

APPLICATION OF LOCALISED UNIFORM CONDITIONING ON TWO HYPOTHETICAL DATASETS

Kathleen Marion Hansmann

A dissertation submitted to the Faculty of Engineering and the Built Environment, University of the Witwatersrand, Johannesburg, in partial fulfilment of the requirements for the degree of Master of Science in Engineering.

Johannesburg 2015

CONTENTS

CONTENTS.....	i
DECLARATION	iii
ABSTRACT.....	iv
ACKNOWLEDGEMENT.....	v
List of Figures	vi
List of Tables	viii
List of Symbols	ix
List of Abbreviations	ix
1. Introduction	1
2. Theory of Uniform Conditioning.....	3
2.1 Linear estimation techniques.....	3
2.2 Non-linear estimation techniques	4
2.3 Support.....	5
2.3.1 Scale and variance.....	6
2.4 Stationarity.....	8
2.5 Uniform conditioning procedure	8
2.5.1 Estimate panel grades.....	9
2.5.2 Normal score transform.....	10
2.5.3 Gaussian anamorphosis	11
2.5.4 Change of support.....	15
2.5.5 Q-T Curves.....	17
2.6 Localised uniform conditioning.....	19
3. Approach.....	20
3.1 Generating simulated realisations	20
3.1.1 Reference distributions for simulation	21
3.1.2 Variograms reproduction of simulated	22
3.1.3 Plot of realisations	24
3.2 Sampling realisations	26
3.3 Exploratory data analysis	27

3.4	Normal scores.....	28
3.5	Variography.....	29
3.6	Panel estimation.....	32
3.6.1	Discretization points	32
3.6.2	Measure of confidence	34
3.6.3	Plots of panel estimate	35
3.7	Bi-Gaussianity.....	36
3.7.1	Madogram variogram ratio.....	36
3.7.2	Proportion effect.....	37
3.7.3	Diffusion model.....	38
3.7.4	Results of testing for bi-Gaussianity	39
3.8	Uniform conditioning	40
3.8.1	Determining change of support coefficients	40
3.8.2	Hermite polynomials.....	43
3.8.3	Calculating conditional SMU distribution	46
3.9	Localisation.....	47
4.	Discussion on Uniform Conditioning	49
4.1	Global grade tonnage assessment	49
4.2	Panel grade tonnage assessment.....	53
4.3	Localisation assessment.....	58
5.	Conclusion.....	60
	REFERENCES.....	62
	Appendix A: Statistics Tables	64
	Appendix B: Slope of Regression	65
	Appendix C: Distribution of Estimated Grades	66
	Appendix D: Model Scatterplots.....	67
	Appendix E: Local Grade Distributions	68

DECLARATION

I declare that this dissertation is my own unaided work. It is being submitted to the Degree of Masters of Science in Engineering to the University of the Witwatersrand, Johannesburg. It has not been submitted before for any degree or examination to any other university.

_____ day of _____ year _____

ABSTRACT

Localised uniform conditioning spatially locates uniform conditioned grades, at a minable scale, within large mining panels. This non-linear estimation method is reviewed and a comparison is made between a normal and log-normal distributed synthetic deposit, to determine the preferred situation where this geostatistical approach would be applicable.

ACKNOWLEDGEMENT

I would like to thank Dr Michael Harley, for tirelessly explaining difficult and simple concepts to me, and for guidance throughout this project. I would also like to thank those who reviewed this work.

List of Figures

Figure 1 Schematic showing support effects, for point samples, SMU and panels	7
Figure 2 Localised uniform conditioning workflow	9
Figure 3 Cdf of grade $Z(x)$ to normal score Y_x transformation	11
Figure 4 First six Hermite polynomials and resultant $Z(x)$ model.....	14
Figure 5 Gaussian anamorphosis experimental data and model	14
Figure 6 Conditional SMU distribuion from a panel grade	17
Figure 7 Q-T Curve for a panel	18
Figure 8 Histogram of Scenario 1's reference distribution.....	21
Figure 9 Histogram of Scenario 2's reference distribution.....	22
Figure 10 Scenario 1 variogram from intermediate unconditional simulation	23
Figure 11 Scenario 2 variogram from intermediate unconditional simulation	23
Figure 12 Scenario 1 model and simulated variograms, showing variogram replication	24
Figure 13 Scenario 2 model and simulated variograms, showing variogram replication	24
Figure 14 Plan view of Scenario 1 simulation (at surface elevation).....	25
Figure 15 Plan view of Scenario 2 simulation (at surface elevation).....	25
Figure 16 Sampling pattern showing locations of vertical pseudo-drillholes for Scenario 1 ..	26
Figure 17 Sampling pattern showing locations of vertical pseudo-drillholes for Scenario 2 ..	27
Figure 18 Scenario 1 sample histogram and cumulative frequency histogram	28
Figure 19 Scenario 2 sample histogram and cumulative frequency histogram	28
Figure 20 Normal scores of Scenario 1 histogram and cdf.....	29
Figure 21 Normal scores of Scenario 2 histogram and cdf.....	29
Figure 22 Scenario 1 experimental variogram, model variogram and parameters	30
Figure 23 Scenario 2 experimental variogram, model variogram and parameters	30
Figure 24 Simulation variogram, for Scenario 1	31
Figure 25 Simulation variogram, for Scenario 2	32
Figure 26 Panel discretisation point analysis for Scenario 1	33
Figure 27 Panel discretisation point analysis for Scenario 2	33
Figure 28 Surface plan view of ordinary kriged panel estimate for Scenario 1.....	35
Figure 29 Surface plan view of ordinary kriged panel estimate for Scenario 2.....	36
Figure 30 Madogram / variogram ratio test for Scenario 1.....	37
Figure 31 Madogram / variogram ratio test for Scenario 2.....	37
Figure 32 Test for proportional effect for Scenario 1 original grades and normal scores	38
Figure 33 Test for proportional effect for Scenario 2 original grades and normal scores	38
Figure 34 Ratio of extreme indicators to test for intrinsic correlation for Scenario 1	39
Figure 35 Ratio of extreme indicators to test for intrinsic correlation for Scenario 2	39
Figure 36 Scenario 1 grouping of panel estimates for R coefficient determination.....	41
Figure 37 Scenario 2 grouping of panel estimates for R coefficient determination.....	42
Figure 38 Scenario 1 sample normal score transform and anamorphosis model.....	44

Figure 39 Scenario 2 normal score transform and anamorphosis model	45
Figure 40 Gaussian anamorphosis model for multiple R/r ratios.....	45
Figure 41 Q-T curve showing average grade calculation at tonnage increment.....	47
Figure 42 Scenario 1 global grade tonnage curves.....	50
Figure 43 Scenario 2 global grade tonnage curves.....	51
Figure 44 Scenario 1 recoverability curve.....	51
Figure 45 Scenario 2 recoverability curve.....	52
Figure 46 The Kriging Oxymoron explaining linear estimation into small versus large block.....	53
Figure 47 Actual versus OK panel scatter plot, showing blocks for LUC analysis.....	54
Figure 48 Grade tonnage curve for panel 1422 (Scenario 1).....	54
Figure 49 Grade tonnage curve for panel 537 (Scenario 1).....	55
Figure 50 Grade tonnage curve for panel 1398 (Scenario 1).....	55
Figure 51 Grade tonnage curve for panel 146 (Scenario 1).....	56
Figure 52 Grade tonnage curve for panel 1216 (Scenario 2).....	56
Figure 53 Grade tonnage curve for panel 917 (Scenario 2).....	57
Figure 54 Grade tonnage curve for panel 632 (Scenario 2).....	57
Figure 55 SMU within panel ranking comparison for Scenario 1 and Scenario 2	58
Figure 56 Slopes of regression for Scenario 1's OK model	65
Figure 57 Slopes of regression for Scenario 2's OK model	65
Figure 58 Global distribution of grades for Scenario 1	66
Figure 59 Global distribution of grades for Scenario 2	66
Figure 60 Scatter plot for actual versus LUC and actual versus OK panel for Scenario 1	67
Figure 61 Scatter plot for actual versus LUC and actual versus OK panel for Scenario 2	67
Figure 62 Normal distribution (Scenario 1) – well estimated panels	68
Figure 63 Normal distribution (Scenario 1) – poorly estimated panels	68
Figure 64 Log-normal distribution (Scenario 2) – well estimated panels.....	68
Figure 65 Log-normal distribution (Scenario 2) – poorly estimated panels	68

List of Tables

Table 1 Descriptive statistics for Scenario 1 and Scenario 2	27
Table 2 Panel block model dimensions.....	32
Table 3 SMU change of support coefficients for calculating ' r '	40
Table 4 Change of support coefficients for Scenario 1.....	42
Table 5 Change of support coefficients for Scenario 2.....	42
Table 6 Hermite coefficients.....	43
Table 7 Example UC model as grades and proportions	46
Table 8 Statistics for Scenario 1 at multiple supports	64
Table 9 Statistics for Scenario 2 at multiple supports	64

List of Symbols

Average co-variance between panel, of support V ("c bar $V V$ ")	$\bar{C}(V, V)$
Average co-variance between SMU, of support v ("c bar $v v$ ")	$\bar{C}(v, v)$
Average variance within a panel, of support V ("gamma bar $V V$ ")	$\bar{\gamma}(V, V)$
Average variance within an SMU, of support v ("gamma bar $v v$ ")	$\bar{\gamma}(v, v)$
Change of support coefficient, for panel	R
Change of support coefficient, for SMU	r
Estimated grade at point location x	$Z^*(x)$
Estimated grade at support V	$Z^*(V)$
Estimated grade at support v	$Z^*(v)$
Gaussian-equivalent estimated grade at location x	$Y^*(x)$
Gaussian-equivalent unknown grade at location x	$Y(x)$
Hermite coefficient ("phi")	ϕ
Hermite polynomial	H
Lagrange multiplier	μ
Large block support i.e. panel	V
Sample variance ("sigma squared")	σ^2
Small block support i.e. SMU	v
Unknown grade at point location x	$Z(x)$
Variogram value at distance h ("gamma h ")	$\gamma(h)$
Weight ("lambda")	λ

List of Abbreviations

Block variance	BV
Coefficient of variation	CoV
Cumulative distribution function	cdf
Discrete Gaussian model	DGM
Disjunctive kriging	DK
Grade tonnage	GT
Grams per ton	g/t
Indicator kriging	IK
Inverse power of distance	IPD
Localised uniform conditioning	LUC
Ordinary kriging	OK
Probability density function	pdf
Selective mining unit	SMU
Simple kriging	SK
Uniform conditioning	UC
Quantitate kriging neighbourhood analysis	QKNA

1. Introduction

Uniform Conditioning (UC) models the conditional distribution of grades of selective mining unit (SMU) support within large blocks (panels), and generates a recoverable resource estimate at an SMU scale. UC is a widely used, but often poorly understood, non-linear estimation technique (Neufeld, 2005). This non-linear method of estimation seeks to address some of the shortfalls of linear estimation, by using a method of change of support to indirectly model a distribution of small block grades instead of directly modelling small blocks themselves.

The conditional grade distribution modelled by UC has unknown locations of ore and waste, and localising the UC estimate places these recoverable small block grades at plausible locations within the large panel. Localised Uniform Conditioning (LUC) is an enhancement to UC which does not change the results of UC, but rather presents the recoverable resource into a more practical format for mine planning.

UC estimates a recoverable resource (tonnes and grade) above multiple cut-off grades in a panel. A recoverable resource is a proportion of an in-situ mineral resource recoverable during mining due to some constraint which could be technology, mining method, machinery or cut-off grade (Vann and Guibal, 1998). With respect to UC, the constraint manifests as the cut-off grade.

The advantage of UC is that it can be used on widely spaced data, across domains that are not strictly stationary, provided that the data is sufficient for a conditionally unbiased estimate of the panel mean grade. A successful UC result relies on a conditionally unbiased panel estimate (Rivoirard, 1994; Vann and Guibal 1998; Assibey-Bonsu, 1998; Assibey-Bonsu and Krige, 1999; Vann et al, 2000; De-Vitry et al., 2007; Deraisme et al, 2008; Deraisme and Assibey-Bonsu, 2011). New developments to the UC procedure means that it also accounts for the information effect (Deraisme and Roth, 2000).

Matheron (1974) made the first reference to UC, after which the technique was first fully described by Remacre (1984) in his PhD thesis. Both of these works are written in French. The first English publications on UC were by Armstrong and Matheron (1986) and Rivoirard (1994), who discuss UC in their papers on disjunctive kriging (DK). Many principles relating to the two methodologies are the same, such as the use of the discrete Gaussian change of support model. Subsequently, UC has been comprehensively described in numerous conference papers and journals.

This project will describe the theory and practice of UC and LUC, through reviewing existing literature, presenting worked examples of UC and LUC estimations, followed by a discussion of the results. Two 'simulated realities', one fairly continuous normal grade distribution and another with short range continuity and a highly skewed log-normal grade distribution, are

compared to determine how well UC and LUC fares in these two extreme cases. Depending on the underlying statistical distribution of the grade data, applying the same estimation method may produce more favourable results in one case than in the other, and this project seeks to understand what types of underlying distributions are more effectively modelled by UC.

2. Theory of Uniform Conditioning

In this chapter, UC and LUC will be discussed from a theoretical perspective. Existing literature written on the topic will also be reviewed.

There are various techniques available for estimating the grade of mineral deposits, which give similar results despite fundamentally different estimation methodologies. Broadly categorised, these techniques are either linear or non-linear estimation techniques. For comparative reasons, a brief overview of the benefits and shortfalls of linear estimation techniques will be given, followed by a similar overview for non-linear techniques, focusing on UC.

2.1 Linear estimation techniques

Linear estimation involves a direct weighted averaging of the surrounding data to calculate the grade at an unknown point or block. This requires an algorithm to calculate appropriate weighting of the data using this averaging technique.

Linear estimation techniques include, but are not limited to, the commonly used inverse power of distance (IPD) and ordinary kriging (OK), where estimates for a point (or block) are calculated from a weighted average of local samples. IPD calculates sample weights based on relative distances between samples and targets. OK calculates weights based on the variogram and the relative spatial arrangement of samples, and the OK algorithm is constructed to minimise the estimation variance. The assignment of weights to samples is independent of the grades of those samples (Vann et al, 2000).

The panel is estimated using the linear OK estimator, shown below:

$$Z^*(x) = \sum_{i=1}^n \lambda_i(x) \cdot Z_i(x)$$
$$\sum_{i=1}^n \lambda_i = 1$$

Where:

$Z^*(x)$ Grade of unknown/estimated point at location x

n Number of samples used in estimating the block

λ Weighting assigned to each sample

Z_i Grade of each sample

The calculation of weights is determined from the covariance between the unknown target (to be estimated) and each sample, as well as the covariance between the samples themselves. The covariance between points is calculated from the modelled variogram, derived from the experimental variogram of the sample data. The solving of these equations is done through solving a system of kriging equations presented as matrices. This takes place automatically in numerous commercial software packages, and the detail of this is not relevant for this example.

The arguments for using OK are well documented in literature, and it is a widely accepted and robust mineral resource estimation technique. The unbiased and spatial assignment of weights leads to the best possible average estimator of a block.

The shortfalls of linear estimation techniques are well understood and documented. These include, in no particular order:

- *Influence of outliers*: The unknown point to be estimated is dependent on samples in the immediate neighbourhood. The effect of this is that a single high valued sample (falling in the extreme “tail” of a distribution) could cause overestimation of locations surrounding this outlier point. Additionally, outliers have an effect on the variogram (increases the nugget effect and sill) and therefore effects the correlation of estimated samples.
- *Conditional bias*: The estimation of the deposit may result in systematically overestimating low grades and under-estimating high grades (or, visa-versa). This effect is seen in OK because of grade smoothing, leading to biases in the grade-tonnage curves and possible miscalculation of resources at grade/cut-off limits. One may find that, on average, the estimated grade values are higher than the “actual”, and the low grade estimates systematically lower than the “actual”. This commonly occurs when the search neighbourhood is ill defined, and can be addressed by optimising the kriging neighbourhood, in particular the slope of regression (Vann et al., 2003).
- *Block sizes*: Block size, relative to the sample spacing, chosen for a kriged estimate can influence the resultant model. Using large blocks, one can overcome the high estimation variance (or precision error) seen in smaller blocks, however this introduces grade smoothing and de-skewing of the underlying distribution. Using small blocks results in a better distribution of grades, but the estimation error can be high leading to imprecise estimates and distorted grade tonnage curves (Vann and Guibal, 1998).

2.2 Non-linear estimation techniques

There are non-linear estimation techniques that attempt to individually address the issues raised above. These techniques include, but are not limited to, log-normal kriging, indicator

kriging (IK), DK and UC. Log-normal kriging attempts to estimate in the presence of a highly skewed distribution; while the use of indicator based methods have been proposed to handle the presence of outliers. DK is powerful but highly complicated method which fully co-kriges multiple indicator grade classes, and this method has been simplified into several derivation methods including UC. UC incorporates a change of support model to condition an estimate from a panel, where a stable and robust estimate can be achieved, into a known distribution for smaller grade blocks.

In UC there are two block sizes of interest: the SMU and the panel. The SMU is considered for the mining extraction of the mineralisation and the panel, or larger volume, is estimated with a linear method. The distribution of SMU grades for a panel is determined, from the change of support model, for a panel grade. As such, an average panel grade is conditioned to a series of cut-off grades and produces estimates of tonnes and grades above these cut-offs i.e. a recoverable resource above cut-off.

2.3 Support

Support refers to the size and shape of the volume that a grade value represents, whether it is the volume of a core sample or the volume of a large mining panel. It is important to be cognisant of the relationship between the volume of a sample in relation to the volume of the SMU, panel or deposit that these samples will be used to estimate. The following example demonstrates the relationship:

Assume a standard HQ core (63.5 mm) sample of 1 m in length. For this example it is assumed that the sample itself is representative, and the core has been split, prepared and sampled to represent the entire volume of the core. This sample represents a volume of 0.0032 m³. Geostatistically, this volume is considered as a point grade.

Considering an SMU block with a size of 10 m x 10 m x 10 m, means that there could be 315 764 unique core samples in this volume. Within a panel (with a size of 50 m x 50 m x 20 m), there could be upwards of 15 million unique core samples in this block.

Of course, it is not necessary or economically viable to sample the entire block to obtain an accurate representation of the mean grade and distribution of the population. A representative subset of the population can describe the statistical properties and be used to characterise the entire population.

The purpose of this demonstration is to emphasise the change in variance and the scale relationship between the support samples taken, and the support of the blocks that are modelled and ultimately mined. The scale at which the data is collected is small, and this is not representative at a larger scale i.e. SMU or panel scale. If it were possible to selectively mine a deposit using core sized shovels, then considering the support at different scales would

not be required. This would allow one to separate tiny volumes of ore from waste material, selectively targeting the very high grade areas.

An SMU represents the smallest volume of ore that logically can be removed as a single unit, as reflected by the selectivity of the mining equipment. Vann and Guibal (1998) define an SMU in geostatistical terms as the minimum support on which mining decisions can be made. In an open pit scenario this is determined by the size of shovels and trucks. In a narrow tabular or a massive underground scenario this is based, in part, on the dimensions of the ore body and geotechnical constraints. The material in an SMU is selectively sampled, usually through core samples as described above. Since the ore body within an SMU cannot be selectively mined, it should be evaluated as a single parcel, despite the support of the samples within it. Sporadic high grade areas within an SMU could exist, but these volumes are too small to separate, and mining through lower grade material would be required to get to these high grades.

A panel can contain several to several dozen SMU, and does not need to be mined as a single unit as each SMU can be individually mined. A panel grade is often expressed as a mean grade for long-term planning purposes, but this does not provide any information about the distribution of grades within the panel. Ideally, a panel should be expressed as a distribution of grades which is the average grade of the SMU within the panel.

In summary, point samples with tiny volumes are used to estimate the grade of the SMU, within even larger panels. The distributions of grade at different supports changes, and this is discussed in the following section.

2.3.1 Scale and variance

The mineral resource estimates of a deposit are reliant on the scale at which that deposit is to be mined. If it were possible to selectively mine a deposit at the scale at which that deposit was sampled, then there would be no requirement for change of support. However, since a deposit is mined at the scale of the SMU, then we should represent the mineral resources at this scale. As the scale of the deposit increases, from point to SMU to panel, the mean remains constant while the shape of the histogram becomes more symmetric, deskewed and there is a reduction in the variance (Figure 1).

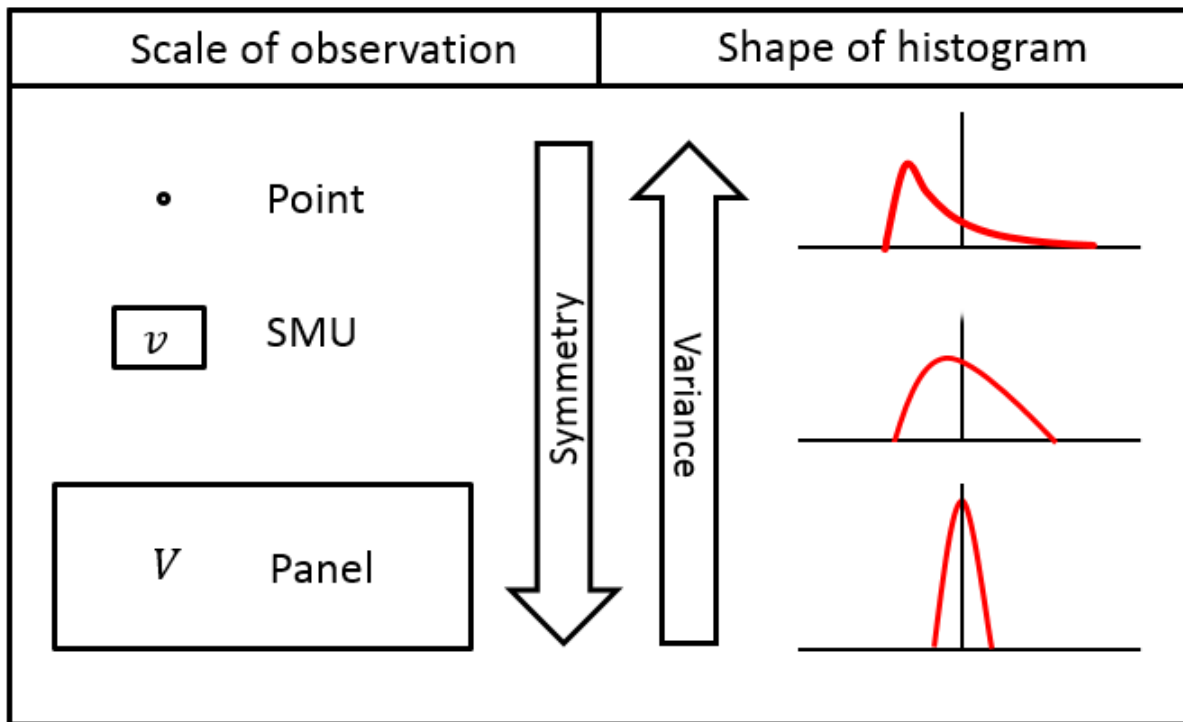


Figure 1 Schematic showing support effects, for point samples, SMU and panels

The relationship between variances is summarised by Krige's relationship, which describes the additivity of variances. There are adaptations to the relationship, but the general formula is:

$$\sigma^2 = \bar{C}(V, V) + \bar{\gamma}(V, V)$$

Where:

- | | |
|----------------------|--|
| σ^2 | Variance of points in deposit |
| $\bar{C}(V, V)$ | Average co-variance of blocks in the deposit |
| $\bar{\gamma}(V, V)$ | Average variance of points in a block |

To explain these further, consider an unknown grade deposit, which has been sampled at closely spaced grid throughout the deposit. The following definitions apply:

- *Point variance*: The point variance is the total variance of the population of point support. Since it is impossible to sample every location in a naturally occurring deposit, the best estimator of the point variance is the variance of the samples. This is the equivalent of the sill of the variogram $\gamma(h)$.
- *SMU variance*: The average variance of the grade of the samples within an SMU e.g. a 10 m x 10m x 10 m block. This value cannot be directly calculated, and a theoretical value is calculated from the average variance (based on the variogram) between the discretisation points with a SMU.
- *Panel variance*: This average variance of mean SMU block grades, within each panel of the deposit e.g. 50 m x 50 m x 20 m block.

Change of support models have been developed which quantify the differences seen in the distribution of the data at different scales. One well established change of support model is the discrete Gaussian model (DGM), as this model provides a good prediction of the mean and variance and has no upper/lower bounding limit (Rivoirard, 1994). The DGM is explored in detail further on in this chapter.

2.4 Stationarity

Neufeld (2005) highlights that domains used for UC should be stationary. This is highlighted through an example where two panels could have identical mean grades but one could contain a highly variable distribution of grades and the other a homogenous grade distribution. Although only local samples would contribute towards the mean grade, all samples within the domain would affect the change of support used.

Some literature states that UC, where compared with DK or other methods that impose strict stationarity, is adaptable to situations where stationarity is not very good (Rivoirard, 1994; Vann, 1998.). This adaptability of UC stems from using a locally varying mean when estimating the panels through OK.

This relaxation of strict stationarity should not be confused with a complete disregard of stationarity. UC is consistent with a diffuse model, where grades transition from one domain to another. Domains should be imposed where there are changes in the style of mineralisation or the variability of the grades.

Another methods of overcoming weak stationarity is to group panels with similar grade variability, and apply a change of support models to these groups. This method is discussed in detail further in this project.

2.5 Uniform conditioning procedure

UC is a non-linear estimation approach that considers a discrete Gaussian change of support model to condition panel grades to a local SMU grade distribution. Alternative change of support approaches, like the log-normal shortcut (David et al., 1977) or the affine correction, are also available, but are not the focus of this study.

A large block can be reasonably well estimated using a linear estimator. It is commonly understood that estimating small blocks relative to the drill spacing is imprecise with linear estimation techniques, particularly as the nugget effect increases (Vann et al., 2000). In such cases the kriging variance (error) will be higher, leading to reduced kriging efficiencies and reduced slopes of regression, resulting in the overall geostatistical confidence in these estimates being less. Increasing the block size effectively reduces the kriging variance, and improves the geostatistical confidence in the resultant block estimate.

Having a higher confidence in this large block (or panel) grade, means that one can confidently assume that the average grade exists in that panel, although the location of that grade within the panel is unknown. Using this local panel mean estimate, one can “condition” the estimate of a local SMU distribution to this panel grade estimate. The result of this is a recoverable resource (proportions and grades above cut-off grade, for a series of cut-off grades).

Localisation occurs at the end of the workflow, where the UC results are converted into a more useable format (particularly for mine planning applications). The spatial locations of SMU blocks within the panel are based upon a local linear estimate of SMU blocks.

A workflow outline is shown in Figure 2. Details of each step are discussed in this chapter.

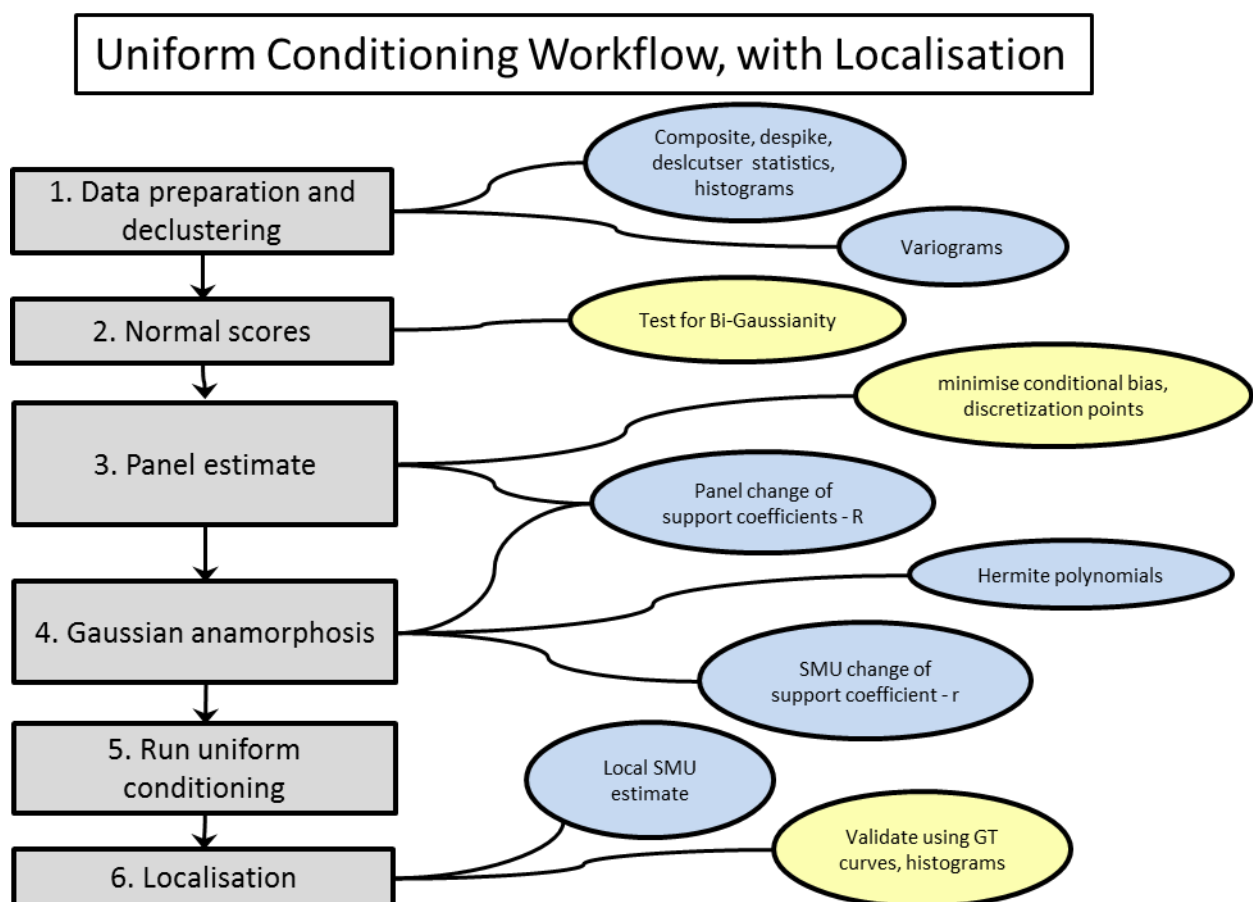


Figure 2 Localised uniform conditioning workflow

UC has a fairly complex workflow, and the selection of parameters requires careful thought and consideration. The effect of poor declustering could result in an incorrect anamorphosis function, which would reduce the performance of UC. Additionally, a poor panel grade estimate has a direct impact on how well uniform conditioning performs.

