDW. Using this method the value at a new location or cell is determined based on an inverse distance weighted average of nearby samples.

The power setting controls the significance of surrounding points. Using a higher power will cause the near points to have a higher relative weight than more distant points.

The radius setting specifies what samples will be used to estimate a value at a new location or cell. A variable radius is used to specify a number of samples as opposed to a fixed distance. For example, the nearest 12 points could be used. When using the variable method, you can also set a maximum distance. So, if 12 points are not available within the maximum distance then less points will be used. When using a fixed method, a fixed radius or distance is defined as opposed to a number of points. When using the fixed method, a minimum number of points can be set. So, if the minimum number of points do not exist within the search radius it will be expanded to include more samples.

Inverse Distance Weighting (IDW)



11

This slide shows an example output from IDW.

Here, bathymetric or water depth point measurements are being interpolated to a raster grid using a power of 2 and a variable search radius with no maximum distance and 12 as the number of points.

If you are not sure what settings to use, I would recommend experimenting with the settings and visualizing the result.

- ❖ Points extrude to the height of their magnitude
- ❖ Spline bends a sheet of rubber that passes through the input points while minimizing the total curvature of the surface
- ❖ Fits a mathematical function to a specific number of near input points while passing through the sample points
- ❖ Surface passes through sample points and has minimum curvature
- ❖ This method is best for interpolating smooth surface that do not change abruptly

Another interpolation method is spline.

Conceptually, point samples are extruded to a height relative to the magnitude of the attribute being interpolated. Spline then bends a sheet of rubber that passes through the input points while minimizing total curvature of the surface. It fits a mathematical function to a specific number of near input points while passing through the sample points. The goal is for the surface to pass through the sample points while minimizing the curvature of the surface.

Generally, this method is best for interpolating smooth surfaces that do not change abruptly. It generally is inappropriate for patterns or phenomena that can change abruptly.

Kriging



What is Kriging?



- ❖ Spatial variability consists of two components:
 - ❖ Global trends
 - ❖ Local spatial autocorrelation
- ❖ Need a model of global trends
- ❖ Need a model of spatial covariance
- ❖ A geostatistical method

We are discussing universal kriging here.

14

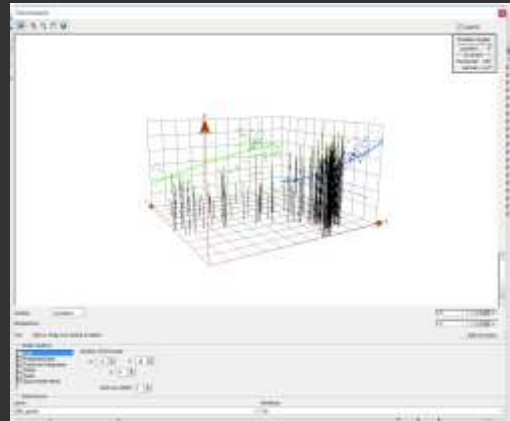
The last interpolation method we will explore is kriging. Kriging is a complex topic. My primary goal here is to provide a brief and conceptual introduction to this topic.

Kriging is a geostatistical interpolation method that models both a global trend using a mathematical function and local patterns using spatial autocorrelation.

More specifically, this form of kriging is termed universal kriging.



- ❖ Global geographic trend approximated with a mathematical model



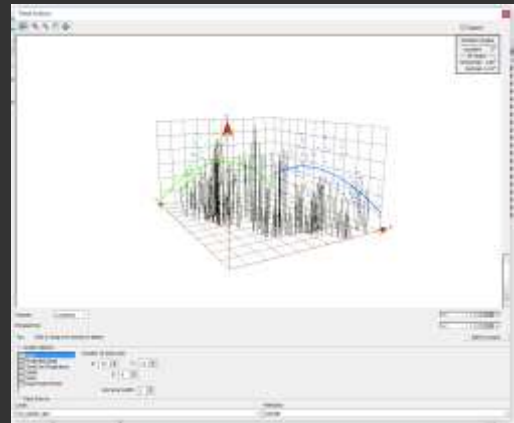
Temperature in Ontario

First, a global trend in the data is approximated using a mathematical function, such as a linear function or a polynomial.

The question here is what mathematical function can best represent the global trend in the data.



❖ Global geographic trend approximated with a mathematical model



Ozone in California

The provided example represents ground level ozone measurements across California.

Each sample point is extruded to a height relative to the ozone measurements. The trends printed on the walls of the graph represent a general pattern in the north-to-south and east-to-west directions. In both directions, the highest values tend to be concentrated in the middle of the extent.

