



UNIVERSITY OF THE
WITWATERSRAND,
JOHANNESBURG

**COMPARISON OF KRIGING NEIGHBOURHOOD ANALYSIS (KNA)
RESULTS FROM TWO DIFFERENT SOFTWARE PACKAGES – A CASE
STUDY**

Martha Munageni Nemupewe Amwaama

670266

School of Mining Engineering
University of the Witwatersrand
Johannesburg, South Africa

SUPERVISOR: Prof C.E Dohm

A research report 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 Mining Engineering.

5 October 2018

Declaration

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

Martha M.N Amwaama

5 Day of October 2018

Abstract

The quality and outcome of any kriging estimation and conditional simulation exercises are dependent on the definition of the kriging neighbourhood parameters that are applied in these processes. It is necessary to minimise the conditional bias that arises from the application of the kriging estimation methodology. This minimisation can be addressed by a kriging neighbourhood analysis (KNA), whereby the optimum kriging neighbourhood parameters are identified prior to executing the actual block kriging estimation.

Conditional bias presents itself in the reality that the real mining blocks grade averages are more variable than the estimated block grades, which is a consequence of the smoothing effect of kriging. The motivation for this research into the KNA outcomes using different software stems from recognising the importance of minimising conditional bias in the application of kriging in resource estimation. Whilst studies on the comparison of software packages have been done in the fields of computer science, bio-technology, geology, and other areas, no evidence was found in the literature review for a comparison of KNA results being carried out using different software packages. This research was embarked on to establish whether the KNA procedures proposed by different software providers would result in the same kriging neighbourhood parameters being selected for the kriging estimation process.

To do the research from a practical aspect, a case study on Rossing Uranium Mine was considered. The results from the KNA using SUPERVISOR geostatistical software and those from using SURPAC geology and mine planning software is compared in this research. The conditional bias measures that need to be optimised are: the slope of regression; kriging efficiency, and the number of negative weights. This is done by analysing the impact of varying the following parameters in the kriging neighbourhood, namely the estimation block size, the minimum and maximum number of samples to be considered, the search range and the configuration of the discretisation points of the block to be estimated; which all have an influence on the afore mentioned conditional bias measures.

The same input data, geological model and block configuration, semi-variogram parameters as well as test locations were used to ensure that the outcomes of the application of the software packages would be comparable, valid and not user introduced. The two block configurations tested are for well-informed blocks and for poorly- informed blocks.

This research study concluded that there is no significant difference between the KNA results produced by SUPERVISOR and SURPAC; the two software packages considered, there are however other differences between the software packages which are not related to KNA. For well-informed blocks the optimised conditional bias measures identified using SUPERVISOR and SURPAC are the same with slight differences in the conditional bias measures for poorly-informed blocks. Differences identified are related to the manner in which the individual Software packages select specific samples for inclusion in the kriging neighbourhood.

*Dedicated to my parents, Mrs. Evelina Amwaama and Mr. Matheus Amwaama, and
my Godchildren, Samuel Tanga-Omuwa Kambidhi and Martha Imorgen Amwaama.*

Acknowledgements

Firstly, I would like to thank my Maker, it is because of His mercy and unconditional love that I was able to enrol for this degree and do it to completion. Rio Tinto's Rossing Uranium Mine is thanked for granting me permission to use their data after months of failing to secure data from the mining industry. Special thanks go to Mrs. Angela Kapapilo who discussed my case with her manager and facilitated the entire process of getting permission for me to use their data. Mrs. Kapapilo is also thanked for her help in validating the data as well as for her timely responses to my emails. I will forever be grateful. The Namibian Student Financial Aid Fund (NSFAF) is thanked for funding my studies.

Prof Christina Dohm, who made the resource evaluation courses so interesting, thereby triggering my interest in the subject of geostatistics. Thank you so much, Prof, for introducing the world of geostatistics to me. Your invaluable input and guidance are highly appreciated. Thank you to Ms. Lisa Chanderman who specifically provided valuable guidance and supervision of the research proposal.

I am also indebted to Wits School of Mining Engineering for providing facilities; among them the SURPAC license and most importantly for buying a Snowden SUPERVISOR license for me. I would also like to acknowledge the conducive environment that the School has provided for its students. There were times when the computer in the Mine Design Laboratory (MDL) was switched off after lectures or due to power failure and I could not access SURPAC. A big thank you goes to Ms. Paskalia Neingo, who always responded to my requests promptly. Prof Bekir Genc is equally thanked for solving the issue of not being able to access SURPAC whenever the MDL computer was switched off.

My parents, siblings, family members and friends are thanked for their continuous support and for their daily prayers. I would like to particularly thank my elder sister, Selma Amwaama, for always being my shoulder to cry on. When things got hectic both socially and academically, she was keen to listen to my whining and advise me accordingly. My little sister Ester Nangolo, your caring and motivation will always be cherished. To my uncle Dr. Silas Hango, it is all because of your inspiration that I decided to pursue this Masters' research. Mrs. Selma Shaanika-Kankondi, thank you so much for your financial support and for being a role model to me, I always draw my inspiration from you.

To Bright, Moshood, Mpeo, Monica, Josee, Tapiwa and all the full-time Mining Engineering Post-graduate students, we were a close-knit group, and this made my stay at Wits fun and memorable. To my 'Pedi friend', Mosima, your check-up calls kept me going, thanks, buddy!

To Ndamona, Josia, Rosalia, Frans and all my fellow Namibians studying at Wits, thank you so much for your constant encouragements.

Table of Contents

Declaration.....	i
Abstract.....	ii
Acknowledgements.....	iv
List of Figures.....	viii
List of Tables	xii
List of Symbols	xiv
List of Acronyms	xv
1 Introduction	1
2 Research Background and Literature Review.....	4
2.1 Review of Approaches for the comparison of software packages in non-mining related fields.....	6
2.2 Review of Approaches for the comparison of software packages in mining related fields.	9
3 Kriging Neighbourhood Analysis (KNA) – Literature Review	11
4 Rossing Uranium Deposit	19
4.1 Local geology.....	21
4.2 Uranium Mineralisation.....	22
5 Statistical and Geostatistical analysis of the Rossing data set.....	24
5.1 Introduction.....	24
5.2 Data validation.....	24
5.2.1 Drill hole Collar Data Validation.....	26
5.2.2 Drill hole Survey Data Validation.....	27
5.2.3 Drill hole Assay Data Validation	27
5.2.4 Geological Data Validation.....	27
5.3 Domaining.....	28
5.4 Compositing.....	33
5.5 Bimodality and Outliers.....	41
5.6 Declustering	51
5.7 Variography	56

6	KNA procedures in SUPERVISOR and SURPAC.....	67
6.1	Measures for minimising conditional bias.....	67
6.2	Identification of well-informed and poorly-informed areas in Rossing Case study area.....	70
6.2.1	SUPERVISOR: Well-informed blocks.....	73
6.2.1.1	SUPERVISOR: Well-informed blocks - Optimum block size	73
6.2.1.2	SUPERVISOR: Well informed blocks - Minimum and Maximum number of samples	76
6.2.1.3	SUPERVISOR: Well-informed blocks – Optimum Search Range.....	77
6.2.1.4	SUPERVISOR: Well-informed blocks – Optimum Block Discretisation	78
6.2.2	SUPERVISOR: Poorly-informed blocks	80
6.2.2.1	SUPERVISOR: Poorly-informed blocks – Optimum Block Size	80
6.2.2.2	SUPERVISOR: Poorly-informed blocks - Minimum and Maximum number of samples	82
6.2.2.3	SUPERVISOR: Poorly informed blocks – Optimum search range	83
6.2.2.4	SUPERVISOR: Poorly-informed blocks – Optimum block discretisation	84
6.3	Kriging Neighbourhood Analysis (KNA) in SURPAC.....	86
6.3.1	SURPAC: Well-informed blocks.....	88
6.3.1.1	SURPAC: Well-informed blocks – Optimum Block Size	88
6.3.1.2	SURPAC: Well-informed blocks - Minimum and maximum number of samples	92
6.3.1.3	SURPAC: Well-informed blocks – Optimum Search Range.....	94
6.3.1.4	SURPAC: Well-informed blocks – Optimum Block Discretisation.....	95
6.3.2	SURPAC: Poorly-informed blocks	95
6.3.2.1	SURPAC: Poorly-informed blocks – Optimum block size.....	95
6.3.2.2	SURPAC: Poorly-informed blocks - Minimum and maximum number of samples	98
6.3.2.3	SURPAC: Poorly-informed blocks – Optimum search range.....	99
6.3.2.4	SURPAC: Poorly-informed blocks – Optimum block discretisation....	100
6.4	Differences between the KNA processes in SUPERVISOR and SURPAC	100

7	Results and Discussions.....	101
7.1	Introduction.....	101
7.2	KN parameters	101
7.3	KN statistics	102
7.4	Block model	104
7.5	Grade estimation	104
7.5.1	OK estimates	104
7.5.2	Validation.....	105
8	Conclusions.....	108
9	References.....	111
	Appendix A: Grade Capping	117
	Appendix B: - U₃O₈ (ppm) histograms, log histograms, and cumulative log- probability plots of 1.5m composites	121
	Appendix C: KNA Results	135

List of Figures

Figure 3.1: Linear regression as a result of the information effect	12
Figure 3.2: Search ellipsoid number versus mean variance, mean slope and proportion of negative kriging weights.	17
Figure 4.1: Location of Rossing Uranium Mine	20
Figure 4.2: Location of other uranium mines and projects near Rossing Uranium Mine with their owner companies highlighted in red.	20
Figure 4.3: Stratigraphic column showing rocks exposed in the Rossing area.....	21
Figure 5.1: Rossing drill holes collar positions.....	25
Figure 5.2: A plan view of the drill hole positions, highlighting the deviation of one drill hole (red).....	26
Figure 5.3: A section view facing North showing mineralization wireframes for Husab (Rossing South) deposit.	29
Figure 5.4: A plan view showing 19 three-dimensional mineralisation domains (distinguished by different colours) for Husab (Rossing South) deposit. 30	
Figure 5.5: SURPAC – A section view showing ore wireframes sections based on a 100ppm uranium cut-off grade and alaskite contacts.....	31
Figure 5.6: SURPAC - 82280 m E section showing ore wireframes sections based on a 100ppm uranium cut-off grade and alaskite contacts.....	32
Figure 5.7: SURPAC - Long section view showing Rossing’s seven mineralisation domains identified in this research.....	32
Figure 5.8: Histogram of sample lengths	34
Figure 5.9: Histogram of sample lengths, zooming in on the 0.0m to 5.5m classes... 35	
Figure 5.10: Rossing Uranium: Scatter plot of sample grades vs sample lengths prior to compositing.	37
Figure 5.11: Histograms of the U_3O_8 (ppm) composites Domain 1 (left) and Domain 2 (right) showing unimodal distributions.	43

Figure 5.12: Histograms of the U_3O_8 (ppm) composites - Domain 3 (left) has a unimodal distribution and Domain 4 (right) displays a bimodal distribution.	44
Figure 5.13: Histograms of the composites U_3O_8 (ppm) for Domain 5 (left) has a bimodal distribution (left) and Domain 6 (right) a unimodal distribution.	45
Figure 5.14: Domain 7: Histogram of the U_3O_8 (ppm) composites depicting bimodality.....	46
Figure 5.15: Domain 3 U_3O_8 (ppm) histogram, and cumulative frequency plot of 1.5m composites before capping (left) with an outlier (circled) and after capping (right).	48
Figure 5.16: U_3O_8 (ppm) histogram (left) and cumulative log-probability plot (right) of 1.5m composites for all domains combined after capping	50
Figure 5.17: Samples and their respective polygons of influence	52
Figure 5.18: Section view facing North showing drill holes displayed by U_3O_8 grades with a high-grade zone circled. Colour coding: 0 - 100ppm = cyan; 100 - 300ppm = blue; 300 - 700ppm = magenta; 700 - 10 000ppm = red	54
Figure 5.19: Declustered mean versus cell volume highlighting the cell volume with the lowest declustered mean.....	55
Figure 5.20: Typical features of a single structured semi-variogram	56
Figure 5.21: Box and whisker plots of each domain and the combined domain.	58
Figure 5.22: Variogram fans in the horizontal direction the purple arrow indicates the direction of maximum continuity	59
Figure 5.23: Variogram fans across strike the purple arrow shows the direction of maximum continuity	59
Figure 5.24: Variogram fans in the dip plane the purple arrow indicates the direction of maximum continuity	60
Figure 5.25: Downhole semi-variogram - Original U_3O_8 (ppm) composites	61
Figure 5.26: Downhole semi-variogram - Normal score transformed data	61

Figure 5.27: Major Direction (1) $00^{\circ} \rightarrow 130^{\circ}$ semi-variogram U_3O_8 (ppm) composites	62
Figure 5.28: Major Direction (1) $00^{\circ} \rightarrow 130^{\circ}$ semi-variogram Normal score	62
Figure 5.29: Semi-major direction $30^{\circ} \rightarrow 220^{\circ}$ semi-variogram U_3O_8 (ppm) composites	63
Figure 5.30: Semi-major direction $30^{\circ} \rightarrow 220^{\circ}$ semi-variogram - Normal score data ..	63
Figure 5.31: Minor direction $-60^{\circ} \rightarrow 220^{\circ}$ semi-variogram U_3O_8 (ppm) composites ..	64
Figure 5.32: Minor direction $-60^{\circ} \rightarrow 220^{\circ}$ semi-variogram – Normal score data	64
Figure 5.33: Semi variogram model for U_3O_8 (ppm) composites	66
Figure 5.34: Normal Score back transformed semi-variogram model for U_3O_8 (ppm)	66
Figure 6.1: Theoretical relationships between the covariance and the semi-variance	67
Figure 6.2: All composite samples available for KNA. A; 3-D view from top B; 2-D view from top C; 2-D view from right D; 2-D front view facing North. .	71
Figure 6.3: Location of well-informed blocks with three internal and surrounding drill holes.	72
Figure 6.4: Location of poorly-informed blocks with one internal and surrounding drill holes.	72
Figure 6.5: Conditional bias performance measures A: SLOR (blue circles) and KE (red triangles) for the range of block sizes tested. B: Sum (red bars) and number (blue bars) of negative weights for the range of block sizes tested.	73
Figure 6.6: Number of samples vs A: SLOR (blue circles) and KE (red triangles) for a 50m x 50m x 15m block, B: selected minimum and maximum number of samples highlighted	76
Figure 6.7: Discretisation configuration optimisation for the 50m x 50m x 15 m blocks	79
Figure 6.8: Block covariance vs discretisation configuration for 50m x 50m x 15m blocks.	80

Figure 6.9: A: Block size optimisation for poorly-informed blocks. B: Negative weights for the tested block sizes.	81
Figure 6.10: Number of samples vs SLOR (blue circles) and KE (red triangles) for a 150m x 200m x 15 m block.....	83
Figure 6.11: SUPERVISOR: Block Covariance versus Discretisation configuration for 150m x 200m x 15m blocks.	85
Figure 6.12: Axis rotation from SUPERVISOR YXZ to SURPAC ZXY axis rotation	86
Figure 6.13: Semi-variogram parameters from a back transformed model that were used in SURPAC.....	87
Figure 6.14: Semi-variogram range ratios that were used in SURPAC.....	87
Figure 6.15: A plan view showing a 100m x 100m x 10m block centred at Y= 120 326, X = 81 997, and Z= 422.	89
Figure 6.16: SURPAC: Number of samples vs SLOR for 50m x 50m x 15 m blocks.	93
Figure 6.17: SURPAC: Optimisation of the minimum and maximum number of samples for the 200m x 150m x 15m blocks.....	98
Figure 7.1: Swath plot by Easting for U ₃ O ₈ grades for all domains combined in well-informed blocks.....	106
Figure 7.2: Swath plot by Northing for U ₃ O ₈ grades for all domains combined in well-informed blocks.....	106
Figure 7.3: Swath plot by Elevation for U ₃ O ₈ grades for all domains combined in well-informed blocks.....	107

List of Tables

Table 3.1: Co & Cd KNA results of four different kriging neighbourhood searches.	15
Table 3.2: Performance comparison of true values vs estimated values of Co and Cd for four scenarios.....	16
Table 5.1: Descriptive statistics of raw (un-composited) sample lengths (m).....	33
Table 5.2: Frequency and cumulative frequency distribution table of sample lengths.	34
Table 5.3: Statistics of individual U ₃ O ₈ assay grades and 1.5m composited grade in the different domains.....	38
Table 5.4: Domain 6 un-composited assay values for drill hole PA61.....	39
Table 5.5: Domain 6 composited assay values for drill hole PA 61	40
Table 5.6: Domain 7 un-composited assay values for drill hole PA 66.....	40
Table 5.7: Domain 7 composited assay values for drill hole PA66	41
Table 5.8: Domain U ₃ O ₈ (ppm) Statistics before and after capping	49
Table 5.9: Statistics of the domains before and after clustering	56
Table 5.10: Accepted double structured spherical semi-variogram parameters	65
Table 6.1: SUPERVISOR: KN measures on 50m x variable Y m x 15m blocks.....	75
Table 6.2: SUPERVISOR: KN measures on variable X m x 50m x 15m blocks.....	75
Table 6.3: SUPERVISOR: well-informed variogram search ranges	77
Table 6.4: SUPERVISOR: Search range optimisation for 50m x 50m x 15m blocks	78
Table 6.5: SUPERVISOR: KN measures on 200m x variable Ym x 15m blocks.....	82
Table 6.6: SUPERVISOR: KN measures on variable Xm x 200m x 15m blocks.....	82
Table 6.7: SUPERVISOR: Search range optimisation for 150m x 200m x 15m blocks	84
Table 6.8: SURPAC ZXY ellipsoid orientation angles	87
Table 6.9: SURPAC: All the tested block sizes and their conditional bias measures	90
Table 6.10: SURPAC: KN measures on variable Ym x 50m x 15m blocks.....	92
Table 6.11: SURPAC: KN measures on 50m x variable Xm x 15m blocks.....	92

Table 6.12: SURPAC: Search range optimisation for 50m x 50m x 15m blocks.....	94
Table 6.13: SURPAC: Discretisation Optimisation of well-informed blocks	95
Table 6.14: SURPAC: Conditional bias measures of poorly-informed block sizes tested.....	96
Table 6.15: SURPAC: Results of KN measures on variable Ym x 200m x 15m blocks.	97
Table 6.16: SURPAC: Results of KN measures on 200m x variable Xm x 15m blocks.	97
Table 6.17: SURPAC: Search range optimisation for 200m x 150m x 15m blocks...	99
Table 6.18: SURPAC: Discretisation optimisation 200m x 150m x 15m blocks.....	100
Table 7.1: KN parameters from SUPERVISOR and SURPAC.....	101
Table 7.2: KN statistics from SUPERVISOR and SURPAC	102
Table 7.3: Block model extents and size.....	104
Table 7.4: Sample search parameters.....	104
Table 7.5: Comparison of block model U ₃ O ₈ averages with composited and declustered mean U ₃ O ₈ averages	105

List of Symbols

Approximately	$\approx$
Betafite	$[(\text{Ca}, \text{Na}, \text{U}, \text{REE})_{16-x}(\text{Nb}, \text{Ti})_{16}\text{O}_{48}(\text{O}, \text{OH}, \text{F})_{8-y} \cdot z\text{H}_2\text{O}]$
Beta-uranophane	$[\text{Ca}(\text{UO}_2)_2(\text{SiO}_3\text{OH})_2 \cdot 5\text{H}_2\text{O}]$
Billion years	Ga
Cadmium	Cd
Calcium carbonate	CaCO_3
Chromium	Cr
Cobalt	Co
Copper	Cu
Degree	$^\circ$
Easting	X
Elevation	Z
Greater than	$>$
Greater than or equal to	$\geq$
Kilometre	Km
Lead	Pb
Less than	$<$
Metre	M
Million years	Ma
Nickel	Ni
Northing	Y
Percentage	%
Plus – minus	$\pm$
tonnes per cubic metre	t/m^3
Tri-uranium octoxide	U_3O_8
Uraninite	UO_2
Uranium grade kilogram per tonne	kg/t
Zinc	Zn

List of Acronyms

Absorption rate constant	K_a
Best Linear Unbiased Estimator	BLUE
Coefficient of Variation	CV
Competent Person's Report	CPR
Correction factor	V_{slope}
Correlation Coefficient	CC
Diamond Drilling	DD
East	E
Elimination constant	K_{el}
Formation	Fm
Identification	ID
Inverse Distance Weighting	IDW
Kriging Efficiency	KE
Kriging Neighbourhood Analysis	KNA
Kriging Variance	KV
Localised Kriging Neighbourhood Analysis	LKNA
Maximum	Max
Mean Square Error	MSE
Minimum	Min
North	N
Nugget effect of the variogram	C_0
Ordinary Kriging	OK
Parts per million	ppm
Protein Data Bank	PDB
Quantile-quantile	Q-Q
Quantitative Kriging Neighbourhood Analysis	QKNA
Range of the variogram	A
Reverse Circulation	RC
Selective Mining Unit	SMU
Sill of the variogram	C_1
Simple Kriging	SK
Slope of Regression	SLOR
Standard deviation	Std dev
Therapeutic Drug Monitoring	TDM
Three-dimensional	3D
Two-dimensional	2D

1 Introduction

Ordinary Kriging produces best linear unbiased block estimates; the kriging weights of the samples have been determined by minimising the variance of the estimation error, this minimum estimation error variance is known as the kriging variance (Journel & Huijbregts, 1978).

The smoothing effect caused by most linear estimation methods such as Ordinary Kriging is the main cause of conditional bias (Yamamoto, 2005). Conditional bias occurs when the expected true values are not equal to the estimated grades. The amount of smoothing is dependent on: the continuity of the mineralisation of a regionalised variable (represented by the modelled variogram), the number of samples that is used in the kriging estimation process as well as the search scheme defined by the kriging neighbourhood (Deutsch, 2007).

The accuracy of kriging estimates, and the quality of conditional simulations are both dependent on the kriging neighbourhood applied in the kriging estimation procedure. A kriging neighbourhood that is, either too restrictive that it excludes correlated samples, or one that is too big that it includes un-correlated samples, can lead to conditional bias (Vann *et al.*, 2003; Khakestar *et al.*, 2013). An important step in Mineral Resource estimation is the minimisation of conditional bias and can be addressed by performing a quantitative kriging neighbourhood analysis, whereby the optimum kriging neighbourhood parameters are identified prior to the actual estimation. (Vann *et al.*, 2003).

Due to the importance of minimising conditional bias in the geostatistical estimation of Mineral Resources this research study into the outcomes of the KNA procedures as applied in the SUPERVISOR geostatistical software package versus that obtained using SURPAC geology and mine planning software package, was embarked on. For comparison purposes, the same information (input data, block configurations, test locations, and semi-variogram parameters) were used in both software packages to guarantee that the differences that might be observed in the results are not user defined.

The literature review on the comparison of software packages revealed that software comparisons have only been done in fields such as geology, computer science, biotechnology, bio-medical but not for the results of the Kriging Neighbourhood (KN) selection parameters using different software packages. From the literature review this research author concluded that most of the published software comparison studies, revealed that at least one software package gives better results than the others.

This research study attempts to answer the following question:

“Is there a significant difference in KN parameters obtained from using the SUPERVISOR KNA procedure and those derived from using the SURPAC KNA procedure when identical input data, semi-variogram model parameters and test locations are used?”

To deliver on the objective of the research under the constraint of identical data it was necessary to first carry out the initial steps that are naturally followed in resource evaluation. This process can broadly be divided into: the consideration of the geological model, logging and assay data validation and an exploratory data analyses applying both statistics and geostatistics, and then modelling the variogram to be used in the estimation. The next step in the estimation is the KNA process wherein the input parameters for the kriging estimation process are analysed, the outcomes thereof define the optimum input for the kriging plan that is finally applied in the block model estimation of grades and tonnages. To answer out the above research question from a practical point of view, a case study on Rossing Uranium Mine is considered.

Below is an outline of the methodology followed for addressing the focus question of this research:

- The determination of the optimal kriging neighbourhood for blocks with different data configurations using SUPERVISOR and SURPAC.
- The analysis and comparison of the conditional bias output measures for the block configurations tested as obtained from applying the Kriging Neighbourhood Analysis procedures in respective packages,

- Application of the identified optimum Kriging parameters and estimating the block model using Ordinary Kriging in both packages,
- Comparing the outcomes of the above exercises and reaching a conclusion about the KN results from the two software packages.

The report is structured as follows:

Chapter 1 provides context to the problem and a justification for the research study, including a discussion on the methodology followed to answer the research question.

Chapter 2 gives context and background to this research study, the application and the choice of software in the mining industry is discussed. The importance of the kriging neighbourhood analysis in resource estimation is detailed. A discussion of the literature review on approaches used to compare software packages in other areas such as computer science, biotechnology, geology, and bio-technology.

Chapter 3 discusses the review of proposed methodologies for defining the kriging search neighbourhood.

The data set used in this research study is sourced from Rossing Uranium Mine. An understanding of the geology of the deposit is a fundamental step in resource estimation and Chapter 4 has been dedicated to the relevant geology.

Statistical and geostatistical analysis of the domained data is done in Chapter 5. Semi-variogram model parameters which were used to determine the KNA were established at this stage.

Chapter 6 explains in detail the Kriging Neighbourhood Analysis (KNA) methodology as applied in SUPERVISOR and SURPAC. The differences between the KNA processes in the two software packages are discussed at the end of Chapter 6. KNA parameters obtained in Chapter 6 were used to estimate block values which were validated against composite mean grades in Chapter 7. Chapter 8 concludes the findings of this research study.

2 Research Background and Literature Review

The use of computer software programs in the mining industry dates back to the mid-1960s when mining companies developed their own in-house reserve modelling and mine planning software packages (Sides, 1997). Ever since then software vendors have been developing such software packages. To date, there are several geological modelling and geostatistical software packages, the more common ones being MICROMINE[®], DATAMINE, ISATIS[®], MINESIGHT[®], SURPAC, SUPERVISOR, VULCAN, and LEAPFROG[®]. Users of these geological modelling and geostatistical estimation software packages are educational institutions, exploration and mining companies, and companies consulting to the Minerals Industry.

Given that a variety of software packages are available for geological modelling and resource estimation, the question is, “how do software users choose which software to purchase?” According to (Dionne, 2015), the main aspects to consider are: a software that fulfils the company’s needs, user-friendliness, efficiency, ability to exchange data between departments, ability to be used with other software packages, learning and customer support from the software vendor, and the monetary value of the software.

Fischer & Getis, (2009) believe companies prefer a software package that comes with all the necessary functionalities within a single package and that it is user-friendly. In case of educational institutions, price matters the most. The software with the most licenses at a reasonable price is often chosen.

The choice of the two software packages used in this research came down to software availability and affordability. The university at which the researcher is enrolled has a SURPAC license and SUPERVISOR was found to be affordable amongst other software packages considered.

A brief overview of how the Kriging Neighbourhood Analysis (KNA) fits in the bigger picture of a resource estimation process follows.

Despite all the advantages that come with the linear estimation techniques, such as ordinary kriging, they have a tendency of reducing the variability of the estimates

(Abzalov, 2016). The variance of the estimates is lower than what the actual variance of the selective mining unit (SMU) supposed to be (Darling & Noble, 2011). This shortcoming is known as the smoothing effect (Yamamoto, 2005; Darling & Noble, 2011; Abzalov, 2016). Imperfect correlation between the variance of the block estimates and the input data is caused by lack of information (Vann *et al.*, 2003; Dohm, 2015b). Smoothing causes overestimation of lower values and underestimation of high values (Yamamoto, 2005).

The smoothing effect that comes with kriging is reduced through a process called kriging neighbourhood analysis (KNA) formalised by Vann *et al.*, (2003) and which is the subject of this research.

Resource estimation methods are classified into two groups; the non-geostatistical and geostatistical methods. The common non-geostatistical methods are classical polygonal, triangulation, cross-sectional, estimation by panels (blocking) and inverse distance weighting (IDW). IDW is more advanced and accurate compared to other non-geostatistical methods (Abzalov, 2016).

Despite the fact that non-geostatistical methods are easy to understand, the accuracy of their estimates cannot be measured (Mucha, 2013) thus, such methods only give estimates without theoretical evaluation of the estimation error. To overcome this drawback, another group of estimation method called geostatistical methods were developed.

What makes geostatistical methods more appealing than non-geostatistical methods is that they recognise and can model the anisotropy, decluster the data, and also create unbiased estimates (Yamamoto, 2000).

Examples of geostatistical estimation methods using kriging are ordinary kriging, simple kriging, indicator kriging, multiple indicator kriging, kriging with a trend, kriging with an external drift, factorial kriging, and non-linear kriging. Kriging estimators give a minimum error variance, an advantage above the non-geostatistical methods (Vann & Guibal, 1998).

The minimum variance of the errors of estimation also known as the kriging variance, indicates the precision of the estimated values to the true values (Rivoirard, 1990). It is for this reason that kriging is also known as the **Best Linear Unbiased Estimator (BLUE)** (Dohm, 2015b).

The discovery of geostatistical methods does not mean that non-geostatistical methods are no longer in use. They continue to be of great use especially during the early stage of mineral projects, where they are used to estimate interim results, and to check the results for the geostatistical estimation methods. The geologic setting, amount of information available, and the characteristics of the geological model are the factors to consider in order to determine which kriging method to use (Rossi & Deutsch, 2013).

Ordinary kriging (OK) is the most commonly used estimation method and is the method applied in this research study. It is also the methodology currently used at Rossing Mine and it provides acceptable block grades. The composite uranium grade distributions have a coefficient of variation (CV) value less than 1 in most of the estimation domains identified in this research, supporting the use of OK for the Rossing uranium resource estimation.

2.1 Review of Approaches for the comparison of software packages in non-mining related fields.

Wahbeh *et al.*, (2011) compared four free data mining toolkits namely WEKA, Tanagra, KNIME and Orange to be used for the classification task in terms of data accuracy and correctness. They did the comparison test by running nine different data sets applying six common classification algorithms in free data mining toolkits. The performance of the algorithms was tested using two test modes. The results showed that none of the data mining toolkits are better than the other for the classification task. The authors believe this was because of the type of the data sets used and the way the algorithms are implemented within the toolkits. In terms of algorithm applicability, all algorithms ran successfully in WEKA, five algorithms ran in Orange and only two ran in Tanagra and KNIME. The rest of the algorithms were unable to run on some data sets. In Orange, one of the algorithms could not run because this algorithm was not

implemented in the toolkit. Whereas for Tanagra and KNIME, the rest of algorithms did not run for any of the following three reasons: The first reason is that the data set is a multi-class instead of binary data. The second reason is that the data set contains discrete values which the algorithm cannot deal with. The third reason is that the algorithm is not available in the data mining toolkit. WEKA was ranked as the best toolkit in terms of algorithm applicability followed by Orange, Tanagra and then KNIME. In terms of classification performance based on the two test modes, WEKA is the best toolkit followed by Orange, KNIME and then Tanagra.

Neef *et al.*, (1994) compared two therapeutic drug monitoring (TDM) software packages (USC*PACK and MW/PHARM) to calculate and determine the reliability of the population pharmacokinetic parameters. The population pharmacokinetics parameters calculated were; elimination constant (K_{el}), correction factor (V_{slope}), absorption rate constant (K_a), and bioavailability. In each software package, three programs were used to calculate the population parameters. The three programs used in USC*PACK are NPEM, BOXES and PC-ADAPT-I whereas KINFIT, KNIBES, and BAYES-ON were used in MW/PHARM. Eight patients with Parkinson's disease who were on levodopa therapy for more five years were used in the study. Data analyzed in this study are from subcutaneous and intranasal administration. Naïve pooling, the standard two-stage method, and NPEM are the three population techniques that were used to analyze the data. Results show that there is a significant difference between K_{el} found by BOXES and BAYES-ON and none of the two K_{el} is represented by any patient from the population. The parameters found from subcutaneous administration data are in accordance with the values from literature. A significant difference was noted between the K_{el} and V_{slope} values from the standard two-stage method and naïve pooling obtained using data from the intranasal compared to those from the subcutaneous administration. The possible reasons that might have caused the differences in the population parameters are however not stated. The K_a value obtained from subcutaneous administration is extremely rapid compared to the one from intranasal administration hence subcutaneous administration data is suitable. In addition, subcutaneous administration has a bioavailability of about 100%. The study

concluded that out of the three population techniques used, NPEM in USC*PACK offers most information about the population and it also gives reliable pharmacokinetic parameter values after subcutaneous administration.

Youssef *et al.*, (2014) compared two academic software packages (LOMETS and I-TASSER) with the commercial software, HyperChem 8.0, for predicting 3D protein structures. Four amino acids sequences of proteins were used in the study. The two academic software packages select the best predictable model differently. In LOMETS, 30 models with the Z-scores above the threshold Z-score cut are taken from nine servers. From the 30 models, the model with the highest Z-score is selected as the best predicted protein model. The method of predicting the best protein structure model in I-TASSER is more complicated. Secondary protein structures are first threaded through a representative Protein Data Bank (PDB) structure library. The fragments that are in aligned regions are removed in order for them to reassemble into full-length models while those that are in unaligned regions are built by ab initio modeling. The coordinates of all clustered structures are then averaged to get the cluster centroids. Fragment assembly simulation is done in order to remove steric clashes and to further refine the models. Another round of simulation is done to extract spatial restrains from the centroids. Structure decoys are clustered at this stage and the structure with the lowest energy is selected. The topological similarity between the protein structure produced by HyperChem 8.0 and those produced by LOMETS and I-TASSER was measured using the Template Modeling Score (TM-score). TM-score ranges between 0 and 1, with a higher value indicating stronger similarity and 1 indicating a perfect match between the two protein structures. TM-score of I-TASSER for the three models out of four tested were higher than those of LOMETS hence I-TASSER is more accurate than LOMETS. In addition, the authors pointed out that their study might not have revealed the full performance of both the academic software because of their limited software proficiency.

2.2 Review of Approaches for the comparison of software packages in mining related fields.

MacCormack & Eyles, (2012) assessed the impact of three different software programs on the accuracy of the 3D geologic model output. The software packages are ArcGIS, RockWorks 2006 and Viewlog. The study was aimed at comparing the output of models interpolated using Inverse Distance Weight (IDW) and Ordinary Kriging (OK) algorithms in the three different software packages. Input data were extracted from four synthetic grids with known point locations so that the variability between actual and interpolated values can be determined. Modelling parameters were kept the same for both interpolation methods for the three software packages so that any difference in the model output would not be due to different modelling parameters. The effect of five variables on the accuracy of the model output was tested. The variables considered were; model complexity (from simple to complex model), sampling patterns (random, regular, and clustered samples), data quantity (100 256 and 676 points), algorithm and software selection. The accuracy of the output 3D models from the three software programs was measured by the root mean square error and bias error. Results revealed that all five variables studied have an influence on the accuracy of the model output. Model complexity was found to have a greater influence on the model output accuracy, followed by sampling pattern, data quantity, algorithm, and finally software selection. The study concluded that although software programs use the same interpolation algorithms, they do not necessarily give the same model output hence this variability should be reflected on when assessing uncertainty associated with subsurface model results. They advise that the software program that is used to model complex subsurface geological environments, interpolate clustered data and when using large quantities of data should be considered as a possible cause of model uncertainty.

The author of this research noted the following observations after reviewing the literature on the comparison of different software packages. Some software packages have similar algorithms which differ in the way in which they were implemented within each software package. One software package provides more information about the population than others despite the fact that the same input information was used in each

software package. It was also noted that some of the software packages compared use different approaches to predict the final results. In summary, one publication concluded that none of the software packages compared proved to be better than the others for data accuracy and correctness classification task. Three reviewed publications on software comparison concluded that at least one software package outperformed the others.

3 Kriging Neighbourhood Analysis (KNA) – Literature Review

KNA is a process of determining the best search volume, by optimising parameters such as block size, minimum and maximum number of samples, search range, and discretisation points, that would minimise conditional bias before kriging estimation and conditional simulation (Vann *et al.*, 2003).

The kriging efficiency (KE) and slope of regression (SLOR) are the two statistical measures that are calculated for each tested block in a model during the KNA test phase. The KE measures how well kriging estimates reflect the true block values (Mining Associates, 2013). The SLOR measures the degree of conditional bias on kriged estimates (Krige, 1996):

The KE and the SLOR are calculated using the formulas below:

$$\text{Kriging Efficiency} = \text{KE} = \frac{\sigma_B^2 - \sigma_K^2}{\sigma_B^2}$$

$$\text{Slope of Regression} = \text{SLOR} = \frac{\sigma_B^2 - \sigma_K^2 + \lambda}{\sigma_B^2 - \sigma_K^2 + 2\lambda}$$

Where, σ_B^2 = block variance, σ_K^2 = kriging variance, λ = LaGrange multiplier

Negative KE values arise when the kriging variance is greater than the block variance (Krige, 1996b). Negative KE values indicate that the kriged estimates poorly estimate the block grades.

Low positive KE values indicate high kriging variance relative to the block variance (very poor, highly biased kriging estimates). High (close to 1 or 100 %) KE values signify low kriging variances (good unbiased kriging estimates) thus the variance of the block estimates is almost equal to the variance of the true block values (Deutsch *et al.*, 2014).

The regression of true values given the estimates is an indication of the conditional bias. The estimated block values for block size V, Z_V^* are the independent variables and are plotted on the x-axis. The true values Z_V (block size V) are unknown and are plotted on the y-axis (Figure 3.1). The variance of the true expected values is always

higher than that of the estimated values due to the smoothing effect of linear estimation methods such as Kriging. This is the reason why the slope of regression is always flatter than the $Z_V = Z_V^*$ line (Deutsch, 2007).

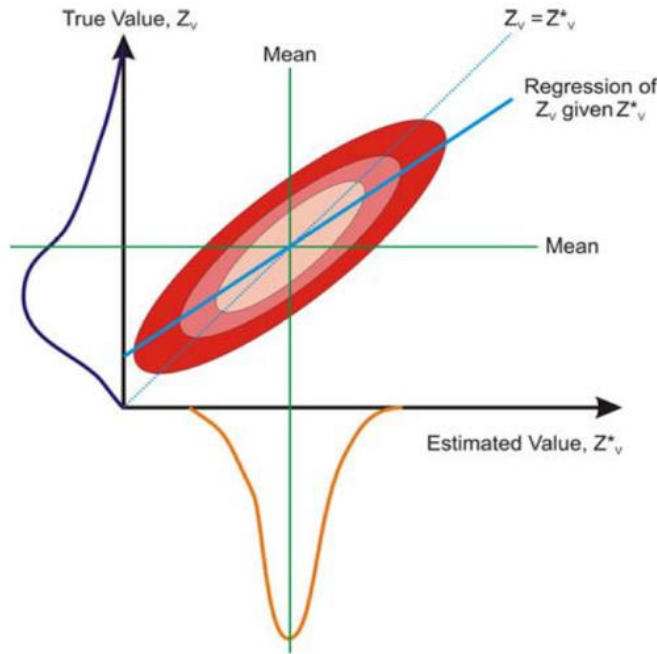


Figure 3.1: Linear regression as a result of the information effect

(Adapted from Deutsch, 2007)

For OK the slope of regression can be calculated from the kriging weights and the covariogram function, detailed in section 6.1 later.

Rivoirard, (1987), Isaaks & Srivastava, (1989) , Vann *et al.*, (2003) and Coombes & Boamah, (2015) have proposed methods of determining the search neighbourhood.

Rivoirard, (1987) pointed out that the two parameters that are crucial in choosing the best search neighbourhood are the weight of the mean for Simple Kriging (SK) and the SLOR. The weight of the mean shows the dependence of kriging on the neighbourhood and it also determines the minimal neighbourhood while the slope of regression determines whether the neighbourhood is large enough. The weight of the mean determines whether SK or OK is to be used to select the minimal neighbourhood. A wrong choice of the estimation method leads to conditional biasedness. If the weight

of the mean is small, either of the two kriging methodologies can be used. However, if the weight of the mean is large, OK cannot be used because it will result in conditional bias. Poor estimation causes a lower slope of regression (< 1). A smaller value of the slope of regression is an indication of conditional bias. To avoid conditional bias, the size of the neighbourhood needs to be enlarged.

Isaaks & Srivastava, (1989) use the following three critical questions in determining the search neighbourhood:

a. Are there enough nearby samples? An answer to this question will determine how wide the search ellipse should be. They believe that the wideness of the search ellipse should be ruled by the geometry of the data set and be based on the orientation and anisotropy ratio of an ellipse.

b. Are there too many samples? Answering this question helps in choosing the maximum number of samples of the search neighbourhood. The authors pointed out two factors that limit the size of the neighbourhood, the first being the slowdown in computational speed with an increase in sample size and the issue of whether the mean and the variance is constant (stationarity) as further samples are being included.

c. Are the nearby samples redundant? Redundancy of samples occur because of clustering and this is not an issue to estimation methods such as OK that accounts for clustering. The use of octant or quadrant search is proposed to reduce sample redundancy for estimation methods that do not account for declustering.

Computational speed is not an issue anymore with current software packages processing thousands of data in a matter of few minutes (Snowden Mining, 2013).

Vann *et al.*, (2003) expanded the criteria that Rivoirard, (1987) used in deciding the best search neighbourhood. In addition to the weight of the mean for SK and the SLOR, they added the distribution of kriging weights and the kriging variance. The distribution of kriging weights was added as another criterion to measure whether the neighbourhood is too restrictive or not. If the search neighbourhood is extended and the added samples are assigned positive weights, then it means that the neighbourhood

is too restrictive. The kriging variance, in addition to the slope of regression, also test the quality of the estimates. The authors demonstrated the methodology with case studies. Blocks with three different data configurations were tested to determine the best search neighbourhood for such block sizes. Blocks tested were classified depending on the number of samples inside or surrounding them. The scenarios they considered are; well-informed blocks (with samples on the sides and inside the block), less well-informed blocks (samples on the sides of the block and fewer samples inside than well-informed blocks) and poorly-informed blocks (with no samples inside and very few on the sides). In each category, the neighbourhood with the highest value of the slope of regression, and a minimal weight assigned to the simple kriging mean, a low kriging variance and the least number of negative weights, preferably $< 5\%$ of the total weight, is chosen as the best kriging neighbourhood.

The amount of negative weight allowed should be less than 5% of the total weight because negative weights can lead to negative estimates which do not make any sense (Szidarovszky *et al.*, 1987).

Negative weights arise when estimating a highly correlated variable (semi-variograms with long ranges and a low nugget effect) or when estimating clustered points (Armstrong, 1998). In such cases, data points which are within the search range but are a little further away from the estimated block are assigned negative weights because they are being screened by nearby data points (Sinclair & Blackwell, 2002).

The KNA results of the three scenarios considered in the-work of Vann *et al.*, (2003) showed that well-informed blocks give better conditional bias results than both the less well-informed and the poorly-informed blocks.