2.5.1 Estimate panel grades

The change of support model between SMU and panels together with the panel grade conditions the SMU grades. The quality of the panel estimate, specifically it being

conditionally unbiased, determines the success of the UC result and this sentiment is expressed by numerous authors.

The panel estimation units must be of the original grade. The estimation can be carried out using any linear estimator, which could be simple kriging (SK) with a locally varying mean or OK.

The size of a panel should be considered relative to the spacing of the sample data. De Vitry et al. (2007) suggested that the panel should be as small as possible to ensure an accurate estimate, but large enough for minimal conditional bias of the estimate. Estimating small panels relative to the sample spacing will result in unreliable estimates; while larger panels will provide a more reliable estimate of the grade. This phenomenon is evident from considering the slope of regression calculated from the same data at different block sizes.

There should be enough SMU within a panel to effectively represent the estimated distribution. De-Vitry et al. (2007) recommend more than 10 SMU per panel, while Rivoirard (1994) recommends at least “several dozen”. The number of SMU within the panel is linked to the resolution of the grade-tonnage relationship, as the number SMU discretise the grade-tonnage curve of a panel (Harley and Assibey-Bonsu, 2007).

In order to obtain an accurate block estimate, several discretisation points in the three principal directions should be used. A block estimate will create an average result of all the point estimates calculated at each discretisation point. A quantitate kriging neighbourhood analysis (QKNA) exercise can be carried out to select the optimal number of discretization points.

It is important to ensure that panel estimate reproduces the grades of the sample data. This can be verified with swath plots, the slope of regression and other statistical indicators of a reliable linear estimation.

2.5.2 Normal score transform

As a Gaussian change of support model is used for UC, the data must be converted to the equivalent Gaussian (or normal score) values. This is done by transforming the cumulative distribution function (cdf) of the original grade $Z(x)$ to a Gaussian probability $Y(x)$ cdf. The transformation is mapped on a percentile to percentile basis, from the original cdf to the Gaussian cdf, as shown in Figure 3. This mapping is then applied to the entire dataset.

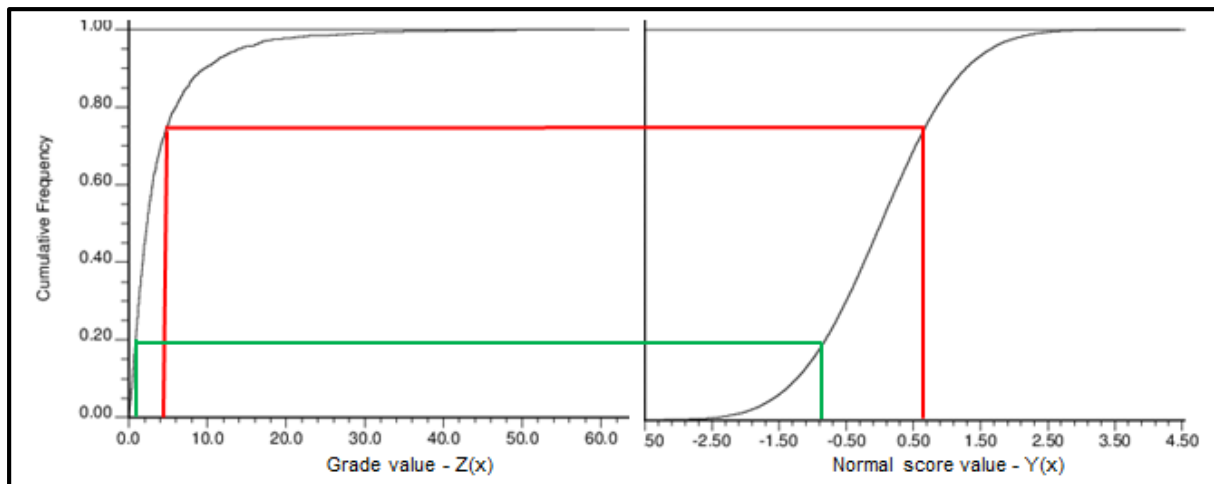


Figure 3 Cdf of grade $Z(x)$ to normal score $Y(x)$ transformation, for a log-normal distribution

Emphasis must be put the correctly applying the normal score transform. If the original grade data is poorly declustered, the normal score transform will not adequately represent the true distribution of the data. As the Gaussian anamorphosis function describes the normal score transformation, an incorrect normal transformation will adversely affect the results of the UC estimation.

2.5.3 Gaussian anamorphosis

The discrete Gaussian change of support model (DGM) uses a Gaussian anamorphosis function to change samples support to a larger support, and is based on an underlying diffusion model (see chapter 3.7.3). Samples $Z(x)$ selected at a point support, are used to krigé blocks at a panel support. This average panel grade is then conditioned to a local SMU distribution.

An anamorphosis function is fitted to the sample data by plotting the normal score transformed values $Y(x)$ on the horizontal, against the original unit point data values $Z(x)$ on the vertical. Considering Figure 3, this is equivalent to plotting equivalent grade values (plotted on the left) and normal score grade values (plotted on the right), for all cumulative frequency pairs. The Gaussian anamorphosis function is modelled by a set of Hermite polynomials, fitted with an accompanying set of Hermite coefficients. A full description of Hermite polynomials is given by Armstrong and Matheron (1986) and Rivoirard (1994).

Hermite polynomials are a set of orthogonal polynomials, which are derivatives of the Gaussian probability density function (pdf). In this respect, they express probabilities. Spatially, each polynomial is uncorrelated with all other polynomials in the set. These polynomials have other applications in physics and calculus, and their relationship to the Gaussian distribution makes them suitable for applications in Gaussian statistics. The applicable properties of a standard Gaussian distribution are that data is distributed symmetrically, with a mean of zero and a variance of one.

Any function can be fitted as a weighted sum of Hermite polynomials (this is similar to Fourier analysis). This allows one to fit the cdf of grade values $Z(x)$, as a function of Hermite polynomials, to the Gaussian $Y(x)$ values, producing a set of coefficients (ϕ_{0-n}) which weight the Hermite polynomials.

The number of coefficients can vary, and the suitable number depends on how well the polynomial set fits the underlying distribution. Neufeld (2005) recommends less than 100 coefficients, although typically 20 to 30 coefficients are usually sufficient. The coefficients have the following properties:

- The first order coefficient is equal to the mean of the grade distribution.
- The sum of the squares of the 1st to nth coefficients is equal to the variance of the point data.
- The contribution of the coefficients to the total variance decreases as the order of the coefficients increases (e.g. a 1st order coefficient accounts for the most variance, followed by the 3rd, 4th etc.).

To calculate the Hermite polynomials, a recursive relationship based on Rodrigues' formula is used (Rivoirard, 1994, Neufeld, 2005). This formula for the 0th, 1st and 2nd order polynomials are:

$$H_0(y) = 1$$

$$H_1(y) = -y$$

$$H_2(y) = \frac{1}{\sqrt{2}}(y^2 - 1)$$

For subsequent Hermite polynomials ($n \geq 2$), the following equation is used:

$$H_{n+1} = -\frac{1}{\sqrt{n+1}}y H_n(y) - \sqrt{\frac{n}{n+1}}H_{n-1}(y)$$

The $Z(x)$ grades are then calculated by multiplying each solved Hermite polynomial (H_n) by a corresponding Hermite coefficients (ϕ_n), depicted in the following formula. The polynomials are solved for Gaussian values $Y(x)$ of an appropriate range (not limited to, but a range could be between -6 and 6).

$$Z(x) = \sum_{i=0}^n \phi_i H_i[Y(x)]$$

Where:

$Y(x)$ Normal score transform of $Z(x)$

n Number of Hermite polynomial terms

ϕ_i Phi; coefficient fitted for each term of polynomial expansion

H_i Hermite polynomial evaluated for Gaussian $Y(x)$ value

Additionally, the variance of the grade $Z(x)$ can be expressed by the sum of the Hermite coefficients, excluding the 0th, which describes the mean of $Z(x)$.

$$\sigma^2_{Z[y(x)]} = \sum_{i=1}^n \phi_i^2$$

To simplify the understanding of the $Z[y(x)]$ grade function, it can also be expressed in the following terms:

$$Z[y(x)] = \phi_0 H_0(y) + \phi_1 H_1(y) + \dots + \phi_n H_n(y)$$

This concept can also be visualised, and an example is given in Figure 4, displaying the first six Hermite polynomial expansions, and the resultant $Z[y(x)]$ model. The grade values are all positive, while the Hermite polynomials of the factors are both positive and negative, while the resultant sum of these is always positive within the bounds where it is defined.

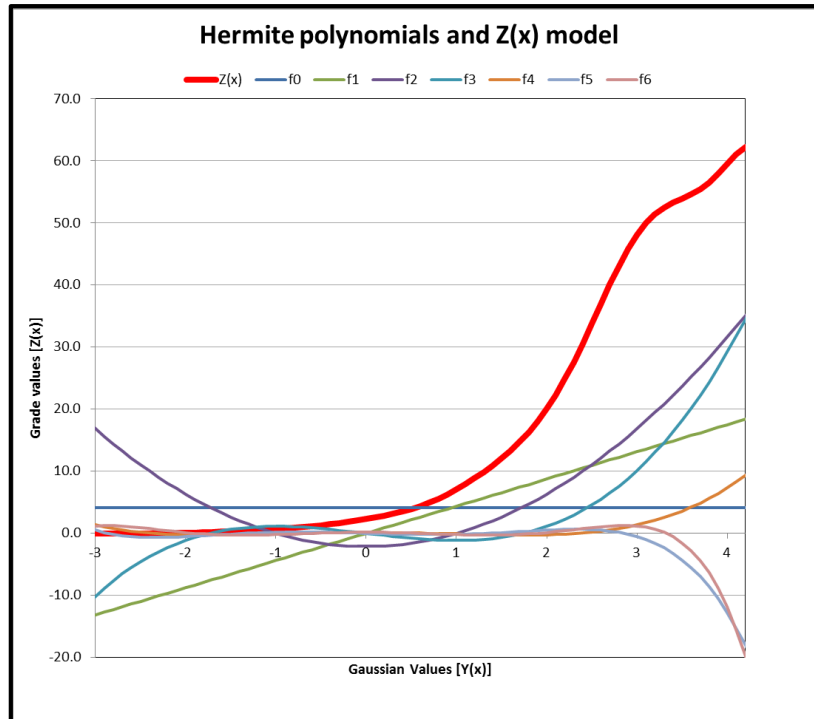


Figure 4 First six Hermite polynomials and resultant $Z(x)$ model

In summary, the anamorphosis function, described in terms of Hermite polynomials, provides a mapping from the Gaussian $Y(x)$ values, to the grade values at a point support. Figure 5 shows a point distribution (i.e. transformation mapping of grade values to equivalent normal score values), and the corresponding Gaussian anamorphosis model fitted to this distribution.

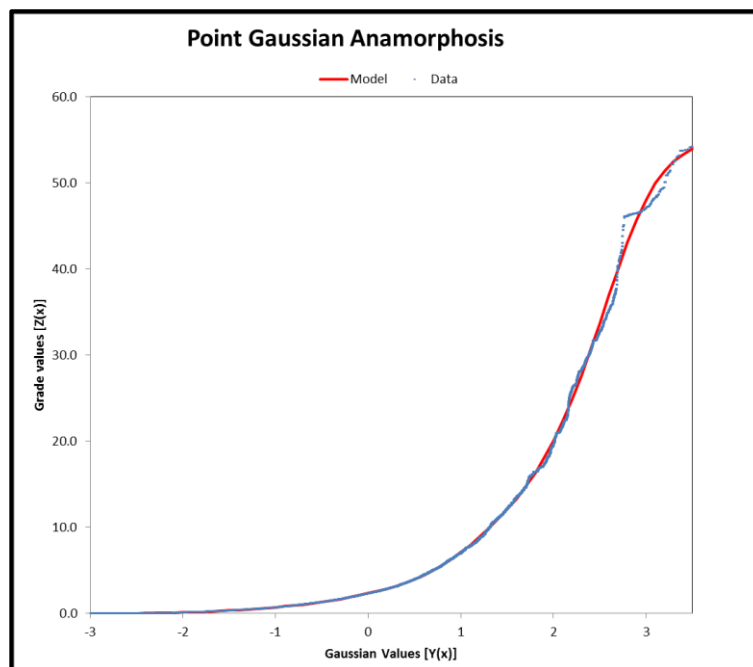


Figure 5 Gaussian anamorphosis experimental data and model

It is important to be aware that the Gaussian properties of the underlying distribution are reflected in this anamorphosis function e.g. At $Y = 0$, the Z value reflects the mean and median value, with 50 % of the data falling above and below this value. Since Gaussian values are

synonymous with probabilities, standard Gaussian probabilities can be applied to this model. Additional observations about the DGM is that the anamorphosis of a normal variable will show a linear model, while a skew/log-normal model will show a curved model (as seen in Figure 5).

A DGM is a suitable tool for the change of support, because of the specific Gaussian properties and probabilistic nature of the model. The model requires an underlying (multi- or bi-) Gaussian distribution of the data.

2.5.4 Change of support

As discussed earlier, the distribution of grades depends on the support. There is correlation between grades seen at a point support and grades seen at panel or SMU support. The DGM uses a ratio of these correlations to adjust the distribution of grades to the required level of support.

Previously, we considered the equation for distribution of grades at a point support i.e. point anamorphosis equation:

$$Z(x) = \sum_{i=0}^n \phi_i H_i[Y(x)]$$

A property of the Hermite polynomials is that the square of the coefficients from the first to the n^{th} coefficient is equal to the variance of the point data, shown in the following equation:

$$\text{Var}[Z(x)] = \sum_{i=1}^n \phi_i^2$$

To account for change of support for an SMU, this point equation above is modified by a support coefficient, r . The r -coefficient is the correlation in Gaussian space between point grades $Y(x)$, and SMU grades $Y(v)$. The following SMU anamorphosis equation is a modification of the point anamorphosis equation, which includes the r -coefficient:

$$Z(v) = \sum_{i=0}^n r \phi_i H_i[Y(v)] \quad [0 \leq r \leq 1]$$

Where:

$Z(v)$ Distribution of grades at an SMU support

r Change of support coefficient at SMU support

Additionally, the variance of grades at an SMU support is given by the equation:

$$Var[Z(v)] = \sum_{i=1}^n r^{2i} \phi_i^2$$

Using the variogram model and Krige's relationship, presented in a different way, we are able to determine the theoretical dispersion variance of grades in an SMU:

$$Var [Z(v)] = Var[Z(x)] - \bar{\gamma}(v, v)$$

Where:

$Var [Z(v)]$	SMU dispersion variance
$Var [Z(x)]$	Total variance, or sill of $\gamma(h)$
$\bar{\gamma}(v, v)$	Average variance of points within a block

Thus, from the theoretical dispersion variance of SMU grades, we are able to calculate the SMU change of support (r -coefficient). The r -coefficient now allows for the calculation of the distribution of Z grades (real space) at an SMU support, from a conditional panel grade.

Similarly to the SMU support, the panel anamorphosis uses the panel support coefficient, R , and is given by the following equation:

$$Z(V) = \sum_{i=0}^n R \phi_i H_i[Y(V)]$$

The dispersion variance of panel grades can be calculated directly from OK estimates of the panel grades. Using this approach, one directly measure the variance of the estimated panel grades from linearly estimating these panels.

Alternatively, if OK was used, the variance of the estimated panel grades $Var[Z(V)^*]$ can be calculated theoretically using the equation below. The variance of the estimated grades may be used as an approximation of the dispersion variance of a panel. If SK was used, then the Lagrange multiplier is not added to the equation.

$$Var[Z(V)^*] = \bar{C}(V, V) - \sigma_k^2 + 2\mu$$

Where:

$\bar{C}(V, V)$	Block variance
σ_k^2	Kriging variance
μ	Lagrange multiplier

Similarly to how the SMU change of support coefficient was calculated, the panel change of support R -coefficient can be calculated from the dispersion variance of panel grades.

The change of support occurs using the ratio of R/r . Using r , which is the correlation between Gaussian point grades $Y(x)$ and SMU grades $Y(v)$, and R , which is a correlation between Gaussian point grades $Y(x)$ and panel grades $Y(V)$, a correlation between Gaussian SMU grades $Y(v)$ and panel grades $Y(V)$ is determined. A distribution of SMU grades is then produced, for a particular panel grade, as a set of proportions and average grades above cut-off. This is illustrated in Figure 6.

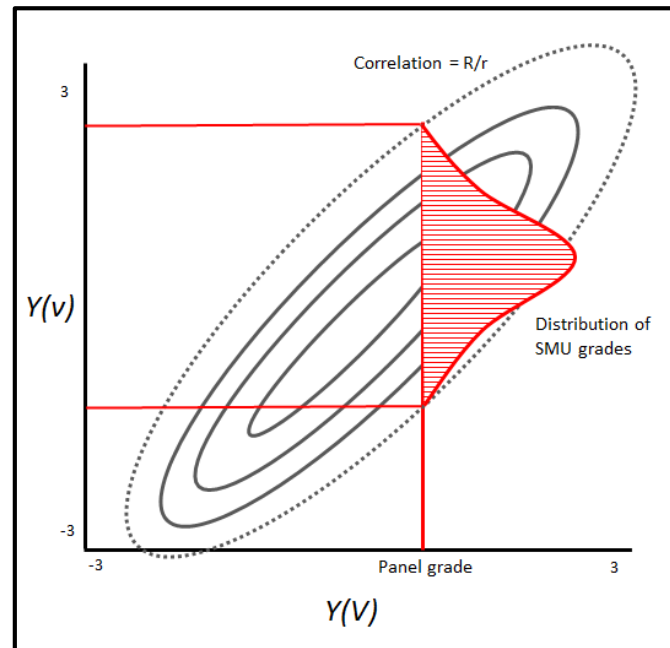


Figure 6 Conditional SMU distribution from a panel grade

The ratio of R/r is a ratio between 0 and 1. A low R/r ratio indicates there a weak correlation between the SMU and panel grades, which is caused by a high nugget variogram and/or short variogram ranges relative to the data spacing. A large R/r ratio is indicative of a strong correlation between SMU and panel grades, which is indicates good grade continuity in the deposit.

Uniform conditioning provides a distribution of SMU grades, based on a panel grade which conditions the SMU distribution. Using the change of support methodology discussed in this section, point-support grades are estimated into a large support (a panel) which provides an indirect determination of grade and proportions above cut-off at an intermediate support (an SMU).

2.5.5 Q-T Curves

Uniform conditioning generates a series of proportions and average grades, at a series of grade cut-offs. While this is insightful information about the grade tonnage distribution, it is not a particularly practical data format for mining engineers. For this data to be more useable for visualising grades and for mine planning purposes, is should be modified into SMU average block grades.

The grade and proportions generated by uniform conditioning can be converted into a metal-quantity tonnage (“Q-T”) curve, as shown in Figure 7. Metal quantity is calculated by multiplying the grade above cut-off by tonnage.

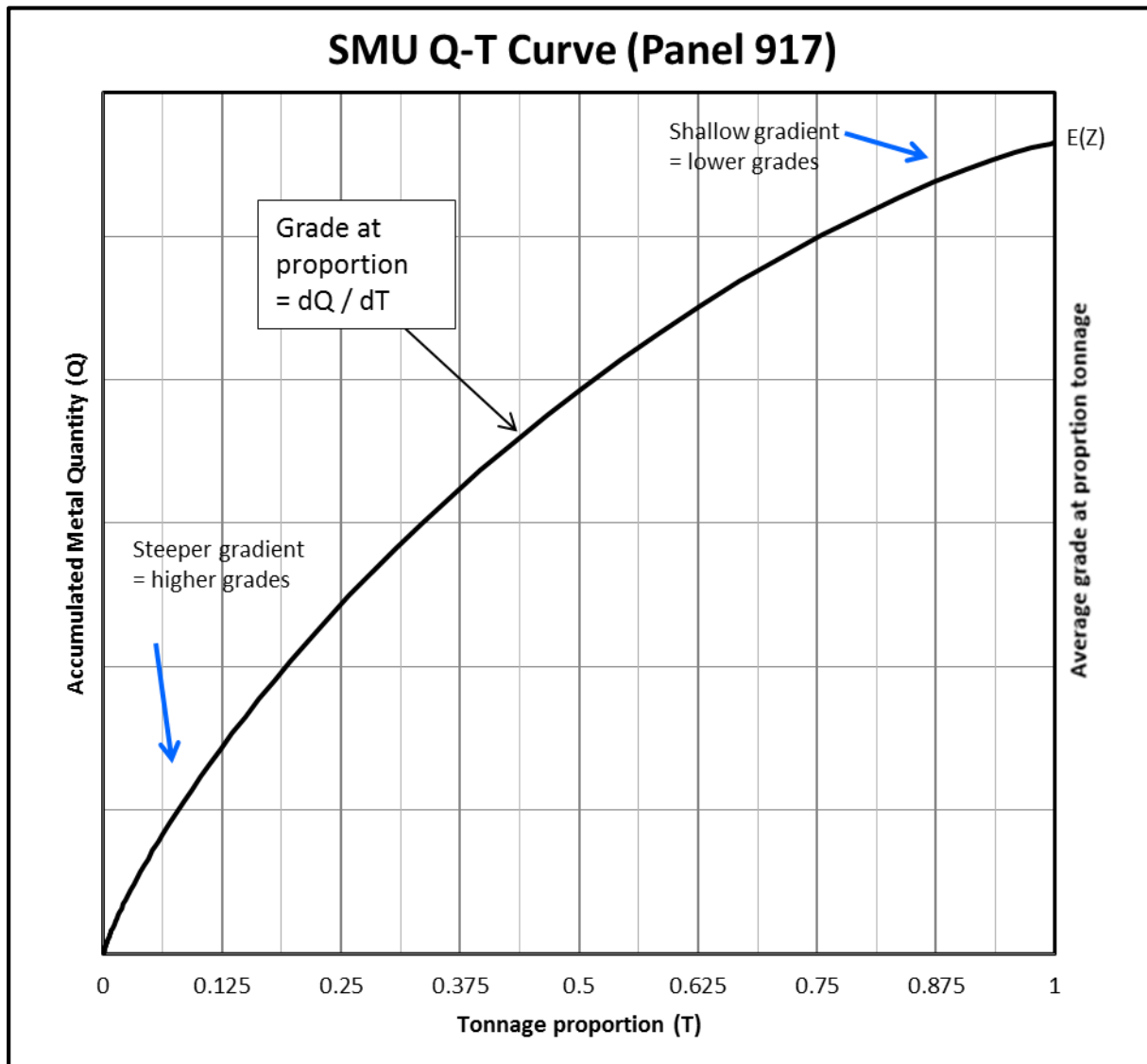


Figure 7 Q-T Curve for a panel

The Q-T curve is a convex curve which shows the grade distribution for the panel. Rivoirard (1994) provides the following information about a Q-T curve:

- At any point on the curve, the slope (i.e. change in metal quantity / change in tonnage) is the grade value at that proportion.
- At any point on the curve, calculating the gradient of a line drawn between the point and the (0,0) axis will provide the average grade above that tonnage proportion.

For any given panel, there are a fixed number of SMU in a panel. By dividing the tonnage proportion into the number of SMU in a panel, one is able to calculate the grade value at that proportion. This enables one to calculate an average grade for each SMU within a panel.

2.6 Localised uniform conditioning

Uniform conditioning calculates the expected proportions of grades above cut-off for a panel. These proportions are then converted to a set of SMU grades, occurring at unknown locations within a panel.

The localisation process maps these SMU grades into appropriate locations within the panel, based on the rank location of OK grades within a model of the equally sized SMU blocks. This results in a direct grade model of SMU blocks, that respect the UC grade distribution for the panel and attempts to respects the estimated location of high and low grades within the panel.

Localised Uniform Conditioning was pioneered by Abzalov (2006). The localisation technique can be applied to any non-linear, large block estimate of a distribution, as demonstrated by Hardtke et al. (2011), who applied this methodology to IK.

The following explain the ranking process in more detail:

- Uniform conditioning is carried out on a panel model using a change of support for a specific SMU size.
- A grade model, of the same block size as the SMU, is estimated using OK or another linear estimation technique. The purpose of this estimation is to identify locally high and low grade SMU.
- Using the number of SMU in the panel, a set of average UC grades for each SMU in the panel is created that honours the proportion above cut-off grade curve. The mean of the set of UC grades is equal to the panel grade.
- The set of UC grades are ranked from lowest to highest grade value. Similarly the grade model SMU, within each panel, are ranked from lowest to highest grade value.
- The set of UC grades are then mapped onto the grade model SMU, and the local UC grades are applied to the model i.e. the SMU with the highest OK grade will be given the highest local UC grade, and so on.

The advantages of localising a UC estimate is that it can be used for mine planning and also benchmarked with other grade estimation techniques.

3. Approach

The objective of this project is to assess the suitability of UC and LUC for two types of grade distributions. The two distributions are distinctly different from one another, and represent two end-members of the range of grade distributions that may typically be seen in mineral occurrences.

The grade distributions are synthetically generated and sampled, to mimic how this would occur in a mineral exploration project. The deposits were then analysed statistically, and evaluated using UC and then LUC. For comparative purposes, OK was also considered. The evaluated models are presented, compared and the results discussed.

The UC approach taken is primarily based on Rivoirard (1994), with inputs from Assibey-Bonsu and Krige (1999), Neufeld (2005), Neufeld and Deutsch (2005), De Vitry et al. (2007) and Deraisme et al. (2008). The LUC approach is based on Abzalov's (2006) and Harley and Assibey-Bonsu's (2007) work on LUC, and Hardtke's (2011) work on localised indicator kriging.

3.1 Generating simulated realisations

Two sets of simulated data were generated, which were used as the base data for the UC and LUC assessments. A single realisation was simulated on a 2 m x 2 m x 2m point grid, over an area sized 800 m x 600 m, with a depth of 200 m, for each scenario.

The data used in this project is hypothetical, and was created to mimic conditions seen in some mineral deposits. For the sake of convention, sample data will be considered as concentration of a mineral occurrence and given the unit grams per ton (g/t). The first and second distributions will be referred to as "Scenario 1" and "Scenario 2" throughout this project.

Scenario 1's grade data was designed as symmetrical, with a low nugget and well defined continuity up to distances of 60 m in the Z direction and 150 m in the X and Y directions. The grade distribution was tested for normality, and was confirmed as being close to normally distributed. The design was used to simulate a set of points that reflect these properties. It is similar to the type of distribution one would find in a porphyry copper mineral occurrence.

Scenario 2's grade data was set up to be approximately log-normally distributed, and the logs of the values were successfully tested for a normal distribution. This distribution has the characteristics of being asymmetrical, strongly positively skewed with a long tail. For spatial variability there is a high nugget effect (approximately 30% nugget), with an assumed short range continuity of approximately 10 m in the Z direction to 30 m in the X and Y directions. The design was used to simulate a set of points for Scenario 2 that reflect these properties, which one might find in highly skewed gold deposit.

The coefficient of variation (CoV) is a normalised measure of dispersion which can be used to compare distributions, and is calculated by dividing the standard deviation over the mean, shown in formula as $\frac{\sigma}{\mu}$. CoV values for Scenario 2 are higher than that of Scenario 1, showing a wider spread of grade values.

3.1.1 Reference distributions for simulation

To create references distributions for the realisations the following methodology was used:

A Gaussian distribution of 2000 data values was generated from the 'Random Number Generator' in Microsoft Excel. Scenario 1's base data was multiplied by a constant 4.6, and another constant value 11.5 was added. The constants 4.6 and 11.5 approximate the standard deviation and mean of Scenario 1's statistics. Remaining negative values were set to zero. A histogram of this distribution is shown in Figure 8 below.

