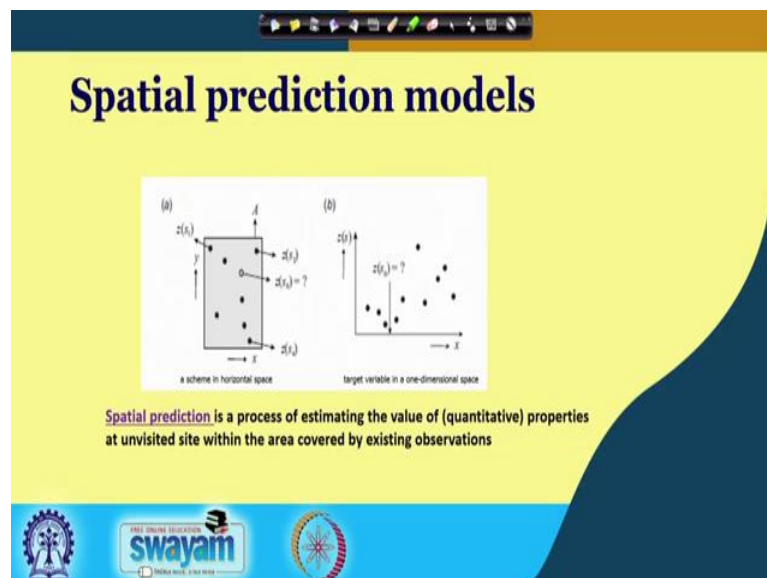


Soil Science and Technology
Prof. Somsubhra Chakraborty
Department of Agricultural and Food Engineering
Indian Institute of Technology, Kharagpur

Lecture - 53
Basics of VisNIR – DRS

Welcome friends to this 3rd lecture of week 11 of Soils Science and Technology and in this lecture we will be finishing of first the Geostatistics and then we will be starting the new sensor that is Diffuse Reflectance Spectroscopy.

(Refer Slide Time: 00:32)

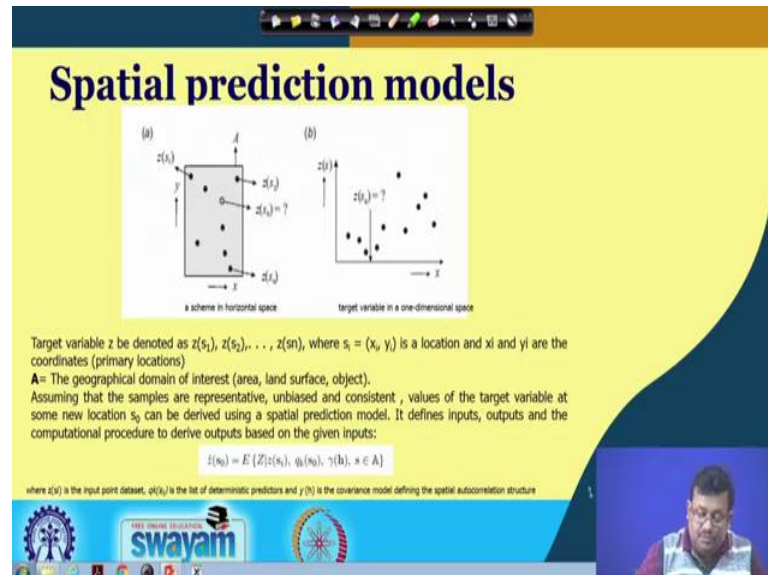


So, in the last lecture we stopped here about the you know in this slide while talking about the spatial prediction model and we remember that the spatial prediction is a process of estimating the value or in terms of quantitative value properties at unvisited side within the area covered by existing observation.

So, if you see here this is our area of interest and this area of interest or in our study area this is an obviously this scheme in the horizontal space. So, obviously there are some observation here like you know denoted by $Z(s_1)$, $Z(s_2)$ and $Z(s_n)$ and we want and observe we want to predict our value at this an observed point that is $Z(s_0)$. Similarly for a target variable in a one directional space you can see here as the target is you know variable distance of you know increases from particular point how this variation changes and you want to predict a value here.

So, we need to consider all these observed points while we are going to predict that particular point. So, let us see how it does how we generally do that.

(Refer Slide Time: 01:45)



So, let us go ahead and see a target variable z can be denoted at $Z(s_1)$, $Z(s_2)$ and up to $Z(s_n)$ where s_i is basically (X_i, Y_i) which is a location and (X_i, Y_i) are the coordinates like primary locations. So, A in our case this is the total area though this is A is our the total geographical domain of interest or in other words, it is the study area. So, assuming that the samples are representative, unbiased and consistent, then values of the target variable at some new location s_0 for example here can be derived using a spatial prediction model.

And it defines both inputs outputs and the computational procedure to derive outputs based on the given inputs. Now, you can see this is the formula through which we can derive you know we can we can define this model where s_0 is an un you know $Z \cap s_0$ is equal to is the value at that unobserved point and un of unobserved point that is s_0 and obviously, this $Z(s_i)$ is the input point data set whereas, the $q_k(s_0)$ is the list of deterministic predictors and this $\gamma(h)$ is basically covariance model defining the spatial autocorrelation structure. Now what is spatial autocorrelation we will see that later on.

