

Resource estimation

ENSG- Marie-Cécile FEBVEY

14/01/2026-15/01/2026



orano

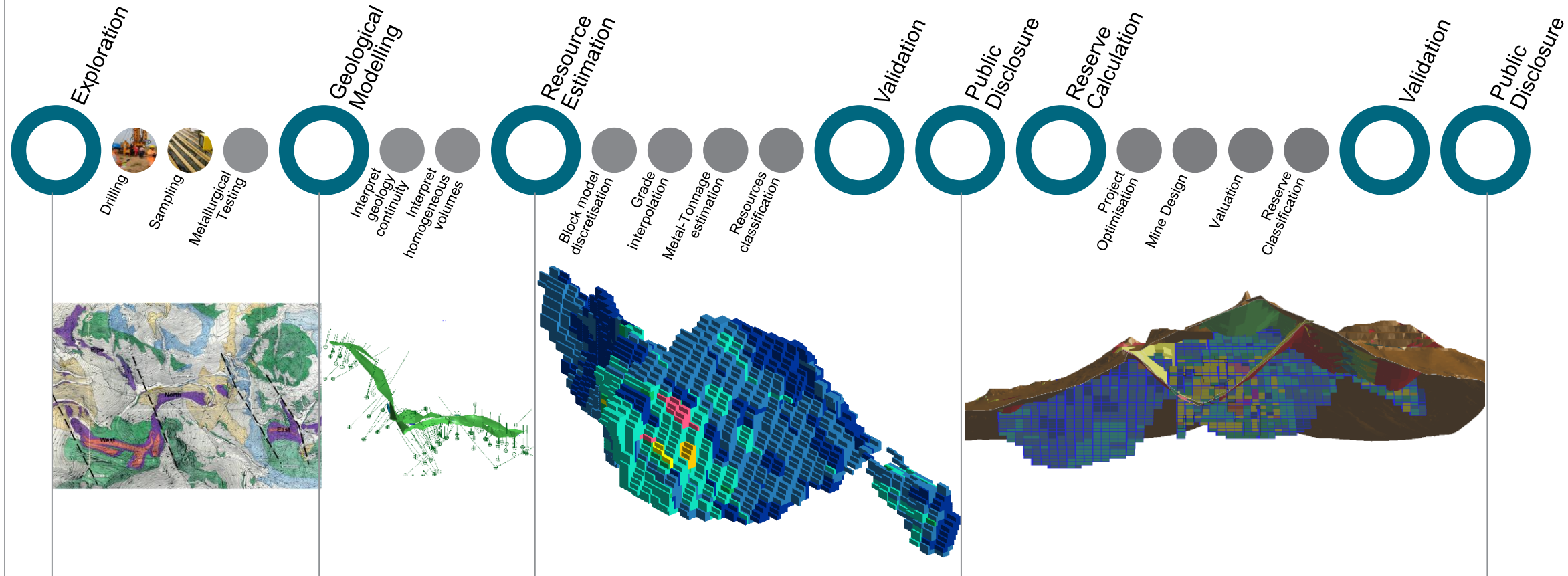
Summary

1. Basic motivations for geostatistics
2. Data preparation
3. Stationarity
4. Necessity of the experimental variogram
5. Variogram model: interpreting spatial continuity
6. Ordinary Kriging
7. Block interpolation: support effect
8. Block interpolation: Block kriging
9. Neighborhood parameters
10. Estimation validation
11. Multivariate Geostatistics
12. Recoverable resources
13. Conclusions

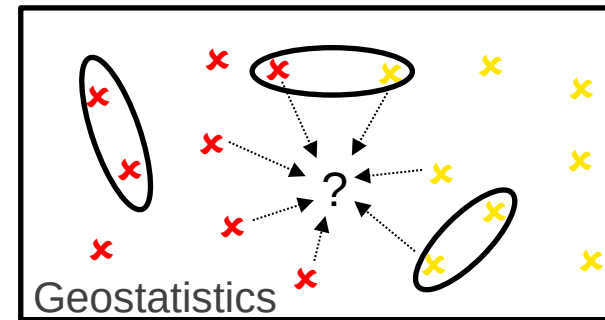
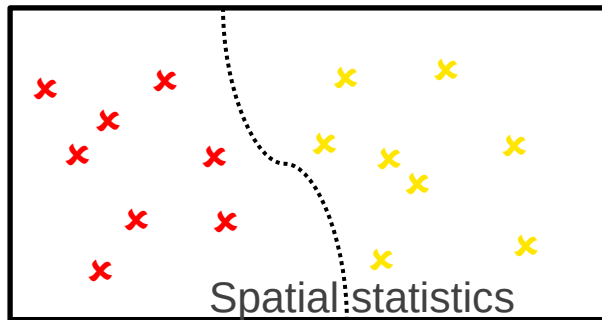
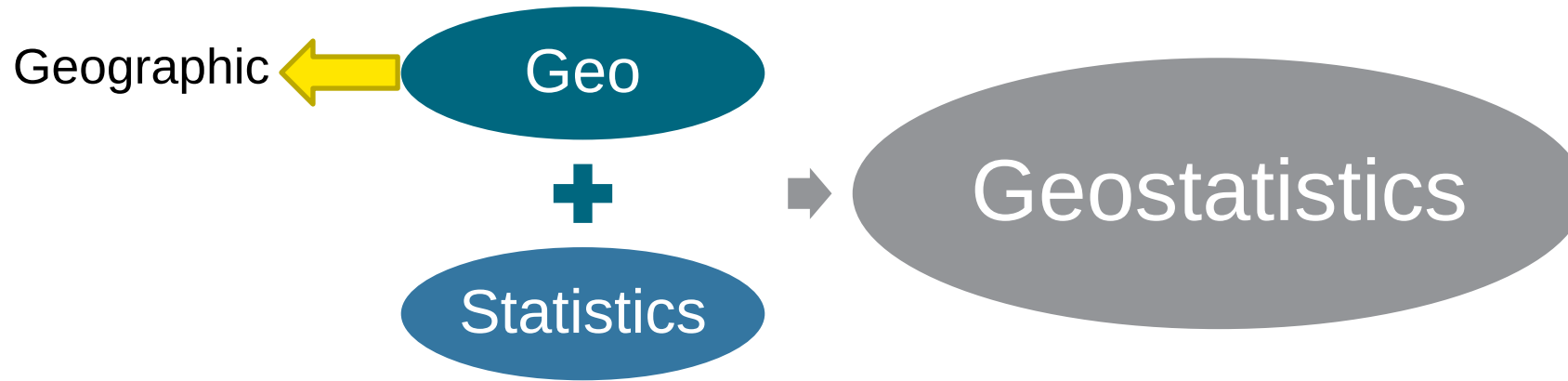
01

Basic motivations for geostatistics

Walk through the Global estimation Process



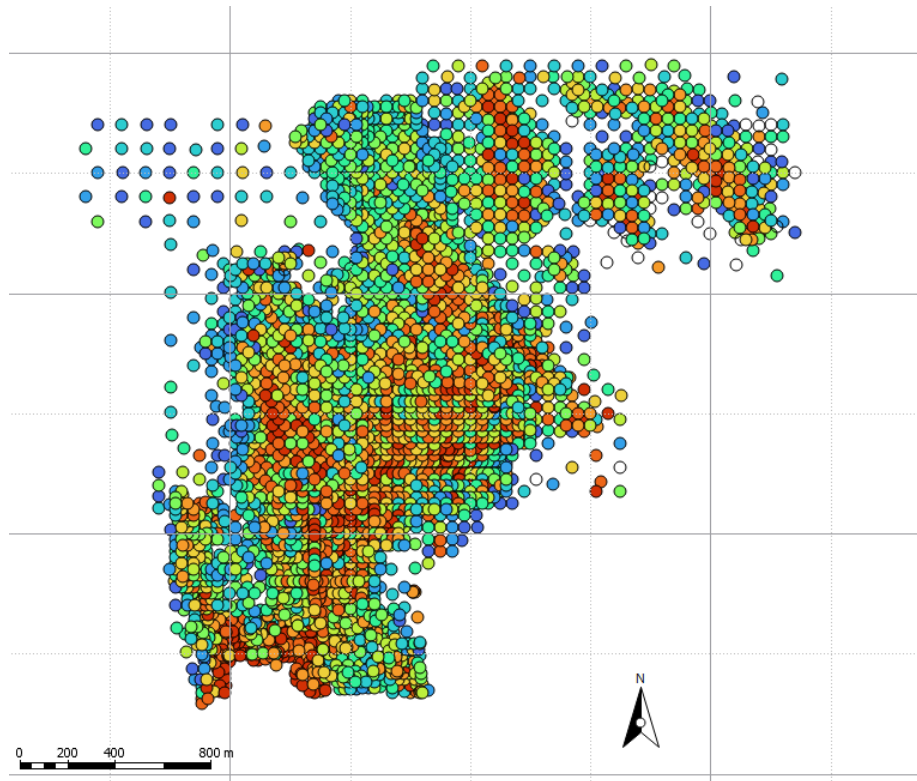
What is geostatistics?



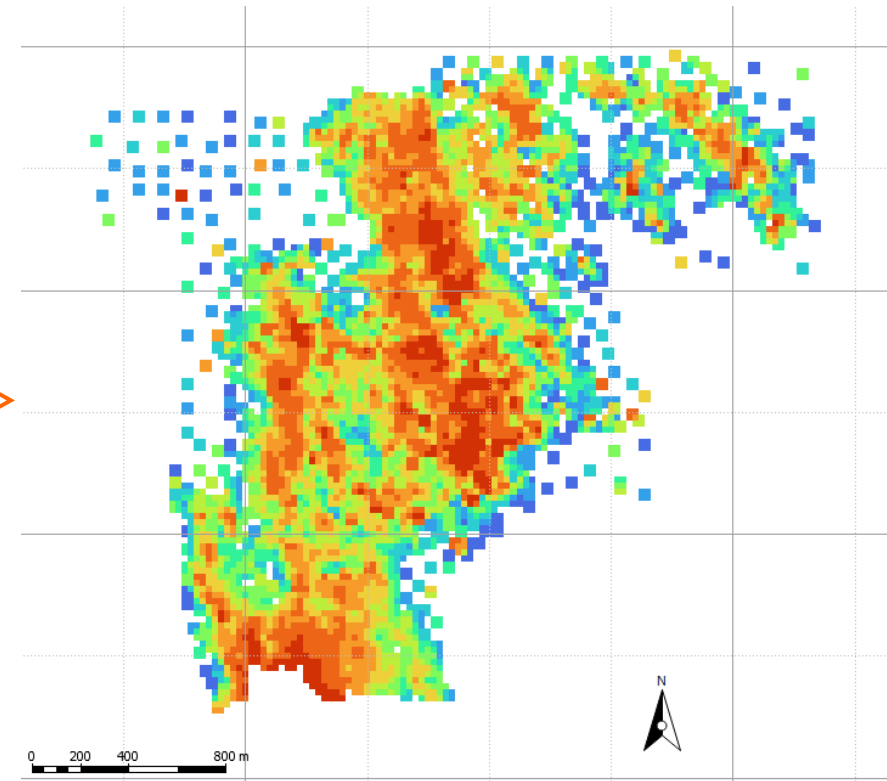
The first definition of geostatistics:

Statistical Processing of Spatial Data

Global purpose: Estimation



Data



Interpolation

Classical estimation: Risk of bias

735	45	125	167
450	337	95	245
124	430	230	460
75	20	32	20

- 16 samples – grade displayed (ppm)
- Cut off grade: 300ppm
- Estimate blocks 1x1
- Comparison of estimation Kriging v/s polygons
- Conventional benefit (profit)

$$B = \sum_{\text{selected blocks}} Q - z_c T$$

Q = metal quantity = grade*T
T = ore tonnage = 1T per block
Zc = cutoff grade

True block grades

505	143	81	207
270	328	171	411
102	220	154	263
101	54	44	155

1. Select the blocks which are above the 300ppm cut off grade
2. Calculate the conventional benefit

- **Ideal benefit** (cut-off of 300ppm) = **344**

$$505-300+328-300+411-300=344$$

- **Ideal benefit: maximum we can get by applying the cutoff on the true unknown block grade**

$$B = \sum_{\text{selected blocks}} Q - z_c T$$

Polygonal (nearest neighbour) estimation

735 505	45 143	125 81	167 207
450 270	337 328	95 171	245 411
124 102	430 220	230 154	460 263
75 101	20 54	32 44	20 155

- Nearest neighbor Grade of block = grade of sample in block
- Actual benefit = 86
- Estimated benefit = 912

1. Select the blocks which are above the 300ppm cut off grade
2. Calculate the conventional benefit

$$B = \sum_{\text{selected blocks}} Q - z_c T$$

(in dark blue: nearest neighbor grade- true grade in grey)

Ordinary Kriging estimation

442 505	190 143	142 81	204 207
354 270	276 328	212 171	279 411
189 102	226 220	216 154	271 263
99 101	81 54	88 44	125 155

- **Estimated benefit =**
442-300 + 354-300 = 196

- **Actual benefit =**
505-300 + 270-300 = 175

1. Select the blocks which are above the 300ppm cut off grade
2. Calculate the conventional benefit

$$B = \sum_{\text{selected blocks}} Q - z_c T$$

(in blue: nearest neighbor grade- true grade in grey)

Comparing the results

Conventional benefit:

$$B = \sum_{\text{selected blocks}} Q - z_c T$$

Q = metal quantity

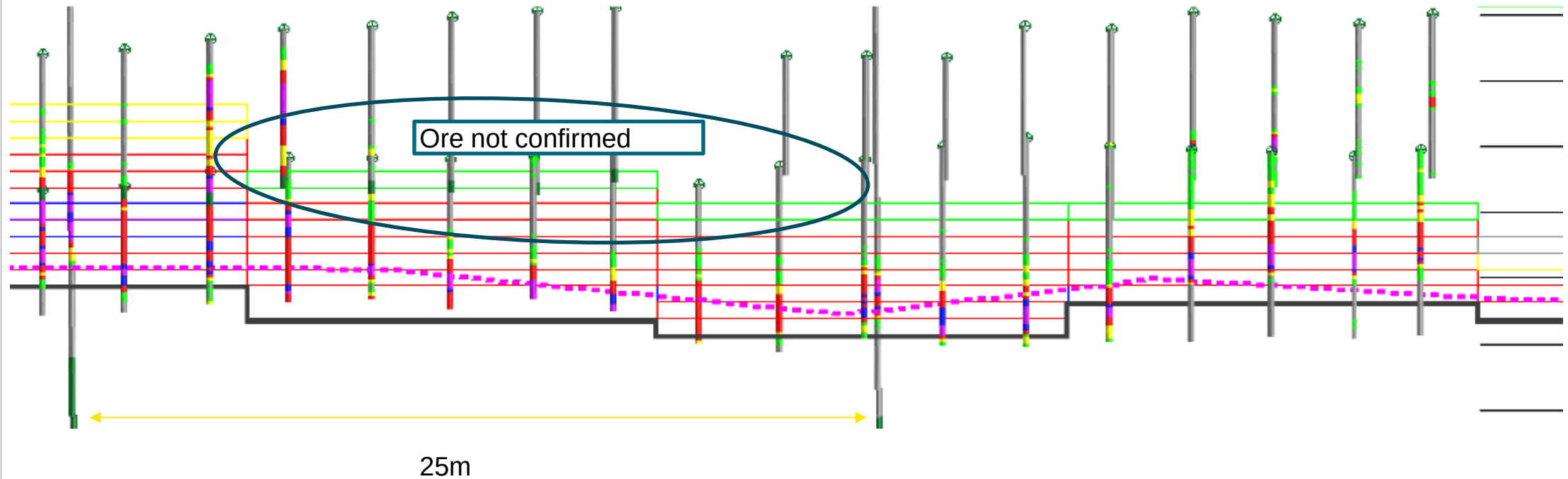
T = ore tonnage

Zc = cutoff grade

Predicted benefit:

	Predicted benefit	Actual benefit
Polygonal Estimate	912	86
Kriging	196	175
Ideal	344	344

Mining operations: specificities



Modelling and estimations at different “scales”, with different datasets (exploration, development, grade control)

Indirect measurements until mining the deposit

02

Data preparation

Why is it so important to validate the data?

- Usually modelling would take between 1 to 3 months to be completed, sometimes more
 - Estimation will take between 4 to 6 months, including reporting
- > Time consuming process, data is the first entry of all the model and estimations



Resource Model

Chemical analysis & Geological interpretation

Sample preparation

Sampling & geological observation

“Resource estimation is a « house of cards », the foundation of which is sampling and geological observation, with a first floor of sample preparation, a second floor of chemical analysis and geological interpretation, and a top floor consisting of geostatistics and the resource model.

”

Long, Parker, & François-Bongarçon

Area of problem	Frequency (%)
Geology, Resource and Reserve estimation	17
Geotechnical analysis	9
Mine design and scheduling	32
Mining equipment selection	4
Metallurgical test work, sampling and scale-up	15
Process plant equipment design and selection	12
Cost estimation	7
Hydrology	4

Source: Monograph 30. Managing Risk in Feasibility Studies.
PL McCarthy 2013

Data preparation

Always validated the data prior to run an estimation

Check the data per drilling campaign

Validation:

- **2D and 3D validation**
- **Topography and drillholes Z**
- **Coordinates**
- **Duplicates samples or drillholes**
- **Analysis validation**
- **QAQC etc**

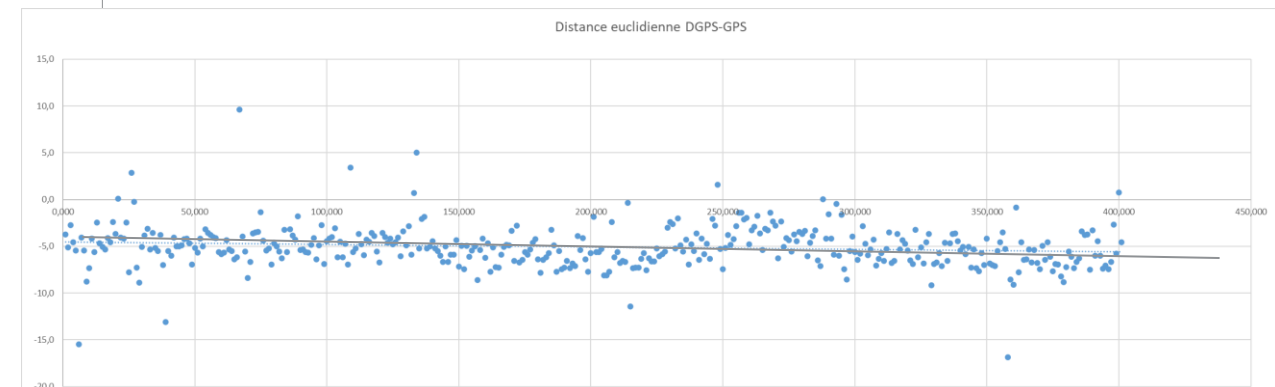
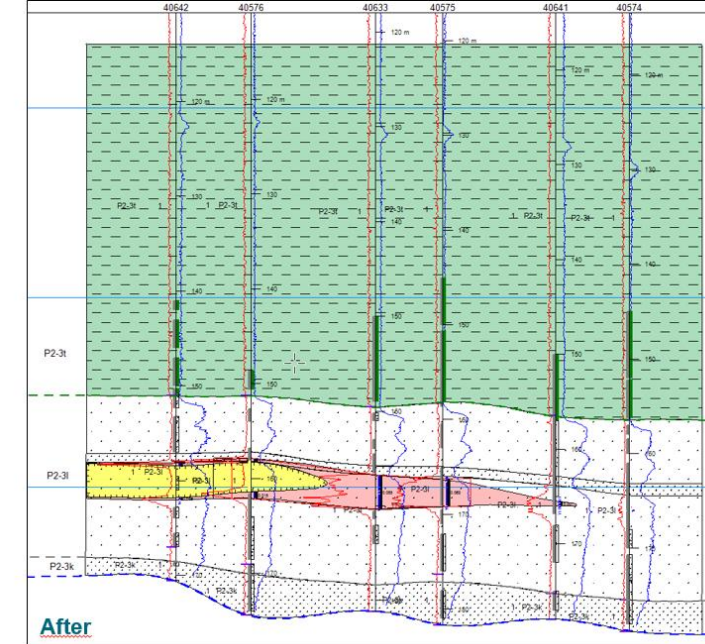
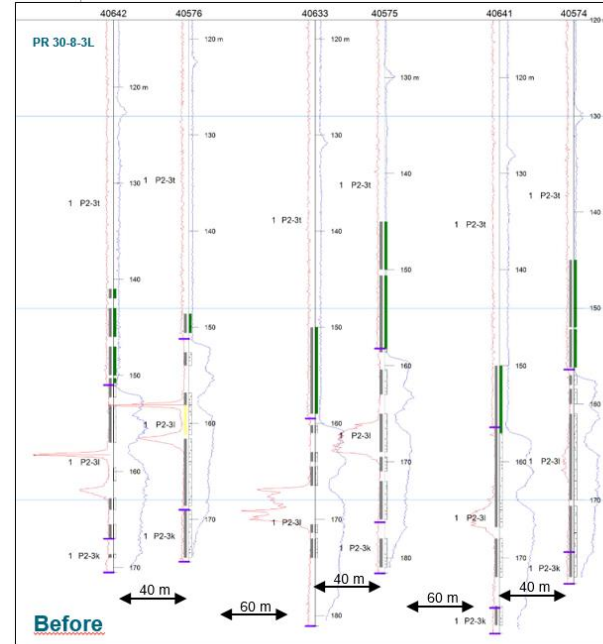
Drillholes data validation: Data preparation

Always validated the data prior to run an estimation

Check the data per drilling campaign

Validation:

- 2D and 3D validation
- Topography and drillholes Z
- Coordinates with coordinates system
- Duplicates samples or drillholes
- Analysis validation
- QAQC etc



Drillholes data validation: Examples of errors

Overlapping intervals

From >= To

Errors (6678)		Current Error (in table d_LithoGIS) Showing holeid 'MS_KD_4075_1' only					
Overlapping Segments (6678)							
d_LithoGIS (6678)							
Warnings (844)							
No samples for collar (844)							
Invalid Values Handling (0)							
No numeric columns							
Ignored	id	holeid	from	to	Lithofacie...	Description	
☐	1	MS_KD_4075_1	387.0	389.5	1	GIS Clay	
☐	2	MS_KD_4075_1	389.5	391.9	6	GIS Clay'ey Sand & unrecognised	
☐	3	MS_KD_4075_1	391.9	404.1	1	GIS Clay	
☐	4	MS_KD_4075_1	404.1	404.5	5	GIS Sandstone	
☐	5	MS_KD_4075_1	404.1	404.5	5	GIS Sandstone	
☐	6	MS_KD_4075_1	404.5	408.7	6	GIS Clay'ey Sand & unrecognised	

Errors (292)		Current Error (in table a_Primary_Assay) Showing holeid 'MS_X_08056___' only					
From depth >= to depth (3)							
a_Primary_Assay (3)							
row 954							
row 7270							
row 7651							
Overlapping Segments (289)							
Ignored	id	holeid	from	to	U_XRF_pct	U_XRF_ER...	UR
☐	7650	MS_X_08056___	503.15	503.65	0.002		
☐	7651	MS_X_08056___	504.0	504.0	0.001		
☐	7652	MS_X_08056___	504.2	504.3	0.003		
☐	7653	MS_X_08056___	504.3	504.5	0.0005		
☐	7654	MS_X_08056___	504.5	504.65	0.002		

Depth exceed from
DH max depth 520.8

Errors (6)		Current Error (in table f_Horizon_Atomgeo) Showing holeid 'MSK23_07_07_B_' only				
From depth >= to depth (1)						
f_Horizon_Atomgeo (1)						
row 371						
Collar max-depth exceeded (4)						
f_Horizon_Atomgeo (4)						
row 110, column to, conflict...						
Ignored	id	holeid	from	to	Horizon	
☐	2969	MSK23_0...	4.8	358.0	Скважина	
☐	2970	MSK23_0...	358.1	391.5	Ынтымак	
☐	2971	MSK23_0...	391.6	420.5	Иканский	
☐	2972	MSK23_0...	420.6	472.7	Уюкский10	
☐	2973	MSK23_0...	472.8	521.0	Канжуганский12	

Drillholes data validation: Examples of errors

1st Database

- “0” depth in collar (removed)
- Negative value in survey (exporting error – removed)
- 47 DH have no survey (35 DH have paper passport, 11 no info found)
- Duplicated intervals in lithology for 6 DH (checked and removed)
- 335 eRa depths were exceeded from collar depths.
- Overlapping intervals – 10355 in eRa (exporting errors – corrected)
- Collar max-depth exceeded – 377 holes in Resistivity
- Corrected files demanded for the second review

2nd Database

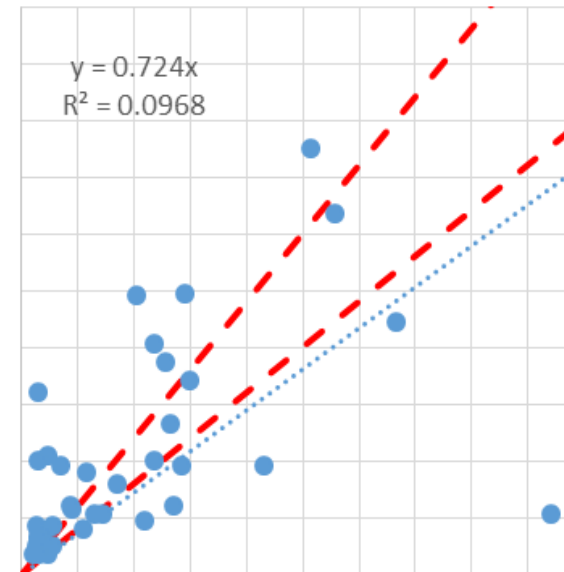
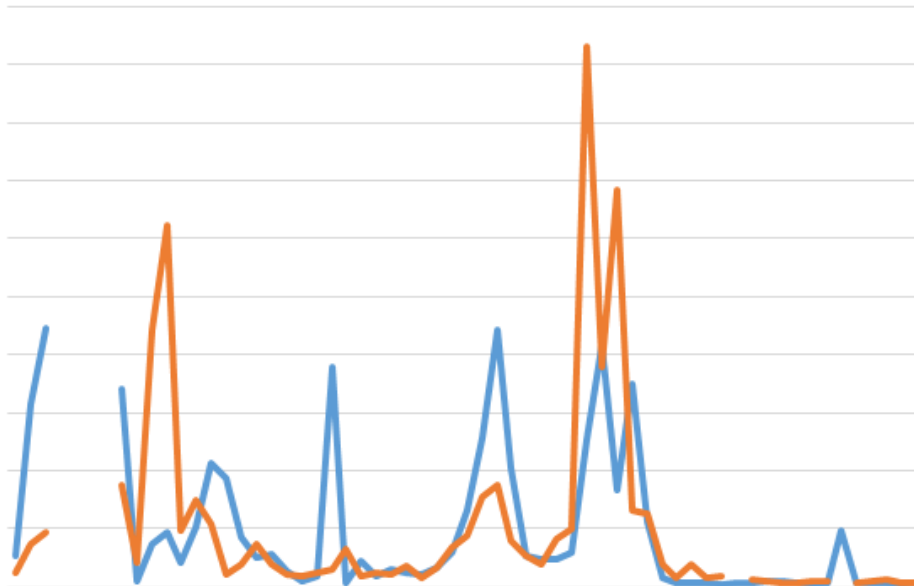
- Overlapping intervals:
 - 788 in construction (export query fixed)
 - 3 in GT (corrected in AtomGeo)
 - 6678 in LithoGIS (export query fixed)
 - 289 in primary assay (not overlaps, Ra analysis date period longer than Uranium)
- Depth = To error – 10 in lithology, 1 in stratigraphy, 14 in assay
- Depth exceeds from actual depths – 4 in lithology, 4 in stratigraphy,
- Missing data checks
- Detail descriptions for codes
- Corrected files demanded for the final review

3rd and 4th Database (1335 DH)

- 3 historical DH was removed proven no available information
- 106 DH have no survey (historical DH corrected as vertical survey, 18 re-drilled or water holes were removed)
- 60 DH construction type was duplicated for intervals (corrected)
- Interval overlaps:
 - 1 in construction (corrected)
 - 3 in dev GT (corrected)
- eRa non-positive values were converted to “0” by Katco.
- From depth $\geq$ to depth error: 1 in stratigraphy (corrected)
- 12 DH belonging to Uyük was excluded.
- 4 missing DH from new GT of historical was included to the database.
- Final corrected database includes 1305 exploration and development DH.