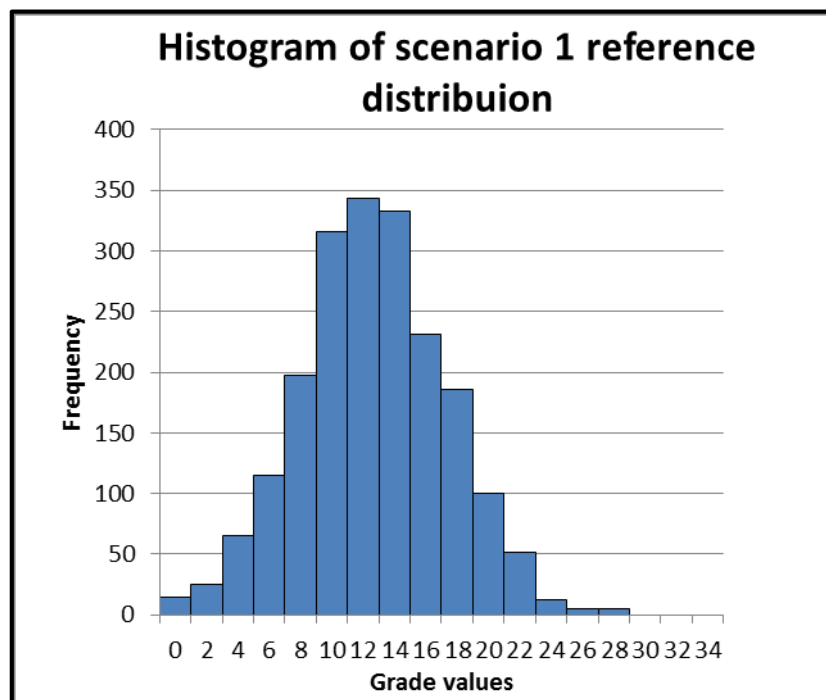


Figure 8 Histogram of Scenario 1's reference distribution

Similarly to above, a Gaussian distribution was generated using the Random Number Generator in Microsoft Excel. The exponent of Scenario 2's base distribution was taken to transform the data into a log distribution. The data was multiplied by a constant (2.48). An additional small constant (0.2) was added to Scenario 2, and remaining all negative values were set to 0. A histogram of the reference distribution can be seen in Figure 9.

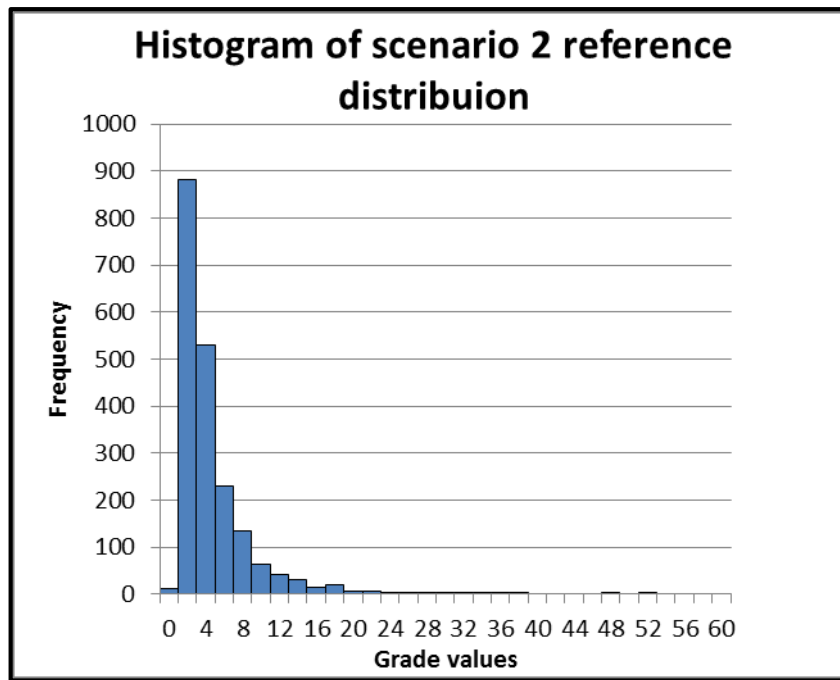


Figure 9 Histogram of Scenario 2's reference distribution

The realisations were checked for histogram reproduction, and the statistics of the simulation reproduced those of the reference distributions.

3.1.2 Variograms reproduction of simulated

Anisotropic variogram models were created in the three primary directions, to describe the spatial variability chosen for each scenario as described in section 3.1. These models were used to direct the variance of points generated in the simulation algorithm.

In order for the simulated realisations to reproduce the variograms, a technique was used where widely spaced unconditional simulations were generated using the variograms, and these simulated points were used to condition the subsequent realisation. The variogram models of the intermediate Gaussian simulation are seen in Figure 10 and Figure 11.

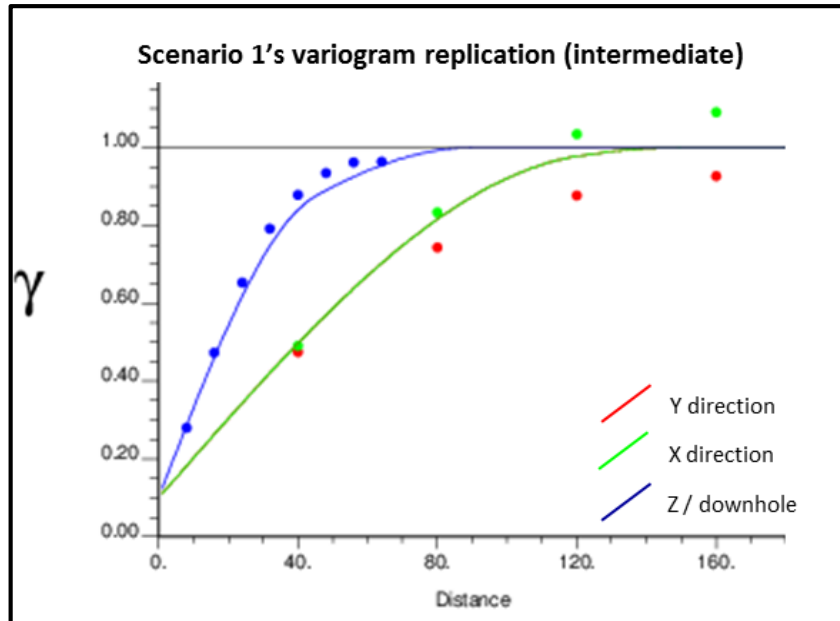


Figure 10 Scenario 1 model and experimental variogram from intermediate (widely spaced) unconditional simulation

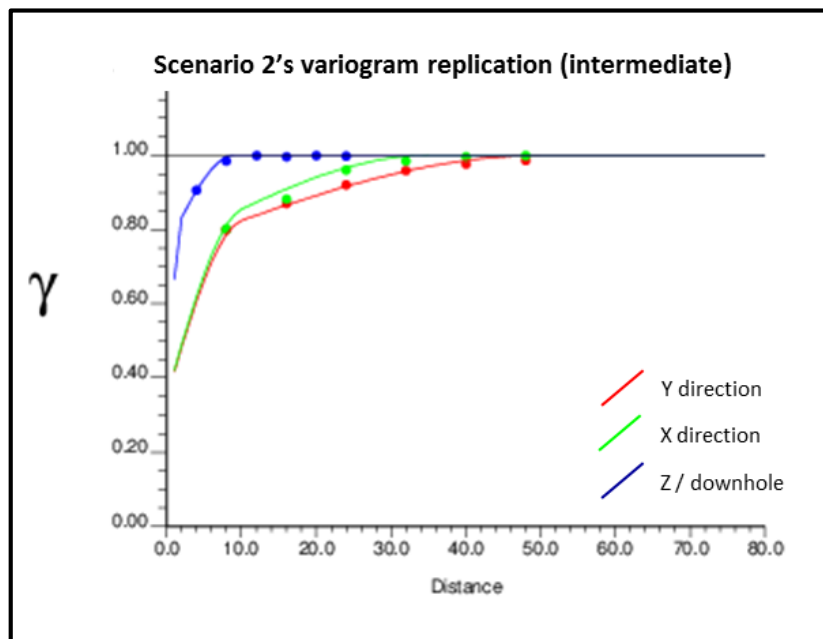


Figure 11 Scenario 2 model and experimental variogram from intermediate (widely spaced) unconditional simulation

The final variogram reproduction of both realisations were checked against the modelled input variograms and were close enough to satisfy the chosen conditions, as shown in Figure 12 and Figure 13 respectively.

However for Scenario 1, the initial variograms (Figure 12) were set up as equal distances in the X and Y directions. The final simulated variogram indicates that some anisotropy is evident, with more continuity seen in the X (red) direction. This effect is caused by the ratio of the X to Y values in the simulated grid, as the grid is 800 m in X and 600 m in Y.

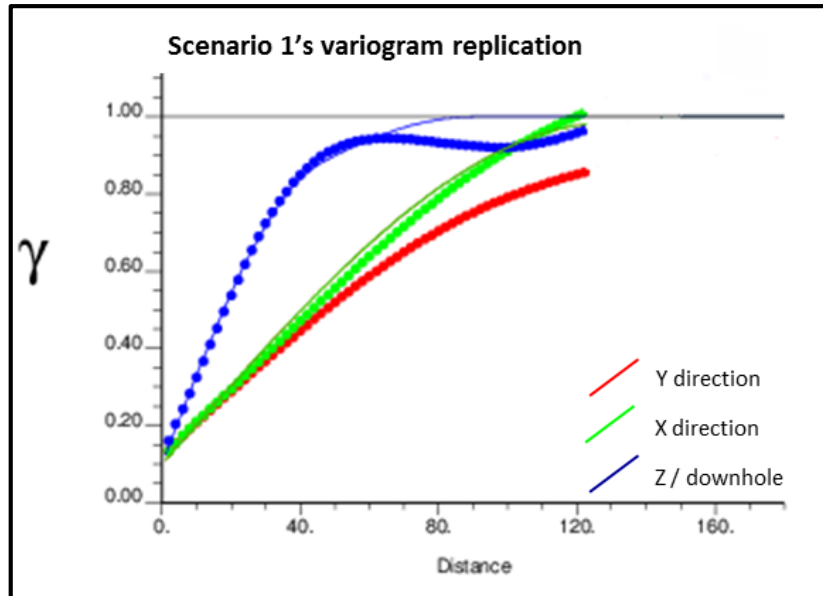


Figure 12 Scenario 1 model and simulated variograms, showing variogram replication

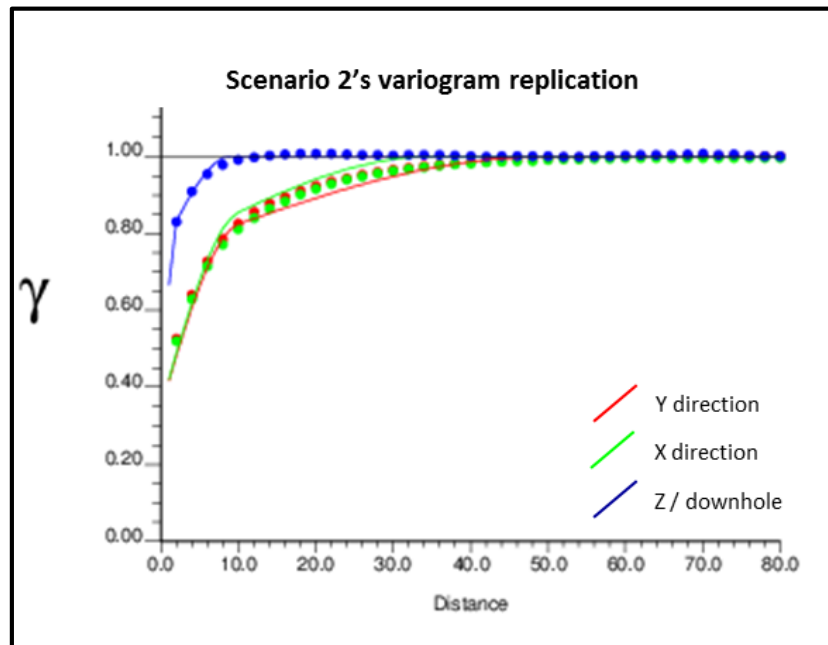


Figure 13 Scenario 2 model and simulated variograms, showing variogram replication

3.1.3 Plot of realisations

A plan view at surface through the realisation for Scenario 1 and Scenario 2 are shown in Figure 14 and Figure 15 respectively. Visually, both realisations turned out as per design criterion. Scenario 1 shows a low nugget effect with good continuity, and Scenario 2 shows a high nugget effect with short continuity. Scenario 1 has a smaller range of grade values than Scenario 2 (approximately half).

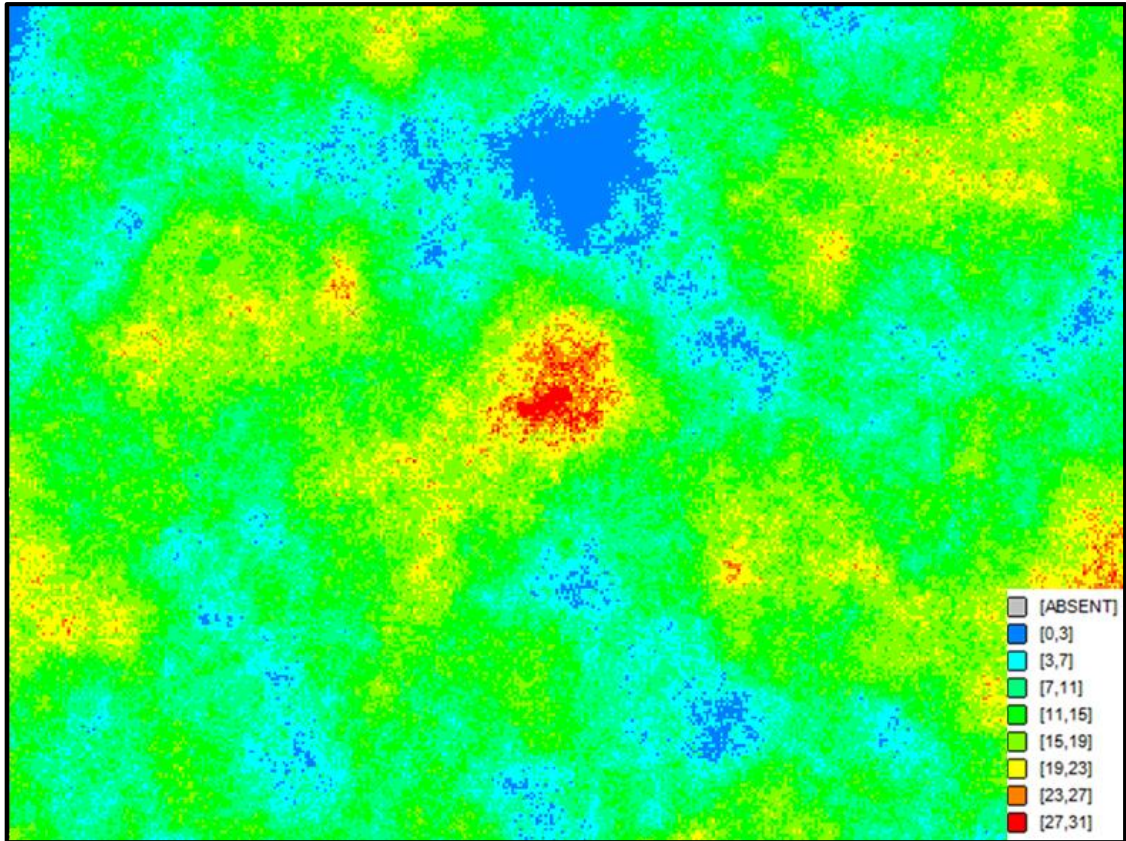


Figure 14 Plan view of Scenario 1 simulation (at surface elevation)

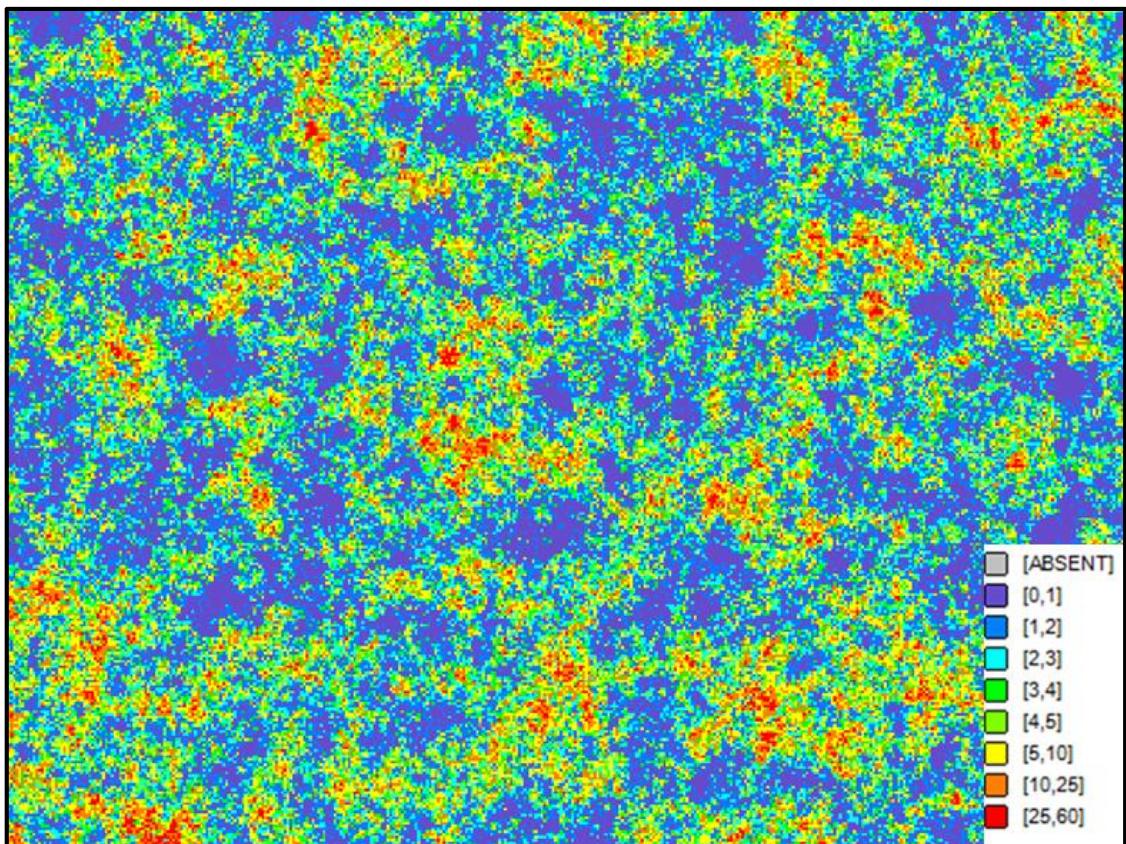


Figure 15 Plan view of Scenario 2 simulation (at surface elevation)

3.2 Sampling realisations

A spatially representative subset of data was taken from both realisations, which makes up the sample database used for this project. The same sampling pattern was used for both deposits.

The simulated “deposits” were sampled with pseudo drillholes on an irregular 40 m x 40 m grid. The grid was tightened to 25 m x 25 m grid in two warm spots of high grade values seen in Scenario 1 seen at surface. This was to mimic an exploration drilling that would typically have taken more samples in presumed higher grade areas. No consideration was taken of the location of Scenario 2’s high grade areas. The simulated model was spaced at a 2 m grid in the Z-direction; these are assumed to represent 2 m down-hole composites.

Although Scenario 1 and Scenario 2 are equally sampled i.e. sample locations are identical, the drill hole spacing relative to the variogram ranges for Scenario 1 are closer than for Scenario 2. For Scenario 1, the samples are spaced at approximately one third the variogram ranges of 120 m in X and Y, while Scenario 2 samples are spaced at approximately the maximum variogram range of 40 m.

A total of 417 pseudo drillholes, each containing 100 composites, were taken over the 600 m x 800 m x 200 m study area. The area is densely sampled, and this drilling grid would be consistent with that of a feasibility stage project. The sampling patterns with grades at surface are shown in Figure 16 and Figure 17 for the two scenarios.

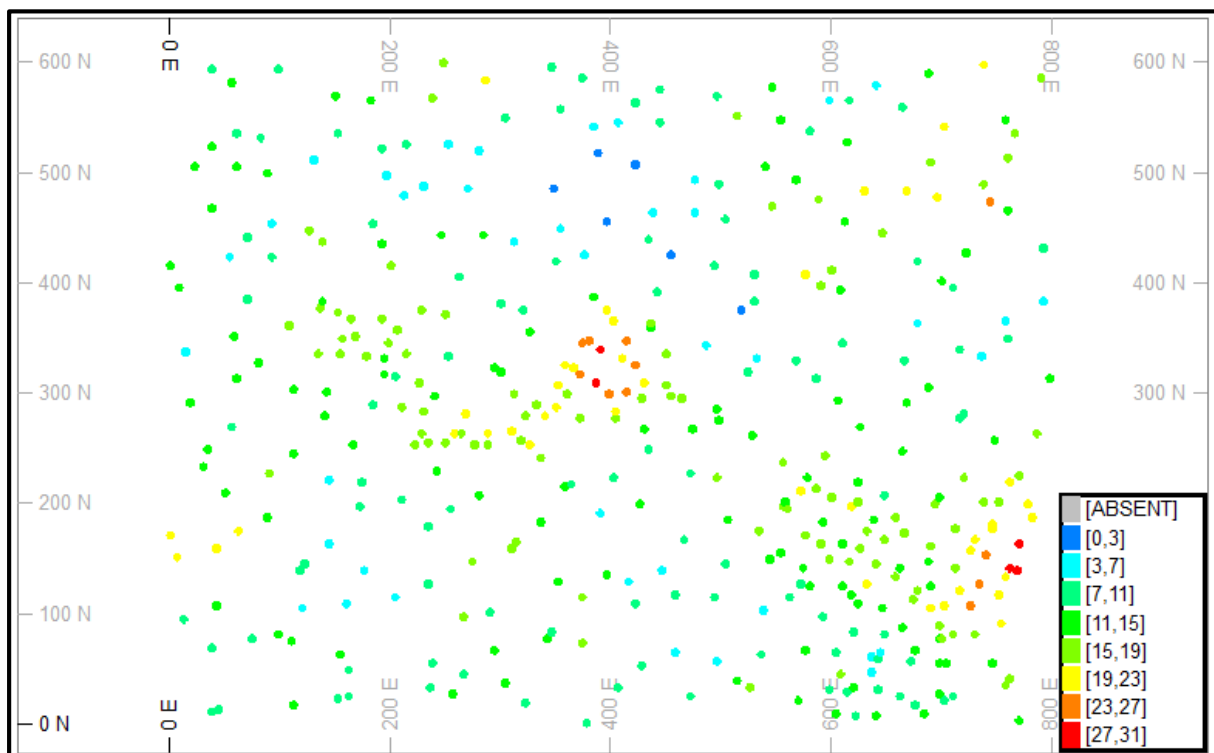


Figure 16 Sampling pattern showing locations of vertical pseudo-drillholes for Scenario 1 (plan view at surface)

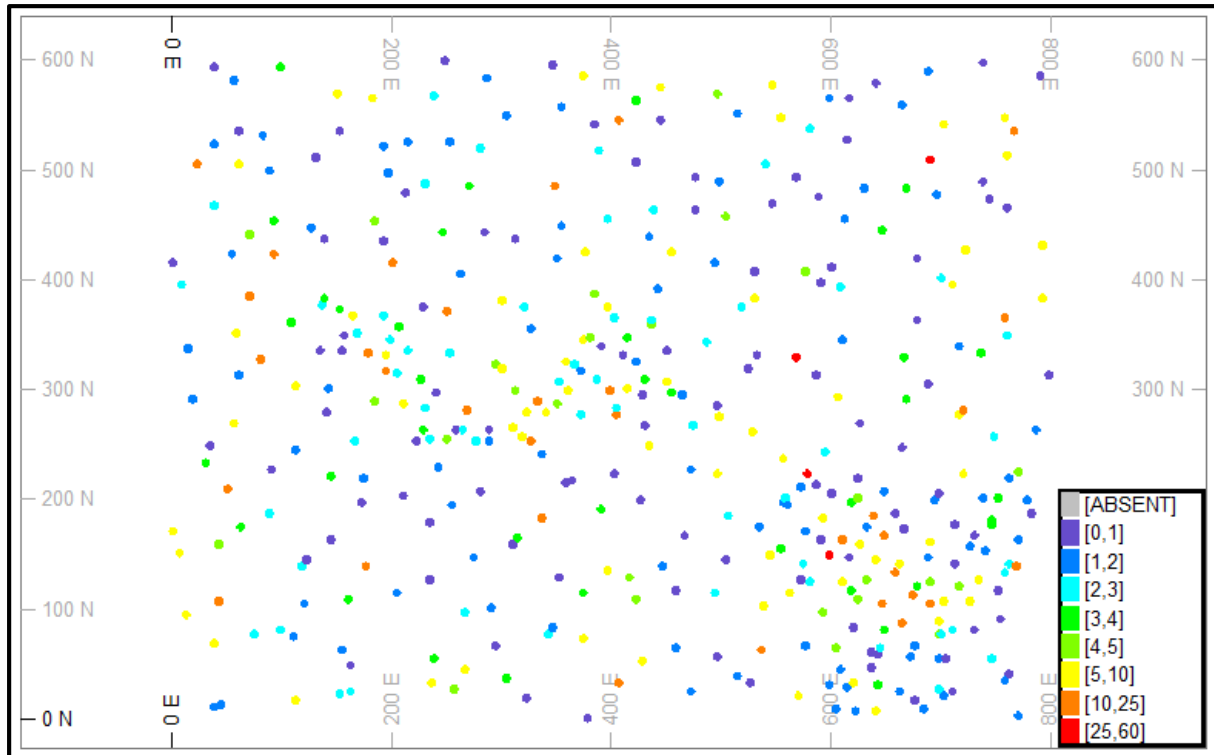


Figure 17 Sampling pattern showing locations of vertical pseudo-drillholes for Scenario 2 (plan view at surface)

3.3 Exploratory data analysis

Exploratory data analyses were carried out for the grade samples of Scenario 1 and Scenario 2, to assess the statistical characteristics of each. Although statistics of both variables are compared, there is no statistical relationship between Scenario 1 and Scenario 2 as the simulations were run independently of each other. Additional statistics of the data sets and model are shown in Appendix A: Statistics Tables, and basic statistics for both variables are presented below in Table 1.

Table 1 Descriptive statistics for Scenario 1 and Scenario 2

	Number of samples	Number of boreholes	Min g/t	Max g/t	Mean g/t	Variance g/t ²	Standard deviation g/t	Skewness	Kurtosis	Coefficient of variation
Scenario 1	41 700	417	0.0	29.6	11.8	21.5	4.6	0.1	0.3	0.4
Scenario 2	41 700	417	0.0	59.3	4.1	30.3	5.5	3.7	19.5	1.3

Scenario 1 has a mean value of 11.8 g/t, with a symmetrical distribution and a standard deviation of 4.6 g/t. There is a low coefficient of variation (0.4), indicating there is low variation due to the bounds of a standard deviation being close to the mean. The histogram shape and cdf plot in Figure 18 appears close to normal in shape.

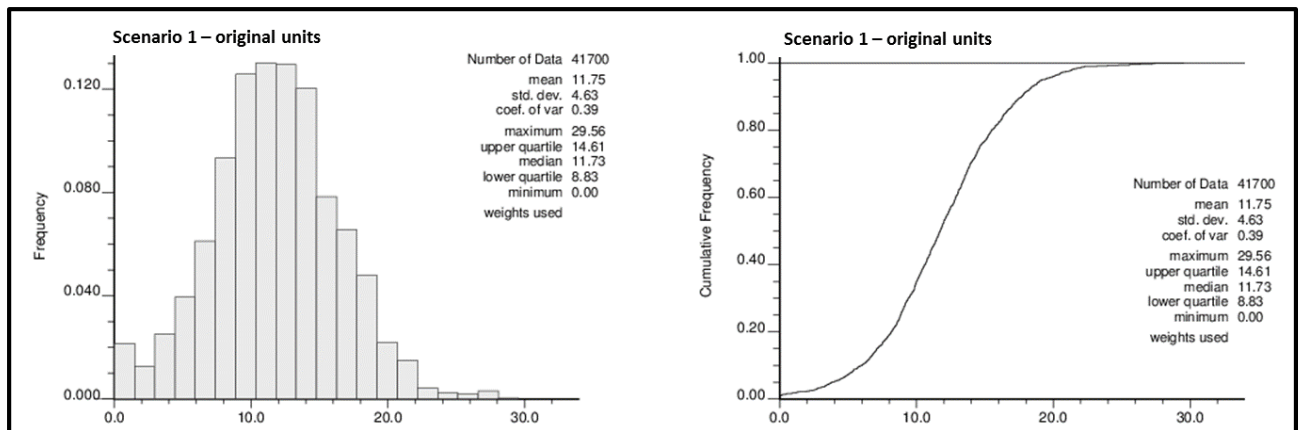


Figure 18 Scenario 1 sample histogram and cumulative frequency histogram

Scenario 2 has a lower mean (4.1 g/t) than Scenario 1, but the range is much wider (almost double) than that of Scenario 1. Scenario 2 is highly positively skewed. The standard deviation is 5.5 g/t, with a coefficient of variation of 1.3 indicating quite a high variability and the possibility of a log-normal distribution. The histogram (Figure 19) shape appears close to log-normal.

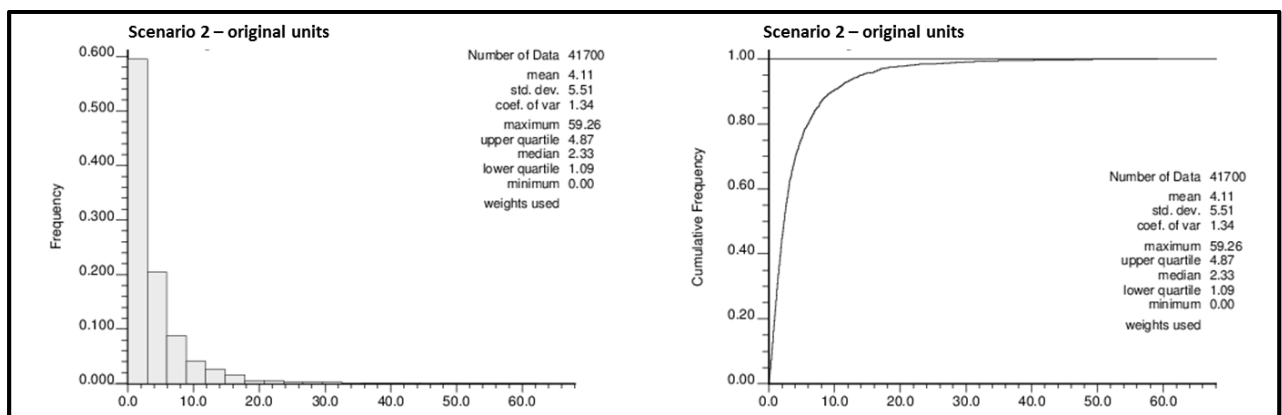


Figure 19 Scenario 2 sample histogram and cumulative frequency histogram

3.4 Normal scores

The discrete Gaussian change of support model requires that sample data be rescaled to equivalent normal score values. In order to ensure a robust normal score transformation, the data was declustered to remove the effects of unequal sampling specifically as high grade areas were targeted during sampling.