(Refer Slide Time: 03:06)

Geostatistical modelling

- **Premise:** One cannot obtain error-free estimates of unknowns (or find a deterministic model)
- **Approach:** Use statistical methods to reduce and estimate the error of estimating unknowns (must use a probabilistic model)
- **Estimation of error:** We need to develop a good estimate of an unknown. Say we have three estimates of an unknown:

Want to estimate unknown, $\hat{T}_0$

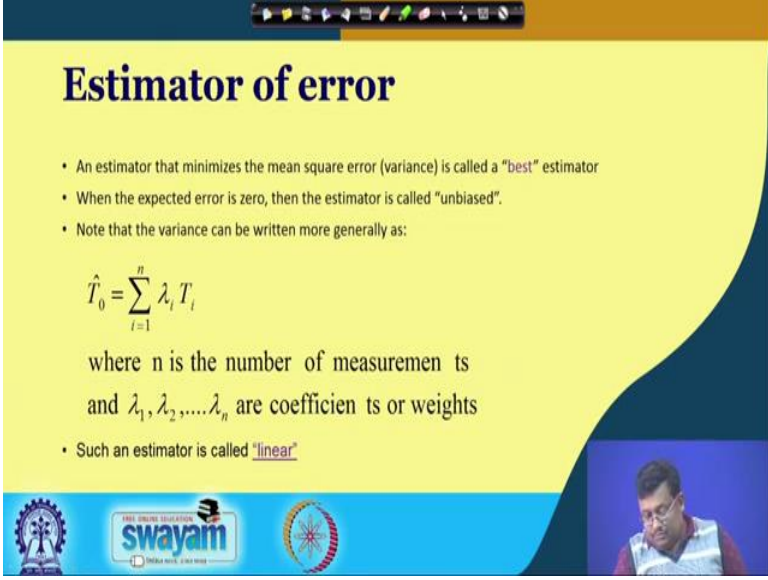
$$\sigma_0^2 = \frac{1}{3}(\hat{T}_0 - T_1)^2 + \frac{1}{3}(\hat{T}_0 - T_2)^2 + \frac{1}{3}(\hat{T}_0 - T_3)^2$$

where σ_0^2 is the mean square error ✓

So, geostatistical modelling let us talk about geostatistical modelling. Geostatistical modelling has the premise that one cannot obtain error free estimates of unknown. So, we have to field find or find a deterministic model. So, there is no possibility. Find a deterministic model and thus our approach would be to use statistical method to reduce the estimate of the error of estimating unknowns. So, it must be a probabilistic model. So, you can see estimation of error here we need to develop a good estimate of unknown say we have a three estimates of an unknown want to estimate the unknown T cap 0.

So, obviously it has to be done by these variance measurements. So, it is basically this σ_0 square is the mean square error.

(Refer Slide Time: 03:56)



Estimator of error

- An estimator that minimizes the mean square error (variance) is called a "best" estimator
- When the expected error is zero, then the estimator is called "unbiased".
- Note that the variance can be written more generally as:

$$\hat{T}_0 = \sum_{i=1}^n \lambda_i T_i$$

where n is the number of measurements
and $\lambda_1, \lambda_2, \dots, \lambda_n$ are coefficients or weights

- Such an estimator is called "linear"

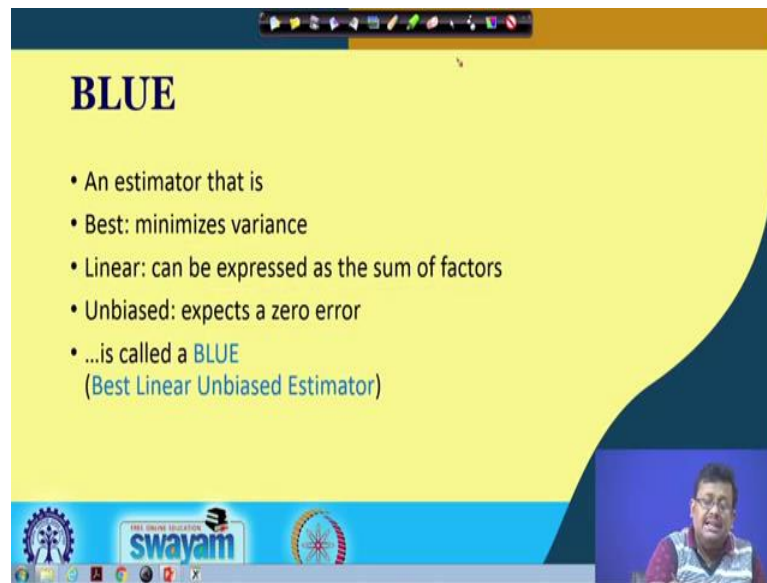
Logos: IIT Bombay, swayam, and a circular emblem.

So, similarly if we go ahead and see the estimation of the error that an estimator that minimizes the mean square error of variance that is the mean square error, obviously that is an variance is called the best estimator. So, this is very very important of the best estimator or the best another you know again the best estimator is the estimator that minimizes the mean squared error of or in other words variance.

So, when the expected error is 0, then the estimator is called unbiased obviously and note the variance can be written the more generally as this format where is the you know summation of $\lambda_i T_i$ for points 1 to n where n is the number of measurements and this λ_1 λ_2 up to λ_n are coefficients or weights. So, such an estimator is called the linear estimator because we are linearly producing we are linearly expressing this you know linearly expressing this value.

So, the variance can be more generally you know expressed as an you know is a summation of this you know individual weights and.