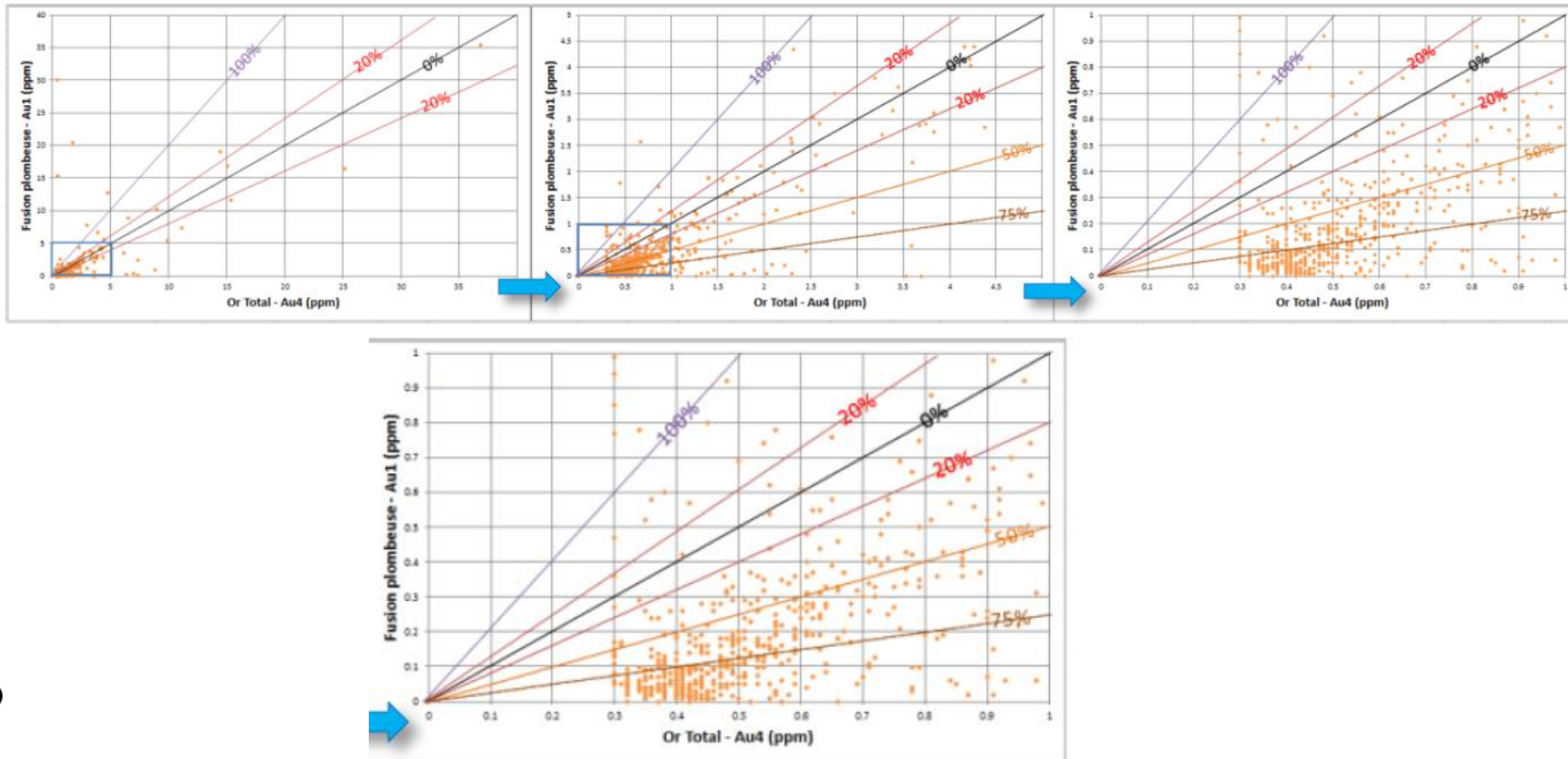
Data preparation

Comparison between diamond drillhole and reverse circulation drillholes (distance smaller than 5m)



Data preparation

Comparison between different analysis types

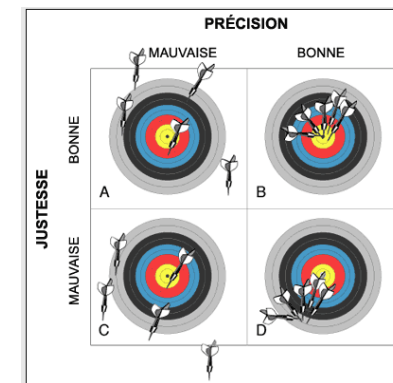
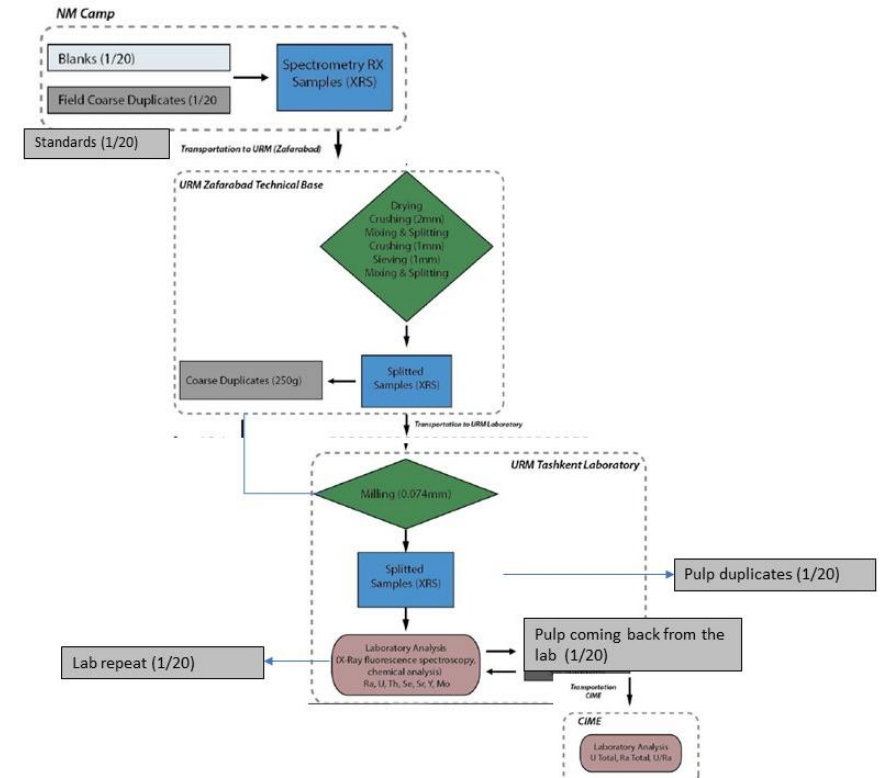


Recommendations for QAQC insertion

Around 20% of QAQC samples should be inserted, meaning:

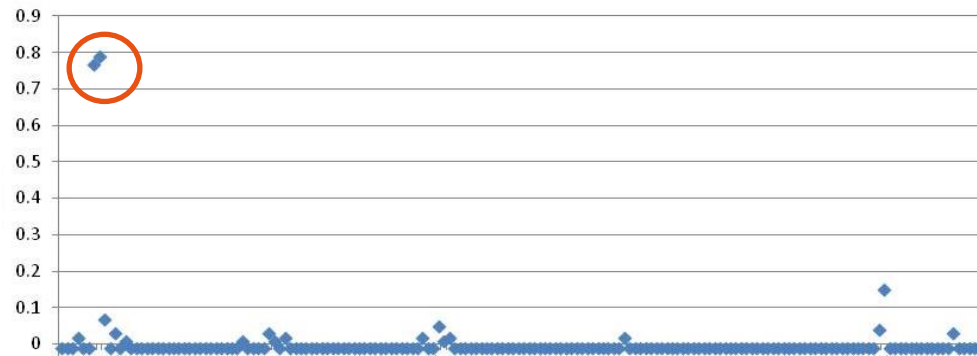
- Blank
- Standard (CRM)
- Duplicates :
 - Field Coarse
 - Coarse duplicate
 - Pulp duplicate
 - Lab repeat
 - Umpire

QAQC ensure that the analysis are precise and accurate

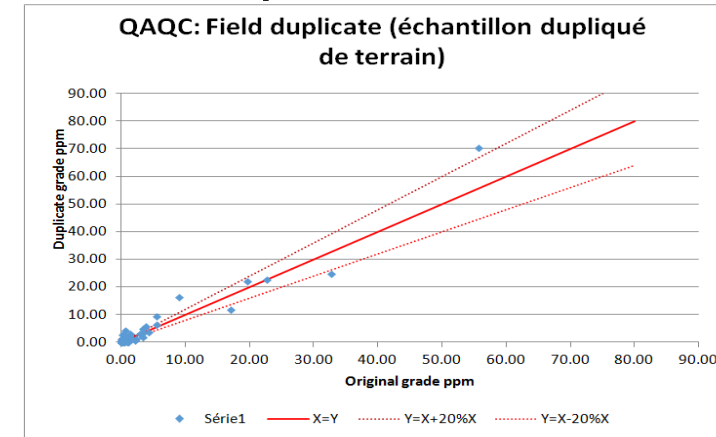


Data preparation : QAQC

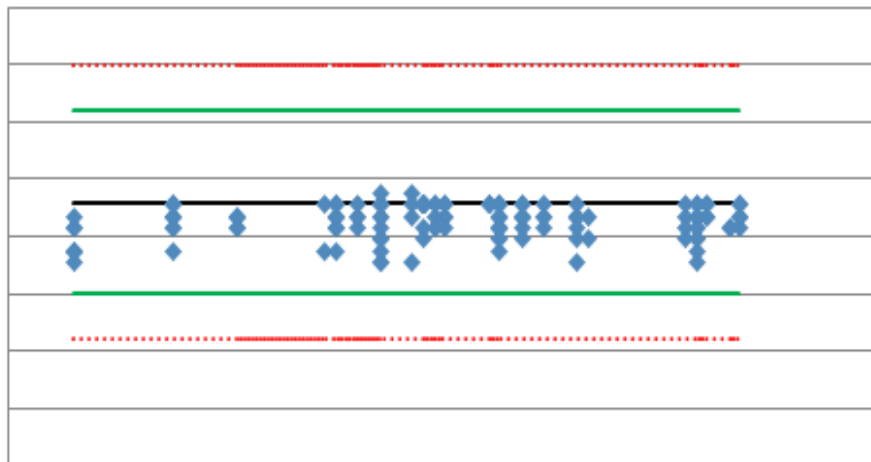
Blanks



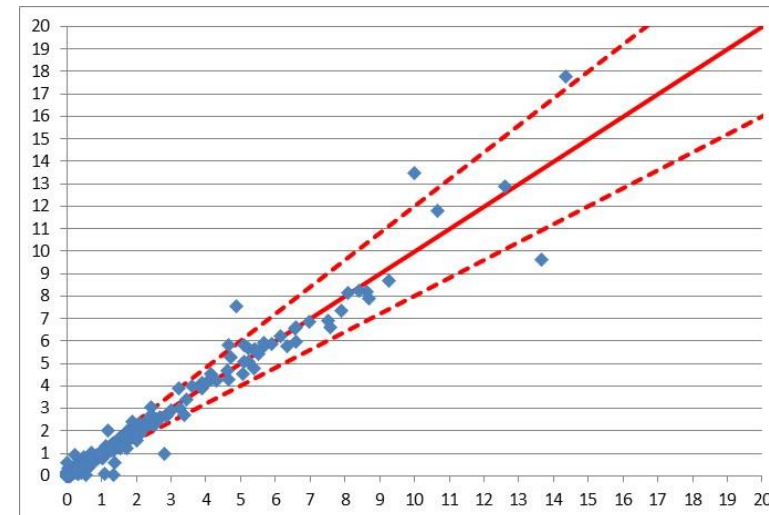
Duplicates



CRM



Umpire



Data preparation

Statistics and geostatistics need good quality data:

1. Additive variables

➤ Data on different supports

Values are not comparable

Samples should be regularized over the same length/volume

2. Representative samples: each sample should have the same probability to be in the dataset as in the entire population

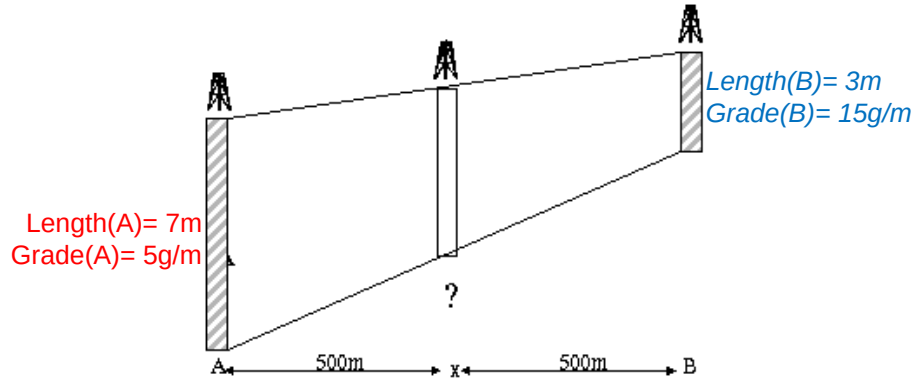
➤ Clustered data

Arithmetic statistics are biased

Irregular sampling patterns should be 'de-clustered'

➤ Issues with the sampling procedure (Fines preferentially lost during drilling or core cutting)

Additivity of variables



- Naïve solution:

$$\text{Grade}(x) = \frac{1}{2} [\text{Grade}(A) + \text{Grade}(B)] = \frac{(5 + 15)}{2} = 10 \text{ g/m}$$

- Length-weighted solution:

$$\text{Accumulation} = \text{Grade} \times \text{Length}$$

$$\text{Accumulation}(A) = 7 \times 5 = 35 \text{ g}$$

$$\text{Accumulation}(B) = 3 \times 15 = 45 \text{ g}$$

$$\text{Length}(x) = \frac{1}{2} [\text{Length}(A) + \text{Length}(B)] = \frac{7 + 3}{2} = 5 \text{ m}$$

$$\text{Accumulation}(x) =$$

$$\frac{1}{2} [\text{Accumulation}(A) + \text{Accumulation}(B)] = \frac{35 + 45}{2} = 40 \text{ g}$$

$$\text{Grade}(x) = \frac{\text{Accumulation}(x)}{\text{Thickness}(x)} = 8 \text{ g/m}$$

Compositing (regularization)

Impact of compositing

Mean is constant

Variance decreases as composite length increases

Support is the term used to denote the volume on which a value is defined

In abstract theory-land, could be an infinitesimal point

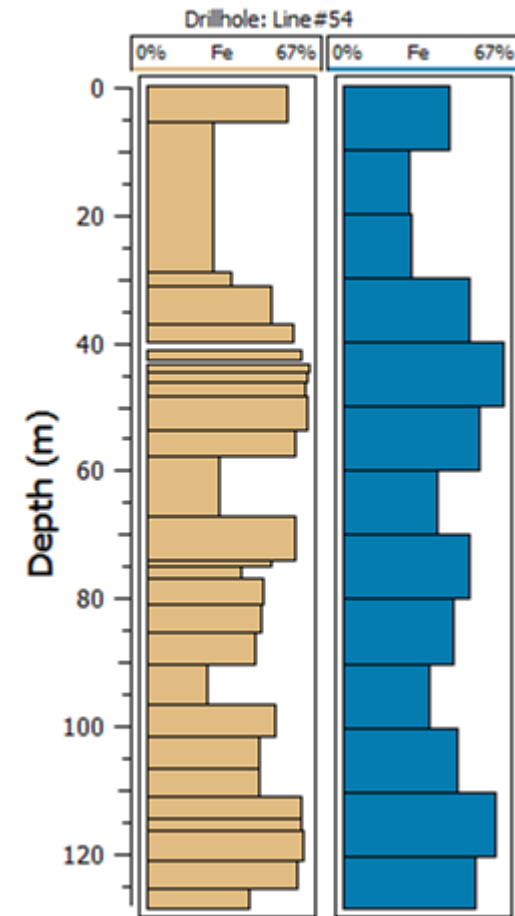
In mining, samples are usually composited to some chosen length

Estimates usually made onto blocks of some chosen size

Support effect

Bigger support: lower variability

Extreme values are averaged



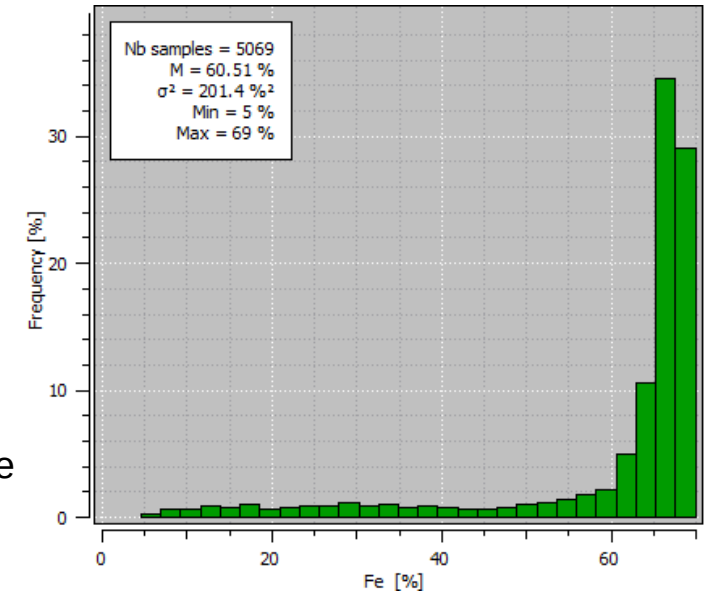
Before and after
compositing

Compositing and statistics

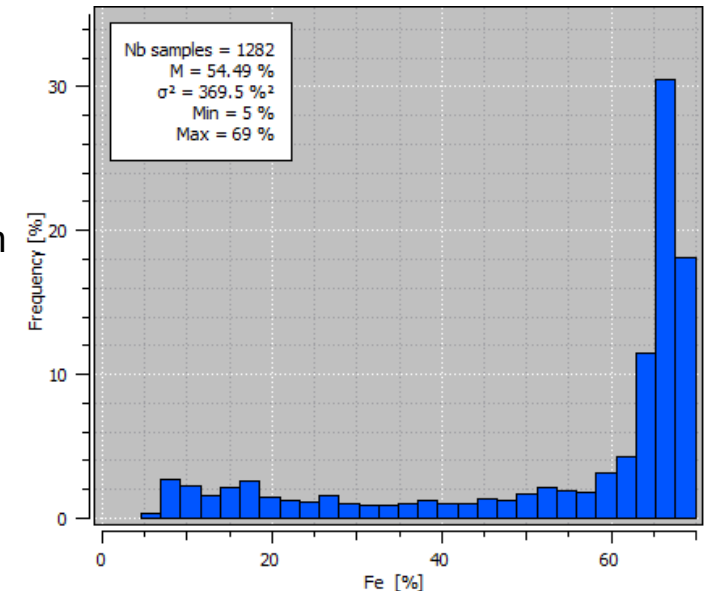
- In the case of a left-skewed histogram, compositing impacts mainly **lower grades**
- Comparing statistics of compositing on 2m, 5m, 10m, 15m

Variable	Count	Minimum	Maximum	Mean	Standard deviation
Fe (no composite)	5069	4.8	69.40	53.97	20.47
Fe (composite on 2m)	9037	4.8	69.30	54.01	20.30
Fe (composite on 5m)	3750	4.8	69.28	54.31	19.94
Fe (composite on 10m)	1910	4.8	69.26	54.40	19.60
Fe (composite on 15m)	1282	4.8	69.18	54.49	19.25

No composite
(raw data)



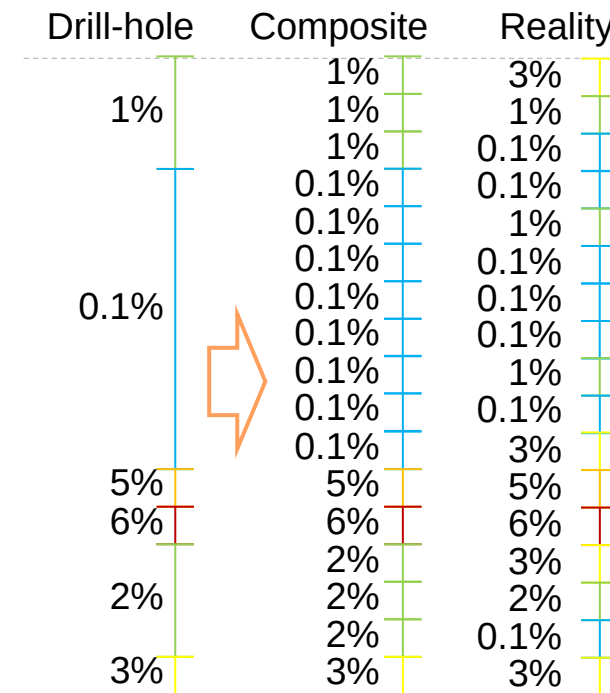
Composite on
15m



Compositing bias

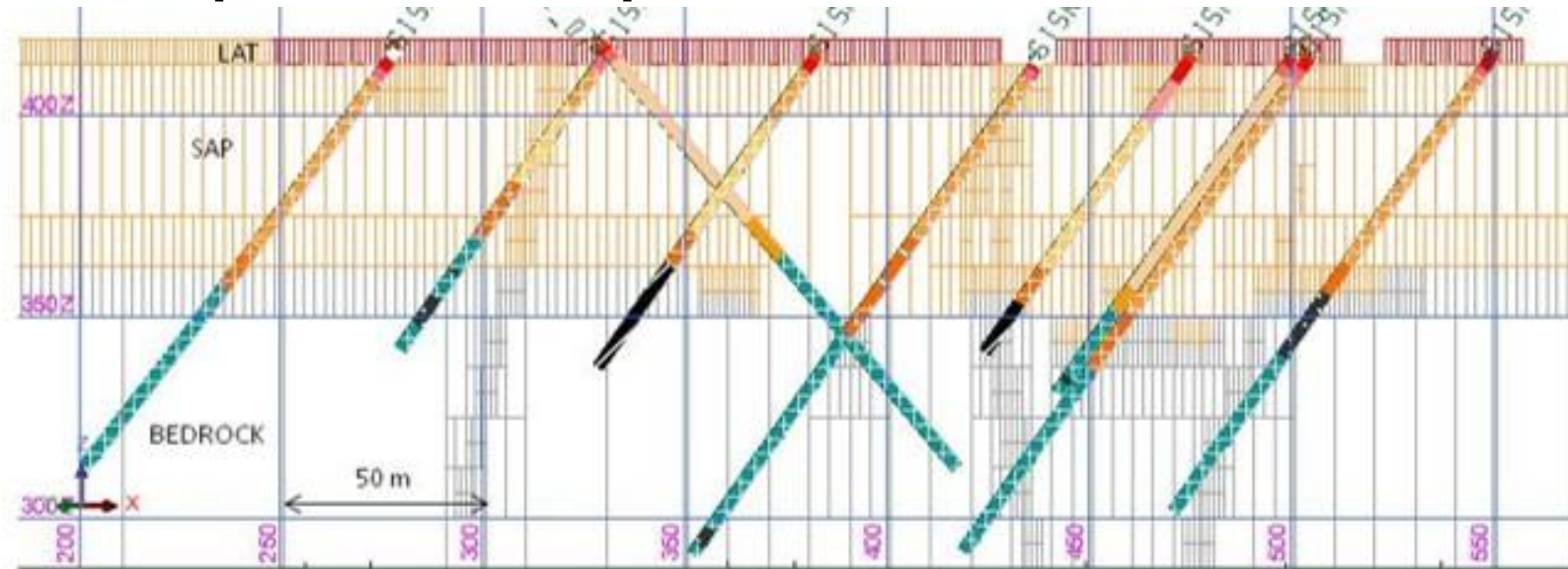
Slicing samples generates inconsistent composite values:

- **False values: not available at this scale**
- **Some extreme values appear several times → Variance may increase**
- **Artificially increases grade continuity**



Clustered samples

- In first stages of mineral exploration, sampling is denser in high-grade locations
- Over-sampling high-grade regions results in biasing the mean: overestimation
- Clustered samples are not representative



Declustering

Objectives:

Unbiasing the average of kriging estimation

Improving the variogram (optional)

Improving the modelling of the distribution through anamorphosis (Simulations / Uniform Conditioning)

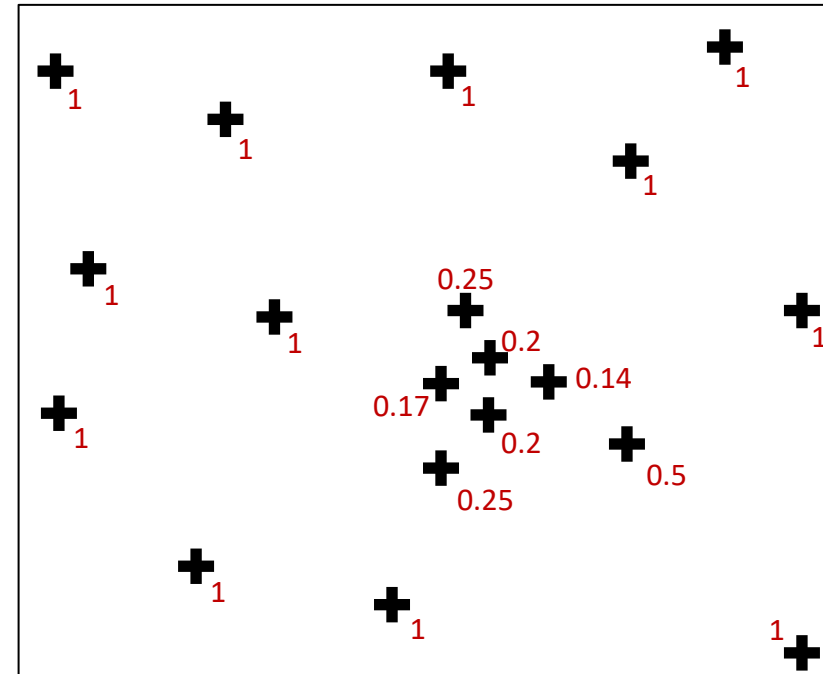
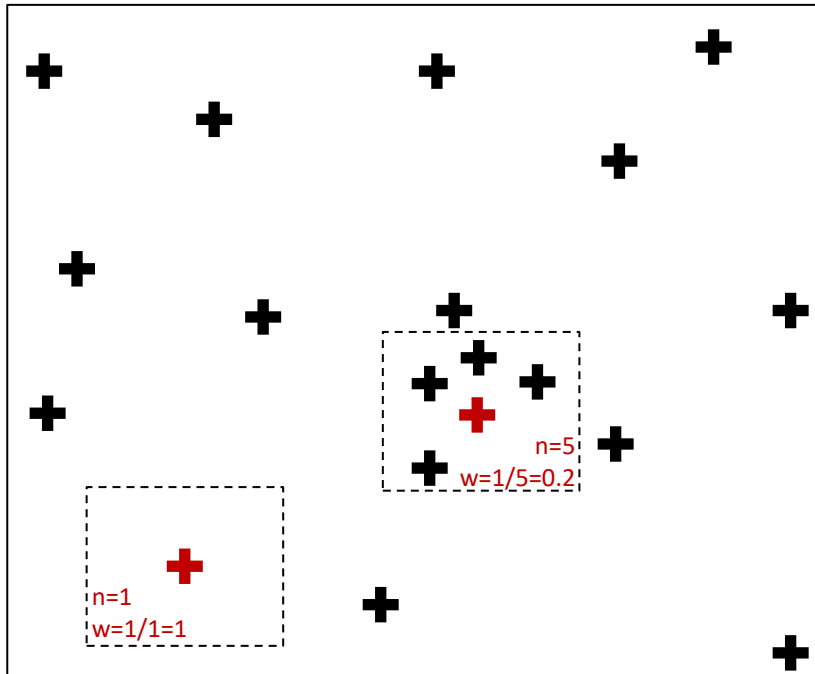
Solution:

A declustering weight is attributed to each sample, which will be used in statistical calculations

The more samples are in a cluster, the lower declustering weight will be attributed

Declustering: moving window

Each sample at the center of a moving window receives the weight of $1/n$

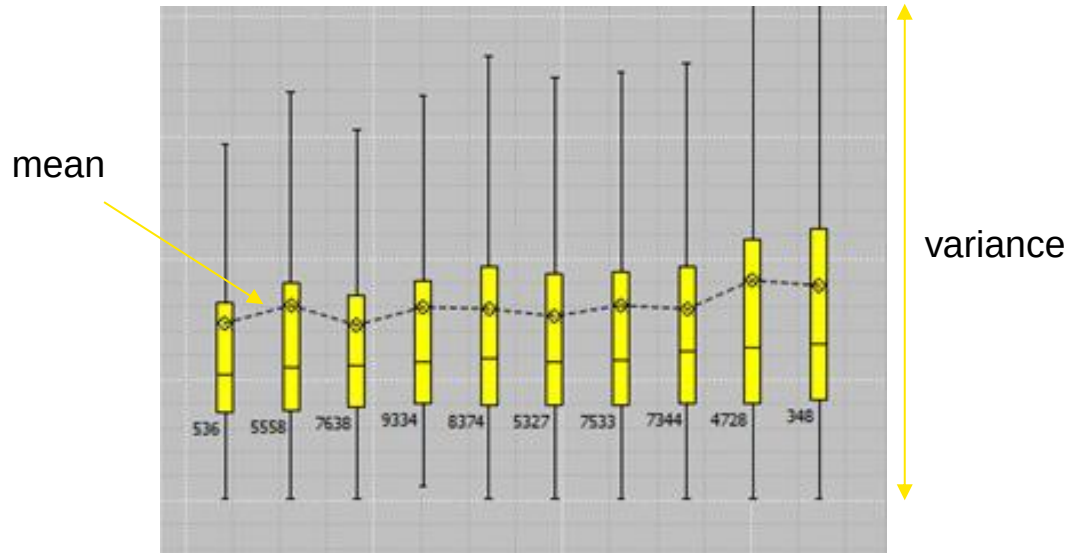


Window

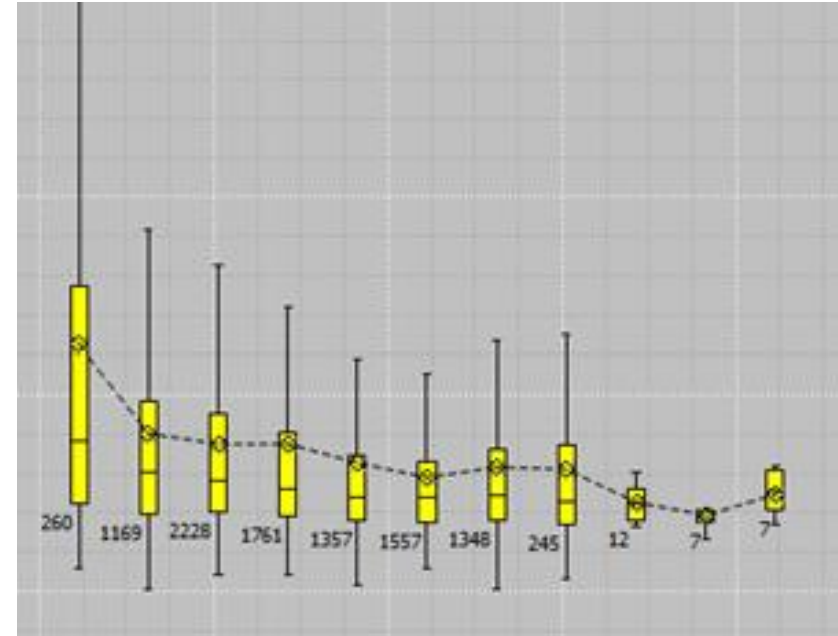
03

Stationarity

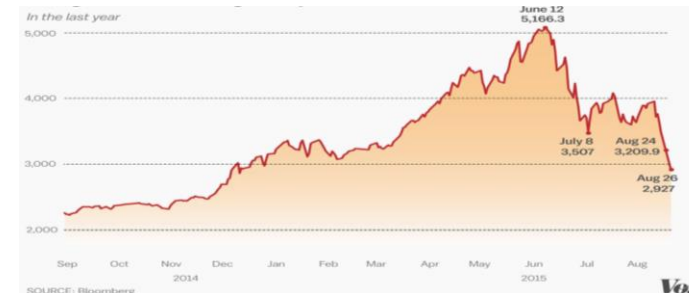
Mean and variance variability in the space



Random function with **constant** mean and variance



Random function with **variable** mean and variance



Stationarity for a regionalized variable

Some characteristics of a **regionalized variable** (average and variance, in practice) are desired to be independent from the position in the space

The difference between two relatively invariant points depends only on the relative position between them, i.e. the distance, and does not depend on their absolute position

Average: $E[Z(x + h) - Z(x)] = 0$ (1st order)

Variance: $Var[Z(x + h) - Z(x)]$ only depends on h (2nd order)