Samples were declustered on a 60 m x 60 m grid using an offset origin, which was chosen by analysing multiple declustering grids sizes for convergence of a declustered mean. Given the equal sampling pattern of the two variable sets, the same weights were applied to both sets.

The grade samples were then transformed to normal scores, using these declustered weights. The normal score transform attempts to create a perfect Gaussian distribution, with a standard deviation of one and the distribution centred on a mean value of zero.

Figure 20 and Figure 21 shows histograms and cdfs for the normal score transformed values for Scenario 1 and Scenario 2 respectively. These can be compared to Figure 18 and Figure 19 respectively (the original grade value histograms and cdfs). The normal score transform changes the shape of the Scenario 2 cdf more significantly, as the values were transformed from a highly skewed distribution to a standard normal distribution.

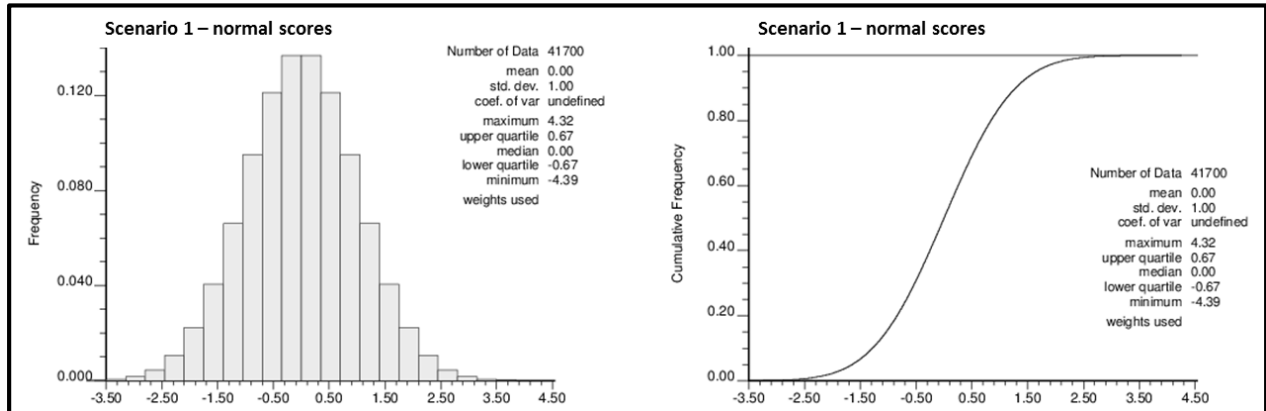


Figure 20 Normal scores of Scenario 1 histogram and cdf

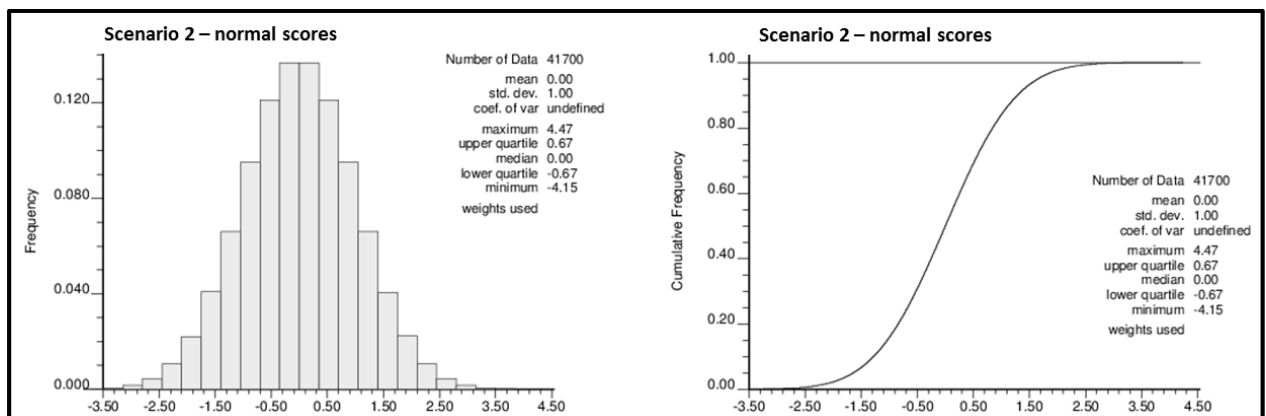


Figure 21 Normal scores of Scenario 2 histogram and cdf

3.5 Variography

Variography of the sampled data was carried out in order to estimate the panel grades. Experimental variograms were modelled in the three principal directions, X, Y and Z, and nested spherical variogram models were fitted to this data. The variograms were not normalised to the variance of the population, so the true variance (rather than relative variance) of each distribution can be viewed on the variograms.

Scenario 1's variogram model showed a slight nugget effect of 12 %, with good continuity in all directions. Slight anisotropy was evident, and possibly some zonal anisotropy seen in the X-direction where the variance does not reach the sill value. In the short range area there is more variance seen in the down-hole variogram. The maximum variogram ranges in all directions range from 170 m to 300 m. The experimental variogram, variogram model and parameters are shown in Figure 22.

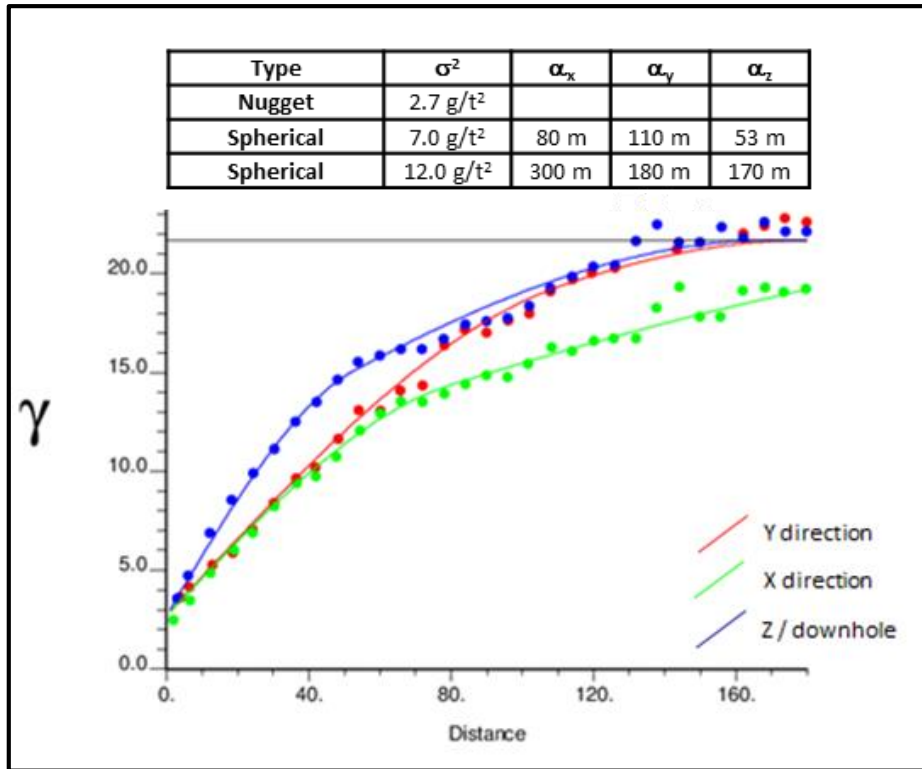


Figure 22 Scenario 1 experimental variogram, model variogram and parameters

Scenario 2's variogram model, in Figure 23, shows a relatively higher nugget effect of 26 % and the variance comes close to the sill at short ranges. Less variance at short ranges is seen in the down-hole variogram, although similar structure is seen in all three principal directions. The final structure is quite long, with variance reaching the sill at 60 m to 90 m.

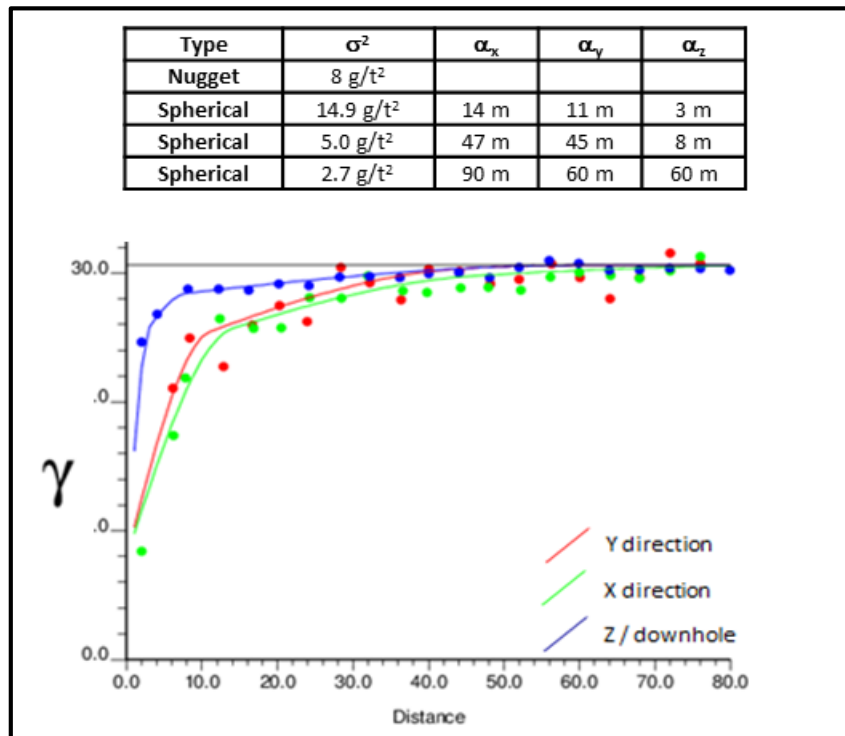


Figure 23 Scenario 2 experimental variogram, model variogram and parameters

A variogram model describes the co-variance relationship between sample pairs in a specific direction. Up to three orthogonal directions can be modelled which captures the spatial variability of the grade seen within the population. The variogram model is based on the sampling information, which is a representative subset of the population. Since this deposit was hypothetical, it is possible to compare the variogram model of the simulation representing the entire population against the variogram models of the sampled subset, to check how successfully the sampling captured the simulation or “reality”.

The impact of the variogram model successfully capturing grade variability of the deposit is reflected in the panel estimate. If the variogram model does not reflect the true grade variance of the deposit, it could lead to either too little or much high grade continuity in the estimated model.

Experimental variograms of the realisations were created for each distribution, which reflect the true variance of the deposit. These experimental variograms are compared with the model variograms created from the re-sampled realisations.

For Scenario 1, the modelled down-hole variogram shows greater continuity than the actual deposit variance. The effect of this could be seen in smoother than expected grade estimates in the Z-direction. In the X and Y directions, the variance modelled is lower than the true variability. The estimated nugget effect and sill value in the X-direction is close to the actual values. The comparison between orthogonal variograms of the simulation, representing “actuals”, and modelled variograms (based on the sample data) for Scenario 1 are seen in Figure 24.

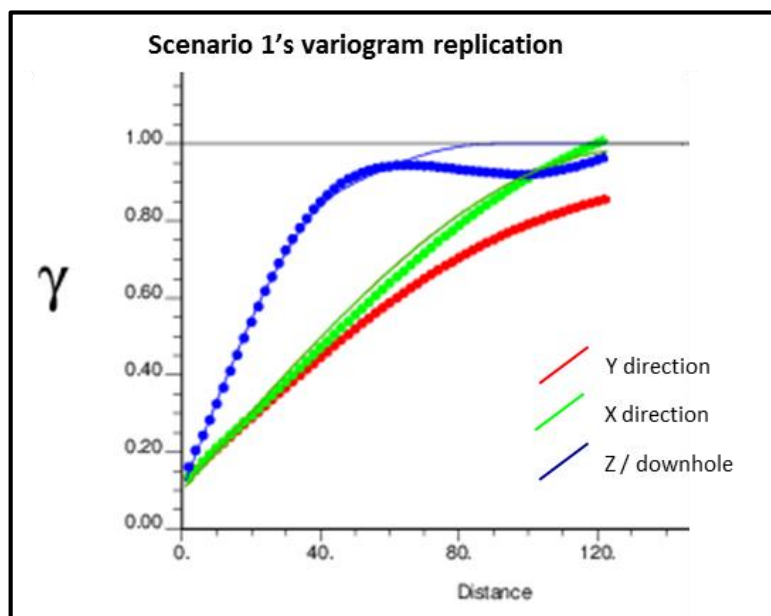


Figure 24 Simulation experimental variograms with sampled model variogram, for Scenario 1

For Scenario 2, the modelled down-hole variogram shows good variance reproduction in all directions. This reproduction indicates that the sampling adequately captures the grade variability of the simulated deposit. A comparison of the orthogonal experimental variograms

of the simulation versus the modelled variograms of the sample data can be seen in Figure 25.

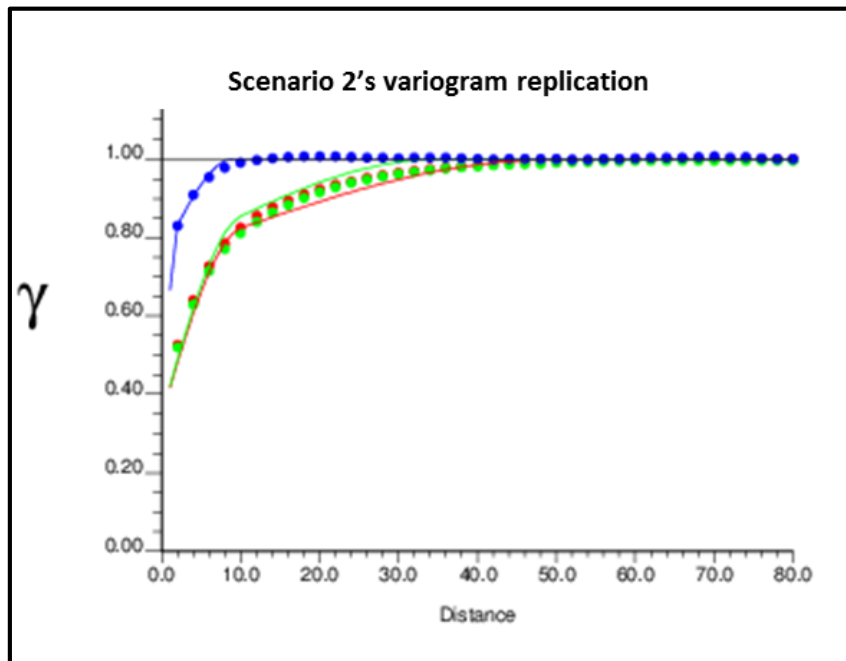


Figure 25 Simulation experimental variograms with sampled model variogram, for Scenario 2

3.6 Panel estimation

UC requires a robust estimate of the panels, which conditions a distribution of SMU grades within that panel. A poor estimate of the panel grade will result in an incorrect distribution of SMU grades in the UC model. In this respect, estimating the panel grade accurately is the most important aspect of the UC work flow.

The panel estimation was carried out using OK. The panel block model contained 1920 panels, and has the physical dimensions shown in Table 2.

Table 2 Panel block model dimensions

	Model Dimension	Block Size	Number of panels
X	600 m	50 m	16
Y	800 m	50 m	12
Z	200 m	20 m	10

3.6.1 Discretization points

The number of panel discretization points in the Z direction was chosen based on the down-hole composite size. Exactly ten times 2 m composite samples fall within the Z-block size; therefore, ten discretization points in the Z direction were used. The QKNA technique (Vann et al., 2003) was used to determine the number of discretisation points in the X and Y directions. This technique considers the change in average point variance within the block $\bar{\gamma}(v, v)$ and the change in the block variance (BV), assessed for an increasing number of

discretisation points in each direction. Once the results have stabilised (i.e. there is little change in $\bar{\gamma}(v, v)$ and BV), the number of discretisation points is chosen.

This analysis also demonstrates Krige's relationship, with the sum of the variance within the blocks and the block variance equal to the total variance of the deposit, or variogram sill $\gamma(h)$. Formulas for Krige's relationship are given in chapter 2.5.4

Such analyses were carried out for Scenario 1 (Figure 26) and Scenario 2 (Figure 27), and the results for both showed a stabilisation at ten discretization points in the X and Y direction. In total, there were one thousand discretisation points per panel as ten points were chosen in the X, Y and Z directions.

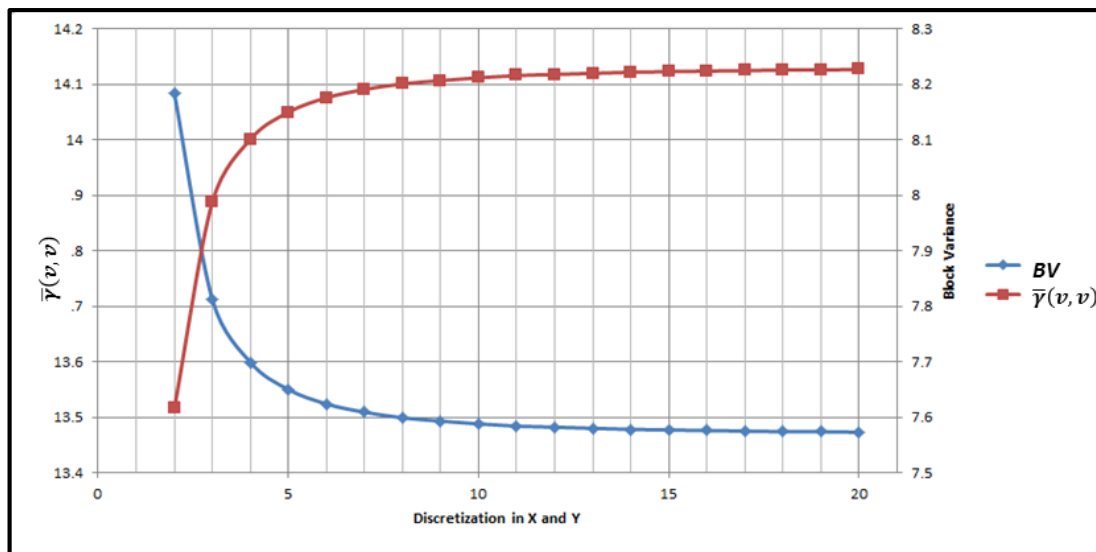


Figure 26 Panel discretisation point analysis for Scenario 1

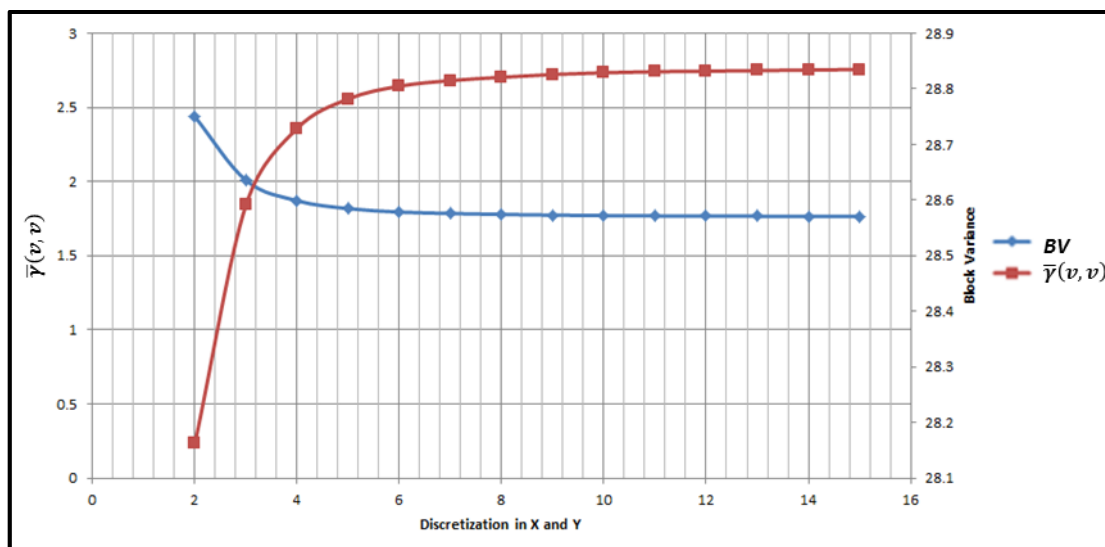


Figure 27 Panel discretisation point analysis for Scenario 2

3.6.2 Measure of confidence

UC has no measure of the confidence of estimates, and thus the only measure of UC's confidence, is confidence in the panel grade estimate (De-Vitry et al., 2007). Typically, the confidence of a mineral resource is made up of information from exploration and sampling, which is used to justify grade continuity (SAMREC Code, 2007). As the deposit data is simulated, the only indicators of confidence that may be applied are geostatistical factors. For the purposes of determining the robustness of this estimate, the slope of regression was used.

The slope of regression is the regression relationship between the true and the estimated block grade, which is an indicator of conditional unbiasedness (Vann et al., 2003). Slopes of regression values fall between 0 and 1, with higher values indicating that the estimated grade is near to the true grade. The slope of regression formula is as follows:

$$\text{Slope of Regression} = \frac{\text{Cov}(Z_v, Z_v^*)}{\text{Var}(Z_v^*)} = \frac{\bar{C}(V, V) - \sigma_k^2 + \mu}{\bar{C}(V, V) - \sigma_k^2 + 2\mu}$$

Where:

$\text{Cov}(Z_v, Z_v^*)$	Covariance of true Z and estimated Z^*
$\text{Var}(Z_v^*)$	Variance of estimated Z^*
$\bar{C}(V, V)$	Block variance / variance of true Z_v
σ_k^2	Kriging variance
μ	Lagrange multiplier

Parameters that influence on the slope of regression are the number of samples found in the search neighbourhood, the block size being estimated and the variogram model. The block size (Table 2) and the variogram model (chapter 3.5) are fixed, leaving the number of samples used for the estimation as the greatest influencer on the resultant slopes of regressions.

Using the QKNA technique, the number of samples used for estimating a block (i.e. samples in the kriging neighbourhood) was optimised. The chosen number of samples attempts to achieve a slope of regression value close to 1 (i.e. conditionally unbiased), without introducing a significant amount of negative weights and without over-smoothing the estimates. The sum of negative weight per block was not allowed to exceed 5 %, and smoothing of the estimate was monitored by assessing the kriging efficiencies. The QKNA process is run iteratively, until a satisfactory result is achieved. The final panel model had Scenario 1's panel model with a mean slope of regression of 0.97, while Scenario 2's panel model had an average slope of regression of 0.72. Plans of the OK models showing slopes of regression are shown in Appendix B: Slope of Regression.

The differences between the mean slopes of regressions for the two models are based on the sample spacing relative to the modelled variograms, as all other parameters are identical for the two scenarios. Scenario 1's variogram displayed significantly higher continuity relative to the sample spacing than Scenario 2, which results in better grade predictions for Scenario 1. For Scenario 1, variogram ranges are 170 m – 300 m, with an average sample spacing of 40 m. Therefore for Scenario 1, there are 4 to 7 samples occurring within the range of the variogram. Conversely for Scenario 2, the variogram ranges are 60 m – 90 m, with the same sample spacing means that 1 to 2 samples occur within the variogram range.

3.6.3 Plots of panel estimate

Figure 28 and Figure 29 show plan views of Scenario 1 and Scenario 2s panel estimates, at surface elevation. Comparing these to the simulated data sets (Figure 14 and Figure 15), the grade smoothing effect and reduction of variance from the OK is evident.

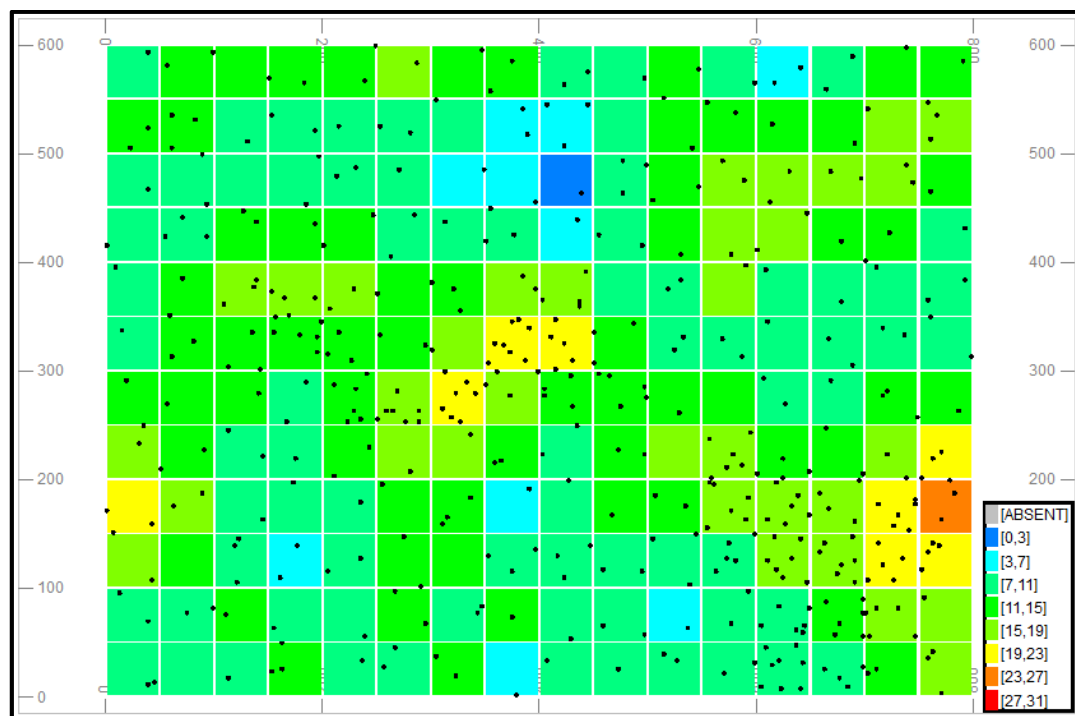


Figure 28 Surface plan view of ordinary kriged panel estimate for Scenario 1

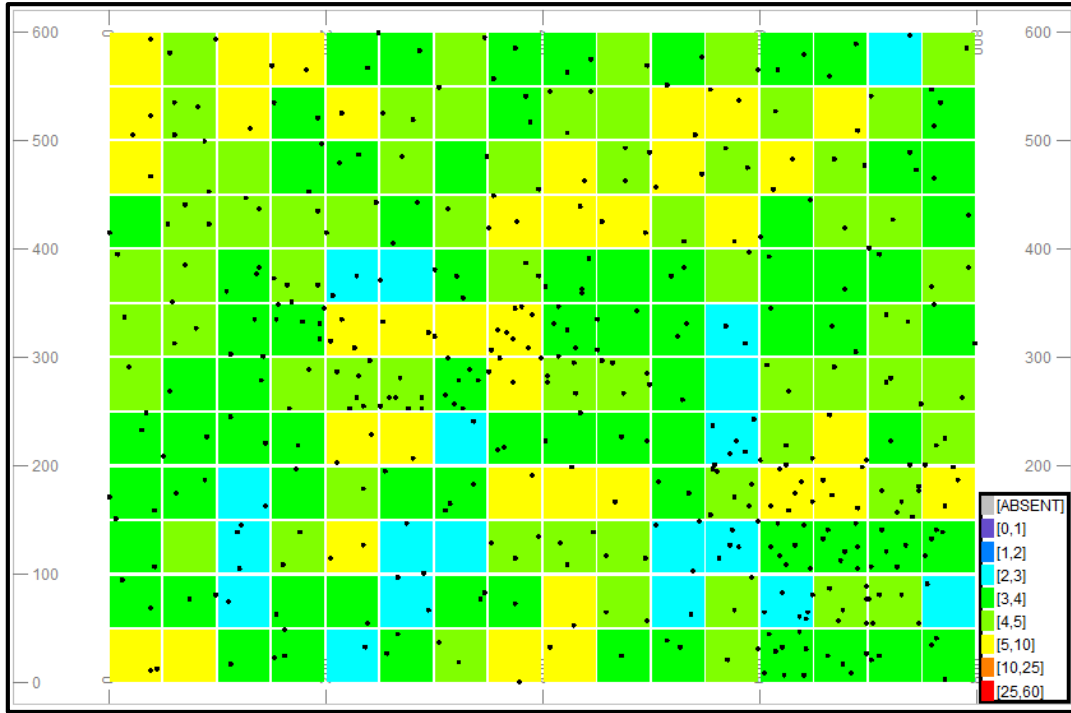


Figure 29 Surface plan view of ordinary kriged panel estimate for Scenario 2

3.7 Bi-Gaussianity

UC relies upon the assumption of bi-Gaussianity of the transformed grade data (Rivoirard, 1994; Humphreys, 1998; Assibey-Bonsu and Krige, 1999). Bi-Gaussianity means that any linear combination of the Gaussian transformed data is also Gaussian. Bi-Gaussianity may also be referred to as bivariate normality. Schofield (1988) describes some practical tests on sample data to check for Bi-Gaussianity. These tests were run on Scenario 1 and Scenario 2's sample data. Further to this, both data sets were tested for consistency with a diffusion model.