Coombes & Boamah, (2015) proposed a methodology of optimising a localised kriging neighbourhood analysis (LKNA). Their method is different from the classical way of determining the kriging neighbourhood analysis in the sense that it localises the search strategy. The classical way of performing a KNA involves testing various block configurations (well informed, moderately, and poorly informed blocks) in a domain to determine the best set of estimation parameters. The best set of estimation parameters

is then used to estimate all the blocks in the domain. In LKNA, each block in a model is treated differently. Instead of applying the best estimation parameters to the entire domain, combinations of all estimation parameters are evaluated for each block. The best estimation parameters for each block are picked and used to estimate each block. Kriging estimates produced using estimation parameters from LKNA are highly likely to be the same as the true estimates, when using LKNA parameters for estimation. However, LKNA requires processing of large files and this reduces computational speed, therefore the authors suggested that it should be built into the estimation process.

Rivoirard's, (1987) approach of defining the search neighbourhood was examined by Boyle, (2010) on a heavy metal soil data set. The dataset consisted of Cd, Co, Cr, Cu, Ni, Pb and Zn, with Cd and Co being selected for the investigation. The reason being, Co is the spatially more continuous variable with the least skewed distribution and Cd is less spatially continuous with a highly skewed distribution.

Four different kriging neighbourhood searches, with one to four rings of samples, were tested for Cd and Co. The KNA test was done on those four different scenarios to determine the optimal kriging neighbourhood. The results are shown in Table 3.1 Boyle, (2010).

Table 3.1: Co & Cd KNA results of four different kriging neighbourhood searches

(Source: Boyle, 2010)

Number of sample rings	Co		Cd	
	Average weight of the mean	Theoretical slope of regression	Weight of the mean	Theoretical slope of regression
1	-0.048	1.044	0.285	0.723
2	0.043	0.980	0.209	0.861
3	0.075	0.981	0.244	0.895
4	0.048	0.992	0.247	0.923

The author concluded that the KNA results for Cobalt favoured a search with one ring of samples because of the low weight of mean and high value of the slope of regression. It was noted that the negative value of the weight of mean for one ring of samples results from large weights being assigned to the samples in this ring and that there is a

redundancy of samples. For Cadmium, four rings of samples were found to be the best kriging neighbourhood to consider.

After obtaining the KNA results, Boyle went on to estimate the points of the data set for Cadmium and Cobalt for all four scenarios. A performance comparison test which involved scatter plots of true and estimated values, the correlation coefficient (CC) between true and estimated values, variance of the histogram residuals (Variance), mean square error (MSE) and actual slope of regression (Slope) between true and estimated values were compared, the results are shown in Table 3.2.

Table 3.2: Performance comparison of true values vs estimated values of Co and Cd for four scenarios.

(Source: Boyle, 2010)

Number of sample rings	Co				Cd			
	CC	MSE	Slope	Variance	CC	MSE	Slope	Variance
1	0.6879	6.681	0.476	6.542	0.1667	0.647	0.125	0.627
2	0.7144	6.214	0.526	6.085	0.1675	0.579	0.105	0.565
3	0.7173	6.178	0.521	6.029	0.1620	0.577	0.099	0.561
4	0.7093	6.312	0.511	6.172	0.1531	0.572	0.090	0.556

The best search neighbourhood from the KNA was compared with those from the performance comparison test. For Cadmium four sample rings were identified as giving the best KNA results, while for a performance comparison one ring of samples was favoured. In a case of Cobalt, one ring of samples was identified as giving the best KNA results whereas, in the performance comparison, one ring of samples performed worst compared to the other three scenario, three rings of samples gave better results. From these results, Boyle, (2010) concluded that the domaining of samples, the stationarity assumption, as well as sampling accuracy, also have an influence on the outcome of the kriging estimates in addition to KNA.

Khakestar *et al.*, (2013) used the slope of regression, the number of negative kriging weights and the kriging variance as measures for choosing the best kriging neighbourhood for reserve estimation of an Iron deposit. Firstly, cross-validation was used to choose optimal sill and range values. They cautioned against the use of the

variogram range in deciding on the kriging neighbourhood since it can be either too short (restrictive) or too long resulting in the inclusion of samples that are spatially uncorrelated. To identify the best search ellipsoid number a graph of the search ellipsoid number versus the mean variance, mean slope and proportion of negative kriging weights was plotted (Figure 3.2).

From this graph search ellipsoid number 9 was identified as the appropriate search, beyond this search an increase in a search ellipsoid size does not seem to improve the slope of regression, decrease the kriging variance, and an increase the proportion of negative weights. The search radius from the variogram range and the suggested search radius as identified in Figure 3.2 were compared using estimation error, the ratio between the mean of kriging variance estimated from the model and mean of squared of differences between actual and estimated value and the correlation coefficient. The suggested search radius was found to be the best kriging neighbourhood.

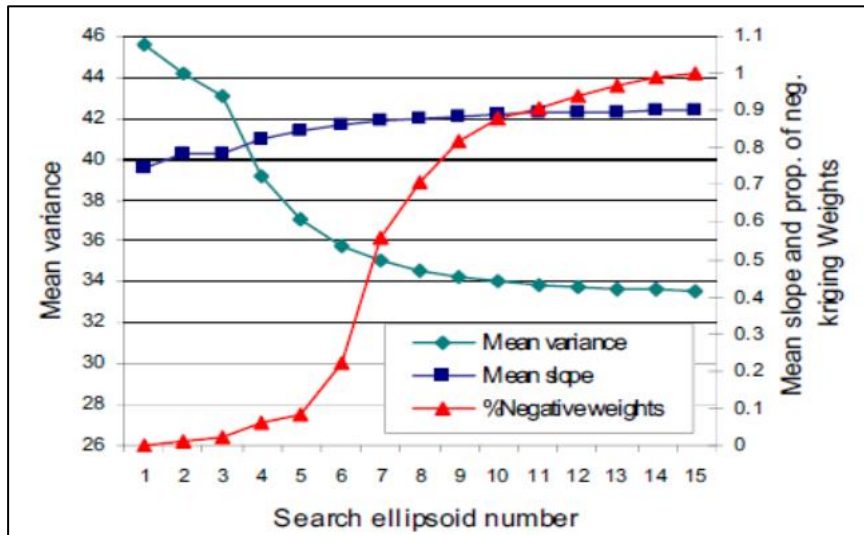


Figure 3.2: Search ellipsoid number versus mean variance, mean slope and proportion of negative kriging weights.

(Source: Khakestar *et al.*, 2013)

From the search neighbourhood methodologies encountered in the literature review, the methodology of Vann *et al.*, (2003) for determining the optimum search

neighbourhood definition, was followed and is reported on in this research. It is also the methodology used in both SUPERVISOR and SURPAC.

4 Rossing Uranium Deposit

An understanding of the geology of the deposit is very important in resource estimation. A resource estimation process that does not consider the geology of the deposit is considered to be a mathematical treatment of samples with no practical value (Stephenson & Vann, 2001).

The dataset used in this project is sourced from Rossing Uranium Mine. Rossing Uranium Mine is an open pit mine situated on the western part of Namibia, about 70 km North-east of the coastal town Swakopmund (Figure 4.1).

Rossing Uranium Mine is run by Rio Tinto which owns most shares (69 %), the rest of the shareholders are; Iranian Foreign Investment company (15 %), Industrial Development Corporation of South Africa (10 %), Namibian government (3 %) and local individuals (3 %) (Rossing Uranium Website, 2016).

Near Rossing Uranium Mine, there are other uranium mines and projects (Figure 4.2). These mines are Langer Heinrich Mine and Husab Mine (then Husab project) which are in operation and Trekkopje Uranium Mine which is on care and maintenance. Etango, Norasa (previously Valencia Project), Deep Yellow and Marenica Project are Uranium projects. Rossing, Husab, Valencia and Etango uranium deposits are classified as intrusive uranium deposits with their mineralisation associated with sheeted leucogranites (alaskites), whereas Langer Heinrich, Trekkopje, and Marenica deposits are surficial uranium deposit within the calcrete-hosted mineralisation. Deep Yellow project consists of a combination of both surficial and intrusive type of deposits.

Whilst this research is on the comparison of KN results from the application of the software packages the researcher considered it appropriate to consider the geological model and data to be used in the Case Study. For the completeness of a resource estimation process on the Case Study area an overview of the local geology of Rossing deposit is therefore reported below.

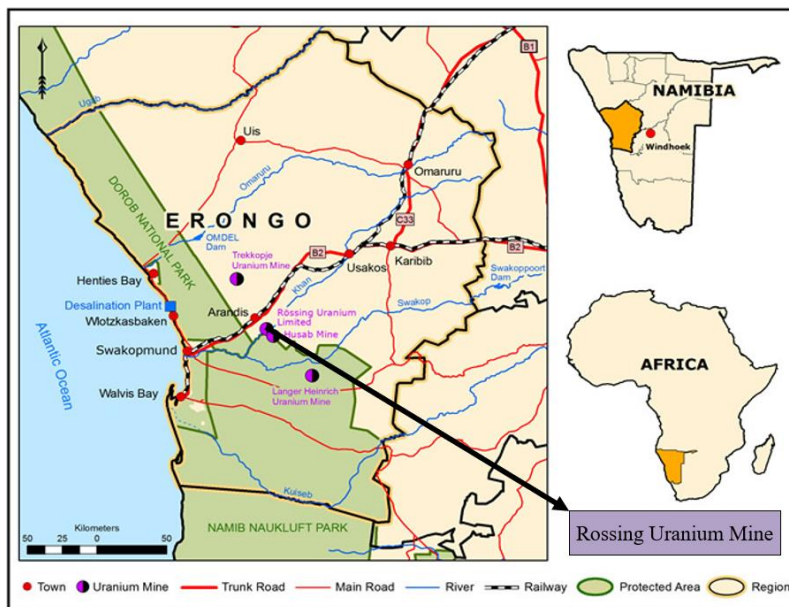


Figure 4.1: Location of Rossing Uranium Mine

(Source: Rossing Uranium Limited, 2011)

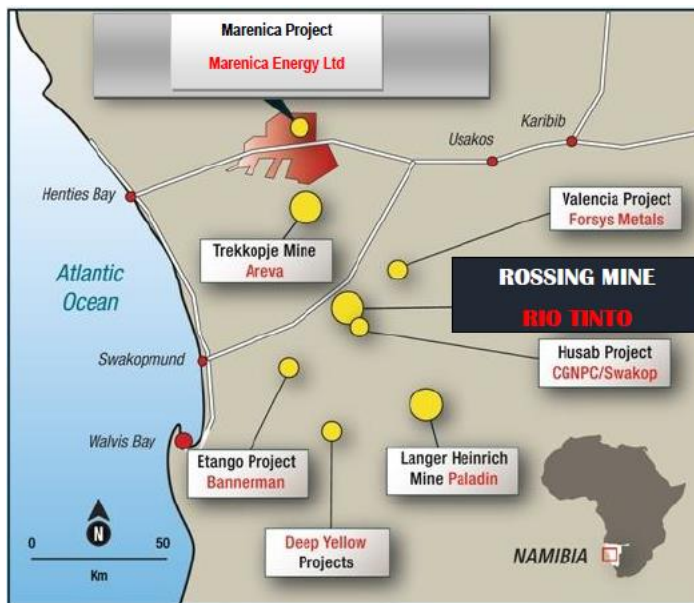


Figure 4.2: Location of other uranium mines and projects near Rossing Uranium Mine with their owner companies highlighted in red.

(Source: Marenica Energy Limited, 2010)

4.1 Local geology

The stratigraphic column of the Rossing area consists of metasediments belonging to the Damara sequence (Anderson & Nash, 1997). These sediments were deposited between 750 and 710 Ma (Abraham, 2009). The Damara sequence is made up of four formations; Etusis, Khan, Rossing and Chuos Formation (Nex & Kinnaird, 1995). However, the Chuos Formation is not exposed in the Rossing area. Most of Rossing's country rocks belong to the Khan and Rossing Formation with minor rocks from Etusis (Abraham, 2009). The oldest rocks in the Rossing area are quartzites, arkoses and minor conglomerates belonging to Etusis formation followed by semi-pelitic lithologies with a carbonate component (Khan Formation). The Rossing Formation, consisting of marble, chert, conglomerates, quartzites and cordierite gneiss, overlies the Khan Formation Figure 4.3 (Kinnaird & Nex, 2007).

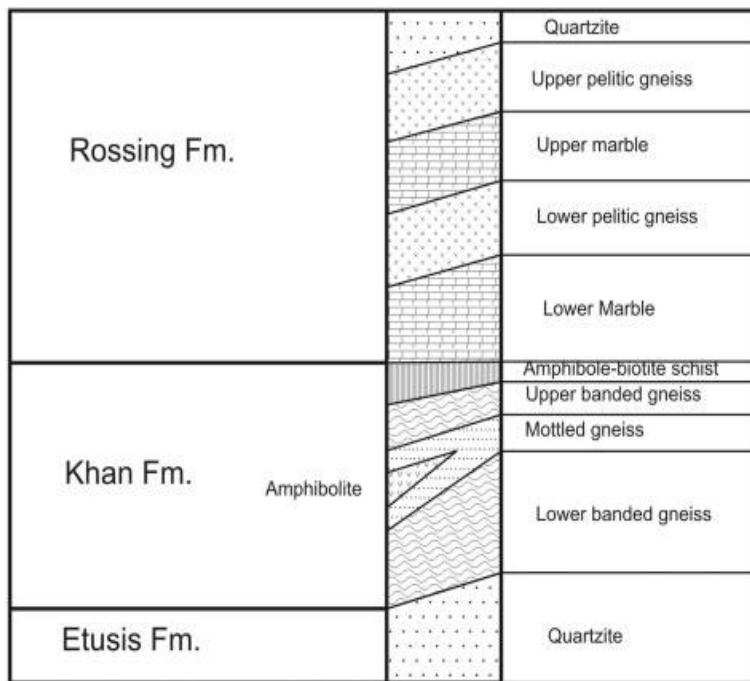


Figure 4.3: Stratigraphic column showing rocks exposed in the Rossing area.

(Source: Kinnaird & Nex, 2007 after Nash, 1971)

From late Proterozoic to early Palaeozoic (650 Ma to 460 Ma) Damara sequence rocks were deformed and metamorphosed to amphibolite-granulite facies during a polyphase deformation and metamorphism event (Anderson & Nash, 1997). Numerous granitization took place after the polyphase deformation event, resulting in syn-tectonic, late-tectonic and post-tectonic granites (Corner & Henthorn, 1978; Roesener & Schreuder, 1992; Abraham, 2009). These granites are vein-like intrusions, fine to coarse grained with high alkali content. They are 458 ± 8 million years in age and are confined to highly metamorphosed areas. Vein-like granite intrusions vary in thickness from few centimetres to several meters (Roesener & Schreuder, 1992). Because of the nature of these intrusions, they are often referred to as alaskites at Rossing. Alaskites are granitic rocks that are rich in alkali-feldspar with little or no mafic minerals (Kinnaird & Nex, 2007).

Rossing's alaskites are classified into six types (A-F) based on texture, colour, grain size and structural settings. Type A is the earliest intruded alaskite and F the latest (Nex & Kinnaird, 1995). Type A, B, and C were affected by the third deformation event (pre-and syn D3), they contain no uranium mineralisation while type D, E and F intrude after the third deformation event (post D3). Of all alaskites, only type D and E contain significant uranium mineralisation (Nex *et al.*, 2001).

Based on this observation, uranium mineralisation is associated with post-tectonic alaskites. Most of the uraniferous alaskites are found at the boundary between Khan and Rossing Formation or Khan and Chuos Formation in cases where the Rossing Formation is very thin or absent (Kinnaird & Nex, 2007). Alaskites show concordant, discordant and gradual relationship with the country rocks. (Roesener & Schreuder, 1992).

4.2 Uranium Mineralisation

Primary uranium mineralisation is confined to post-tectonic alaskites between the Khan and Rossing Formation. Secondary uranium mineralisation is widespread, thus it is not confined to any lithological unit or position (Nex & Kinnaird, 1995).

The most common primary uranium mineral in the alaskites is uraninite (UO_2), in some areas betafite $[(\text{Ca}, \text{Na}, \text{U}, \text{REE})_{16-x}(\text{Nb}, \text{Ti})_{16}\text{O}_{48}(\text{O}, \text{OH}, \text{F})_{8-y} \cdot z\text{H}_2\text{O}]$ can be the most common uranium phase. The most common secondary uranium phase in the Central zone of the Damara mobile belt is beta-uranophane $[\text{Ca}(\text{UO}_2)_2(\text{SiO}_3\text{OH})_2 \cdot 5\text{H}_2\text{O}]$ (Ministry of Mines and Energy, 2010).

Uraninite appears as inclusions or along cracks of quartz, feldspar, and biotite, or it can be found as free grains in host rocks (Bowden *et al.*, 1995).

At Rossing Mine, 50% of the uranium occurs as uraninite, 5% betafite and 45% occur as secondary uranium phases such as coffinite, beta-uranophane, and U-phosphates (Kinnaird *et al.*, 2011). Alaskite at Rossing mine host all the primary and most of the secondary mineralisation. Some of the secondary mineralisation can also be found in country rocks (Roesener & Schreuder, 1992; Abraham, 2009).

5 Statistical and Geostatistical analysis of the Rossing data set

5.1 Introduction

After the geology of the Rossing deposit was reviewed and understood, the Rossing data set was validated. Data validation was done in Microsoft Excel and also in SURPAC. This chapter gives the details of the data validation, orebody domaining, compositing, statistical and geostatistical analyses of the Rossing dataset.

Domaining and compositing of the validated data set were done in SURPAC because SUPERVISOR does not have functions to domain and composite data.

It is impossible, and it would be very expensive to drill the entire deposit hence in most cases less than one billionth of the volume of a deposit is drilled (Rossi & Deutsch, 2013). To estimate grades at un-sampled locations, one must make use of the few drilled samples.

5.2 Data validation

A reliable estimation is based on a clean data set (Dohm, 2015b), thus data validation must be done before any data analysis is done. The database was checked for negative assay values, inconsistencies in the naming of sample ID, inconsistent values for missing values, errors in survey information, data duplication, missing and overlapping sampling intervals.

The Rossing Uranium Mine dataset considered consists of 75 drill holes drilled on an irregular grid covering an area of about 800m (N) x 700 m (E) as shown in Figure 5.1.

For confidential purposes, the mine geologist shifted the easting and northing so that the data plot far from the location of the actual project. The data set consists of four files; assay, collar, survey, and geology file. Data validation was done for each file.

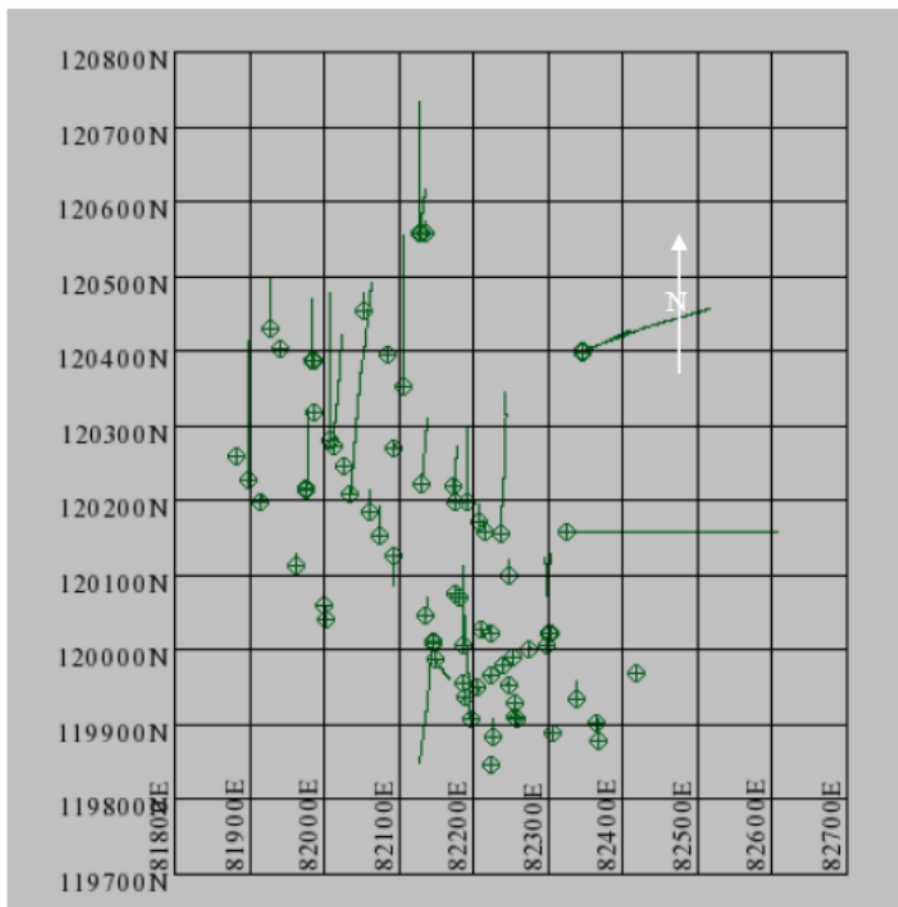


Figure 5.1: Rossing drill holes collar positions

5.2.1 Drill hole Collar Data Validation

The collar file had 75 drill holes with the following information: drillhole ID, easting, northing, reduced level (elevation), depth, and drill type. Three holes did not have reduced level and depth values and were thus discarded. In addition, one hole, identified in Figure 5.2 (red trace), deviated from the desired drilling orientation and as its drilling orientation could not be resolved, the mine geologist advised that the hole should be deleted from the collar file.

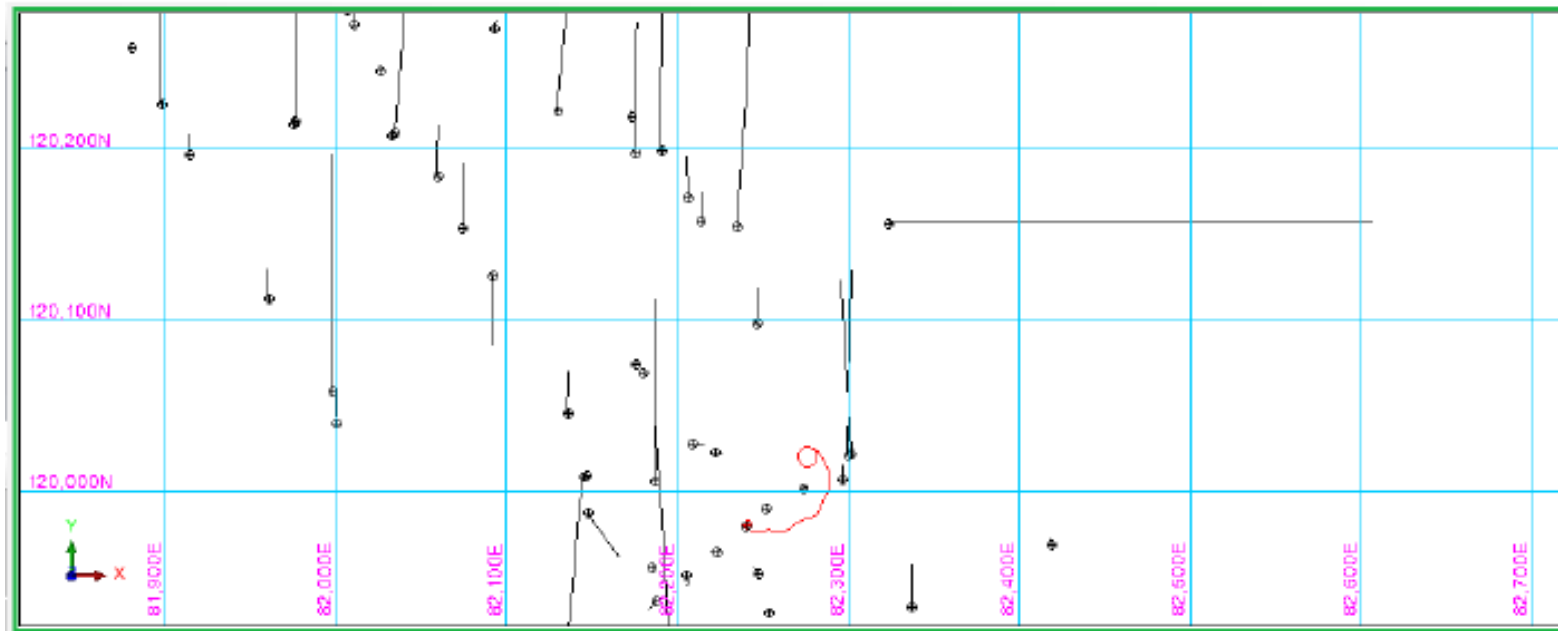


Figure 5.2: A plan view of the drill hole positions, highlighting the deviation of one drill hole (red)

The number of drill holes in the collar file reduced from 75 to 71. There are 62 Diamond Drilling (DD) drill holes and the remaining nine are Reverse Circulation (RC) drill holes. U_3O_8 (ppm) assays were done on samples collected from all the DD and none of the RC drill holes were chemically analysed for U_3O_8 .

5.2.2 Drill hole Survey Data Validation

The Survey file had 31125 individual records with the following information: holeID, depth, azimuth, dip, and survey type per record. When importing the survey file into SURPAC, two drill holes produced errors. These errors could be due to incorrect depth, azimuth, and dip values. The mine geologist was contacted with regards to these two drill holes and it was advised that they should be removed from the database. Therefore, the database was reduced to a total of 21874 individual records.

5.2.3 Drill hole Assay Data Validation

The Assay file consists of 4190 individual records with record fields being holeID, sample from, sample to, $CaCO_3$ and U_3O_8 . The assay values for $CaCO_3$ and U_3O_8 are in kg/t, the values were however, converted to ppm for ease of use. Three drill holes did not have the required sample “from” and sample “to” information, or U_3O_8 assay values, hence they were removed from the assay file. The removed drill holes accounted for 53 individual sample records, hence, the number of individual sample records in the assay file reduced to 4137. The file was checked for missing, overlapping, or mixed up intervals and all sample intervals were found to be acceptable.

5.2.4 Geological Data Validation

The geology file consists of 1531 individual sample records. The record fields for this file are; holeID, sample from, sample to, rock type and lithostratigraphy. Drill hole (PA06D26) was found to have a missing interval between 75.2 and 91 m. Upon enquiry with the mine geologist such intervals are replaced with ‘unidentified’ code -9999. Six holes, namely PA06D02, PA06D27, PA06D29, PA06D30, PA06D34 and PA06D35 had mixed up intervals. These intervals were placed in their respective correct

positions. When the geology file was imported into SURPAC, inconsistent naming of the drill hole ID was identified as an issue. An error message identified that drill holes ID PA06D026 and PA06D027 do not exist in the collar table. PA06D26 and PA06D27 were also captured as PA06D026 and PA06D027 respectively. The “sample from”, sample to, rock type and lithostratigraphy fields of these boreholes were compared and were found to be the same. Thus, PA06D26 and PA06D27 were maintained whereas PA06D026 and PA06D027 were deleted from the Geology file. The final geology file consists of 1458 records.

5.3 Domainining

Domainining of samples also has an influence on the outcome of the kriging estimate in addition to KNA (Boyle, 2010). According to Glacken & Snowden, (2001), if mineralisation corresponds to a geological unit, the geological model should be used to define domains.

The Rossing Uranium Mine Competent Person’s Report (CPR) is not available in the public domain and could not be reviewed. Husab deposit, located 8 km south of Rossing has the same geology and mineralization as the Rossing deposit. The publicly available technical report of Husab deposit (Rossing South) was therefore reviewed due to its relevance to this research.

An overview of the geological interpretation and modelling of Husab deposit is given below, and the same approach was used as a guide for the domainining that was done for the Rossing Case study used in this research project.

Mineralisation of Husab deposit is confined to sheeted leucogranites locally referred to as alaskites and the uranium deposit domainining was based on alaskite units. Alaskite contacts were used to guide sectioning and 3D modelling of mineralization shapes during modelling (Figure 5.3). The mineralisation model was based on a cut-off grade of 75ppm for grade continuity purposes. The 75ppm cut-off grade represents the natural uranium cut-off grade shown in the drill holes. The final 3D mineralisation domains for Husab deposit are shown in Figure 5.4. (Inwood & Hill, 2009).

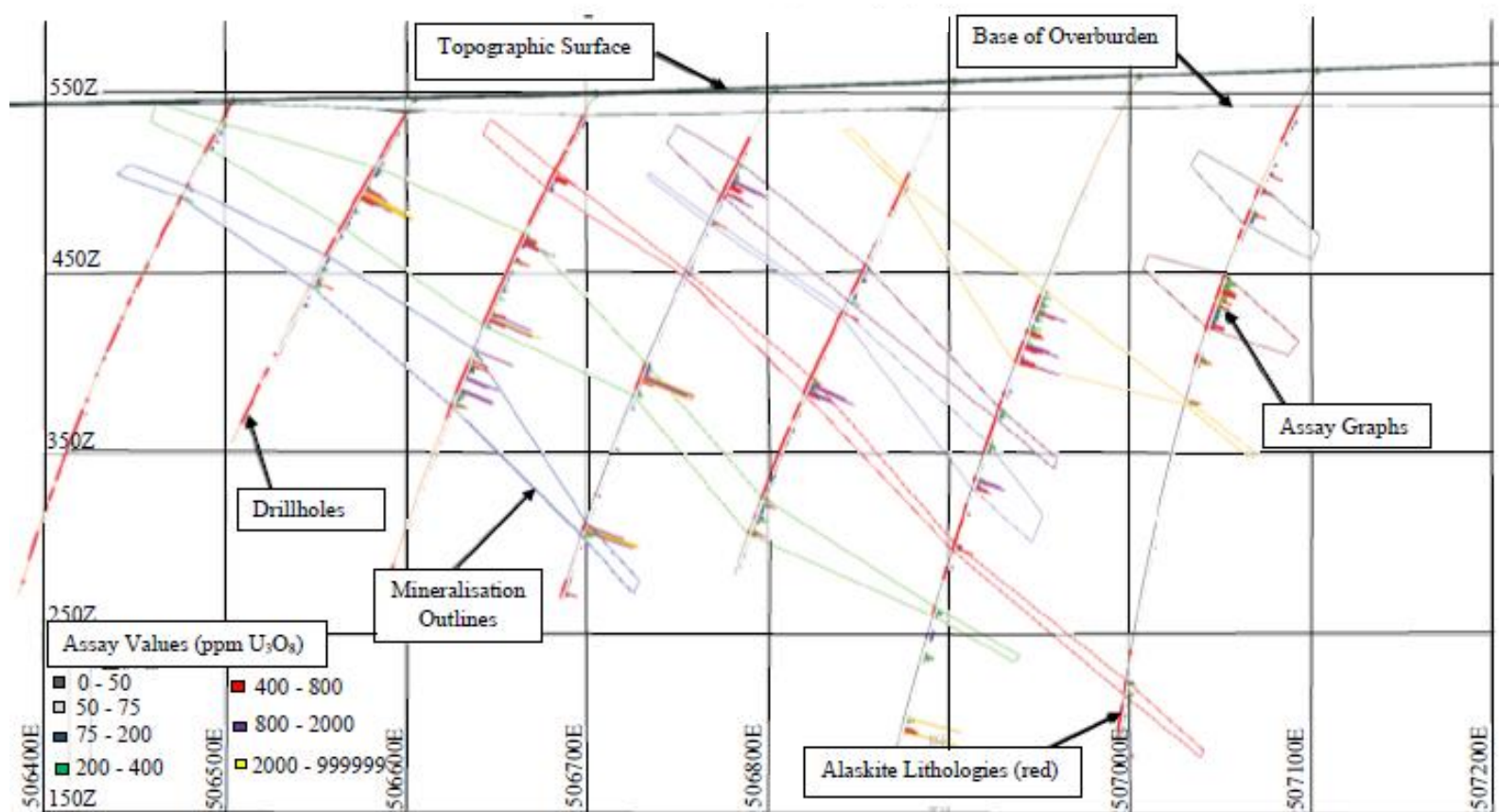


Figure 5.3: A section view facing North showing mineralization wireframes for Husab (Rossing South) deposit.

(Source: Inwood & Hill, 2009)

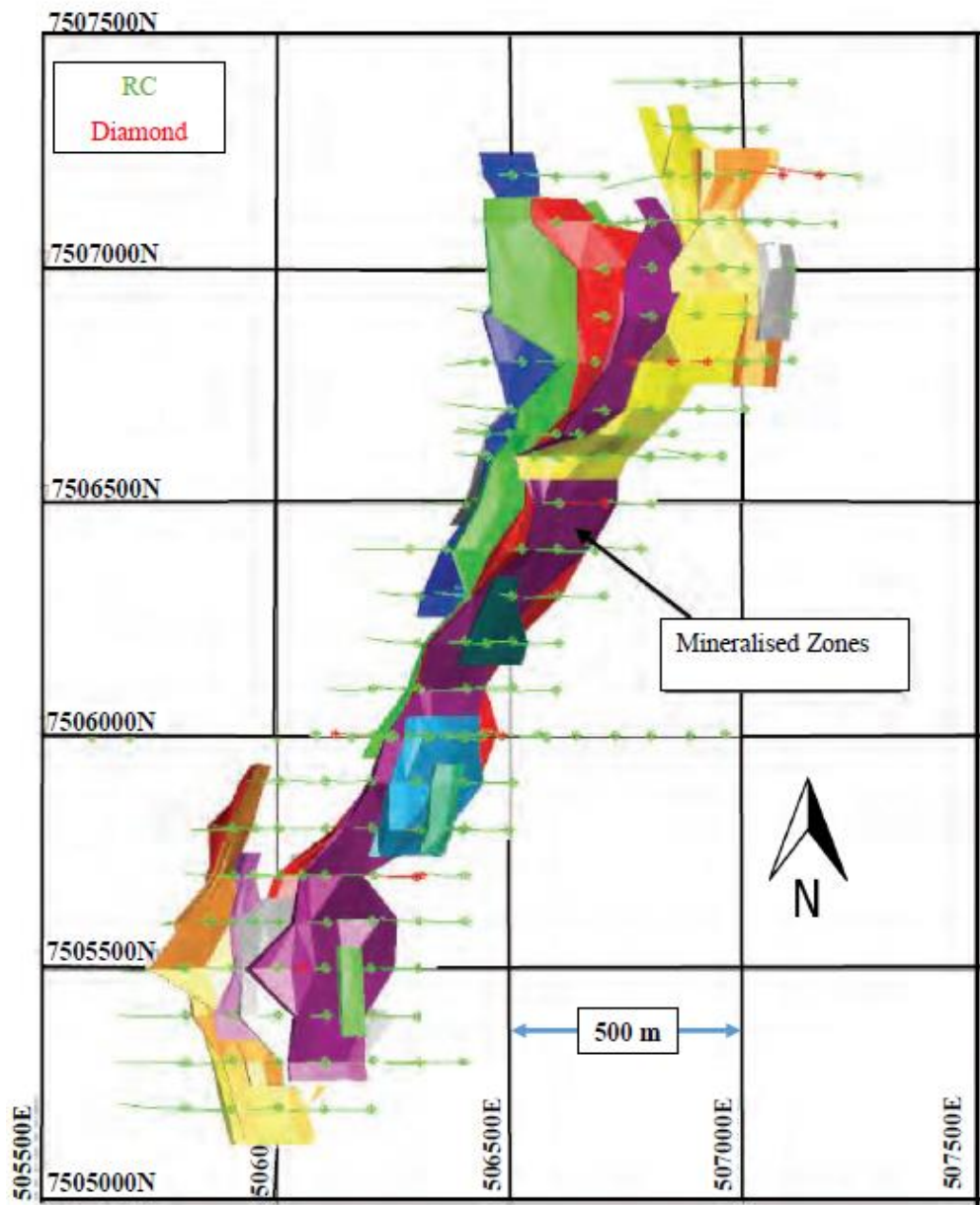


Figure 5.4: A plan view showing 19 three-dimensional mineralisation domains (distinguished by different colours) for Husab (Rossing South) deposit.

(Source: Inwood & Hill, 2009)

In this research, Rossing drill holes are displayed based on the alaskite lithologies (pink) and uranium grades. The colour coding for Uranium assay grades in Figure 5.5 and Figure 5.6 as follow: 0 – 100ppm = cyan; 100 – 300ppm = blue; 300 – 700ppm = magenta; 700 – 10 000ppm = red.

The cut-off grade at Rossing is 100ppm. For grade continuity purposes, alaskite contacts and uranium grades above a cut-off grade of 100ppm were used to guide the 3D modelling of the mineralisation shapes, Figure 5.5 and Figure 5.6 show examples of the mineralisation wireframes based on this information.

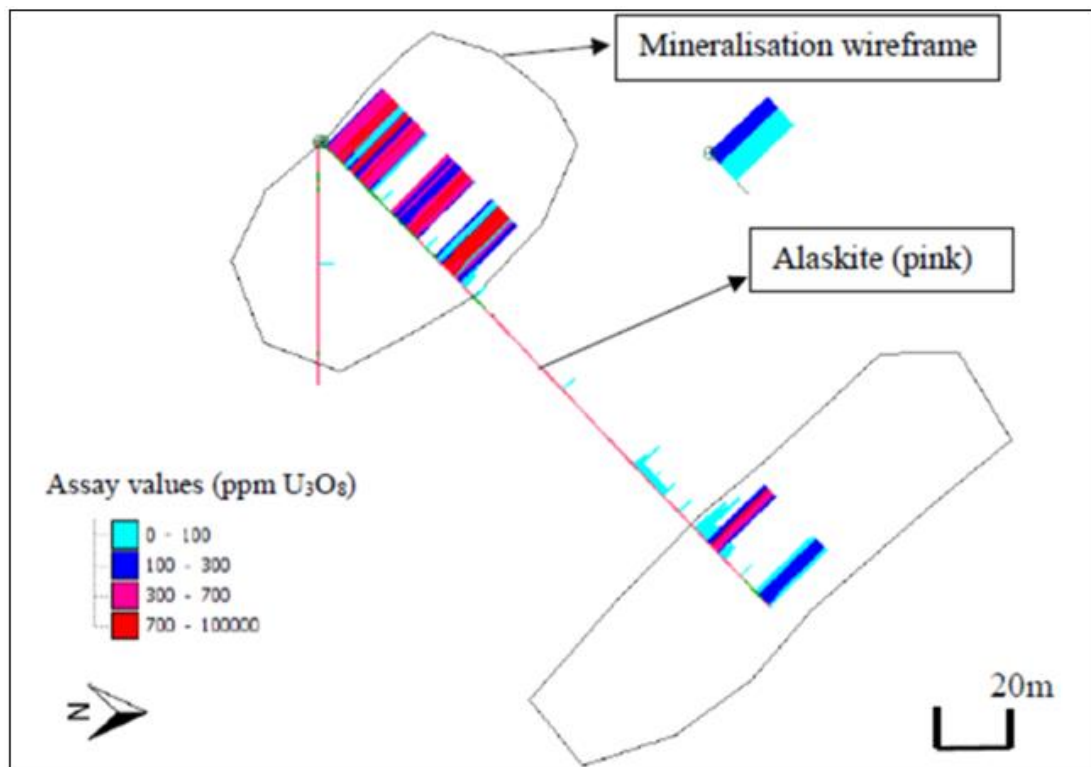


Figure 5.5: SURPAC – A section view showing ore wireframes sections based on a 100ppm uranium cut-off grade and alaskite contacts.

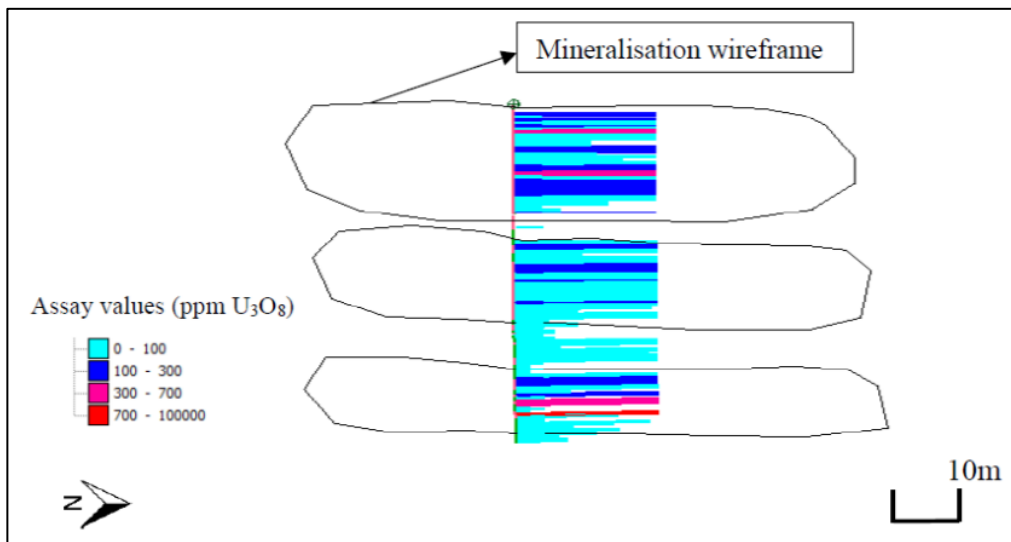


Figure 5.6: SURPAC - 82280 m E section showing ore wireframes sections based on a 100ppm uranium cut-off grade and alaskite contacts.

Wireframes that are on the same level in different drill holes were joined to form ore domains, the final seven mineralisation domains identified using this process are shown in Figure 5.7. Alaskites thickness vary from a few centimetres to several meters, which explains why there are so many separate three-dimensional domains (Figure 5.7).

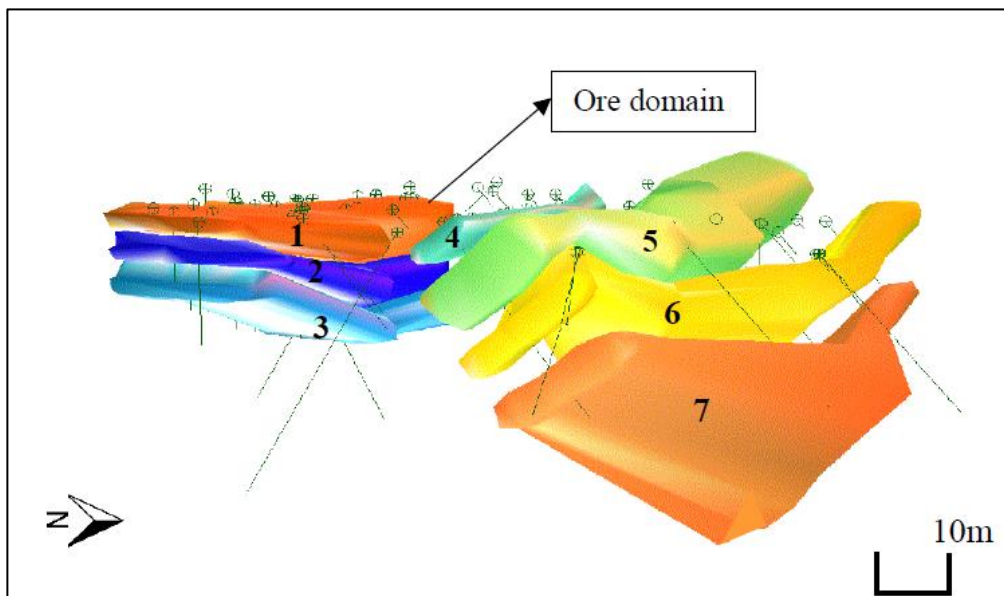


Figure 5.7: SURPAC - Long section view showing Rossing's seven mineralisation domains identified in this research.