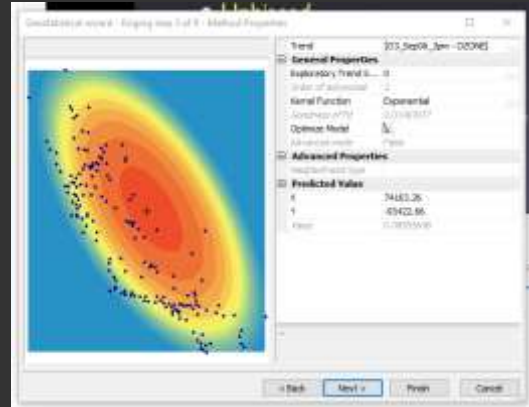
So, this global trend could be approximated using a second-order polynomial function.



Improving upon the estimate made by the global trend



- ❖ Reduce the residual by modeling a **local trend**
- ❖ Explore the spatial dependency of the residuals
- ❖ Model **spatial autocorrelation** using a **semivariogram**



Ozone in California

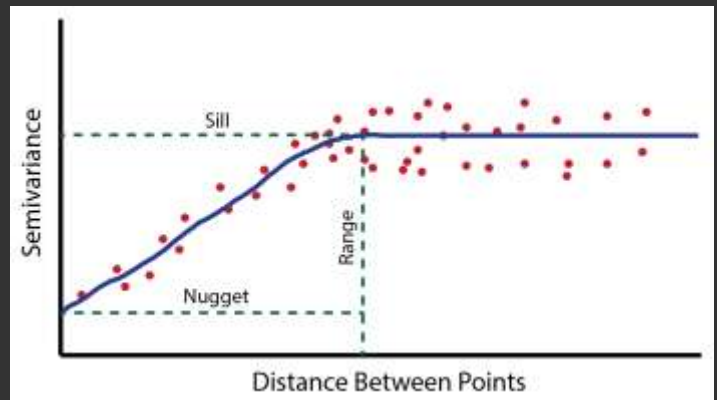
This slide shows the global trend in the ground level ozone data as a polynomial.

This is very generalized. The next step is to improve the estimate by assessing and incorporating local spatial autocorrelation in the residuals or error term.

This can be accomplished with the semivariogram.

Review of Semivariogram

- ❖ **Semivariance** = measure of the spatial dependence between two observations as a function of the distance between them
- ❖ **Semivariance** = squared difference between the two values
- ❖ **Semivariogram** = a graph of how semivariance changes as the distance between observations changes



18

The following slides offer a review of the semivariogram, which was discussed in the prior module.

A semivariogram is a graphical representation of the concept of spatial autocorrelation.

On a semivariogram, the x-axis represents distance between sample points, or lag. In the graph space, each point represents a pair of samples or observations. The x-coordinate is the lag or distance between the samples. The y-axis is the semivariance, which is a measure of the spatial dependency between two samples or observations. High semivariance suggests that the two observations or samples have very different values whereas low values of semivariance suggest they have very similar values.

Based on Tobler's First Law of Geography, we would expect near objects to have more similar values and further objects to have less similar values. In other words, near objects are expected to have low semivariance whereas samples that are more separated are expected to have high semivariance. This is the concept of spatial autocorrelation.

The semivariogram shown here represents simulated data. Real data are likely to show a more complex or noisy pattern.

The next slide will explain key features of the semivariogram.

Review of Semivariogram

Range: distance where model first flattens out

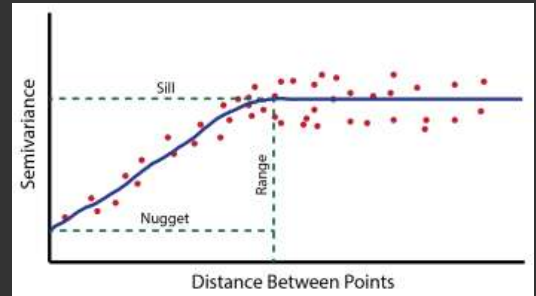
Sill: value of semi-variance at range

Nugget: semi-variance at distance = zero

Partial Sill: Sill minus nugget

*Theoretically, at zero separation distance (lag = 0), the semivariogram value is 0. However, at an infinitesimally small separation distance, the semivariogram often exhibits a nugget effect, which is some value greater than 0.