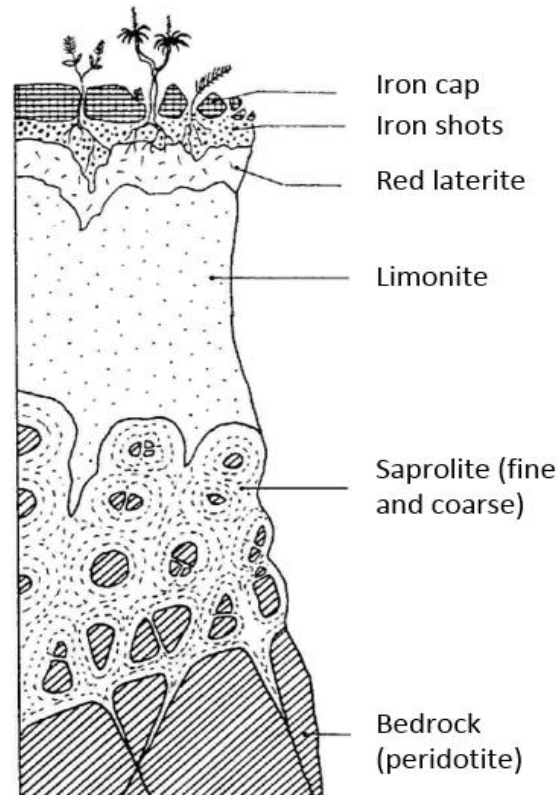
Following this double-hypothesis of stationarity, it is possible that an experimental variogram be estimated using statistic variance of the differences between the data pairs:

$$\gamma^*(h) = \frac{1}{2N(h)} \sum_{i=1}^{N(h)} [z(x_i) - z(x_i + h)]^2$$

Stationarity in domains

Not desired to have abrupt changes in the mean or variance as a function of location:

Solution is domaining (or zonation)



Coarse Saprolite

Nb samples : 6283

Mean : 1.79

Std. Dev. : 0.34

Muddy Saprolite

Nb samples : 8332

Mean : 1.25

Std. Dev. : 0.28

Density

Coarse Saprolite

Nb samples : 6283

Mean : 4.05

Std. Dev. : 1.73

Muddy Saprolite

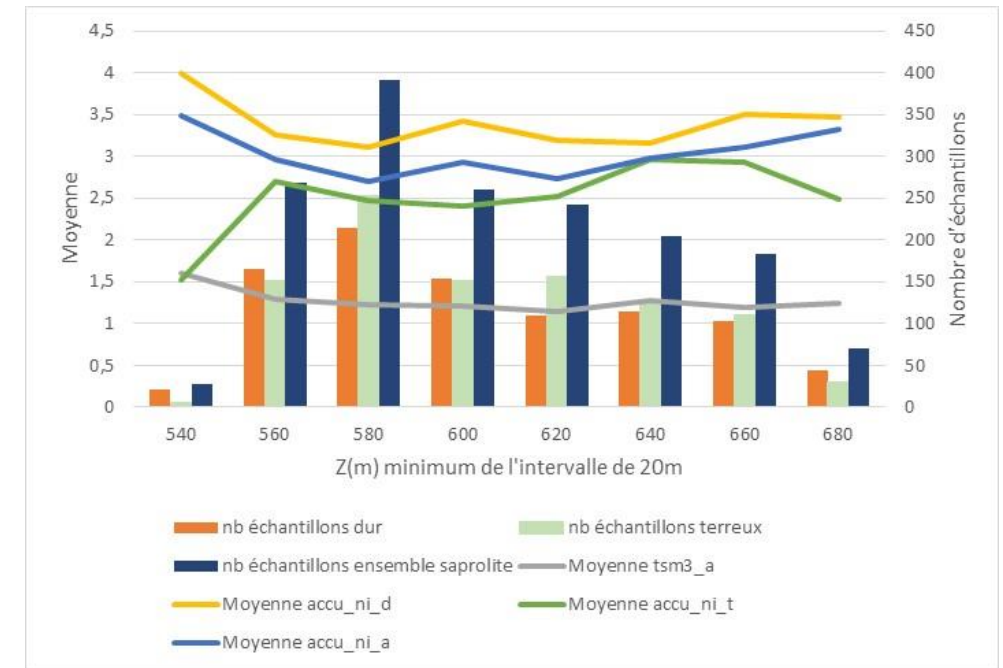
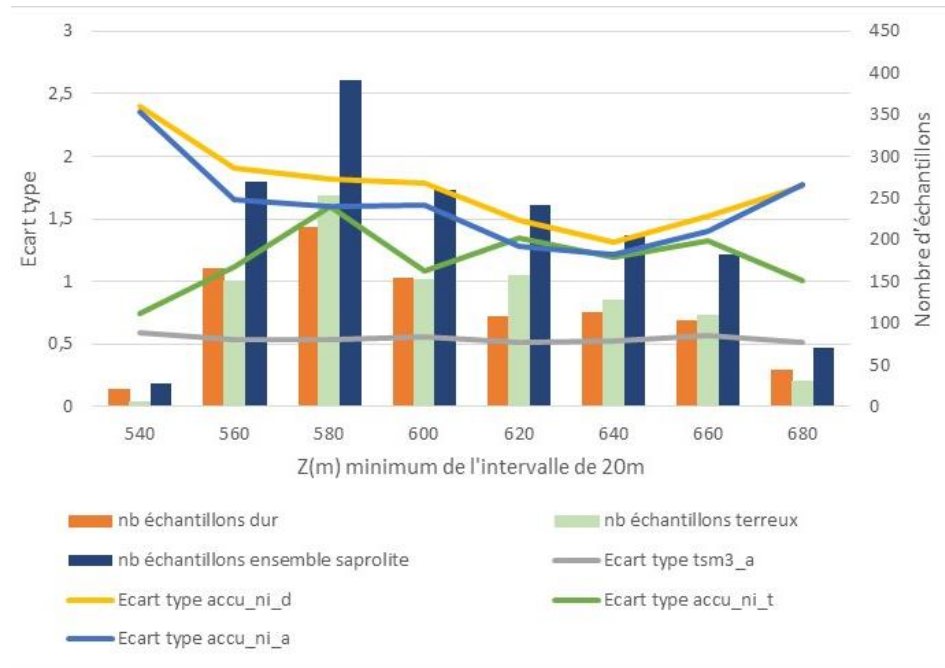
Nb samples: 8332

Mean : 2.7

Std. Dev. : 1.09

Stationarity in practice

Swath plots



What are domains?

Domains are geographical selection of deposits applied to both input data and on grids. It can be the geological modelling but can also take into account other properties such as granulometry, porosity, permeability etc...

They have different (geo)statistical properties:

Mean

Variance

Variogram

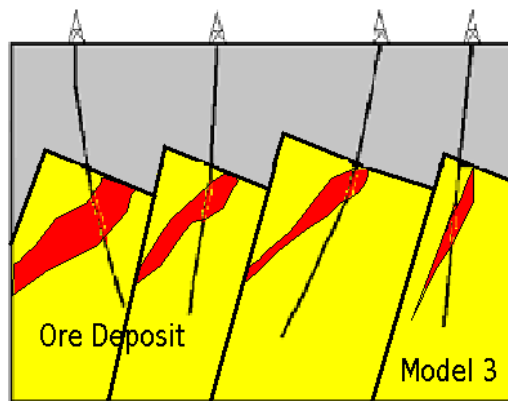
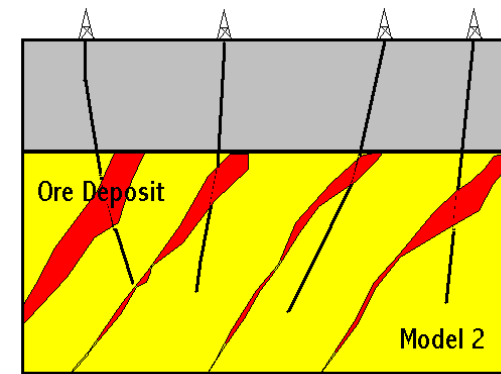
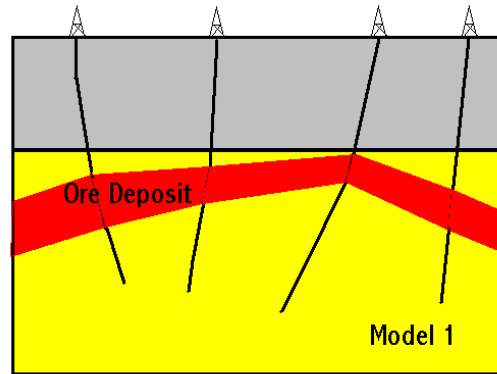
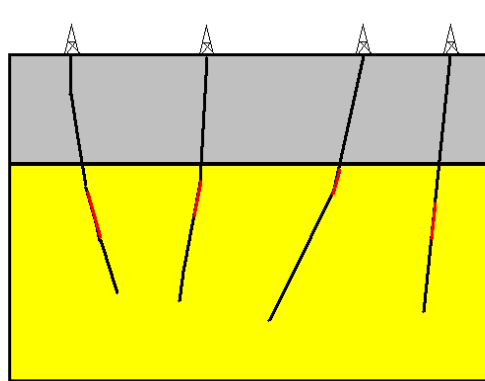
Anisotropy

Behavior at the border...

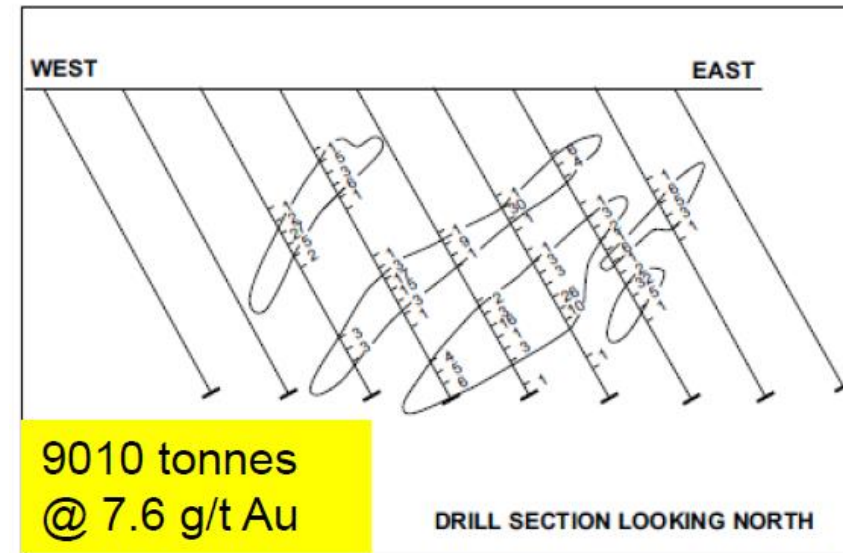
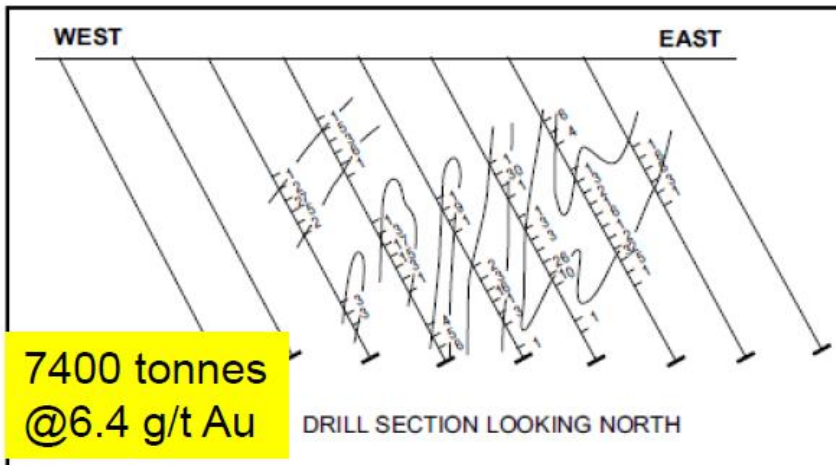
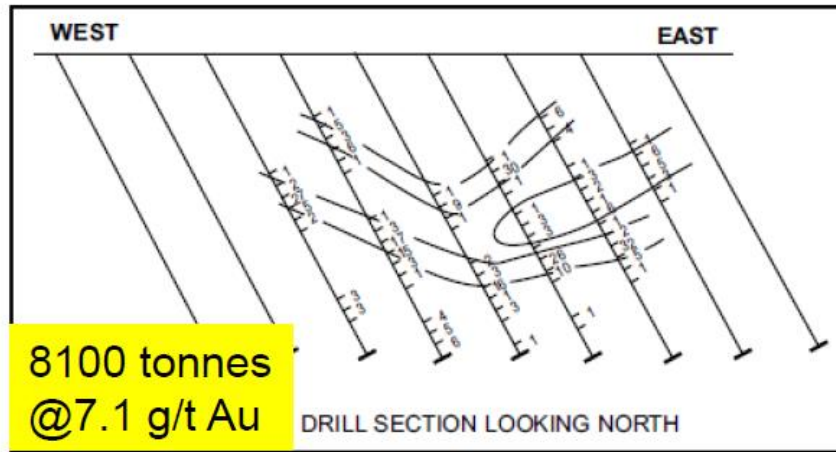
An early stage in a geostatistical project is to partition the study area into (approximately) stationary domains

What is a Geometric Model ?

- Geological Concepts
- Need to Fit Data
- Account for Structural Geology



Geological continuity and uncertainty

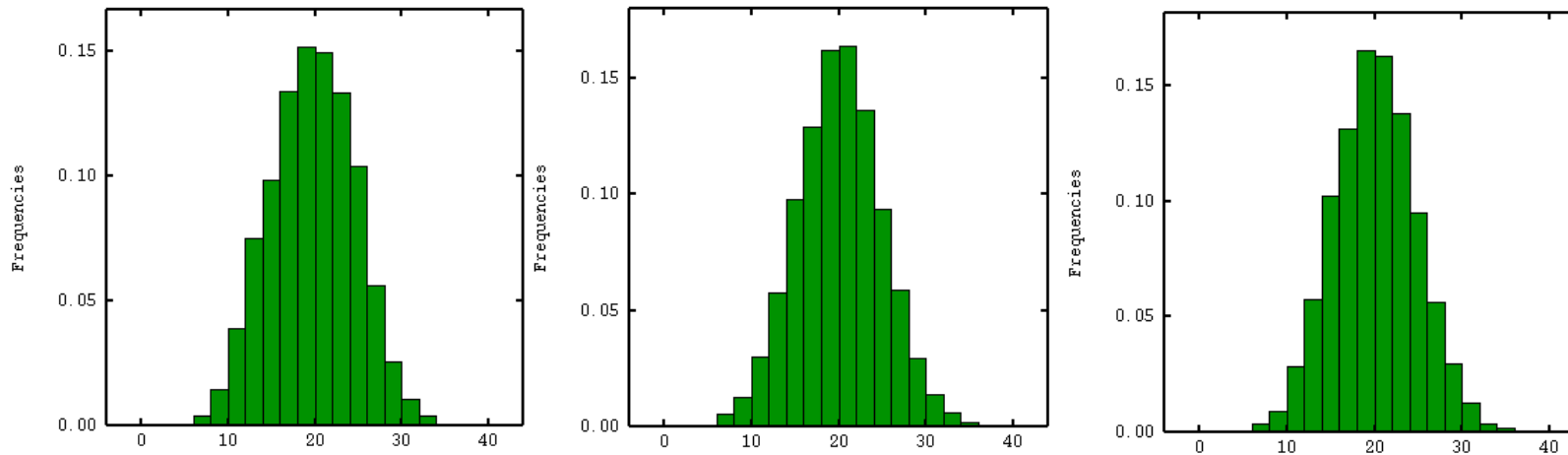


04

Necessity of the experimental variogram

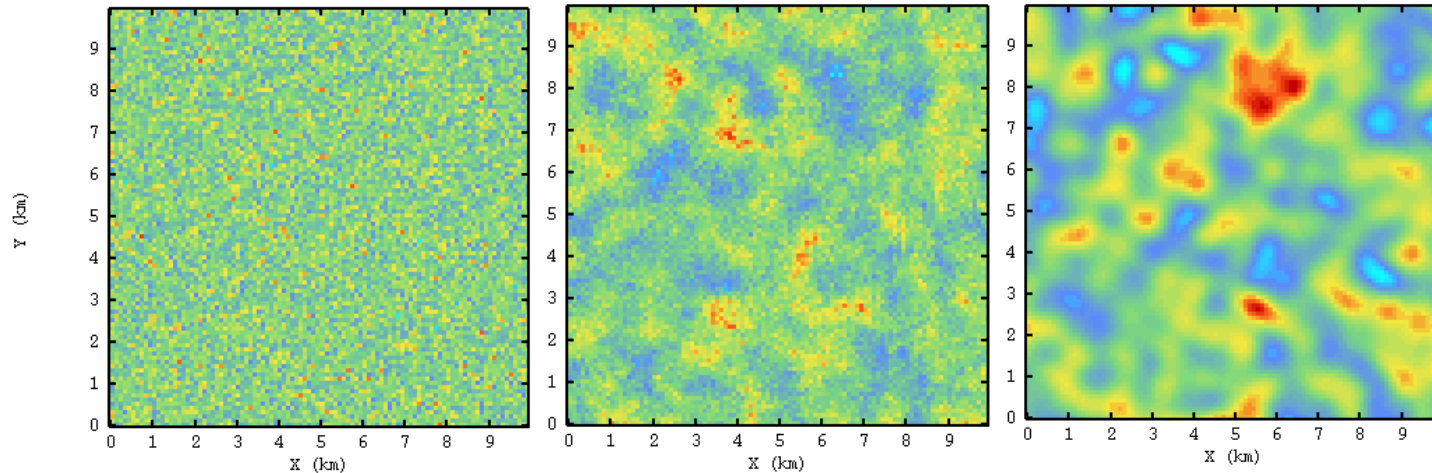
Why do we need variograms?

Three variables. Means 19.7, 20.0, 20.0. Standard deviations 4.82, 4.73, 4.87. Histograms:



Why do we need variograms?

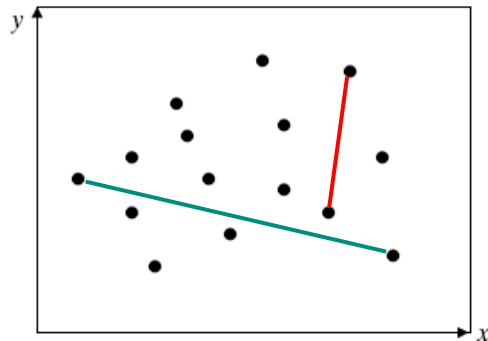
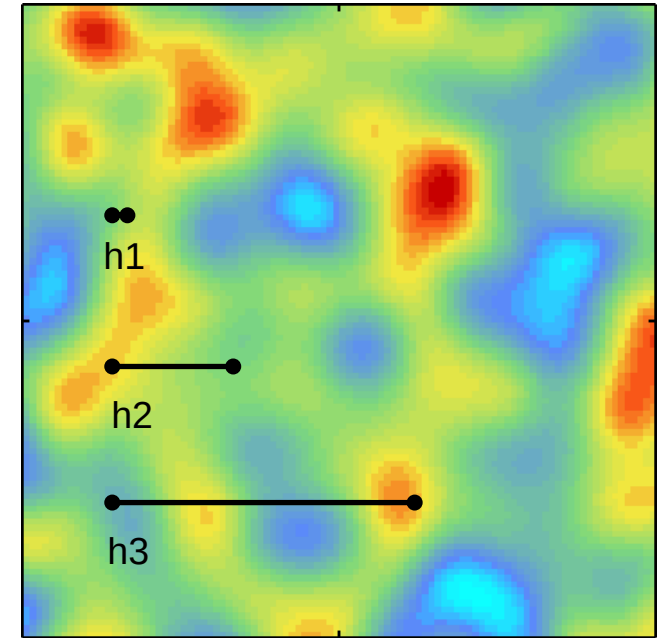
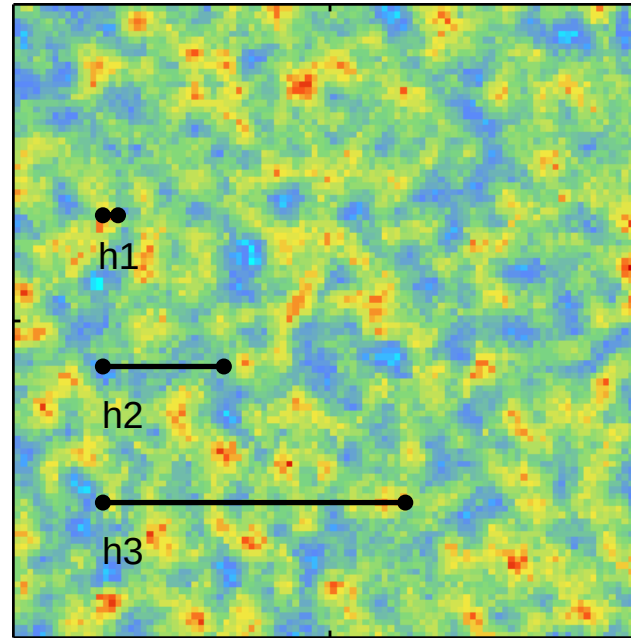
But color maps of the three variables look nothing alike!



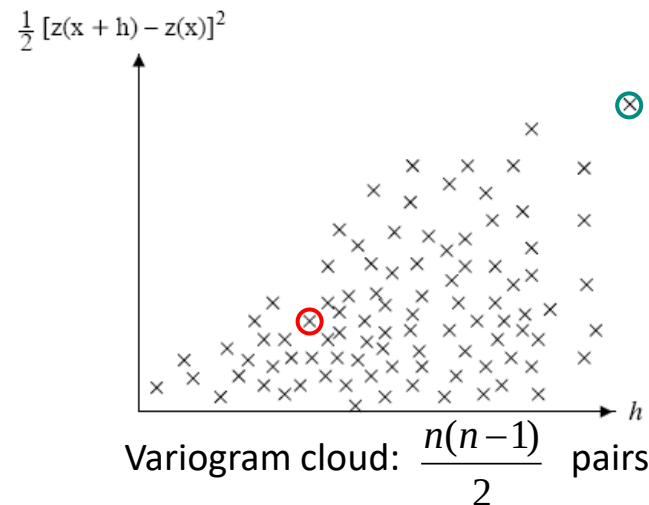
We need something more than just the **histogram** to have some hope of describing the spatial structure of the variable: the **variogram**.

Spatial Model

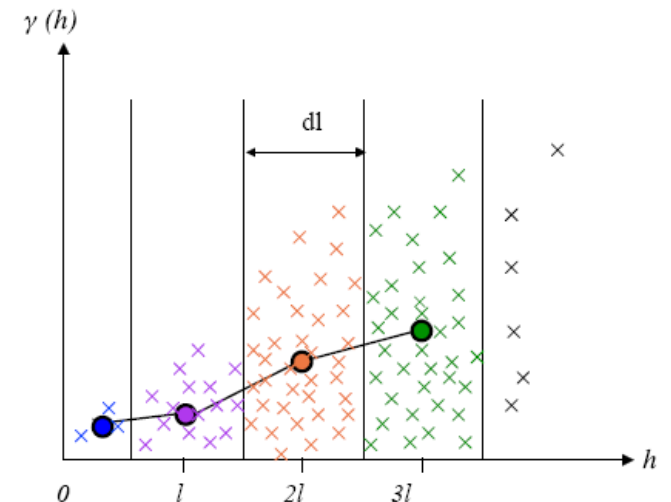
Spatial models are linked to the relationship between data at different distances.



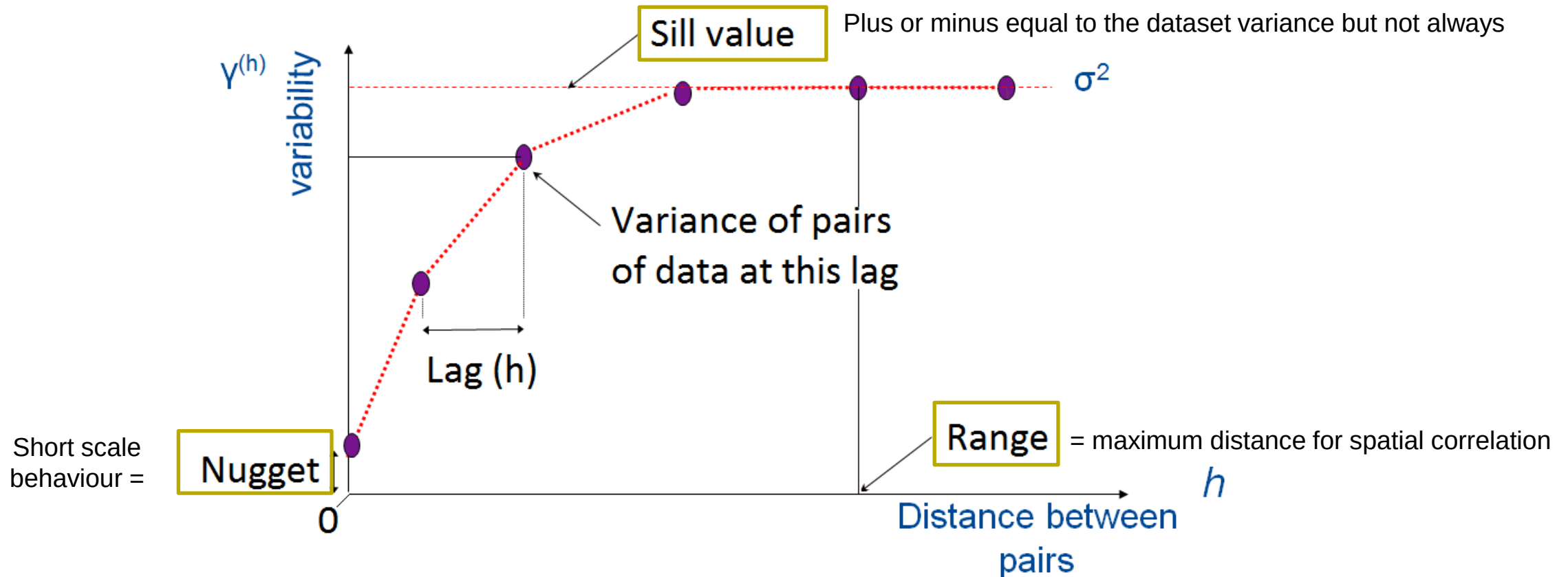
Base map of n points



Variogram cloud: $\frac{n(n-1)}{2}$ pairs



Schematic of a variogram

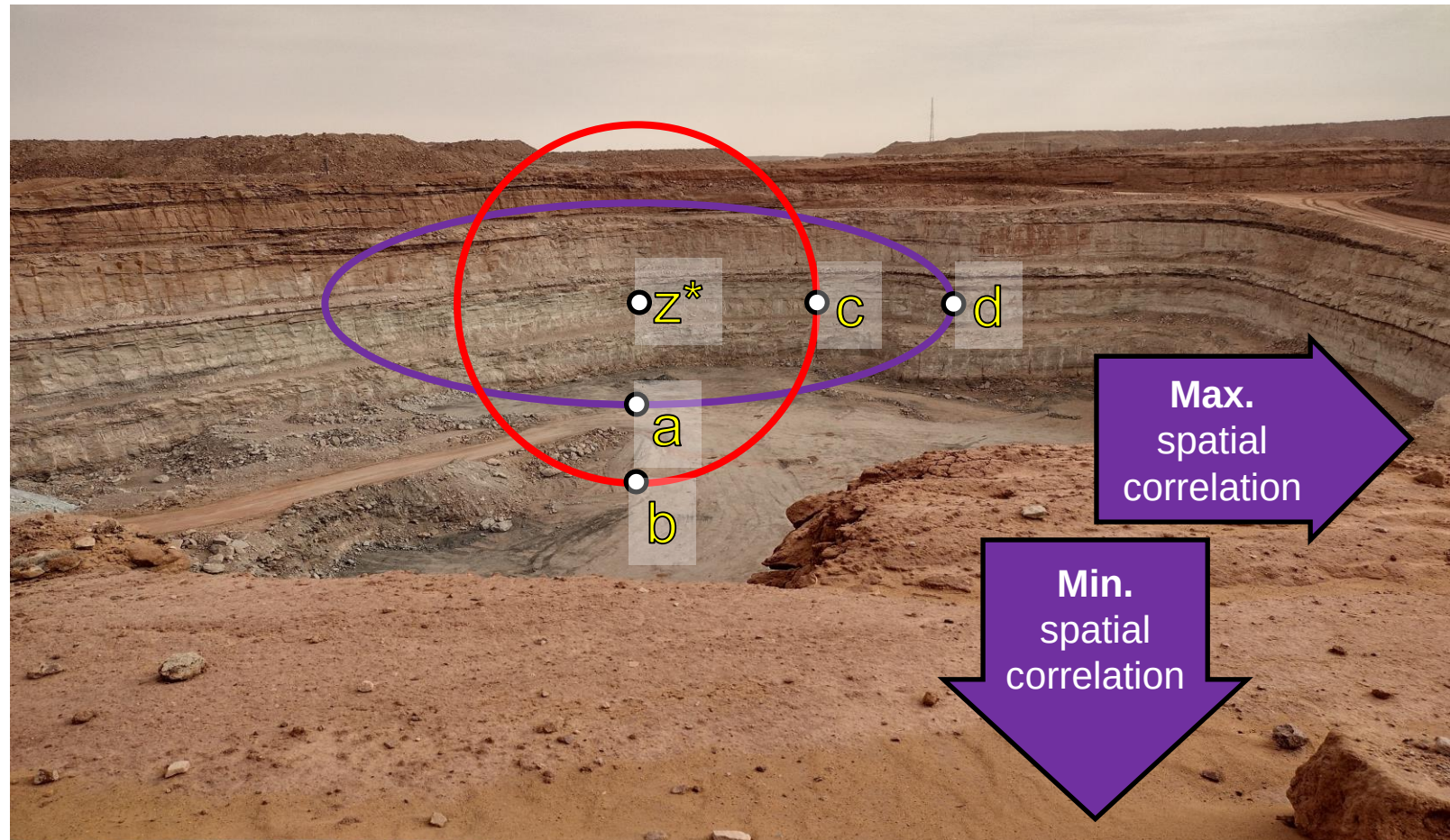


Understand **spatial correlation** between points in space, for a given property.

Anisotropy/Directional variogram

Which sample (a, b, c or d) look at z^* the most?