Seven domains in Figure 5.7 represent different U₃O₈ mineralisation zones. Domain 1, 2, 3, and 7 are in a low-grade zone with their average mean composite grades below 160ppm whereas the rest of the domains (Domain 4, Domain 5 and Domain 6) are in a high-grade zones with average mean composite grades above 160ppm. Statistics of the mineralisation domains is given in detail in the next section.

5.4 Compositing

Compositing is a process of combining samples of various sample lengths to the same sample length (Sinclair & Blackwell, 2002; Abzalov, 2016). Grade values of composited samples are weighted by length or by both length and density in cases where density is variable (Darling & Noble, 2011; Rossi & Deutsch, 2013). If samples of varying lengths are used, then longer samples would contribute more weight than shorter ones creating biasedness. To avoid this biasedness, all sample length should be combined to a common length so that they have the same weight when used in calculating block values (Hustrulid & Kuchta, 2006). In addition, compositing reduces the variability of assay values resulting in much better and useful semi-variograms (Rossi & Deutsch, 2013).

Rossing sampling is done at variable sample intervals, from 0.1m to 83.82m. An analysis of the sample lengths was carried out to determine the most occurring sample length, to be used for compositing. Descriptive statistics (Table 5.1), a histogram (Figure 5.8) and frequency, and cumulative frequency distribution table (Table 5.2) of sample lengths were carried out to get a feel for the sample length data.

Table 5.1: Descriptive statistics of raw (un-composited) sample lengths (m)

Statistic	Value	Statistic	Value
Mean	1.97m	Kurtosis	150.18
Standard Error	0.07m	Skewness	10.60
Median	1.45m	Range	83.72m
Mode	1.50m	Minimum	0.1m
Standard Deviation	3.66m	Maximum	83.82m
Sample Variance	13.36m ²	Number of samples	2926

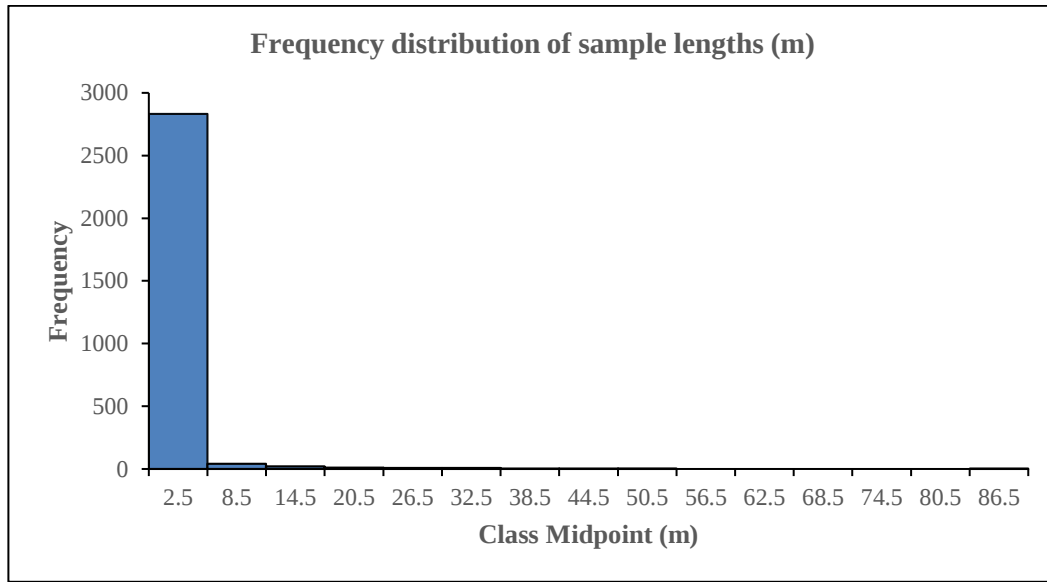


Figure 5.8: Histogram of sample lengths

Table 5.2: Frequency and cumulative frequency distribution table of sample lengths.

Lower Limit	Upper Limit	Class Midpoint	Class Frequency	Relative Frequency	Cumulative Frequency	Relative Cumulative %
0.00	5.50	2.50	2832	0.9678	2832	96.79
5.50	11.50	8.50	42	0.0143	2874	98.22
11.50	17.50	14.50	21	0.0071	2895	98.94
17.50	23.50	20.50	10	0.0034	2905	99.28
23.50	29.50	26.50	7	0.0023	2912	99.52
29.50	35.50	32.50	7	0.0023	2919	99.76
35.50	41.50	38.50	2	0.0006	2921	99.83
41.50	47.50	44.50	1	0.0003	2922	99.86
47.50	53.50	50.50	3	0.0010	2925	99.97
53.50	59.50	56.50	0	0.0000	2925	99.97
59.50	65.50	62.50	0	0.0000	2925	99.97
65.50	71.50	68.50	0	0.0000	2925	99.97
71.50	77.50	74.50	0	0.0000	2925	99.97
77.50	83.50	80.50	0	0.0000	2925	99.97
83.50	89.50	86.50	1	0.0003	2926	100.00

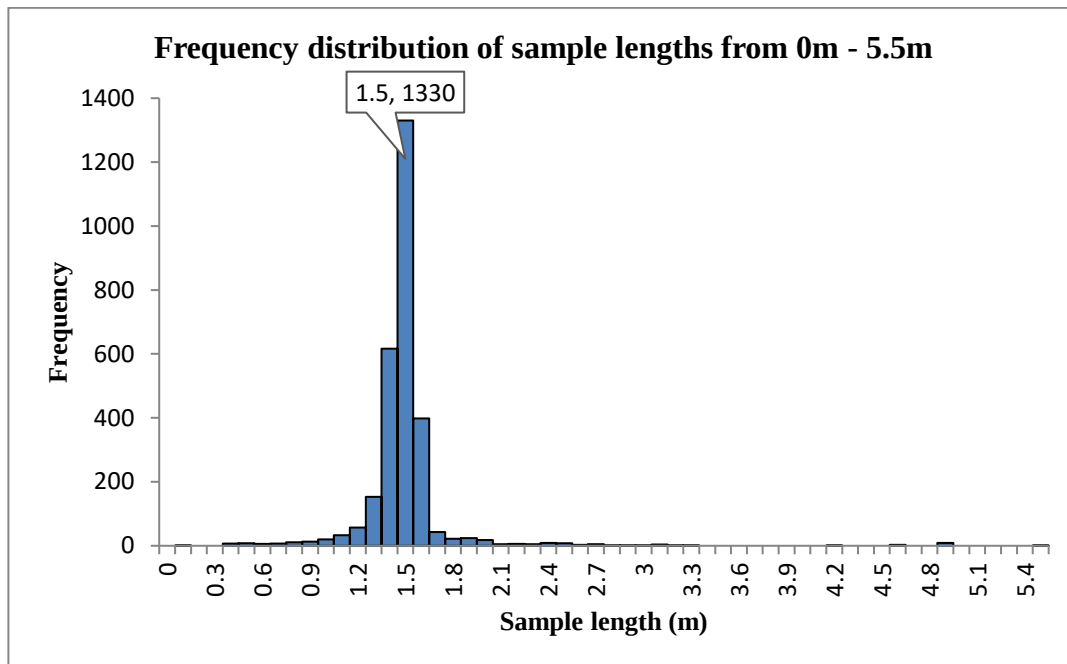


Figure 5.9: Histogram of sample lengths, zooming in on the 0.0m to 5.5m classes.

About 97 % of the sample lengths are between 0.0m to 5.5m (Table 5.2). To determine the most common sample length, the class interval between 0.0m to 5.5m was zoomed into (Figure 5.9). The sample length with the highest frequency is 1.5m confirming the modal sample length of 1.5m from Table 5.1. The appropriate compositing length for this dataset is thus 1.5m.

Compositing can be done either by bench compositing, downhole compositing or ore-zone compositing depending on the nature of the mineralisation, mining method and orientation of the drill holes (Darling & Noble, 2011). Bench and downhole compositing are the most common methods used in resource estimation (Rossi & Deutsch, 2013). In bench compositing, the composite sample length is equal to the bench height (Hustrulid & Kuchta, 2006). Downhole compositing is done by compositing samples of equal length from the collar of the drill hole or the top of the first assayed interval to the bottom of the drill hole (Dassault Systemes GEOVIA, 2014). The use of bench compositing for a database with a mixture of vertical, sub-vertical and inclined drill holes is being cautioned by Glacken & Snowden, (2001) and

Rossi & Deutsch, (2013). Using bench compositing for a database with a mixture of drill holes drilled at various directions could result in composites of various sample lengths (Glacken & Snowden, 2001).

The database used in this research study consists of a mixture of vertical and inclined drilled drill holes and a fixed length of 1.5m downhole compositing method was used.

SUPERVISOR lacks a compositing function, therefore compositing of drill hole samples was done in SURPAC. A minimum allowable percentage of composite sample length must be specified during compositing. Composite samples with lengths greater than or equal to the nominated minimum percentage are saved as “string 1”. The rest of the composite samples with lengths less than the nominated minimum percentage are saved as “string 2”. The minimum percentage of composite length is determined based on the correlation between sample lengths and assay values.

If there is a correlation between sample length and assay values, a higher minimum percentage of composite length is used whereas a lower minimum percentage of composite length is used when there is no correlation (Rossi & Deutsch, 2013).

A plot of the uranium sample grades versus sample lengths (Figure 5.10) shows no correlation thus 50% was used as the minimum percentage of composite length.

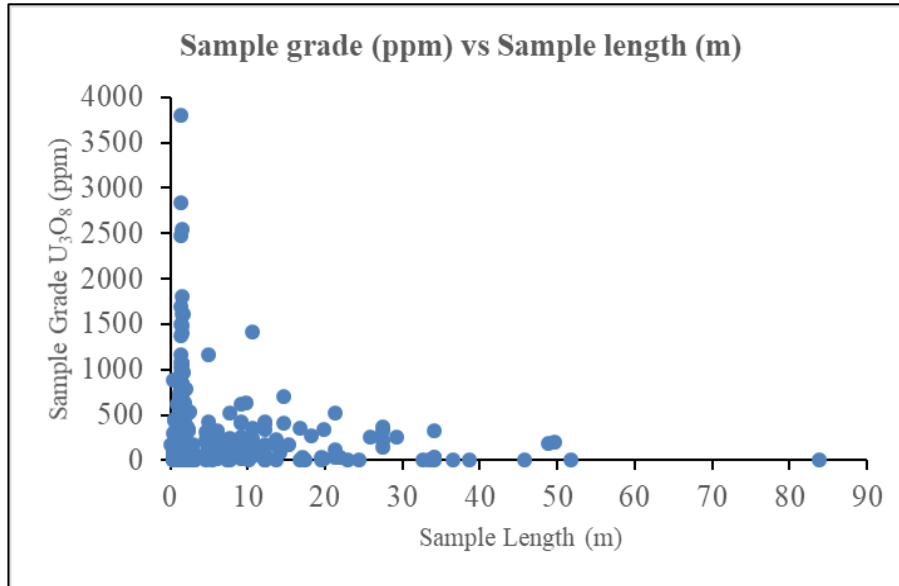


Figure 5.10: Rossing Uranium: Scatter plot of sample grades vs sample lengths prior to compositing.

About 98% of the composite samples are $\geq 50\%$ of the minimum sample length (0.75m). Due to limited number of samples, both “string 1” and “string 2” were used for geostatistical analysis. Besides, shorter length composites make up only 2.3% of the total composite samples and they are assumed to not have a significant influence on the overall uranium grades.

After compositing, the sample length should have reduced from the original variable lengths to the composite length, and the U_3O_8 variance should also reduce. The mean U_3O_8 grades of both raw (un-composited) and composited samples should be similar. This is to be expected because compositing should not change the mean grades, otherwise it would have created biasness. Therefore, it is good practice to compare the mean of the raw sample lengths with the mean of the composited sample lengths for the domains.

Comparison of the raw sample lengths and grade with composited sample lengths and grades is shown in Table 5.3.

Table 5.3: Statistics of individual U₃O₈ assay grades and 1.5m composited grade in the different domains

Domain	Raw Data				Composite Data			
	Number of samples	Mean sample length (m)	U ₃ O ₈ (ppm) mean	U ₃ O ₈ (ppm) ² variance	Number of composite	Mean sample length (m)	U ₃ O ₈ (ppm) mean	U ₃ O ₈ (ppm) ² variance
1	263	2.34	128.23	18224.16	302	1.50	117.93	14978.15
2	177	2.51	100.18	14623.98	170	1.49	102.12	13287.74
3	144	2.36	165.54	114599.00	137	1.49	167.73	68194.30
4	106	2.64	282.86	110161.30	166	1.48	244.32	53533.62
5	533	2.48	306.22	50947.19	796	1.49	272.46	49728.47
6	264	2.33	265.22	72473.27	383	1.50	235.25	49302.01
7	66	4.02	95.65	20220.66	150	1.49	140.19	12673.94

As expected, for all domains, the variance of composited data has decreased because of change of support (Table 5.3). Differences in the mean of the assay grades for domain 1, 4, 5, 6 and 7 (Table 5.3) are observed, which are substantial but not alarming when the detailed compositing methodology is investigated.

The differences are due to the compositing process itself and the way in which composites in partial composite lengths are treated in SURPAC. In the case of the raw (un-composited) data, every sample is considered to have exactly the same weight, and its length is effectively ignored in the calculation of the mean grade value. For example, there is no difference between a sample of 1.5m at 400ppm and a sample of 15m at 300ppm, the naïve average grade of the two samples is 350ppm (it is called a naïve average because all the samples are considered to have the same impact as in Table 5.4). Whereas when these two samples are composited at 1.5m, there will be one 1.5m sample for the first case with a value of 400ppm and ten 1.5m samples with the same grade value of 300ppm for the second case and the weighted average of the composites would be $(400+10 \times 300)/11 = (400+3000)/11 = (3400)/11 = 309\text{ppm}$.

An example of how compositing is treated in SURPAC is shown below.

PA 61 is one of the drill holes that intersected domain 6. Sample intervals and their respective uranium assay values for drill hole PA 61, taken from the assay file, are shown in Table 5.4 below.

Table 5.4: Domain 6 un-composited assay values for drill hole PA61

Sample from (m)	Sample to (m)	Sample length (m)	U ₃ O ₈ (ppm)
77.72	83.82	6.10	181.00
83.82	91.44	7.62	0.00
91.44	108.20	16.76	9.00
Naïve Average			63.33

Domain 6 intersects drill hole PA 61 at a depth of 81.90 to 107.24m. The starting depth for compositing is 81.90m. Composite 1 (Table 5.5) is assigned an assay value of 181ppm because its sample interval falls between 77.72 and 83.82m (Table 5.4). The sample interval for composite 2 begins at 83.4m (81.90m + 1.5m) and ends at 84.9m (83.4m + 1.5m), however, the first 0.42m (83.4 to 83.82m) of this interval falls between 77.72 to 83.82m hence this portion is assigned an assay value of 181ppm. The second portion of composite 2's sample interval is 1.08m (83.82 to 84.9m) since it falls between 91.44 to 108.2m, it is assigned a value of 0ppm. The sum of the weighted average of these two portions is then the final assay value for composite 2 (Table 5.5). The rest of the composites were assigned an assay value for a sample interval in which they fall.

The weighted average of the composited samples is 19.26ppm (Table 5.5), about 70% lower than that of un-composited samples which is 63.33ppm (Table 5.4). This is caused by four composite samples (composite 3, 4, 5 and 6) which are assigned a composite value of zero and composites 8 to 17 which all have the same value of 9ppm. The composite average grade is based on a weighted average and the un-composited grades are calculated from the sample grades only ignoring their individual lengths. This explains why in domains 1, 4, 5 and 6 the composited mean is lower than that of their corresponding raw means.

Table 5.5: Domain 6 composited assay values for drill hole PA 61

Depth from (m)	Depth to (m)	Sample length (m)	Sample value x Interval proportion (ppm)	Weighted value (ppm)	Composite #
81.90	83.40	1.50	181 x (1.5/1.5)	181.00	1
83.40	83.82	0.42	181 x (0.42/1.5)	50.68	2
83.82	84.90	1.08	0 x (1.08/1.5)		
84.90	86.40	1.50	0 x (1.5/1.5)	0.00	3
86.40	87.90	1.50	0 x (1.5/1.5)	0.00	4
87.90	89.40	1.50	0 x (1.5/1.5)	0.00	5
89.40	90.90	1.50	0 x (1.5/1.5)	0.00	6
90.90	91.44	0.54	0 x (0.54/1.5)	5.76	7
91.44	92.40	0.96	9 x (0.96/1.5)		
92.40	93.90	1.50	9 x (1.5/1.5)	9.00	8
93.90	95.40	1.50	9 x (1.5/1.5)	9.00	9
95.40	96.90	1.50	9 x (1.5/1.5)	9.00	10
96.90	98.40	1.50	9 x (1.5/1.5)	9.00	11
98.40	99.90	1.50	9 x (1.5/1.5)	9.00	12
99.90	101.40	1.50	9 x (1.5/1.5)	9.00	13
101.40	102.90	1.50	9 x (1.5/1.5)	9.00	14
102.90	104.40	1.50	9 x (1.5/1.5)	9.00	15
104.40	105.90	1.50	9 x (1.5/1.5)	9.00	16
105.90	107.40	1.50	9 x (1.5/1.5)	9.00	17
Weighted Average				19.26	

The composited mean grade of domain 7 is higher than that of the un-composited samples (Table 5.3). The sample lengths and their corresponding assay values for individual un-composited samples for drill hole PA66 are presented in Table 5.6.

Table 5.6: Domain 7 un-composited assay values for drill hole PA 66

Sample from (m)	Sample to (m)	Sample length (m)	U ₃ O ₈ (ppm)
167.63	251.45	83.82	5
251.45	259.07	7.62	109
259.07	265.16	6.09	327
Naïve Average			147

Domain 7 intersects drill hole PA66 at a depth of 248.02 m to 265.16 m. The same compositing method described for drill hole PA 61 is used, however, drill hole PA66 is in a high-grade zone. Individual un-composited samples in domain 7 have high uranium values therefore, the corresponding 1.5 m composite samples are also assigned

high uranium values. One un-composited sample (7.62m) with a uranium grade of 109ppm (Table 5.6) after compositing resulted in four composite samples each with a grade of 109ppm (Table 5.7). In addition, another un-composited sample (6.09m) with a uranium grade of 327ppm (Table 5.6) resulted in four composite samples each with a grade of 327ppm (Table 5.7), explaining the higher composited mean in domain 7.

Table 5.7: Domain 7 composited assay values for drill hole PA66

Depth from (m)	Depth to (m)	Sample length (m)	Sample value x interval proportion (ppm)	Weighted value (ppm)	Composite #
248.02	249.52	1.50	5 x (1.5/1.5)	5.00	1
249.52	251.02	1.50	5 x (1.5/1.5)	5.00	2
251.02	251.45	0.43	5 x (0.43/1.5)	79.19	3
251.45	252.52	1.07	109 x (1.07/1.5)		
252.52	254.02	1.50	109 x (1.5/1.5)	109.00	4
254.02	255.52	1.50	109 x (1.5/1.5)	109.00	5
255.52	257.02	1.50	109 x (1.5/1.5)	109.00	6
257.02	258.52	1.50	109 x (1.5/1.5)	109.00	7
258.52	259.07	0.55	109 x (0.55/1.5)	247.07	8
259.07	260.02	0.95	327 x (0.95/1.5)		
260.02	261.52	1.50	327 x (1.5/1.5)	327.00	9
261.52	263.02	1.50	327 x (1.5/1.5)	327.00	10
263.02	264.52	1.50	327 x (1.5/1.5)	327.00	11
264.52	265.16	0.64	327 x (1.5/1.5)	327.00	12
Weighted Average				173.35	

5.5 Bimodality and Outliers

The U₃O₈ grade distributions of all composite domains were checked for bimodality and the presence of and outliers. A distribution is said to be bimodal if it has two distinct peaks. Bimodality is often caused by two different domains being represented as a single domain (Dassault Systemes GEOVIA, 2014).

The U₃O₈ (ppm) distributions of composite data for domain 1 and 2 (Figure 5.11), domain 3 (on the left of Figure 5.12), and domain 6 (on the right of Figure 5.13) display

a single peak (unimodal distribution). Bimodal distribution in domain 4 (on the right of Figure 5.12), domain 5 (on the left of Figure 5.13) and domain 7 (Figure 5.14).

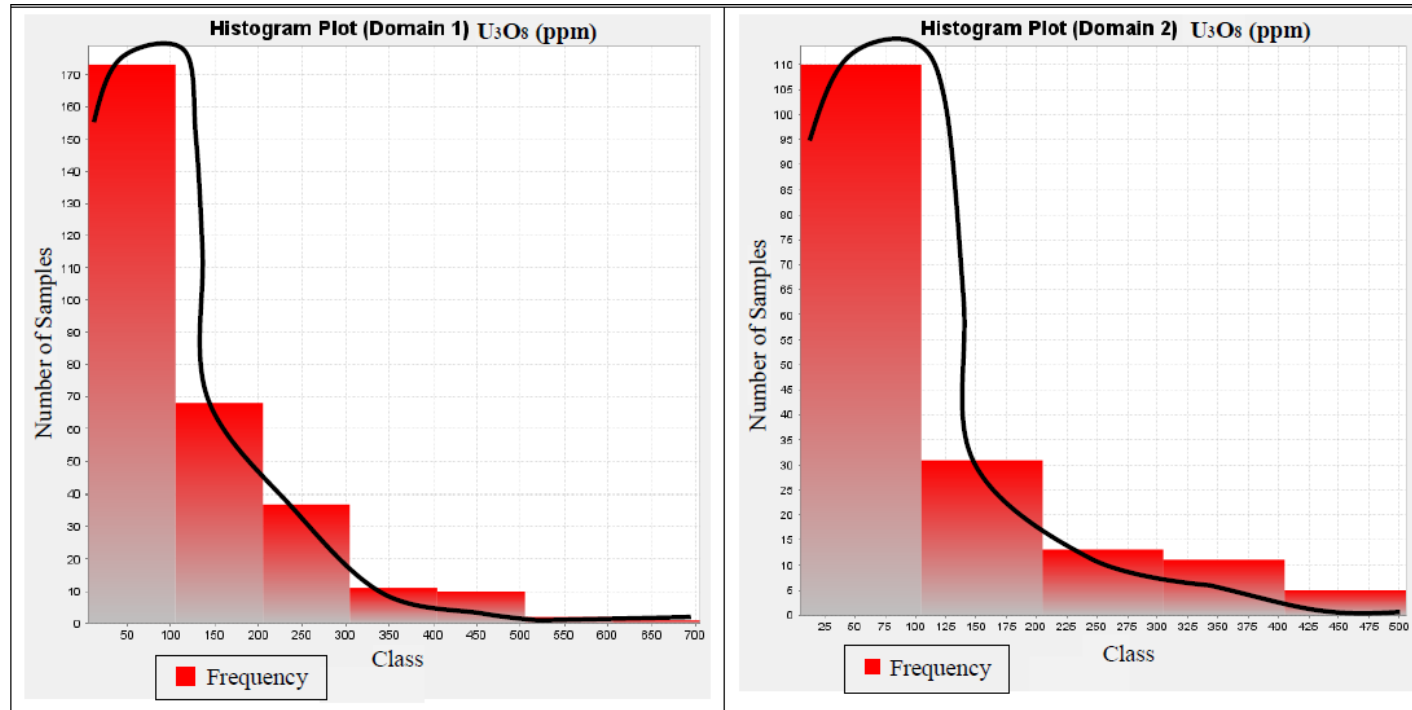


Figure 5.11: Histograms of the U_3O_8 (ppm) composites Domain 1 (left) and Domain 2 (right) showing unimodal distributions.

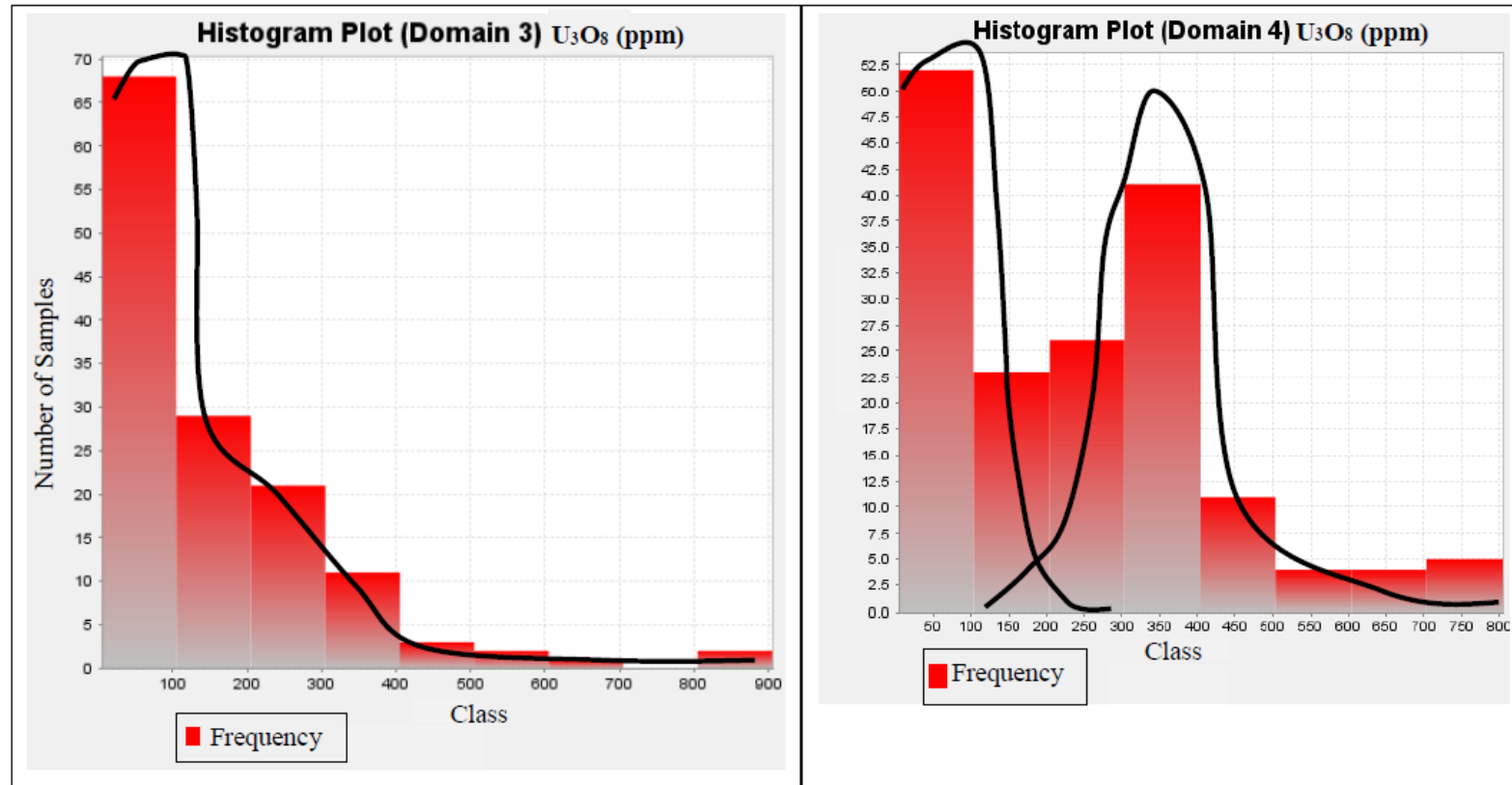


Figure 5.12: Histograms of the U₃O₈ (ppm) composites - Domain 3 (left) has a unimodal distribution and Domain 4 (right) displays a bimodal distribution.

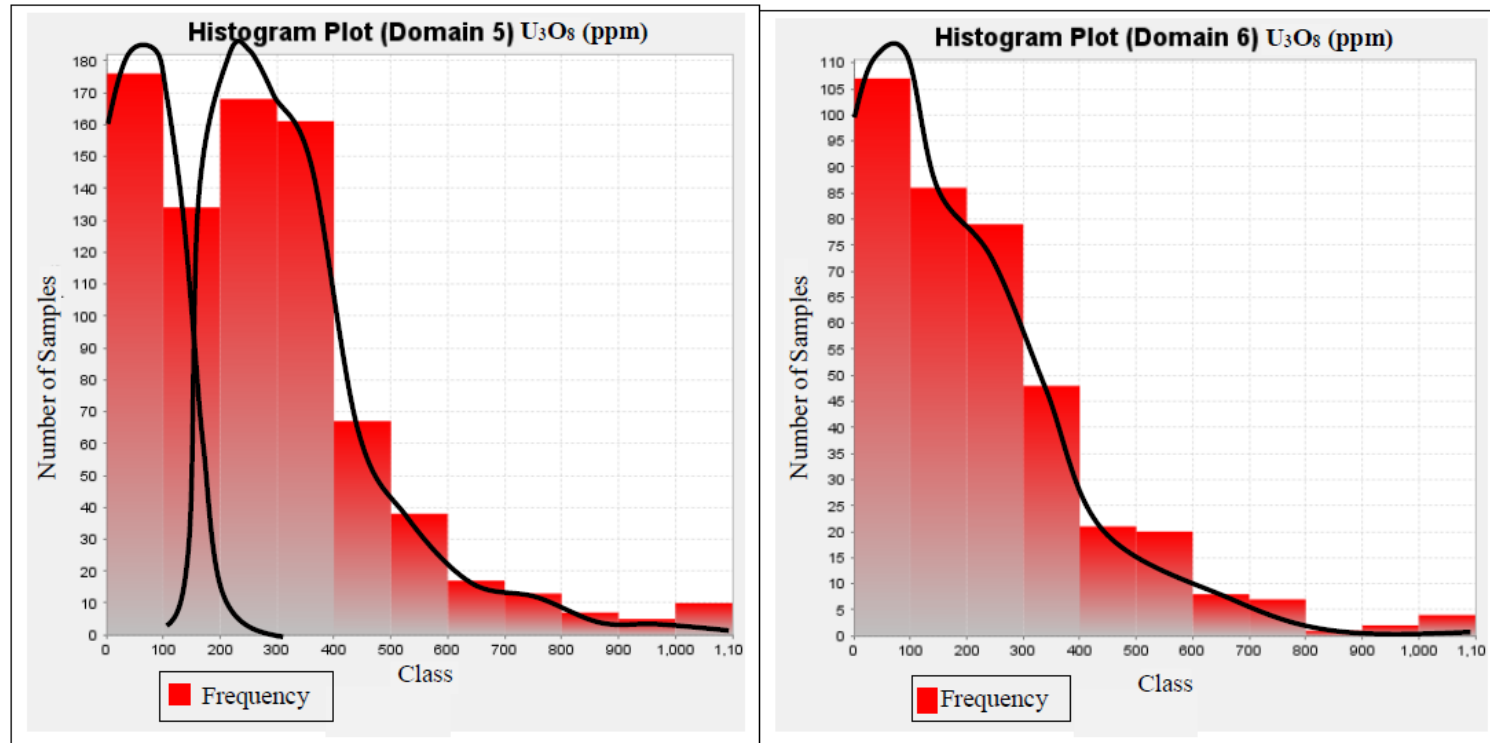


Figure 5.13: Histograms of the composites U₃O₈ (ppm) for Domain 5 (left) has a bimodal distribution (left) and Domain 6 (right) a unimodal distribution.

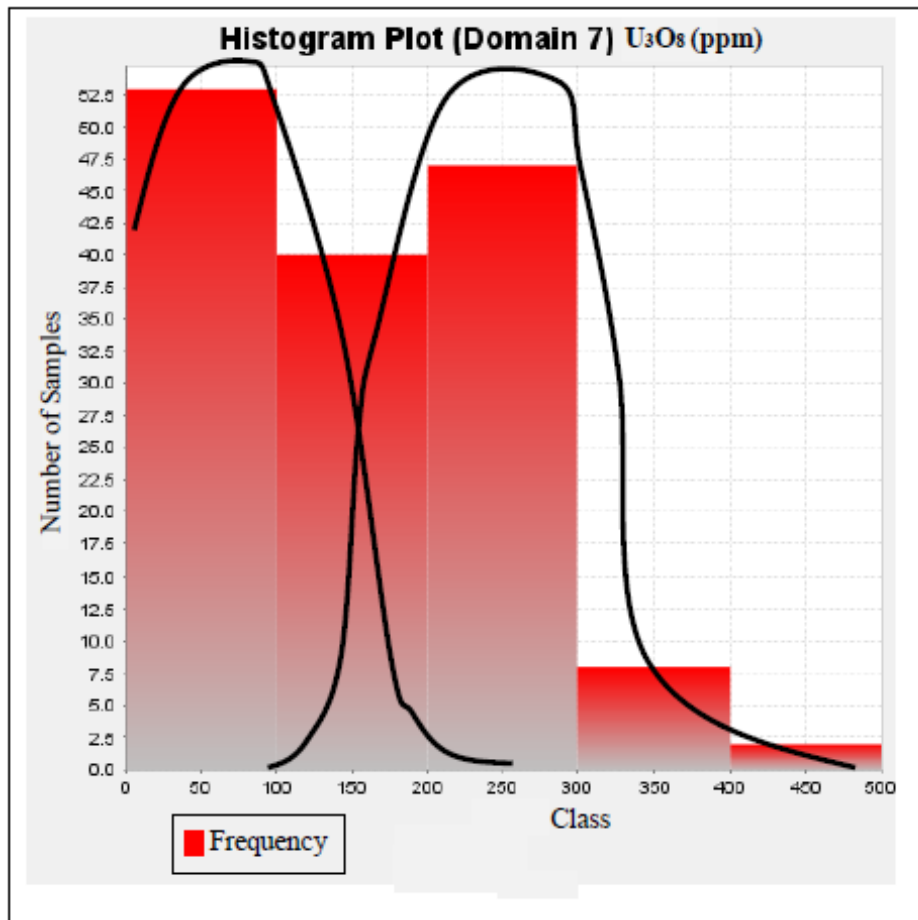


Figure 5.14: Domain 7: Histogram of the U₃O₈ (ppm) composites depicting bimodality

The immediate interpretation that bimodality is caused by mixing of two different distributions is being cautioned by Clark & Garnett, (1974), thus the domains with bimodal distributions were scrutinised. It was discovered that bimodality was caused by the way compositing is done in SURPAC. For example domain 4 grade distribution shown in Figure 5.12: The second modal interval is in the 300 – 400ppm class interval. One of the drill holes intersecting domain 4 has an uncomposited raw sample length of 34.13m at 330ppm. Compositing was done at 1.5m and a raw sample length of 34.13m resulted in about 22 samples ($34.13\text{m}/1.5\text{m} \approx 22$) at 1.5m with the same uranium grade value of 330ppm. As a consequence pronounced peak in the class interval between 300 and 400ppm is observed.

An outlier is an observation which is far from the rest of the data values. Outliers are known for introducing variability in the mean estimate and cause an increase in the variance and covariance (Sinclair & Blackwell, 2002). In addition, high-grade outliers can lead to overestimation of the estimated grades hence they need to be capped.

In this research, outliers were determined by reviewing U_3O_8 (ppm) histograms and their respective cumulative frequency plots of composite data. Sample points which were found to be far from the rest of the data points were treated as outliers. Such points were identified by a significant break at the top end of the cumulative distribution of U_3O_8 (ppm) of the composite data. Domain 1 and 2 do not have outliers hence capping was not done on them. Domain 3 U_3O_8 (ppm) cumulative frequency plot of 1.5m composites (Figure 5.15) shows a break in distribution at 855ppm, hence, 855ppm was used as the capping value. The capping values for other domains were determined in the same manner. U_3O_8 (ppm) histograms and cumulative frequency plot of 1.5m composites of domain 4, 5, 6, and 7 are shown in Appendix A.

Before capping, the cumulative frequency plot increases and then flattens off. After capping, instead of the cumulative frequency plot flattening off, a sudden increase in the cumulative frequency plot of the cap value class interval is observed. The sudden increase in the cumulative frequency plot is as a result of an increase in the frequency of the cap value class interval due to capping (Figure 5.15).

Capping was done in this research to remove few high-grade values from the domains. If high grade values are not capped, they can cause jagged variograms that are difficult to model. Modelling of jagged variograms can lead to unrealistic block estimates. Capping can also lead to underestimation of block estimates if it is done severely. In this research, less than 2% of the U_3O_8 composites were capped and for the combined domain the capped percentage was 0.4%.

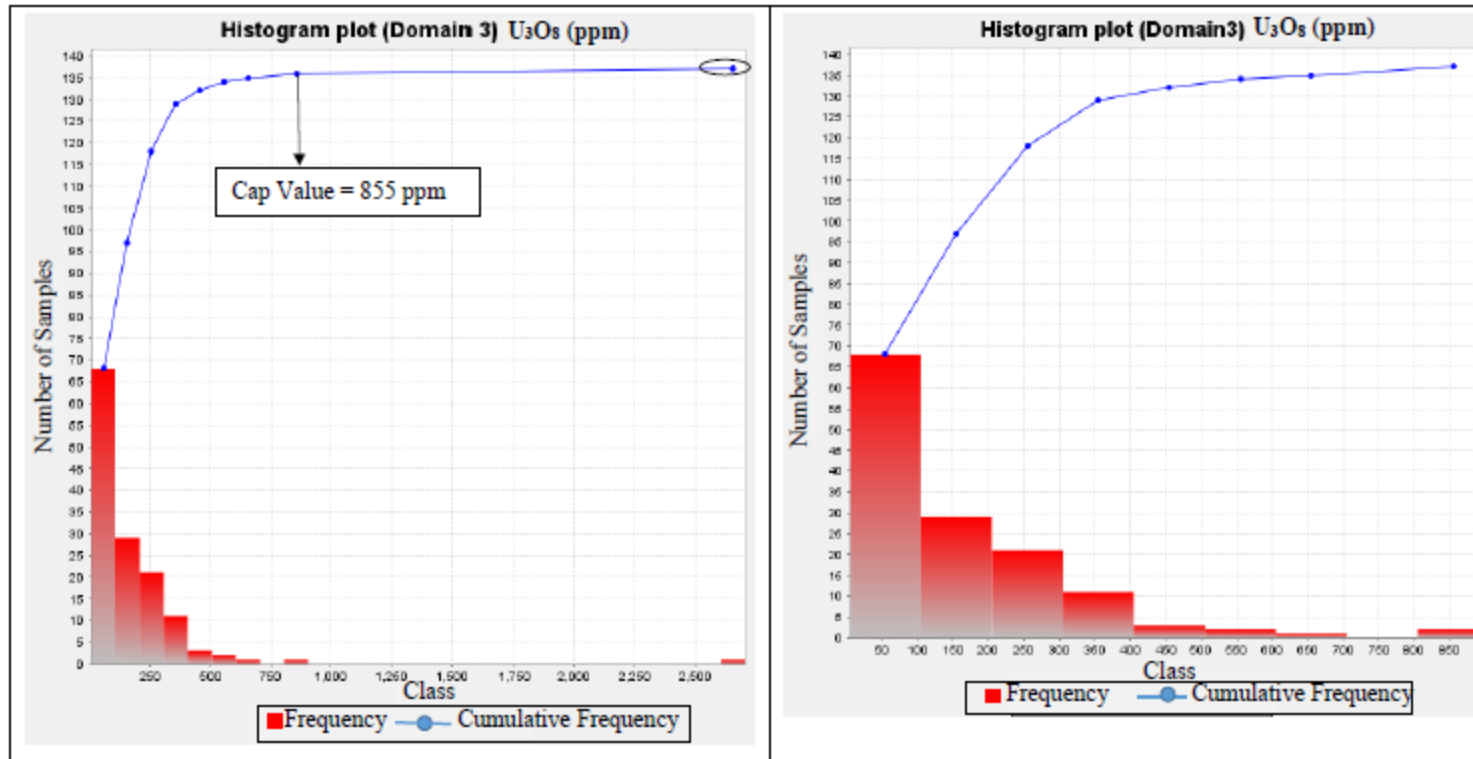


Figure 5.15: Domain 3 U₃O₈ (ppm) histogram, and cumulative frequency plot of 1.5m composites before capping (left) with an outlier (circled) and after capping (right).

The statistics of each domain before and after capping is shown in Table 5.8.

Table 5.8: Domain U₃O₈ (ppm) Statistics before and after capping

Domains statistics before capping							Domains statistics after capping			
Domain	Count	Min (ppm)	Max (ppm)	Mean (ppm)	Variance (ppm) ²	CV	Mean (ppm)	CV	Cap Value (ppm)	% Cap
1	302	5	640.00	117.93	14928.55	1.04	117.93	1.04	None	-
2	170	5	467.74	102.12	13209.58	1.13	102.12	1.13	None	-
3	137	5	2684.68	167.73	67696.53	1.55	154.38	1.02	855	0.7
4	166	4	1816.12	244.32	53211.12	0.94	233.52	0.79	754	1.8
5	796	0	1411.00	272.46	49666.00	0.82	269.34	0.78	1050	0.9
6	383	0	1627.23	235.25	49173.29	0.94	231.68	0.88	1050	0.5
7	150	0	773.36	140.19	12589.44	0.80	138.02	0.74	450	0.7
All	2104	0	2684.68	211.27	44759.22	1.00	207.38	0.92	1050	0.2

Capping resulted in a slight reduction in the mean of the affected domains.

Domain 3 has the highest difference in the mean confirming that the one very high-grade outlier of 2684.68ppm had a significant influence. The mean, variance and CV all reduced after capping was introduced. Most of the domains have coefficients of variations of less than one after grade capping which shows that OK is a suitable estimation method for the dataset.

The capped U₃O₈ (ppm) histogram of 1.5m composites for all combined domains shows a positively skewed distribution on the left of (Figure 5.16). Positively skewed distribution indicates many low uranium grades and few high-grade values. U₃O₈ (ppm) histograms of composite data for each domain (Appendix B) are also positively skewed.

If the cumulative log probability plot is not a straight line then it means that the distribution is not normal (Babish, 2000). A cumulative log-probability plot that curves upwards signifies a positively skewed distribution (Rundel, 2012).

The cumulative U_3O_8 (ppm) log-probability plot for all domains on the right of (Figure 5.16) and those of individual domains (Appendix B) all curve upwards confirming the positively skewed nature of the U_3O_8 distributions even after capping the grades in the Case Study.