*The nugget effect can be attributed to measurement errors or spatial sources of variation at distances smaller than the sampling interval or both. Measurement error occurs because of the error inherent in measuring devices. Natural phenomena can vary spatially over a range of scales. Variation at microscales smaller than the sampling distances will appear as part of the nugget effect. Before collecting data, it is important to gain some understanding of the scales of spatial variation.



19

First, the distance at which the semivariance stops increasing is the range. This value can be determined by extrapolating from the fitted semivariance curve to the x-axis. Samples that are closer together than this distance are spatially autocorrelated. Samples that are further than this distance are not spatially autocorrelated.

The sill is the semivariance at the range or the semivariance at the distance where spatial autocorrelation is no longer present.

You would expect the semivariance at a distance or lag of zero to be zero. Or, you would expect samples at the same location to have the same value. However, this is not always the case for real data. This is known as the nugget effect. This arises due to measurement error or spatial sources of variation at distances smaller than the sampling unit.

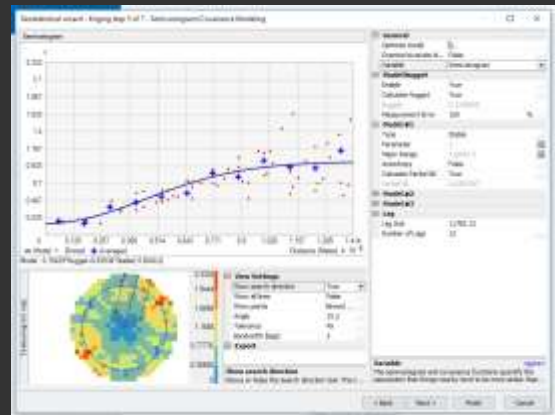
The partial sill is the semivariance obtained when the nugget is subtracted from the sill.

Again, the semivariogram can be thought of as a graphic representation of spatial autocorrelation. In this example, near features tend to have a lower semivariance, or are more correlated, whereas more distant features tend to have a greater semivariance or are less correlated. At a certain distance, the

range, samples are no longer spatially autocorrelated and the semivariance stops increasing. So, near features have more similar values than observations that are more separated. In other words, spatial autocorrelation is present.



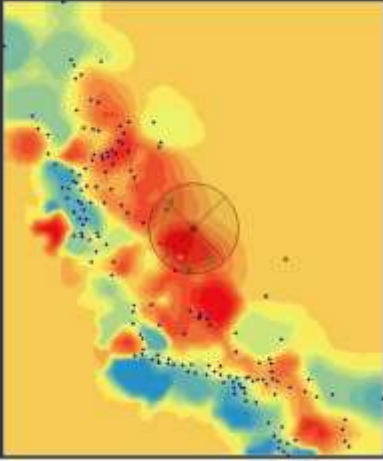
- ❖ Directional differences in spatial dependence
- ❖ Can incorporate this into the kriging model by exploring the semivariogram



Ozone in California

It is also possible to incorporate measures of directionality in our analysis of the semivariogram since there may be directional difference in spatial dependence.

This directional influence is known as anisotropy and can be explored using the semivariogram.



Ozone in California

1. Exploratory statistical analysis of the data
2. Semivariogram modeling
3. Creating the surface
4. Exploring a variance surface (optional)

❖ **Kriging** is most appropriate when you know there is a spatially correlated distance or directional bias in the data

Performing kriging involves multiple steps including exploring the data, semivariogram modeling, surface generation, and assessing error and variance.

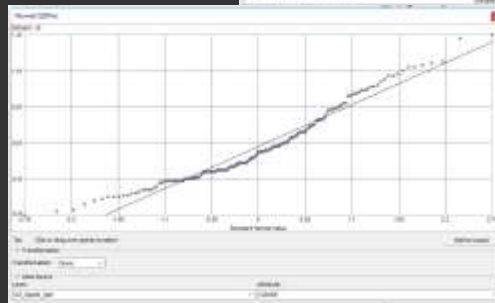
Kriging tends to outperform other interpolation methods when there is strong spatial autocorrelation and/or directional bias in the data.



Things to Consider



1. Is the data **normally distributed**?
2. Is there **skewness**?
3. Is there **kurtosis**?
4. Does the **QQ Plot** suggest normality?