3.7.1 Madogram variogram ratio

Under bi-Gaussian conditions, the ratio of the normal score madogram over the normal score variogram is constant (Verly et al., 1986; Schofield, 1988). The formula depicting this relationship is as follows:

$$\frac{\gamma_{1(h)}}{\sqrt{\gamma_2(h)}} = \frac{1}{\sqrt{\pi}}$$

The ratio was plotted for Scenario 1 (Figure 30) and Scenario 2 (Figure 31) to beyond the range of the variogram, and satisfactorily shows that the transformed data are approximately bi-Gaussian.

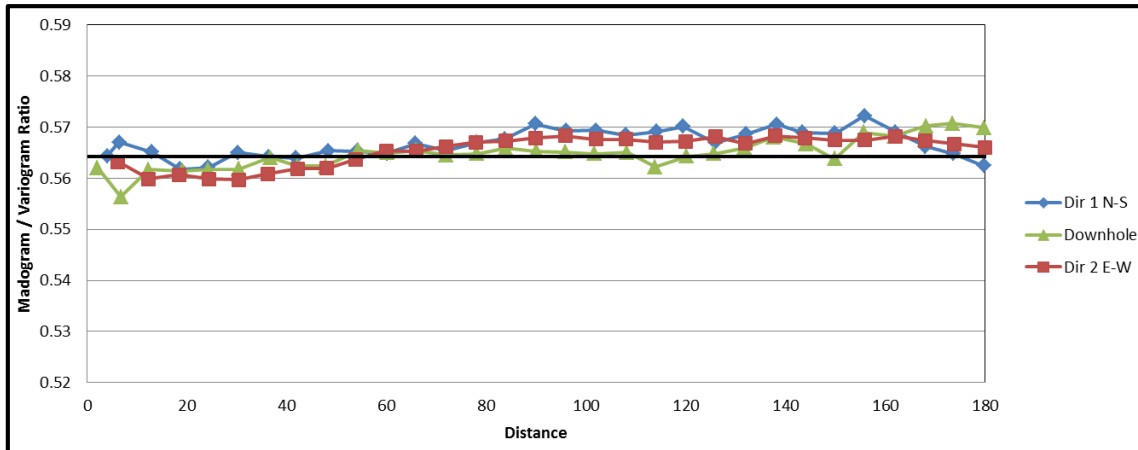


Figure 30 Madogram / variogram ratio test for Scenario 1

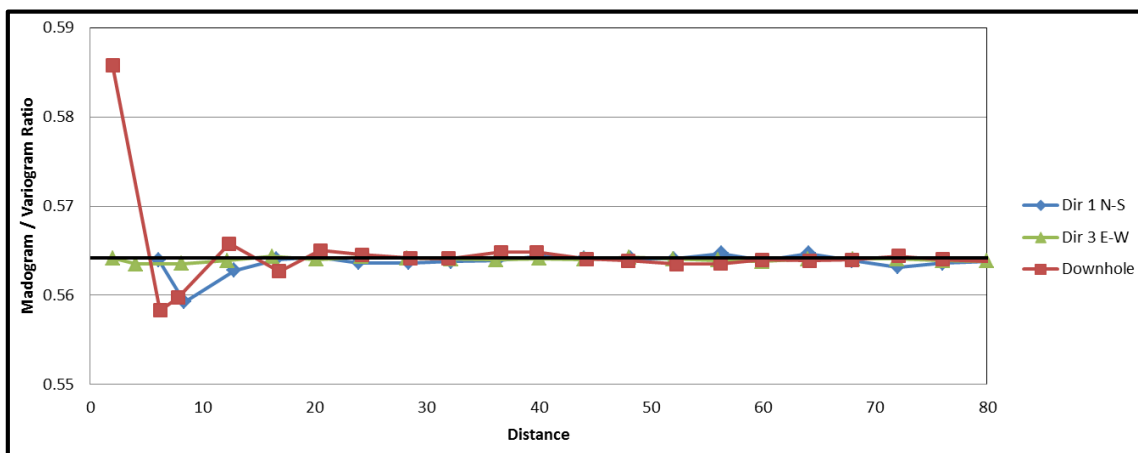


Figure 31 Madogram / variogram ratio test for Scenario 2

3.7.2 Proportion effect

The proportion effect occurs when there is a linear or other non-random relationship between the mean and the variance of a sample set. Testing for the proportional effect is a test for multivariate normality, which is not strictly a requirement for UC. However, proving multivariate normal also supports bi-Gaussianity. The test involves proving that no proportional effect exists for the normal score transforms of the data.

Both the original unit and the normal score data was tested for the proportional effect by comparing the standard deviation of the data against the mean of the data within large panels of 50 m x 50 m x 20 m. Tests results for Scenario 1 are shown in Figure 32, where no proportional effect was evident in the original unit data or the normal scores of this data. This is consistent with multivariate Gaussian conditions.

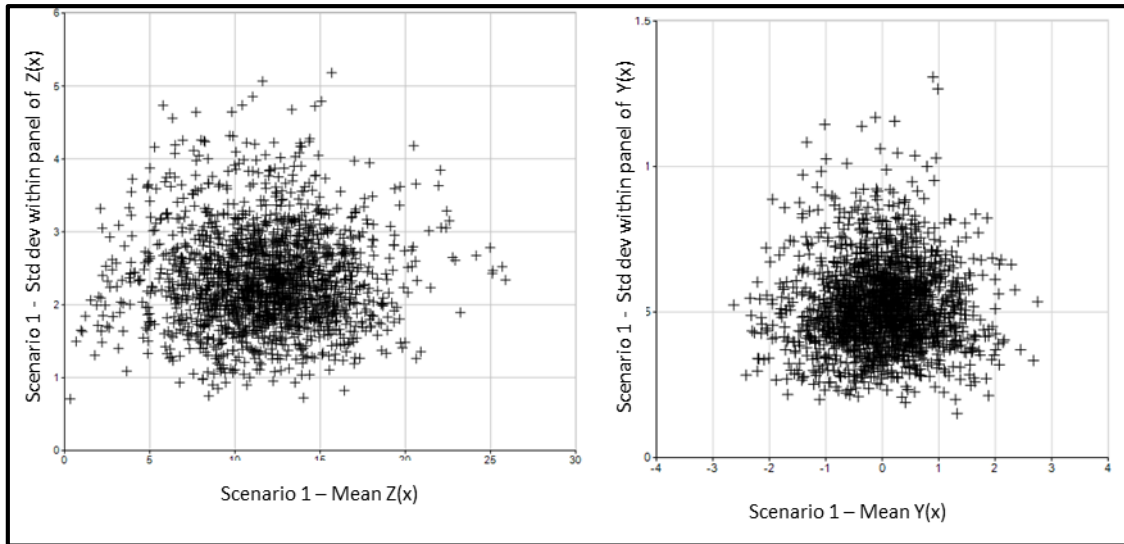


Figure 32 Test for proportional effect for Scenario 1 original grades (left) and normal scores (right)

Proportional effect tests were also run on Scenario 2, shown in Figure 33. A proportional effect exists for the original unit data. This effect was removed when taking the normal scores of this data, which is sufficient to support multivariate normality.

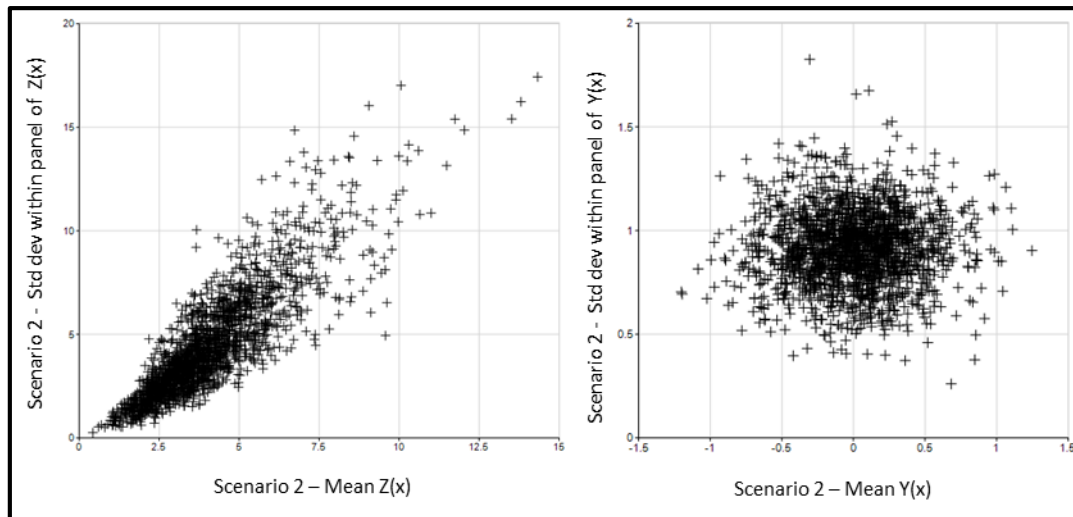


Figure 33 Test for proportional effect for Scenario 2 original grades (left) and normal scores (right)

3.7.3 Diffusion model

The Gaussian model is suitable for conditions with an edge effect or a diffusion model, where there is a continuous transition between neighbouring zones (Rivoirard, 1994; Vann and Guibal, 1998). To test for a diffuse model, one can test for intrinsic correlation of the data. A positive result (where intrinsic correlation is evident) proves the existence of a mosaic model (i.e. a model with no edge effect). Conversely, disproving intrinsic correlation proves the existence of a diffuse model.

To test for intrinsic correlation, the ratios of extreme indicators classes of the normal score data are compared. Intrinsic correlation would result in a correlation between these

indicators pairs. The ratio of the 25th percentile and 75th percentile were compared, in three principal directions (X, Y and Z), shown in Figure 34 and Figure 35.

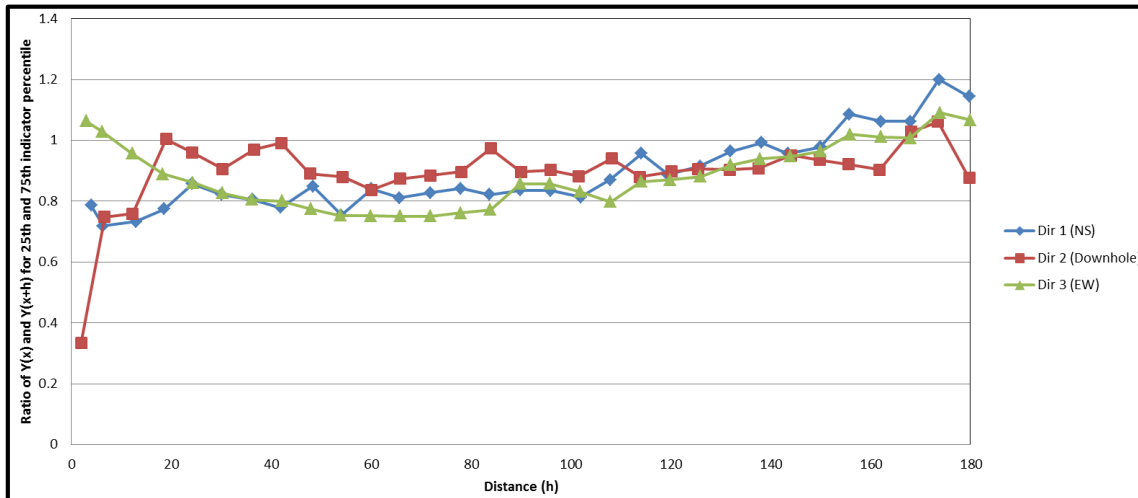


Figure 34 Ratio of extreme indicators to test for intrinsic correlation for Scenario 1

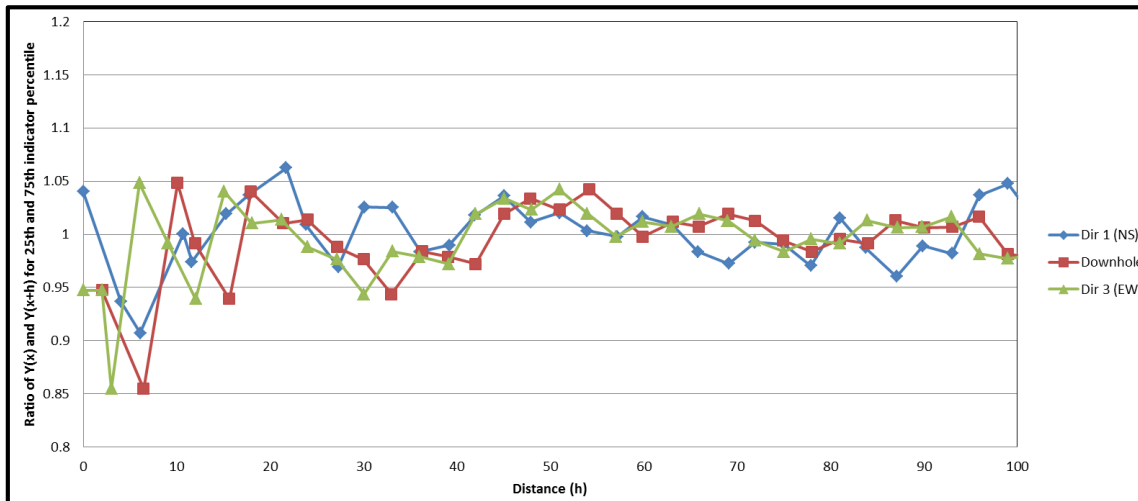


Figure 35 Ratio of extreme indicators to test for intrinsic correlation for Scenario 2

Intrinsic correlation would be evident if all variograms were proportional to each other, and such a result is consistent with a mosaic model with no edge effect. Neither Scenario 1 nor Scenario 2 display a correlation between the 25th percentile and 75th percentile indicators, which proves no intrinsic correlation. By deduction, this proves a diffusion model, which is suitable for uniform conditioning.

3.7.4 Results of testing for bi-Gaussianity

The grade samples proved consistent with an underlying conditions of bi-Gaussianity and multi-Gaussianity of the normal score transformed grades.

As the data was sampled from realisations using a sequential Gaussian simulation, it was expected that both Scenario 1 and Scenario 2 met the conditions of bi-Gaussianity required for Uniform Conditioning. While the histogram parameters use to guide the sequential

Gaussian simulation for Scenario 1 and Scenario 2 were intended to generate realisations of normal and a log-normal data, the normal scores of both of these data sets have a Gaussian distribution. The normal score transform of any realisation generated using Sequential Gaussian Simulation has a Gaussian distribution.

Despite this prediction, the test for bi-Gaussianity and multi-Gaussianity were executed to prove that the data sets were consistent.

3.8 Uniform conditioning

In preparation for running uniform conditioning, a set of Hermite polynomial coefficients were fitted to Scenario 1's grade distribution and Scenario 2's grade distribution, and change of support coefficients were determined. The uniform conditioning was then carried out, which produced conditional SMU proportions and grades above cut-off from panel estimates. Finally, these were converted into a set of SMU grades for a panel.

3.8.1 Determining change of support coefficients

Recall that the change of support coefficient: R for panel change of support, and r for SMU change of support. R is the correlation between Gaussian equivalent point grades $Y(x)$ and Gaussian equivalent panel grades $Y(V)$. Similarly, r is the correlation between Gaussian equivalent point grades $Y(x)$ and Gaussian equivalent SMU grades $Y(v)$.

The r -coefficient is calculated from the theoretical block variance of the SMU and the Hermite coefficients (see chapter 2.5.4 for formulas). The block variance was calculated from the sill minus the average dispersion variance within an SMU block (where the dispersion variance was calculated from the average variogram value between multiple discretization points in an SMU and the sill value was determined from the variogram). A GSlib program (gammabar.exe) was used to calculate this relationship and the results were verified against results generated in Datamine Studio 3, and results from the two packages were identical. The SMU block variances used are shown in Table 3.

Table 3 SMU change of support coefficients for calculating ' r '

Variable	Scenario 1	Scenario 2
Sample variance (sill)	21.7	30.6
Dispersion variance in SMU	4.22	24.45
Block variance in SMU	17.48	6.15

The R -coefficient can be theoretically calculated in a similar manner as the r -coefficient (see chapter 2.5.4 for formulas), using the panel block variance and Hermite coefficients. However, instead of using a theoretical block variance for the determination of R , a preferred approach was to take a direct measurement of the variance of the estimated panel grades and use this value for the panel block variance. A further improvement on this was to group the panel

grades according to how well the blocks were estimated, and use the direct grade variance of these groups of estimates.

The third approach should result in the best UC result, as consideration is taken for the panel information effect (i.e. how well panels are estimated). When there is more (or less) certainty in the estimate, the resultant UC will reflect this by using an appropriate distribution of SMU grades.

Any geostatistical indicator of “goodness” of the estimator can be used to group the panel grade estimates e.g. slope of regression, kriging variance or kriging efficiency. The indicator of “goodness” is plotted against panel grades in order to make a decision about grouping.

For Scenario 1, panel grades were plotted against kriging variance. Three groups were chosen, with grouped grade variances of 6.17 g/t², 11.20 g/t² and 13.39 g/t² (Figure 36). These values are the grade variances at panel support used in the calculation the R-coefficients.

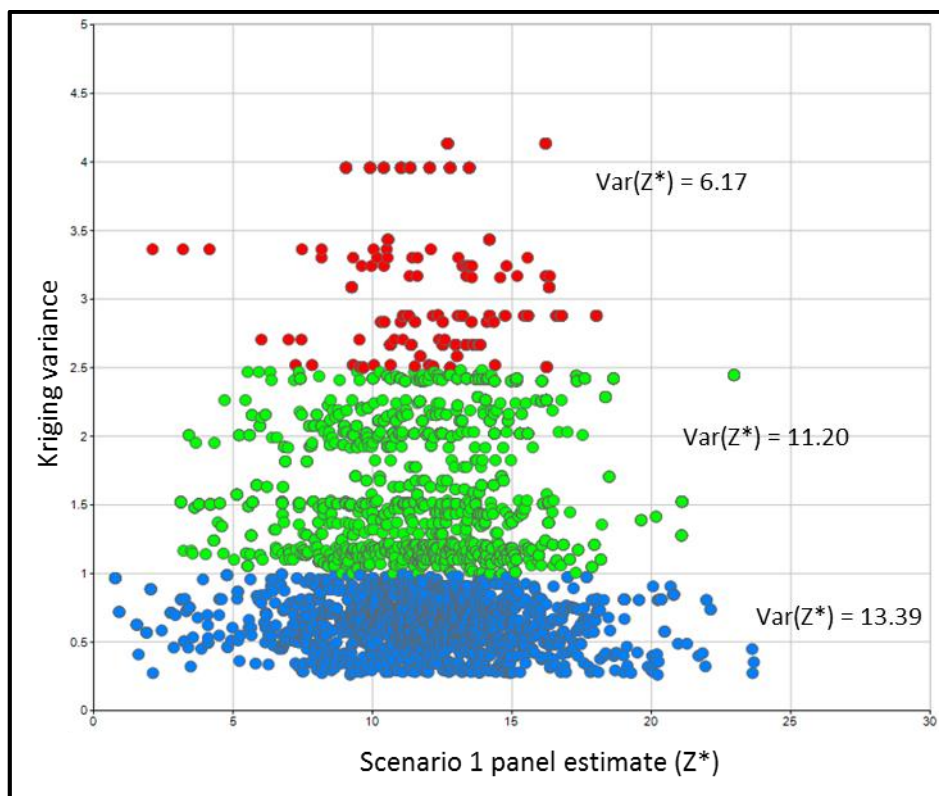


Figure 36 Scenario 1 grouping of panel estimates for R coefficient determination

Similarly, as for Scenario 1, Scenario 2 panel estimates were grouped (Figure 37). These estimates were grouped according to slopes of regression, and the variance of these grouped grade estimates was 1.04 g/t², 0.93 g/t² and 0.91 g/t², which were used in the calculation of R-coefficients.

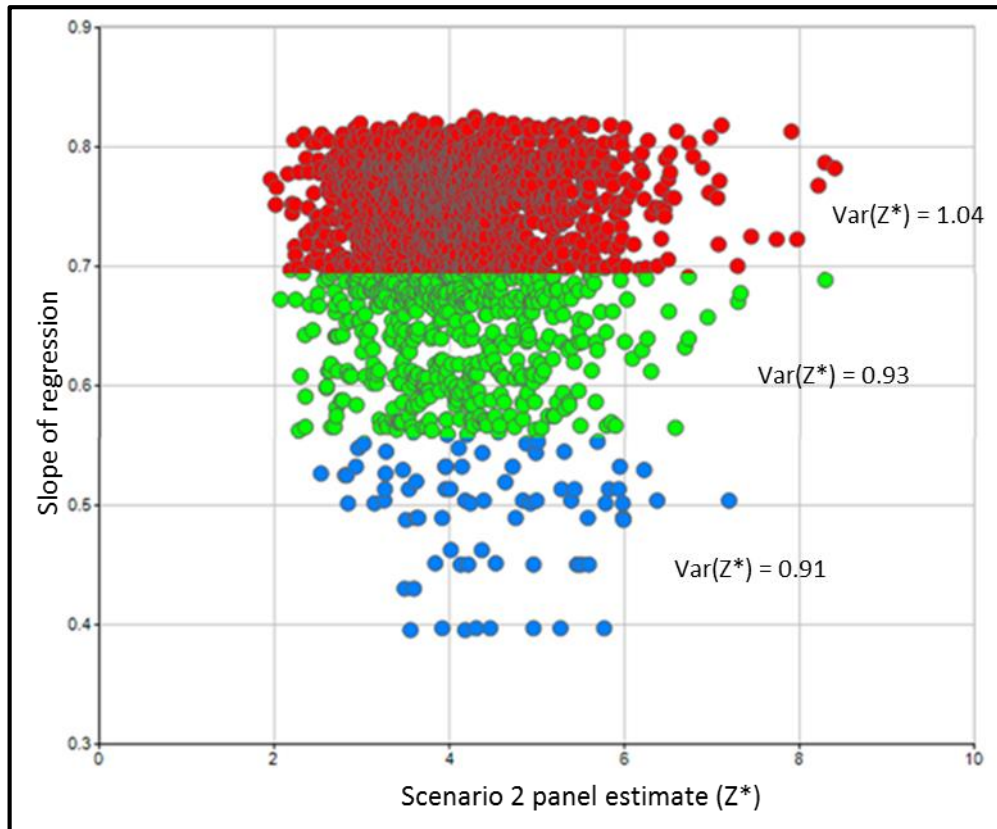


Figure 37 Scenario 2 grouping of panel estimates for R coefficient determination

Using the appropriate SMU and panel variances, respectively calculated theoretically and directly from the data, r -coefficients and R -coefficients are determined for each distribution. A GSLib executable (preUC.exe) was used to generate a set of Hermite polynomial coefficients from the normal score transform, as well as the r -coefficients and R -coefficients, from the provided variances. The r -coefficients and R -coefficients for the grouped data is shown in Table 4 and Table 5.

Table 4 Change of support coefficients for Scenario 1

	Scenario 1 – R Panel change of support coefficient	Scenario 1 – r^* SMU change of support coefficient	Change of support ratio R/r^*
Group 1	0.791	0.903	0.876
Group 2	0.723	0.903	0.801
Group 3	0.537	0.903	0.595

Table 5 Change of support coefficients for Scenario 2

	Scenario 1 – R Panel change of support coefficient	Scenario 1 – r^* SMU change of support coefficient	Change of support ratio R/r^*
Group 1	0.230	0.503	0.457
Group 2	0.217	0.503	0.431
Group 3	0.215	0.503	0.427

3.8.2 Hermite polynomials

Using the above data, a set of thirty Hermite coefficients were generated using a GSLib executable (preUC.exe) to fit the Hermite polynomials to the distributions. Any function can be fitted by Hermite coefficients as a weighted sum of Hermite polynomials, and these are calculated through a Fourier analysis.

Both sets of coefficients are shown in Table 6.

Table 6 Hermite coefficients

	Hermite coefficients (Scenario 1)	Hermite coefficients (Scenario 2)
ϕ_0	11.75	4.11
ϕ_1	-4.62	-4.38
ϕ_2	0.07	2.98
ϕ_3	-0.08	-1.38
ϕ_4	0.12	0.22
ϕ_5	0.17	0.31
ϕ_6	-0.01	-0.33
ϕ_7	-0.04	0.12
ϕ_8	-0.08	0.06
ϕ_9	-0.02	-0.12
ϕ_{10}	0.08	0.06
ϕ_{11}	0.03	0.04
ϕ_{12}	-0.05	-0.09
ϕ_{13}	-0.02	0.03
ϕ_{14}	0.02	0.06
ϕ_{15}	0.00	-0.07
ϕ_{16}	0.01	-0.01
ϕ_{17}	0.01	0.07
ϕ_{18}	-0.03	-0.03
ϕ_{19}	-0.01	-0.04
ϕ_{20}	0.04	0.05
ϕ_{21}	0.01	0.01
ϕ_{22}	-0.03	-0.04
ϕ_{23}	0.00	0.01
ϕ_{24}	0.02	0.03
ϕ_{25}	0.00	-0.02
ϕ_{26}	0.00	-0.01
ϕ_{27}	0.01	0.02
ϕ_{28}	-0.01	0.00
ϕ_{29}	-0.01	-0.02
ϕ_{30}	0.01	0.00

Each set of Hermite polynomial coefficients (shown in Table 6) were verified in three ways. The first check was that the ϕ_0 coefficient is equal to the mean of the sample data. The second check is that the sum of the squared coefficients (excluding the ϕ_0 coefficient) equals the sample variance. The third check involved expanding the Hermite polynomial (see Rodriguezes formula in chapter 2.5.3) to display the Gaussian anamorphosis model, and check that this is consistent with the point sample normal score transform.

The normal score transform of the sample data is plotted Gaussian anamorphosis function in Figure 38 for Scenario 1. The shape of Scenario 1's model is typical of a normal distribution, as there is a near linear mapping from the original units to Gaussian units. For low (<3 g/t) and high grade (> 21g/t) values, there is a deviation of the Gaussian anamorphosis model from the point data, whereas there is a very good fit for rest of the grade distribution (3 – 21

g/t). The poor fit corresponds with areas where there are fewest grade samples, and conversely the best fit corresponds with areas where there are the most grade samples as can be confirmed with the sample histogram in Figure 18.

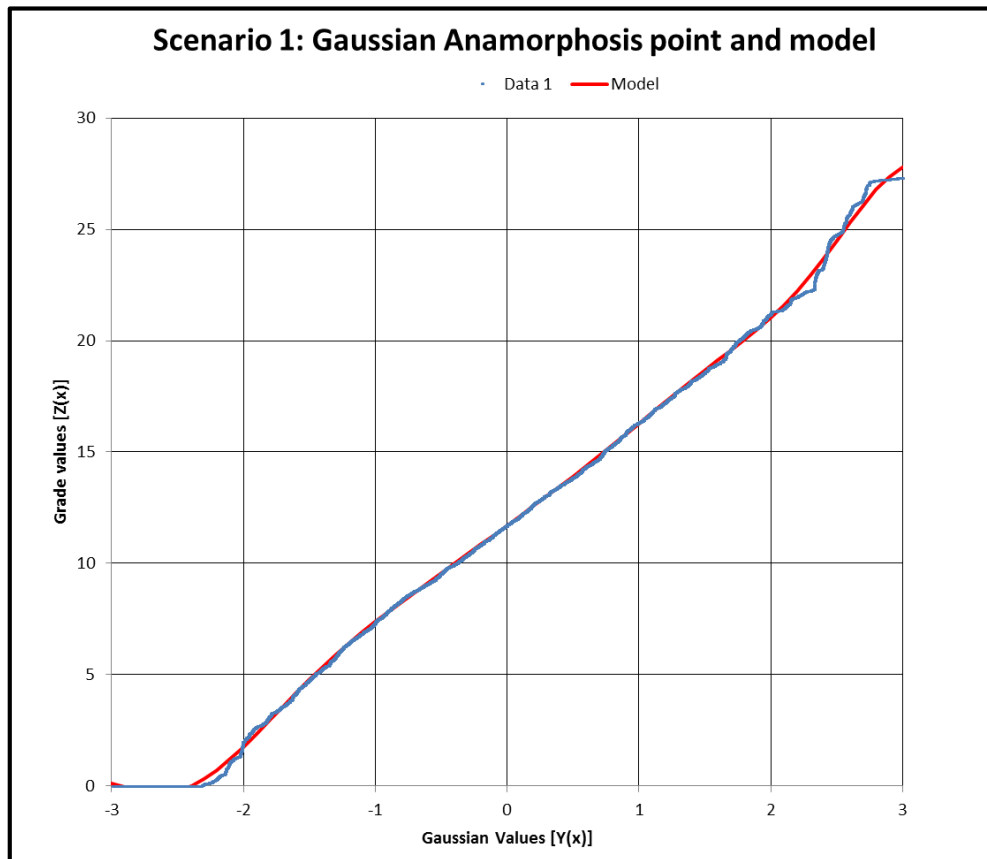


Figure 38 Scenario 1 sample normal score transform and Gaussian anamorphosis model

The normal score transform of the sample data is plotted Gaussian anamorphosis function in Figure 39 for Scenario 2. The curved shape of Scenario 2's model is typical of a log-normal distribution transformation to Gaussian units.

In Scenario 2, the Gaussian anamorphosis model fits the sample data well for low and medium grade values (<20 g/t) where there is sufficient data. The Gaussian anamorphosis model does not fit the sample data well for high grade values (> 20 g/t) where there are fewer data. This is confirmed by referring to Scenario 2's histogram in Figure 19.