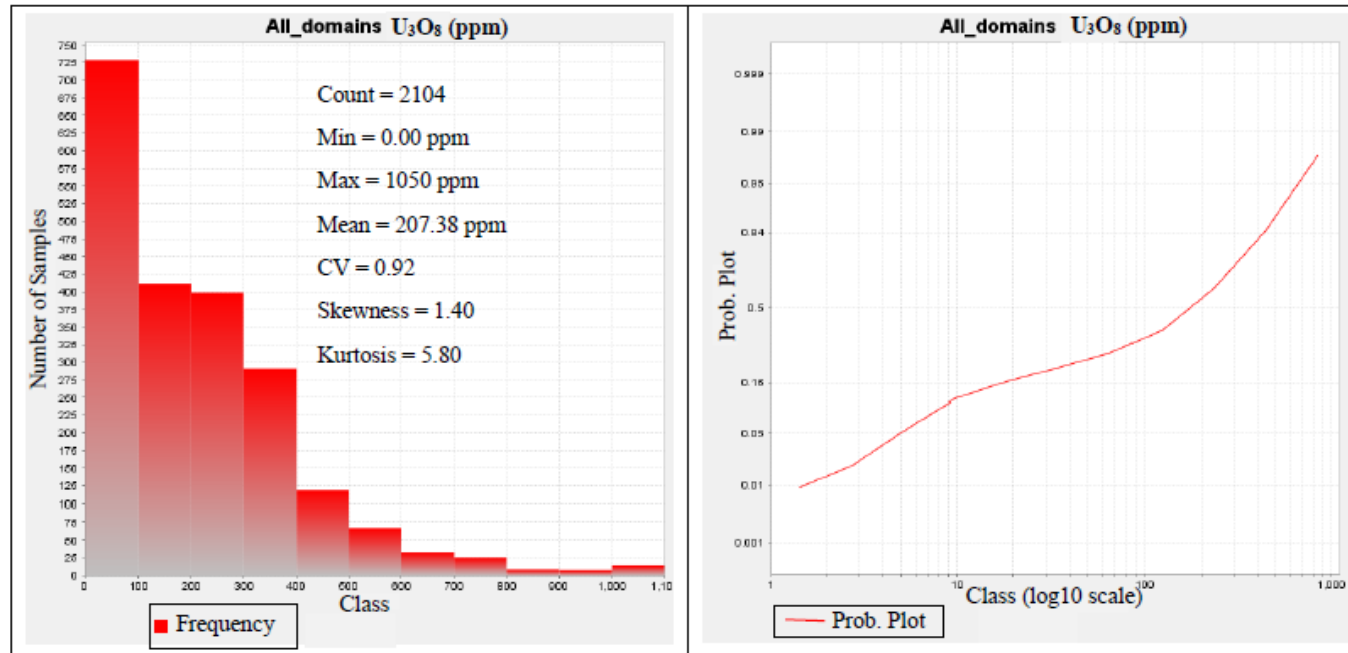


Figure 5.16: U_3O_8 (ppm) histogram (left) and cumulative log-probability plot (right) of 1.5m composites for all domains combined after capping

5.6 Declustering

Resource estimation relies heavily on the use of samples to interpolate assay values of the variable of interest at un-sampled locations. To get estimates which are close to the true values, collected samples should be unbiased in order for the samples to represent the population (Sarma, 2009).

One of the assumptions made when applying statistical estimation in resource evaluation is that samples are taken randomly and independently with respect to their values and not necessarily their position (Dohm, 2015b).

Sampling is done either by probability sampling or by non-probability sampling. Samples are randomly sampled during probability sampling whereas, in non-probability sampling, samples are preferentially sampled (Sarma, 2009). In most cases, samples are collected by non-probability sampling (Goovaerts, 1997).

The Rossing samples used in this research study were collected by preferential sampling. Preferential sampling is done because of the geologic intuition, previous samples indicating areas of interest such as high-grade areas (Pyrz & Deutsch, 2003), and conditions of accessibility (Goovaerts, 1997).

Samples collected by probability (randomly) sampling represent their population, however, non-probability collected samples might not be representative of their population due to clustering. It is for this reason that non-probability samples should undergo a process called declustering. Declustering is a process of removing the clustering effect by assigning weights to all points in the dataset (Dassault Systemes GEOVIA, 2015).

During declustering, samples in densely sampled areas are assigned less weights than those in sparsely sampled areas (Goovaerts, 1997).

Polygonal and cell declustering method are the two common methods of declustering. In polygonal method, a polygon of influence surrounding each sample is drawn by connecting all the midpoints line between each sample and its neighbour (Figure 5.17) (Isaaks & Srivastava, 1989).

The weight (w_j) of each sample (j) is calculated by dividing the area of its polygon ($area_j$) by the sum of the total area of all polygons in the area of interest.

$$w_j = \frac{area_j}{\sum_{j=1}^n area_j} \times n$$

Where n = number of data (Pyrz & Deutsch, 2003)

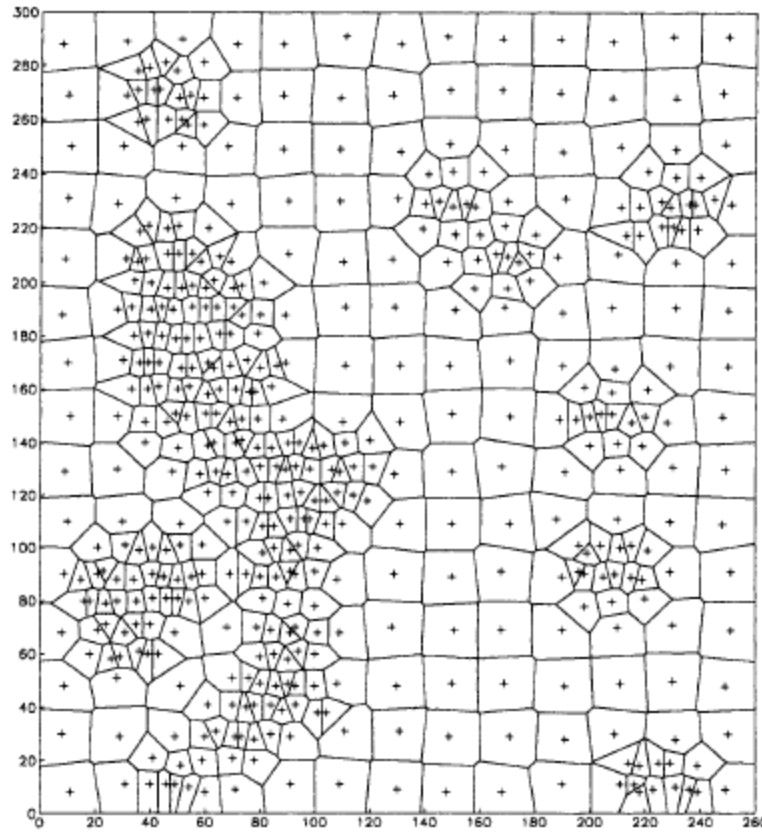


Figure 5.17: Samples and their respective polygons of influence

(Source: Isaaks & Srivastava, 1989)

Clustered samples tend to have very small polygons compared to the sparse samples and as a result, they are assigned lesser weights. Sparsely samples get more weights because of their larger polygon areas.

For cell declustering method, the area of interest is divided into rectangular blocks (cells). The weight (w_j) of each sample (j) is estimated based on the cell in which it is located. The sum of the weight of samples in a cell is normalised to 1 hence samples

are assigned weights as a ratio of 1 to the total number of samples in the cell (Abzalov, 2016).

$$W_j = \frac{\frac{1}{n_i}}{\text{number of cells with data}} \times n$$

Both SUPERVISOR and SURPAC use the cell declustering method.

A plot of various cell sizes versus their respective declustered means is created to determine the cell size that would be used for declustering. The cell size with the lowest declustered mean (when preferential sampling was done in high-grade areas) or highest declustered mean (low-grade preferential sampling) is chosen (Isaaks & Srivastava, 1989; Rossi & Deutsch, 2013).

The tricky part of the cell declustering method is choosing the correct cell size for declustering (Sinclair & Blackwell, 2002).

If the chosen cell size is too small, then each sample will fall in one cell hence all samples will be given equal weights of 1. Similarly, if the chosen cell size is large, all samples will fall into one cell and they are all given equal weights (Isaaks & Srivastava, 1989; Rossi & Deutsch, 2013).

If sampling is done on a regular grid, the correct cell size for declustering is the cell size with the same size as the sample spacing (Isaaks & Srivastava, 1989). In a case where sampling was done on a non-regular grid, the cell size approximately equal to the spacing of the sparse drilled samples is chosen (Rossi & Deutsch, 2013). In a case where none of the previous cases exist the correct cell size would be the one with the minimum/maximum declustered mean (Sinclair & Blackwell, 2002).

In the research study area, drilling was done on an irregular grid and there is no fixed spacing for the sparse drilled samples. Samples were preferentially drilled in high-grade areas (Figure 5.18).

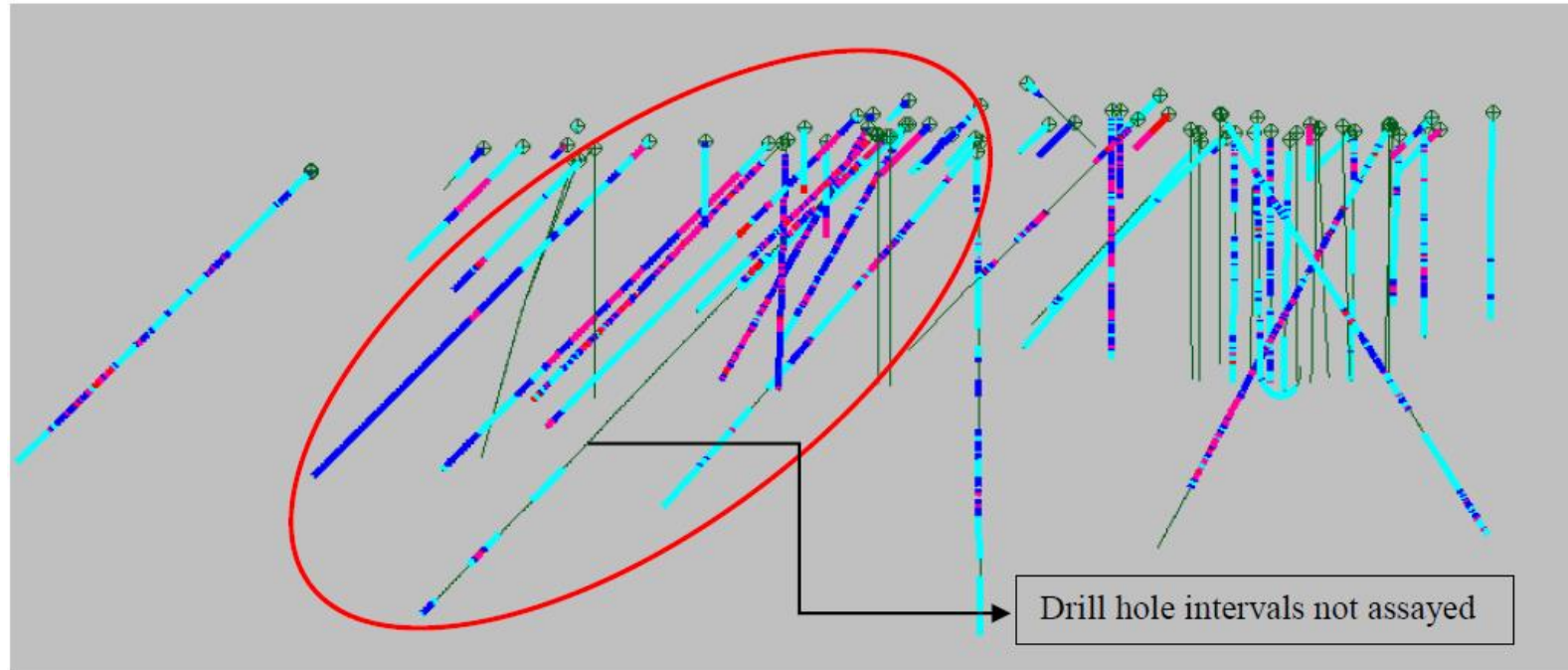


Figure 5.18: Section view facing North showing drill holes displayed by U_3O_8 grades with a high-grade zone circled. Colour coding: 0 - 100ppm = cyan; 100 - 300ppm = blue; 300 - 700ppm = magenta; 700 - 10 000ppm = red

The sample cell size to be used for declustering is the cell size with a minimum declustered mean. From the plot of declustered mean versus cell volume (Figure 5.19), a cell volume of 160 000m³ which corresponds to a 100m (X) x 40m (Y) x 40m (Z) has the minimum declustered mean of 171ppm and thus used for declustering the composite data.

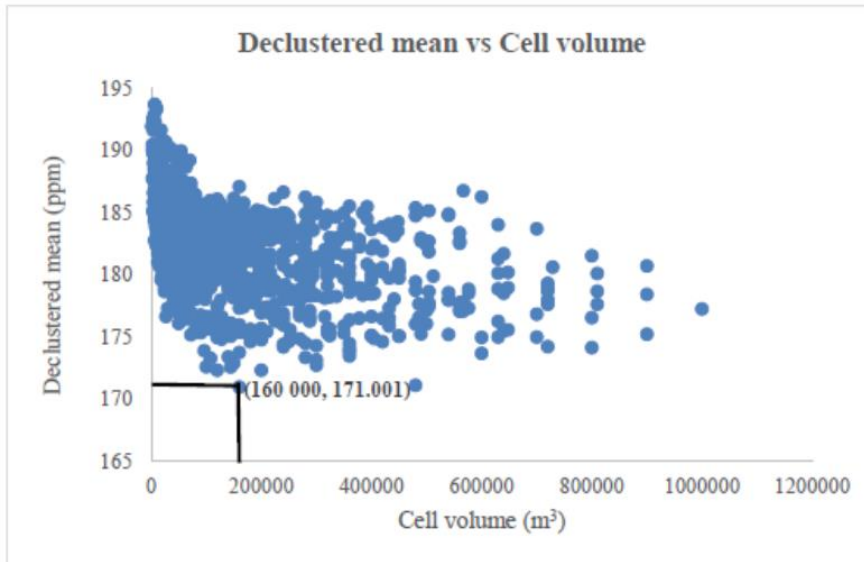


Figure 5.19: Declustered mean versus cell volume highlighting the cell volume with the lowest declustered mean.

The results of the declustering are shown in Table 5.9. The degree of clustering is low since the difference in mean for all the domains are within 2% (Table 5.9). Rossing drill holes are sparsely drilled and that explains why the degree of clustering is low.

Table 5.9: Statistics of the domains before and after clustering

Domains statistics before declustering				Domains statistics after declustering			
Domain	Count	Mean (ppm)	CV	Count	Mean (ppm)	CV	Mean Diff (%)
1	302	117.93	1.04	287	120.28	1.03	-2
2	170	102.12	1.13	166	103.73	1.12	-2
3	137	154.38	1.02	133	155.62	1.01	-1
4	166	233.52	0.79	159	238.17	0.77	-2
5	796	269.34	0.78	786	270.76	0.77	-1
6	383	231.68	0.88	378	230.76	0.87	0
7	150	138.02	0.74	147	137.01	0.74	1
All	2104	207.38	0.92	2056	209.38	0.91	-1

5.7 Variography

A semi-variogram is a model that measures the spatial correlation between sample values and this correlation is a function of the distance between the samples (Rendu, 1978). A typical single structured semi-variogram comprises a nugget effect (C_0), a sill (C_1), and a range (a) as shown in Figure 5.20.

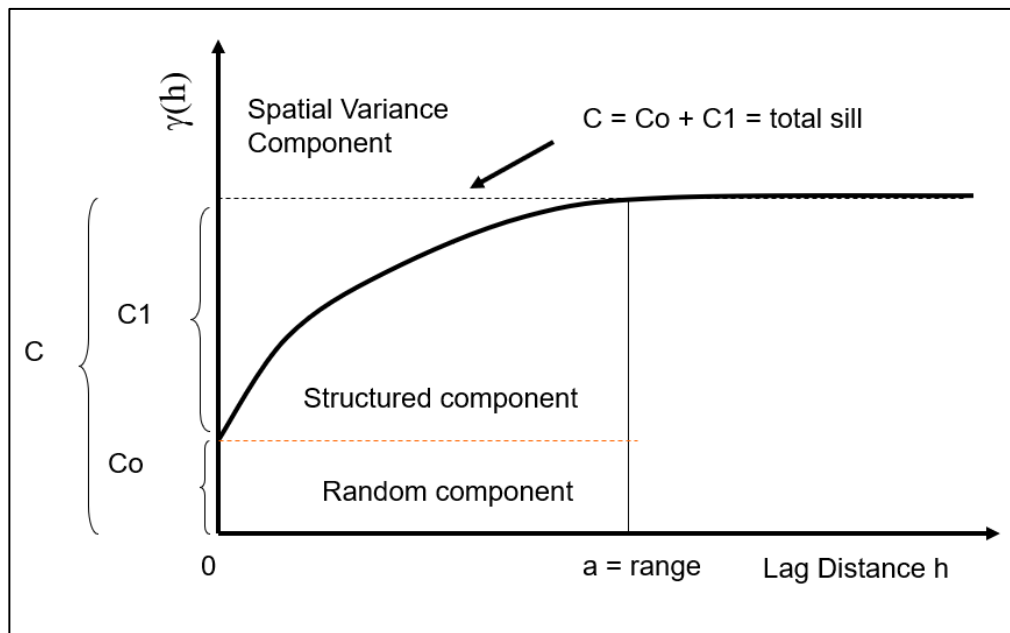


Figure 5.20: Typical features of a single structured semi-variogram

(Source: Dohm, 2015a)

A semi-variogram always starts at zero (on the x-axis), it then rises with the increase in the distance up to a certain level where it flattens off. The vertical axis of the semi-variogram represents the variability of samples whereas the horizontal axis is the distance between samples. The distance (on the x-axis) between zero and where the semi-variogram model flattens out is called the range (a). Samples that are within the range are spatially correlated whereas those beyond the range are not. The value (on the y-axis) at which the semi-variogram model flattens off is called the sill (C_1). Nugget effect, C_0 , is the random component which represents the variability between samples and all the measurement errors (Dohm, 2015a).

In this research, the semi-variogram analysis was done for each domain on composited capped grades. Due to few samples in each domain, it was impossible to derive at meaningful semi-variograms per domain. The few samples in domains is attributed to the fact that Rossing mineralisation is hosted by alaskites which vary in thickness from a few centimetre to hundreds of meters thick. In addition, the relatively sparse drilling coverage in the Case Study area also contributed to smaller number of samples per domain.

All domains were plotted on a box and whisker plot to see if any of the domains can be combined. A box and whisker plot is a plot that provides visual representation of the statistics and dispersion of a data set (Supervisor, 2017). From the box and whisker plot, one can determine which of the domains have positively skewed shapes. The grade distribution in all the domains have different statistics and dispersions (Figure 5.21).

The technical reports of Husab deposit also revealed an issue of poor variography due to too few samples in the domains, which was resolved by combining samples in all domains for calculating and modelling the semi-variograms.

In an attempt to improve on the variography it was decided to pool and carry out the variography analysis on the combined domain data. Semi-variogram analyse for this research was done in SUPERVISOR and the resulting models were applied in

SURPAC, comparing variography in the packages did not form part of the scope of this research.

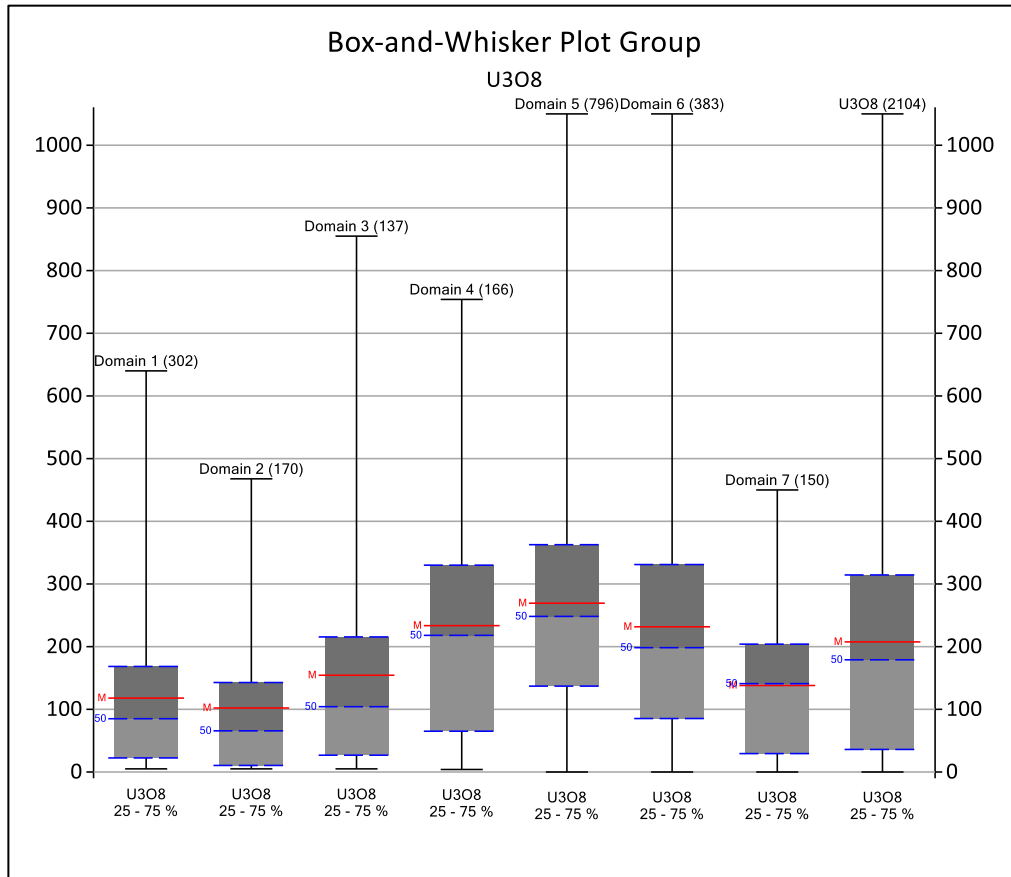


Figure 5.21: Box and whisker plots of each domain and the combined domain.

Due to the sparsity of data and the positive skewed nature of uranium grade distributions U_3O_8 (ppm) composites resulted in noisy semi-variograms. The composite uranium grades were transformed to a normal distribution with a mean of zero and standard deviation of one in using the Normal score option in SUPERVISOR to obtain improved semi-variograms, these transformed values are unitless.

The SUPERVISOR Normal score variogram fans were used to determine the direction of maximum continuity in the along strike (horizontal) direction (Figure 5.22), across strike direction (Figure 5.23) and in the dip-plane direction (Figure 5.24). The direction of maximum continuity in each direction is indicated by a purple arrow.

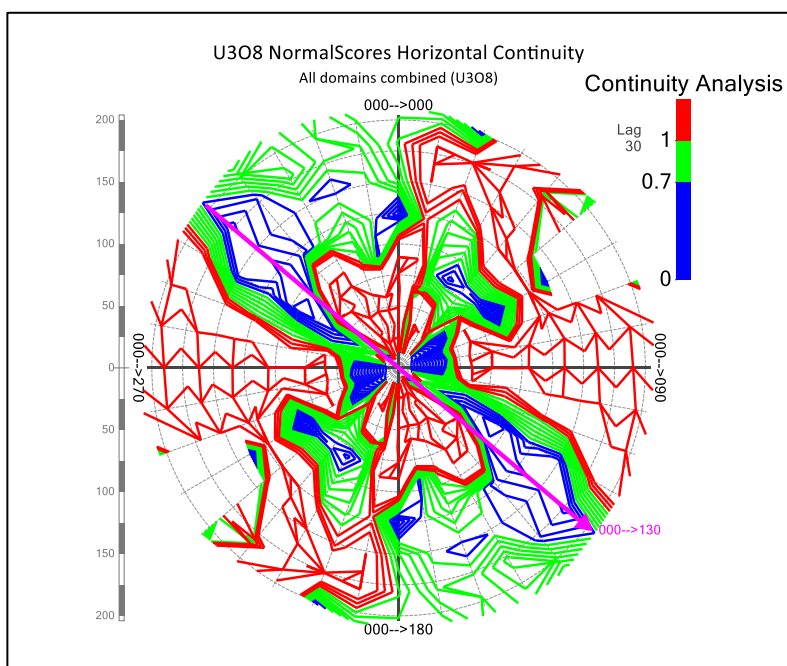


Figure 5.22: Variogram fans in the horizontal direction the purple arrow indicates the direction of maximum continuity

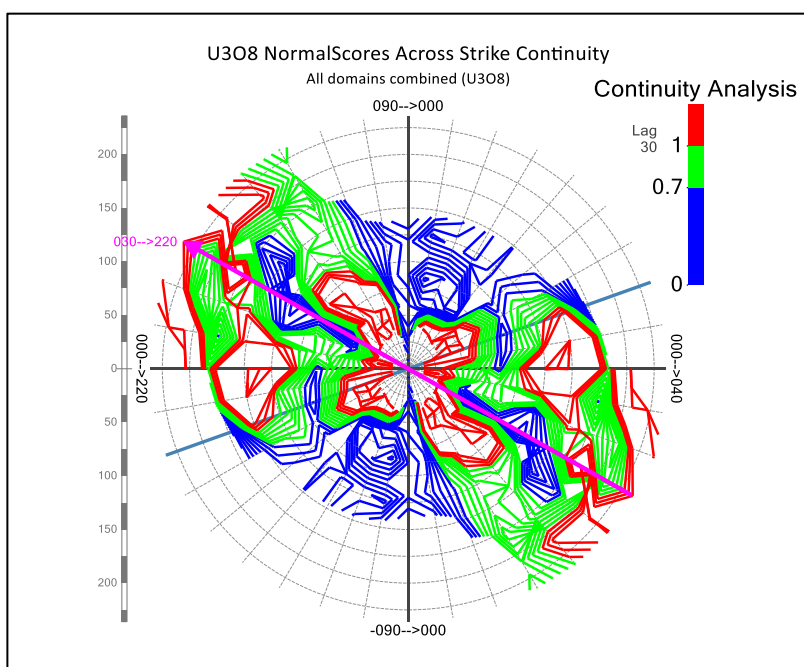


Figure 5.23: Variogram fans across strike the purple arrow shows the direction of maximum continuity

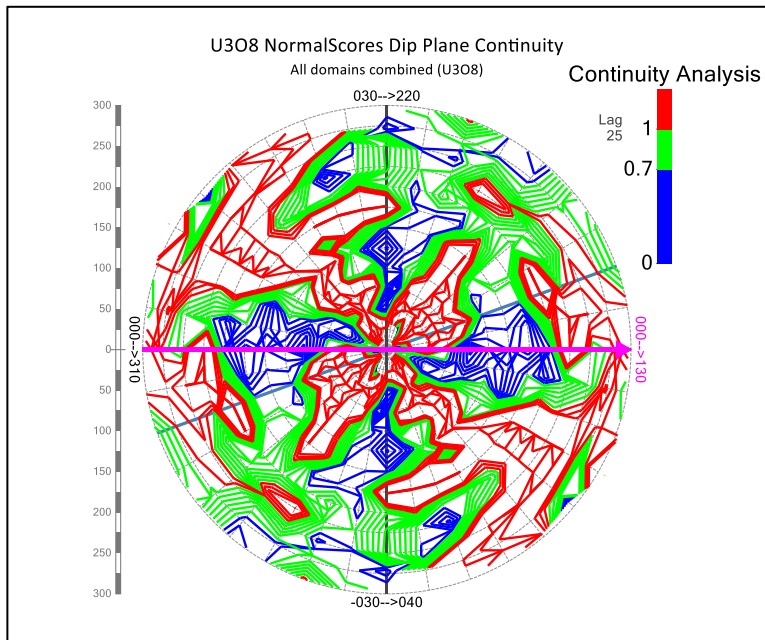


Figure 5.24: Variogram fans in the dip plane the purple arrow indicates the direction of maximum continuity

Variography analyses and modelling were done on both the original U_3O_8 (ppm) composites and the U_3O_8 normal score transformed data composites.

Figure 5.25 and Figure 5.26 show the experimental and fitted downhole variograms from which the nugget effect had been estimated for both the original 1.5m composites and the normal score transformed 1.5m U_3O_8 data.

The experimental and fitted semi-variogram models for original U_3O_8 composites and normal scores data are shown in Figures 5.27 to 5.32.

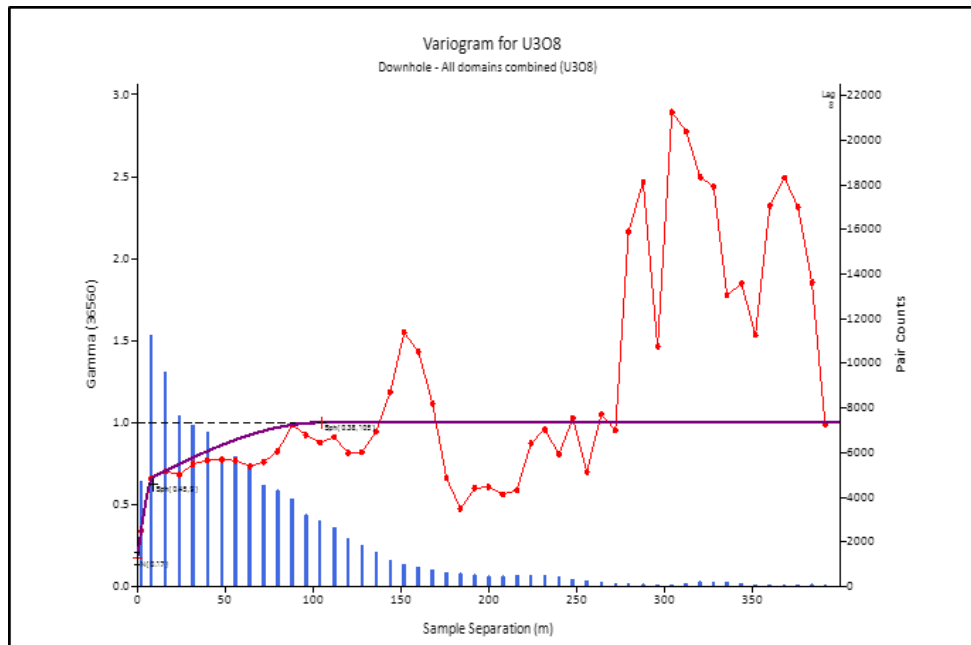


Figure 5.25: Downhole semi-variogram - Original U₃O₈ (ppm) composites

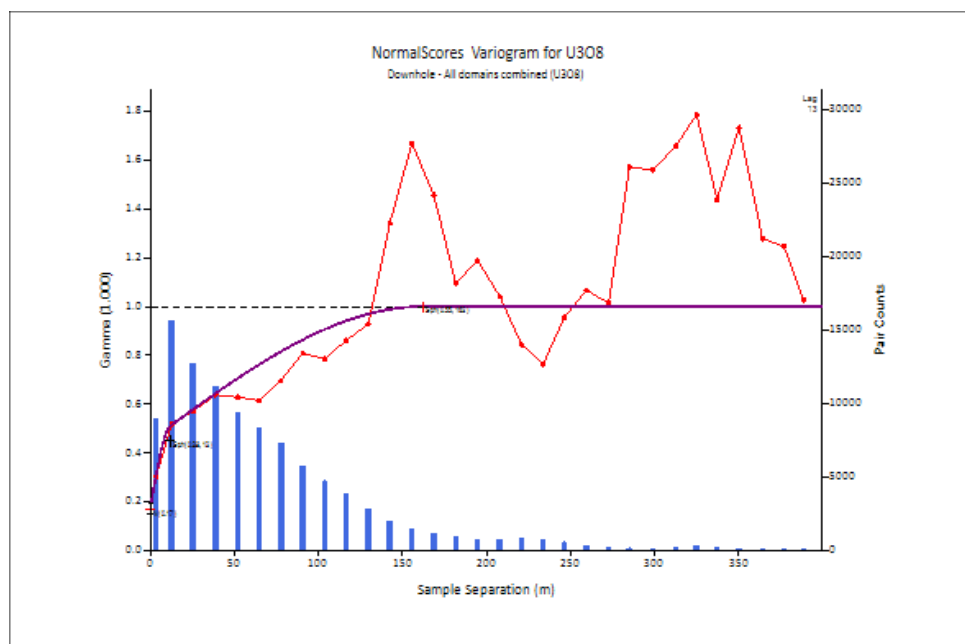


Figure 5.26: Downhole semi-variogram - Normal score transformed data

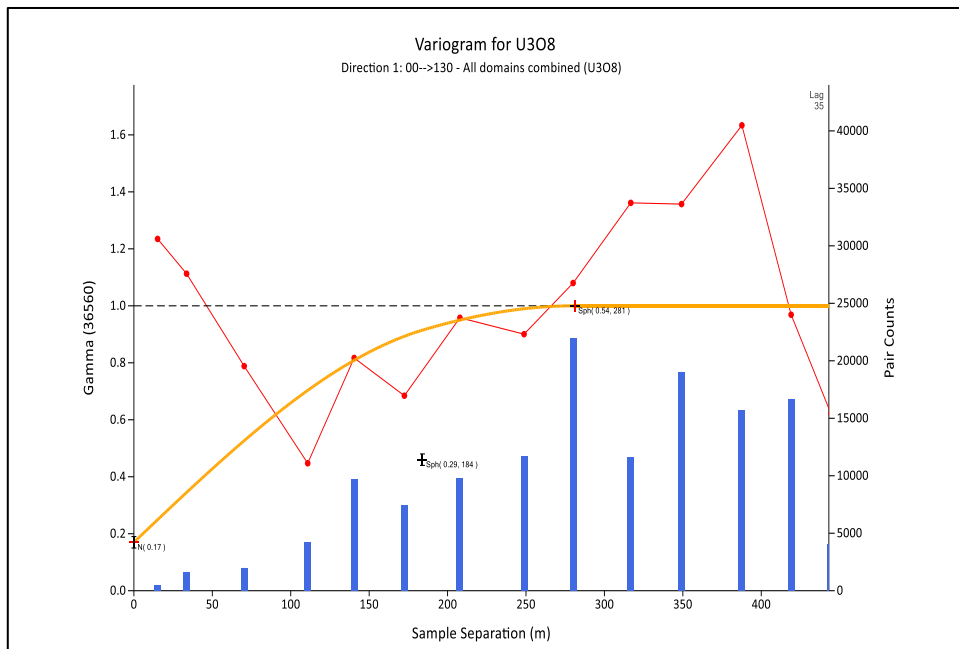


Figure 5.27: Major Direction (1) 00°→130°semi-variogram U₃O₈ (ppm) composites

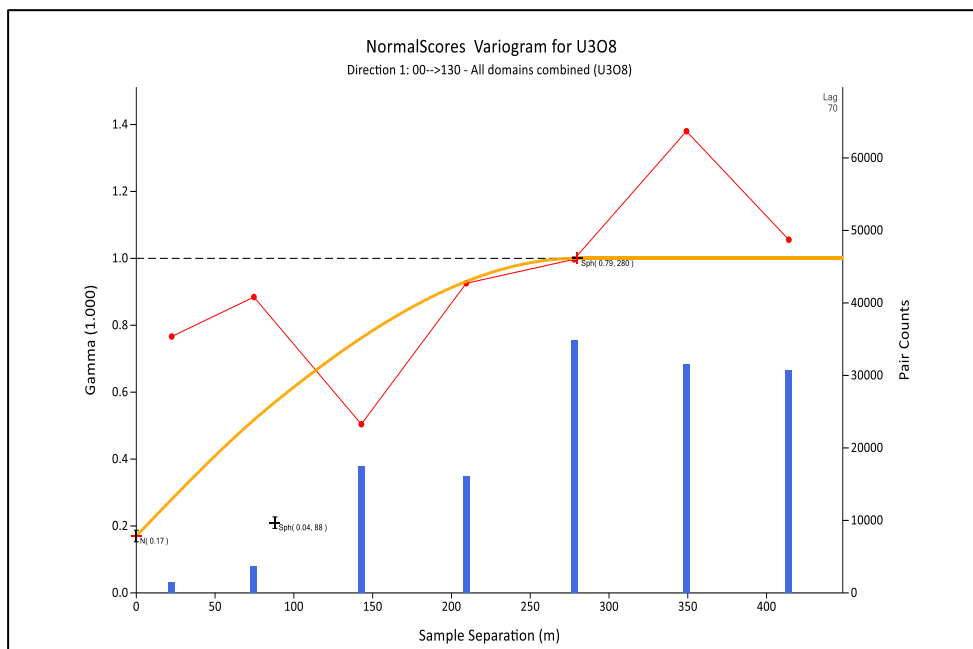


Figure 5.28: Major Direction (1) 00°→130°semi-variogram Normal score

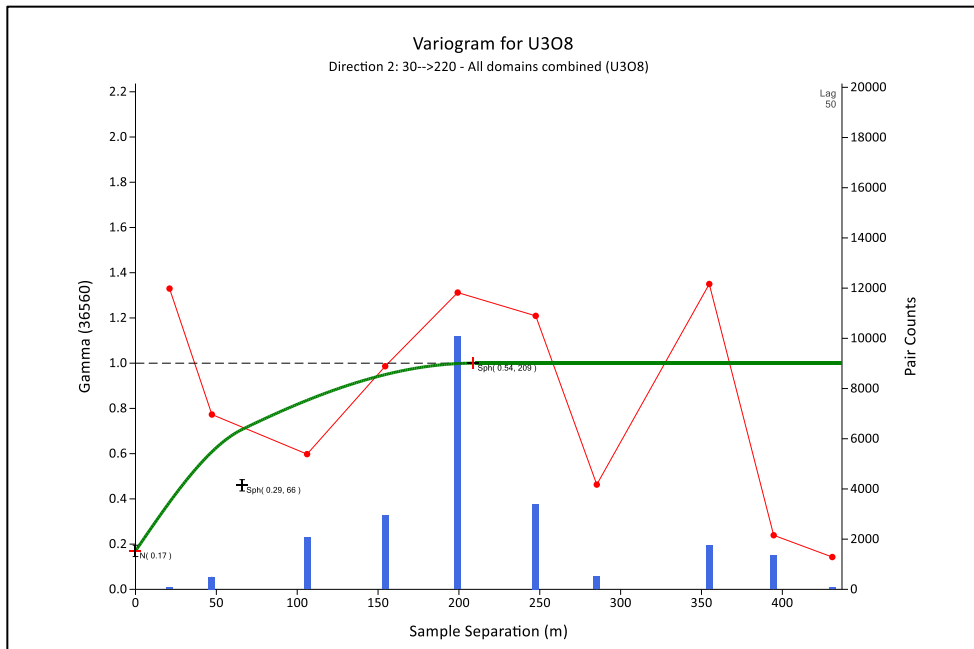


Figure 5.29: Semi-major direction $30^{\circ} \rightarrow 220^{\circ}$ semi-variogram U_3O_8 (ppm) composites

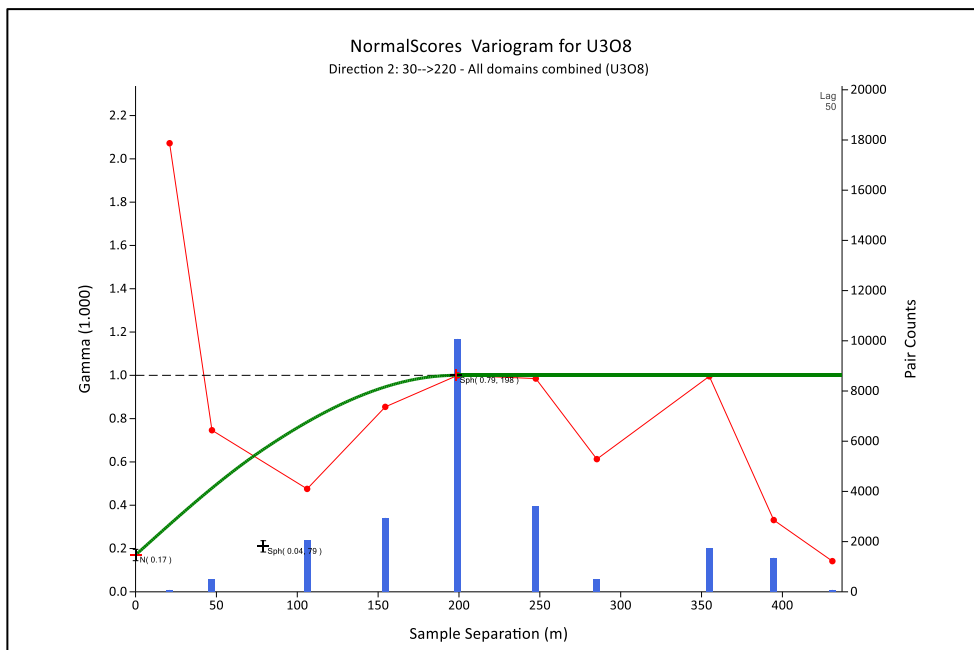


Figure 5.30: Semi-major direction $30^{\circ} \rightarrow 220^{\circ}$ semi-variogram - Normal score data

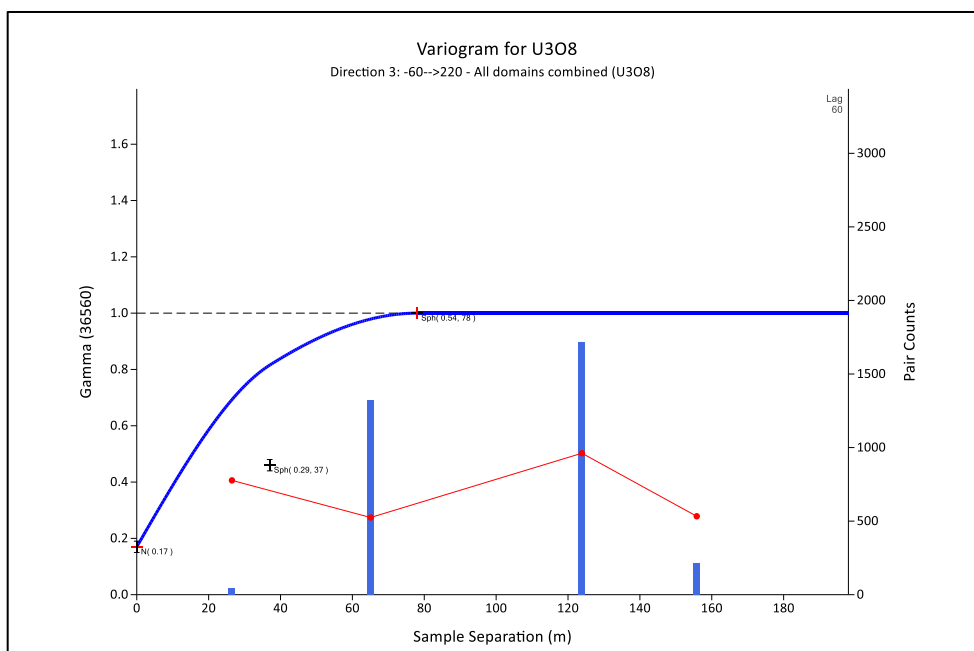


Figure 5.31: Minor direction -60°→220° semi-variogram U₃O₈ (ppm) composites

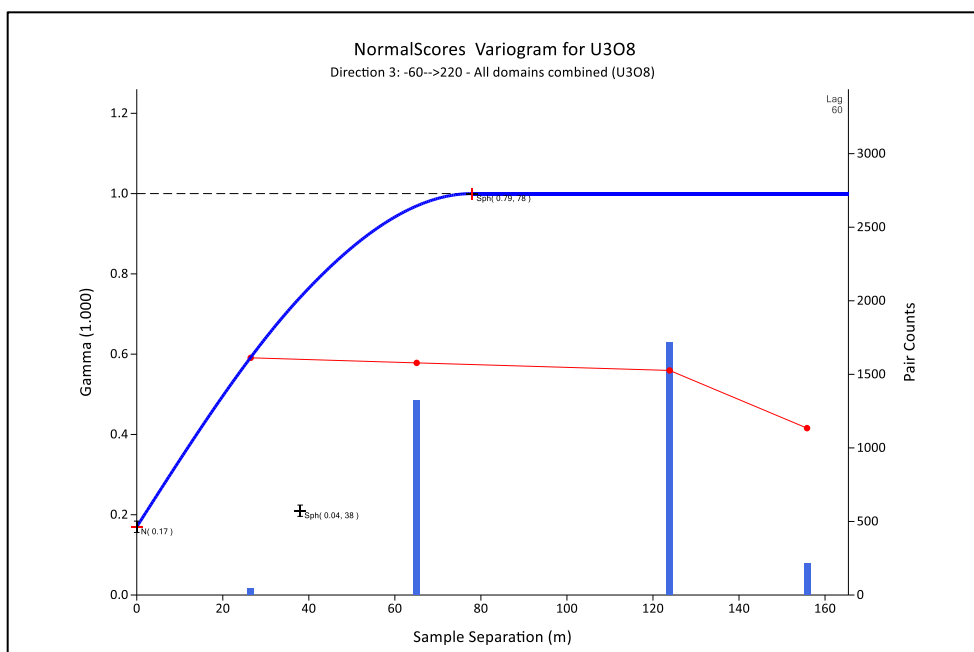


Figure 5.32: Minor direction -60°→220° semi-variogram – Normal score data

The semi-variogram model for the U_3O_8 (ppm) composites is shown in Figure 5.33 and the U_3O_8 (ppm) normal score back transformed semi-variogram model in Figure 5.34. Although the semi-variogram modelling of the original U_3O_8 (ppm) composites and normal scores data was done in the same directions of continuity, the back transformed variogram models are different.