Need of directional variogram → can be seen via a variogram map

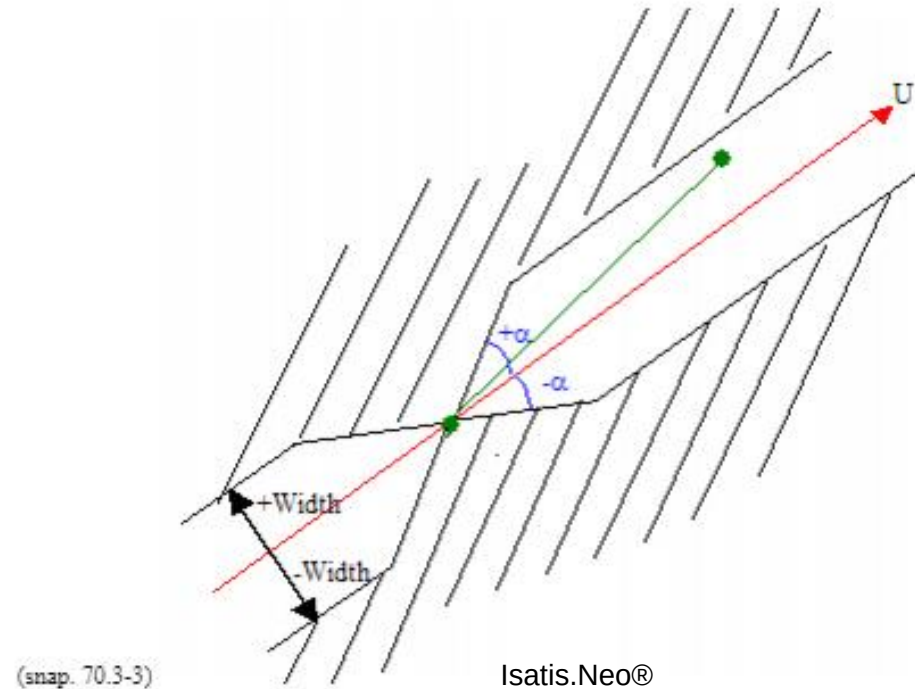


Directional variogram, 2D

For each **direction (U)**, the pairs are selected within an **angular sector** defined by:

slicing width

angular tolerance



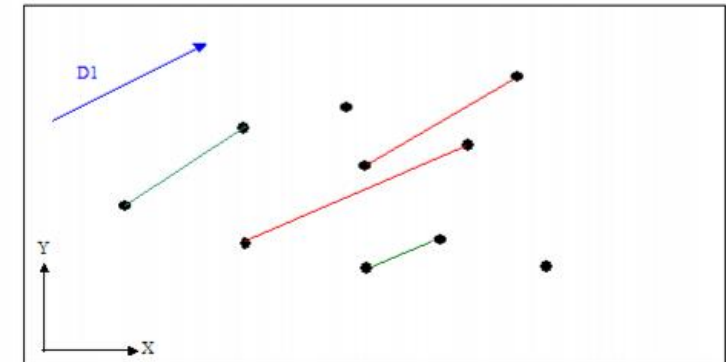
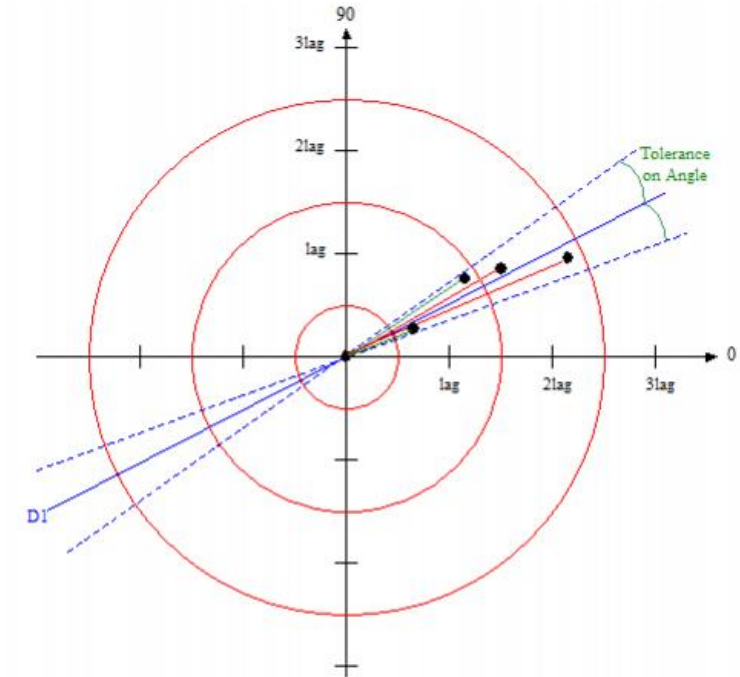
Directional variogram, 2D

Classes of distance :

lag distance
maximum distance (number of lags)
lag tolerance

Classes of directions (sectors):

reference direction
number of sectors
angular tolerance



(snap. 70.3-2)

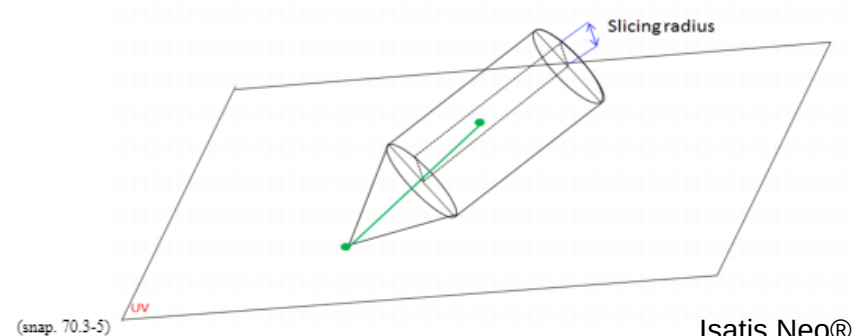
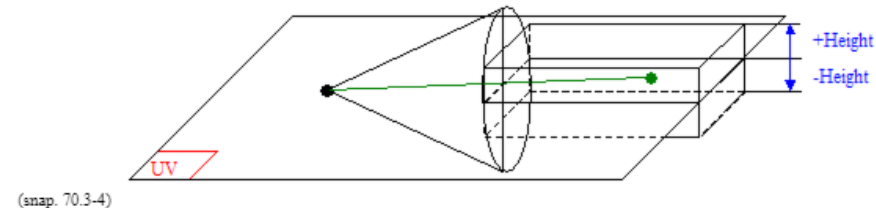
Directional variogram, 3D: slicing height

For each **direction (U)**, the pairs are selected within an angular sector defined by:

slicing width (as in 2D)

angular tolerance (as in 2D)

slicing height (specific for 3D)



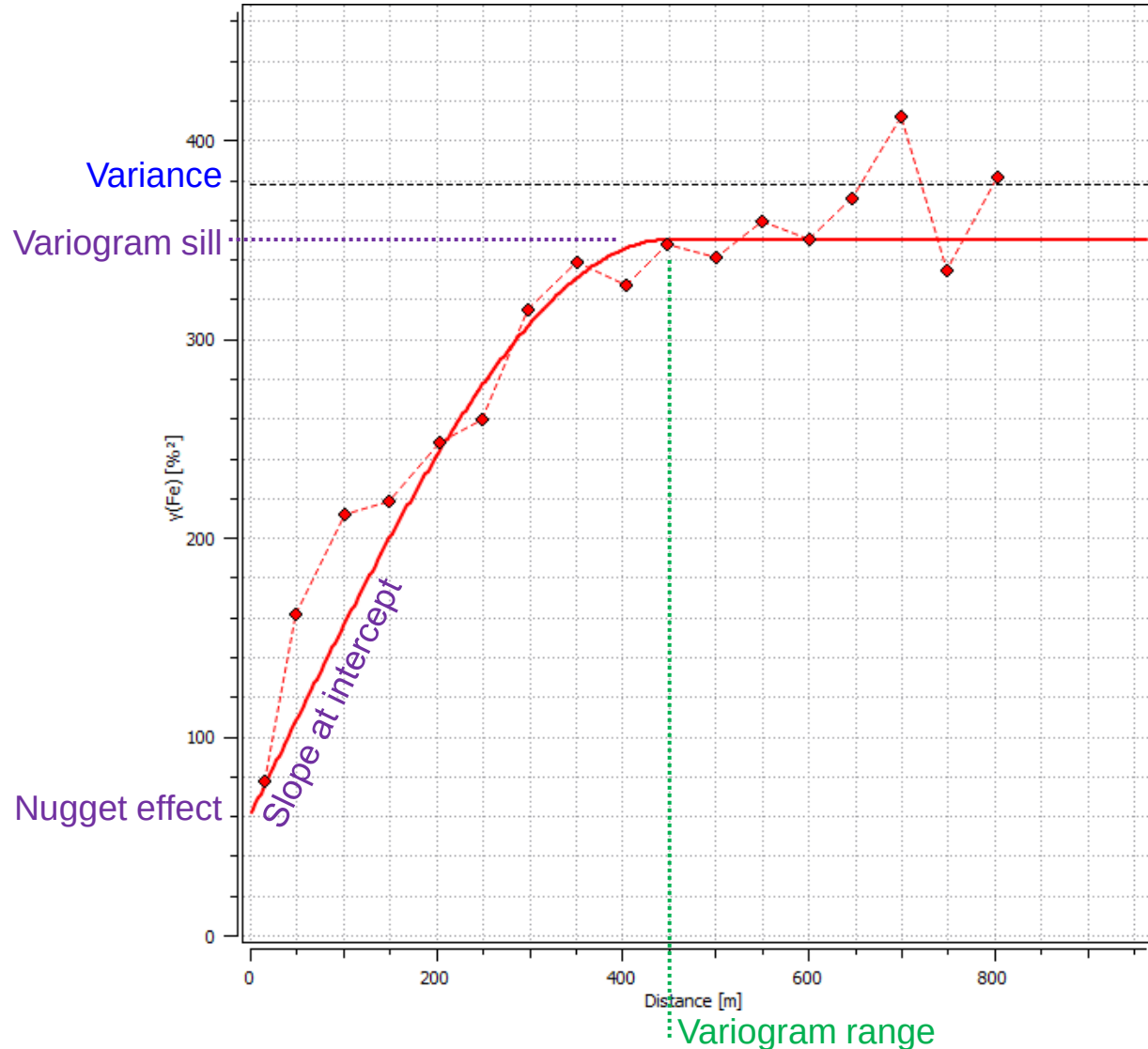
Isatis.Neo®

05

**Variogram model:
Interpreting spatial
continuity**

Variogram model parameters

- Variogram increases with distance up to a **sill**
- Variogram **sill** should be close to the dataset **variance**
- **Range** is the distance at which **sill** is reached
- It is the distance beyond which no spatial correlation



A focus on the nugget effect

Nugget effect is a **random component** (**not spatialized**) of the regionalized variable

Possible origins of the nugget effect

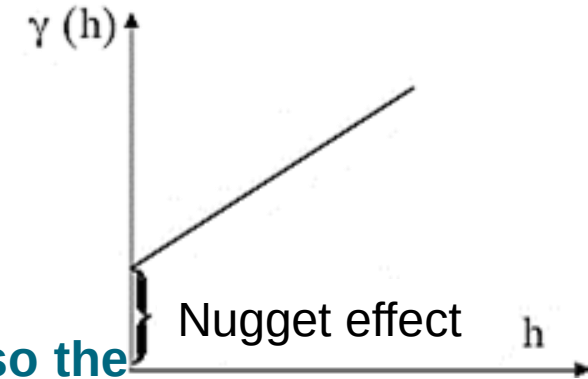
By sampling, e.g. soil, the **environment is perturbed**, so the second sample at the same location will be different from the first sample

Short scale heterogeneity (gold nugget)

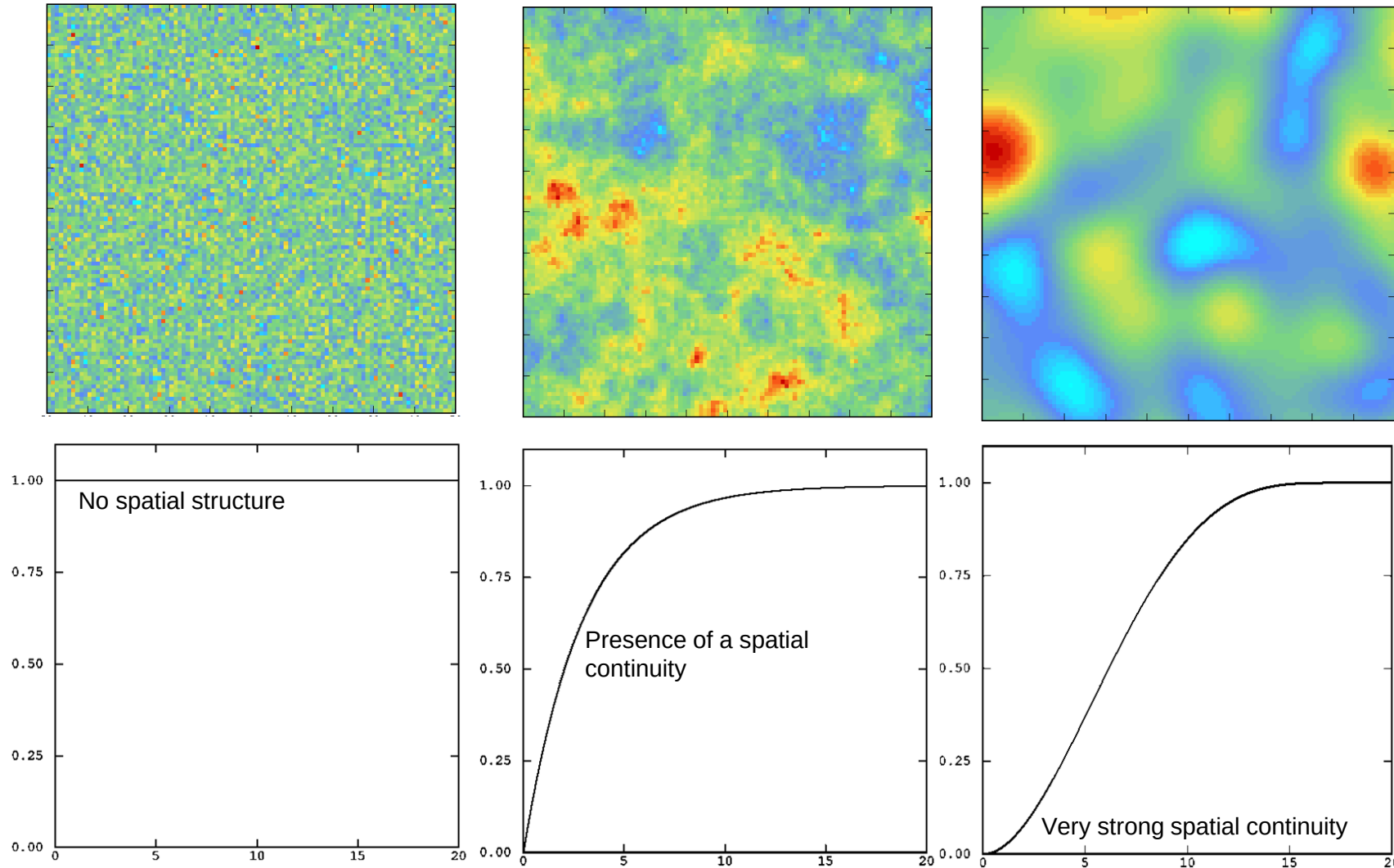
Uncertainties of **laboratory** analyses

While the measurement mechanism is **stochastic**, e.g. some geophysical measurements, the repetitive measurements will not be identical (metrology)

Errors and **uncertainties** of positioning system (GPS or mapping imprecision), two points are supposed to be close while they are not close in the reality

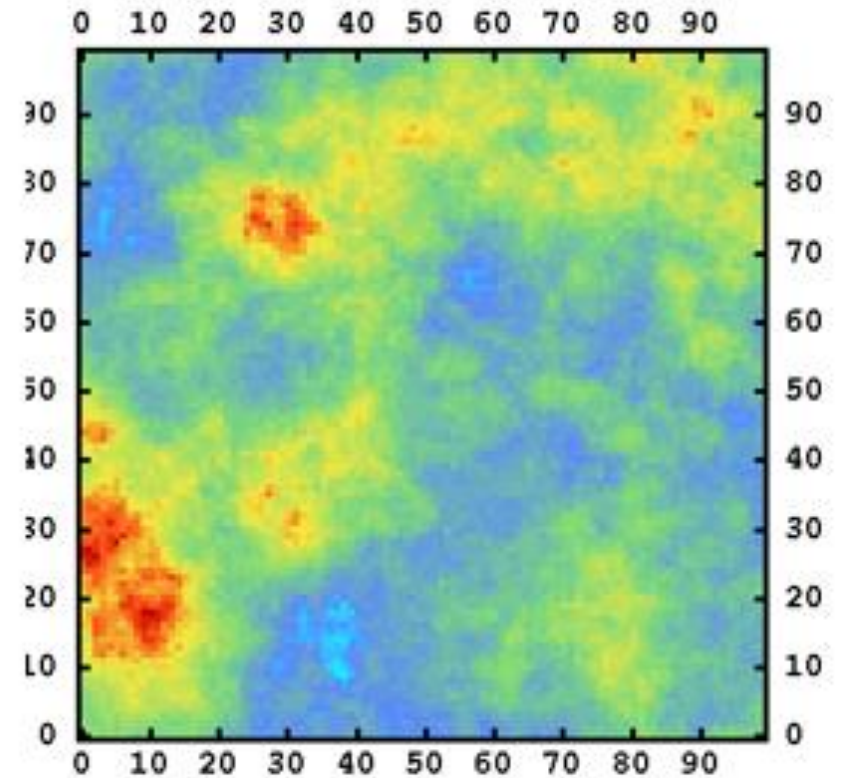
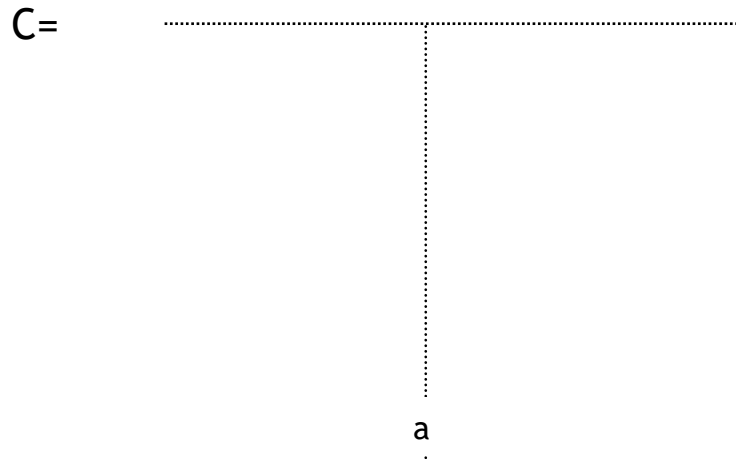


Variograms of three examples



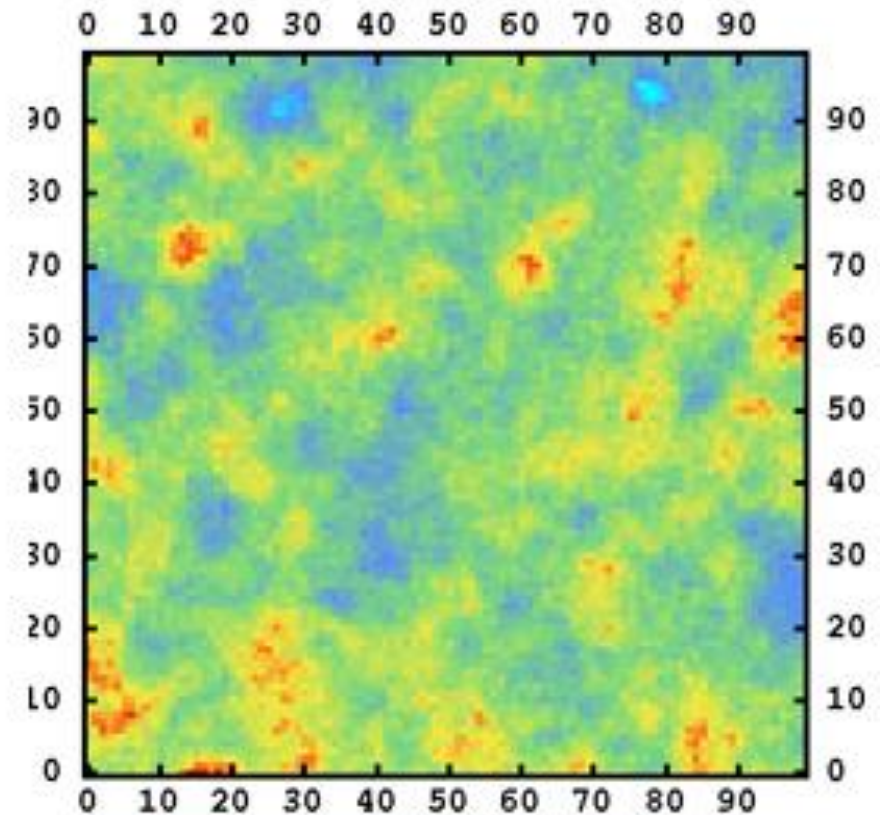
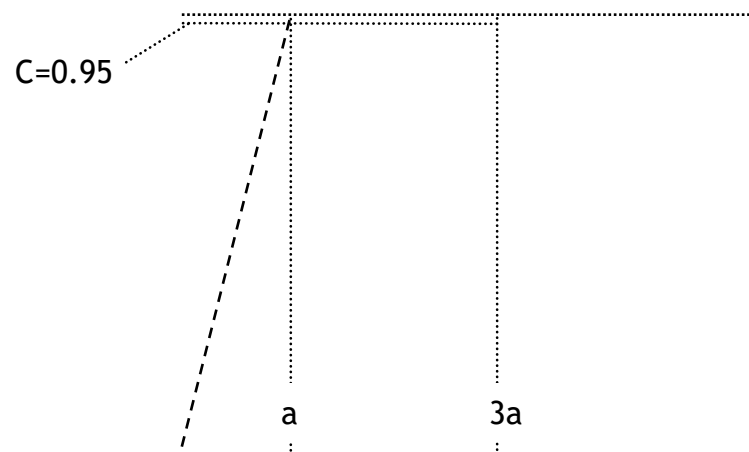
Spherical variogram

$$\gamma(h) = \begin{cases} C \left(\frac{3h}{2a} - \frac{1}{2} \frac{h^3}{a^3} \right) & \text{if } 0 \leq h \leq a \\ C & \text{if } h \geq a \end{cases}$$

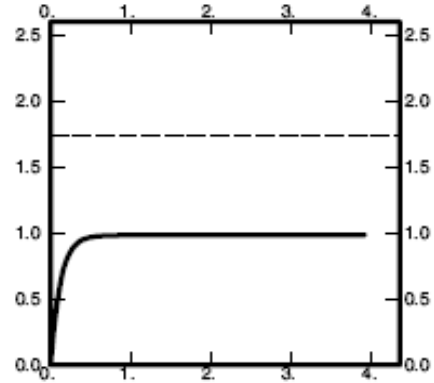


Exponential variogram

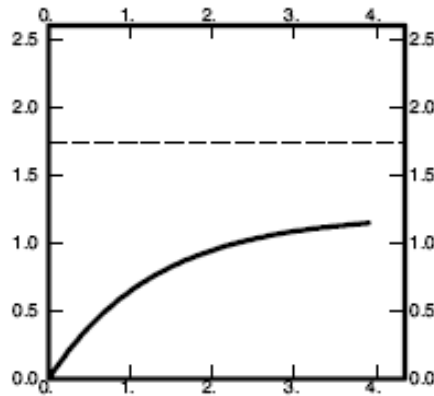
$$\gamma(h) = C \left(1 - \exp \left(-\frac{h}{a} \right) \right)$$



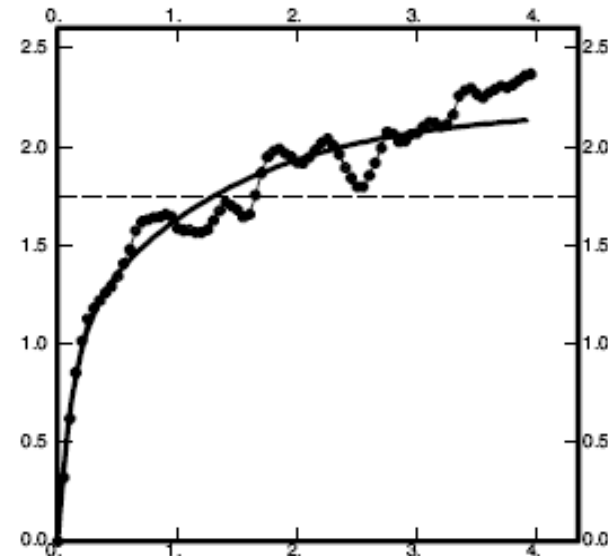
Variogram with nested structures



Short-range structure



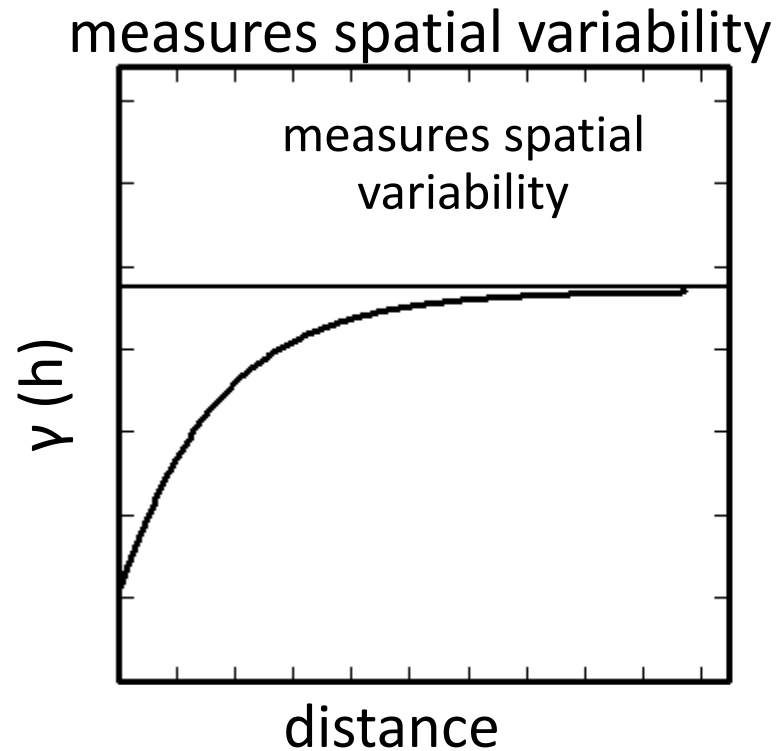
Long-range structure



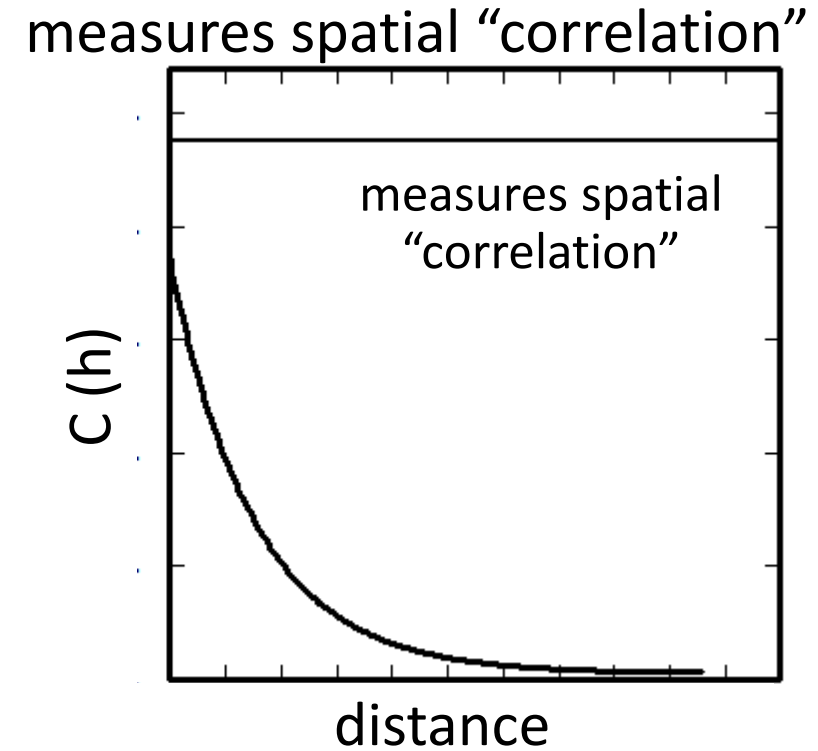
Variogram model
with 2 nested structures

Variograms and covariance

Variogram



Covariance



If stationary: $C(h) = C(0) - \gamma(h)$

06

Ordinary kriging

Geostatistical interpolation: Kriging

Linear Kriging aims to interpolate the variable as a linear combination of the sample value:

$$Z_0^* = \sum_{\alpha} \lambda_{\alpha} Z_{\alpha}$$

Z_0^* : interpolated value (unknown)

λ_{α} : weights at location α (determined by solving the kriging equation)

Z_{α} : sample values at location α (known)

Kriging aims to interpolate the **sample values** taking the spatial structure/variability into account by referring directly to the variogram **model**

The variogram model, a mathematical function, is modeled from the experimental variogram

The Kriging algorithm uses the variogram model to define the weights:

$$\lambda_{\alpha} = \text{function}(\text{distance}, \text{variogram})$$

Ordinary Kriging

In matrix notation, the Ordinary Kriging system is (in variogram)

$$\underbrace{\begin{bmatrix} \gamma_{11} & \cdots & \gamma_{1n} & \mathbf{1} \\ \vdots & \ddots & \vdots & \vdots \\ \gamma_{n1} & \cdots & \gamma_{nn} & \mathbf{1} \\ \mathbf{1} & \cdots & \mathbf{1} & 0 \end{bmatrix}}_{\text{Variogram values between data points, pair by pair}} \times \underbrace{\begin{bmatrix} \lambda_1 \\ \vdots \\ \lambda_n \\ -\mu \end{bmatrix}}_{\text{Weights}} = \underbrace{\begin{bmatrix} \gamma_{01} \\ \vdots \\ \gamma_{0n} \\ \mathbf{1} \end{bmatrix}}_{\text{Variogram values between data and target point}}$$

Kriging Estimate:

$$Z_0^* = \sum_{\alpha} \lambda_{\alpha} Z_{\alpha}$$

Kriging Variance:

$$\text{Var}(Z_0 - Z_0^*)_{\min} = \sigma_K^2 = - \sum_{\alpha} \lambda_{\alpha} \gamma_{0\alpha} + \mu$$

What type of Kriging?