(Refer Slide Time: 05:16)



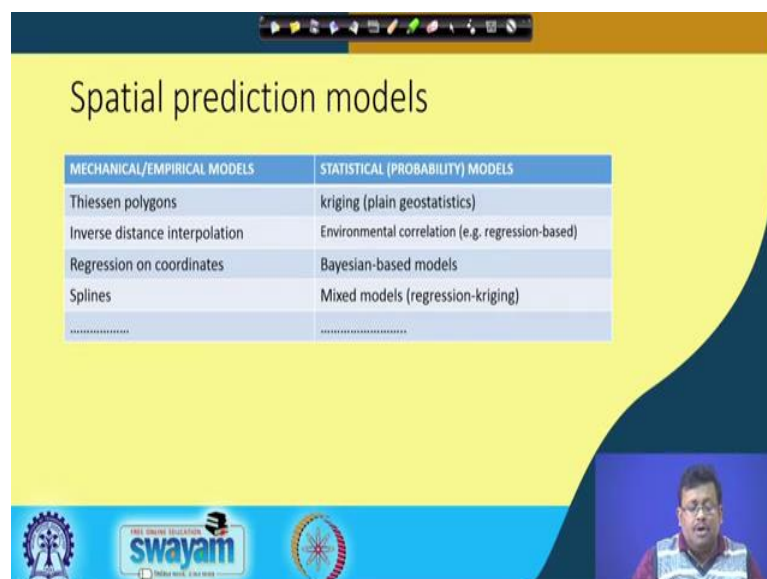
BLUE

- An estimator that is
- Best: minimizes variance
- Linear: can be expressed as the sum of factors
- Unbiased: expects a zero error
- ...is called a **BLUE**
(Best Linear Unbiased Estimator)

The slide features a yellow background with a dark blue curved shape on the right. At the bottom, there is a blue banner with the 'swayam' logo and a small video inset of a presenter on the right.

So, that is why it is called BLUE. So, BLUE is basically a short form of Best Linear Unbiased Estimator. So, you know Blue stands for B stands for which minimizes the variance and L stands for which can be expressed as the terms as the sum of factors and U stands for un unbiased that is expects a zero error. So, it is called a BLUE estimator. So, similarly we will see that Kriging which is an important geostatistical interpolation is a BLUE estimator.

(Refer Slide Time: 05:53)



Spatial prediction models

MECHANICAL/EMPIRICAL MODELS	STATISTICAL (PROBABILITY) MODELS
Thiessen polygons	kriging (plain geostatistics)
Inverse distance interpolation	Environmental correlation (e.g. regression-based)
Regression on coordinates	Bayesian-based models
Splines	Mixed models (regression-kriging)

The slide features a yellow background with a dark blue curved shape on the right. At the bottom, there is a blue banner with the 'swayam' logo and a small video inset of a presenter on the right.

So, let us see what are the different types of spatial prediction model. So, there are two different types of spatial predict categories of spatial prediction model. One is called the mechanical model or empirical models; another is statistical model or probabilistic model. Now in the mechanical model there are several types of mechanical model. One is Thiessen Polygon, then Inverse Distance Interpolation, then Regression on Coordinates Splines so on and so forth and also in the statistical or probabilistic model you can see kriging which is a plain geostatistics. Secondly, environmental correlation which is basically regression based and then Bayesian based models and finally, mixed model which is regression kriging.

We will talk about regression kriging later on while we will be discussing digital soil mapping. So, you see there are two types of models.

(Refer Slide Time: 06:44)

MECHANICAL/EMPIRICAL MODELS	STATISTICAL (PROBABILITY) MODELS
Arbitrary or empirical model parameters are used. No estimate of the model error is available and usually no strict assumptions about the variability of a feature exist	the model parameters are commonly estimated in an objective way, following the probability theory. The predictions are accompanied with the estimate of the prediction error.
	Drawback: input dataset usually need to satisfy strict statistical assumptions

So, what is the difference between this Mechanistic Mechanical model as well as the Statistical model. Remember that these mechanical models are you know arbitrary or empirical model parameters. You know this model this in this mechanical model we generally use arbitrary or empirical model parameters and no estimate of the model error is available and usually no strict assumption about the variability of a feature exist.

So, these are the features of mechanical model, however in case of statistical or probabilistic model you see that the model parameters are commonly estimated in a objective way following the probability theory and the predictors are accompanied by

you know with the estimate of the prediction error. So, you will see not only the prediction, but also you will see side by side there you know estimate of their error.

So, obviously there are some drawbacks. Drawbacks is input data set usually need to satisfy strict statistical assumptions.

(Refer Slide Time: 07:39)

Inverse distance interpolation

- A value of target variable at some new location can be derived as a weighted average:

$$\hat{z}(s_0) = \sum_{i=1}^n \lambda_i z(s_i)$$
- where λ_i is the weight for neighbour i . The sum of weights needs to equal one to ensure an unbiased interpolator.
- Matrix form is:

$$\hat{z}(s_0) = \lambda_0^T \cdot z$$
- The simplest approach for determining the weights is to use the inverse distances from all points to the new point:

$$\lambda_i(s_0) = \frac{1}{\sum_{j=1}^n \frac{1}{d(s_0, s_j)^\beta}} \quad \beta > 1$$
- Where $d(s_0, s_j)$ is the distance from the new point to a known sampled point and β is a coefficient that is used to adjust the weights. This way, points which are close to an output pixel will obtain large weights and that points which are farther away from an output pixel will obtain small weights. The higher the β , the less importance will be put on distant points. The remaining problem is how to estimate β objectively so that it reflects the inherent properties of a dataset.

So, this is the drawback of this probabilistic model. So, let us see one example for each. So, Inverse Distance Interpolation which is an important mechanistic model or mechanical model for interpolation, so you see in the inverse distance interpolation in short form we call it IDI or we some time call it IDW also. So, a value of target variable at some new location can be derived as a weighted average.

So, you can see just like we have shown that then be estimated can be you know variance can be shown as a you know as a summation of individual weights and the points just like λ_i , T_i if you remember in the last couple of slides we have talked about. So, similarly a value of a target variable and some new location here can be expressed as an weighted average. So, here λ is a weight for the neighbour i and the sum of weights needs to be equal to 1, should be equal to 1 to ensure an unbiased estimator. So, it should be equal to 1 you remember that and matrix form is obviously this is the matrix form.