The nugget effect of the normal score back transformed semi-variogram model is 19% which is slightly higher than that of the observed 1.5m U_3O_8 (ppm) composite semivariogram models. Within the first structure of this back transformed semi-variogram the behaviour of the variability in first and second directions is fairly isotropic with range of 88m and 79m respectively. In the third direction ranges are of the same order in both cases (38m and 37m) respectively. A more significantly difference is the reduction in the sill of the back transformed semi-variogram model it is 8% versus 29% for the semi variogram model for U_3O_8 (ppm) composites. Meaning that in carrying out the normal score transformation less spatial variability and higher spatial correlation has been identified within the first range which was not that clear from the experimental semi- variogram of the observed 1.5m U_3O_8 (ppm) composite data.

Hence the spatial correlation for the third structure of the semi variogram models differ and is more for the back transformed. The second ranges of the two semivariogram models are similar. The scalar change of the normal score transformation assisted in improving the modelling of the individual variograms. Table 5.10 summarises accepted double structured semivariogram model for the estimation of the U_3O_8 (ppm)

Table 5.10: Accepted double structured spherical semi-variogram parameters

		Structure 1		Structure 2	
Direction	Nugget	Sill	Range (m)	Sill	Range (m)
(1) $00^\circ \rightarrow 130^\circ$	0.19	0.08	88	0.73	280
(2) $30^\circ \rightarrow 220^\circ$			79		198
(3) $-60^\circ \rightarrow 220^\circ$			38		78

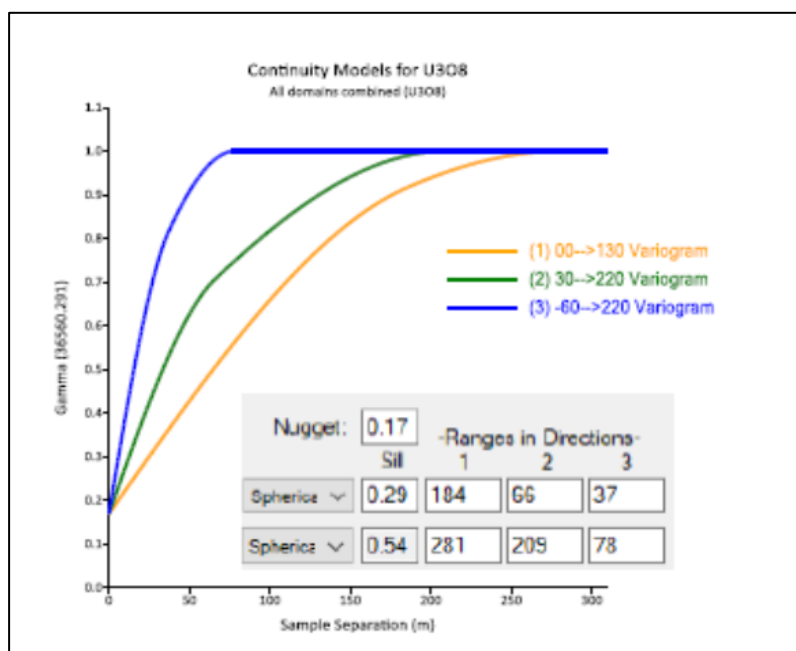


Figure 5.33: Semi variogram model for U_3O_8 (ppm) composites

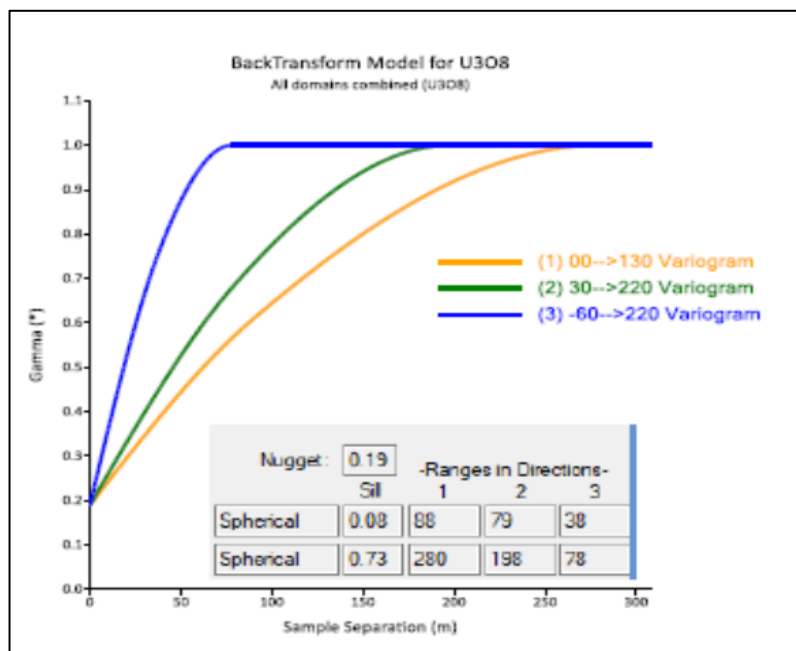


Figure 5.34: Normal Score back transformed semi-variogram model for U_3O_8 (ppm)

6 KNA procedures in SUPERVISOR and SURPAC

6.1 Measures for minimising conditional bias

The minimisation of conditional bias can be addressed by a kriging neighbourhood analysis, whereby the optimum kriging neighbourhood parameters are theoretically identified prior to the actual estimation. In such an analysis the slope of regression (SLOR) of the true block values conditioned by the estimated block values, the kriging efficiency (KE) and number of negative weights are all assessed in an attempt to minimise conditional bias. These measures are affected by the following parameters in the kriging neighbourhood plan: the estimation block size, the minimum and maximum number of samples to be considered, the search range, the configuration of the sample points and the block discretisation.

The KE and SLOR are theoretically calculated using the:

- covariogram model, which is a user input derived from the semi-variogram model
- sample locations, as observed in the orebody and captured in the data file used
- size and position of the test block relative to the sample positions, also user defined
- number of kriging equations, that are a function of the minimum and maximum number of samples, which are user input
- kriging weights come from solving the kriging equations

The graphical and the theoretical relationships between the covariance $C(h)$ and the semi variance $\gamma(h)$ for any distance h is given in Figure 6.1 below (Dohm, 2015a)

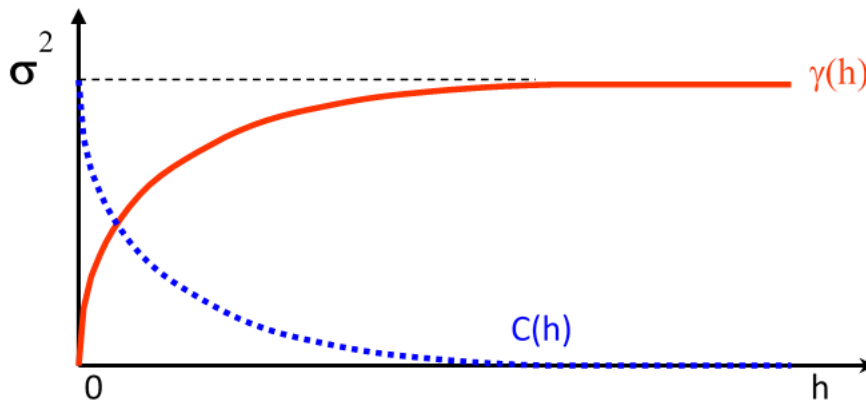


Figure 6.1: Theoretical relationships between the covariance and the semi-variance

$C(h) = C - \gamma(h)$ where C = Sill of the variogram

$C(0) = \sigma^2$ = population variance = $C_0 + C_1 + C_2$ = sum of the semi-variogram sills

Covariance functions for v = sample support and V = block support

$C(v, v) = \{ C_0 + C_1 + C_2 \} - \gamma(v, v) = \sigma^2 - \gamma(v, v)$ covariance relationship of the samples

$C(V, V) = \{ C_0 + C_1 + C_2 \} - \gamma(V, V) = \sigma^2 - \gamma(V, V)$ covariance relationship of discretised points in the block

$C(v, V) = \{ C_0 + C_1 + C_2 \} - \gamma(v, V) = \sigma^2 - \gamma(v, V)$ covariance relationship of sample v and every discretised point in block V

Slope of regression of the true block values given the estimated blocks values expressed in terms of the covariogram is given by

$$b = \frac{\text{Cov}\{Z_v, Z_v^*\}}{\sigma_{Z_v^*}^2} = \frac{\sum_{i=1}^n \lambda_i \bar{C}(v_i, V)}{\sum_{i=1}^n \sum_{j=1}^n \lambda_i \lambda_j C(v_i, v_j)}, \text{ where } \lambda_i \text{ for } i = 1, 2, \dots, n \text{ are the kriging}$$

weights, and n the number of samples in the kriging neighbourhood

Kriging Efficiency in terms of the covariogram is given by

$$KE = \frac{\bar{C}(V, V) - \sigma_K^2}{\bar{C}(V, V)}$$

Where σ_K^2 is the kriging variance and in terms of the covariogram is given by

$$\sigma_K^2 = \bar{C}(V, V) - \sum_{i=1}^n \lambda_i \bar{C}(v_i, V) + \lambda$$

Prior to comparison of results, an overview of how Kriging Neighbourhood Analysis as done in SUPERVISOR and as done in SURPAC is discussed. The procedures for carrying out KNA in SURPAC are stated in the Geostatistical tutorial manual. In SUPERVISOR the KNA guidelines are explained in the User's manual and a demonstration video.

The KNA procedure in SUPERVISOR

The estimation method used in this research is OK hence KNA is done specifically for OK. The KNA process in SUPERVISOR provides the kriging efficiency (KE) and the slope of regression (SLOR) as output measures of conditional bias. In addition, the number of negative weights, block covariance, estimated grade, estimate variance (kriging variance) and LaGrange multiplier are also provided.

- The first step involves testing blocks of variable sizes. The block size with the highest KE and SLOR and a few or no negative weight is chosen.
- The minimum and a maximum number of samples for the chosen block size are determined in the second step. The minimum number of samples with the highest KE and slope of regression is selected. The maximum number of samples is the value at which an increase in the number of samples does not cause much improvement in KE and SLOR values.
- After the number of samples have been identified, the best search range is determined. Various search ranges are tested including the one obtained from the semi-variogram. The optimal search range is the one with the highest KE and SLOR.
- Lastly, discretisation points are determined by plotting a graph of the block covariance against the discretisation points. The point at which there is no downward trend in the block covariance is considered suitable.

The KNA procedure in SURPAC

The following procedure is followed when testing various block sizes for the kriging to choose the best block size neighbourhood.

- First the X, Y, and Z coordinates of the test location is selected.
- After that, the block dimensions (X, Y, and Z) are chosen.
- A block model is created using an origin such that the coordinates of the centroid of the block at the test location in the model are the same as the coordinates of the test location.

- Ordinary kriging is performed within the constraint to estimate a value for the block centred at the test location. This is done to solve the kriging equations. An OK output file is created with the values for Kriging Efficiency (KE), conditional bias slope (SLOR), estimated grade, kriging variance, block variance, and LaGrange multiplier.
- The block size with the highest values of KE; highest conditional bias slope, lowest kriging variance, and no or minimum number of negative weights is chosen.
- The procedure explained above (KNA in SURPAC) is also used in determining the minimum and a maximum number of samples, search radius and discretization points.

6.2 Identification of well-informed and poorly-informed areas in Rossing Case study area

The Rossing data base consists of few and scattered drill holes (Figure 6.2) with different drilling orientations which made the block location selection process for a KNA analysis difficult, hence only two block configurations could be tested i.e. well-informed blocks, and poorly-informed blocks. This process was carried out in SUPERVISOR and the same locations were considered in SURPAC.

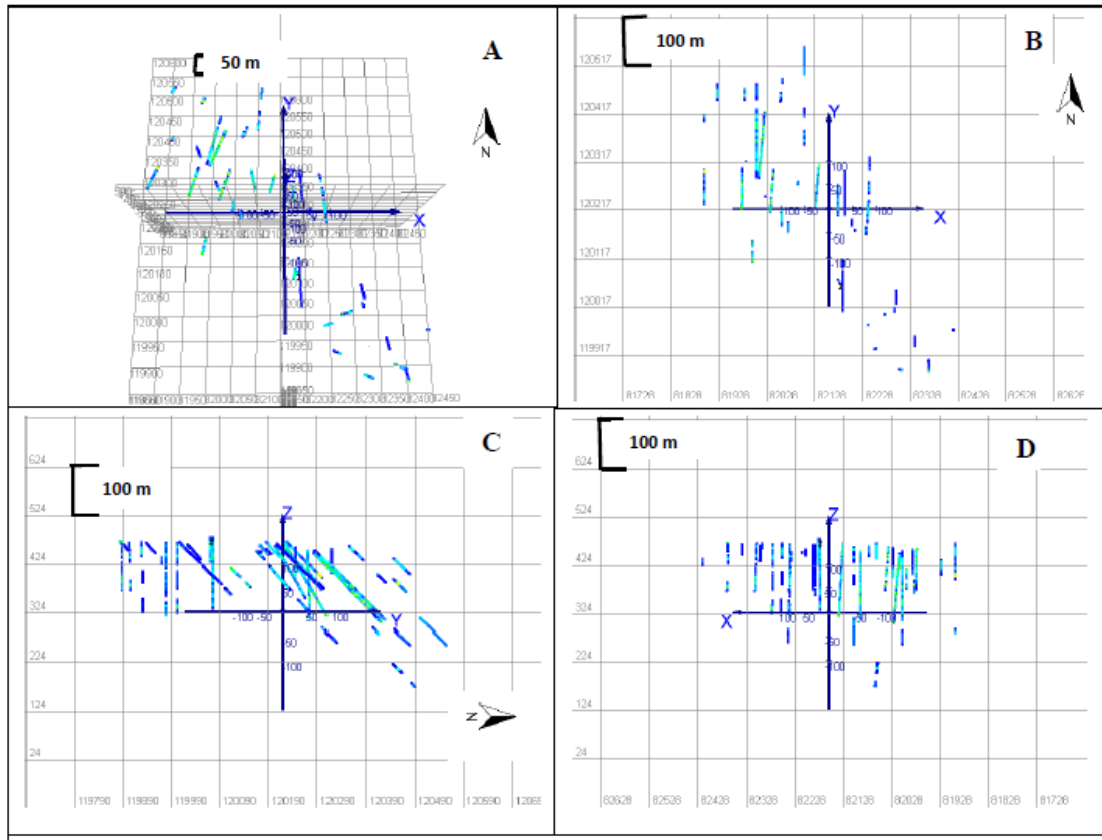


Figure 6.2: All composite samples available for KNA. **A;** 3-D view from top **B;** 2-D view from top **C;** 2-D view from right **D;** 2-D front view facing North.

Well-informed blocks, have three internal drill holes with some surrounding drill holes whereas poorly-informed blocks have one internal drill hole with surrounding drill holes. Analysis for well-informed blocks was done at $Y = 120\,326$, $X = 81\,997$, $Z = 422$ (Figure 6.3) while poorly-informed blocks were performed at $Y = 119\,941$, $X = 82\,306$, $Z = 448$ (Figure 6.4).

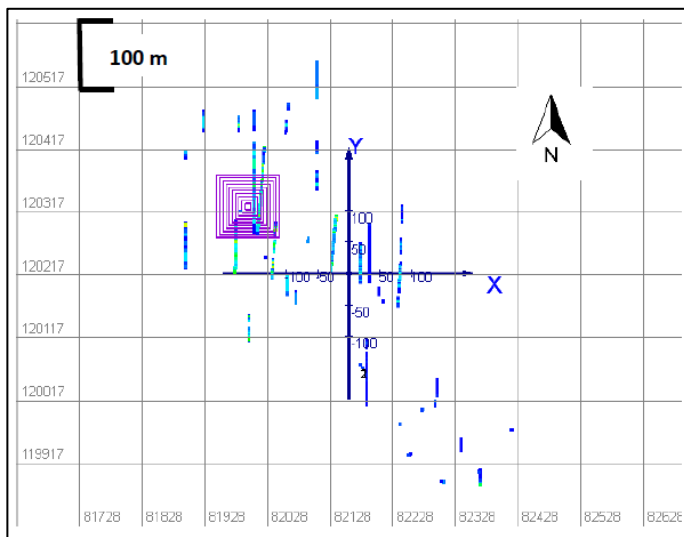


Figure 6.3: Location of well-informed blocks with three internal and surrounding drill holes.

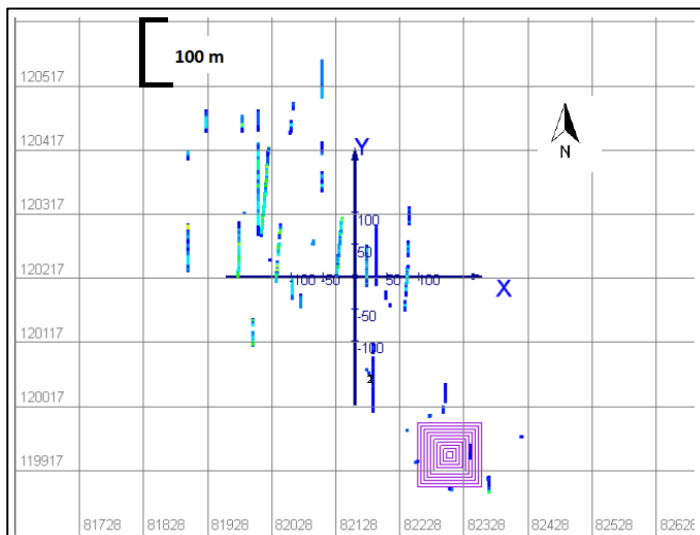


Figure 6.4: Location of poorly-informed blocks with one internal and surrounding drill holes.

6.2.1 SUPERVISOR: Well-informed blocks

6.2.1.1 SUPERVISOR: Well-informed blocks - Optimum block size

Numerous block sizes ranging from 10m (X) x 10m (Y) x (Z) increasing steps of 10m in the X and Y directions up to 100m (X) x 100m (Y) x (Z), and in the Z direction for Z = 10, 15, 20 m, were tested. The mining bench height at Rossing is 15 m, therefore, elevation (Z) was tested at 10 and 20 m as well. Table C1 shown in Appendix C shows the detailed results for the slope of regression, kriging efficiency, number and proportion of negative weights of the block size range considered. This information is graphically summarised in Figure 6.5. Each drill hole was restricted to a maximum of six samples so that no individual drill hole can have an overly-large influence on the estimate.

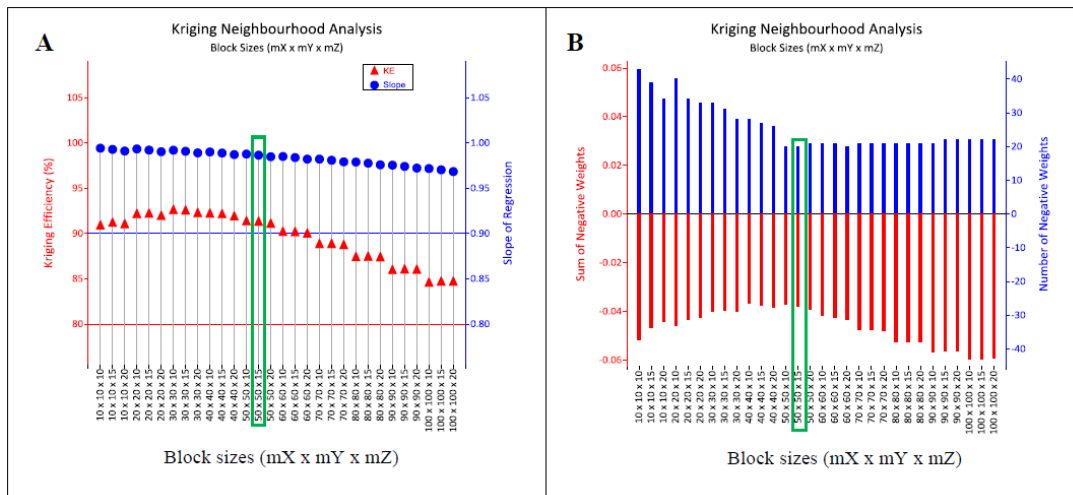


Figure 6.5: Conditional bias performance measures **A:** SLOR (blue circles) and KE (red triangles) for the range of block sizes tested. **B:** Sum (red bars) and number (blue bars) of negative weights for the range of block sizes tested.

A preferred block size would be one with a SLOR close to 1 and a KE around 100 %. In practice, such a block is rarely attained hence any block size with values above the blue dotted line (SLOR = 0.9) and red dotted line (KE = 80 %) with no or a low proportion of negative weights is considered optimal (Snowden Mining, 2013). In this

instance, all block sizes have SLOR and KE values above the dotted lines, however all block sizes have negative weights (Figure 6.5).

Negative weights arise when estimating a highly correlated variable (a semi-variogram with long ranges and a low nugget effect) or when estimating clustered points (Armstrong, 1998).

The drilling data in the study area is relatively sparse, leading to a low nugget effect long ranges of continuity being modelled (Table 5.10). The negative weights that are observed arise from a combination of these two factors and are also related to the size of the block considered.

Including many negative weights in the estimation process will result in negative estimates which practically do not exist (Szidarovszky *et al.*, 1987).

Smaller block sizes have high number of negative weights compared to the rest of the blocks. (Table C1, Appendix C). The number of negative weights decrease as the block size increase to medium size blocks and slightly increase with the increase in block sizes (Figure 6.5 B). The semi-variogram has a nugget effect of 0.19 with the first and second sill values of 0.08 and 0.73 respectively. The first weights are assigned to 27% $((0.19 + 0.08) \times 100)$ of the semi-variogram range and the rest are assigned to the second range (73%). Smaller block sizes have high number of negative weights because the closer samples are being screened by distant samples within the second range and because of clustered samples. For larger block sizes, negative weights are caused by clustered samples.

Block sizes of 10m x 10m x Zm to 40m x 40m x Zm can be ignored because they have high proportion of negative weights. Amongst blocks with low proportion of negative weights, the block size search narrowed down to three sizes, 50m x 50m x 10m, 50m x 50m x 15m and 60m x 60m x 20m (Table C1). A 60m x 60m x 20m block was also ignored because of the low SLOR and KE values. There is not much improvement in the KE, SLOR and proportion of negative weights of the 50m x 50m x 10m and the

50m x 50m x 15m block sizes (Figure 6.5 A and B) and the 50m x 50m x 15m block size was chosen as being optimal direction size. The choice of this block size is also confirmed by the comparisons that follow.

The Y and X values of the 50m x 50m x 15m block size were varied to see if any of the resultant blocks would improve the results. The results of these varied blocks are presented in Table 6.1 and Table 6.2.

Table 6.1: SUPERVISOR: KN measures on 50m x **variable Y m** x 15m blocks

Block size 50m x Ym x 15m	Number of samples	SLOR: Slope of regression	KE: Kriging efficiency	No. of negative weights	Proportion of negative weights
50 x 10 x 15	98	0.989	0.915	30	0.306
50 x 20 x 15	98	0.989	0.917	29	0.296
50 x 30 x 15	98	0.988	0.918	26	0.265
50 x 40 x 15	98	0.988	0.917	25	0.255
50 x 50 x 15	98	0.987	0.913	20	0.204

Table 6.2: SUPERVISOR: KN measures on **variable X m** x 50m x 15m blocks

Block size Xm x 50m x 15m	Number of samples	SLOR: Slope of regression	KE: Kriging efficiency	No. of negative weights	Proportion of negative weights
10 x 50 x 15	98	0.990	0.912	24	0.245
20 x 50 x 15	98	0.990	0.915	22	0.224
30 x 50 x 15	98	0.989	0.917	23	0.235
40 x 50 x 15	98	0.988	0.916	22	0.224
50 x 50 x 15	98	0.987	0.913	20	0.204

Four of the blocks that were varied in the Y dimension besides having high SLOR and KE values, also have higher proportions of negative weights compared to the proportion of negative weights of the 50m x 50m x 15m block. Similarly, blocks that were varied in the X dimension all have high SLOR and KE values but have higher proportions of negative weights in comparison to the ones of the 50m x 50m x 15m block except for the 10m x 50m x 15m block which has a slightly lower KE.

From the above analysis on block size the 50m (X) x 50m (Y) x 15m (Z) block is considered the optimal block size for well-informed blocks.

6.2.1.2 SUPERVISOR: Well informed blocks - Minimum and Maximum number of samples

To determine the optimum number of samples for the 50m x 50m x 15m block size a plot of the SLOR and KE against a minimum number of two and a maximum number of 98 samples to be considered in a kriging neighbourhood were analysed. The details of this analysis appear in Appendix C Table C2; for legibility, the interval of the number of samples in Figure 6.6 A was increased to four and the graph shown in Figure 6.6 B is only up to 64 samples in the kriging neighbourhood.

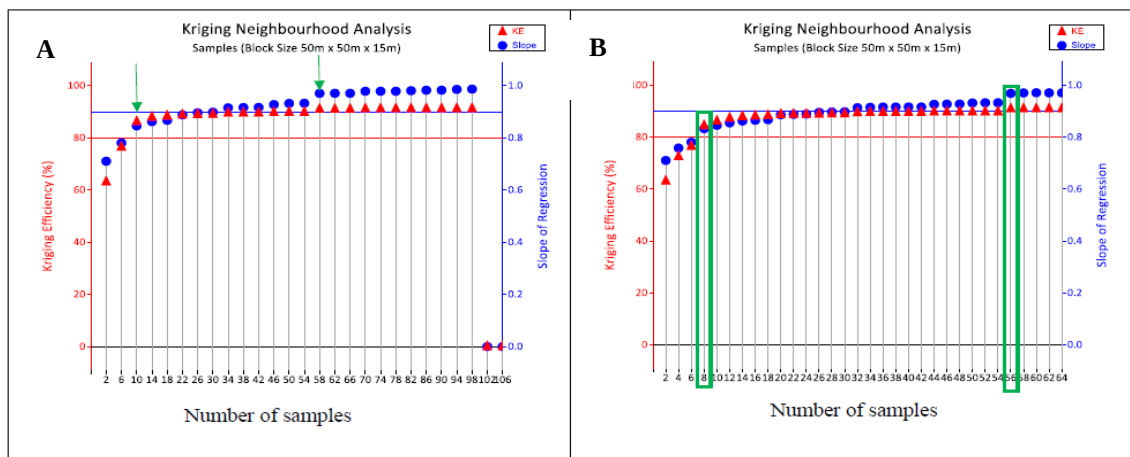


Figure 6.6: Number of samples vs **A:** SLOR (blue circles) and KE (red triangles) for a 50m x 50m x 15m block, **B:** selected minimum and maximum number of samples highlighted

A significant increase in the SLOR and KE is observed as the number of samples increases from 2 to 8 (Figure 6.6 B). The number of samples beyond 8 have KE values above the kriging efficiency dotted line (red) and less than 8 have KE values below the kriging efficiency dotted line (Figure 6.6 B).

The maximum number of samples to consider in the kriging neighbourhood is a value beyond which there is no significant improvement in the SLOR and KE. Beyond the 56th sample, increasing the number of samples has negligible improvement impact on the SLOR and KE see Figure 6.6 B with supporting data in Appendix C Table C2.

A minimum of 8 samples and a maximum of 56 samples are to be used in the kriging estimation neighbourhood of well-informed 50m x 50m x 15m blocks.

6.2.1.3 SUPERVISOR: Well-informed blocks – Optimum Search Range

Vann *et al.*, (2003) and Khakestar *et al.*, (2013) showed that it is not a good idea to define the neighbourhood using the semi-variogram range. The semi-variogram range might either be too restrictive that it excludes correlated samples or too long that it includes un-correlated samples.

Therefore, to identify the optimum search ranges, the semi-variogram range was shortened and extended to see if there is any improvement in the SLOR and KE values. Five scenarios were tested. In scenario 1 and 2 the search ranges considered were at 60 % and 80 % of the semi-variogram range. Scenario 3 represents the range of the semi-variogram, and in the last two scenarios, the search ranges were increased to 1.2 and 1.4 times that of the semi-variogram range. These ranges are presented in Table 6.3. Instead of using the maximum number of samples, 56, identified in section 6.2.1.2, a total number of correlated samples (98) were used in determining the optimal ellipsoid search range. Determination of the search range using a maximum of 56 samples resulted in the same SLOR and KE values for all search range scenarios. This is due to limited samples as a results of the sparsely drilled coverage of the Rossing Deposit.

Table 6.3: SUPERVISOR: well-informed variogram search ranges

Range ratio	Direction 1	Direction 2	Direction 3
0.60	168.00	118.80	46.80
0.80	224.00	158.40	62.40
1.00	280.00	198.00	78.00
1.20	336.00	237.60	93.60
1.40	392.00	277.20	109.20

It is evident from Table 6.4 that the number of samples for the 0.6 range ratio tested is 60, which is more than 56.

Table 6.4: SUPERVISOR: Search range optimisation for 50m x 50m x 15m blocks

Range Ratio	Number of samples	SLOR: Slope of regression	KE: Kriging efficiency	No. of negative weights	Proportion of negative weights
0.60	60	0.971	0.910	9	0.092
0.80	84	0.981	0.912	20	0.204
1.00	98	0.987	0.913	20	0.204
1.20	98	0.987	0.913	20	0.204
1.40	98	0.987	0.913	20	0.204

Shortening of the semi-variogram range to 60 and 80 % resulted in the exclusion of 38 and 14 samples respectively. This explains why the two scenarios have the lowest SLOR and KE. Extension of the semi-variogram range did not improve the results (Table 6.4) because all 98 samples have been utilised so there are no further samples to include. The range ratio of 1.00 (semi-variogram) range was chosen. This is what was expected because the semi-variogram range was used to determine the optimum 50m x 50m x 15m block size.

6.2.1.4 SUPERVISOR: Well-informed blocks – Optimum Block Discretisation

Discretisation is the process of calculating point-block average values (Chanderman, 2015). The maximum value for discretisation in SUPERVISOR is 10. Discretisation points that are compatible with the 50m x 50m x 15m block dimensions were tested. The tested discretisation points in the X and Y dimensions are 2, 5 and 10. In the Z dimension, the tested discretisation points are 3, and 5. Combinations of compatible discretisation configurations were tested on the optimal block size of 50m x 50m x 15m and the KE (%) & SLOR results for these are shown in Figure 6.7 with the details summarised in Table C3, Appendix C.

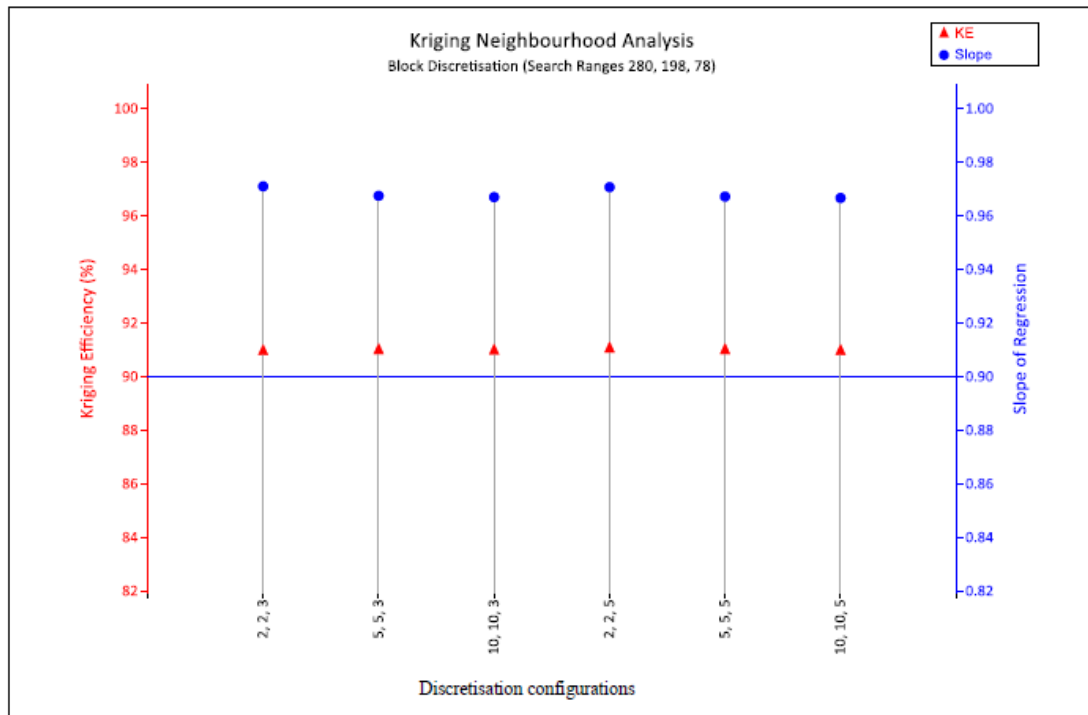


Figure 6.7: Discretisation configuration optimisation for the 50m x 50m x 15 m blocks

It is difficult to pinpoint a suitable discretisation configuration for the 50m x 50m x 15m block from Figure 6.7 because of small difference in the KE and SLOR values of the discretisation configurations tested. It is for this reason that the block covariance versus discretisation configurations was plotted (Figure 6.8). The discretisation configuration at which there is no significant downward trend in the block covariance is selected.

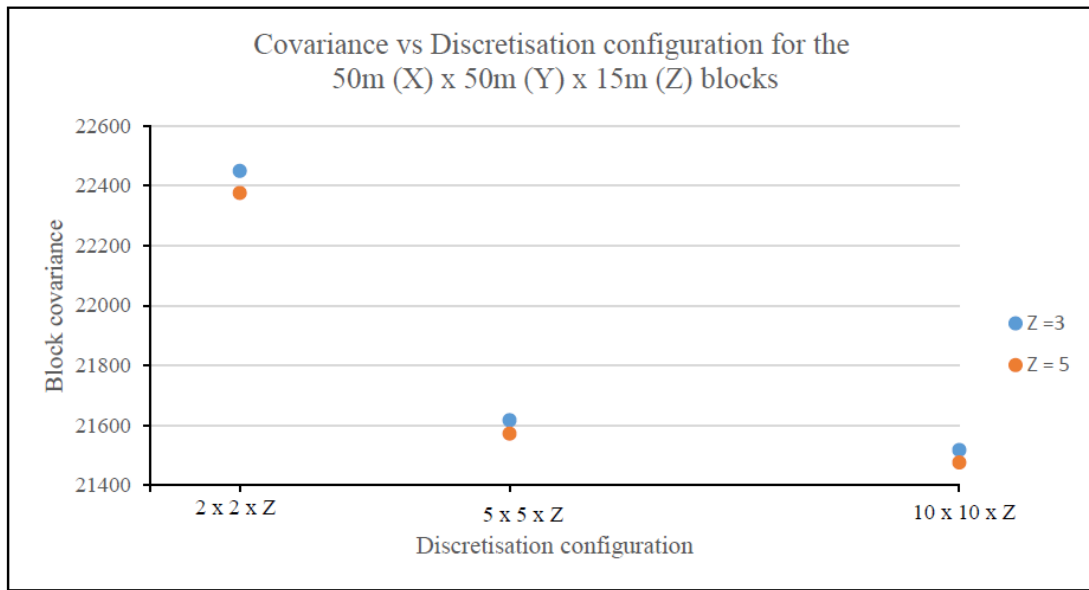


Figure 6.8: Block covariance vs discretisation configuration for 50m x 50m x 15m blocks.

The above plot shows that from 5x5xZ, the covariance does not show a significant downward trend. The 5x5x3 configuration has however better SLOR and KE values to minimize conditional bias compared to the ones for the 5x5x5 configuration (Table C3, Appendix C). Because of these two reasons, appropriate discretisation points for the 50m x 50m x 15m blocks is 5x5x3.

6.2.2 SUPERVISOR: Poorly-informed blocks

6.2.2.1 SUPERVISOR: Poorly-informed blocks – Optimum Block Size

It is not worth testing for block sizes which proved to be too small for well-informed areas. Block sizes tested for poorly informed areas are from 50m x 50m x Zm (Z = 10, 15 and 20m) where the X and Y dimensions are varied in multiples of 50 (Figure 6.9).

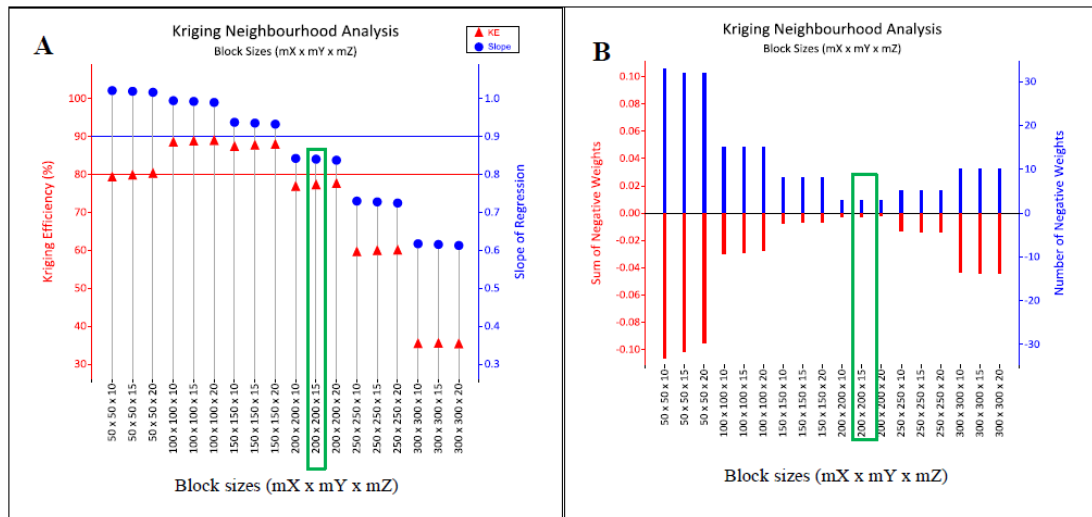


Figure 6.9: A: Block size optimisation for poorly-informed blocks. **B:** Negative weights for the tested block sizes.

Blocks of size 50m x 50m x Z to 150m x 150m x 15m have high SLOR and KE with some significant proportion of negative weights (Figure 6.9). The 250m x 250m x Zm and 300m x 300m x Zm block sizes have low SLOR and KE coupled with proportions of negative weights. The 200m x 200m x Zm block sizes were selected because of the low proportion of negative weights (Figure 6.9 B and Table C4). The improvement in the SLOR and KE values of blocks of size 200m x 200m with different Z values (Figure 6.9 A), is insignificant, therefore a block of size 200m x 200m x 15m is considered optimal for poorly-informed blocks.

Results of varying the Y dimension and X dimensions of a 200m x 200m x 15m block are respectively given in Table 6.5 and Table 6.6.

Table 6.5: SUPERVISOR: KN measures on 200m x variable Ym x 15m blocks.

Block size Xm x Ym x Zm	Number of samples	SLOR: Slope of regression	KE: Kriging efficiency	No. of negative weights	Proportion of negative weights
200 x 50 x 15	78	0.947	0.828	9	0.115
200 x 100 x 15	78	0.930	0.855	9	0.115
200 x 150 x 15	78	0.895	0.838	6	0.077
200 x 200 x 15	78	0.840	0.771	3	0.038

Table 6.6: SUPERVISOR: KN measures on variable Xm x 200m x 15m blocks

Block size Xm x Ym x Zm	Number of samples	SLOR: Slope of regression	KE: Kriging efficiency	No. of negative weights	Proportion of negative weights
50 x 200 x 15	78	0.937	0.800	14	0.179
100 x 200 x 15	78	0.916	0.825	7	0.090
150 x 200 x 15	78	0.883	0.816	3	0.038
200 x 200 x 15	78	0.840	0.771	3	0.038

None of the blocks that were varied in the Y dimension gave better conditional bias measures (SLOR, KE and proportion of negative weights) than the 200m x 200m x 15m block (Table 6.5). Variation of blocks in the X dimension resulted in the 150m x 200m x 15m block size which has higher SLOR and KE compared to those of the 200m x 200m x 15m block (Table 6.6). The 150m (X) x 200m (Y) x 15m (Z) block is selected as the best block size for poorly-informed blocks.

6.2.2.2 SUPERVISOR: Poorly-informed blocks - Minimum and Maximum number of samples

The optimal number of samples that can be used in the kriging estimation of the 150m x 200m x 15m block size are determined from a plot of kriging efficiency and slope of regression against the number of samples (Figure 6.10).