Simple Kriging (SK)

Order 2 stationarity

Known mean (m)

No condition on weights

$$Z_0^* = \sum_{\alpha} \lambda_{\alpha} Z_{\alpha} + \lambda_{mean} \times m$$

Problem with SK: m is generally unknown, it is replaced by an estimate m^*

Note: λ_{mean} is a derived quantity, not directly solved for:

$$\lambda_{mean} = 1 - \sum \lambda_{\alpha}$$

Ordinary Kriging (OK)

Intrinsic stationarity

Unknown mean

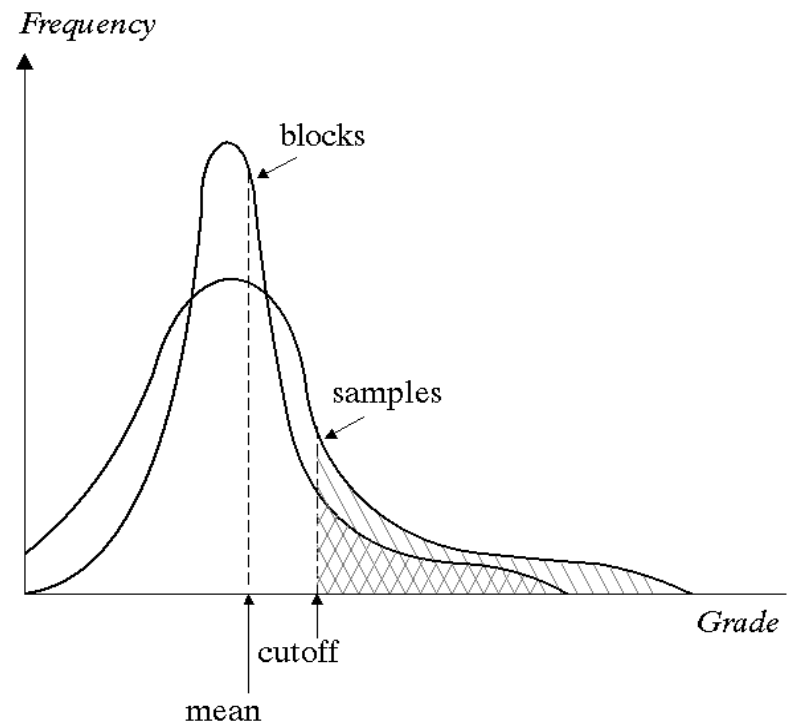
Sum of weights must be equal 1

$$Z_0^* = \sum_{\alpha} \lambda_{\alpha} Z_{\alpha}$$
$$\sum_{\alpha} \lambda_{\alpha} = 1$$

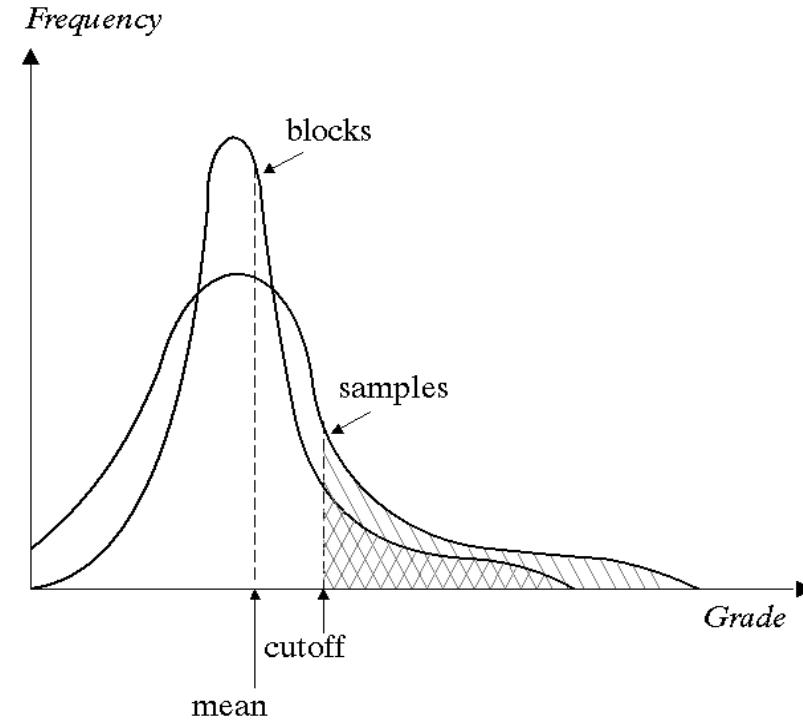
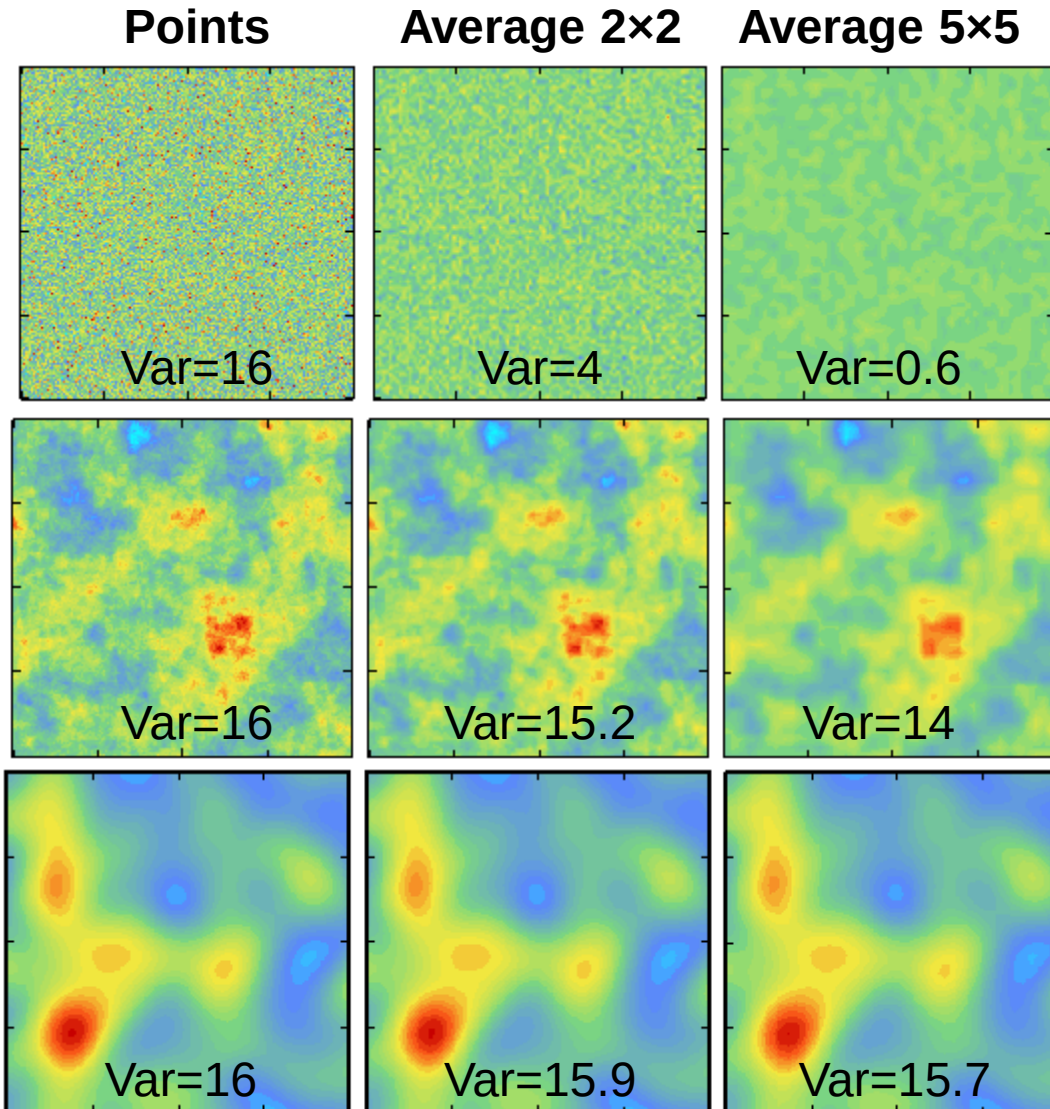
07

**Block interpolation :
Support effect**

Support effect



Averaging effect



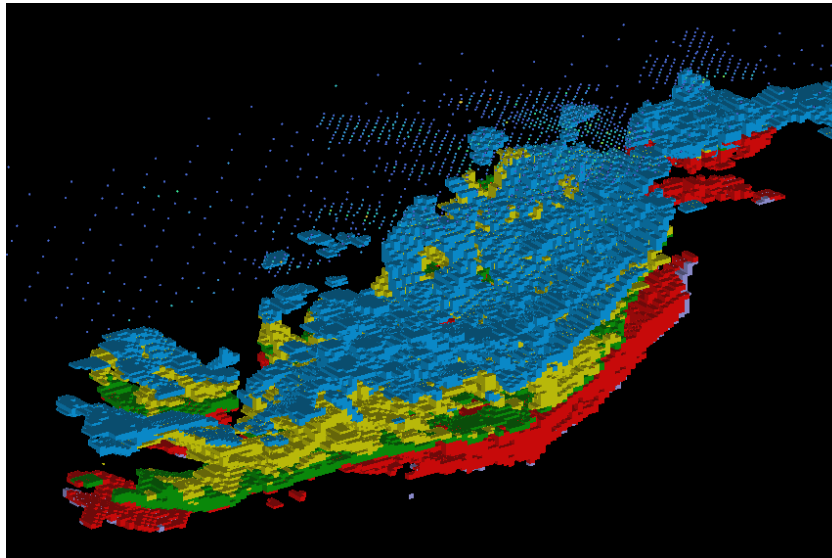
As support increases:
Variance decreases
Mean does not change significantly

Notion of block

Samples (core, cutting etc) can be approximated by points as their volumes are not very important

Nevertheless, when considering a mineralization, the support has to be taken into account

To consider this support effect, the space is divided into blocks, usually rectangles : **the block model**



Block model cell size?

It depends in:

- the estimation method
- and the spatial continuity of the variable:

For kriging, inverse distance or nearest neighbor, the block model cell size is usually equal to the **drillhole spacing** : trying to estimate a smaller support is not recommended as the variability at this scale is unknown

If the variable is very continuous (long range variogram, no or small nugget effect), the cell size can decrease until half of the drillhole spacing

Usually the mining team wants an estimation on the **selective mining unit (SMU)** → **non linear geostatistics**

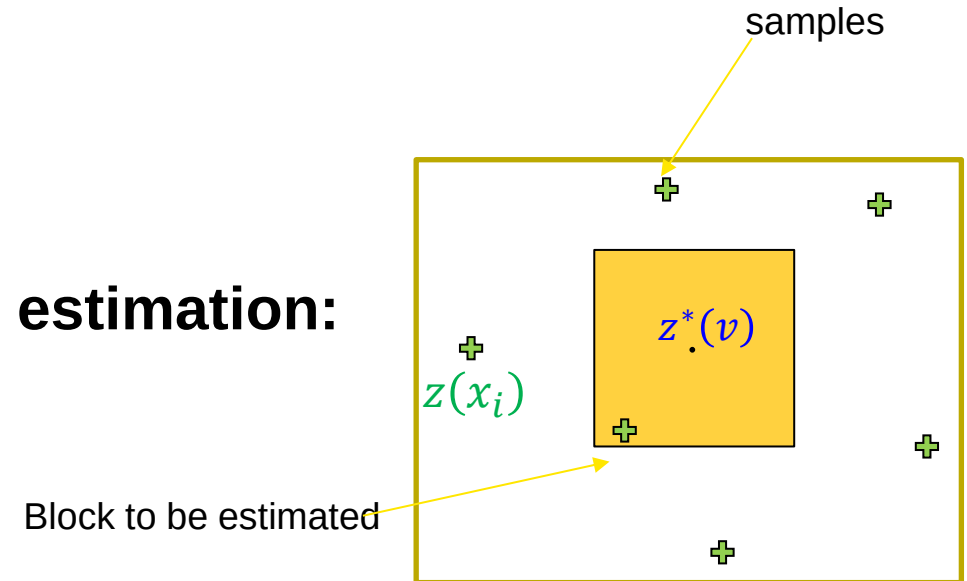
08

Block interpolation
Block kriging

Block kriging

The estimate of an average block grade is not simply the point estimate of the grade at the block centroid

The larger the support, the smoother is the estimation:
Larger blocks result more averaging

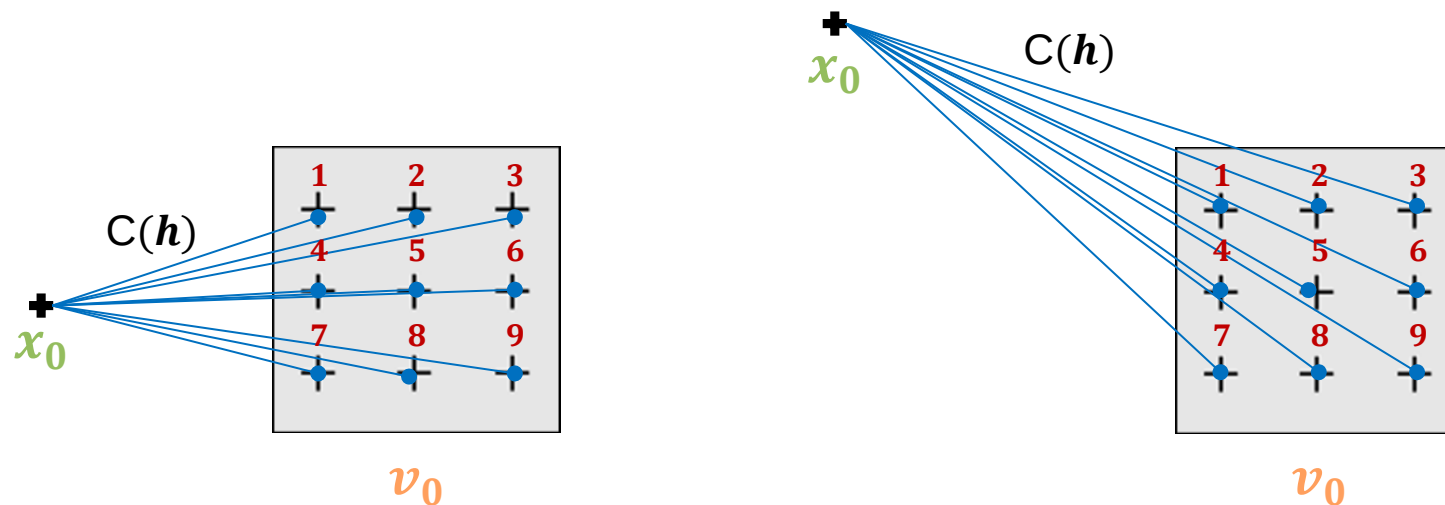


Goal: estimate an average value $z^*(v)$ of the block v from the available punctual data $z(x_i)$

$$z^*(v) = \sum_{i=1}^n \lambda_i z(x_i)$$

Block Kriging

The variability between a sample and a block is given by a **point-block covariance**. It is approximated by discretizing the block and averaging the covariance values from sample to discretized points:



$$C(x_0, v_0) = \frac{1}{9} \sum_{i=1}^9 C(x_0, x_i)$$

Ordinary Kriging (blocks)

In matrix notation, the block Kriging system is:

$$\underbrace{\begin{bmatrix} \gamma_{11} & \cdots & \gamma_{1n} & 1 \\ \vdots & \ddots & \vdots & \vdots \\ \gamma_{n1} & \cdots & \gamma_{nn} & 1 \\ 1 & \cdots & 1 & 0 \end{bmatrix}}_{\text{Variogram values between data points, pair by pair}} \times \underbrace{\begin{bmatrix} \lambda_1 \\ \vdots \\ \lambda_n \\ \mu \end{bmatrix}}_{\text{Weights}} = \underbrace{\begin{bmatrix} \bar{\gamma}(x_1, v) \\ \vdots \\ \bar{\gamma}(x_n, v) \\ 1 \end{bmatrix}}_{\text{Average variogram values between data and target block}}$$

09

Neighborhood Parameters

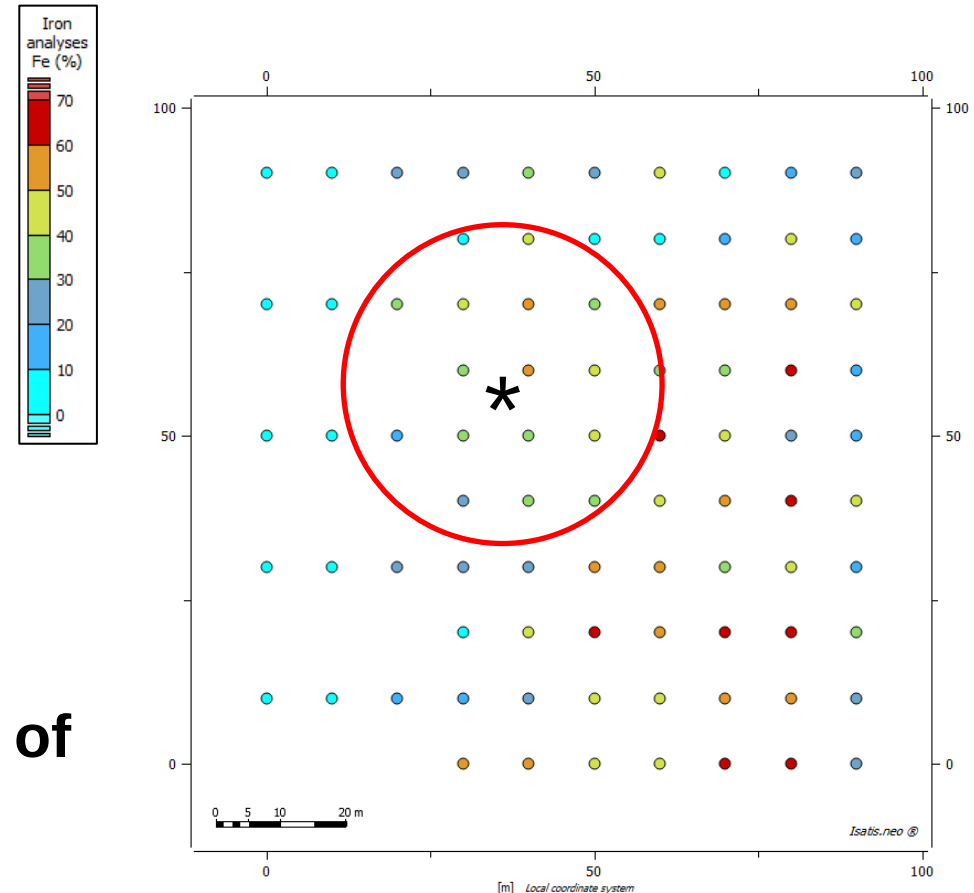
Two neighborhood types

Unique (global):

- All samples are used for each target interpolation
- Applicable only for small datasets (computation time scales badly)

Moving:

- For each target interpolation, a limited set of **surrounding** samples will be used for interpolation



How to choose the neighborhood parameters?

- Usually, the neighborhood ellipse size is based on the variogram range
- But if the area is densely drilled or if the variable is very continuous, the ellipse range can be smaller
- Take into account stationnarity
- Depends in the estimation method used
- Uniform conditioning or simulations required to use more samples and bigger ellipsoid than ordinary kriging
- Some rule of thumbs
 - Minimum sample used at least 3
 - Can put a constraint in the number of samples used per drillholes to make sure that several drillholes are used in the estimation

Nested neighborhood

- To ensure that all the domain will be estimated, a nested neighborhood can be used
 - First run can have a restricted neighborhood with constraints in the ellipsoid size and number of samples used
 - Second run can be less restrictive and use samples further from the estimated block
 - Third run is even less restrictive
 - Even infinite neighborhood can be used
-
- Remember that the estimation quality will be poorer when the constraints are eased → « bull eyes » might appear

Kriging neighborhood analysis

Quantities to consider are:

Kriging standard deviation or Kriging efficiency KE

Slope of regression Z/Z^*

Correlation Z/Z^*

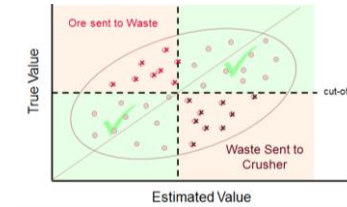
Weight of the mean from Simple Kriging

Sum of positive weights

Computation time

$$KE = 1 - \frac{\sigma_K^2}{\sigma_V^2}$$

Kriging variance
Theoretical block variance



Neighbourhood parameters that can be varied:

Search radii of the ellipsoid

Rotation of the ellipsoid

Use of anisotropic distances

Number of angular sectors (used to try to have samples from all directions surrounding the target block)

Maximum number of samples per sector

Minimum number of samples

Each estimated block will have its own estimate quality (based on the surrounding samples), but it is infeasible to optimize the neighbourhood for every block.

Instead, look at averages over the domain, some subset of the domain, or a representative slice.

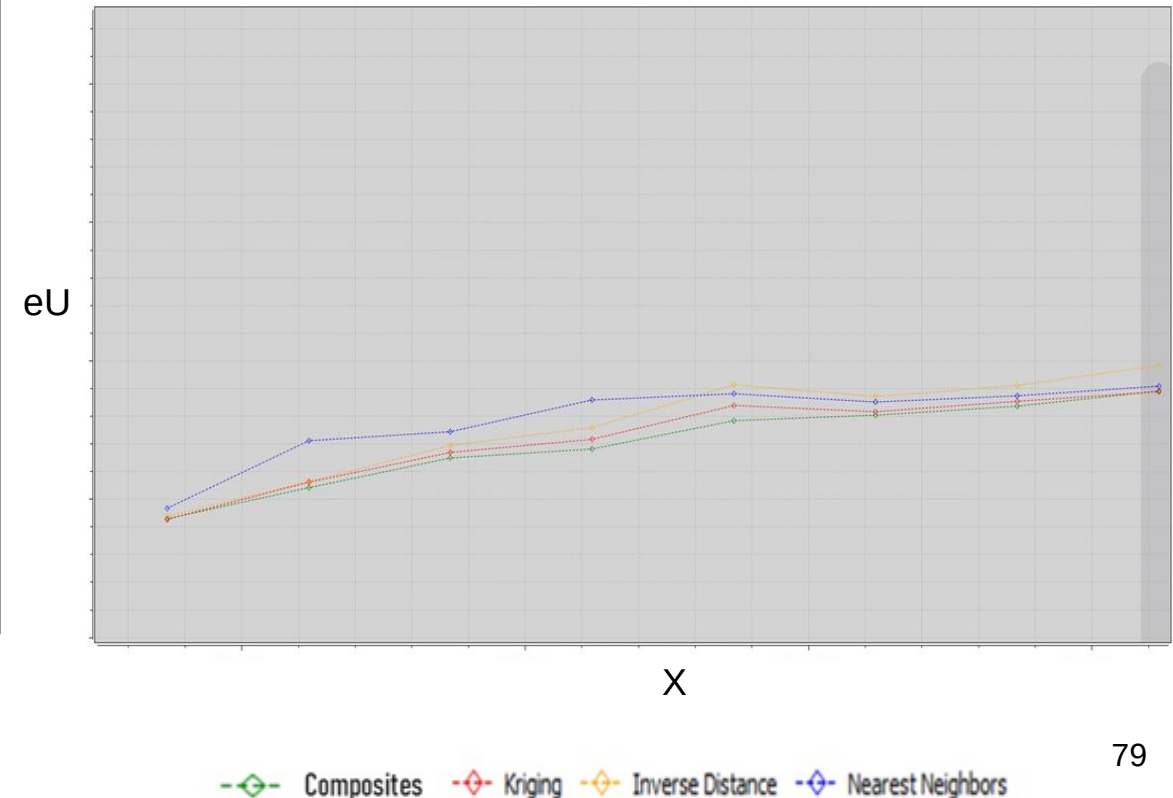
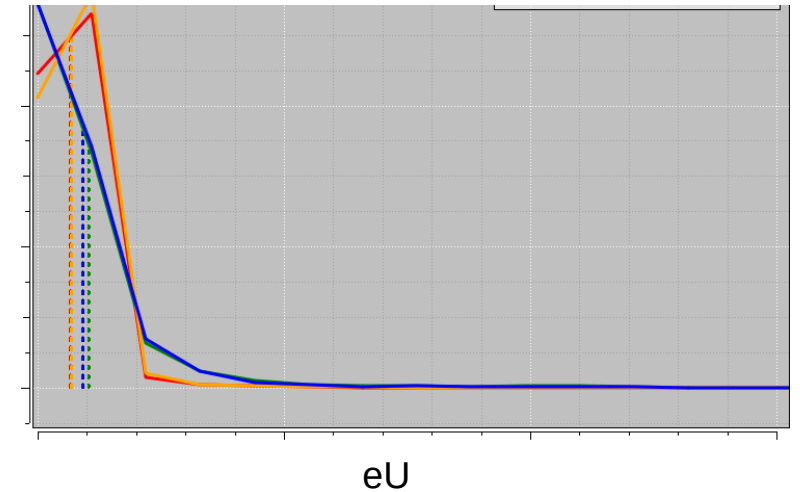
10

Estimation validation

Estimation validation

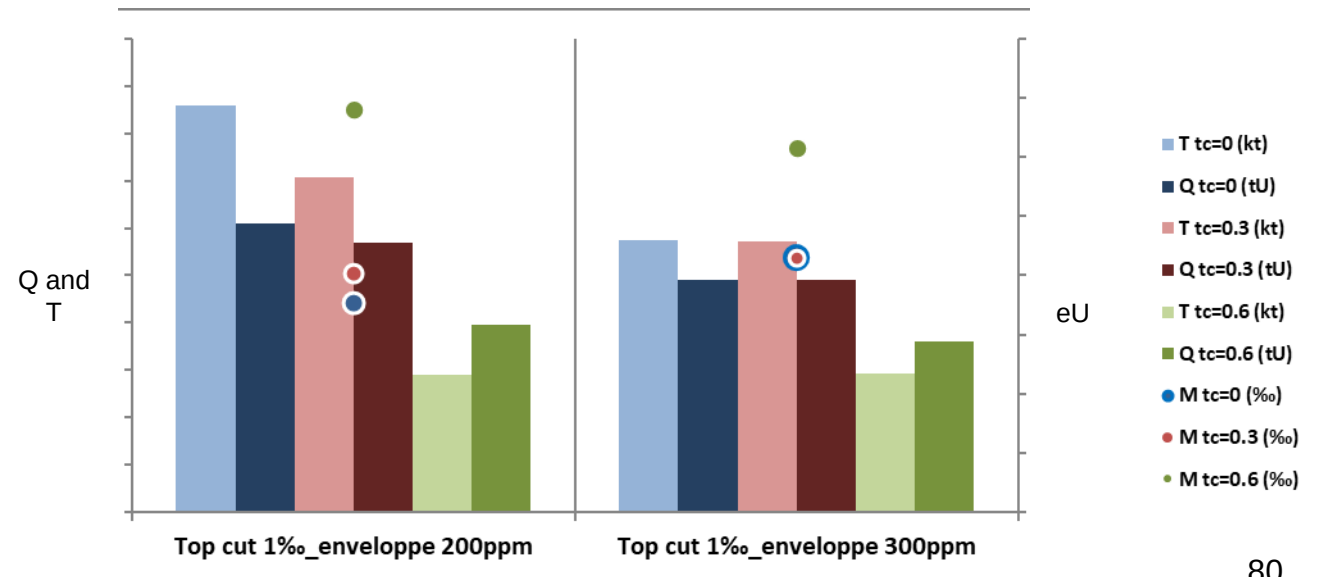
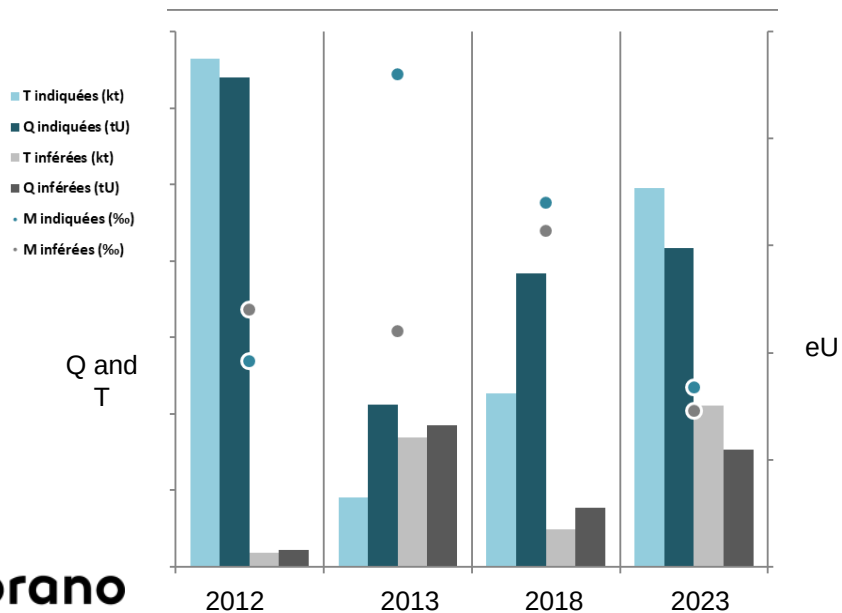
How to validate the estimation?