There is the simplest approach for determining the weight is to use the inverse distances from all points to a new point. So, you see we have we can assign any value at a

particular point based on its nearby neighbouring points, however we have to assign some particular values of weight to this individual neighbouring points. So, the simplest approach for determining the weight is to use the inverse distance from the all points to the new point. That means, those points which are close to the unobserved points or the point of interest we will have better or higher weights than that of the points which are more further away from the point of interest.

So, you can see here the λ_i you can be it can be calculated by $1/d^\beta$ where d is the distance from the new point of a known sampled point and β is the coefficient that is used to adjust the weights. So, this is the way points which are close to an output pixel. We will obtain large weights as I have talked about and the points which are further away.

So, for example if this is a particular point, so the points which are closed by will have better weight than that of the points which are having which are located far away from the these points. So, in that way the you know large weights and you know the you know the points which are closed to output pixel will obtain large weights and that points which are further away from an output pixel will obtain small weights. So, the higher the β , the less importance will be put into that put on distance points.

So, the remaining problem is to how to estimate beta objectively. So, that reflects the inherent properties of a data set. Now, the problem is the objectively you know the major problem for IDW is to calculate this β objective we cannot do that in you know in this inverse distance interpolation. That is the major problem of inverse distance interpolation. Although beta gives us the relative distribution of weightage for the points which are nearby to this point of interest or output pixel as compared to those points which are further away from the output pixel, however there is no objective way to determine these beta. So, this is what we call inverse distance interpolation.

(Refer Slide Time: 11:38)

Kriging

- For many decades been used as a synonym for geostatistical interpolation.
- It originated in the mining industry in the early 1950's as a means of improving ore reserve estimation.
- A standard version of kriging is called ordinary kriging (OK).

$$Z(s) = \mu + \epsilon'(s)$$

- μ = constant stationary function (global mean)
- $\epsilon'(s)$ = spatially correlated stochastic part of variation.

$$\hat{Z}(s_0) = \sum_{i=1}^n w_i(s_0) \cdot z(s_i) = \lambda_0^T \cdot z$$

λ_0 = vector of kriging weights (w_i)
 z = vector of n observations at primary locations
In a way, kriging can be seen as a sophistication of the inverse distance interpolation

swayam

Now, what is Kriging? Kriging's another name is Geostatistical Interpolation and this is more scientific and we can say it is a sophistication of IDW or IDI. So, it originated in the mining industry in the early 1950s as a measure of improving one you know ore reserve estimation. As you can see this is a Kriging map I have produced here. So, a standard version of Kriging is called ordinary kriging. Obviously, there are several types of kriging you see Universal kriging, Indicator kriging, Co kriging. You will see several types of kriging and then, regression kriging, however this is there you know we will be talking about only the simplest ordinary kriging.

So, a standard version of kriging can be termed as a ordinary kriging and this ordinary kriging can be represented here where this is an value of an you know of a pixel can be represented as a combination of is μ and this ϵ dash s . So, basically μ represents the constant stationary function or global mean where μ you know ϵ dash s basically spatially correlated stochastic part of the variation or in other words, this ordinary kriging can take this mathematical formula where again you can see this taking this summation of $\lambda_i T_i$ which we have seen previously. So, it is basically also we can see it in case of IDI.

So, here λ_0 is basically vector of kriging weights or W_i and Z is basically vector of n observation at primary locations and in a way kriging can be seen as a sophistication of the inverse distance interpolations.

(Refer Slide Time: 13:22)

Variogram

Problem of IDI: to determine how much importance should be given to each neighbor

Variogram — differences between the neighbouring values

$$\gamma(h) = \frac{1}{2} E \left[(z(s_i) - z(s_i + h))^2 \right]$$
$$\hat{\gamma}(h) = \frac{1}{2N(h)} \sum_{i=1}^N \left[z(s_i) - z(s_i + h) \right]^2$$

$z(s_i)$ = value of target variable at some sampled location
 $z(s_i + h)$ = value of the neighbour at distance $(s_i + h)$

1. Measures the variability of data with respect to spatial distribution
2. Specifically, looks at variance between pairs of data points over a range of separation scales

The graph shows a curve starting at the origin and leveling off, representing the variogram. The x-axis is labeled 'distance' and the y-axis is labeled 'variogram'. Data points are plotted along the curve.

swayam

So, obviously why it is called Inverse Sophistication of Inverse Distance Interpolation? Because of this feature called Variogram because you remember in case of inverse distance interpolation we calculated beta for assigning the weights, but there was no objective way to calculate the beta.

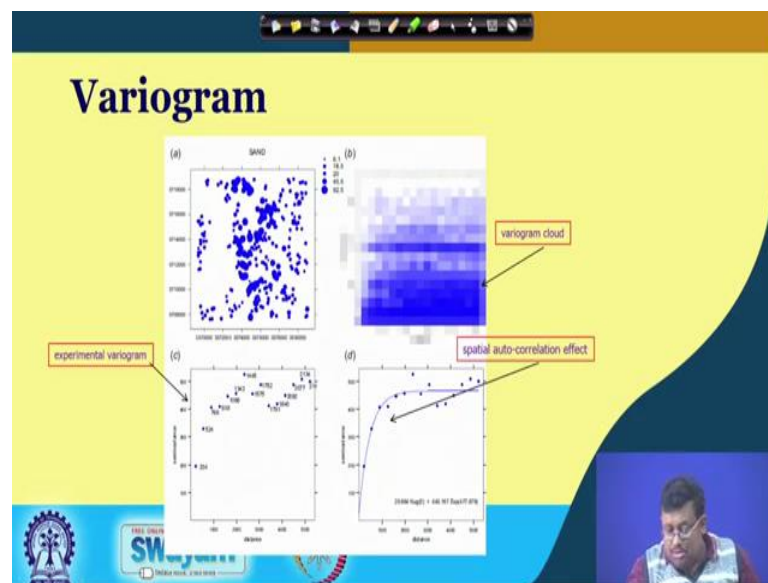
However, these weights in the kriging interpolation can be calculated objectively by using a using a specific model we call it Variogram. Now, in the variogram in the problem of IDI used to determine how much importance should be given to each neighbouring points. We cannot assign that we did not know that. So, variogram can you know variogram can solve that problem. So, variogram can be expressed is a difference between the neighbouring values. So, here h is basically lag or in the x axis if we have the distance between the point pairs and in the y axis basically we are putting the variance.