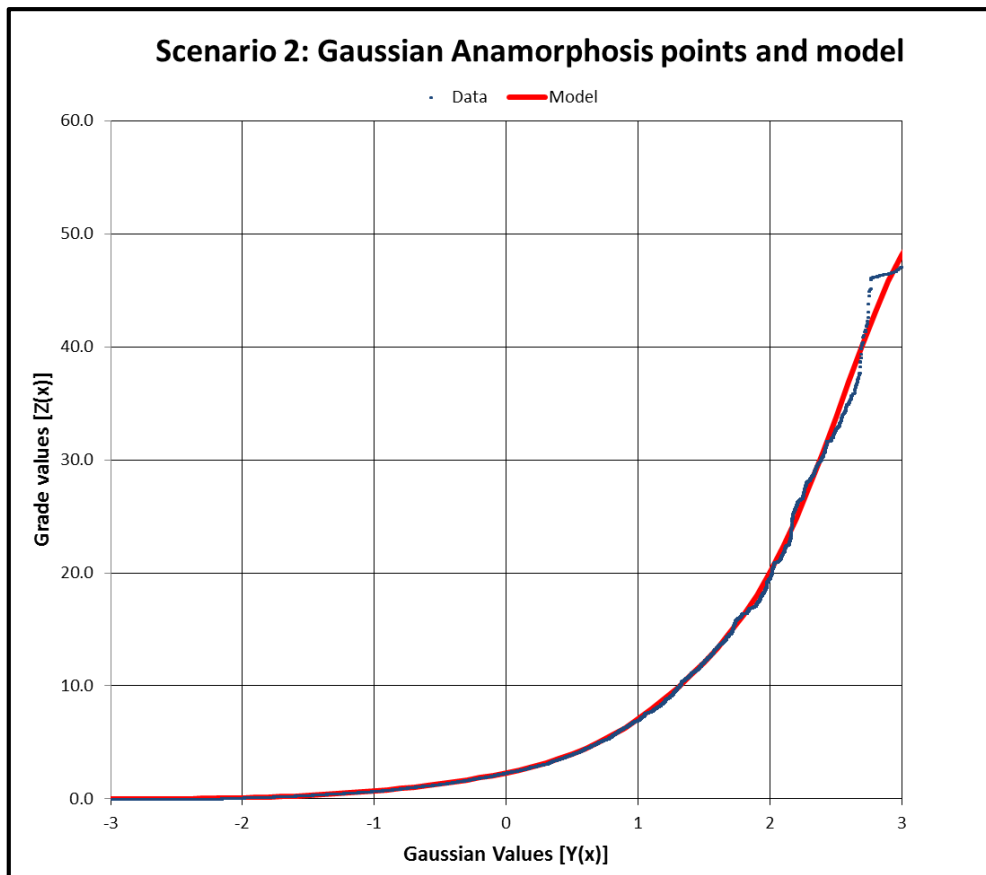


Figure 39 Scenario 2 sample normal score transform and Gaussian anamorphosis model

The shape of the Gaussian anamorphosis model is changed by the R/r ratio, and flattens as the support effect increases i.e. as the R/r ratio decreases. Figure 40 shows the Hermite polynomials, plotted for R/r ratios representing point support and change of support coefficients groupings as seen in Table 4 and Table 5.

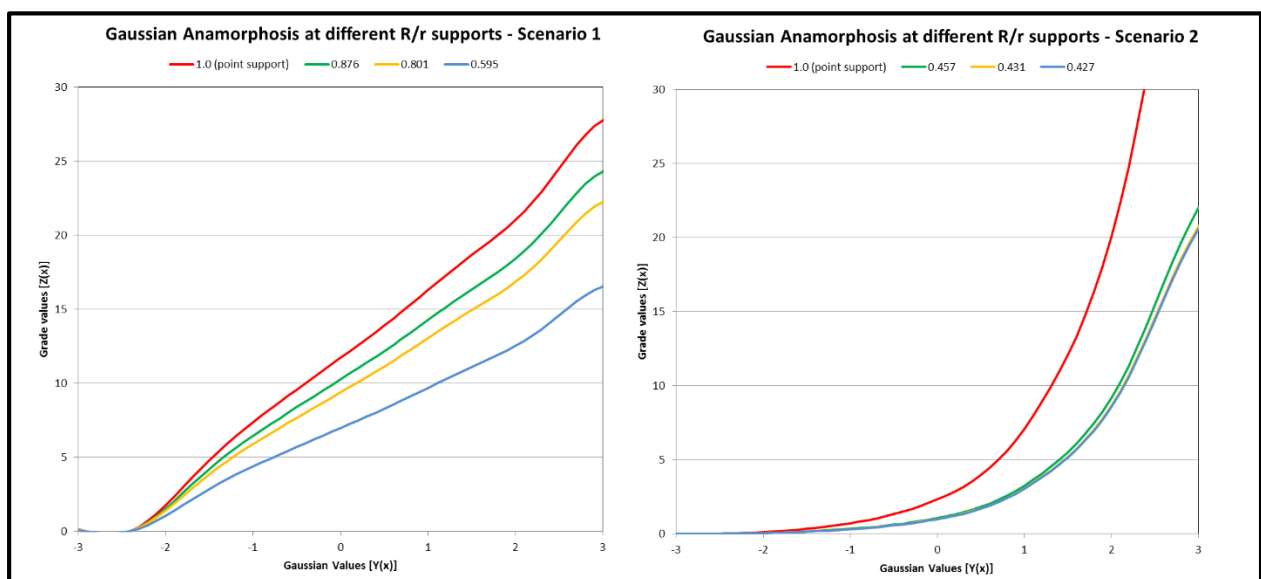


Figure 40 Gaussian anamorphosis model for multiple R/r ratios

3.8.3 Calculating conditional SMU distribution

As stated earlier; the correlation coefficients are as follows:

- r is the correlation between point grades $Y(x)$ versus SMU grades $Y(v)$;
- R is the correlation between point grades $Y(x)$ versus panel grades $Y(V)$.

The UC change of support uses the ratio of R/r to calculate the distribution of conditional SMU grades $Y(v)$ from an estimated panel grade $Y(V)$. R/r is the correlation between $Y(v)$ and $Y(V)$, and these correlation values are shown in Table 4 and Table 5. The schematic shown in Figure 6 shows the $Y(v)$ to $Y(V)$ relationship, and how the distribution of $Y(v)$ values are conditional on a panel grade.

The UC procedure was executed using a GSLib executable (uc.exe). The inputs into this program are a set of panel grades (see chapter 40), Hermite coefficients, as well as SMU change of support coefficient r and Panel change of support coefficients R . The GSLib program generates a series of grades and proportions above a specified cut-off grade. Metal above the cut-off grade is calculated from the product of proportion and grade. These proportions and grades above cut-off are determined from the conditional $Y(v)$ distribution.

An example of a UC model with grades and proportions above cut-off is given in Table 7. The example shows proportions and grades above cut-off calculated for cut-off grades between 0 g/t to 20 g/t, in 0.5 g/t increments. Values between cut-off grades of 1.5 g/t to 19.5 g/t were omitted in order to keep the example simplistic.

Table 7 Example UC model as grades and proportions

PanelEst	Tc:0.00	Tc:0.50	Tc:1.00	...	Tc:20.00	Mc:0.00	Mc:0.50	Mc:1.00	...	Mc:20.00
4.87	1	0.91	0.79	...	0.00	4.87	5.32	6.10	...	0.00
10.8	1	0.99	0.98	...	0.05	10.8	10.8	10.9	...	19.6
17.3	1	1	0.96	...	0.21	17.3	17.3	17.5	...	32.1

Scenario 1 and Scenario 2's UC models were processed in 0.25 g/t cut-off increments, from a cut-off grades of 0 g/t to 28.0 g/t. The range of the cut-off grades should reflect areas of interest in the block model.

The format of a model as proportion above cut-off format is challenging to work with, as it is not compatible with general mine planning and scheduling software. A script was written to convert the model from the GSLib format into a series of SMU grades for each panel, using the Q-T curve methodology. The Q-T curve shows the logic used to calculate average grades for each SMU within a panel, assuming that there are 50 SMU within a panel (Figure 41).

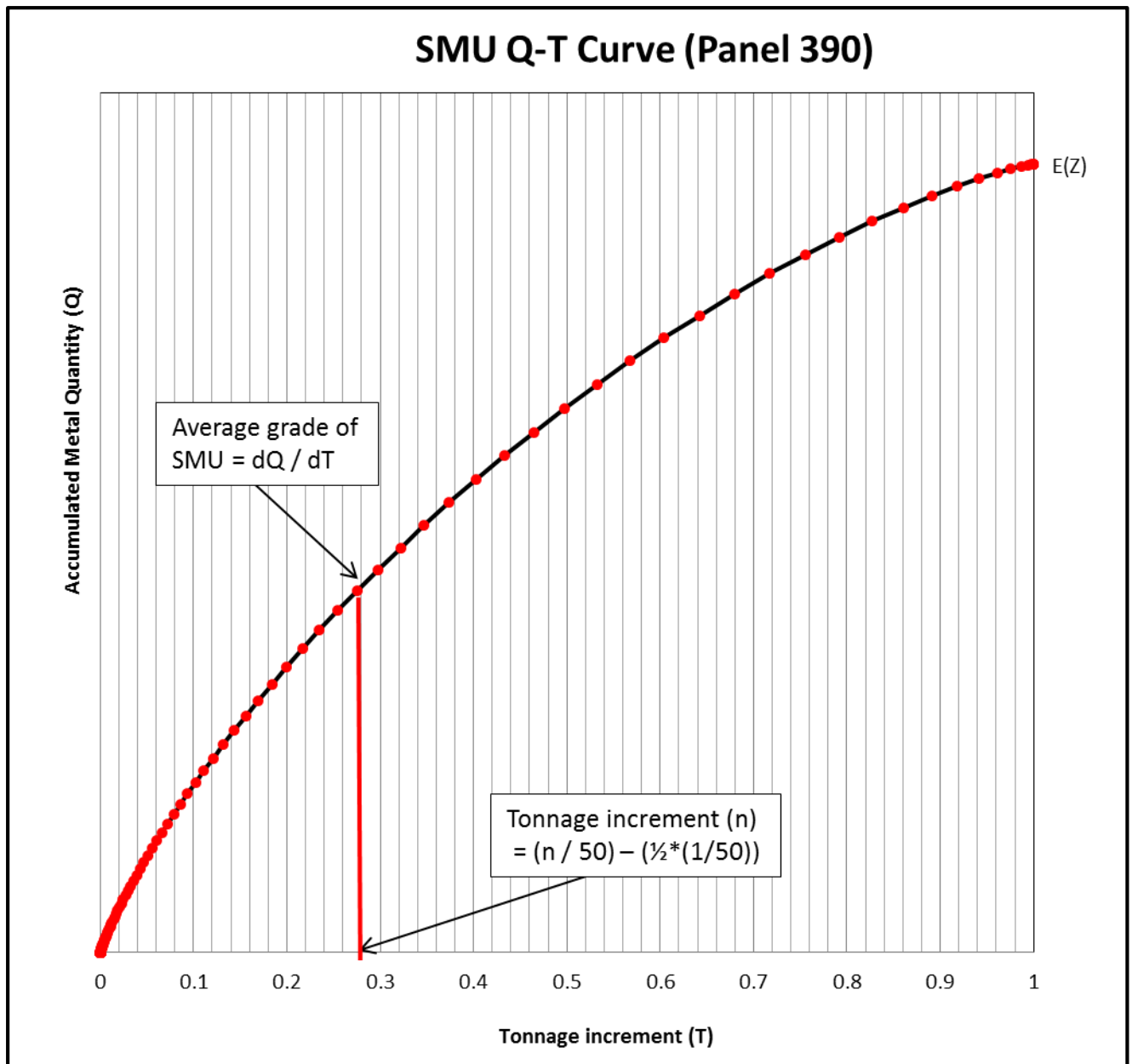


Figure 41 Q-T curve showing average grade calculation at tonnage increment

Finally, as there were three panel change of support coefficients (group 1, group 2 and group 3) for each distribution (Scenario 1 and Scenario 2), the UC procedure was run for each group and the final models merged based on the grouping of the panel estimates.

3.9 Localisation

Localising a UC model spatially positions the set of SMU grades within a panel, based on rankings provided by a local linear estimate. The advantage of having the spatial location of recoverable UC grade estimates is that the model provides a realistic assessment of recoverable resources, as well as honouring local grade variations. Additionally, as the LUC estimate is on a SMU block size, this estimate can be easily compared with other SMU

estimates. The localisation of UC estimates can be considered a practical enhancement to UC, which creates a more usable result.

The ranking of SMU grades is based on an OK estimate at SMU sized blocks. Within a panel, each SMU block is given a ranking from 1 to n, where n is 50 as there are 50 SMU per panel, and in order of highest to lowest kriged grade. The ranking is then applied to the set of SMU grades for a given panel.

The localisation procedure for this project was written into a JavaScript program that calculated the set of mean SMU grades for a panel from the UC model, ranked these grades from lowest to highest and allocated mean SMU grades to the spatial location of the equivalent-ranked OK grades.

The result of this localisation process is that the highest grade SMU, ranked according to the OK estimate, are mapped the highest average grade blocks, according to the UC conditional distribution. And visa-versa for the lowest grade blocks.

The success of the localisation is dependent on the data available for the ranking (Abzalov, 2006). Despite the seeming advantage of having recoverable grades on smaller block sizes, the LUC grades are not more accurate than the surrounding data. As such the localisation process cannot presume to know more about the distribution of grades within the panel than what the available data provides.

4. Discussion on Uniform Conditioning

The purpose of this study was to assess the performance of UC and LUC on normally distributed data with good grade continuity and reasonable data coverage relative to the variogram ranges (Scenario 1) and log-normally distributed data with poor data coverage relative to the variogram ranges and a high nugget effect (Scenario 2). GT curves show the optimistic presentation of grades and tonnages that are extractable material at given cut-off grades. The performance of UC is globally measured by how closely the UC estimate predicts the actual GT curves, and this is presented in terms of how well the estimate conforms to the GT curve of the actual simulated model. OK results are also compared with the UC results, as a benchmark for linear estimators.

The GT curves for selected panels, some of which were well estimated and some of which were poorly estimated, were assessed. This was to demonstrate how the effect of getting the panel grade right/wrong affects the resultant distribution of grades within a panel.

It is assumed throughout that the density of the model is 1000 kg/m³, making the tonnage and volume equal.

4.1 Global grade tonnage assessment

Figure 42 shows simulated (composited to an SMU scale), UC and OK (estimated at an SMU scale) grade tonnage curves for Scenario 1. It is of no consequence whether the UC or LUC grades and tonnages are displayed, as the localisation does not impact on the grade tonnage results. The UC prediction of tons and grades is very close to the actual simulated grades, and shows an improvement on the OK grade tonnage curve.

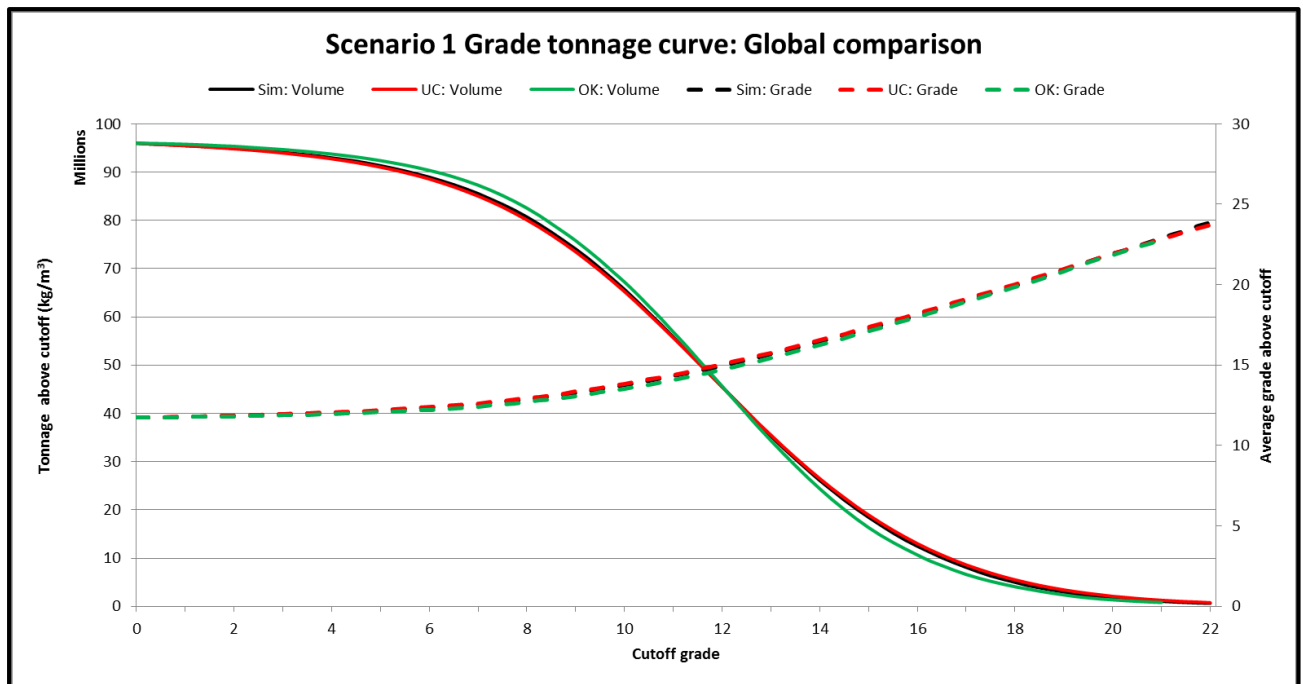


Figure 42 Scenario 1 global grade tonnage curves

In the case of normally distributed data where there is good data coverage (relative to the variogram ranges), as is the scenario for Scenario 1, OK performs adequately for determining recoverable resources. Slopes of regression for the OK model were generally close to 1, indicating a low conditional bias (which is reflected in the OK results being close to the actual).

However, UC outperforms OK in terms of estimating a recoverable resource, and is closer to the simulated reality. At relatively low cut-off grades, there is slightly less tonnage than predicted for the both the OK and UC model, but the UC results will be closer to the actual.

Figure 44 shows the grade tonnage curves for the log-normal distribution (Scenario 2), comparing the simulation (composited to SMU sized blocks), UC and OK (estimated at an SMU block size) grade tonnage results. The simulated reality shows a steep decline in tonnage or volume as the cut-off grade increases. OK generates a poor estimation of the grade and tonnage extractable for any cut-off grade. This poor adherence to the simulated GT curve can be explained by a conditional bias, which was expected, as the slopes of regression of Scenario 2's panel estimate were poor. UC gives a better result than OK, but the resultant estimation of grades and tonnage does closely conform to the reality.

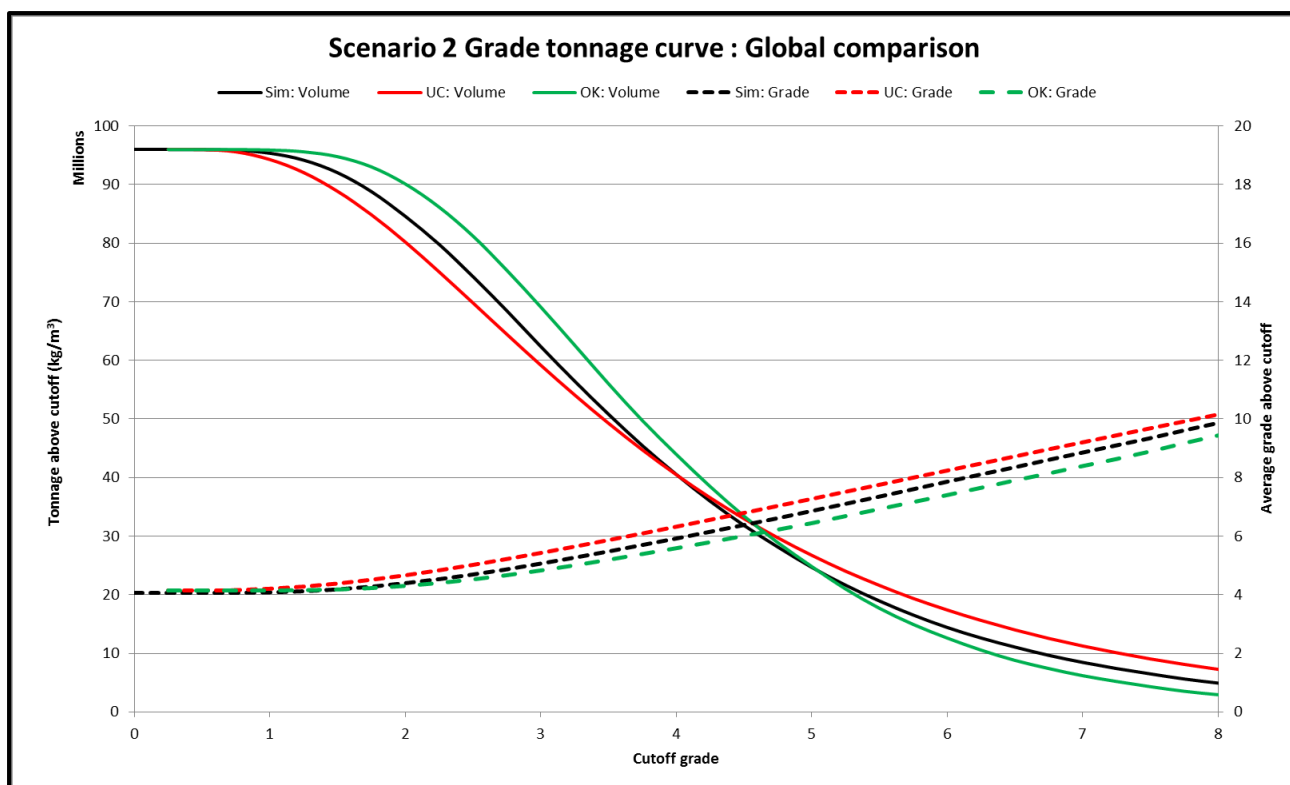


Figure 43 Scenario 2 global grade tonnage curves

A recoverability curve for Scenario 1 is shown in Figure 44, which illustrates the average grade of the volume of recoverable material. The graphic shows that as the selectivity increases i.e. high grade areas are targeted, the average grade of planned material will be slightly higher than the OK model predicts. The UC grade is closer to the reality or simulated grade.

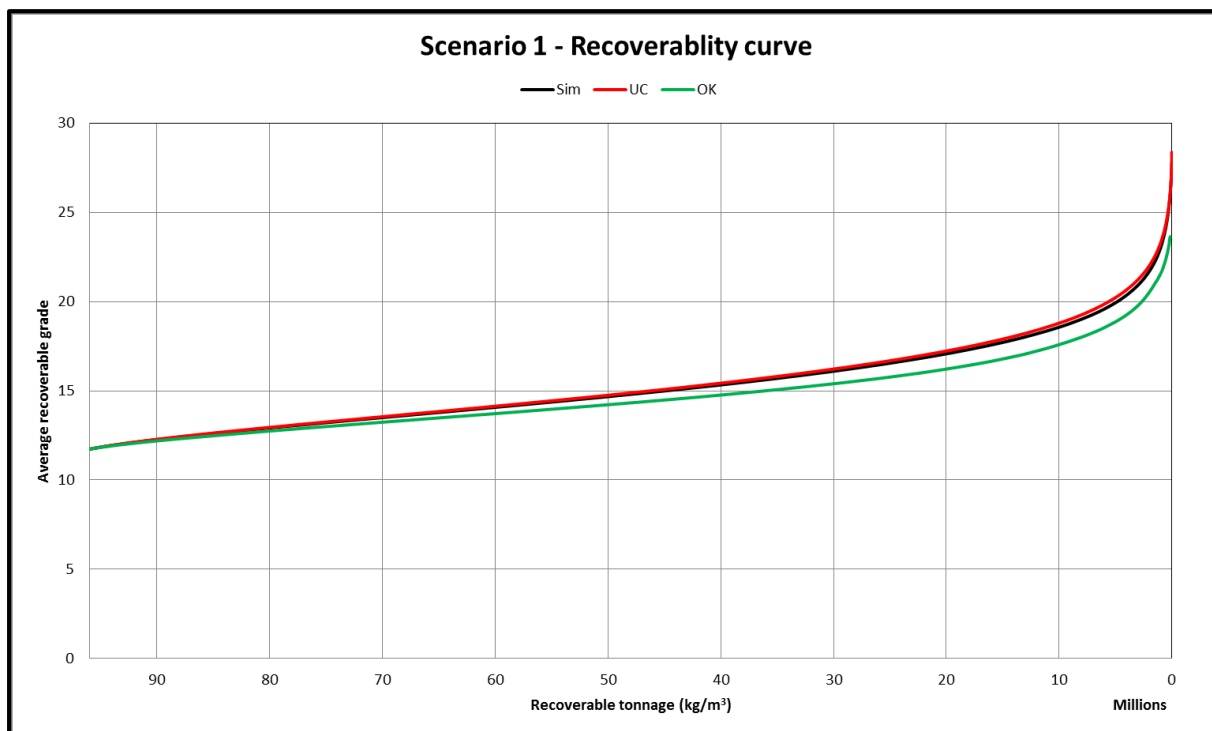


Figure 44 Scenario 1 recoverability curve

A similar recoverability curve was plotted (Figure 45) for the log-normal distribution (Scenario 2). This shows a significant discrepancy in predicted tonnage and grades (OK and UC) from the actual; however UC is closer to the actual grades and out performs the linear estimator OK.

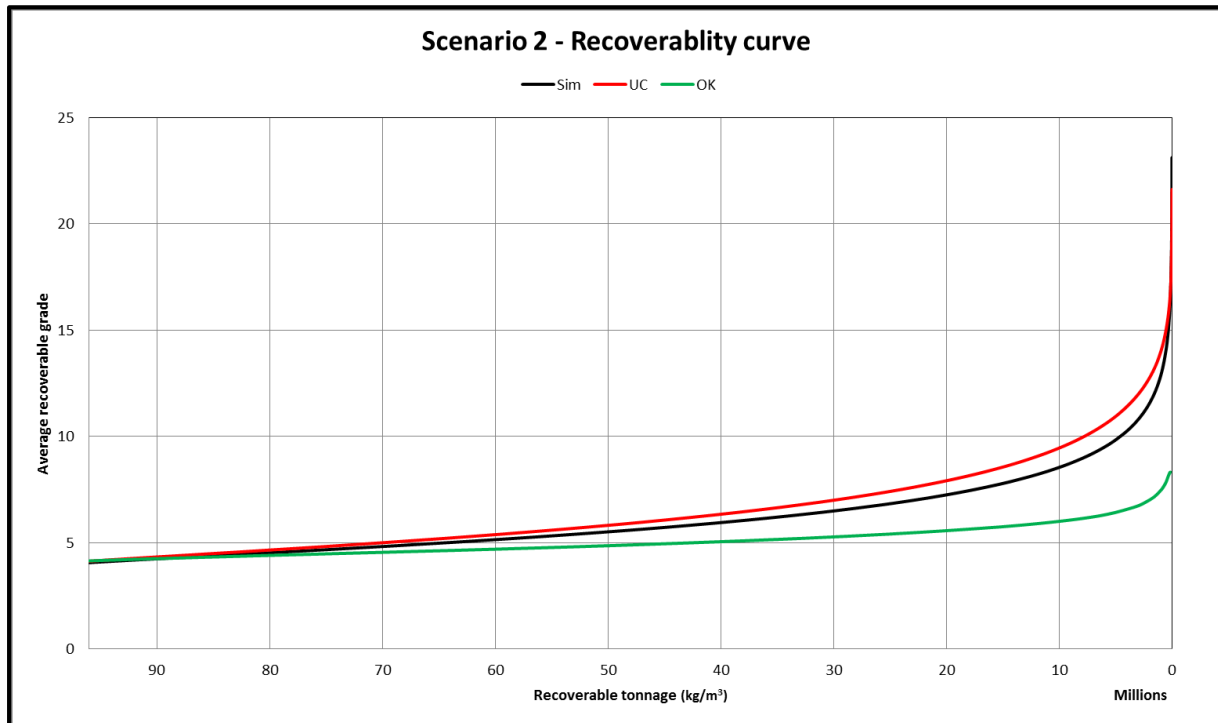


Figure 45 Scenario 2 recoverability curve

Additional plots showing the distributions of grades are given in Appendix C: Distribution of Estimated Grades, and cross plots comparing the UC versus actual and OK versus actual grades are given in Appendix D: Model Scatterplots.

At low cut-off grades, a linear estimated model frequently shows an overestimation of volume (or payable ground), which is commonly referred to as the “vanishing tonnes” problem (which is seen when mining commences and less material is recovered than was predicted). This is caused by a conditional bias in the estimate, which reflected by a higher estimation variance and a low slope regression in the estimated result. This phenomenon is amplified by high nugget effect and small block sizes used for estimation.

In order to resolve a conditional bias, it is possible to estimate grades into larger blocks. However, estimating into larger blocks is likely to produce an over-smoothed histogram, or too much “average” material and does not provide the distinction in grades required to select blocks during mine planning. This is the Kriging Oxymoron (Isaaks, 2005), which surmises that a Kriged estimate cannot be conditionally unbiased and accurate at the same time. The schematic in Figure 46 summarises this.