- Comparison of the sample statistics and the estimation statistics (mean, variance, minimum, maximum etc)
- Comparison between different estimation methods (nearest neighbor, inverse distance, ordinary kriging)
- Comparison between a 2D estimation and a 3D one
- Swath plots
- Cross sections validation, maps
- Looked at the kriging variance maps
- Compare the estimation with the production to see what the differences are



Estimation sensitivities

- Sensitivities can be run for the estimation, to see if the estimation result is very sensitive to some parameters
- Removing of some samples used for the experimental variogram calculation to see if the variogram model is still making sense
- Top capping the variable to see the impact in term of metal, mean grade and tonnage
- Testing several neighborhood parameters



11

Multivariate geostatistics

Objective of Multivariate Geostatistics

Objective: to improve the quality of the estimation by using the auxiliary variable(s)

Multivariate data types

Multi-element assays (Au, Ag, Cu etc)

Top, bottom and thickness of a layer

Thickness, accumulation and 2D grade

Geophysical data (e.g. depth in seismic marker in wells)

Multivariate Geostatistics

Uses relationships between variables

Secondary variable can be sampled at same or different points

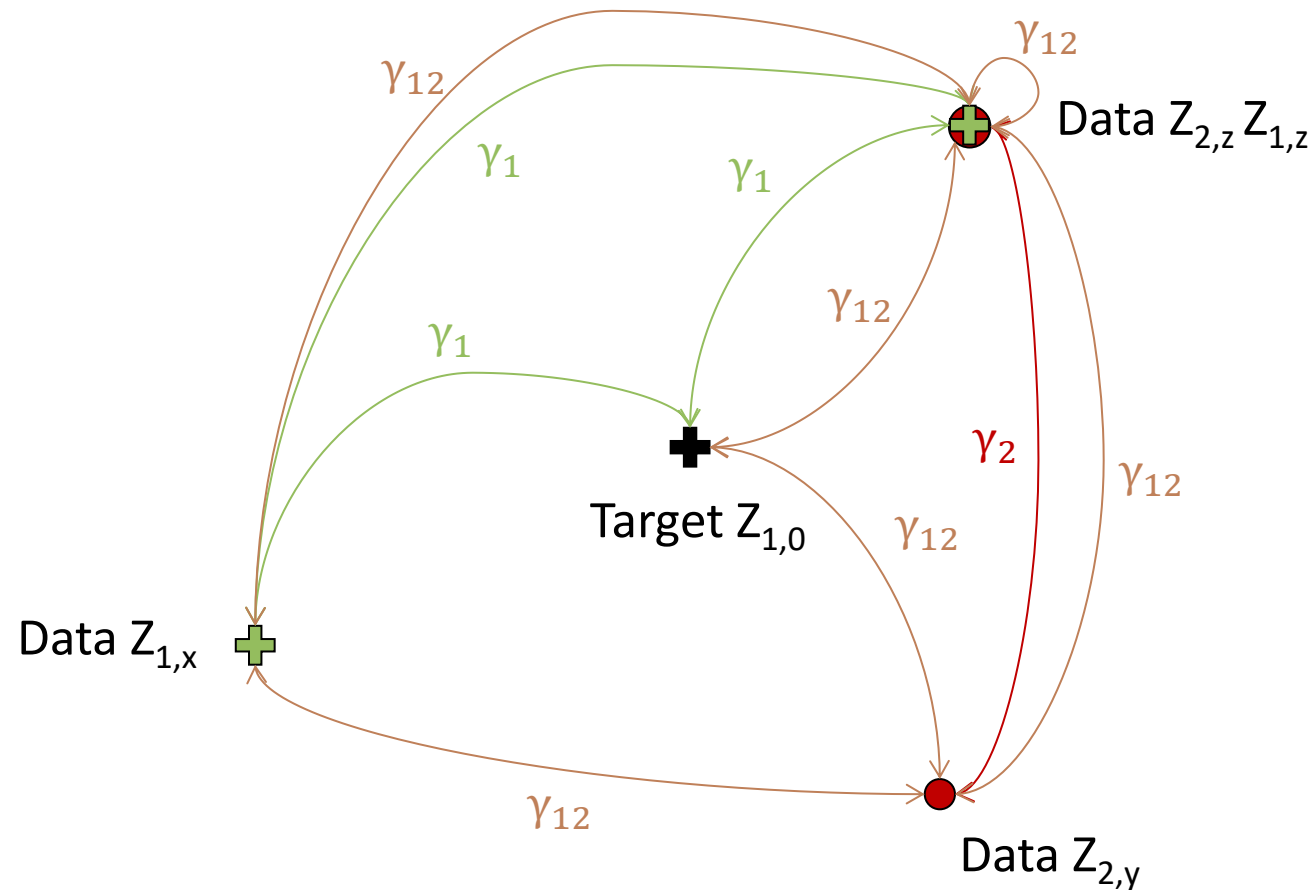
Improves consistency within estimated variables (e.g. layer's top, bottom & thickness)

Partly substitute for missing data

Simulate several variables jointly

Multivariate estimation

- It requires knowing the links between the target variable $Z_{1,0}$ and in between every data $Z_{i,h}$:



Variogram Modeling

The relation between the variable will be described with the cross variogram. To be able to model the variogram, a hypothesis is taken:

The **linear model of coregionalisation (LMC)**: all simple and cross variograms must be linear combinations of same basic structures

Structure 1:

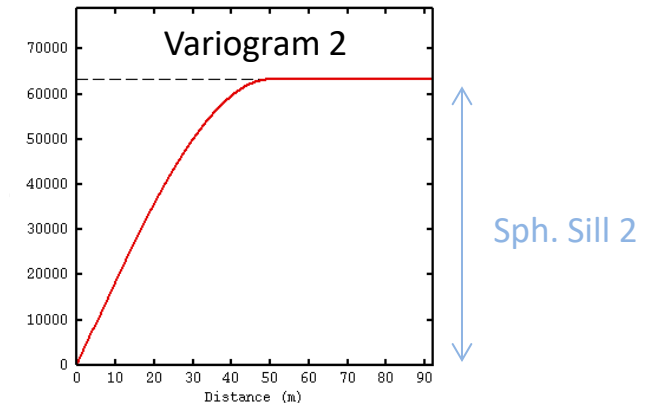
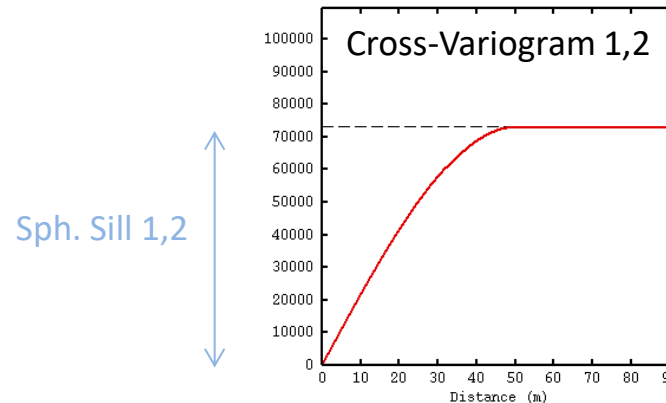
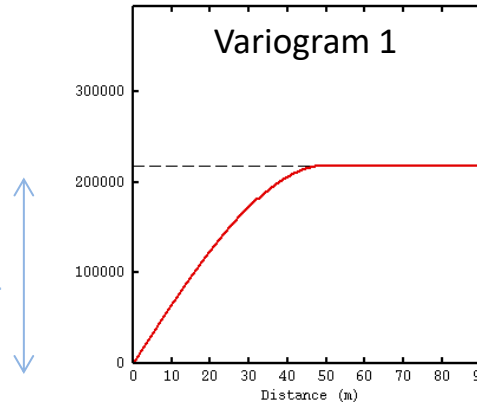
Spherical – 50m

Sill Matrix 1:

$$\begin{bmatrix} \text{Sill 1} & \text{Sill 1,2} \\ \text{Sill 1,2} & \text{Sill 2} \end{bmatrix}$$

Sph. Sill 1

Sph. Sill 1,2



Ordinary Co-Kriging

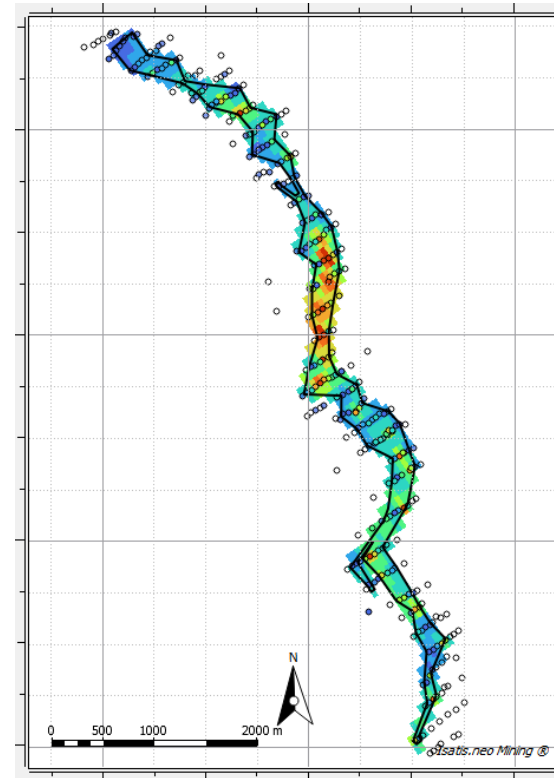
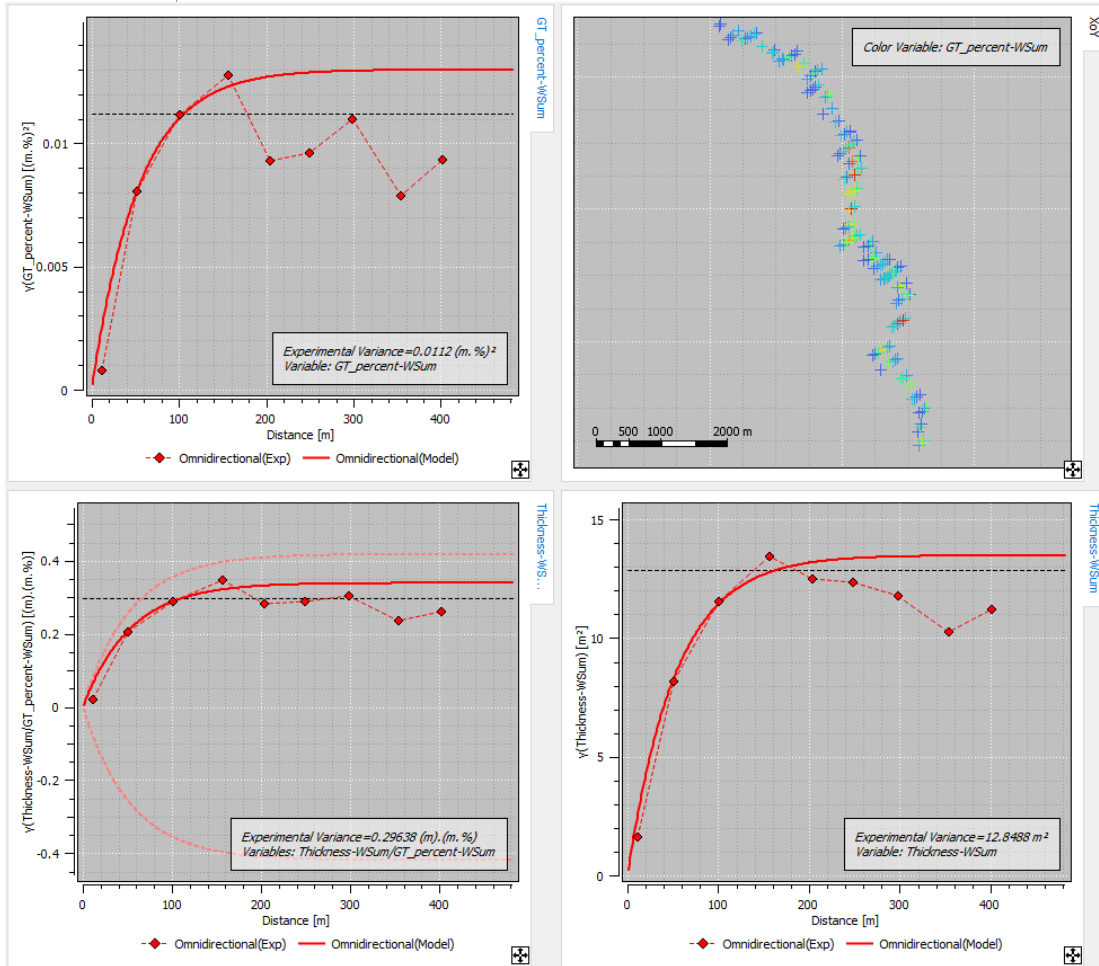
Principle: Estimate Z_1 using a linear combination of Z_1 and Z_2 measured at scattered data points

$$Z_1^* = \sum_{\alpha} \lambda_1^{\alpha} Z_1^{\alpha} + \sum_{\alpha} \lambda_2^{\alpha} Z_2^{\alpha}$$

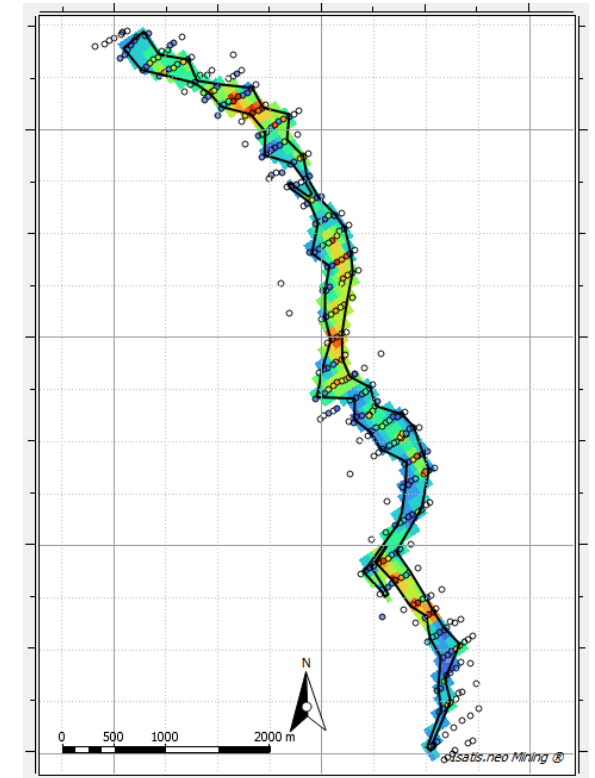
The ordinary co-Kriging system is

$$\left\{ \begin{array}{l} \sum_{\beta} \lambda_1^{\beta} C_{\alpha\beta}^{Z_1} + \sum_{\beta} \lambda_2^{\beta} C_{\alpha\beta}^{Z_1 Z_2} + \mu_1 = C_{\alpha 0}^{Z_1} \quad \forall \alpha \\ \sum_{\beta} \lambda_1^{\beta} C_{\alpha\beta}^{Z_1 Z_2} + \sum_{\beta} \lambda_2^{\beta} C_{\alpha\beta}^{Z_2} + \mu_2 = C_{\alpha 0}^{Z_1 Z_2} \quad \forall \alpha \\ \sum_{\alpha} \lambda_1^{\alpha} = 1 \\ \sum_{\alpha} \lambda_2^{\alpha} = 0 \end{array} \right.$$

Example of co kriging



Metal quantity



Thickness

12

Recoverable resources

Recoverable resources

Consider a **selection block** v , with grade $Z(v)$

Definition:
The recoverable resources at cut-off z_c are defined as

Ore

$$T(z_c) = 1_{Z(v) \geq z_c}$$

With 1 beginning an indicator taking the value 1 when $Z(v) \geq z_c$, and 0 elsewhere

Metal

$$Q(z_c) = Z(v) \times 1_{Z(v) \geq z_c}$$

(multiplied by block tonnage = volume*density expressed in tons)

Recoverable resources

Recoverable resources at cut-off z_c for N blocks v
(panel V or whole deposit, assuming constant density)

Ore proportion

$$T(z_c) = \frac{1}{N} \sum_{i=1}^N 1_{Z(v_i) \geq z_c}$$

Metal

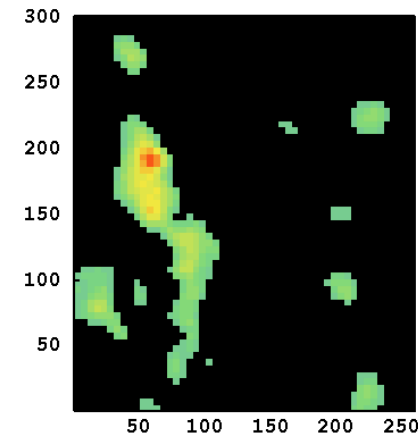
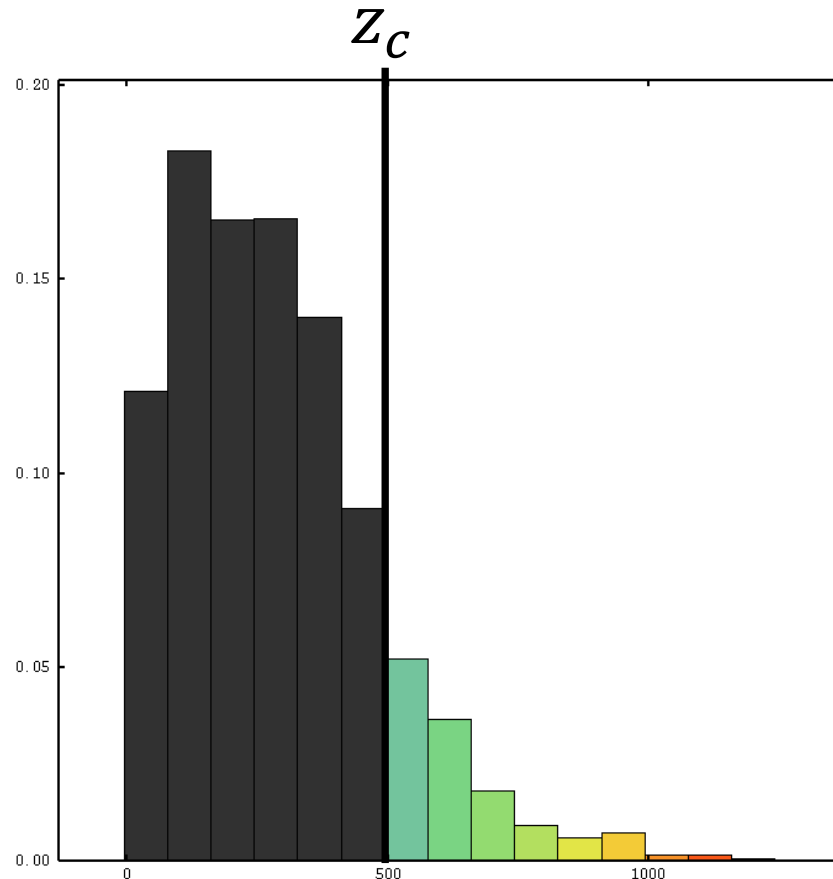
$$Q(z_c) = \frac{1}{N} \sum_{i=1}^N Z(v_i) \times 1_{Z(v_i) \geq z_c}$$

Grade:

$$m(z_c) = \frac{\sum_{i=1}^N Z(v_i) \times 1_{Z(v_i) \geq z_c}}{\sum_{i=1}^N 1_{Z(v_i) \geq z_c}} = \frac{\sum_{Z(v_i) \geq z_c} Z(v_i)}{\text{Nb}[Z(v_i) \geq z_c]} = \text{mean}[Z(v_i) | Z(v_i) \geq z_c]$$

Recoverable resources

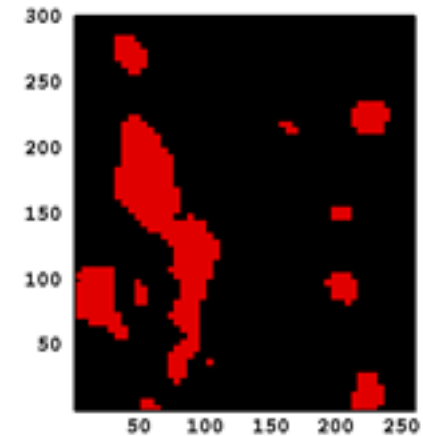
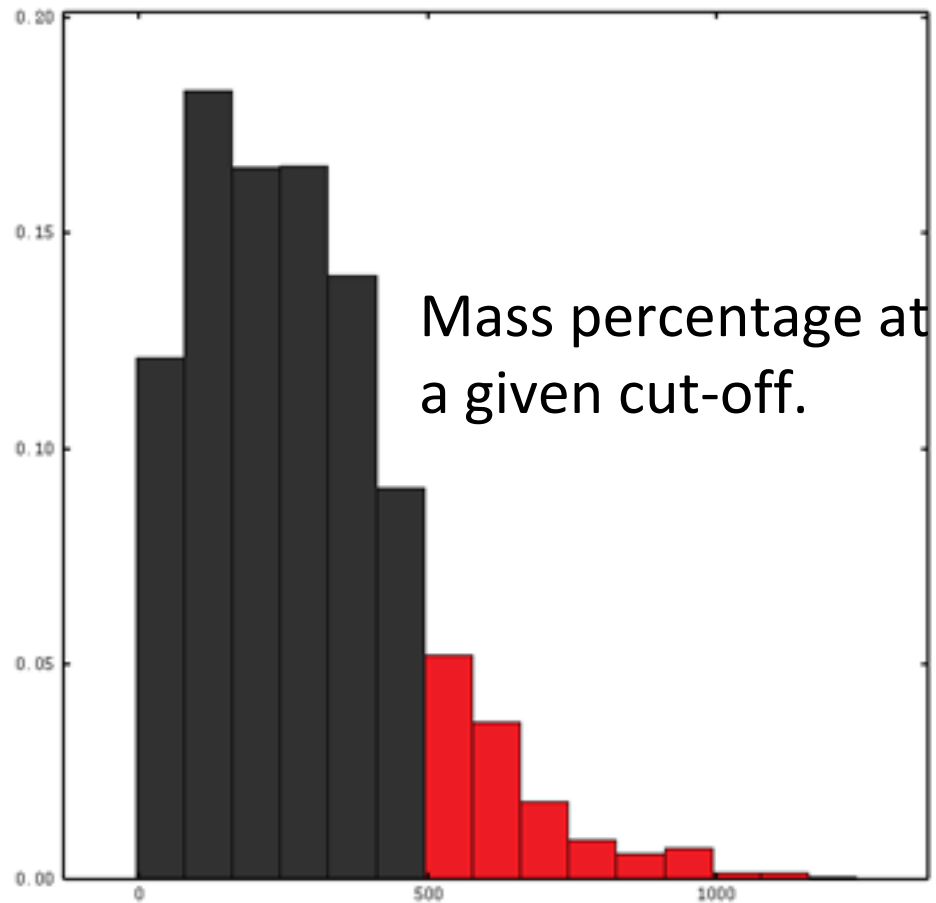
After grade estimation, the blocks whose grades are above cutoff are considered to be recovered, and we can define new variables on this basis.



e.g., $z_c = 500$

Recoverable resources

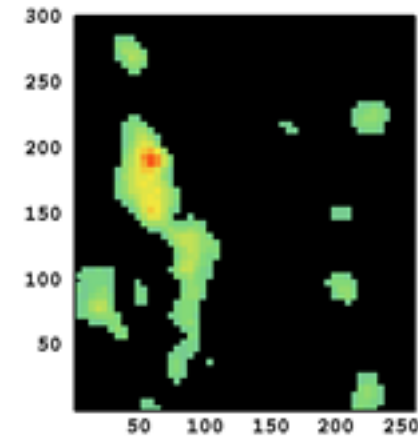
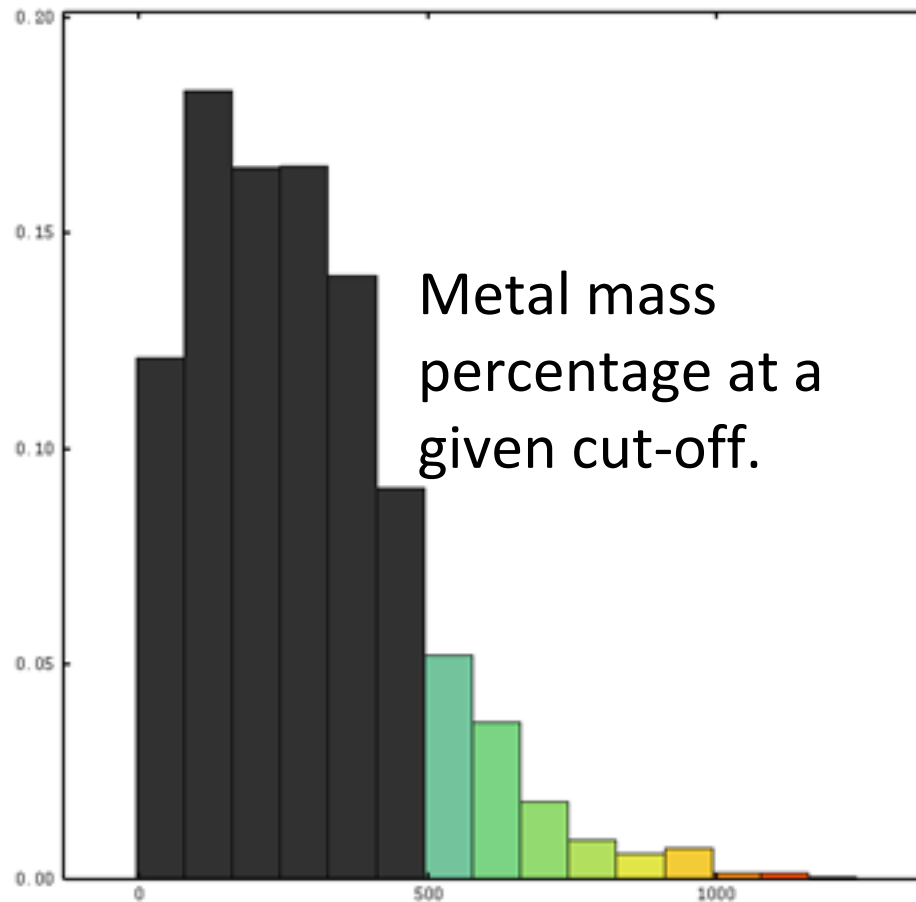
T : tonnage above cutoff, $T = E[\mathbb{I}_{Z \geq z_c}]$



$T = 15.16\%$

Recoverable resources

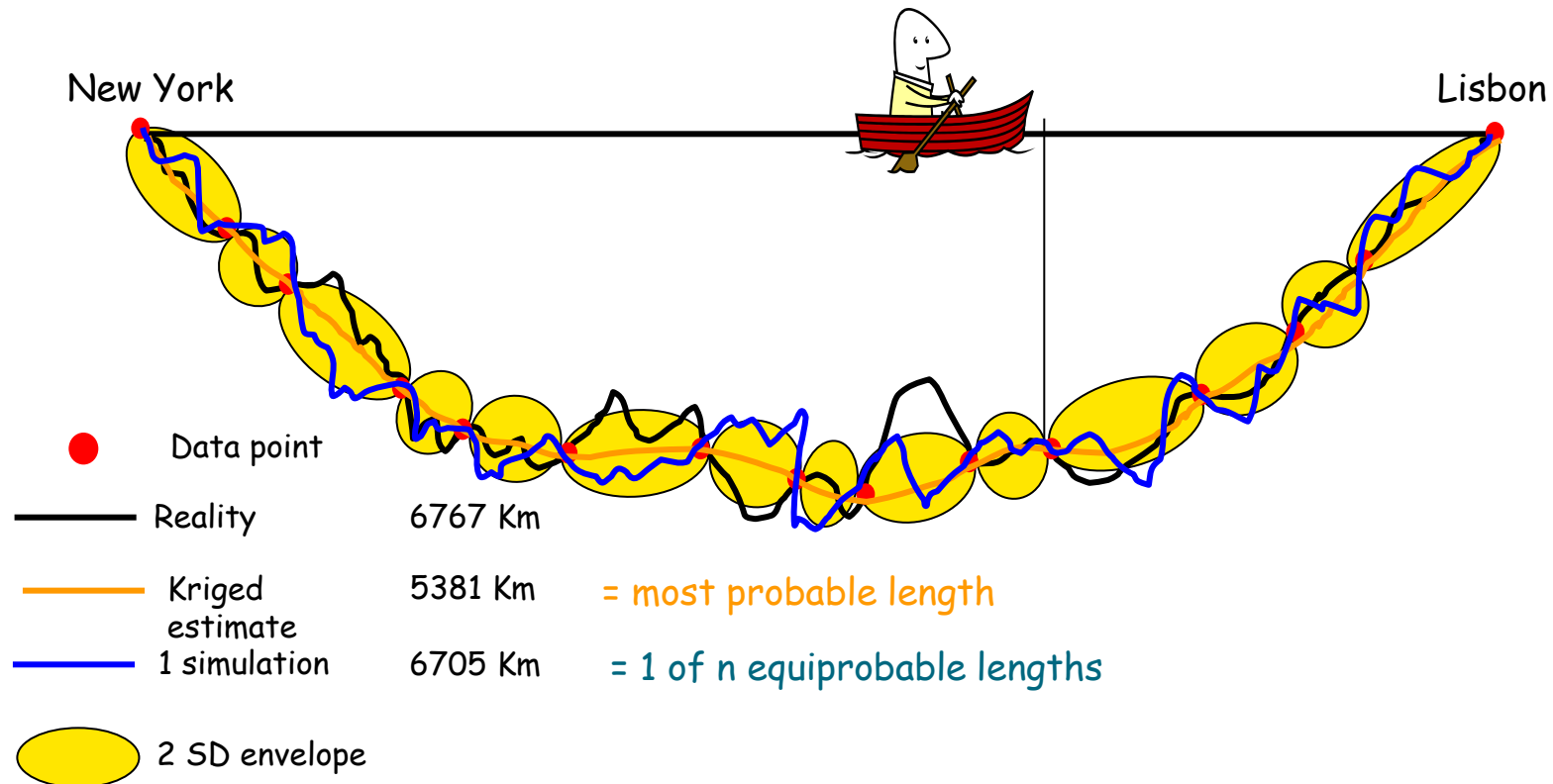
Q : metal quantity above cutoff; $Q = E[Z \mathbb{I}_{Z \geq z_c}]$



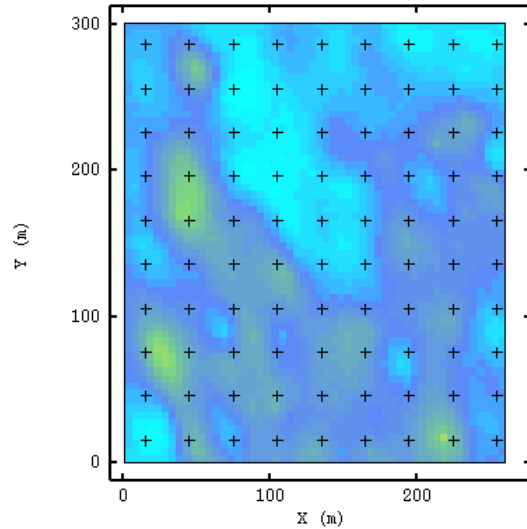
$Q = 0.0097\%$

Kriging issue

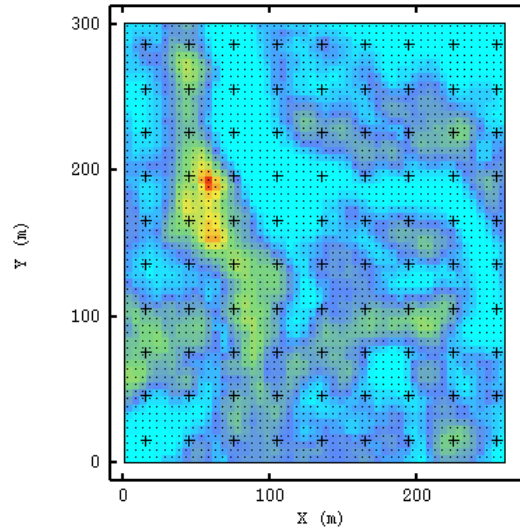
What is the length of the phone cable?



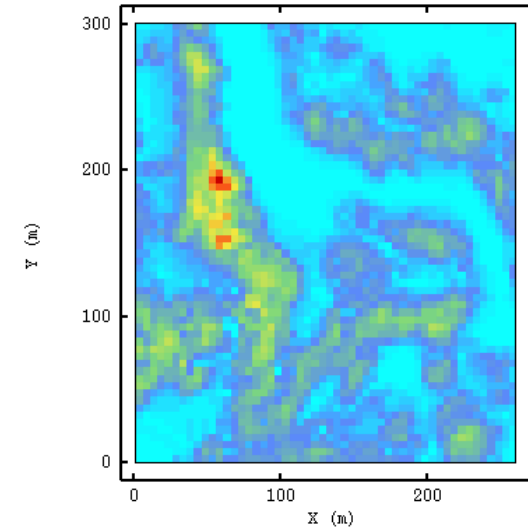
Information Effect



Exploration Data



Grade Control +
Exploration Data



Reality

An increasing level of information takes the model closer to the reality.