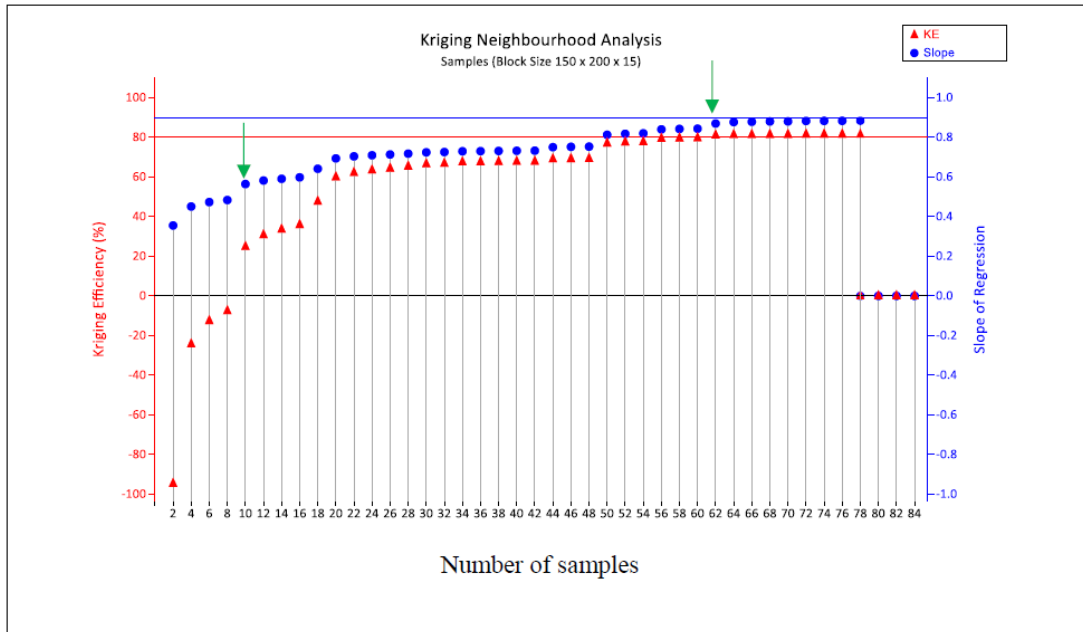


Figure 6.10: Number of samples vs SLOR (blue circles) and KE (red triangles) for a 150m x 200m x 15 m block

All cases with less than 10 samples gave negative KE and very low SLOR values (Figure 6.10 and Table C5). Negative KE values imply that the kriging variance is greater than the block variance (Krige, 1996b). Negative KE and low SLOR suggests that a minimum number of samples for poorly-informed blocks should not be less than 10 samples. A significant increase in the SLOR is observed as the number of samples increase from 8 to 10 (Figure 6.10). A minimum of 10 samples was considered fair for the 150m x 200m x 15m block size.

Increasing the number of samples in the kriging neighbourhood beyond 42 result in negative weights (Table C5, Appendix C) being observed. As the number of samples increase to 62, no significant increase in the SLOR is observed until all 78 samples are included (Figure 6.10). The maximum number of samples of the 150m x 200m x 15m block size is therefore selected as 62.

6.2.2.3 SUPERVISOR: Poorly informed blocks – Optimum search range

As with the well-informed blocks, the semi-variogram range was shortened to 20% and 40% and extended by 20% and 40%. As with well-informed blocks, all correlated

samples were used instead of 62 samples. The results of the optimisation of the search ranges are presented in Table 6.7.

Table 6.7: SUPERVISOR: Search range optimisation for 150m x 200m x 15m blocks

Range ratio	Number of samples	SLOR: Slope of regression	KE: Kriging efficiency	No of negative weights	Proportion of negative weights
0.60	60	0.843	0.797	3	0.050
0.80	78	0.883	0.816	3	0.038
1.00	78	0.883	0.816	3	0.038
1.20	78	0.883	0.816	3	0.038
1.40	78	0.883	0.816	3	0.038

The first scenario (range ratio = 0.60) recorded the lowest KE, SLOR and a high proportion of negative weights (Table 6.7). This is because 18 correlated samples were excluded from the search volume due to the restriction on the range. The KE and SLOR values become constant as from a 0.80 range ratio because all correlated samples have been included in the longer search. A search range corresponding to a range ratio of 0.8 of the semi-variogram range ($a_1 = 224.00\text{m}$, $a_2 = 158.40\text{m}$, $a_3 = 62.40\text{m}$) is favoured.

6.2.2.4 SUPERVISOR: Poorly-informed blocks – Optimum block discretisation

Ten series of compatible discretisation configurations were tested on the 150m x 200m x 15m block size as shown in Figure 6.11.

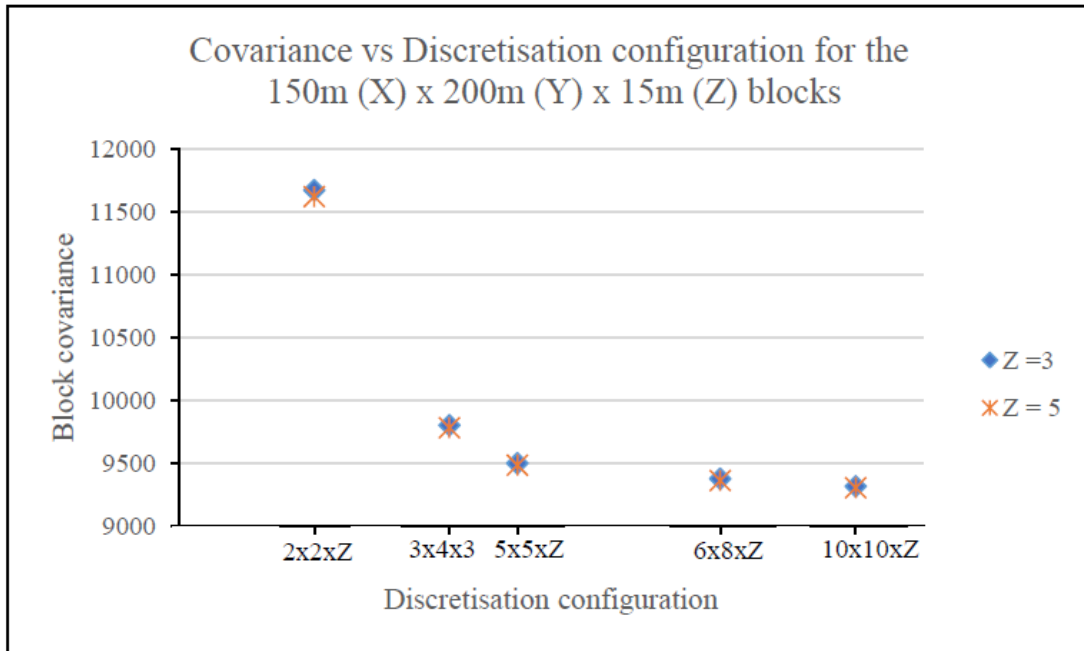


Figure 6.11: SUPERVISOR: Block Covariance versus Discretisation configuration for 150m x 200m x 15m blocks.

The block covariance decreases significantly up until 5x5xZ discretisation points. From 5x5xZ onwards, the decrease in block covariance is minimal (Figure 6.11). The two discretisation configurations (5x5x3 and 5x5x5) seems to have the same conditional bias measures, however the SLOR for 5x5x3 is 0.859 whereas for 5x5x5 it is 0.858 (Table C6, Appendix C). It is because of this slight difference in the SLOR that the 5x5x3 discretisation configuration has been identified as the preferred for the 150m x 200m x 15m blocks

6.3 Kriging Neighbourhood Analysis (KNA) in SURPAC

To obtain comparable KNA results in both software packages, the same domains, semi-variogram parameters test locations, and block configurations for well-informed and poorly informed areas were used.

SUPERVISOR rotation axis is different from that of SURPAC, thus axis rotation needed to be done so that the same search ellipsoid is used. Ellipsoid in SUPERVISOR is defined by YXZ rotation whereas in SURPAC it is defined either by ZXY/ZYX (mathematical rotations) or by ZXX/ZYZ (geological rotations).

Axes names and the order of rotation differentiate mathematical from geological rotation. Mathematical rotation requires the user to insert the bearing of the 1st axis, plunge of the 2nd axis, and a dip of the 3rd axis. Geological rotation on the other hand requires an input of a bearing of the 1st axis, dip of the 2nd axis, and a pitch of the 3rd axis (Dassault Systemes Geovia, 2015).

It can be seen from Figure 6.12 that the axis rotation from SUPERVISOR to SURPAC is defined by ZXY (mathematical rotations).

```
Surpac ZXY Rotation for
NormalScores Dip Plane (90°)

1[ 000-->130 ]-Y  2[ -030-->040 ]-X  3[ -060-->220 ]-Z
Surpac ZXY angles are [ 130.000, 0.000, 30.000 ]
```

Figure 6.12: Axis rotation from SUPERVISOR YXZ to SURPAC ZXY axis rotation

Based on the ZXY axis rotation, the order of data entry on the ellipsoid orientation form is as follows:

Table 6.8: SURPAC ZXY ellipsoid orientation angles

ZXY and ZYX Rotations		
Order of rotation	Axis Name	Value
1st Axis	Bearing	130°
2nd Axis	Plunge	0°
3rd Axis	Dip	30°

A back transformed semi-variogram model (with the parameters shown in Figure 6.13) was exported from SUPERVISOR to SURPAC file format (SURPAC.vgm).

	Sill	-Ranges in Directions-		
		1	2	3
Spherical	0.08	88	79	38
Spherical	0.73	280	198	78

Figure 6.13: Semi-variogram parameters from a back transformed model that were used in SURPAC

	1	2	3	
				(Structure)
	1.114	1.414	0	(Major)
	2.316	3.59	0	(Minor)

Figure 6.14: Semi-variogram range ratios that were used in SURPAC.

SUPERVISOR calculated the range ratios that were used in SURPAC (Figure 6.14). The range ratios of structure 1 and 2 were calculated as follow:

Structure 1

$$\text{Major} = \frac{\text{Range of major axis}}{\text{Range of semi-major axis}} = \frac{88}{79} = 1.114$$

$$\text{Minor} = \frac{\text{Range of major axis}}{\text{Range of minor axis}} = \frac{88}{38} = 2.316$$

Structure 2

$$\text{Major} = \frac{\text{Range of major axis}}{\text{Range of semi-major axis}} = \frac{280}{198} = 1.414$$

$$\text{Minor} = \frac{\text{Range of major axis}}{\text{Range of minor axis}} = \frac{280}{78} = 3.59$$

6.3.1 SURPAC: Well-informed blocks

6.3.1.1 SURPAC: Well-informed blocks – Optimum Block Size

The same block sizes that were tested in SUPERVISOR (10m (Y) x 10m (X) x Zm up to 100 m (Y) x 100m (X) x Zm where Z = 10, 15, 20 m) were tested at Y = 120 326, X = 81 997, Z = 422. Again, a maximum of six samples from each drill hole was used. For each block size, a block model centred at the coordinates of the tested location was created. Figure 6.15 shows a 100m x 100m x 10m block size centred at the coordinates of the tested location.

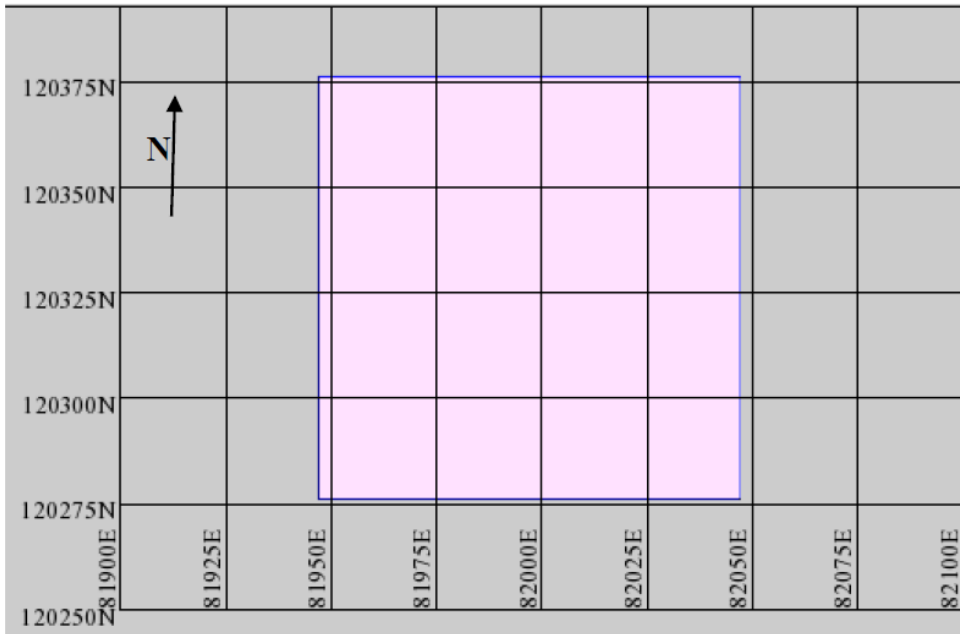


Figure 6.15: A plan view showing a 100m x 100m x 10m block centred at Y= 120 326, X = 81 997, and Z= 422.

Ordinary kriging was performed within a constraint to estimate the value of the block centred at the coordinates of the tested location. A report file with the block parameters was generated. An example of how block size analysis was done is shown below: A 100m x 100m x 10m block size is used to illustrate the block analysis process.

1. The analysis was performed at location: Y = 120 326, X = 81 997, Z = 422
2. Block dimensions: Y = 100, X = 100, Z = 10
3. Block model origin (centroid): Y = 120 276, X = 81 947, Z = 417

Block model origin values are calculated by subtracting half of the block dimension from the coordinates at which block analysis is performed. The centroid of the block is of interest because all information about the block is stored at the centre as attributes.

4. Ordinary kriging within constraint: Y < 120 376, X < 82 047, Z < 427

5. The output is a report file with an estimated grade, kriging variance, twice standard deviation, block variance, kriging efficiency, conditional bias slope which is the SLOR and LaGrange multiplier.

The conditional bias measures (SLOR, KE, KV, and negative weights) for tested block sizes are summarised in Table 6.9.

Table 6.9: SURPAC: All the tested block sizes and their conditional bias measures

Block size Ym x Xm x Zm	SLOR: Slope of regression	KE: Kriging efficiency	KV: Kriging variance	Number of negative weights	Proportion of negative weights	Number of samples
10 x 10 x 10	0.994	0.908	0.068	43	0.439	98
10 x 10 x 15	0.993	0.911	0.064	39	0.398	98
10 x 10 x 20	0.991	0.910	0.063	34	0.347	98
20 x 20 x 10	0.993	0.921	0.056	40	0.408	98
20 x 20 x 15	0.992	0.922	0.055	34	0.347	98
20 x 20 x 20	0.990	0.919	0.055	33	0.337	98
30 x 30 x 10	0.992	0.926	0.050	33	0.337	98
30 x 30 x 15	0.991	0.925	0.050	31	0.316	98
30 x 30 x 20	0.989	0.922	0.051	28	0.286	98
40 x 40 x 10	0.990	0.922	0.050	28	0.286	98
40 x 40 x 15	0.989	0.921	0.050	27	0.276	98
40 x 40 x 20	0.987	0.918	0.050	26	0.265	98
50 x 50 x 10	0.988	0.913	0.053	20	0.204	98
50 x 50 x 15	0.987	0.913	0.052	20	0.204	98
50 x 50 x 20	0.985	0.910	0.053	21	0.214	98
60 x 60 x 10	0.985	0.901	0.056	21	0.214	98
60 x 60 x 15	0.984	0.901	0.056	21	0.214	98
60 x 60 x 20	0.982	0.900	0.056	20	0.204	98
70 x 70 x 10	0.982	0.888	0.061	21	0.214	98
70 x 70 x 15	0.981	0.888	0.060	21	0.214	98
70 x 70 x 20	0.979	0.887	0.059	21	0.214	98
80 x 80 x 10	0.979	0.874	0.064	21	0.214	98
80 x 80 x 15	0.978	0.874	0.063	21	0.214	98
80 x 80 x 20	0.976	0.873	0.063	21	0.214	98
90 x 90 x 10	0.976	0.859	0.067	22	0.224	98
90 x 90 x 15	0.974	0.860	0.066	22	0.224	98
90 x 90 x 20	0.972	0.860	0.065	22	0.224	98
100 x 100 x 10	0.972	0.845	0.069	22	0.224	98
100 x 100 x 15	0.970	0.846	0.068	22	0.224	98
100 x 100 x 20	0.968	0.846	0.067	22	0.224	98

The slope of regression decreases with the increase in block size. There is no clear relationship between kriging efficiency and block sizes, however high kriging efficiencies (0.91) are associated with blocks of size 10m x 10m x Zm to 50m x 50m x Zm (Table 6.9). Smaller block sizes (10m x 10m x Zm to 20m x 20m x) and bigger block sizes (60m x 60m x Zm to 100m x 100m x Zm) have kriging variance of ≥ 0.06 whereas blocks of sizes between 30m x 30m x Zm and 50m x 50m x Zm have kriging variance below 0.06

All block sizes have negative weights with smaller block sizes (10m x 10m x Zm to 40m x 40m x Zm) having high proportions of negative weights ≥ 0.3 whereas block sizes (50m x 50m x Zm to 100m x 100m x Zm) have proportions of negative weights of about 0.2. Smaller blocks sizes have high number of negative weights because closer samples are being screened by distant samples within the second range. From these observations, the suitable block size for the tested area is likely to be one of the 50m x 50m x Zm block size.

All block sizes with proportions of negative weights ≥ 0.3 were ruled out from the search. For all the block sizes with low proportions of negative weights (≈ 0.2), blocks of size 50m x 50m x Zm stood out to be blocks of choice because of their high kriging efficiency and slope of regression, low proportions of negative weights as well as low kriging variance.

The block of size 50 m (Y) x 50m (X) x 15m (Z) with the slope of regression of 0.987, kriging efficiency of 0.913, a kriging variance of 0.052 and 0.204 proportions of negative weights was chosen as the optimal block size (Table 6.9). Variations of the X and Y dimensions of the 50m x 50m x 15m was done to see if any of the resultant blocks would give better results (Table 6.10 and 6.11).

Table 6.10: SURPAC: KN measures on variable Ym x 50m x 15m blocks.

Block size Ym x Xm x Zm	SLOR: Slope of regression	KE: Kriging efficiency	KV: Kriging variance	No. of negative Weights	Proportions of negative weights	No. of samples
10 x 50 x 15	0.989	0.915	0.056	30	0.306	98
20 x 50 x 15	0.989	0.917	0.053	29	0.296	98
30 x 50 x 15	0.988	0.918	0.052	27	0.276	98
40 x 50 x 15	0.988	0.917	0.051	25	0.255	98
50 x 50 x 15	0.987	0.913	0.052	20	0.204	98

Table 6.11: SURPAC: KN measures on 50m x variable Xm x 15m blocks.

Block size Ym x Xm x Zm	SLOR: Slope of regression	KE: Kriging Efficiency	KV: Kriging variance	No. of negative Weights	Proportion of negative weights	No. of samples
50 x 10 x 15	0.990	0.912	0.057	24	0.245	98
50 x 20 x 15	0.990	0.915	0.054	22	0.224	98
50 x 30 x 15	0.989	0.917	0.052	23	0.235	98
50 x 40 x 15	0.988	0.916	0.052	22	0.224	98
50 x 50 x 15	0.987	0.913	0.052	20	0.204	98

Blocks that were varied in the Y dimensions (Table 6.10) have higher proportions of negative weights compared to those varied in X dimension (Table 6.11). Negative weights are due to the drilling pattern; affecting the data configuration for the different block sizes; as well as the variogram model. All resultant blocks have slightly higher KE and SLOR, but they also have higher proportions of negative weights in comparison to those of the 50m x 50m x 15m block size.

The 50m x 50m x 15m was selected as the best block size for well-informed areas; the same as that identified for SUPERVISOR.

6.3.1.2 SURPAC: Well-informed blocks - Minimum and maximum number of samples

To assess the optimum number of samples to be used in the Kriging neighbourhood of a 50m x 50m x 15m block, the same process of increasing the number of samples in steps of two samples up to 98 samples, as used before was applied in SURPAC. The

SLOR, KE and kriging variance as well as the number and proportion of negative weights were recorded for each scenario and appear in Table C7 of Appendix C.

The numerical information in Table C7 makes it difficult to identify the minimum and maximum number of samples for the 50m x 50m x 15m block. A graph of the SLOR versus the number of samples is therefore plotted (Figure 6.16)

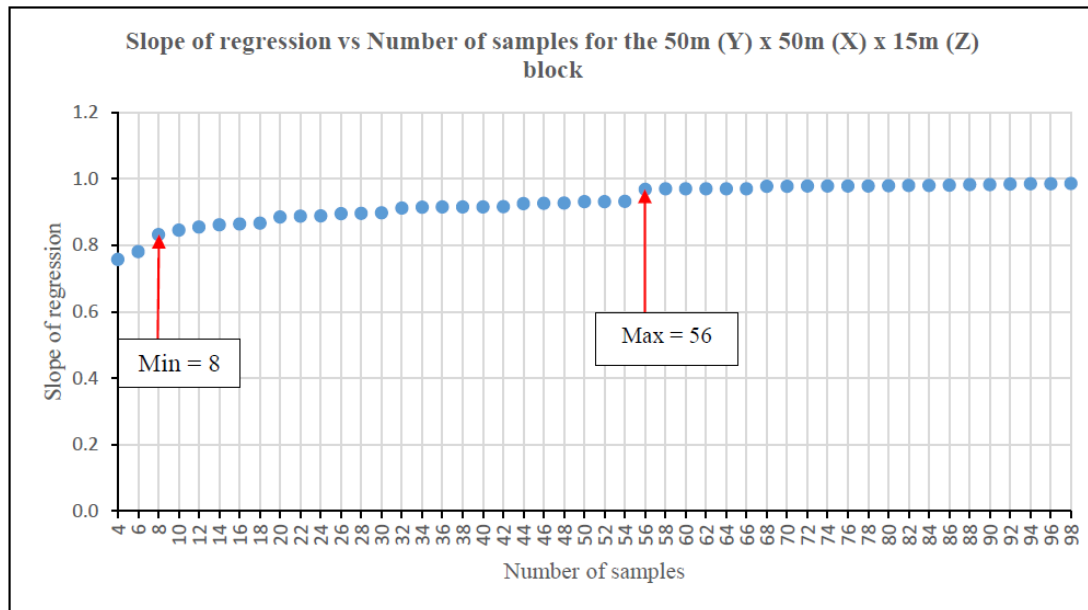


Figure 6.16: SURPAC: Number of samples vs SLOR for 50m x 50m x 15 m blocks.

The minimum number of samples is the lowest number of samples on the graph that corresponds to a significant increase in the SLOR. The SLOR increase with the increase in the number of samples, however, as the number of samples increase from 6 to 8, a significant increase in the SLOR is evident (Figure 6.16). Therefore, 8 is the minimum number of samples for the 50m x 50m x 15m block size.

There is a slight increase in the SLOR as the number of samples increase from 8 to 54 samples. From the 54th to the 56th sample, an increase in the SLOR is significant. (Figure 6.16). As the number of samples increase from 56 onwards, an increase in the SLOR is no longer significant. It is for those reasons that 56 is considered the maximum number of samples for the 50m x 50m x 15m block size for well-informed blocks.

6.3.1.3 SURPAC: Well-informed blocks – Optimum Search Range

Search ranges were considered by varying the semi-variogram range. The semi-variogram range was shortened to a range ratio of 0.60 and 0.80 and extended by a range ratio of 1.2 and 1.4 to see if it was too restrictive or not. Should better SLOR, KE and lower number of negative weights be obtained from the shortened search range, then it means that the semi-variogram is too long. An improvement in the SLOR and KE with additional positive weights beyond the semi-variogram range indicates that it is restrictive.

As explained earlier in the report, all correlated samples (98) were used in determining the optimal search range.

The SLOR, KE, kriging variance (KV) and negative weights for the five scenarios are shown in Table 6.12. RR = range ratio, a1 = range of the 1st structure, a2 = range of 2nd structure.

Table 6.12: SURPAC: Search range optimisation for 50m x 50m x 15m blocks

RR	a1	a2	SLOR	KE	KV	No. of negative weights	Proportion of negative weights	No. of samples
0.60	52.80	168.00	0.971	0.910	0.054	9	0.092	60
0.80	70.40	224.00	0.981	0.912	0.053	20	0.204	84
1.00	88.00	280.00	0.987	0.913	0.052	20	0.204	98
1.20	105.60	336.00	0.987	0.913	0.052	20	0.204	98
1.40	123.20	392.00	0.987	0.913	0.052	20	0.204	98

The KE and SLOR improve with the increase in the search range up to the semi-variogram range. The first two scenarios recorded low KE and SLOR because not all correlated samples were included in the search range (Table 6.12). However, there is no significant increase in KE and KV of the first two scenarios and the KE and KV of the semi-variogram range. Increasing the search range beyond the semi-variogram range did not improve the SLOR, KE, KV and the proportion of negative weights. This is because all correlated samples were included in the search range. The semi-variogram range was chosen as the best search range.

6.3.1.4 SURPAC: Well-informed blocks – Optimum Block Discretisation

In SURPAC the maximum number of discretisation in any dimension is 5, the software does not allow any value higher than 5, therefore, only four scenarios could be tested as shown in Table 6.13. This limitation meant that it was impossible to test exactly the same discretisation options in SURPAC as in SUPERVISOR. The discretisation points with high SLOR, KE and low KV is considered optimum for the discretisation configurations considered Table 6.13.

Table 6.13: SURPAC: Discretisation Optimisation of well-informed blocks

Discretisation YxXxZ	Slope of regression	Kriging Efficiency	Kriging Variance	No. of negative weights	Proportion of negative weights
2 x 2 x 3	0.971	0.910	0.055	10	0.179
5 x 5 x 3	0.968	0.910	0.053	8	0.143
2 x 2 x 5	0.971	0.911	0.055	10	0.179
5 x 5 x 5	0.967	0.910	0.053	8	0.143

Higher discretisation configurations 5x5x3 and 5x5x5 are appealing compared to smaller discretisation points because of their better conditional bias measure. A 5x5x3 configuration has the highest SLOR, lowest KV and proportion of negative weight hence it is the suitable discretisation points for the 50m x 50m x 15m blocks.

Higher discretisation points improve conditional bias results, but they increase the computation time for estimation (Vann *et al.*, 2003; Abzalov, 2016). This might be the reason why discretisation configuration is limited to 5 in SURPAC.

6.3.2 SURPAC: Poorly-informed blocks

6.3.2.1 SURPAC: Poorly-informed blocks – Optimum block size

Poorly-informed blocks lack sampling information compared to well-informed blocks, because they do not have internal samples. Samples that will be used in estimating uranium grades for poorly-informed blocks would be from drill holes surrounding the blocks. This means that the tested block sizes should be bigger than the ones for well-

informed areas. Block sizes tested at Y = 119 941, X = 82 306, Z = 448 are 50m x 50m x Zm to 300m x 300m x Zm (Z = 10, 15, and 20m).

OK was performed on each tested block size with a restriction of six samples from each drill hole. Table 6.14 presents all the tested block sizes and conditional bias measures

Table 6.14: SURPAC: Conditional bias measures of poorly-informed block sizes tested.

Block size Ym x Xm x Zm	SLOR: Slope of regression	KE: Kriging efficiency	KV: Kriging variance	No. of negative weights	Proportion of negative weights	No. of samples
50 x 50 x 10	1.022	0.792	0.127	33	0.423	78
50 x 50 x 15	1.020	0.796	0.122	32	0.410	78
50 x 50 x 20	1.017	0.801	0.117	31	0.397	78
100 x 100 x 10	0.996	0.885	0.051	15	0.192	78
100 x 100 x 15	0.994	0.888	0.050	15	0.192	78
100 x 100 x 20	0.992	0.890	0.048	15	0.192	78
150 x 150 x 10	0.940	0.876	0.041	9	0.115	78
150 x 150 x 15	0.938	0.879	0.039	9	0.115	78
150 x 150 x 20	0.935	0.881	0.038	9	0.115	78
200 x 200 x 10	0.846	0.772	0.055	2	0.026	78
200 x 200 x 15	0.844	0.776	0.054	2	0.026	78
200 x 200 x 20	0.841	0.780	0.052	3	0.038	78
250 x 250 x 10	0.733	0.601	0.075	5	0.064	78
250 x 250 x 15	0.731	0.604	0.073	5	0.064	78
250 x 250 x 20	0.728	0.605	0.072	4	0.051	78
300 x 300 x 10	0.620	0.360	0.097	10	0.128	78
300 x 300 x 15	0.619	0.361	0.095	10	0.128	78
300 x 300 x 20	0.616	0.360	0.094	10	0.128	78

All block sizes tested have significant proportions of negative weights except for the 200m x 200m x Zm block sizes. Block sizes smaller than 200m x 200m x Zm have high SLOR and KE with considerable proportions of negative weights. Block sizes bigger than the 200m x 200m x Zm have lower SLOR and KE values with higher

proportion of negative weights compared to those of the 200m x 200m x Zm block sizes (Table 6.14).

The difference between SLOR and KE values of the 200m x 200m with different Z values is not significant hence the 200m x 200m block size with a Z = 15m (Rossing's current bench height) block sizes is optimal for poorly-informed blocks.

The optimality of the 200m x 200m x 15m was tested by varying its Y (Table 6.15) and X dimensions (Table 6.16).

Table 6.15: SURPAC: Results of KN measures on variable Ym x 200m x 15m blocks.

Block size Ym x Xm x Zm	Slope of regression	Kriging Efficiency	Kriging Variance	No. of negative Weights	Proportion of negative weights	No. of samples
50 x 200 x 15	0.951	0.835	0.062	9	0.115	78
100 x 200 x 15	0.934	0.861	0.046	9	0.115	78
150 x 200 x 15	0.898	0.843	0.044	5	0.064	78
200 x 200 x 15	0.844	0.776	0.054	2	0.026	78

Table 6.16: SURPAC: Results of KN measures on 200m x variable Xm x 15m blocks.

Block size Ym x Xm x Zm	Slope of Regression	Kriging efficiency	Kriging variance	No. of Negative weights	Proportion of negative weights	No. of samples
200 x 50 x 15	0.938	0.802	0.070	14	0.179	78
200 x 100 x 15	0.918	0.828	0.054	7	0.090	78
200 x 150 x 15	0.886	0.820	0.049	2	0.026	78
200 x 200 x 15	0.844	0.776	0.054	2	0.026	78

Variation of the Y dimension of the 200m x 200m x 15m block size confirmed that none of the blocks gave better conditional results than the 200m x 200m x 15m block size (Table 6.15). However, the 200m x 150m x 15m block size has the same proportion of negative weights as the 200m x 200m x 15m block size, but it has higher SLOR, KE and KV values (Table 6.16).

The optimal block size for poorly informed areas is 200m (Y) x 150m (X) x 15m (Z).

6.3.2.2 SURPAC: Poorly-informed blocks - Minimum and maximum number of samples

To determine the number of samples for minimum and the maximum number of samples for 200m x 150m x 15m block size, hundred samples were tested. The minimum and the maximum number of samples tested varied between two and hundred at intervals of two samples. The conditional bias measures of the minimum and maximum number of samples are shown in Table C8, Appendix C. The slope of regression and the number of samples were plotted to determine the minimum and maximum number of samples (Figure 6.17).

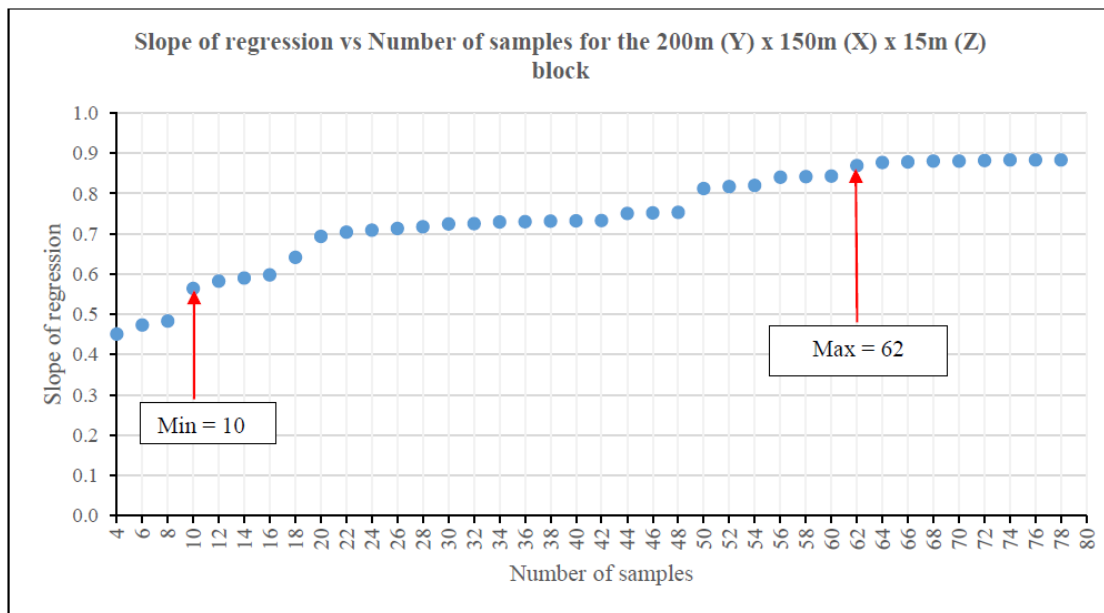


Figure 6.17: SURPAC: Optimisation of the minimum and maximum number of samples for the 200m x 150m x 15m blocks.

A significant increase in the slope of regression is recorded when the number of samples increase from 8 to 10 (Figure 6.17). An improvement from negative to positive KE values is also observed as the number of samples increase from 8 to 10 (Table C8). Based on these observations, the minimum number of samples for the 200m x 150m x 15m block size is 10.

There are three points on the plot (Figure 6.17) where an increase in the number of samples does not result in much improvement in the SLOR. These points are at the 20, 50 and 62 number of samples, and 62 was selected to be the maximum number of samples for the 200m x 150m x 15m block size for two reasons. The first reason is that the 200m x 150m x 15m block size is poorly informed and using a smaller value for the maximum number of sample would mean that there would not be enough samples to estimate the block values and this leads to conditional bias. Secondly, 62 number of samples has higher SLOR, KE and a lower KV values than the values for 20 and 50 number of samples (Table C8).

6.3.2.3 SURPAC: Poorly-informed blocks – Optimum search range

The same search ranges that were tested as for well-informed blocks were used. The tested scenarios and the corresponding conditional bias measures are shown in Table 6.17.

Table 6.17: SURPAC: Search range optimisation for 200m x 150m x 15m blocks

Range ratio	a1 (m)	a2 (m)	SLOR	KE	KV	No. of negative weights	Proportion of negative weights	No. of samples
0.60	52.80	168.00	0.847	0.802	0.055	2	0.033	60
0.80	70.40	224.00	0.886	0.820	0.049	2	0.026	78
1.00	88.00	280.00	0.886	0.820	0.049	2	0.026	78
1.20	105.60	336.00	0.886	0.820	0.049	2	0.026	78
1.40	123.20	392.00	0.886	0.820	0.049	2	0.026	78

An increase in the SLOR and the KE is observed from a search range ratio of 0.60 to 0.80 because not all correlated samples were included in the search (Table 6.17). After the 0.80 range ratio, the SLOR, KE, KV and the proportion of negative weights are the same. At a search range ratio of 0.8 all 78 correlated samples have been reached so there are no more samples beyond this range.

For poorly informed areas, the semi-variogram range is too long hence the suitable search range for the 200m x 150m x 15m block size is 0.8 times that of the semi-variogram range, with a1 = 70.40 m, a2 = 224.00 m.

6.3.2.4 SURPAC: Poorly-informed blocks – Optimum block discretisation

All compatible points in the X, Y and Z dimensions of the 200m x 150m x 15m were tested. Conditional bias results of the tested discretisation points are recorded in Table 6.18.

Table 6.18: SURPAC: Discretisation optimisation 200m x 150m x 15m blocks

Discretisation (Y x X x Z)	Slope of regression	Kriging Efficiency	Kriging Variance	No. of negative weights	Proportion of negative weights
2 x 2 x 3	0.899	0.776	0.086	1	0.016
4 x 3 x 3	0.865	0.818	0.054	1	0.016
5 x 5 x 3	0.859	0.824	0.050	0	0.016
2 x 2 x 5	0.898	0.779	0.084	1	0.016
4 x 3 x 5	0.865	0.819	0.053	1	0.016
5 x 5 x 5	0.858	0.824	0.049	0	0.016

The only discretisation points, from Table 6.18, that is capable of minimising conditional bias in the 200m x 150m x 15m blocks is the 5x5x3 because it has the lowest kriging variance and the highest kriging efficiency.

6.4 Differences between the KNA processes in SUPERVISOR and SURPAC

The KNA steps in SUPERVISOR and SURPAC are the same, there are however a few technical differences in the KNA application:

- The KNA process in SURPAC takes longer than the KNA process in SUPERVISOR. This is simply because SUPERVISOR allows the KNA to be run on multiple test blocks whereas SURPAC allows only one block to be tested at a time.
- Another difference is in maximum number of discretisation points in any direction, the highest discretisation number in SUPERVISOR is ten while it is five in SURPAC.
- In addition to the slope of regression, kriging efficiency and the number of negative weights, SURPAC also provides the kriging variance as opposed to SUPERVISOR which provides the variance of the estimate.

7 Results and Discussions

7.1 Introduction

In this chapter the outcomes of the KNA procedures obtained from SUPERVISOR and SURPAC are discussed. The KN parameters were used as input to estimate block model U_3O_8 grades and determine the mean U_3O_8 grades of the domains. The estimated block model U_3O_8 grades were validated by comparing them with the composited U_3O_8 grades.

7.2 KN parameters

The KN parameters as identified for the tested block; well-informed and poorly-informed blocks in Supervisor and SURPAC are summarised in Table 7.1.

Table 7.1: KN parameters from SUPERVISOR and SURPAC.

	Well-informed blocks		Poorly-informed blocks	
	SUPERVISOR	SURPAC	SUPERVISOR	SURPAC
Block size	50 x 50 x 15	50 x 50 x 15	150 x 200 x 15	150 x 200 x 15
Min no. of samples	8	8	10	10
Max no. of samples	56	56	62	62
Search range (dir1, dir2, dir3)	280, 198, 78	280, 198, 78	224, 158.4, 62.4	224, 158.4, 62.4
Discretisation	5 x 5 x 3	5 x 5 x 3	5 x 5 x 3	5 x 5 x 3

Both software packages resulted in the same kriging parameters being selected after performing the KNA procedures for both the well-informed and poorly-informed blocks. A 50m x 50m x 15m block size was identified as the optimal block size, for well-informed areas, with a minimum of 8 samples and a maximum of 56 samples. As expected, the search range test confirmed that the semi-variogram range is suitable for well-informed blocks with the 5x5x3 discretisation.

A block of size 150m x 200m x 15m with a minimum of 10 and a maximum of 62 samples at a search range 80% of the semi-variogram range was considered optimal for poorly-informed blocks. The maximum number of 62 samples were included when the search range was 80% of the semi-variogram range. The 0.8 range ratio of the semi-

variogram is equivalent to a major axis distance of 224.00 m, a semi-major distance of 158.40 m and minor distance of 62.40 m. Discretisation points suitable for the poorly-informed blocks were found to be 5x5x3.

Well-informed blocks have a minimum of 8 and a maximum of 56 samples whereas poorly-informed blocks have a minimum number of 10 samples and a maximum of 62 samples. In the final kriging a maximum of 48 samples were used in estimating block grades for well-informed blocks while 52 samples are needed to estimate poorly-informed blocks. For poorly-informed blocks the maximum no of samples increased by 4 compared to the well-informed blocks. A much larger search neighbourhood is required for poorly-informed block as opposed to well-informed blocks because there are fewer samples close to the blocks being estimated.

7.3 KN statistics

KN parameters (Table 7.1) were used as input to carry out an Ordinary Kriging estimation U_3O_8 grade block model on the selected two block sizes. The conditional bias measures of the tested block configurations are presented in Table 7.2.

Table 7.2: KN statistics from SUPERVISOR and SURPAC

	Well-informed blocks		Poorly-informed blocks	
	SUPERVISOR	SURPAC	SUPERVISOR	SURPAC
Slope of regression	0.968	0.968	0.855	0.859
Kriging efficiency	0.910	0.910	0.821	0.824
No. of negative weights	8	8	1	0
Proportion of negative weights	0.143	0.143	0.016	0.000

Since the same test locations, input data, semi-variogram parameters, and block configurations were used to estimate block model U_3O_8 grades, the resultant conditional bias measures from the KNA procedures are expected to be the same.

It is clear from Table 7.2 that both software packages lead to the same KN conditional bias measures for well-informed blocks, and slightly different measures for poorly-informed blocks; SURPAC performed slightly better in this case.

For poorly-informed blocks, slight differences have been observed in every step of the KNA process:

- i. Variation of the 200m x 200m x 15m block in the Y and X dimensions in SUPERVISOR resulted in blocks with lower KE and SLOR Table 6.5 and Table 6.6 respectively values compared to the results of the same blocks obtained from SURPAC Table 6.15 and Table 6.16 respectively.
- ii. During the search range step, differences in the SLOR and KE results from the two software packages were observed (Table 6.7 and 6.17).
- iii. For discretisation configuration of the 200m x 150m x 15m block size, differences are evident in Table 6.18 and Table C6, Appendix C.

As the same input data, semi-variogram parameters, estimation method (OK), block locations and configurations were used in both software packages, these differences might be caused by the manner in which samples are selected for inclusion in the kriging neighbourhood analysis.

A restriction of a maximum of six samples per drill hole was introduced. The way SURPAC choose the six samples to include in the KNA might be different from how SUPERVISOR does it and this could be the reason why slight differences in the KNA statistics of poorly informed blocks were observed. Unfortunately, the sample selection algorithm for estimation is not explained by either of the two software packages.

As to be expected, well-informed blocks have better KN parameters for the minimisation of the conditional bias than the poorly-informed blocks. A better correlation between estimated and actual block U_3O_8 grades is expected for well-informed blocks than for the poorly-informed blocks, because of the abundance of data in this kriging neighbourhood.

7.4 Block model

SURPAC version 6.7.1 (x64bit) was used to create a block model with the extents presented in Table 7.3. From the KNA an estimation block size of 50m x 50m x 15m was used. The block model was constrained below the topographic surface which was created using drillhole collars. Attributes that were included in the block model output are; estimated U_3O_8 assay grade, number of samples used, number of drill holes used, block variance, kriging efficiency, number of negative weights, conditional bias slope, LaGrange multiplier and an alaskite specific gravity of 2.6 t/m³.

Table 7.3: Block model extents and size

	Easting (X)	Northing (Y)	Elevation (Z)
Minimum coordinates (m)	81 870	119 790	70
Maximum coordinates (m)	82 440	120 650	530
User block size (m)	50	50	15

7.5 Grade estimation

7.5.1 OK estimates

The Ordinary kriging estimation carried out in SURPAC is based on the semi-variogram parameters (Table 5.10) and the optimum Kriging Neighbourhood parameters identified from the KNA process above; these are the search orientation and associated search ranges as well as the minimum and maximum no of samples and the maximum number of samples per borehole (Table 7.4). The discretisation of 5 x 5 x 3 were used to estimate U_3O_8 grades into the block model.

Table 7.4: Sample search parameters

Domains	Search orientation			Search Range (m)			Number of samples		
	Bearing	Plunge	Dip	Dir 1	Dir 2	Dir 3	Min	Max	Max/hole
All	130°	0°	30°	280	198	78	8	56	6

7.5.2 Validation

Block model estimation for all the domains and block model validation was done using SURPAC. The OK block estimates were validated using the following two methods:

1. Comparison of the mean U_3O_8 grades of the domains with the composited mean U_3O_8 grade and declustered mean U_3O_8 grade.
2. Graphical comparison of the easting (X-coordinate), northing (Y-coordinate), and elevation (Z coordinate) sections of the block model with the 1.5m composited U_3O_8 grade.