So, this variogram is basically a graph in a simple term it is basically a graph between the distance between the point pairs and the variance of the values between the point pairs. So, you can see the you know here this is the mathematical formula of a variogram or you know or the same of the variogram here you know $z(s_i)$ is basically the value of a target variable at some sample location, where $Z(s_i) + h$ is the value of neighbour at distance $s_i + h$ and using this formula n is the number of point pairs. So, using this formula using this formula we can see you know how these variances you know we by

using this formula we can plot this you know we can plot this graph and this is called semi-variogram.

So, basically what is the importance of the semi-variogram or variogram? It basically measures the variability of data with respect to spatial distribution specifically look at the variance between the pairs of data points over a range of separation scales or h.

(Refer Slide Time: 15:30)



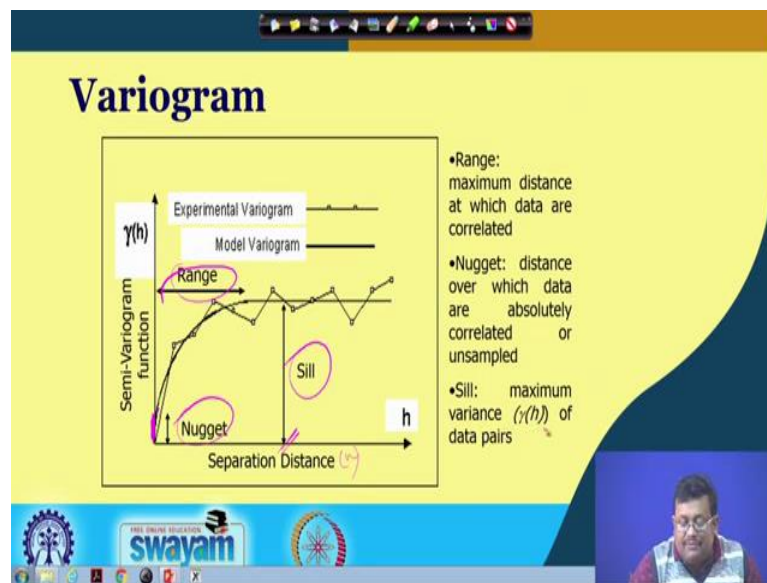
So, you can see here this is the variation you know this is the you know experimental variogram. So, basically in the in this graph, it shows in the x axis the distance in the y axis. Obviously, the obviously these are some data clouds and these are the data pairs. So, in the x axis we are putting the distance and the y axis would have putting the semi variance. Now, you know how to calculate the semi variance we have just seen in the last slide.

And then we are fitting a we are fitting a line by least square method and this is called the you know empirical semi-variogram and these basically shows spatial autocorrelation effect.

Now what is Spatial Auto-correlation? This is very important. Now you see here as the distance between the point pairs increases in the x axis, the semi-variance is continuously increasing up to a certain point and then it is getting you know it is reaching a plateau. So, you see up to this point for example up to this point it this semi-variance is

increasing. So, up to this point there is a spatial dependence between the points. In other words, you will see that if the nearby points are more spatially correlated, then those points which are further away and this relationship is called the spatial dependence or spatial autocorrelation and this spatial autocorrelation structure can be captured by using a variance using a variogram. So, there are several other terms we will see that later on.

(Refer Slide Time: 17:15)



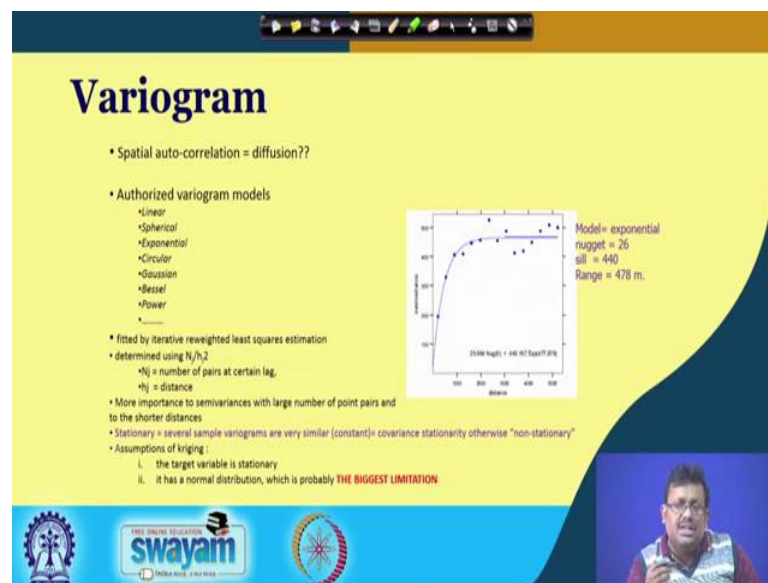
So, a variogram this is the variogram or semi-variogram function. You can see how this semi-variogram or variance we have calculated. I have shown you the I have already shown you the formula. Now you see in the x axis basically we call lag or separation distance we generally termed we generally expressed is in terms of h .