Unknown information leads to errors in the model.

Example of variance calculation

Isaaks & Srivastava 1989

26	27	30	33
30	34	36	32
43	41	42	46

Mean of the point value: 35

Variance of the point value :
$$= \frac{1}{12} \sum_{n=1}^{12} (value - mean)^2 = 40 \text{ ppm}^2$$

Variance of the block
$$= \frac{1}{4} \sum_{n=1}^4 (value - mean)^2 = 2.5 \text{ ppm}^2$$

33	34	36	37
----	----	----	----

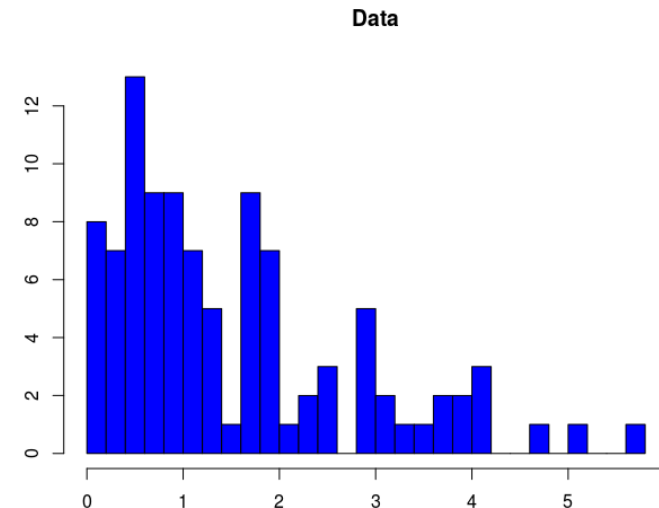
Variance of point values within blocks
$$= \frac{1}{12} \sum_{n=1}^{12} (value \text{ pts} - value \text{ block})^2 = 37.5 \text{ ppm}^2$$

Selectivity: 100 data

1.78	1.26	1.03	1.10	0.79	0.57	1.72	3.88	3.78	1.08
0.94	3.41	4.80	1.88	0.43	0.23	1.14	1.86	1.69	0.49
0.62	2.26	2.31	1.00	0.28	0.18	0.44	0.28	0.31	0.13
1.93	4.12	1.88	0.94	0.66	0.41	0.56	0.12	0.17	0.14
3.29	4.10	3.61	1.67	2.60	1.11	0.54	0.07	0.10	0.13
5.77	4.16	3.14	1.30	1.73	0.74	0.70	0.49	0.45	0.48
2.52	1.66	1.65	1.29	1.39	1.05	1.51	1.63	0.62	0.44
0.82	0.65	0.38	0.47	0.58	0.83	1.99	2.86	1.26	0.83
0.37	0.79	0.62	2.50	2.90	2.92	3.98	2.18	0.95	0.81
0.25	0.97	1.14	5.03	3.13	2.93	1.83	1.97	1.65	2.97

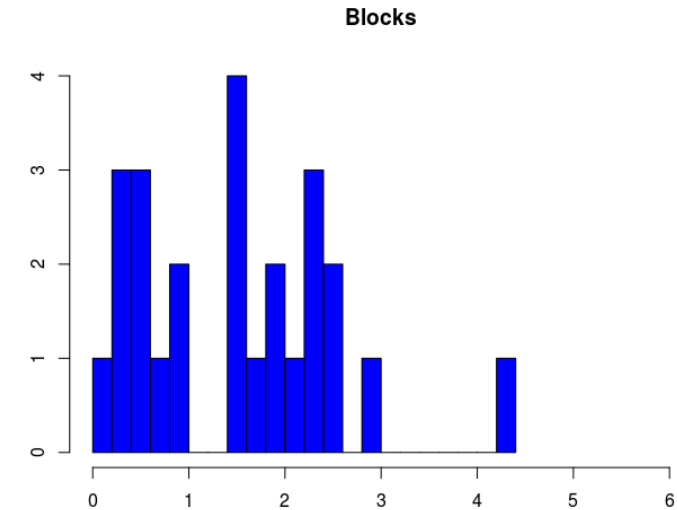
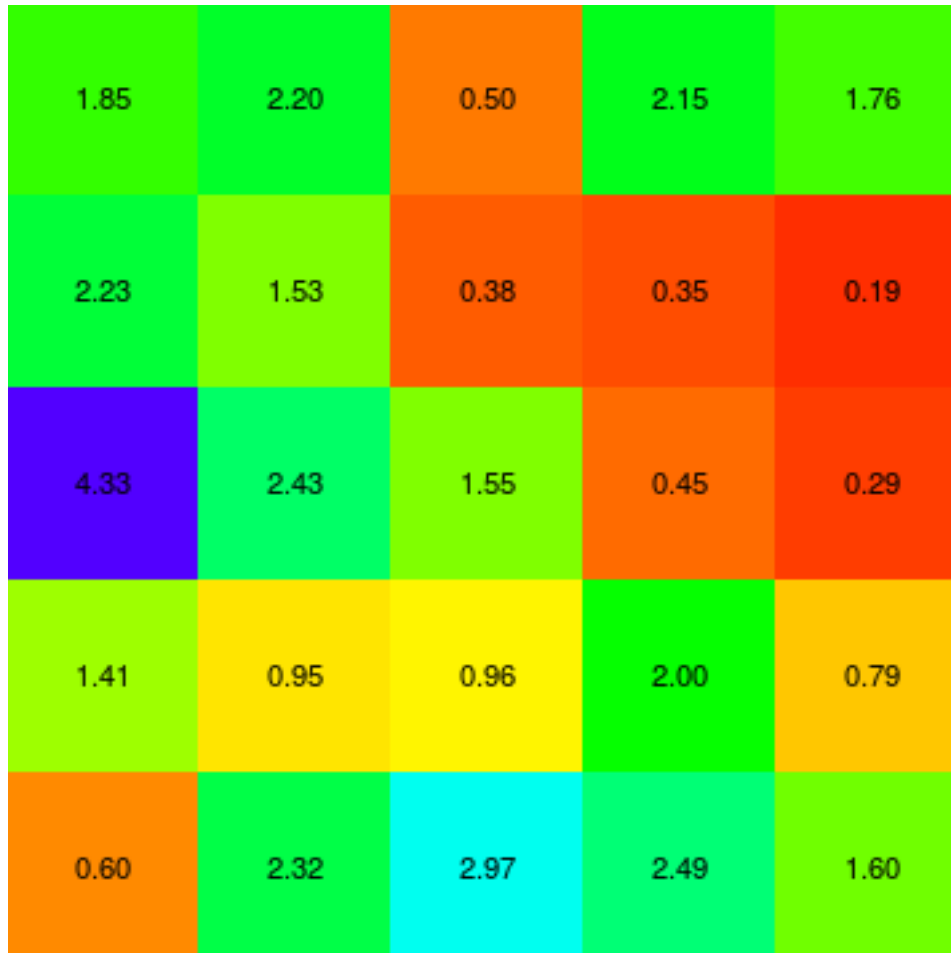
A fictitious deposit V, exhaustively known, is divided into 25 identical blocks with same tonnage (support v).

Each block is divided into 4 identical land plot or data



Mean=1.53
Variance=1.61

Selectivity: 100 data



Mean=1.53
Variance=0.97

Selectivity: variances

The mean grade of the deposit is 1.53 (%)

The dispersion variance of samples in the deposit is 1.61

The dispersion variance of blocks in the deposit is 0.97

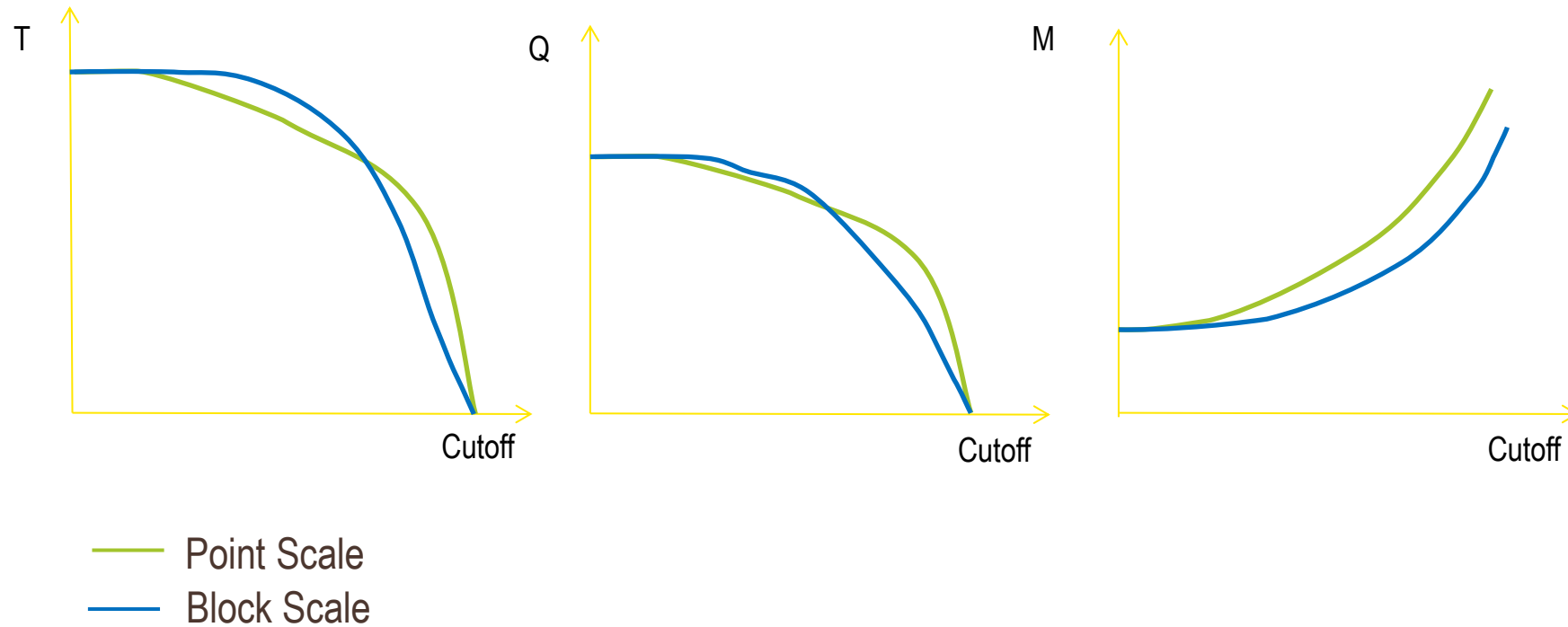
The dispersion variance of a sample within a block (mean of dispersion variance over the blocks) is 0.64

Additivity of variances:

$$1.61 = 0.64 + 0.97$$

$$s^2(O|V) = s^2(O|v) + s^2(v|V)$$

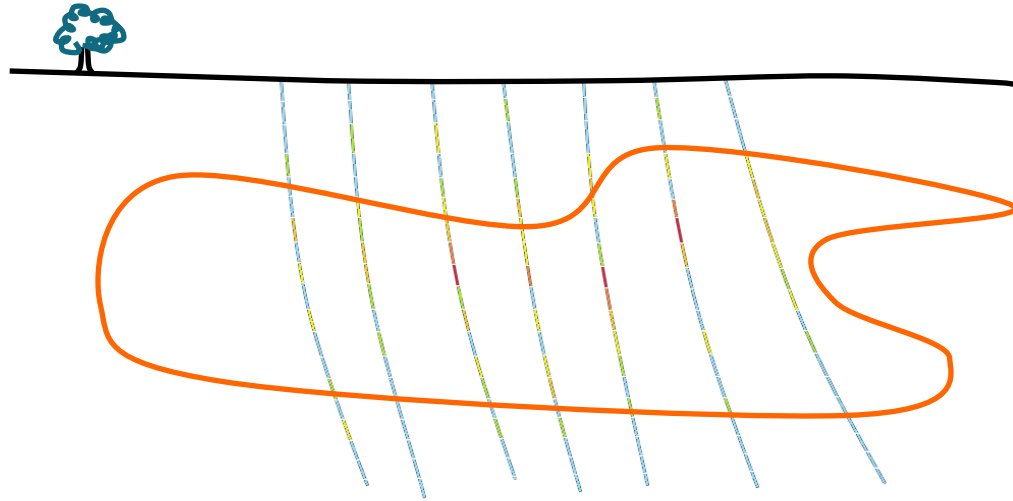
Support Effect



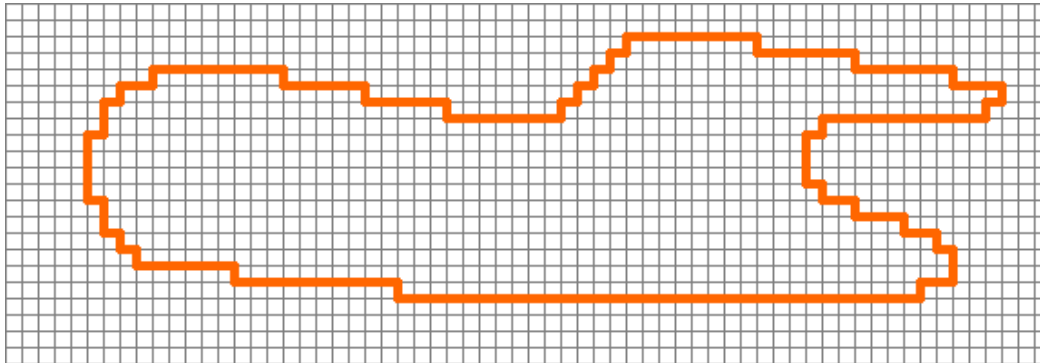
12

**Global change of
support**

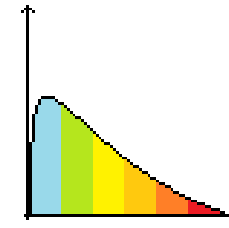
Global support change



« Point » support: Data

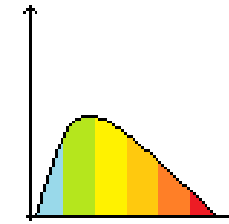


Block support: SMU



Data distribution

Change of support model



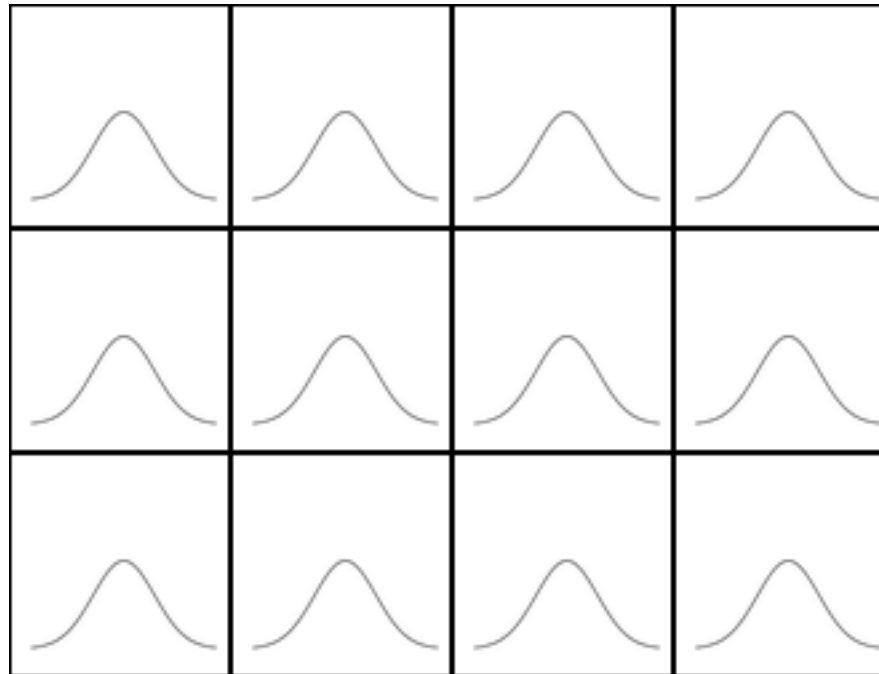
SMU distribution



Q, T, M & B

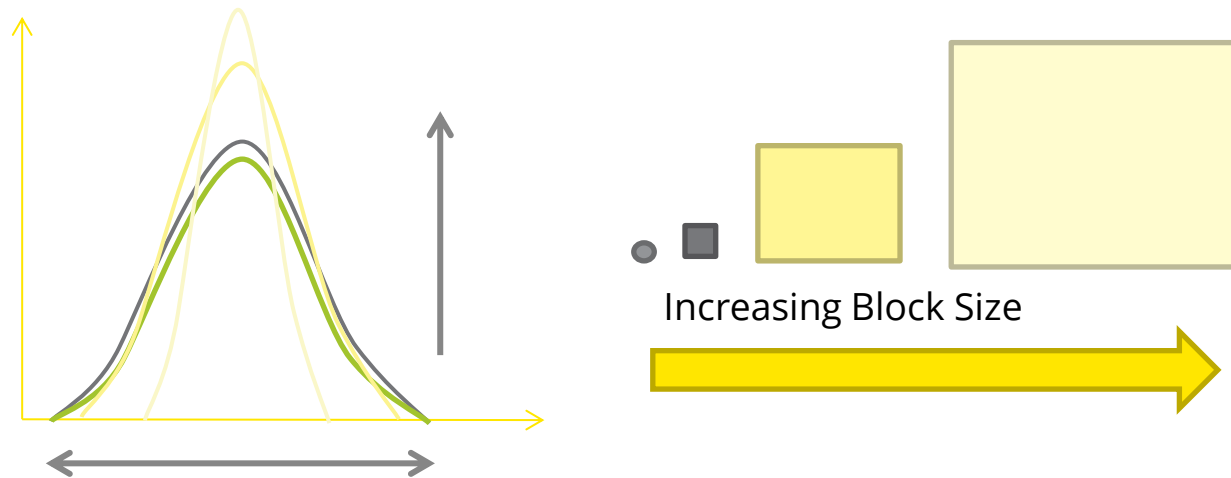
Discrete Gaussian Model

The Discrete Gaussian Model is so named because we think of a domain as consisting of a discrete number of blocks, each having a Gaussian (block) value.



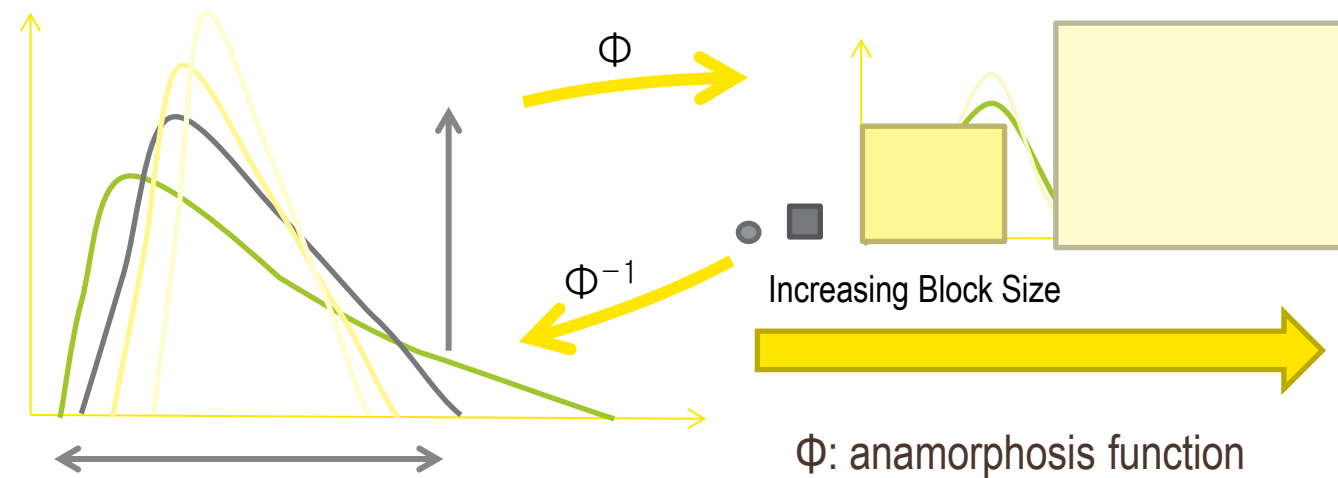
Support Effect

For a Gaussian-distributed grade ...
it is very simple:



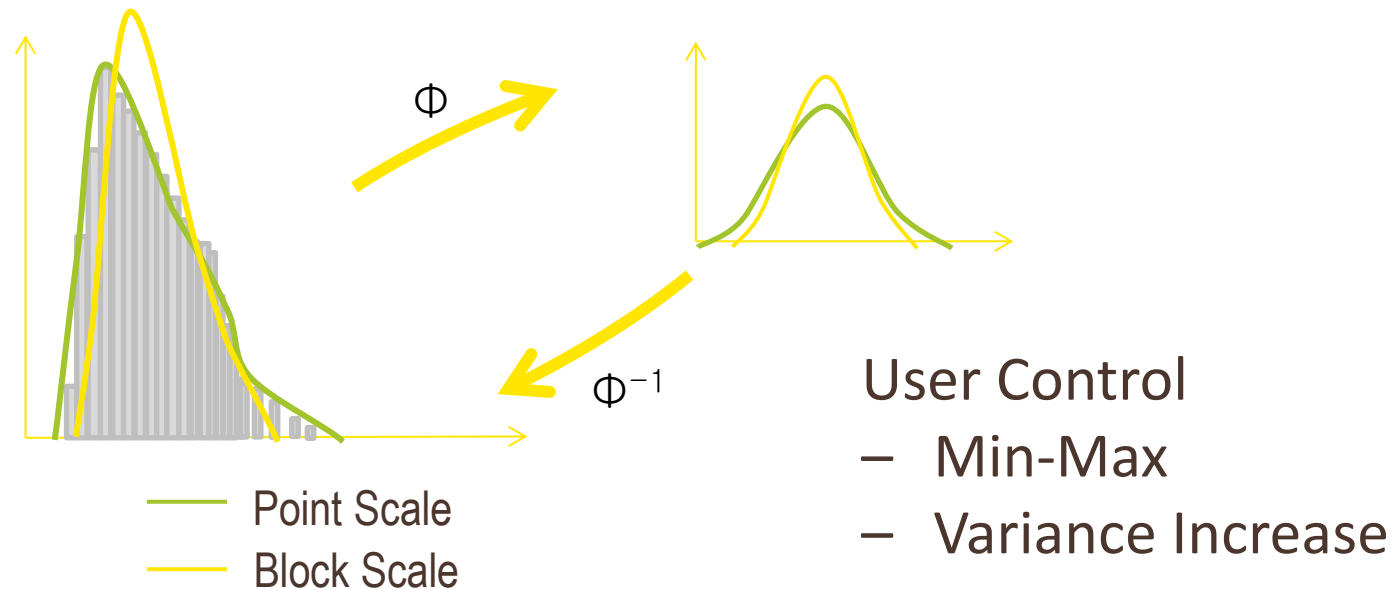
Support Effect

For a non Gaussian-distributed grade...
it is more complex



Support Effect

In practice



Gaussian anamorphosis

Much of the theory is based on converting back and forth between raw and Gaussian values, via a function called the *Gaussian anamorphosis*.

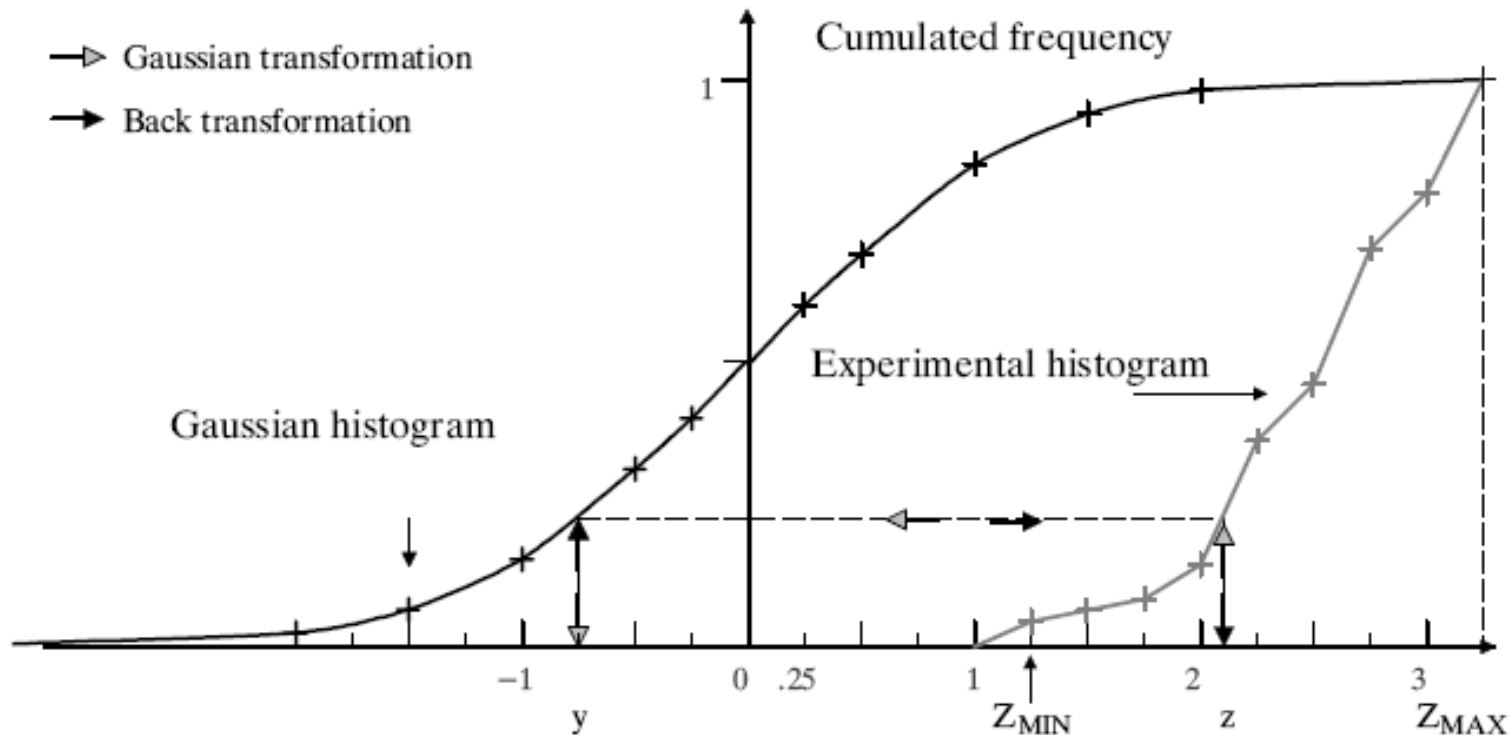
Knowing the Gaussian anamorphosis of a variable is equivalent to knowing its histogram. And if you know the histogram, then you can calculate tonnages and metal above a cutoff.

Gaussian anamorphosis

Let $G(x)$ be the standard Gaussian cumulative distribution function.

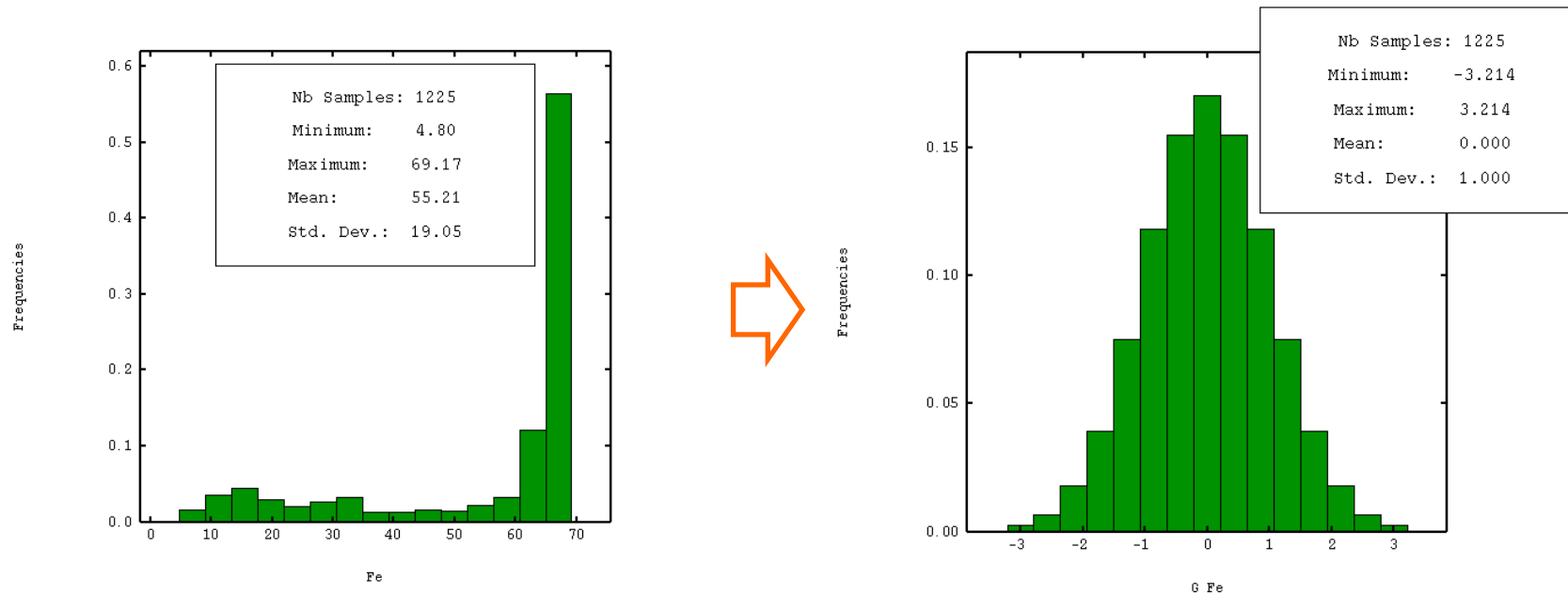
Sort the n raw values z_i in ascending order: $z_1 < z_2 < \dots < z_n$.

Then define Gaussian values $y_i = G^{-1}\left(\frac{i}{n} - \frac{1}{2n}\right)$.