Table 7.5: Comparison of block model U_3O_8 averages with composited and declustered mean U_3O_8 averages

Domain	Block model average grade (ppm)	Composite mean (ppm)	Declustered mean (ppm)	Difference (%)	
				BM to composite mean	BM to declustered mean
1	104.06	117.93	120.28	-13	-16
2	100.09	102.12	103.73	-2	-4
3	122.96	154.38	155.62	-26	-27
4	250.64	233.52	238.17	7	5
5	218.98	269.34	270.76	-23	-24
6	207.78	231.68	230.76	-12	-11
7	155.06	138.02	137.01	11	12
Combined	175.94	207.38	209.38	-18	-19

In Domain 2 and 4 the differences between the block U_3O_8 estimates and the composited and declustered U_3O_8 grades are the least. For the remainder of the domains the differences between the block U_3O_8 estimates and the composited and declustered U_3O_8 grades (Table 7.5) are substantial. The differences between block grades and composited as well as declustered grades are due to the limitation of a maximum of six samples per drill hole during the estimation of block grades. Etango Uranium Project and Husab deposit technical reports also recorded differences in the mean grades between block estimates and composite samples as high as -27 %.

The swath plots generally show good correlation between block U_3O_8 grades and estimated U_3O_8 grades (Figures 7.1, 7.2 and 7.3). Significant differences in mean are observed on the easting swath plot (Figure 7.1) as well as on the northing swath plot

(Figure 7.2). These differences are due to the effect of limiting six samples per drill hole. The deviation in mean around 120550N (Figure 7.2) is caused by the edge of the extent of the block model.

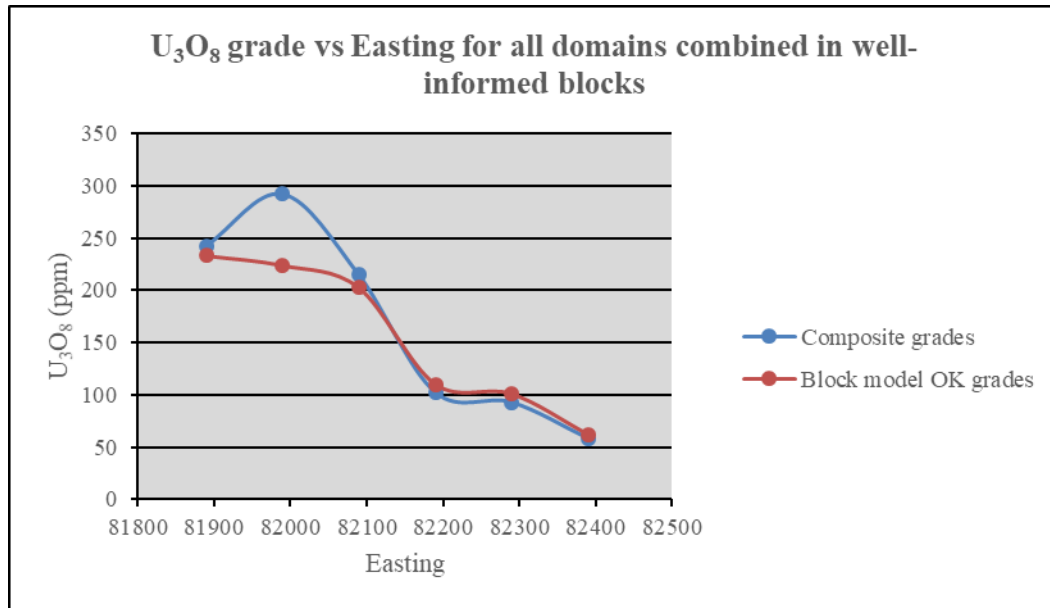


Figure 7.1: Swath plot by Easting for U₃O₈ grades for all domains combined in well-informed blocks.

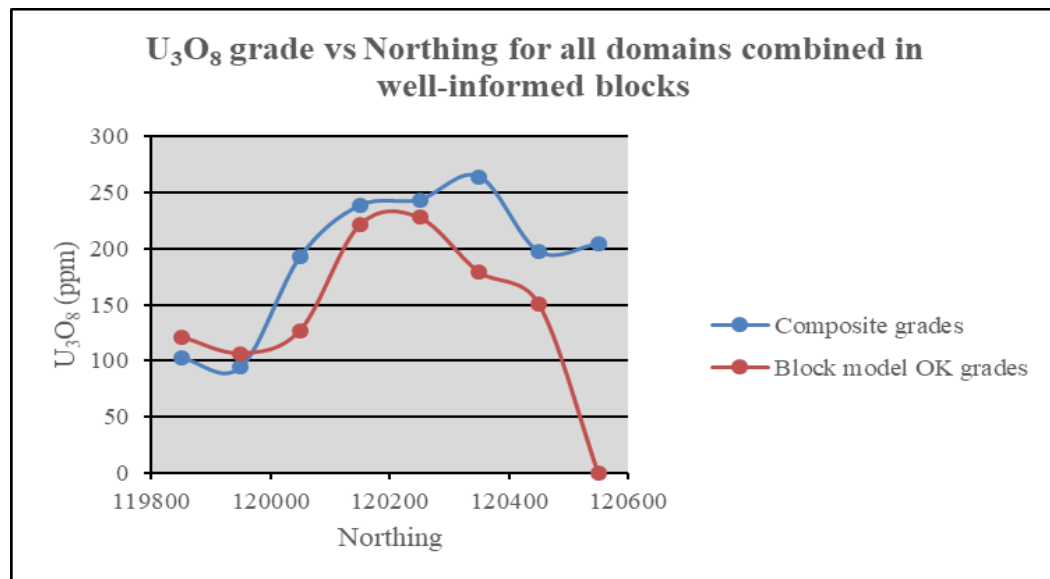


Figure 7.2: Swath plot by Northing for U₃O₈ grades for all domains combined in well-informed blocks.

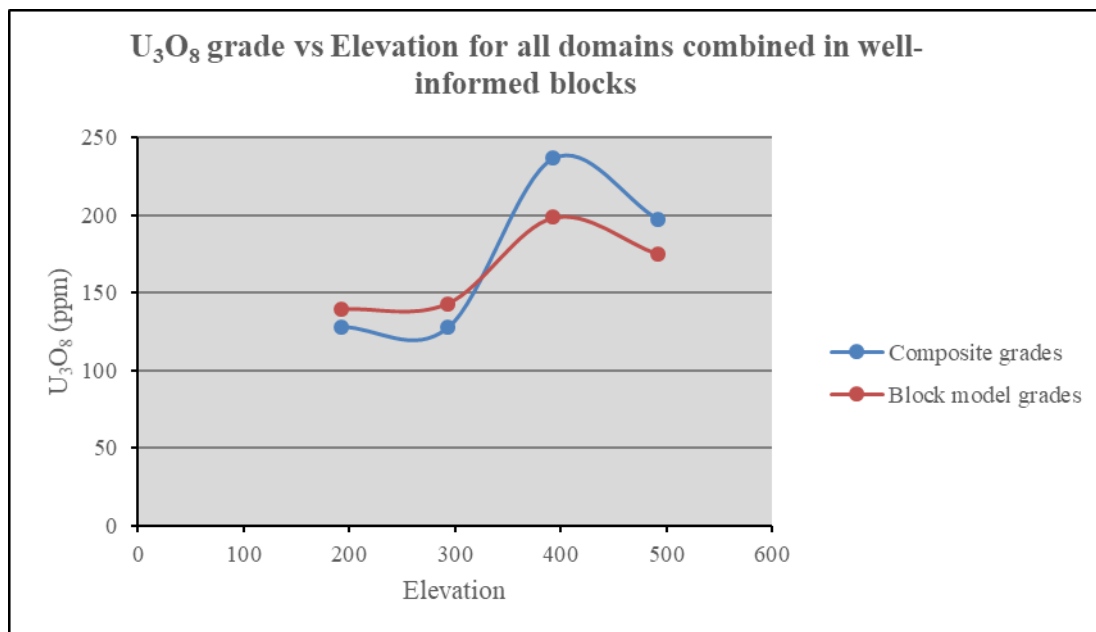


Figure 7.3: Swath plot by Elevation for U₃O₈ grades for all domains combined in well-informed blocks.

8 Conclusions

The kriging search neighbourhood has a huge influence on the accuracy of the kriging estimate. It is common practice to perform a Kriging Neighbourhood Analysis to identify the kriging parameters to be used to minimise the conditional bias which arises from the inherent smoothing of the kriging estimation methodology.

The impact of conditional bias and the need to minimise it by carrying out a Kriging Neighbourhood Analysis gave rise to the research question:” *Is there a significant difference in KN parameters obtained from using the SUPERVISOR KNA procedure and those derived from using the SURPAC KNA procedure when identical input data, semi-variogram model parameters and test locations are used?*” and thereby fill the gap in the literature for a comparison of the of Kriging Neighbourhood Analysis outcomes using different software packages.

For comparison sake, all relevant information required to carry out the KNA were identical in the application of both SUPERVISOR and SURPAC. They are; geological domains, semi-variogram parameters, block configurations, and test locations. Block configurations tested were for well-informed blocks (with three drill holes inside and others surrounding the blocks) and poorly-informed blocks (with one internal drill hole and surrounding drill holes).

The KN parameters compared were; block size, minimum and maximum number of samples, search range, and discretisation points. Scenarios of various block sizes, minimum and maximum number of samples ranging from 1 to 98 samples for well-informed blocks and 1 to 78 for poorly-informed blocks, search ranges shorter and longer than the semi-variogram range, and a series of discretisation points were tested.

The slope of regression, kriging efficiency, and the number of negative weights were the criteria used to identify the best neighbourhood. In addition to those criteria, kriging variance was also considered in SURPAC because SUPERVISOR does not provide it. For each scenario, the neighbourhood with the highest slope of regression, kriging

efficiency and lowest kriging variance and the least number of negative weights were identified.

Both software packages gave the same KN results for the tested data configurations. The KN results for well-informed blocks were found to be; a block of size 50m x 50m x 15m, a minimum of 8 and a maximum of 56 samples, a semi-variogram search ranges of $a_1 = 280\text{m}$, $a_2 = 198\text{m}$, and $a_3 = 78\text{m}$ with a 5x5x3 block discretisation.

The 150m x 200m x 15m block size was identified to be optimal for poorly-informed areas with a minimum of 10 and a maximum of 62 samples. The optimal search range for poorly-informed blocks was found to be 20% shorter than the semi-variogram range ($a_1 = 224\text{ m}$, $a_2 = 158.40\text{ m}$ and $a_3 = 62.40\text{ m}$) with a 5x5x3 block discretisation.

The conditional bias measures (the slope of regression, kriging efficiency and number of negative weights) for well-informed areas in both software packages were the same. However, for poorly-informed blocks, the conditional bias measures were slightly different. SURPAC gave better conditional bias measure than SUPERVISOR. The differences in conditional bias measures might be caused by the way each software choose the samples for estimation. It is not explained how each software package choose samples for estimation.

Validation of the block estimates was done by visual comparison of the block and composite average grades as well as by graphical comparison of the X, Y, and Z dimensions of the block. The correlation between visual comparison of the block U_3O_8 average grades and composite U_3O_8 grades for most domains (Table 7.5) was generally poor due to the restriction of six samples per drill hole. Graphical comparison of the sections of the blocks (in terms of easting, northing and elevation) with the composite U_3O_8 grade gave better correlations (Figures 7.1, 7.2 and 7.3).

Whilst there are technical differences between SUPERVISOR and SURPAC which is to be expected as SUPERVISOR is a geostatistical software package and SURPAC a geology, resource estimation and mine planning software, but these differences did not form part of the scope of this research study. Due to inherent differences between them,

some aspects of the research were carried out in more relevant software and the output applied in the other software, this was done to ensure that identical input was used in the KNA procedures of the two software packages.

The geological modelling was done in SURPAC and the outcome and conclusion were applied in SUPERVISOR as were the identification of the locations of well-informed and poorly informed blocks.

The variography analyses were carried out in SUPERVISOR and the variogram models derived from that exercise were also applied in SURPAC.

A difference in the KNA processes in the two software packages is the maximum number of discretisation points. SUPERVISOR allows a maximum of 10 points in the block X, Y and Z dimension whereas SURPAC allows a maximum of 5. SURPAC also provides the kriging variance in addition to other conditional bias measures while SUPERVISOR does not give kriging variance.

Most of the general studies that have been considered in the literature review on the comparison of software packages revealed that at least one software package is better than the other in terms of result accuracy.

This study reveals that although the difference in accuracy (caused by the methodology of data selection in each software) is small, SURPAC gave better conditional bias measures for poorly-informed blocks than SUPERVISOR.

This study concludes that there is no significant difference in the KNA results produced by SURPAC and SUPERVISOR as the same kriging neighbourhood parameters were identified in both packages. The outcomes were the same for well-informed blocks with a negligible difference for the poorly-informed blocks thus the research question has been answered.

9 References

Abraham, I.M. (2009) *Geology and spatial distribution of uranium mineralisation in the SK anomaly area, Rossing area, Namibia*. MSc dissertation. [Online]. Johannesburg, South Africa, University of the Witwatersrand. Available from: <http://wiredspace.wits.ac.za/handle/10539/8283> [Accessed: 19 August 2016].

Abzalov, M. (2016) *Applied Mining Geology*. 1st ed. 2016 edition. New York, NY, Springer.

Anderson, H. & Nash, C. (1997) Integrated lithostructural mapping of the Rossing area, Namibia using high resolution aeromagnetic, radiometric, Landsat data and aerial photographs. *Exploration Geophysics*. [Online] 28 (2), 185. Available from: doi:10.1071/EG997185.

Armstrong, M. (1998) *Basic Linear Geostatistics*. 1998 edition. Berlin ; New York, Springer.

Babish, G. (2000) *Geostatistics without tears: a practical guide to geostatistics, variograms and kriging*. [Online]. Available from: https://www.google.co.za/search?dcr=0&source=hp&ei=tggbWu6GJYGxU_erhOAF&q=geostatistics+without+tears&oq=geostatistics+wit&gs_l=psy-ab.1.0.35i39k1j0l4j0i22i30k1l5.1375.5870.0.7864.18.12.0.0.0.0.778.2695.2-1j2j1j1j1.6.0....0...1.1.64.psy-ab..12.6.2690.0..0i67k1.0.bSb_As6USDA [Accessed: 26 November 2017].

Bowden, P., Herd, D. & Kinnaird, J.A. (1995) The significance of uranium and thorium concentrations in pegmatitic leucogranites (alaskites), Rossing mine, Swakopmund, Namibia. *Communications of the Geological Survey of Namibia*. 10, 43–49.

Boyle, C. (2010) Kriging neighbourhood analysis by slope of regression and weight of mean – evaluation with the Jura data set. *Mining Technology*. [Online] 119 (2), 49–58. Available from: doi:10.1179/037178410X12741755140804.

Chanderman, L. (2015) *3D geological modelling and mineral resource estimate for the FE2 gold deposit, Sadiola mine, Mali*. MSc research report. [Online]. Johannesburg, South Africa, University of the Witwatersrand. Available from: <http://wiredspace.wits.ac.za/handle/10539/20173> [Accessed: 14 November 2016].

Clark, I. & Garnett, R.H.T. (1974) *Identification of multiple mineralization phase by statistical methods*. 83, A43.

Coombes, J. & Boamah, P. (2015) Local Kriging Neighbourhood Optimisation. In: *Africa Australia Technical Mining Conference*. 2015 Adelaide, Australia.

Corner, B. & Henthorn, D.I. (1978) *Results of a palaeomagnetic survey undertaken in the Damara Mobile Belt, South West Africa, with special reference to the magnetisation of the uraniferous pegmatitic granites*. [Online]. Available from: http://www.iaea.org/inis/collection/nclcollectionstore/_public/10/429/10429919.pdf [Accessed: 11 November 2016].

Darling, P. & Noble, A.C. (2011) Mineral Resource Estimation. In: *SME Mining Engineering Handbook*. 3 edition. Society for Mining, Metallurgy, and Exploration. pp. 203–217.

Dassault Systemes Geovia (2015) *Surpac Version 6.7.1*. Vancouver, British Columbia, Canada.

Dassault Systemes GEOVIA (2014) *Geological Database Tutorial - Manual*.

Dassault Systemes GEOVIA (2015) *Geostatistics Tutorial- Manual*.

Deutsch, C.V. (2007) The Slope of Regression for Kriging Estimators. *Centre for Computational Geostatistics, Department of Civil and Environmental Engineering, University of Alberta*. 311-1-311–314.

Deutsch, J.L., Szymanski, J. & Deutsch, C.V. (2014) Checks and measures of performance for kriging estimates. *Journal of the Southern African Institute of Mining and Metallurgy*. 114 (3), 223–223.

Dionne, Y. (2015) *9 Aspects To Consider When Choosing A Mining And Geology Software*. [Online]. 2015. Available from: <http://www.promine.com/blog/9-aspects-to-consider-when-choosing-a-mining-and-geology-software> [Accessed: 2 August 2016].

Dohm, C., E. (2015a) *MINN 7006 Geostatistical methods in mineral resource evaluation lecture notes*. University of the Witwatersrand, Johannesburg.

Dohm, C., E. (2015b) *MINN 7007 Statistical evaluation of mineral resources lecture notes*. University of the Witwatersrand, Johannesburg.

Fischer, M.M. & Getis, A. (2009) Google-Books-ID: PuzojwEACAAJ. *Handbook of Applied Spatial Analysis: Software Tools, Methods and Applications*. Springer Berlin Heidelberg.

Glacken, I.M. & Snowden, D.V. (2001) Mineral Resource Estimation. *ResearchGate*. [Online] 23. Available from: https://www.researchgate.net/publication/264869156_Mineral_Resource_Estimation [Accessed: 19 August 2016].

Goovaerts, P. (1997) *Geostatistics for Natural Resources Evaluation*. 1 edition. New York, Oxford University Press.

Hustrulid, W.A. & Kuchta, M. (2006) *Open Pit Mine Planning and Design, Two Volume Set, Second Edition*. 2 edition. London, CRC Press.

Inwood, N. & Hill, M. (2009) *National Instrument 43-101 Technical Report, Rossing South- January 2009 Resource Update, Husab Project, Namibia*.

Isaaks, E.H. & Srivastava, R.M. (1989) *An Introduction to Applied Geostatistics*. 1 edition. New York, Oxford University Press.

Journel, A.G. & Huijbregts, C.J. (1978) *Mining Geostatistics*. London, Academic Press Inc.

Khakestar, M.S., Madani, H., Hassani, H. & Moarefvand, P. (2013) Determining the best search neighbourhood in reserve estimation, using geostatistical method: A case study anomaly No 12A iron deposit in central Iran. *Journal of the Geological Society of India*. [Online] 81 (4), 581–585. Available from: doi:10.1007/s12594-013-0074-y.

Kinnaird, J., Nex, P. & Freemantle, G. (2011) The occurrence of betafite in pegmatitic sheeted uraniferous granites of the central zone of the Damara Orogen. In: *Contributions to the 5th international symposium on granitic pegmatites*. [Online]. 2011 Mendoza, Argentina, Geological Association of Argentina. pp. 115–117. Available from: http://www.minsocam.org/MSA/Special/Pig/PIG_articles/PEG2011_Contributions.pdf [Accessed: 10 August 2017].

Kinnaird, J.A. & Nex, P.A.M. (2007) A review of geological controls on uranium mineralisation in sheeted leucogranites within the Damara Orogen, Namibia. *Applied Earth Science*. [Online] 116 (2), 68–85. Available from: doi:10.1179/174327507X167091.

Krige, D.G. (1996a) A basic perspective on the roles of classical statistics, data search routines, conditional biases and information and smoothing effects in ore block valuations. In: *Proceedings of the Regional APCOM, Slovenia*. [Online]. 1996 p. Available from: <http://saimm.org.za/Conferences/DanieKrige/DGK41.pdf> [Accessed: 18 November 2016].

Krige, D.G. (1996b) *A practical analysis of the effects of spatial structure and of data available and accessed, on conditional biases in ordinary kriging*. In: 1996 Wollongong, Australia. pp. 799–810.

MacCormack, K.E. & Eyles, C.H. (2012) Assessing the impact of program selection on the accuracy of 3D geologic models. *Geosphere*. [Online] 8 (2), 534–543. Available from: doi:10.1130/GES00732.1.

Marenica Energy Limited (2010) *Marenica Uranium Project Namibia*. [Online]. Available from: <http://webcache.googleusercontent.com/search?q=cache:zkCull2whnQJ:webcast.viosstream.com/Download.axd%3Fviocast%3D3015%26auth%3Dbbe7e990-e36f-473a-9b77-e0d5b09264c8%26type%3DDeckPDF%26deck%3D1192+&cd=8&hl=en&ct=clnk&gl=za> [Accessed: 19 June 2017].

Mining Associates (2013) *Resource estimation and SurpacTM*. [Online]. Available from: https://www.google.com/search?sclient=psy-ab&client=firefox-b-ab&q=Resource+estimation+ans+surpac+mining+associates&oq=Resource+estimation+ans+surpac+mining+associates&gs_l=serp.3...35406.40451.0.41761.18.18.0.0.0.461.5038.2-15j2j1.18.0....0...1c.1.64.psy-ab..0.14.3869...33i160k1j33i21k1.AkpwzHAIhsw&pbx=1&bav=on.2,or.&bvm=bv.139250283,d.ZGg&biw=1525&bih=360&ech=1&psi=uZ0uWIy2IObBgAb-5KvQDw.1479450051826.3&ei=uZ0uWIy2IObBgAb-5KvQDw&emsg=NCSR&noj=1&gfe_rd=cr [Accessed: 18 November 2016].

Ministry of Mines and Energy (2010) *Strategic Environmental Assessment for the central Namib Uranium Rush*. [Online]. Available from: <http://www.saiea.com/uranium/14chap4march2011.pdf> [Accessed: 17 June 2017].

Mucha, A.J. (2013) *Statistical and geostatistical methods of resources estimation of bituminous coals and lignites deposits*. In: [Online]. 15 October 2013 Ostrava, Czech Republic. p. Available from: https://www.researchgate.net/profile/Jacek_Mucha/publication/259952477_Statistical_and_geostatistical_methods_of_resources_estimation_of_bituminous_coals_and_lignites_deposits/links/00b7d52eac5366a84c000000/Statistical-and-geostatistical-methods-of-resources-estimation-of-bituminous-coals-and-lignites-deposits.pdf [Accessed: 25 August 2017].

Neef, C., Jelliffe, R.W., van Laar, T., Loohuis, T., (1994) MEDINFO'92 Satellite Open Conference on Estimation and Control of Patient Variability for Individualized Antibiotic Therapy Comparison of two software programs to be used for the calculation of population pharmacokinetic parameters. *International Journal of Bio-Medical Computing*. [Online] 36 (1), 143–150. Available from: doi:10.1016/0020-7101(94)90107-4.

Nex, P.A.M. & Kinnaird, J.A. (1995) Granites and their mineralisation in the Swakop River area around Goanikontes, Namibia. *Communications of the geological survey of Namibia*. 10, 51–56.

Nex, P.A.M., Kinnaird, J.A. & Oliver, G.J.H. (2001) Petrology, geochemistry and uranium mineralisation of post-collisional magmatism around Goanikontes, southern

- Central Zone, Damaran Orogen, Namibia. *Journal of African Earth Sciences*. [Online] 33, 481–502. Available from: doi:10.1016/S0899-5362(01)00096-3.
- Pyrcz, M.J. & Deutsch, C.V. (2003) Declustering and debiasing. *Newsletter*. 19, 25.
- Rendu, J.-M. (1978) *An introduction to geostatistical methods of mineral evaluation* /. Johannesburg :, South African Institute of Mining and Metallurgy,.
- Rivoirard, J. (1990) Introduction to disjunctive kriging and nonlinear geostatistics. *Centre de geostatistique*. 89.
- Rivoirard, J. (1987) Two key parameters when choosing the kriging neighborhood. *Mathematical geology*. 19 (8), 851–856.
- Roesener, H. & Schreuder, C. (1992) *Uranium, 7.1-1-7.1-55*. In: *Geological Survey of Namibia. The Mineral Resources of Namibia, 1.1-1-9-17*. [Online]. Available from: https://www.google.com/search?q=uranium-geological+survey&ie=utf-8&oe=utf-8&client=firefox-b-ab&gfe_rd=cr&ei=9PgYWI2nI-2o8wfrR9r3wBA [Accessed: 1 November 2016].
- Rossi, M.E. & Deutsch, C.V. (2013) *Mineral Resource Estimation*. 2014 edition. New York, Springer.
- Rossing Uranium Ltd (2003) *Report to Stakeholders, Rossing Uranium Ltd*. [Online]. Available from: <http://www.rossing.com/files/publications/2003-rossing-report-to-stakeholders.pdf> [Accessed: 25 April 2017].
- Rossing Uranium Ltd (2016) *Report to Stakeholders, Rossing Uranium Ltd*. [Online]. Available from: http://www.rossing.com/files/rossing_stakeholder_report2016.pdf [Accessed: 25 April 2017].
- Rossing Uranium Website (2016) *Rössing Shareholders*. [Online]. 2016. Available from: <http://www.rossing.com/shareholding.htm> [Accessed: 2 November 2016].
- Rundel, C. (2012) *Lecture 9: Plots, Graphs, and Pictures - Statistics 10*. [Online]. Available from: <https://www2.stat.duke.edu/courses/Spring12/sta10.1/Lectures/Lec9.pdf> [Accessed: 2 November 2017].
- Sarma, D.D. (2009) *Geostatistics with Applications in Earth Sciences*. 2nd ed. 2009 edition. Dordrecht, Springer.
- Sides, E.J. (1997) Geological modelling of mineral deposits for prediction in mining. *Geologische Rundschau*. [Online] 86 (2), 342–353. Available from: doi:10.1007/s005310050145.

Sinclair, A.J. & Blackwell, G.H. (2002) *Applied Mineral Inventory Estimation*. Cambridge ; New York, Cambridge University Press.

Snowden Mining (2013) *Supervisor V8 Demo 2013*. [Online]. Available from: <https://www.youtube.com/watch?v=l6E9DUukDdE> [Accessed: 18 November 2016].

Stephenson, P.R. & Vann, J. (2001) Common sense and good communication in Mineral Resource and Ore Reserve estimation. *ResearchGate*. [Online] 23. Available from: https://www.researchgate.net/publication/253729084_Common_sense_and_good_communication_in_Mineral_Resource_and_Ore_Reserve_estimation [Accessed: 21 August 2016].

Supervisor (2017) *Supervisor for windows. [Software] Version 8*. Perth, Australia, Snowden.

Szidarovszky, F., Baafi, E.Y. & Kim, Y.C. (1987) Kriging without negative weights. *Mathematical geology*. 19 (6), 549–559.

Vann, J. & Guibal, D. (1998) Beyond Ordinary Kriging—An overview of non-linear estimation. In: *Proceedings of a one day symposium: Beyond Ordinary Kriging*. [Online]. 1998 p. Available from: <http://gaa.org.au/pdf/bok%20keynote.pdf> [Accessed: 24 August 2017].

Vann, J., Jackson, S. & Bertoli, O. (2003) Quantitative kriging neighbourhood analysis for the mining Geologist-A description of the method with worked case examples. *The Australasian Institute of Mining and Metallurgy*. 1–10.

Wahbeh, A., H., Al-Radaideh, Q., A., Al-Kabi, M., N. & Al-Shawakfa, E., M. (2011) *a comparison study between some data mining tools over some classification methods - Google Search*. [Online]. 2011. Available from: https://www.google.co.za/?client=firefox-b#q=a+comparison+study+between+some+data+mining+tools+over+some+classification+methods&gfe_rd=cr [Accessed: 12 August 2016].

Yamamoto, J.K. (2000) An alternative measure of the reliability of ordinary kriging estimates. *Mathematical Geology*. 32 (4), 489–509.

Yamamoto, J.K. (2005) Correcting the Smoothing Effect of Ordinary Kriging Estimates. *Mathematical Geology*. [Online] 37 (1), 69–94. Available from: doi:10.1007/s11004-005-8748-7.

Youssef, B.A., Rashwan, S. & Redwan, E.M. (2014) Comparison of Two Academic Software Packages for Protein Structure Prediction. *Journal of Bioinformatics and Intelligent Control*. [Online] 3 (2), 143–146. Available from: doi:10.1166/jbic.2014.1074.

Appendices

Appendix A: Grade Capping

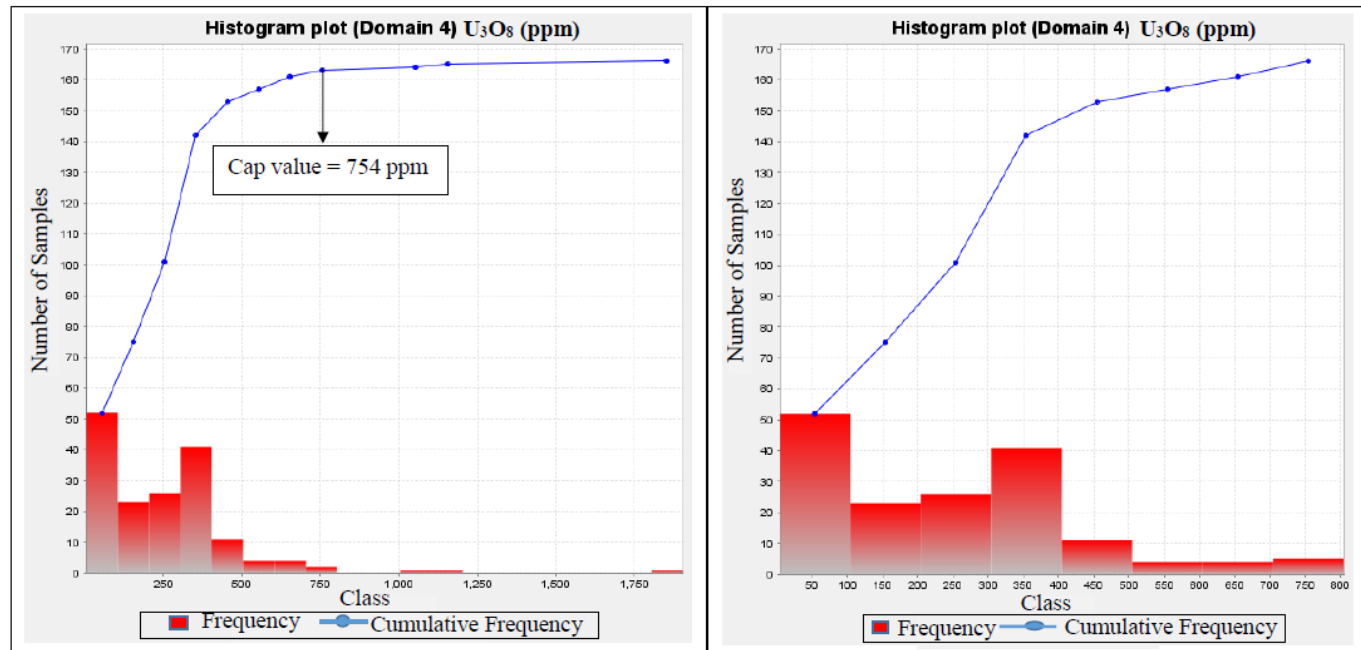


Figure A1: Domain 4 U₃O₈ (ppm) histogram, and a cumulative frequency plot of 1.5m composites before (left) and after (right) capping.

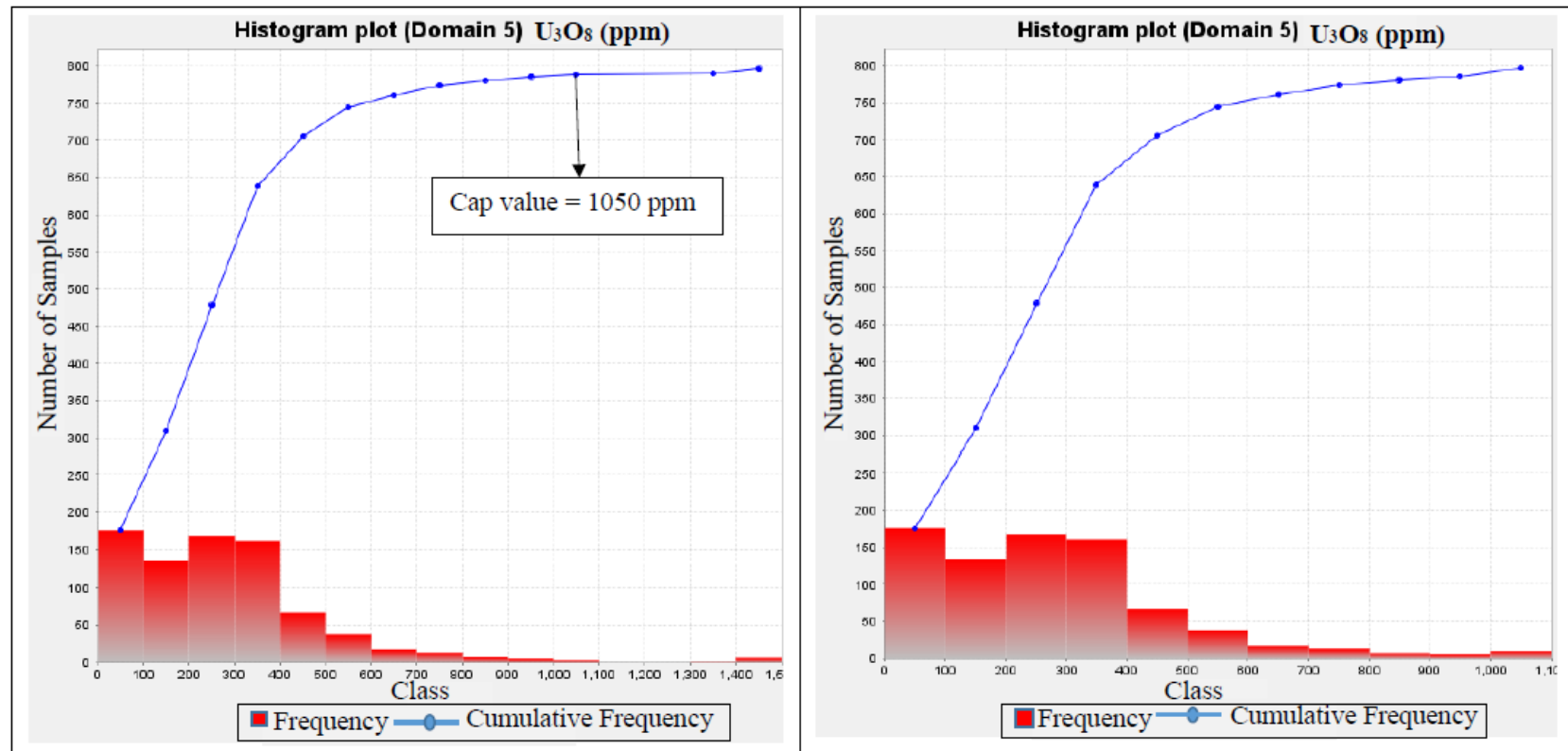


Figure A1: Domain 5 U₃O₈ (ppm) histogram, and a cumulative frequency plot of 1.5m composites before (left) and after (right) capping.

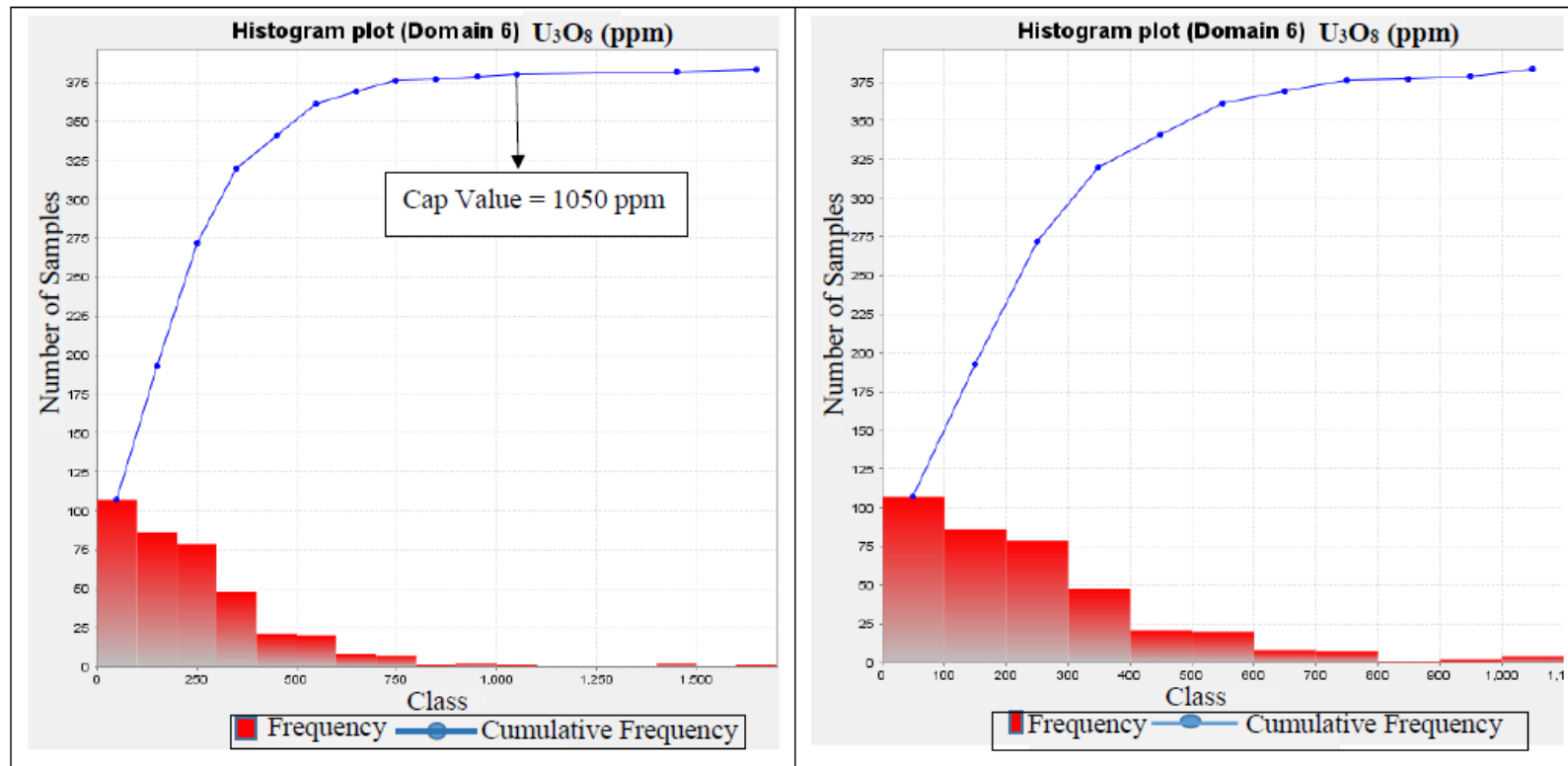


Figure A2: Domain 6 U₃O₈ (ppm) histogram, and a cumulative frequency plot of 1.5m composites before (left) and after (right) capping.

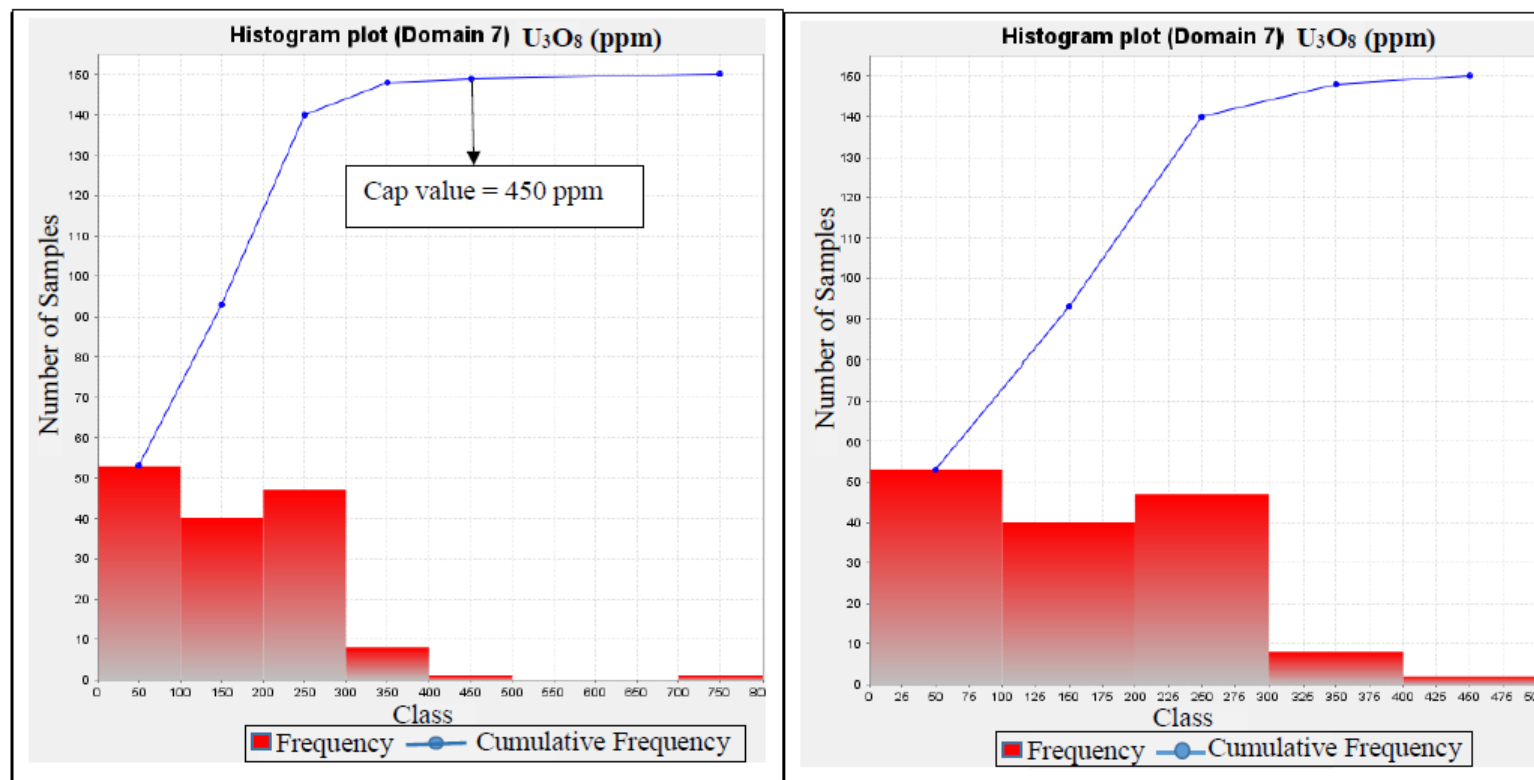


Figure A3: Domain 7 U_3O_8 (ppm) histogram, and a cumulative frequency plot of 1.5m composites before (left) and after (right) capping.