Ozone in California

22

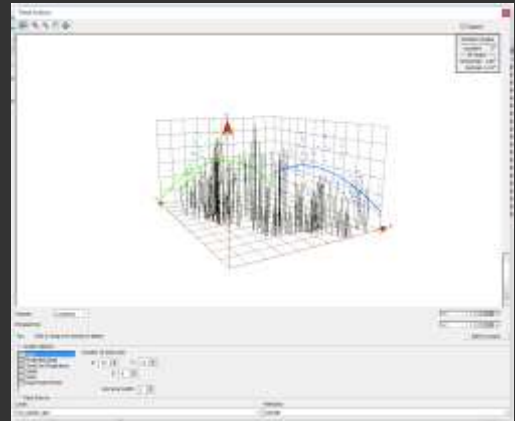
When exploring your data, it is important to consider if the attribute being interpolated is normally distributed.

Is there skewness? Is there kurtosis? What does the QQ plot suggest?

If the data deviate strongly from normal, they may need to be transformed. For example, a log, square root, or square transformation could be applied.



5. What is the **global trend**?



Ozone in California

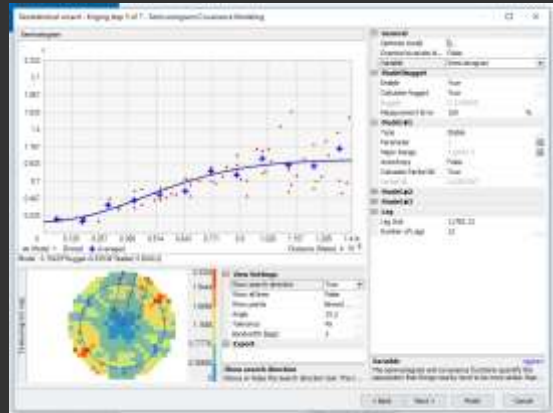
You need to explore the global trends. What type of mathematical model would best approximate this trend?



Things to Consider



6. What does the **semivariogram** tell you about **spatial autocorrelation**?
7. Should you consider **anisotropy**?



Ozone in California

You then need to explore the local patterns using the semivariogram.

It is also important to consider if there is a directional influence or whether anisotropy exists in the data.



Things to Consider



An analysis of your data will allow you to make decisions relating to:

1. If data transformations should be applied
2. Best function to fit
3. Sill, range, and nugget settings
4. Directional considerations (anisotropic)
5. Best search direction and width

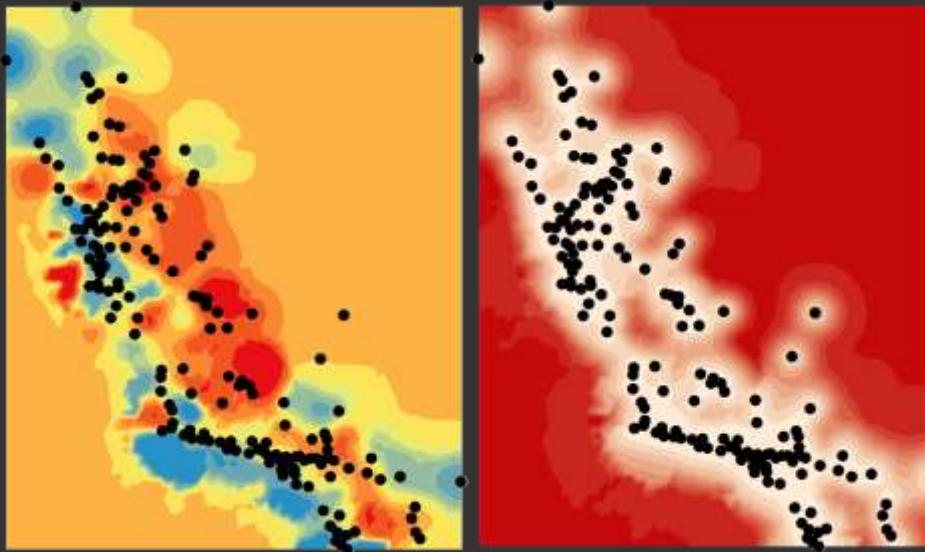
25

Working through the data exploration and analysis will allow you to determine if the variable needs to be transformed in order to make the distribution more normal.

It will allow you to decide what global trend exists and what mathematical function best approximates it.

When assessing local variability using the semivariogram, you will be able to determine what settings to use to model spatial autocorrelation as the sill, range, and nugget.

You will also be able to determine if there are any directional trends and how to incorporate them into the model.