So, you see that as the separation distance of point pairs increases, this semi-variogram or semi-variance increases up to certain point and then, it reaches the plateau. So, that means and this you know it does not starts from a origin; it starts from certain points and these distance is called the Nugget. At this nugget is basically denotes this Nugget basically denotes the pure noise or measurement error and the total variance or maximum variance is called the Sill. So, this is called Nugget which is the basically you know the measurement error or pure noise which you cannot measure and this is the Total Variance or Sill.

So, this is the Total Variance or Sill at the distance separated separate distance separation distance at which these semi-variograms take a flat shape is called the Range. So, you see these are the model wave, these are the experimental variogram and this is the model variogram. This model can take several forms like you know it can be spherical, it can be exponential, it can be circular, it can take Bessel and all other forms and these are basically fitted based on least square estimates just like we fit you know model linear regression model using least square estimate. So, this is basically Variogram. So, range is the maximum distance at which data are correlated.

So, up to this maximum distance this data will be correlated and after that data will be independent. Nugget is the distance over which data are absolutely correlated or unsampled and then, Sill is the maximum variance of the data pairs. So, now I hope these terms are clear to you.

(Refer Slide Time: 19:19)

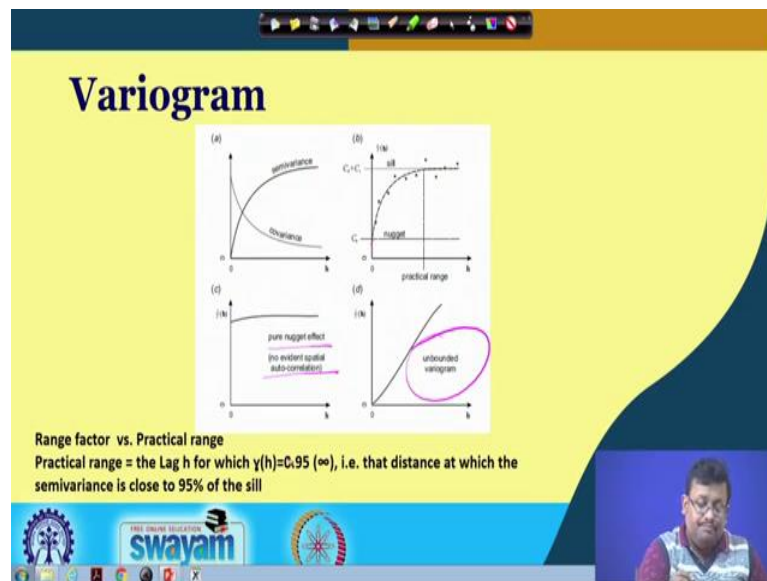


Now, there are different types of variogram like you know linear, spherical, exponential, circular, gaussian, bessel power and so on and so forth. You can see these range up you know this is an exponential model of semi-variogram and obviously, this fitted through you know through least square estimation.

So, what are the assumptions of Kriging? Assumptions of kriging says that the target variable is stationary. That means, right now we are measuring the stationary variable for example we are if we measure the pH that is the stationary variable that is you know and

finally it has a normal distribution which is probably the biggest limitation. So, another you know problem for kriging is that, it assume the data comes from the normal distributions. So, if the data is not normally distributed, we have to make some transformations, some logarithmic transformation to make it first normal and then, we use the data for you know Kriging Interpolation.

(Refer Slide Time: 20:19)

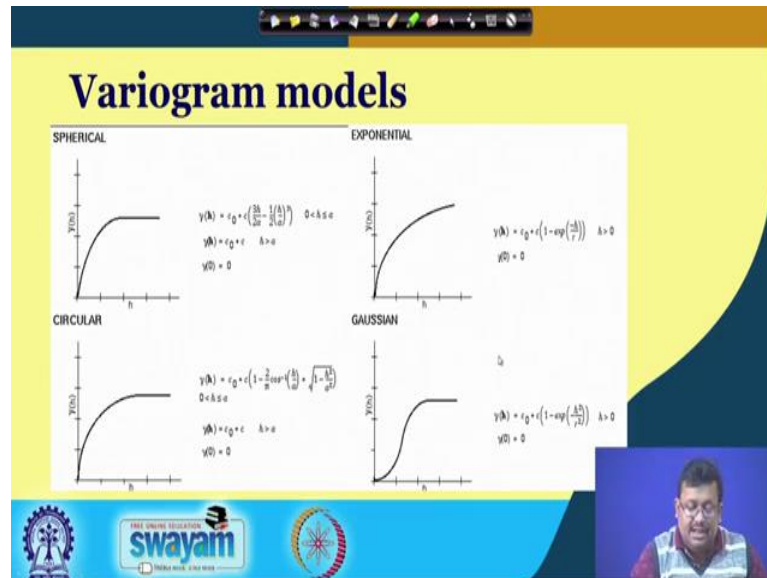


So, these are some drawbacks for kriging. These are some you know you know relation of kriging you know kriging you know these are some relation between the variogram and how this variance semi-variance and covariance you know relate to each other. So, you can see here this semi-variance and covariance are showing the opposite ends as the semi-variance is increasing along with the increasing distance of the data pairs or lag distance, thus covariance is decreasing.

Obviously the covariance shows how these two variable changes along with you know how covariance is a measure of variation between two variables and it obviously decreases when the distance between these points are increasing and this is again the nugget, this is another sill, this is another this is the nugget effect. This c_0 and c_0 plus c_1 is the sill and this is the practical range this is called the pure nugget effect. That means, we cannot see any sill. Remember that the higher the sill, that means higher the spatial dependence however here we are getting maximum amount of nugget effect, we are not getting the sill. So, sill is very less as compared to nugget.

So, we are getting pure nugget effect. So, no evident of spatial autocorrelation we are getting here. This is an example of unbounded variogram.