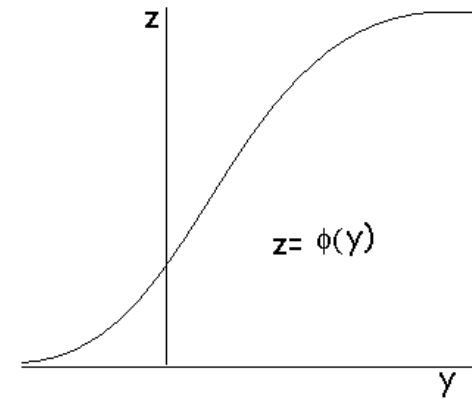
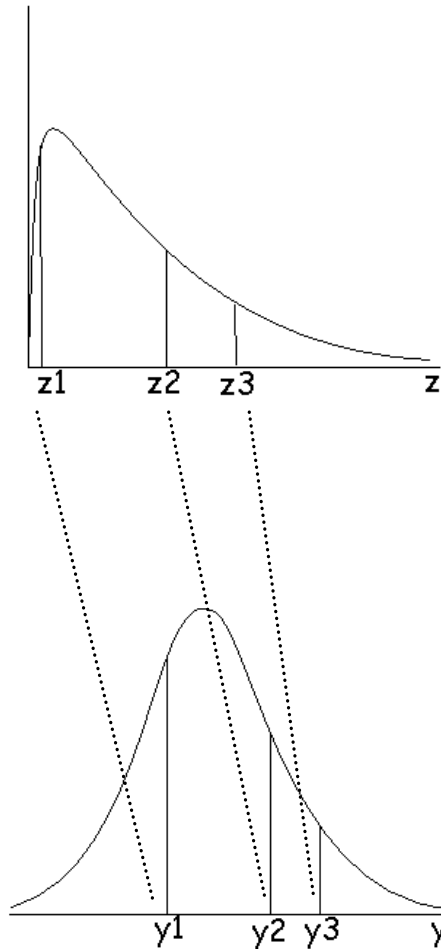
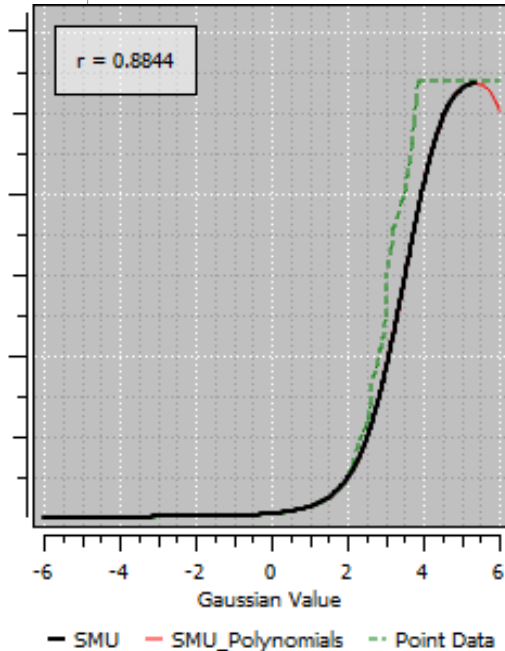
Gaussian Anamorphosis

A continuous variable Z could be considered as the transform of a standard Gaussian variable Y (mean 0, variance 1), through an anamorphosis function Φ



Gaussian anamorphosis

Interpolating between the mapped values defines an anamorphosis function Φ , so that $Z = \Phi(Y)$.



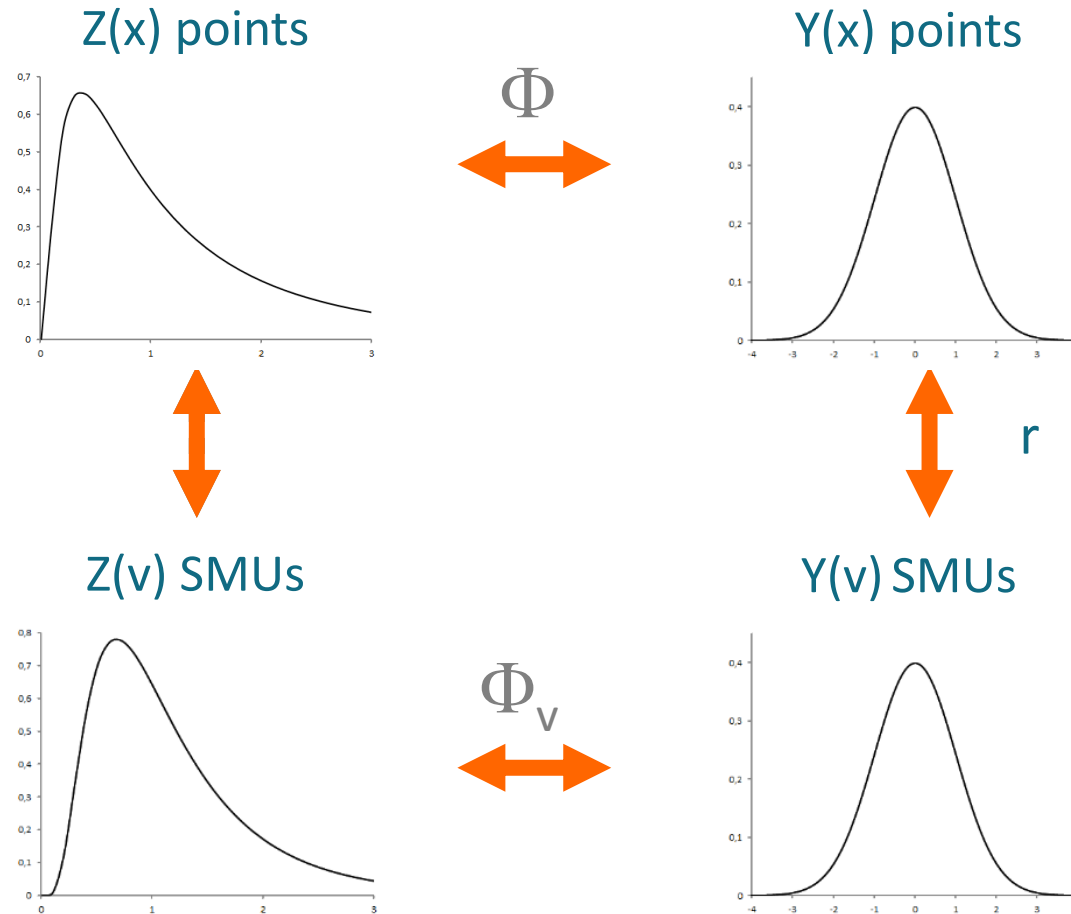
Gaussian anamorphosis

Modelling the anamorphosis function is equivalent to modelling the (full) histogram of the raw data.

Implicitly, such modelling assumes **strict stationarity**: the distribution is assumed to apply over the whole domain.

It may also be important to use declustering weights – if e.g. drilling is denser in high-grade regions, then the unweighted mean will be artificially high.

Discrete gaussian model



r characterizes the relationship between gaussian point variables and gaussian SMU variables, that is assumed to be bigaussian.

It is called the **coefficient of support change**.

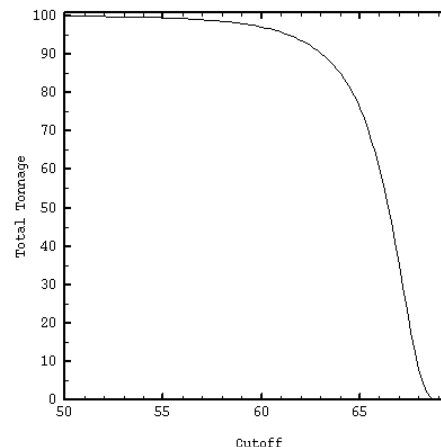
The SMU anamorphosis Φ_v can be calculated using r and the point anamorphosis Φ .

Discrete Gaussian Model

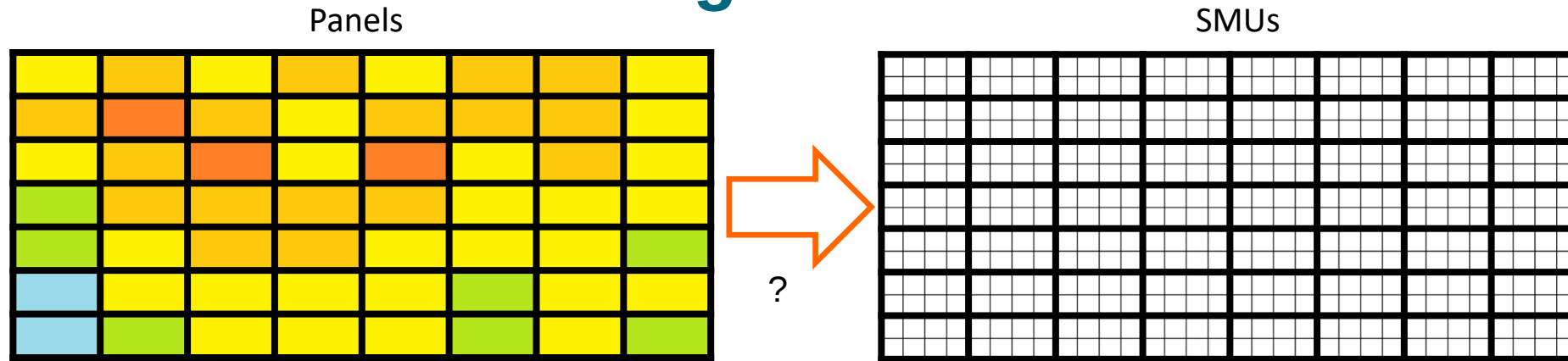
Once r has been computed, we have the block anamorphosis, which means we know the histogram of the block values.

Lots of useful quantities can be calculated from the histogram!

These statistics are all defined globally – e.g. a tonnage above a cutoff over the whole domain, rather than panel by panel.



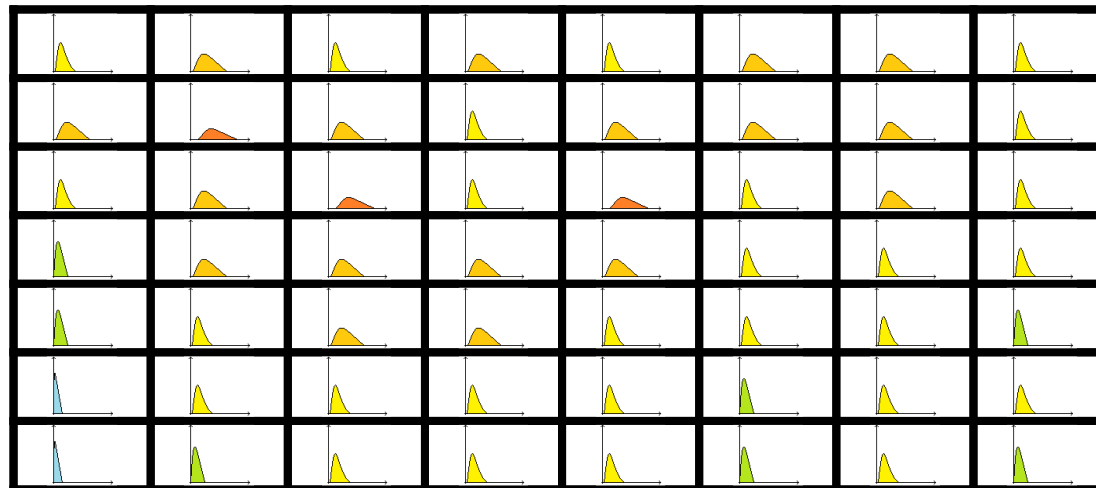
Uniform Conditioning overview



- Quality estimation
- Consistent with drillholes spacing
- Inappropriate size to calculate recoverable resources

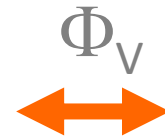
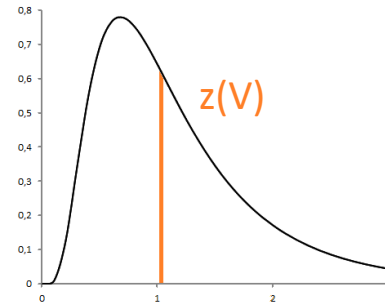
- Adequate size for recoverable resources
- Non-consistent with drillholes spacing

Use the panels
estimation to infer
the distributions

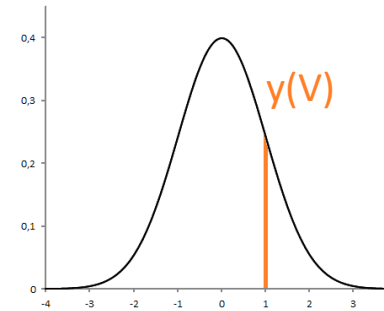


SMU distribution - Panel relationship

Z(V) panel

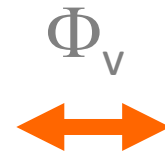
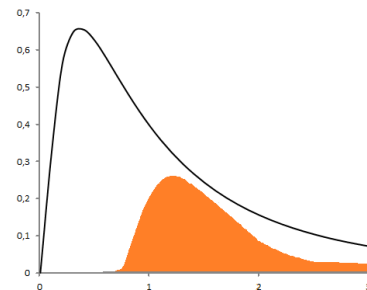


Y(V) panel

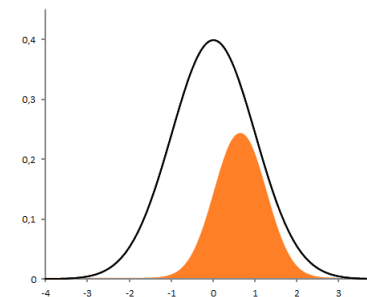


Direct link between
Gaussian distributions

Z(v) SMUs

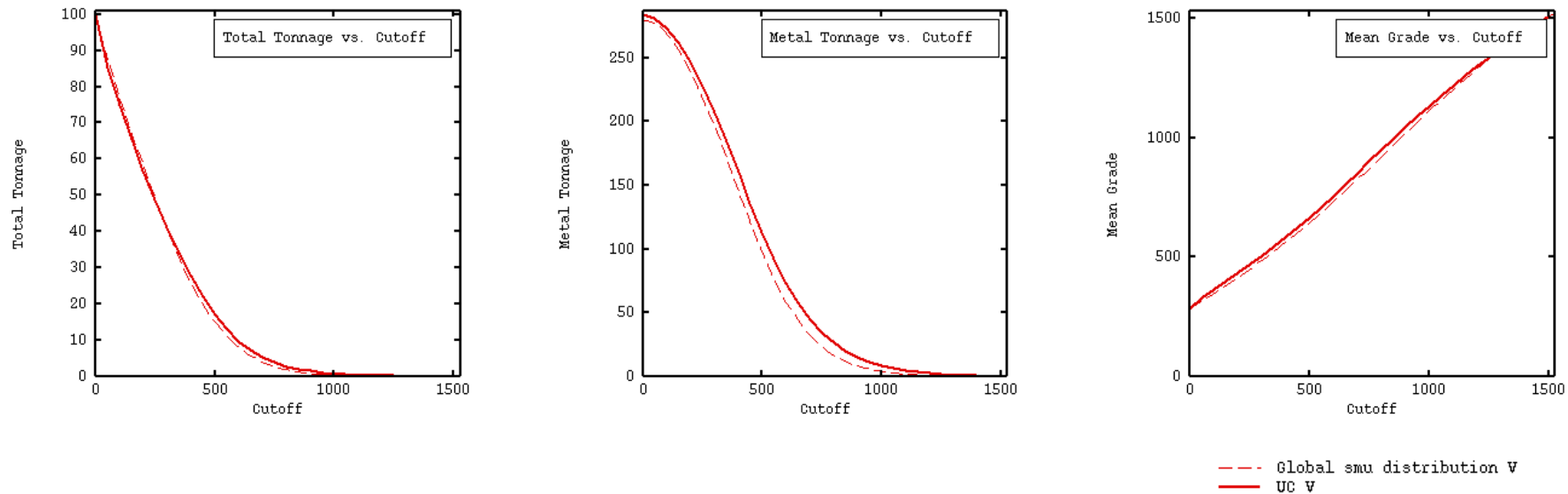


Y(v) SMUs

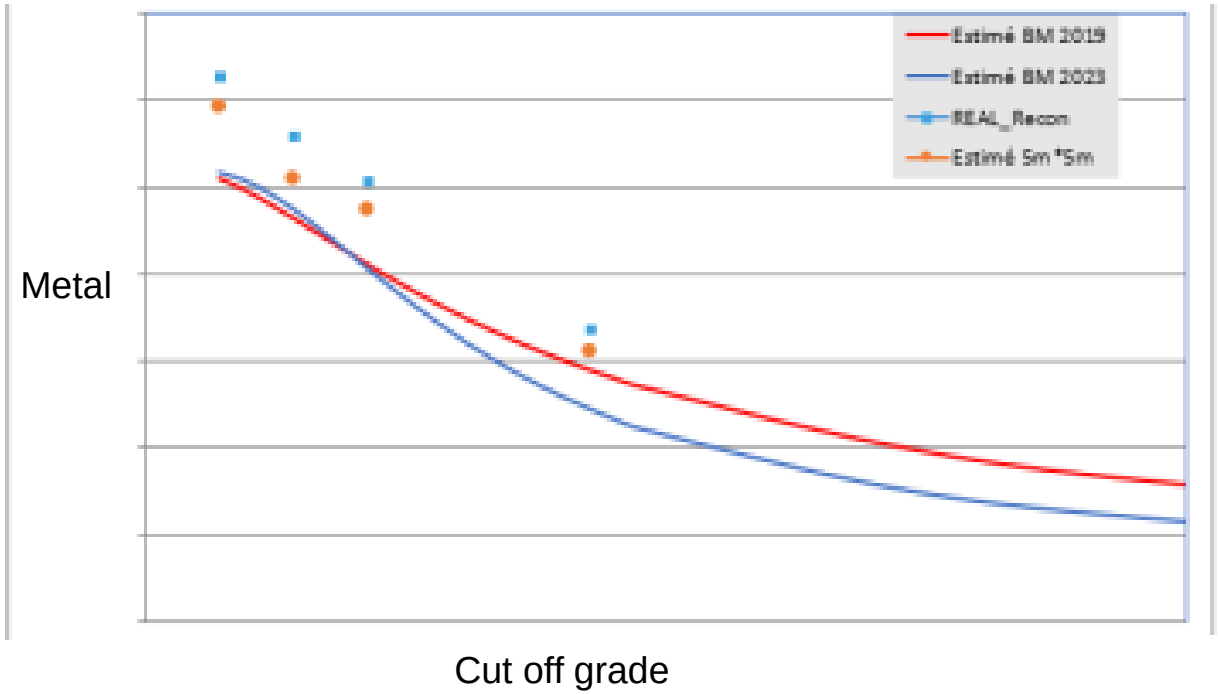
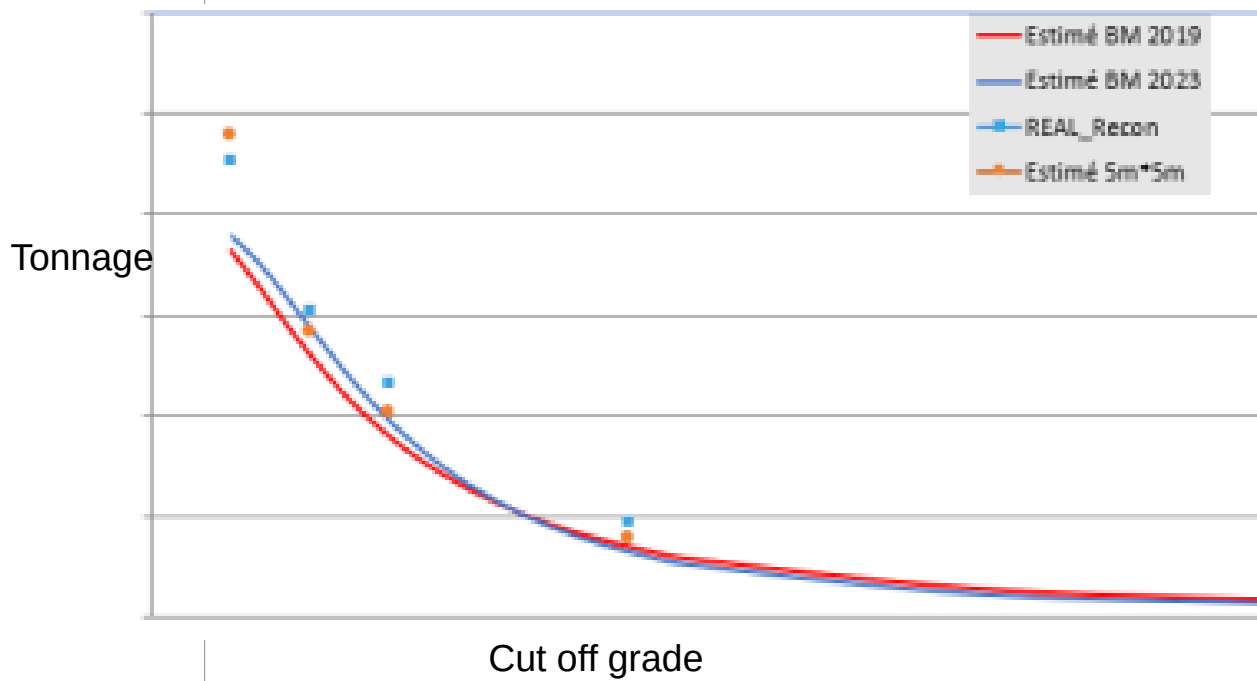


Results

Results of UC can be compared to the theoretical SMU distribution from the Global Change of Support. Their Grade Tonnage Curves should match, although some differences due to non-stationarity and clustering are acceptable.



Grade tonnage curves



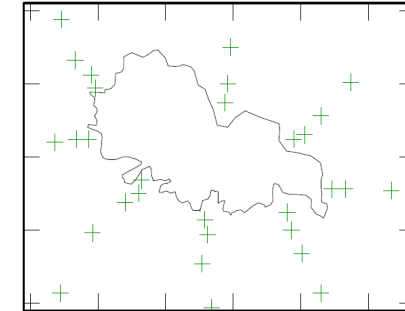
Another method: Conditional simulations

Data: Bathymetric survey of Yeu Island

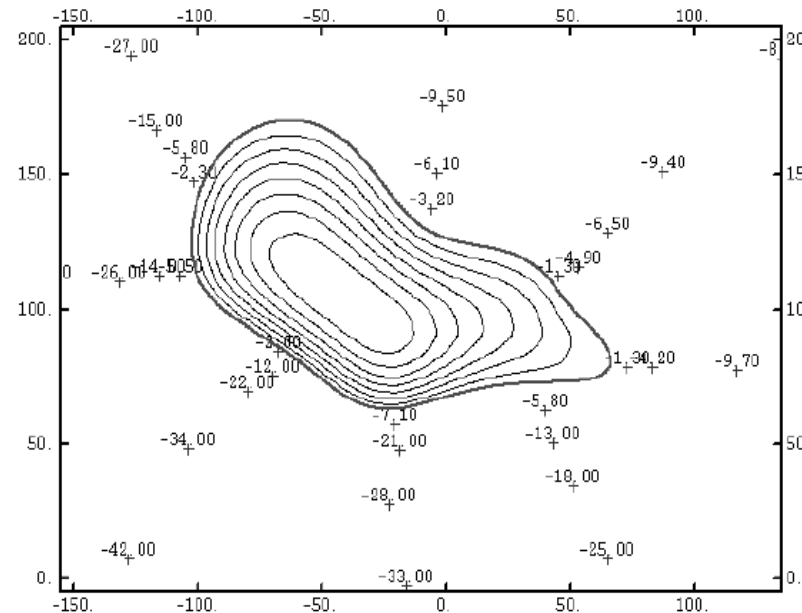
Objective: simulate surface area

Analogue with metal tonnage

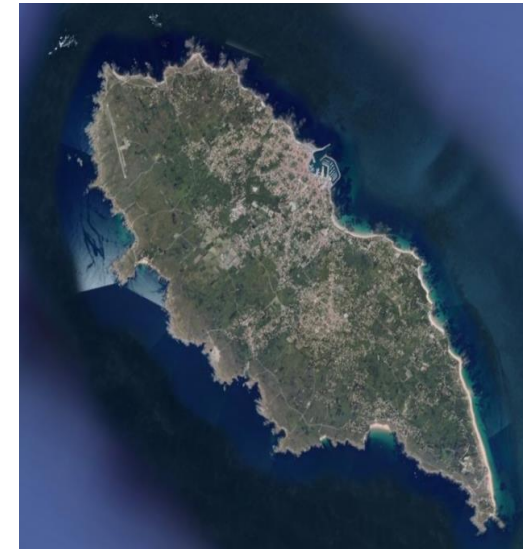
Data locations



Kriged estimate

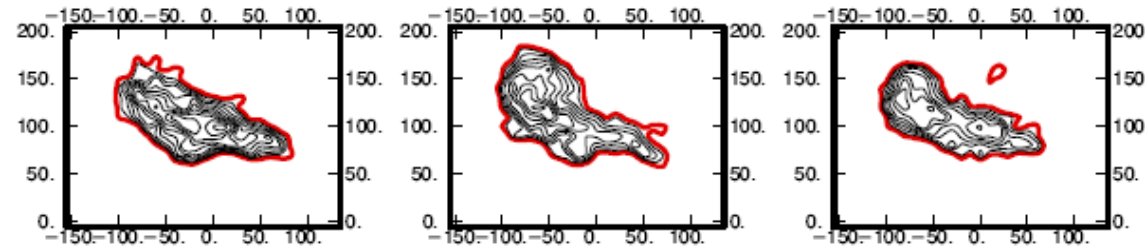


Reality



Simulation results

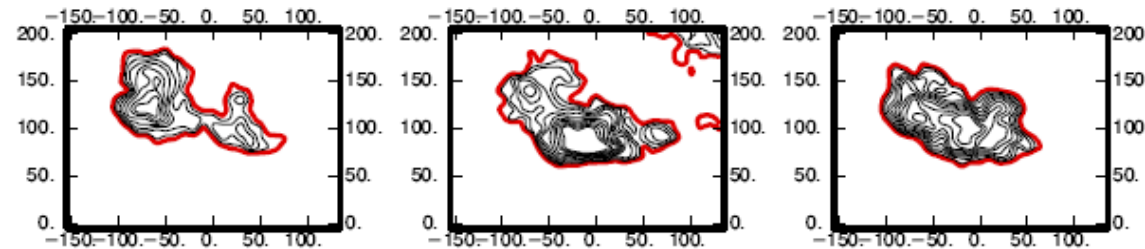
**Some of the
50
equiprobable
realizations
simulated**



Simulation #1

Simulation #2

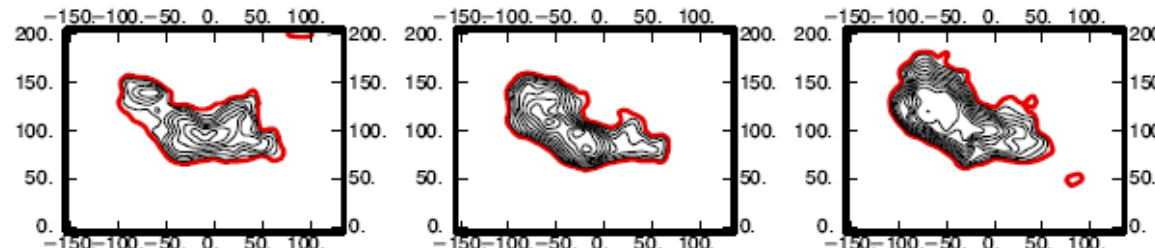
Simulation #3



Simulation #4

Simulation #5

Simulation #6



Simulation #7

Simulation #8

Simulation #9

Simulation results

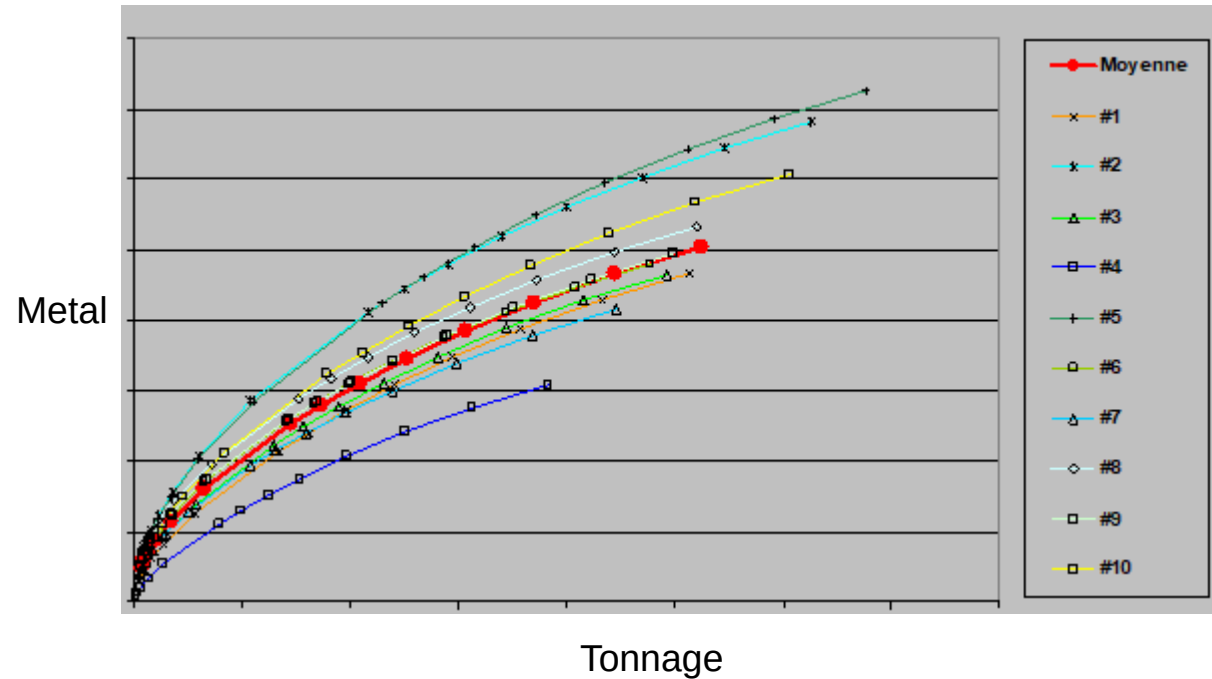


Grade simulated with in red: mean of the
simulations,
In dash line: mediane

Tonnages simulated with in red: mean of
the simulations,
In dash line: mediane

**One issue with the simulations is to be able to work with several
realisations for reserves calculation**

Grade tonnage curves



Scenario reduction algorithm can be used to select different scenario and analyse them more in details

13

Conclusions

Conclusions

Resources estimation:

- **Validated dataset**
- **Additive variable**
- **Ordinary kriging in a block scale, using a variogram model and with an appropriate cell size of the block model**
- **Need a proper validation of the estimation**

Conclusions

Multivariate geostatistics

- **Use the relationship between the variables**
- **Need a co variogram to be fitted**

Recoverable resources

- **Consider non linear variables**
- **Will require to use non linear geostatistics**



orano

Giving nuclear energy its full value