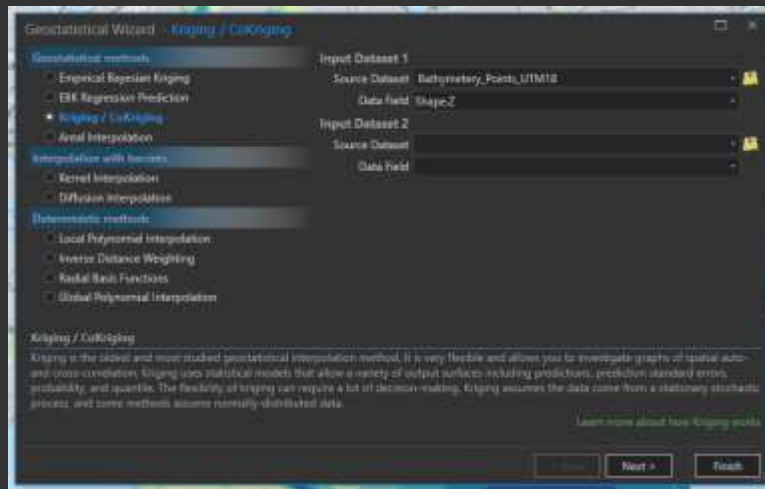
Ozone in California

26

This slide shows the results for the kriging Interpolation of ground level ozone in California.

The first image represents the ozone predictions and the second model represents the predicted variance values.

There is more certainty or less variance for the predictions near the sample points. There is more uncertainty further from the sample points.



ArcGIS Pro provides the Geostatistical Analyst Extension that can guide you through the process of performing a variety of analyses, including kriging.



Types of Kriging



Method	Mean assumption	Stationarity	Uses covariates?	Output	Typical use
Ordinary	Unknown, constant	Yes	No	Point estimate	General-purpose interpolation
Simple	Known, constant	Yes	No	Point estimate	Simulation; known-mean processes
Universal	Polynomial trend in coords	Residuals only	No	Point estimate	Data with spatial drift or trend
Kriging w/ External Drift	Trend via covariate	Residuals only	Yes	Point estimate	Temp. mapped using elevation, etc.
Regression kriging	Regression on covariates	Residuals only	Yes	Point estimate	Soil/environmental mapping
Co-kriging	Unknown, constant	Yes	Yes (co-variables)	Point estimate	Multiple correlated variables
Indicator	Unknown, constant	Yes	No	Probability	Risk mapping; non-Gaussian data
Probability	Unknown, constant	Yes	No	Probability	Improved threshold exceedance
Disjunctive	Unknown, constant	Yes	No	Probability / nonlinear	Nonlinear functionals; ore grade
Block	Unknown, constant	Yes	No	Areal average	Mining blocks; land-use averages
Lognormal	Unknown, constant (log scale)	Yes (log scale)	No	Point estimate	Skewed data: pollutants, ore grades
Bayesian	Prior-informed	Flexible	Optional	Posterior distribution	Small samples; full uncertainty

28

Again, this is just a broad overview of kriging. Kriging is a complex topic with many variants. It is generally covered in more detail in an advanced spatial analysis or spatial statistics course. This slide provides a simple table that summarizes different kriging variants.

Method = name of the kriging variant.

Mean assumption = this is one of the most fundamental distinctions between kriging types. Kriging models a spatial field as a mean component plus a random, spatially correlated residual. This column describes what assumption is made about that mean. Is it known in advance (simple kriging)? Unknown but treated as a constant (ordinary kriging)? Or does it vary across space as some kind of trend (universal, KED, regression kriging)?

Stationarity = a process is stationary if its statistical properties (mean, variance, spatial correlation structure) don't change depending on where you are in the study area. Most kriging methods require this, at least for the residual component after any trend is removed. "Residuals only" means the raw data aren't stationary, but once the trend is subtracted the leftover variation is assumed to be. Bayesian kriging is marked flexible because the prior can be set up to relax stationarity assumptions more readily.

Uses covariates? = indicates whether the method incorporates auxiliary or secondary variables beyond the primary observation locations. KED and regression kriging, for example, use an external variable like elevation or land cover that is known everywhere. Co-kriging uses secondary variables that are only sampled at points, like the primary variable. Most classical kriging methods use no covariates at all.

Output = what the method actually produces. Most methods give a point estimate (a predicted value at an unsampled location, along with a kriging variance). Indicator, probability, and disjunctive kriging produce probability estimates — the chance that the true value exceeds some threshold. Block kriging produces an average over an area rather than a single point. Bayesian kriging produces a full posterior distribution, which implicitly contains far more information than a single estimate and variance.

Typical use = illustrative examples of the kinds of real-world problems each method is best suited for, just to ground the more abstract properties in something concrete.



This is the end of this lecture module.

Please return to the West Virginia View
Webpage for additional content.



Thanks! Hope you found this useful.