(Refer Slide Time: 21:50)



So, these are some important terms. There are several models for variograms. You can see spherical model, exponential model, circular model, Gaussian models and these are there in mathematical functions.

(Refer Slide Time: 22:01)

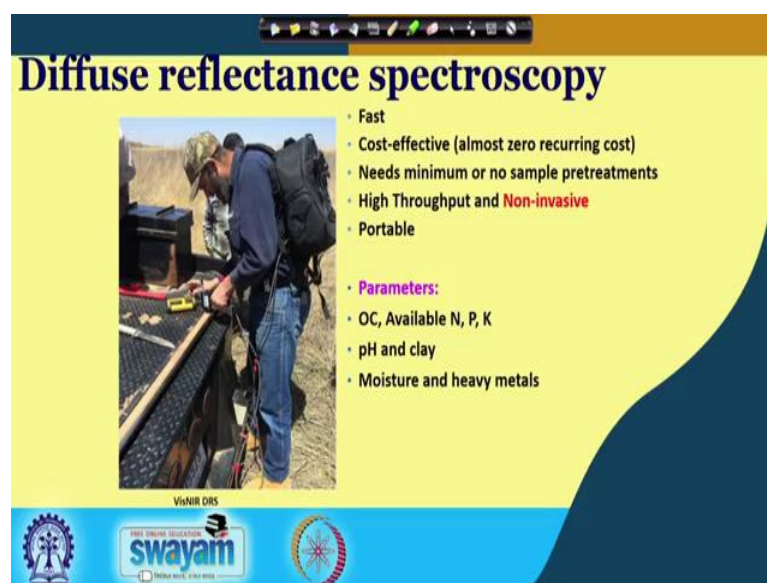


And these are this is actually kriging outputs. So, based on this model estimation based on the spatial dependence which is modelled through this semi-variogram or variogram,

kriging basically interpolates the value in some unsampled locations and you can see using different types of kriging, we are producing these you know interpolated maps.

This universal kriging map, this is an exponential kriging map, this is circular kriging map you can see here and these are clay points. So, guys we have finished this geostatistics.

(Refer Slide Time: 22:40)



Diffuse reflectance spectroscopy

- Fast
- Cost-effective (almost zero recurring cost)
- Needs minimum or no sample pretreatments
- High Throughput and **Non-invasive**
- Portable

Parameters:

- OC, Available N, P, K
- pH and clay
- Moisture and heavy metals

VisiIR DRS

swayam

And let us move ahead and start a new topic that is the basic of diffuse reflectance spectroscopy. Now diffuse reflectance spectroscopy has become a very very you know important tool now a days in the hand of soil science and in the hand of soil scientist because now a days we are using this for extensive analysis of different soil features, both the soil physical properties as well as soil you know chemical properties as well as soil biological properties and we will see that later on.

Now, why we require diffuse reflectance spectroscopy? Now you know that we use different types of you know standard methods in our laboratory for measurement of different soil properties starting from you know for pH, we use the pH meter for EC, we use the EC meter, Electrical Conductivity meter for measurement of organic carbon, we use Standard Walkley Black methods, sometime we use the Loss on Ignition method also or in advance cases we use this CHNS analyzer, however all though these methods are very much accurate, these are time consuming and then costly and also produces some caustic wastes to the environment.

So, we require some alternative for you know we require some alternative tools for measurement of this particular soil properties and diffuse reflectance spectroscopy is one of these tools which can offer cost effective, rapid and non-destructive measurement of soil analysis or soil parameters. So, what are the advantages of this technique? First of all first of all it is very fast as you can see here. This technique is very very fast. You can literally take a reading you know you can take a you know you can literally take a scan of a soil sample using only you know 5 to 10 second.

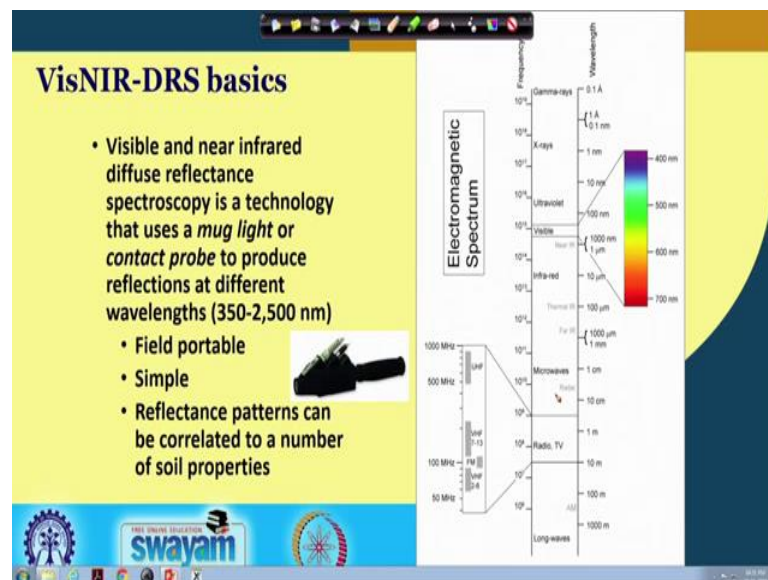
So, you can see this is the diffuse reflectance spectrometer. You can literally carry it in your backpack with a handheld probe and using this handheld probe, you can take the reading at any desired points of the soil. Now, you see they have isolated a core from the soil surface and from the soil and they are just measuring using this handheld probe and these handheld probe will transmit the results into through a Bluetooth to this computer or you know PDA and this PDA will store all the results here. So, what are the benefits of using this diffuse reflectance spectroscopy? First of all this is fast and you can literally take the scan of a particular soil sample within 5 to 10 seconds.