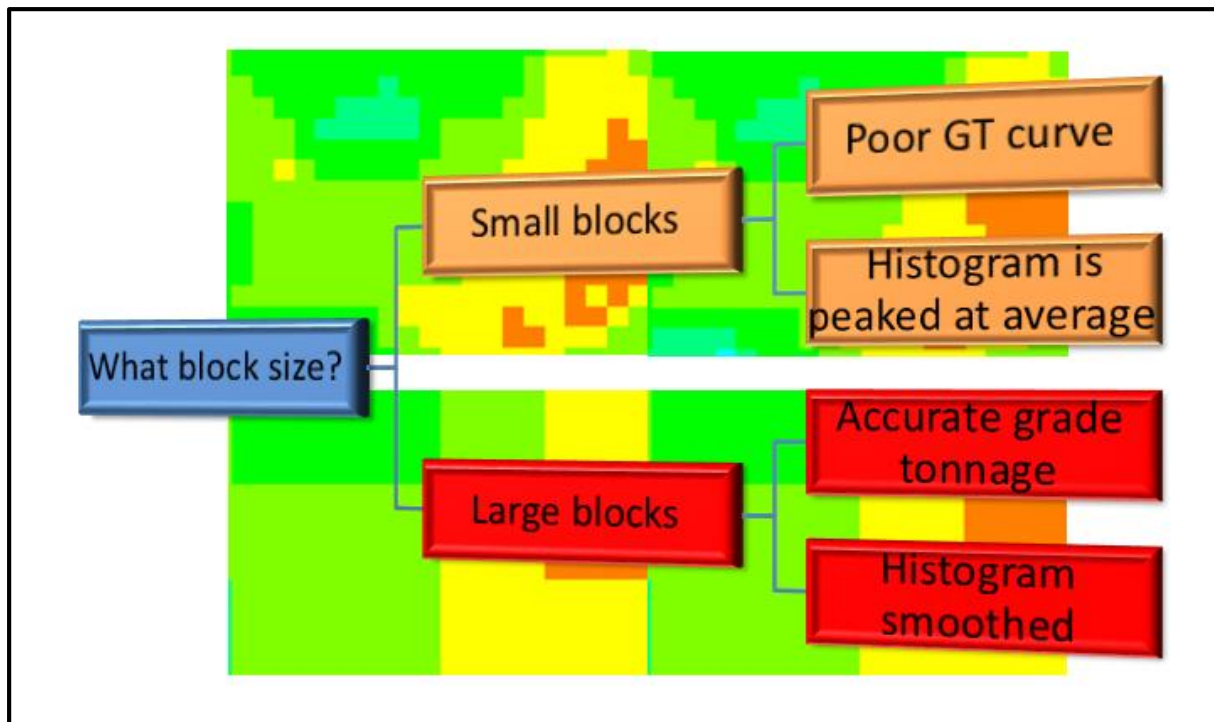


Figure 46 The Kriging Oxymoron explaining linear estimation into small blocks versus large block

UC uses the ‘conditionally unbiased’ large block estimator to condition a distribution of small blocks, thereby maintaining the reliable grade-tonnage curves and applying a conditional distribution to obtain an accurate histogram of small block (SMU) grades. This attempts to satisfy the apparent contradiction seen by the Kriging oxymoron.

Vann and Guibel (1998) cited non-linear estimation methods, like UC, to reduce GT-distortions, and this is confirmed by what is seen in both distributions.

Globally, UC performed better for the normal distribution than for the log-normal distribution. This is as a result of the low conditional biases seen in the normal distribution’s panel estimate, and the significant conditional biases seen in the log-normal distribution’s panel estimate. Conditional biases can be identified by slopes of regression. Using the slope of regression as an indicator of goodness of a panel estimate, one can predict how accurate a UC result will be in predicting the GT of a model.

4.2 Panel grade tonnage assessment

Scatter plots between the OK panel estimated grades and the simulated grades averaged per panel were created for both Scenario 1 and Scenario 2 in Figure 47. These plots highlight how closely the estimate panel grades reproduce the simulated grades. Panels representing well estimated blocks and poorly estimated blocks were selected.

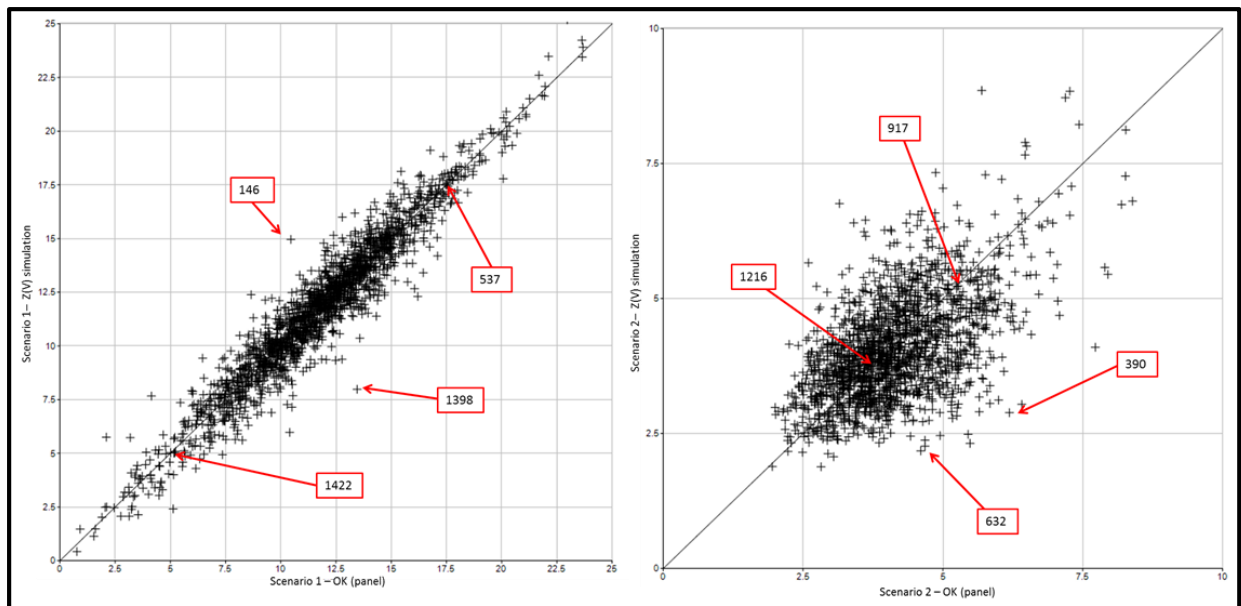


Figure 47 Actual versus OK panel scatter plot, showing blocks for LUC analysis for Scenario 1 (left) and Scenario 2 (right)

Well estimated panels occur close to the 1:1 line between actual and estimated values, while poorly estimated panels fall far from the line. The relationship of how well the panel blocks are estimated is captured in the slope of regression, with blocks where the actual and estimated values are near equal having a high regression slope; panels where these grades are different have a low regression slope. A low average slope of regression is indicative of a conditional bias, which is evident in scenario's 2 OK panel estimate.

Considering a well estimated block in Scenario 1's model (Figure 48), one can see how closely the UC estimated tonnage conforms to the simulated tonnage. In this example, the OK result is on par with the UC result.

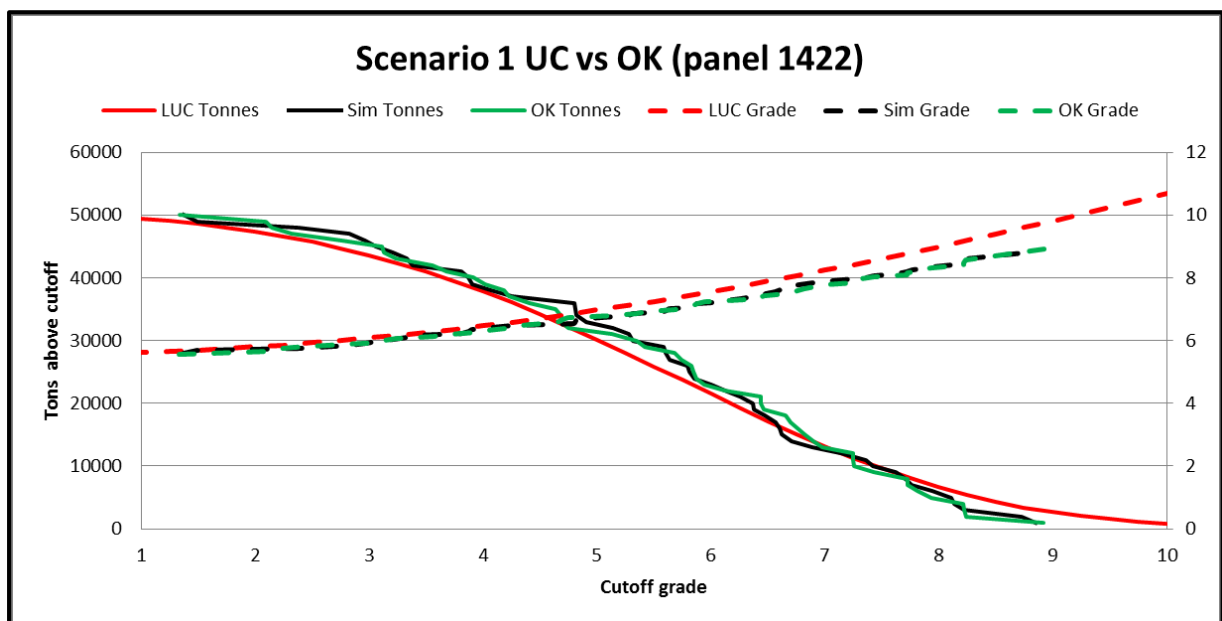


Figure 48 Grade tonnage curve for panel 1422 (Scenario 1)

In a second Scenario 1 example of a well estimated block (Figure 50), one can see how OK provided the correct mean grade, but UC was better at predicting the distribution of grades and tonnage within the panel.

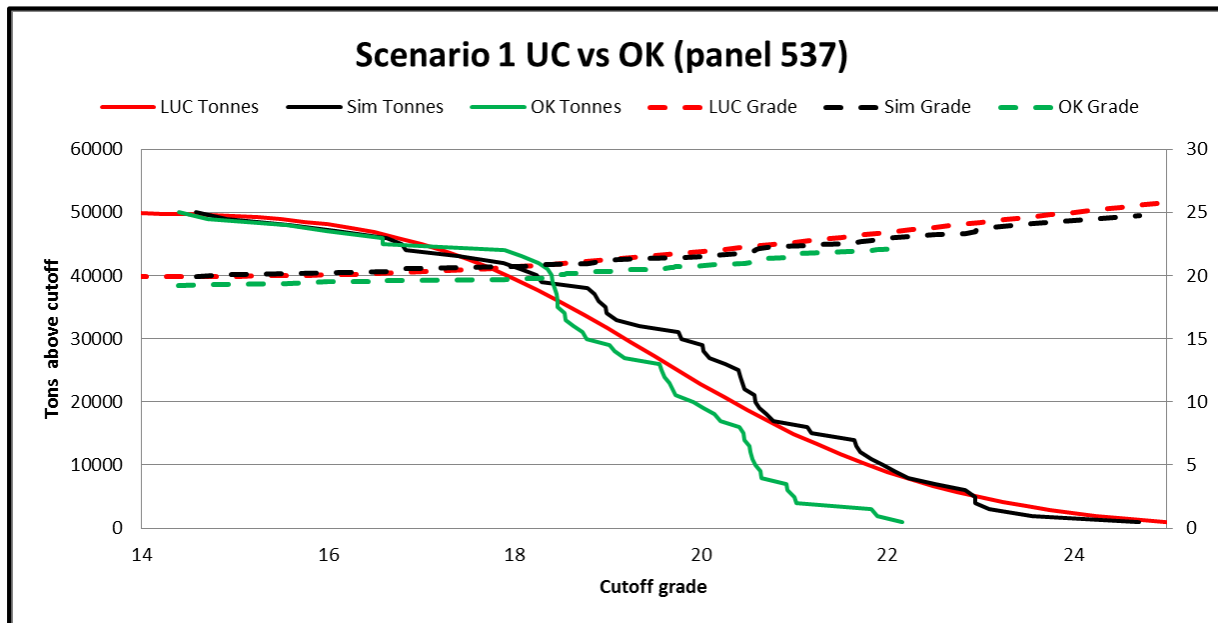


Figure 49 Grade tonnage curve for panel 537 (Scenario 1)

The panel in Figure 50 shows an example of a poorly estimated block. Since the panel grade was poorly estimated, the UC distribution of grades and tonnage are incorrect. In this example the OK estimated panel grade was higher than the actual, resulting in UC over estimating the tonnage for this panel.

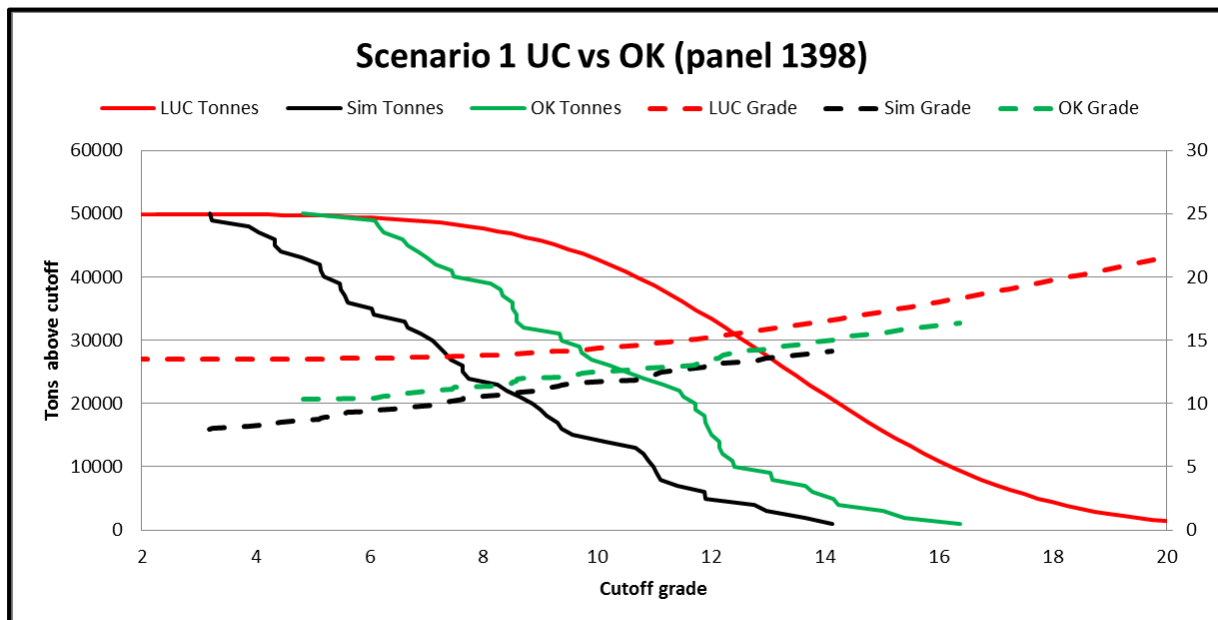


Figure 50 Grade tonnage curve for panel 1398 (Scenario 1)

Another example of poor estimation is shown in Figure 51, where the panel estimated grade was lower than the actual, resulting in an under-estimation of tonnages within the panel.

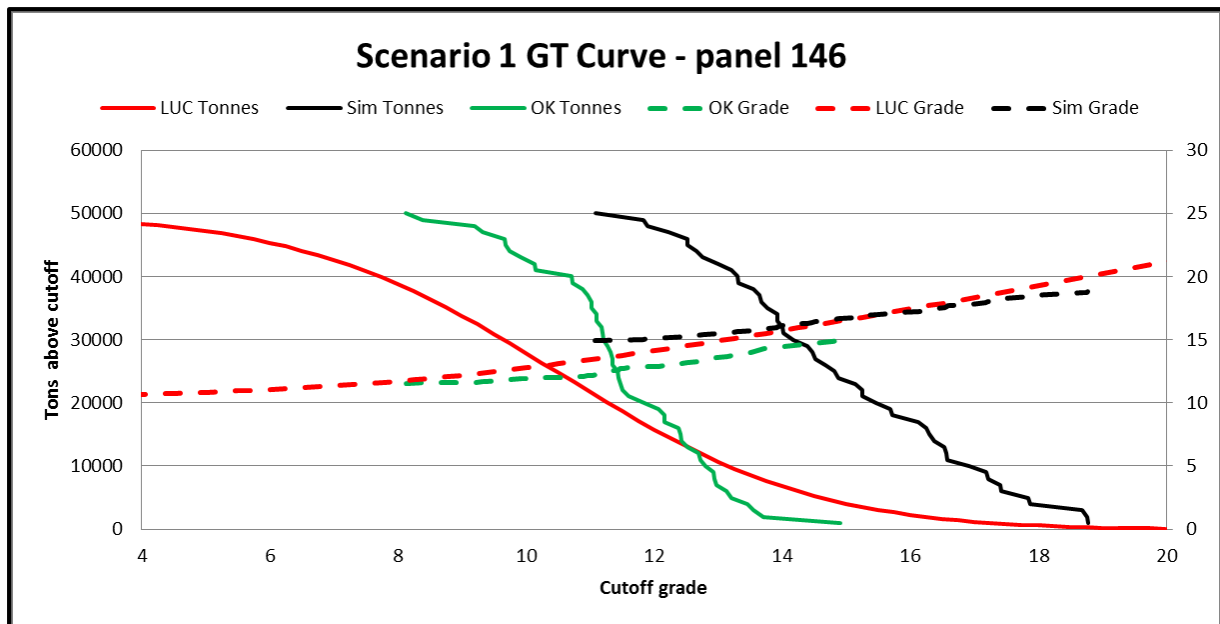


Figure 51 Grade tonnage curve for panel 146 (Scenario 1)

Taking a look at Scenario 2, where the overall panel estimate was less robust than with Scenario 1. The example shown in Figure 52 shows a good UC prediction of grades and tonnage. The UC result out-performs the OK result. The OK provided an accurate mean grade, but the result is over-smoothed and has a conditional bias which resulted in an inaccurate prediction of grade and tonnage.

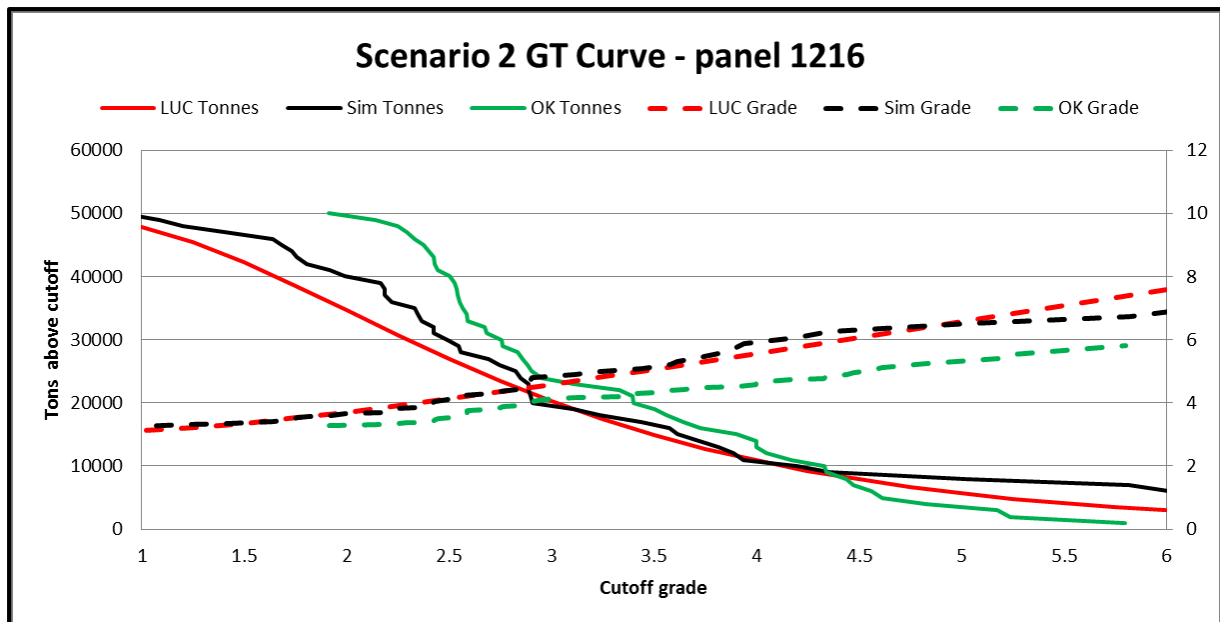


Figure 52 Grade tonnage curve for panel 1216 (Scenario 2)

Another well estimated block is shown in Figure 53. The OK panel mean grade is correct, but a conditional bias has caused an overestimation of OK predicted tons. The UC slightly under estimates the tonnage at low cut-off grades, and over estimates the grades.

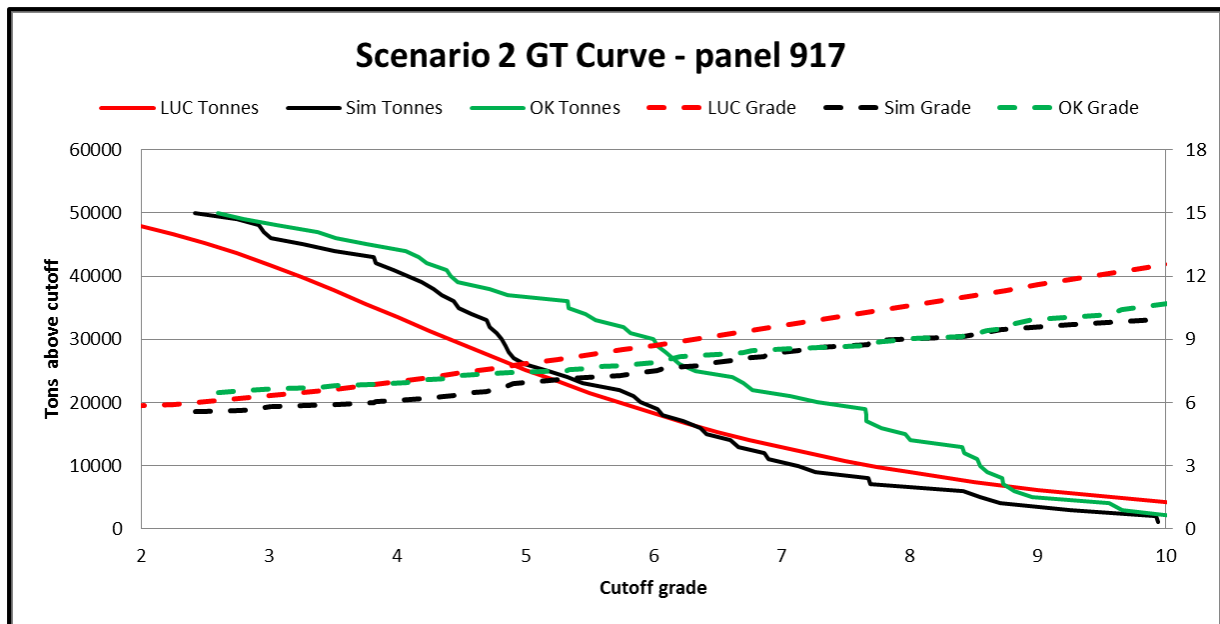


Figure 53 Grade tonnage curve for panel 917 (Scenario 2)

An example of a poorly estimated block is shown in Figure 54, where an inaccurate panel estimate has resulted in a gross overestimation of tonnages.

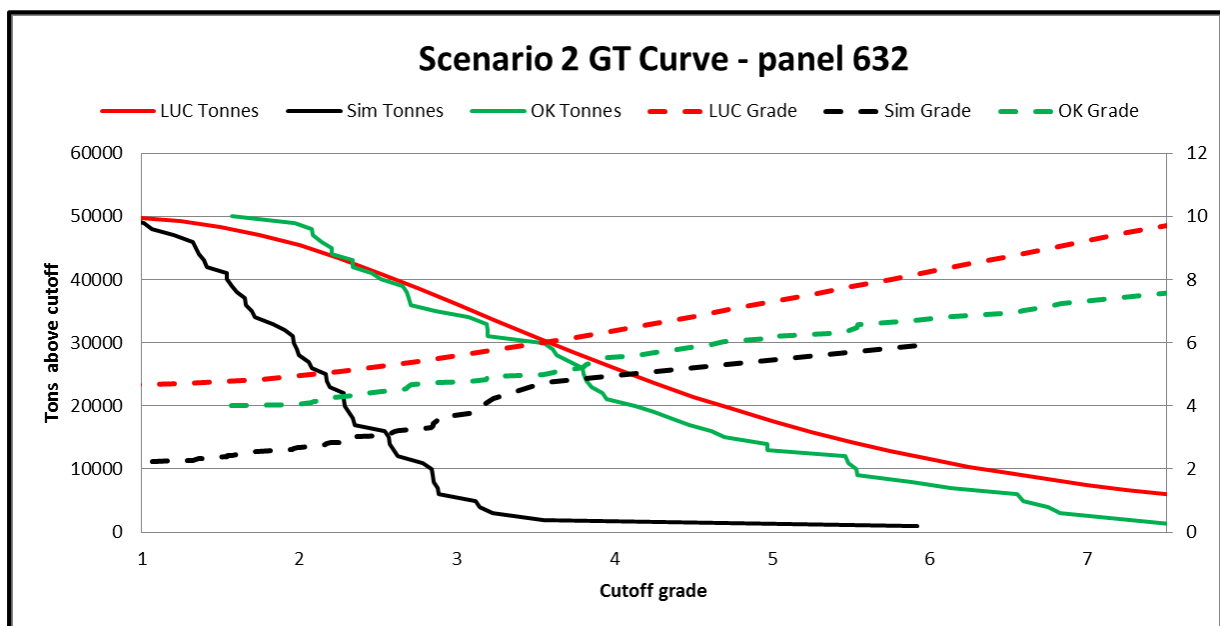


Figure 54 Grade tonnage curve for panel 632 (Scenario 2)

Additional plots showing the distribution of grades for the selected SMU are shown in Appendix E: Local Grade Distributions.

From these local GT curves it is evident that if the panel grade estimate is correct, the UC accurately predicts the distribution of SMU grades. UC applied to the normal distribution was superior at approximating the distribution than the UC applied to the log-normal distribution.

4.3 Localisation assessment

The localisation of the UC result places individual SMU grades (derived from the SMU grade tonnage curve within the panel) at specific locations within the panel, based on the estimated grades of ordinary kriged SMU. To validate how accurately the OK ranking was, the ranked values from the SMU were plotted against the ranked SMU from the simulation (representing the actuals) in Figure 55. The coloured contours represent 10 % percentile intervals, from the 90th percentile to 10th percentile contours. So the x-axis plots the true ranking of the SMU within the panel, whereas y-axis represents the estimated ranking.

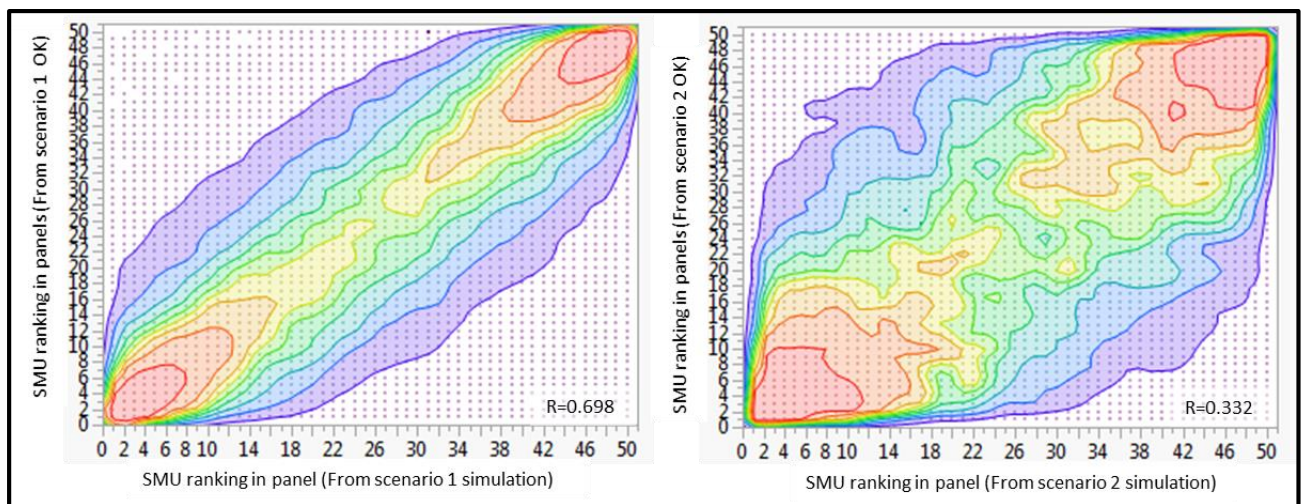


Figure 55 SMU within panel ranking comparison for Scenario 1 (left) and Scenario 2 (right)

Scenario 1's localisation placement was quite successful, with a good visual and statistical correlation between the simulated SMU ranking and the OK ranking. A perfect rank matching would result in a single straight line with a 1:1 relationship between actual and predicted rankings. Scenario 2's localisation is more variable. There is a weaker correlation between the OK ranked order and the simulated ranked order (actual), which is due to the data spacing relative to the variogram range for this distribution.

Limitation to the success of localisation depends largely on the reliability of the OK SMU estimate. However the smoothing and inaccuracy of this estimate is the prime motivation to use UC in preference to linear estimates. If the OK SMU estimate provides a good spatial representation of the local grades, then the location of the UC grades within the panel will be more accurate. This confirms Abzalov's (2006) findings that the localisation success is depending on available data amongst other factors.

If the data is closely spaced enough to provide accurate localisation, then it is also likely that the data are sufficiently closely spaced for a linear estimation to accurately predict the model grade value. In this circumstance, the benefit of using a non-linear UC estimator over a linear estimator is not as significant as the benefit seen with widely spaced data. This is evident in Scenario 1's grade tonnage predictions where the OK and UC results were very similar; and in

Scenario 2's grade tonnage predictions where there is a significant improvement in the UC over OK results.