It is cost effective. That means, almost zero recurring cost. It runs through battery. So, you require only the battery to charge and you can use it, you do not require any consumable. It does not require any other you know caustic chemicals and it is high throughput and non-invasive. This is very important. High throughput means you can literally take the scan of a soil sample and you can literally save a spectra of a particular sample and that spectra can be used for predicting soil properties or predicting a multiple soil properties. So, from a single spectra, you can predict a multiple of soil properties and that is why this is called high throughput.

And also it is called non-invasive because it is not destructing the soil. It can you can use this soil sample later on for some further analysis also. So, that is called non-invasive and you can see it is portable in nature. So, due to these all these things all these you know plus points we are now using this spectroscopic sensor or diffuse reflectance spectroscopy sensor extensively in the soil science discipline. What are the parameters we can measure? These are some of the parameters we can measure like organic carbon available, nitrogen, phosphate, potassium, then pH and clay and you know other sand silt and then, moisture and heavy metals.

So, you can see range of parameters you can measure through this property. This you know diffuse reflectance spectro radiometer and also remember that, these diffuse reflectance spectroradiometer can also measure different microbiological properties, different biological properties, different physical properties of the soil also.

(Refer Slide Time: 27:26)



So, why we call it visible to near infrared diffuse reflectance spectroscopy or in a short form we call it VisNIR-DRS.

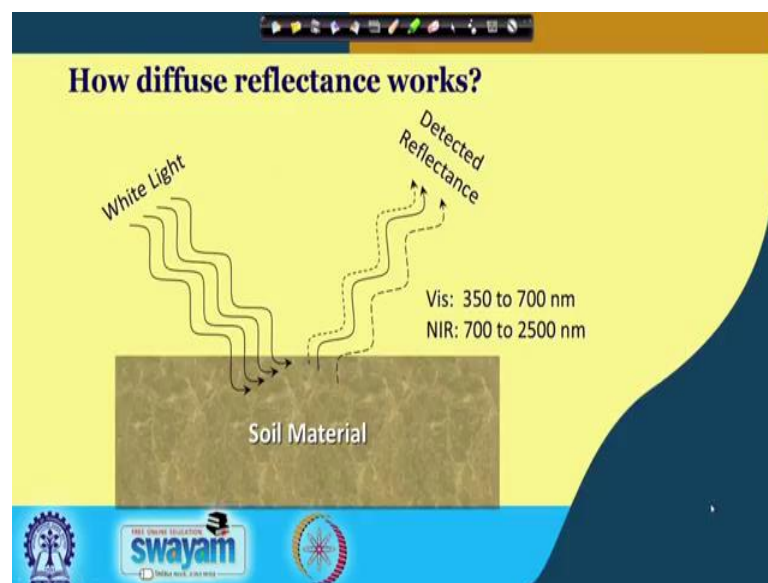
So, remember I hope that you know what is the Electromagnetic Spectrum? Now, the electromagnetic spectrum basically consist of several regions starting from gamma ray, then x-ray, then ultra violet, then visible infrared, microwaves, radio waves and long waves. So, the gamma rays and x rays are you know highly penetrating rays because they have high energy. That means, they have high frequency; however as we go down we will see visible range which is varying from 400 to 700 nanometer and then we see the near infrared range which basically varies from 700 to 2500 nanometer roughly and after that there is a mid infrared and far infrared and all these so on and so forth.

So, these visible and near infrared diffuse reflectance spectroscopy is a technique that uses mug light or contact probe to produce reflection at different wavelengths. So, basically what happens these instrument as you see in the last slide. So, let me go back to the last slide, so that it become you know more clear. So, if you go to the last slide you see it is an handheld probe.

So, these handheld probe has a light source, halogen light source and it has got it is also connected you know through a fiber optic cable to the detector which is presented inside this inside this main spectroradiometer. So, these fiber optic cable or you know these you know halogen light source basically shines a light or produce a white light and these light basically get reflected from soil surface and the reflection is basically detected in the 350 to 2500 nanometer range.

So, that is why it is called visible to near infrared diffuse reflectance spectroscopy. So, basically it is called visible to near infrared diffuse reflectance spectroscopy because it can measure the reflections from 350 to 2500 nanometer range which roughly covers the visible and as well as the near infrared range, ok. So, basically it covers from this visible to near infrared range.

(Refer Slide Time: 30:06)



So, if we go ahead and see how this diffuse reflectance works.

So, as you know that soil material has a rough surface and when there is an white light hits the any material, some amount is getting reflected and this detected you know and this reflected radiation can be detected by these detectors, which is present in this VisNIR-DRS and it can detect from 350 to 2500 nanometer. So, that is why we are calling it VisNIR-DRS.

(Refer Slide Time: 30:38)



Now, the question comes why it is called Diffuse Reflectance. Now there are two types of reflection generally occurs. First of all the first of all it is called Specular Reflection; another is called Diffuse Reflection. Now the specular reflection here generally occurs for a very plane surface like mirror surface or water surface, however the diffuse reflection always occurs from very rough surface.

Now, in the specular reflection always occurs when there is an you know the incidence, the angle of incidence is always equal to the angle of reflection whereas, in case of diffuse reflection since it is occurring from a rough surface like you know like any powder surface like soil, the angle of incidence is not equal to the angle of reflection. So, this is the difference between you know you know specular reflection and diffuse reflection. Remember that the specular reflection occurs from vary plane surface like water or mirror surface and in case of diffuse reflection, it generally occurs from fine particle and powder and rough surfaces likes soil and coal.

So, guys let us stop here and we will continue from here. We will continue from here and we will talk about different applications and their basics of diffuse reflectance spectroscopy in the next lecture and we will finish it in the next lecture.

Thank you very much.