While LUC is a useful addition to UC, it does not improve the accuracy of the UC estimate and the localisation algorithm cannot predict the placements of SMU beyond the available data. This is the main challenge: one cannot simultaneously know the local mean and the local variability from limited local data. The single largest contribution of the localisation approach is to present a recoverable resource estimate in a more accessible and immediately useful format for mine planning.

5. Conclusion

UC is a non-linear estimation method which estimates a grade and tonnage distribution within a large mining panel above a series of cut-off grades. LUC is an add on to UC, that presents a UC estimate in SMU blocks where the SMU grades within a panel honour the UC grade distribution and are arranged to reflect the local grade pattern.

This project sought to test the effectiveness of UC and LUC, by applying the estimation method in two scenarios based on synthetically generated data. The data was generated from a single realisation of sequential Gaussian simulation, and the first scenario was designed to be normally distributed with good grade continuity, as one might see in a porphyry copper mineral deposit, and the second scenario was designed as log-normally distributed data with a high nugget and short range continuity, as one might see in a gold or iron mineral deposit. The performance of the approaches was measured by applying UC and LUC to a sampled subset of the realisations and compared to the un-sampled realisation which represents the actual. Both scenarios were also estimated with OK, as a benchmark for linear estimation.

UC is applicable for estimating data that has a bi-Gaussian distribution, which means that any linear combination of the normal score transforms of the data is Gaussian. Both scenario 1 and scenario 2 confirmed this behaviour, which was expected as a Gaussian simulation method was used for the generation of data in both scenarios. In this study, the mostly significant difference between the two data sets was the amount of data relative to the ranges of the variogram. The normally distributed data has on average three data points within the range of variogram; whereas the log-normal data has on average one data point within the variogram range.

A GT curve is an optimistic presentation of the extractable grades and tonnages of material for a set of cut-off grades. By demonstrating how closely the estimated GT curve predicts the actual grade tonnage, one can determine the success of the estimator.

From this study it is seen that UC performs well in terms of global GT estimation when there is an underlying normal distribution and there is sufficient data falling within the range of the variogram model, which results in low conditional biases. In such conditions, the linear estimator OK also produces an accurate GT assessment. Thus when there is sufficient data coverage within the variogram ranges and a low conditional bias, there is only a slight benefit offered by UC for a global GT predictions, as the estimated results from the linear estimator OK are also reasonable.

UC performed better than OK when predicting the grades and tonnage of log-normally distributed data with poor data coverage in the ranges of the variogram. In these circumstances, OK performed poorly due to conditional bias which may be amplified by a high

nugget effect and/or small blocks. It is concluded that when there is poor data coverage within the ranges of the variogram, UC is better at predicting the grades and tonnage of material above cut-off than the linear estimator OK.

The estimated distribution of grades within a panel that is predicted by UC is close to the actual distribution if the estimated mean panel grade is correct. Conversely, if the estimated mean panel grade is incorrect, then the distribution of grades predicted by UC is wrong. The slope of regression is a good indicator of how close the estimated panel grade is to the actual. Even if the slopes of regression of all the estimated panels are, on average, close to 1, the slope of regression of an individual panel should be considered as a confidence indicator for the UC predicted distribution of grades within that panel.

The performance of LUC is dependent on the reliability of the OK SMU estimate that is used to spatially allocate SMU grades within a panel. If there is sufficient data within the ranges of the variogram, the SMU estimate, used for localisation, will reliably estimate the local grade patterns. Therefore the localisation process does a good job of placing the SMU grades, which are derived from the UC grade distribution, within a panel. Conversely if there is insufficient data within the variogram ranges, the localisation process will do a poor job of predicting the spatial arrangements of SMU grades within a panel. While LUC is a useful addition to UC, it does not improve the accuracy of the UC estimate and the localisation algorithm cannot predict the placements of SMU beyond the available data.

REFERENCES

- Abzalov, M.Z., 2006. Localised uniform conditioning (LUC): A new approach for direct block modelling, *Mathematical Geology*, Volume 38, No 4, 393-411.
- Armstrong, M and Matheron, G, 1986. Disjunctive kriging revisited, Part 1, *Mathematical Geology*, Vol 18(8), 711-728.
- Assibey-Bonsu, W. and Krige, D.G., 1999. Use of direct and indirect distributions of selective mining units for estimation of recoverable resource/reserves for new mining projects, AOPCOM Symposium 99, Colorado School of Mines.
- Assibey-Bonsu W., 1998. Use of uniform conditioning technique for recoverable resource/reserve estimation in a gold deposit, Proc Geocongress '98, Geol. Soc. S. Afr., South-Africa, 68-74.
- David M., Dagbert, M. and Belisle J.M., 1977. The practice of porphyry copper deposit estimation for grade and ore-waste tonnages demonstrated by several case studies, 15th APCOM symposium, Brisbane, Australia.
- Deraisme, J. and Assibey-Bonsu, W., 2011. Localised uniform conditioning in the multivariate case - An application to a porphyry copper gold deposit, 35th APCOM Symposium, Wollongong, Australia.
- Deraisme, J., Rivoirard, J. and Carrasco Castelli, P., 2008. Multivariate uniform conditioning and block simulations with discrete gaussian model: Application to the Chuquicamata deposit, Geostats 2008, Santiago Chile.
- Deraisme, J. and Roth, C., 2000. The information effect and estimating recoverable reserves, Geovariance white paper.
- De-Vitry, C., Vann, J. and Arvidson, H., 2007. A guide to selecting the optimal method of resource estimation for multivariate iron deposits, Proceedings of iron ore conference, Perth, Australia, 67-77.
- Hardtke, W., Allen L. and Douglas I., 2011. Localised indicator Kriging, 35th APCOM Symposium, Wollongong, Australia.
- Harley, M., and Assibey-Bonsu, W., 2007. Localised uniform conditioning: How good are the local estimates?, 33rd APMCOM Symposium, Chile.
- Humphreys, M, 1998. Local recoverable estimation: A case study in uniform conditioning on the Wandoo Project for Boddington Gold Mine. Proceedings of a one day symposium: Beyond Ordinary Kriging, Australia, 63-75.

Isaaks, E., 2005. The Kriging Oxymoron: A conditionally unbiased and accurate predictor (2nd Edition), Quantitative Geology and Geostatistics, Geostatistics 2004, Banff, volume 14, 363-374.

Matheron, G, 1974. Les fonctions de transfert des petits panneaux. Technical Report N-395, Centre de Geostatistique, France.

Neufeld, C and Deutsch, CV, 2005. Calculating Recoverable Reserves with Uniform conditioning, Proceedings of IAMG '05: GIS and Spatial Analysis, vol 2, 1065-1070.

Neufeld, C.T, 2005. Guide to Recoverable Reserves with Uniform Conditioning, Guidebook Series Volume 4, Centre for Computational Geostatistics (CCG), Canada

Remacre, AZ, 1984. L'estimation du R'ecup'erable local, le conditionnement uniforme. PhD thesis, Ecole Nationale Sup'erieure des Mines de Paris, 1984.

Rivoirard J, 1994. Introduction to disjunctive kriging and nonlinear geostatistics, Centre de Geostatistique, Ecole des mines, France

SAMREC code, 2009. The South-African minerals code for reporting of exploration results, mineral resources and mineral reserves, 2009 edition.

Schofield, N., 1988. Ore reserve estimation at the enterprise Gold mine, Pine Creek, Northern Territory, Australia. Part 1: structural and variogram analysis, CIM Bulletin volume 81, No 909, 56-66.

Vann, J., Jackson S. and Bertoli, O., 2003, Quantitative Kriging Neighbourhood Analysis for the Mining Geologist – A Description of the Method with Worked Case Examples, 5th international Mining Geology Conference, Australia.

Vann, J and Guibal, D, 1998. Beyond ordinary kriging – An overview of non-linear estimation, Mineral Resource and Ore Reserve Estimation - The AusIMM Guide to Good Practice, Australia, 6-23.

Vann, J., Guibal, D. and Harley, M, 2000. Multiple indicator kriging: is it suited to my deposit?, 4th International Mining Geology Conference, Australia, 9-17.

Verly, G. and Ramani, R. V. (chairperson), 1986. Multigaussian kriging - a complete case study, Application of Computers and Operations Research in the Mineral Industry, United States, volume 19, 283-298.

Appendix A: Statistics Tables

Statistics are presented, in Table 8 and Table 9, at three levels of support (point, SMU and panel), for simulation (actuals), LUC and OK estimation. The following are observed:

- Mean value stays approximately constant for different supports
- Measures of variance (including variance, standard deviation and range) decrease for larger supports.
- For Scenario 1 (a normal population), skewness stays approximately constant as support increases
- For Scenario 2 (log-normal population), skewness tends towards 0 for increasing levels of support
- Coefficient of variation decreases for larger supports

Table 8 Statistics for Scenario 1 at multiple supports

Model	Support	Number of Samples	Min	Max	Mean	Variance	Standard Deviation	Skewness	Kurtosis	Coefficient of Variation
Actual (sim)	point	12 000 000	0.0	30.0	11.7	19.5	4.4	0.00	0.25	0.38
Samples	point	41 700	0.0	29.6	11.8	21.5	4.6	0.05	0.30	0.39
LUC	SMU	96 000	-3.7	28.4	11.7	16.3	4.0	0.01	0.33	0.34
OK	SMU	96 000	-0.1	27.1	11.7	13.3	3.6	0.00	0.51	0.31
Actual (sim)	SMU	96 000	0.0	27.8	11.7	15.5	3.9	-0.02	0.37	0.34
OK	Panel	1 920	0.8	23.7	11.7	11.2	3.3	-0.01	0.53	0.28
Actual (sim)	Panel	1 920	0.4	25.2	11.7	11.8	3.4	-0.01	0.48	0.29

Table 9 Statistics for Scenario 2 at multiple supports

Model	Support	Number of Samples	Min	Max	Mean	Variance	Standard Deviation	Skewness	Kurtosis	Coefficient of Variation
Sim (point)	point	12 000 000	0.0	60.0	4.1	29.5	5.4	3.75	20.41	1.34
Samples	point	41 700	0.0	59.3	4.1	30.3	5.5	3.68	19.50	1.34
LUC	SMU	96 000	0.4	21.6	4.1	5.9	2.4	1.45	2.87	0.59
OK	SMU	96 000	0.3	19.8	4.1	3.0	1.7	1.22	2.81	0.42
Actual (sim)	SMU	96 000	0.2	23.2	4.1	4.4	2.1	1.50	3.88	0.51
OK	Panel	1 920	2.0	8.4	4.2	0.9	1.0	0.65	0.77	0.23
Actual (sim)	Panel	1 920	1.9	8.9	4.1	0.9	0.9	0.83	1.47	0.23

Appendix B: Slope of Regression

The slope of regressions was plotted for the panel estimates, for Scenario 1 and Scenario 2. The high slopes of regression for Scenario 1 are reflected in the closeness of the OK GT curve to the actual, while the moderate slopes of regression for Scenario 2's are reflected in the poor GT curve adherence of the estimate to the reality. Slopes of regression are an indication of conditional bias.

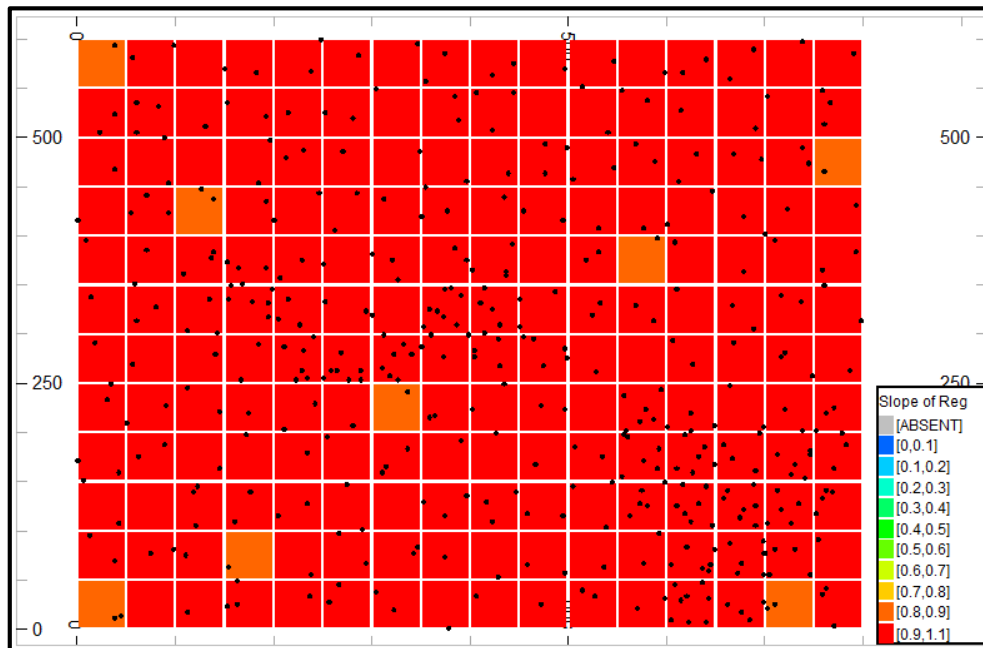


Figure 56 Slopes of regression for Scenario 1's OK model

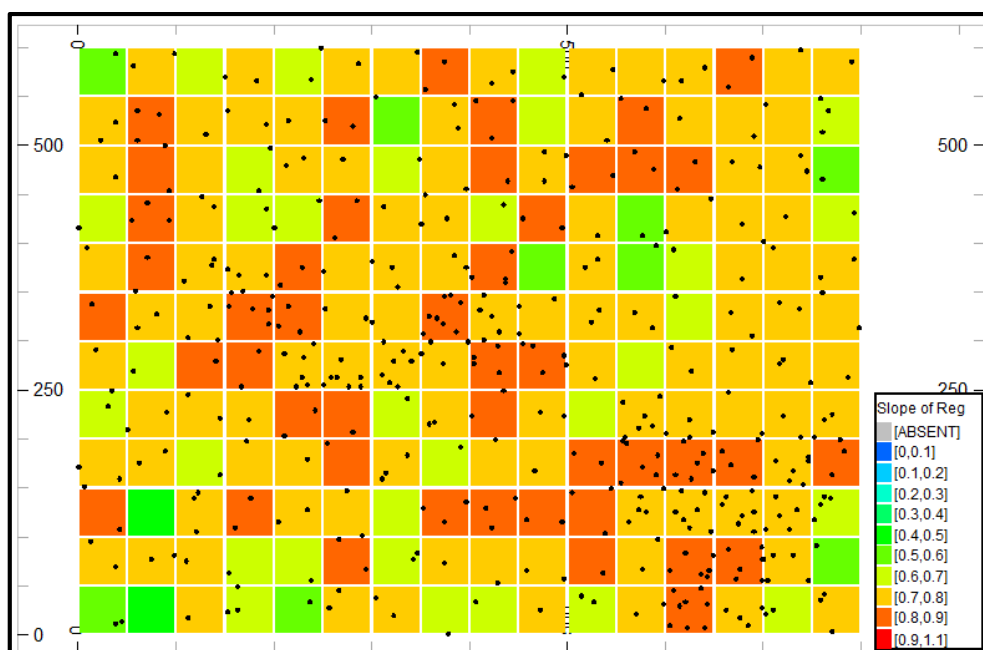


Figure 57 Slopes of regression for Scenario 2's OK model

Appendix C: Distribution of Estimated Grades

Distribution frequencies of the actual (simulated), ordinary kriged (OK) and uniform conditioned (UC) models were plotted in Figure 58 and Figure 59 for comparison. The smoothing effect of the OK is highlighted, with a higher 'peak' of grade values. For the normally distributed grades, UC closely predicts the grade distribution. For the log-normal distribution, UC slightly underestimated the grades around the mean.

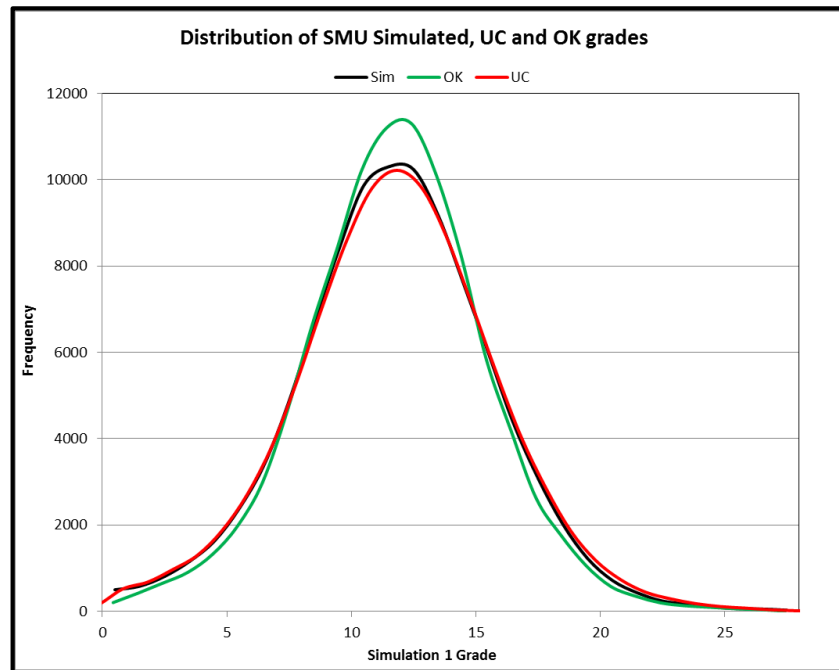


Figure 58 Global distribution of grades for Scenario 1

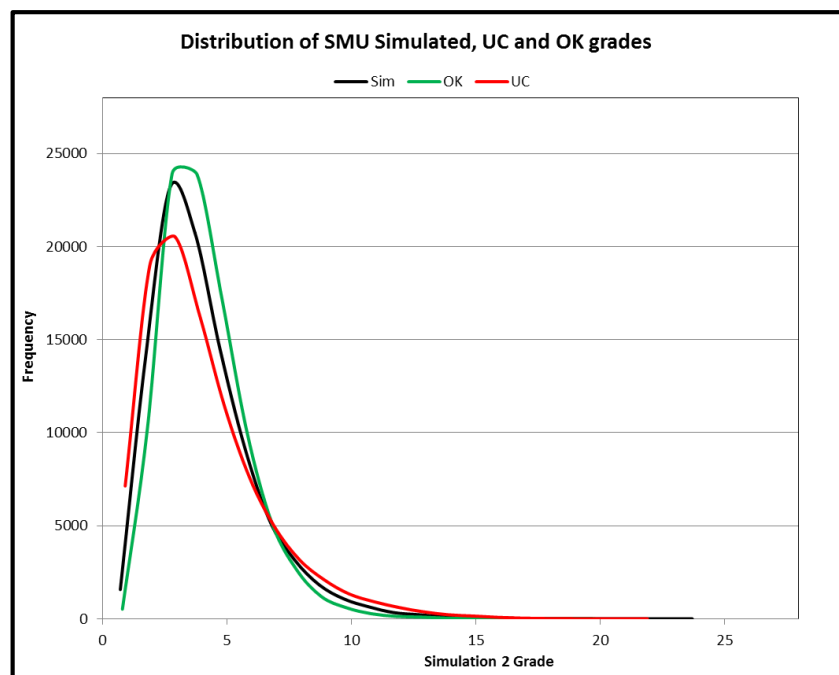


Figure 59 Global distribution of grades for Scenario 2

Appendix D: Model Scatterplots

Scatter plots showing the localised uniform conditioned grades versus simulation (at SMU scale (left), and OK panel estimated grade versus actual simulated grades (right). The scatter plot on the right is reflects the conditional bias for the OK panel estimate as shown in Figure 60 and Figure 61.

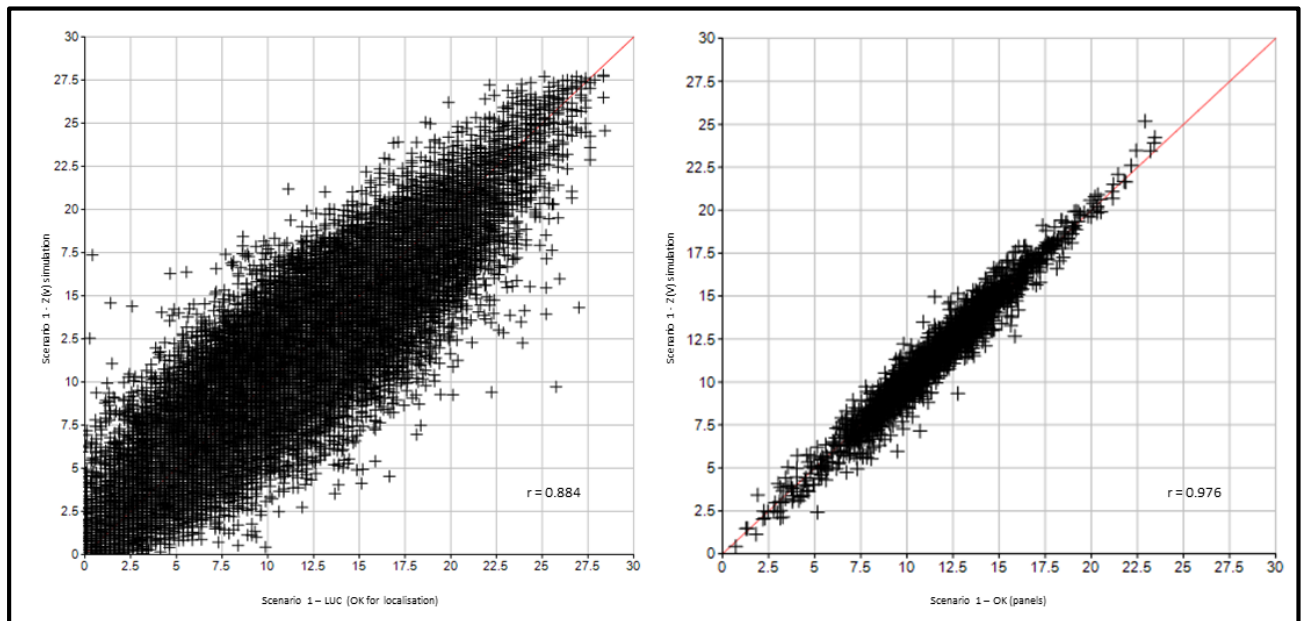


Figure 60 Scatter plot for actual versus LUC (left); and actual versus OK panel (right), for Scenario 1

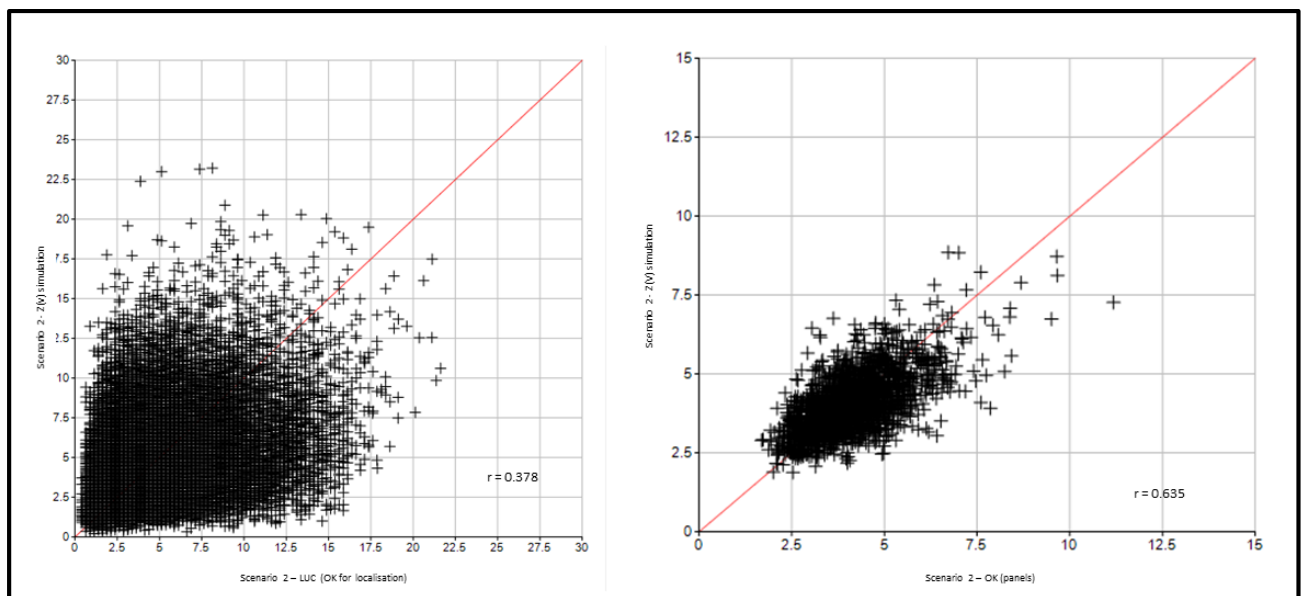


Figure 61 Scatter plot for actual versus LUC (left); and actual versus OK panel (right), for Scenario 2

Appendix E: Local Grade Distributions

Plots of the grade distribution, showing distribution of simulated, LUC and OK (in SMU blocks) are shown in Figure 62 to Figure 65. The location of these panels is given in Figure 47.

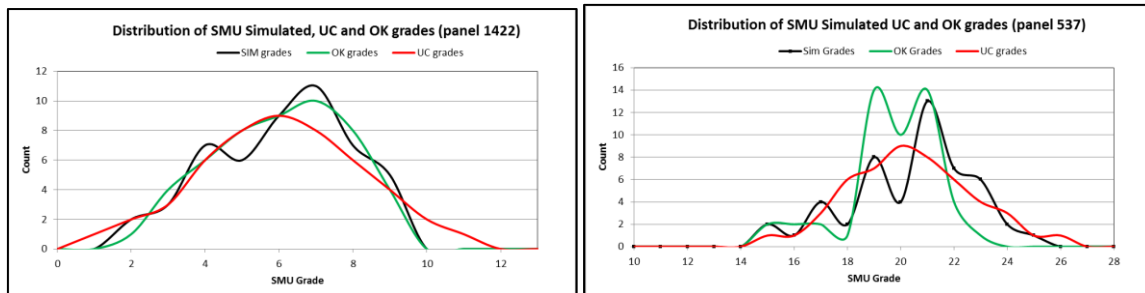


Figure 62 Normal distribution (Scenario 1) – well estimated panels

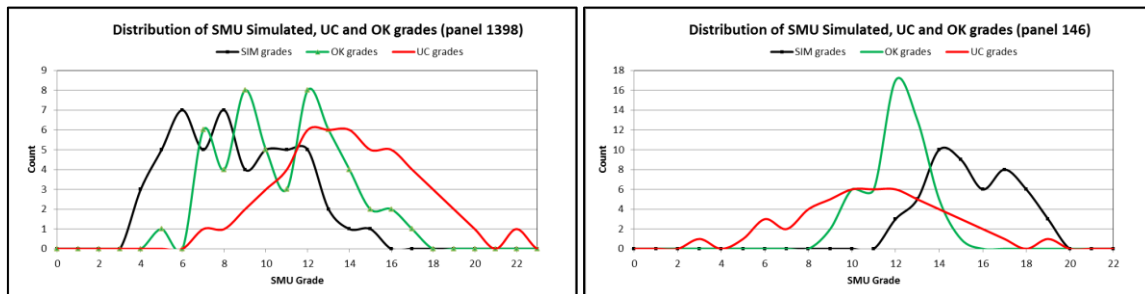


Figure 63 Normal distribution (Scenario 1) – poorly estimated panels

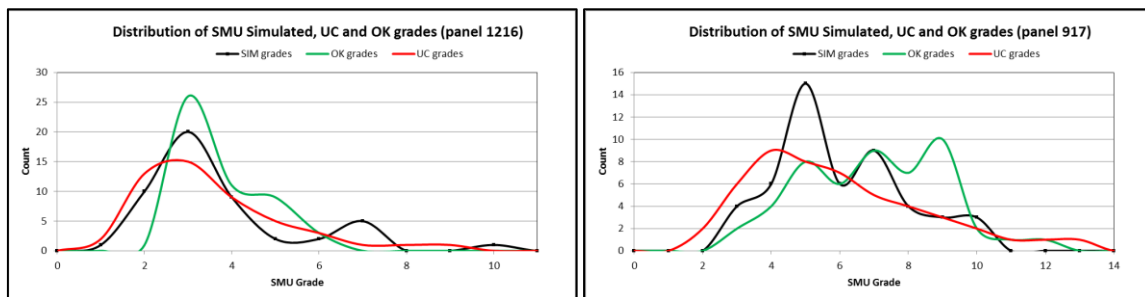


Figure 64 Log-normal distribution (Scenario 2) – well estimated panels

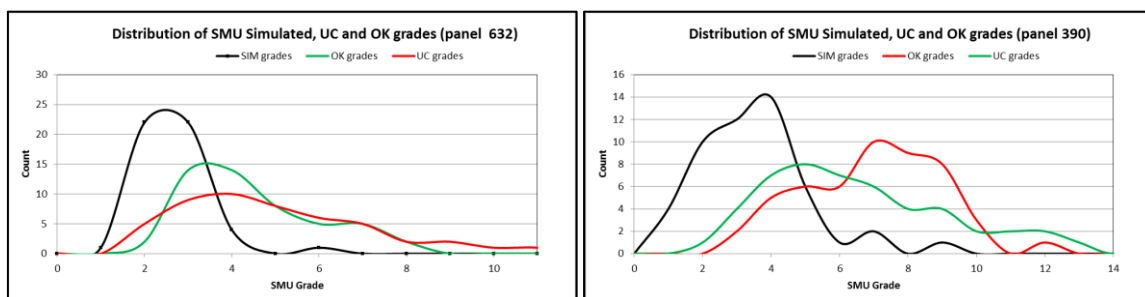


Figure 65 Log-normal distribution (Scenario 2) – poorly estimated panels