Appendix B: - U_3O_8 (ppm) histograms, log histograms, and cumulative log-probability plots of 1.5m composites

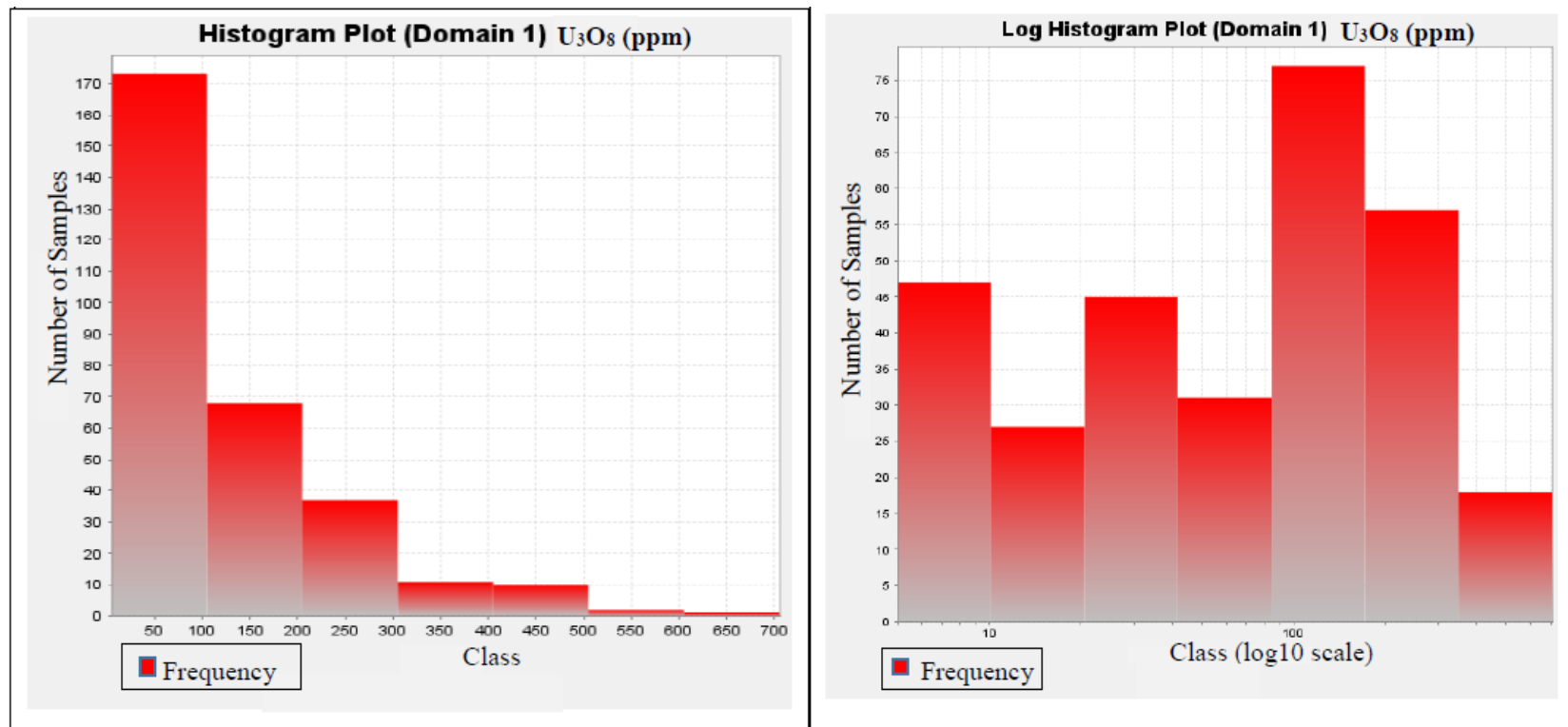


Figure B1: Domain 1 U_3O_8 (ppm) histogram (left), and the histogram of the $\log_{10}$ (U_3O_8) (right) of 1.5m composites

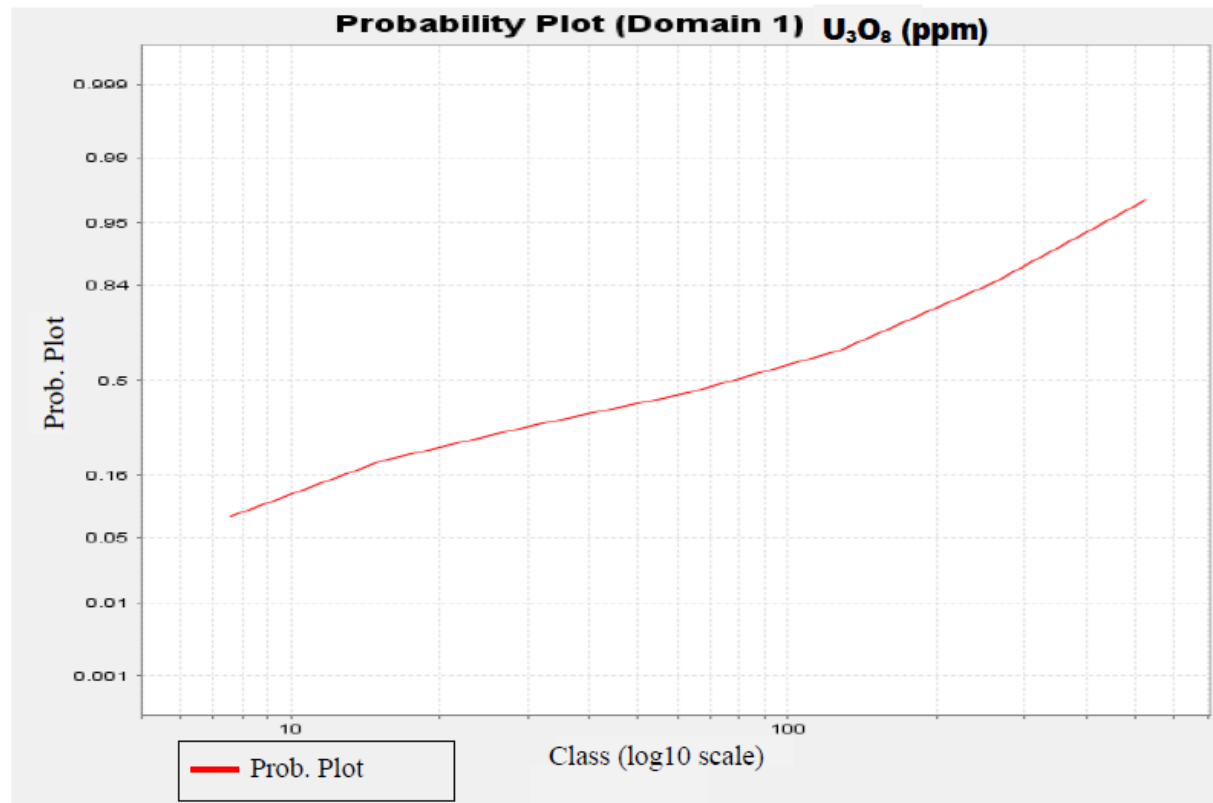


Figure B2: Domain 1- cumulative log-probability plot of 1.5m composites.

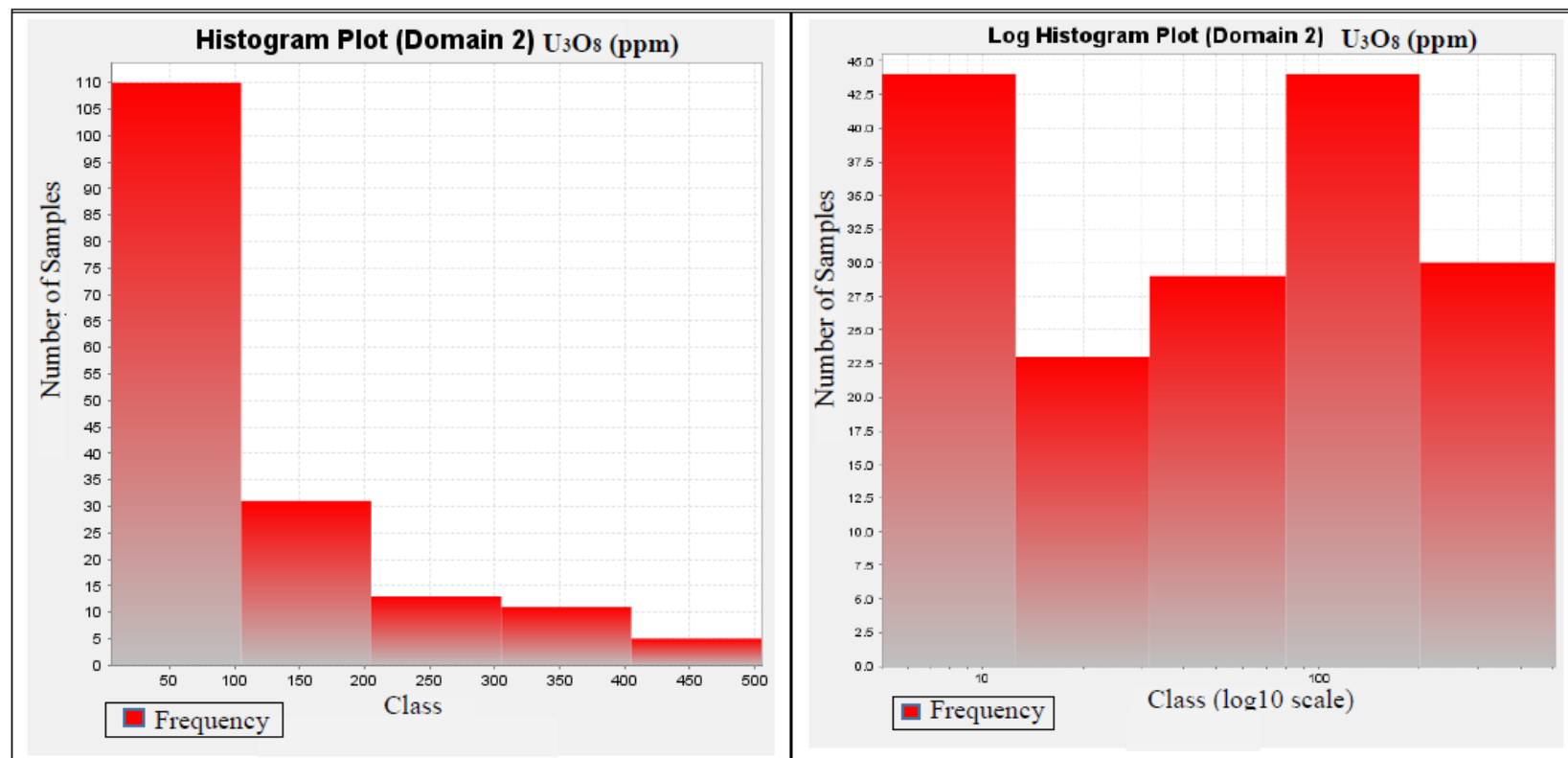


Figure B3: Domain 2 U_3O_8 (ppm) histogram (left), and the histogram of the $\log_{10}(U_3O_8)$ (right) of 1.5m composites.

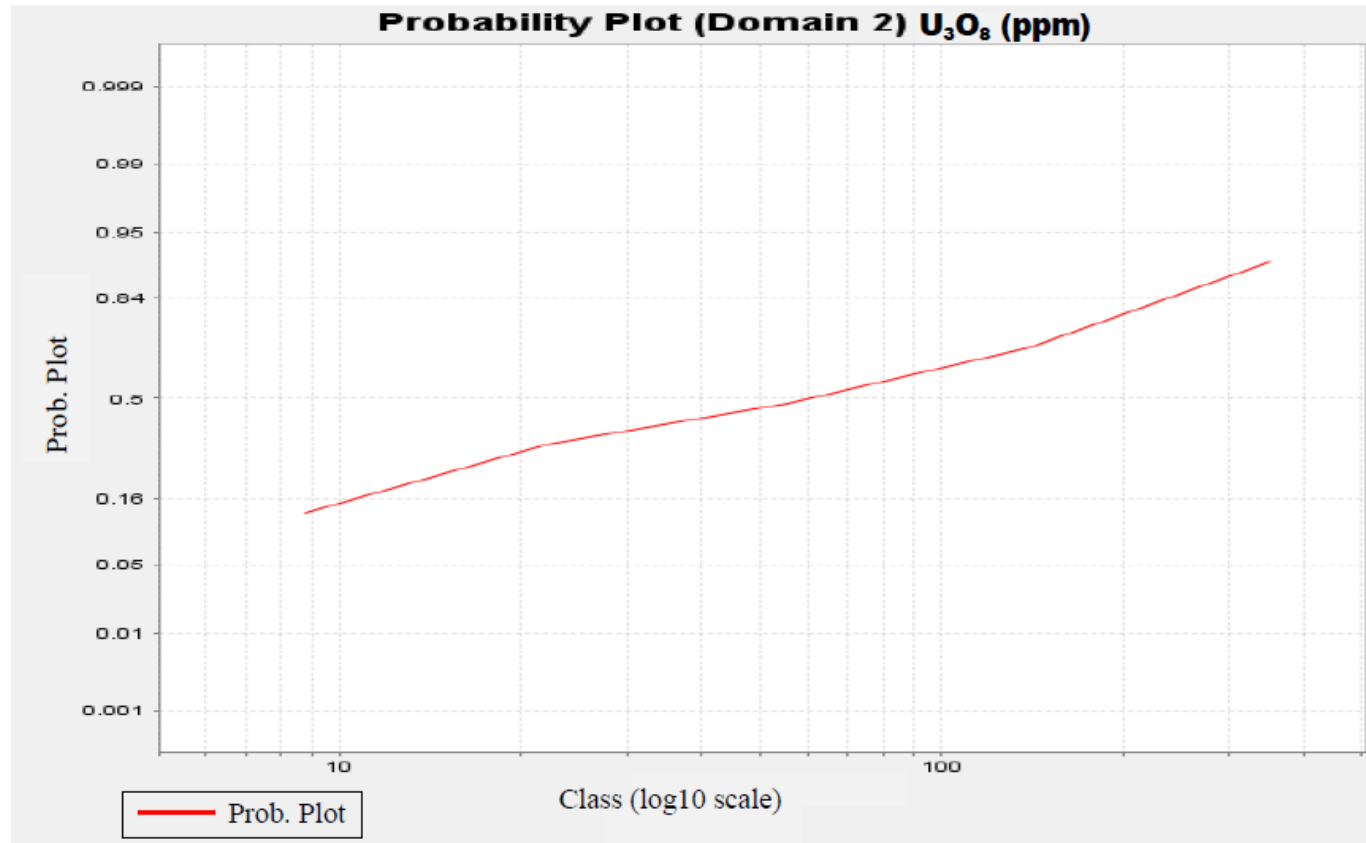


Figure B4: Domain 2 - cumulative log-probability plot of 1.5m composites.

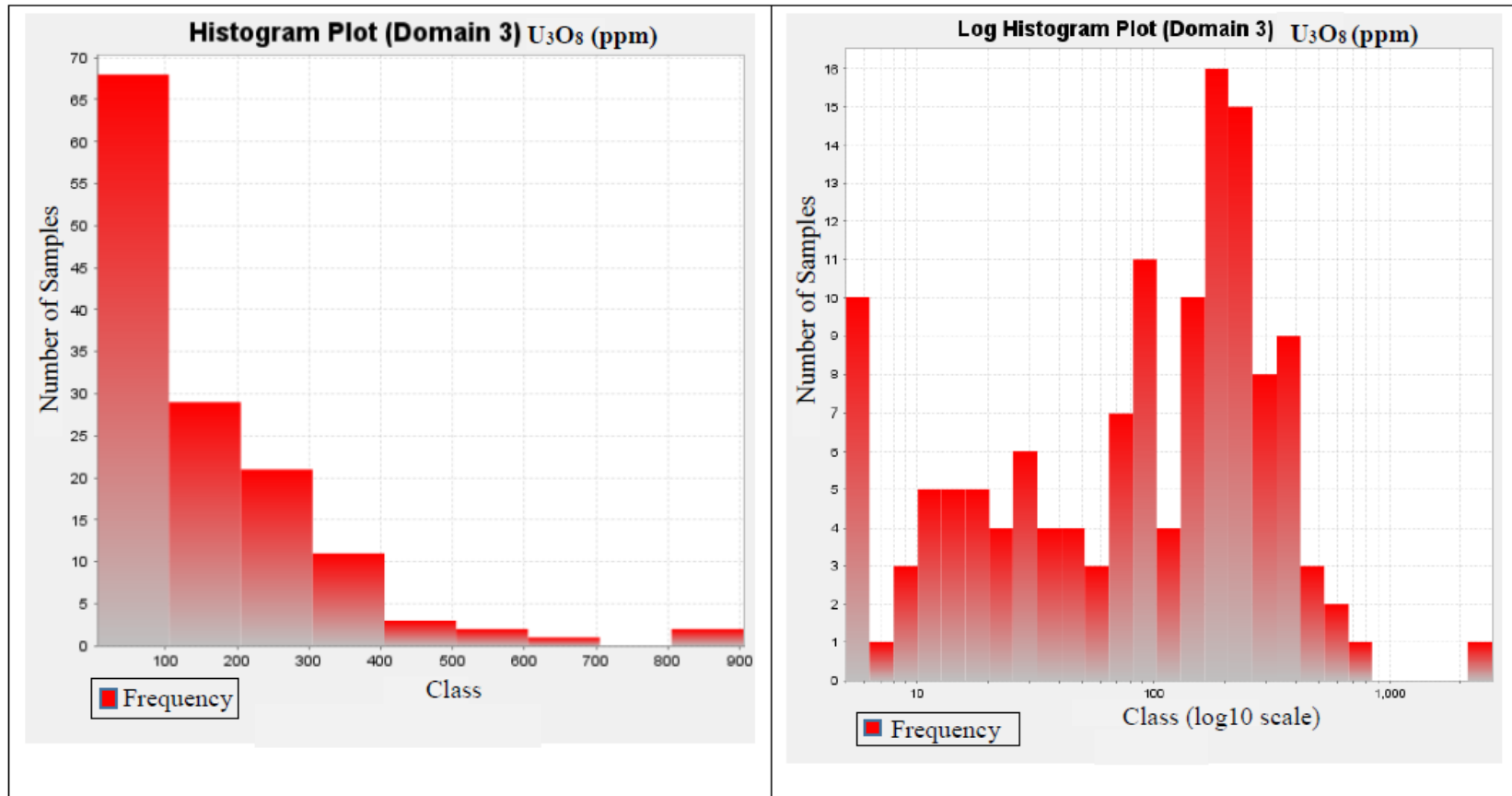


Figure B5: Domain 3 U_3O_8 (ppm) histogram (left), and the histogram of the $\log_{10}(U_3O_8)$ (right) of 1.5m composites.

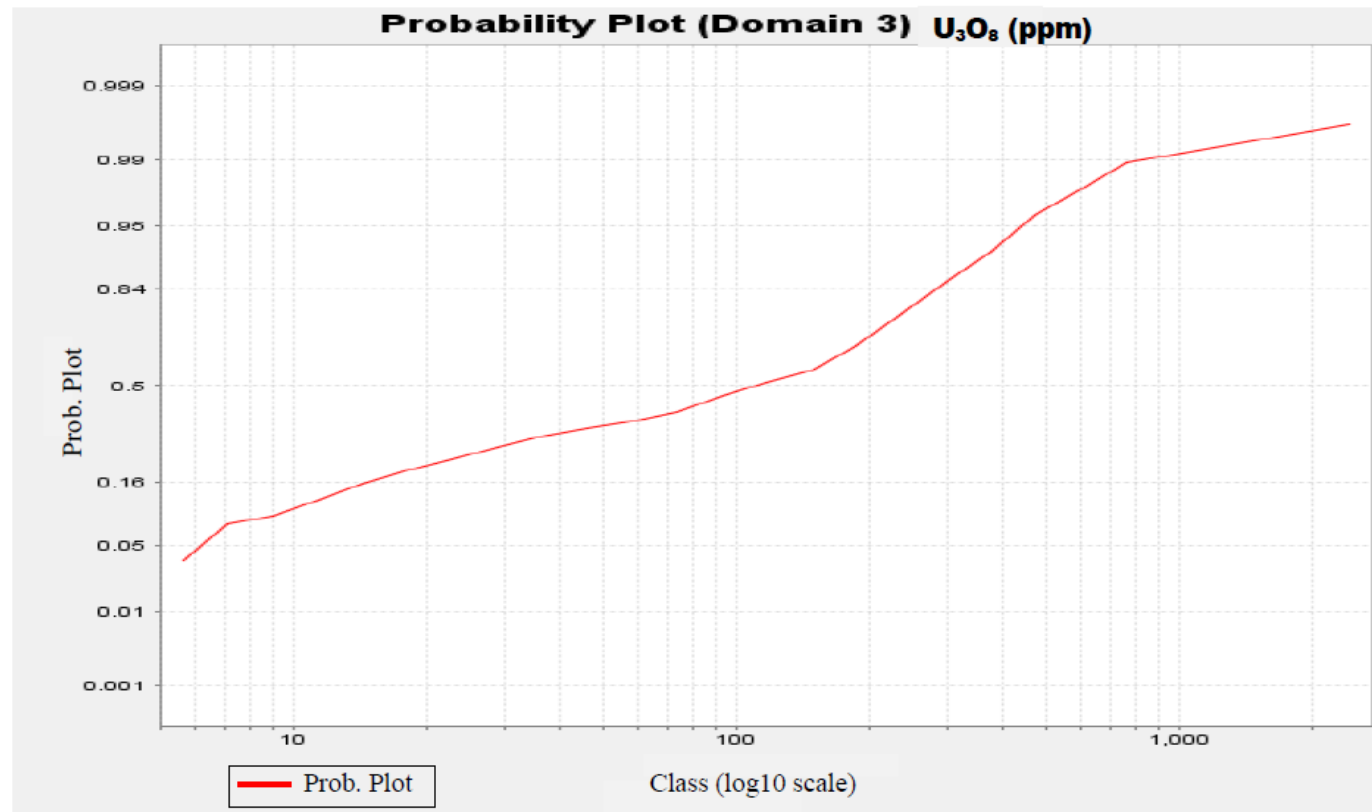


Figure B6: Domain 3 - cumulative log-probability plot of 1.5m composites.

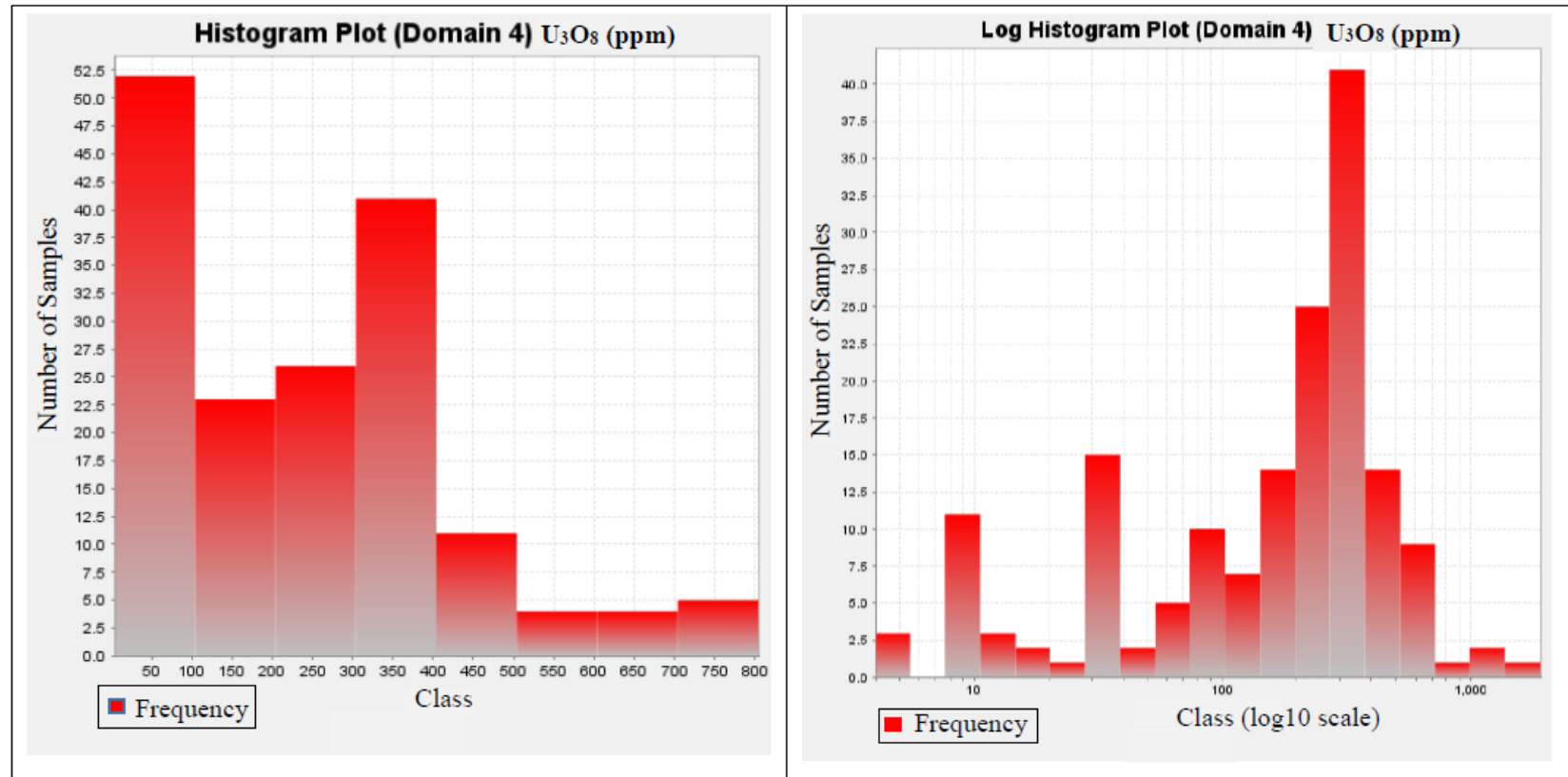


Figure B7: Domain 4 U_3O_8 (ppm) histogram (left), and the histogram of the $\log_{10}$ (U_3O_8) (right) of 1.5m composites.

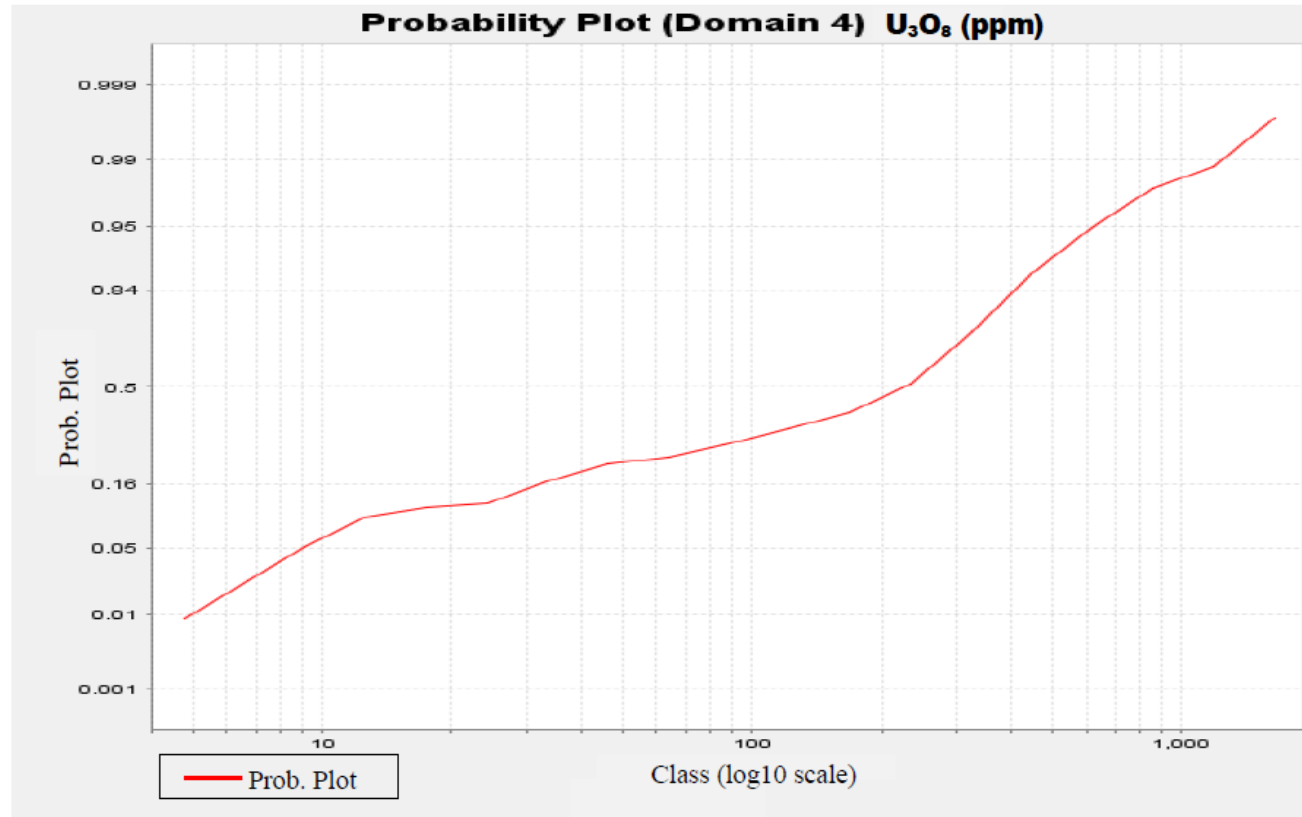


Figure B8: Domain 4 - cumulative log-probability plot of 1.5m composites.

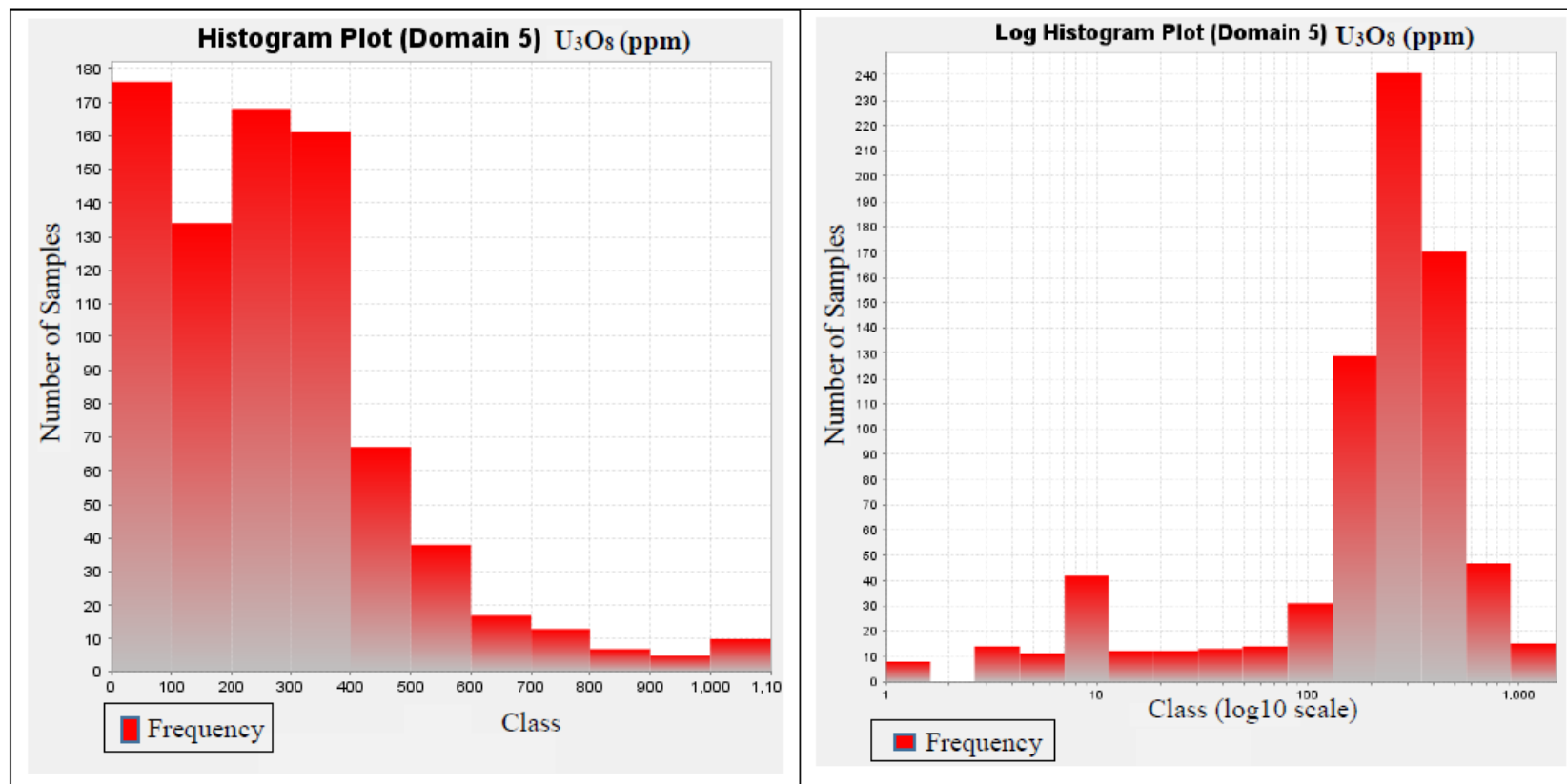


Figure B9: Domain 5 U₃O₈ (ppm) histogram (left), and the histogram of the log₁₀ (U₃O₈) (right) of 1.5m composites.

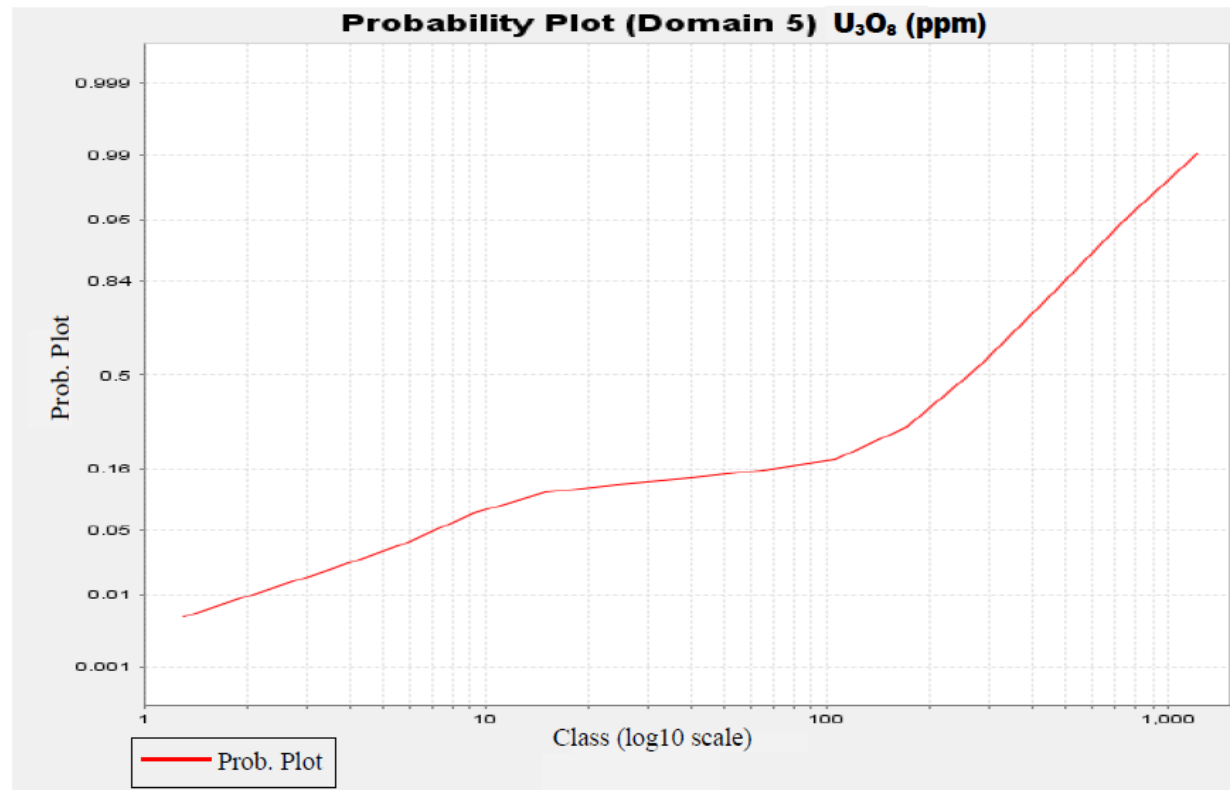


Figure B10: Domain 5 - cumulative log-probability plot of 1.5m composites.

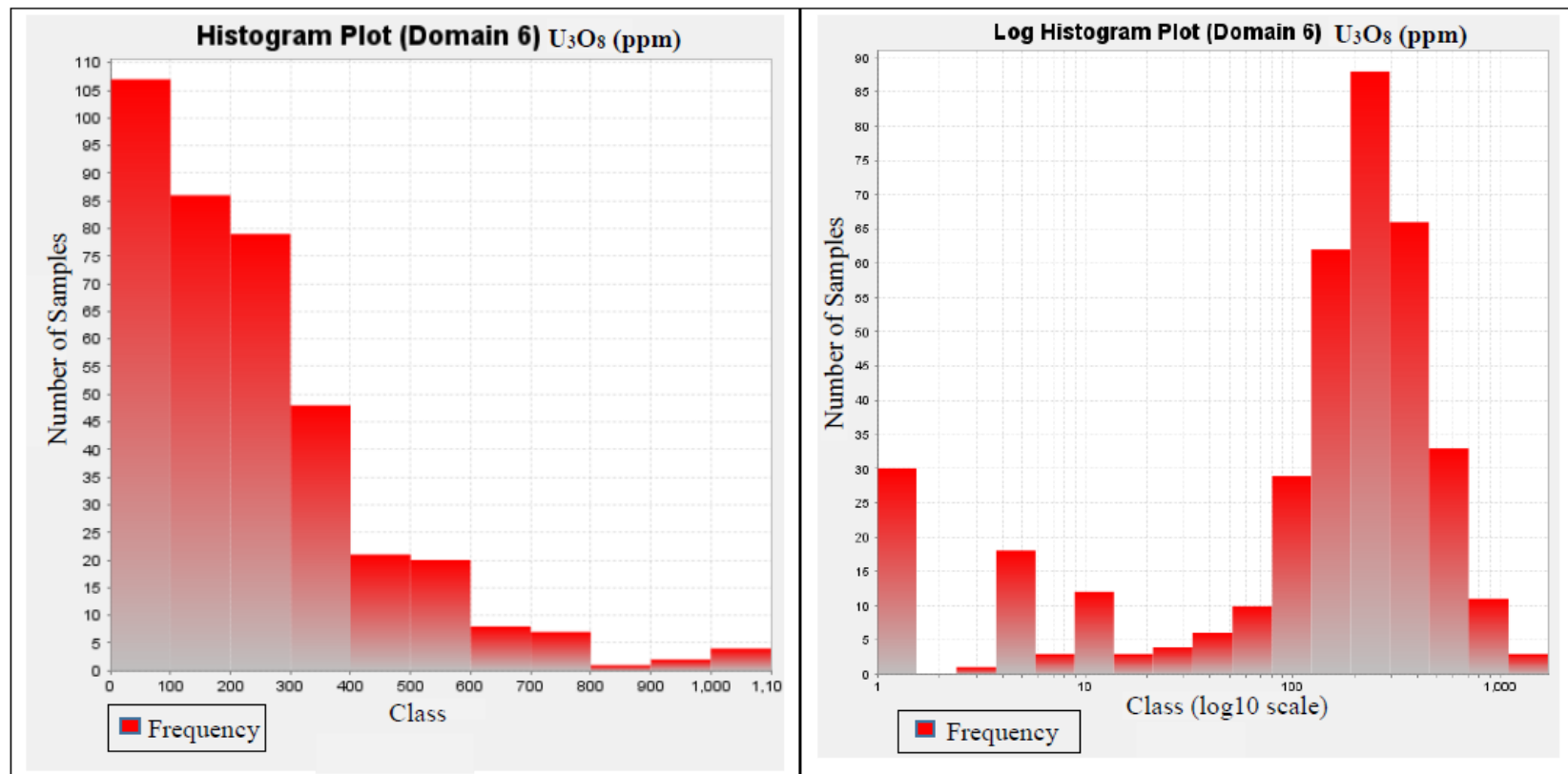


Figure B11: Domain 6 U_3O_8 (ppm) histogram (left), and the histogram of the $\log_{10}(U_3O_8)$ (right) of 1.5m composites.

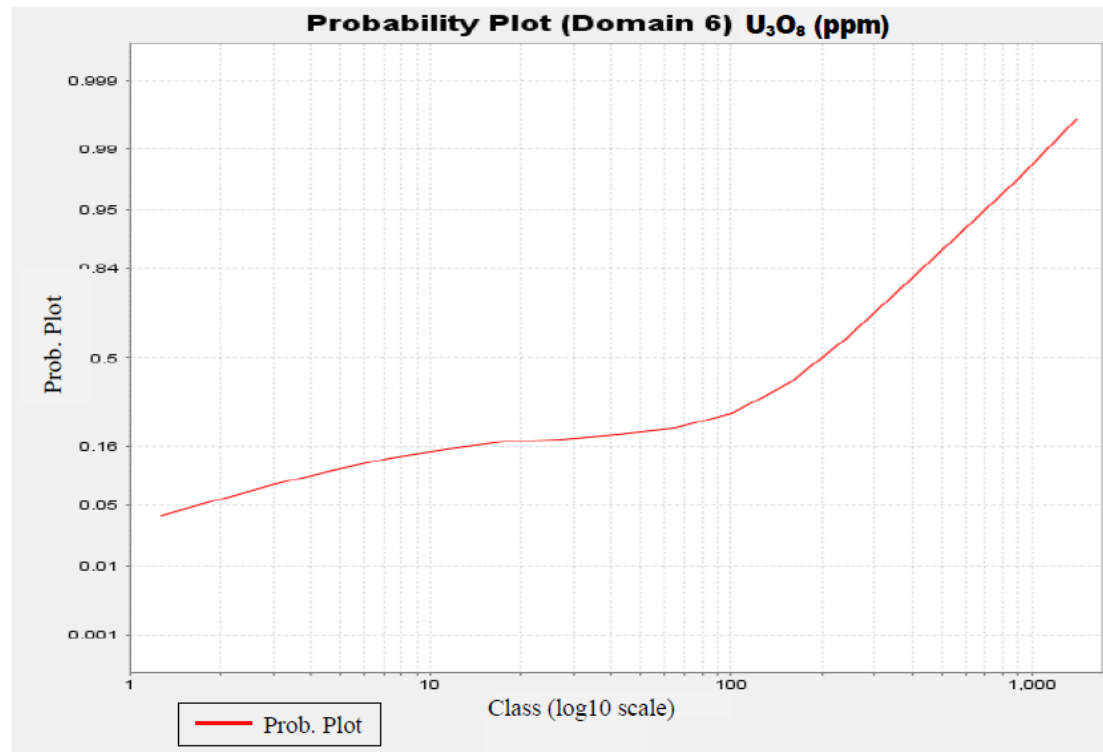


Figure B12: Domain 6 - cumulative log-probability plot of 1.5m composites.

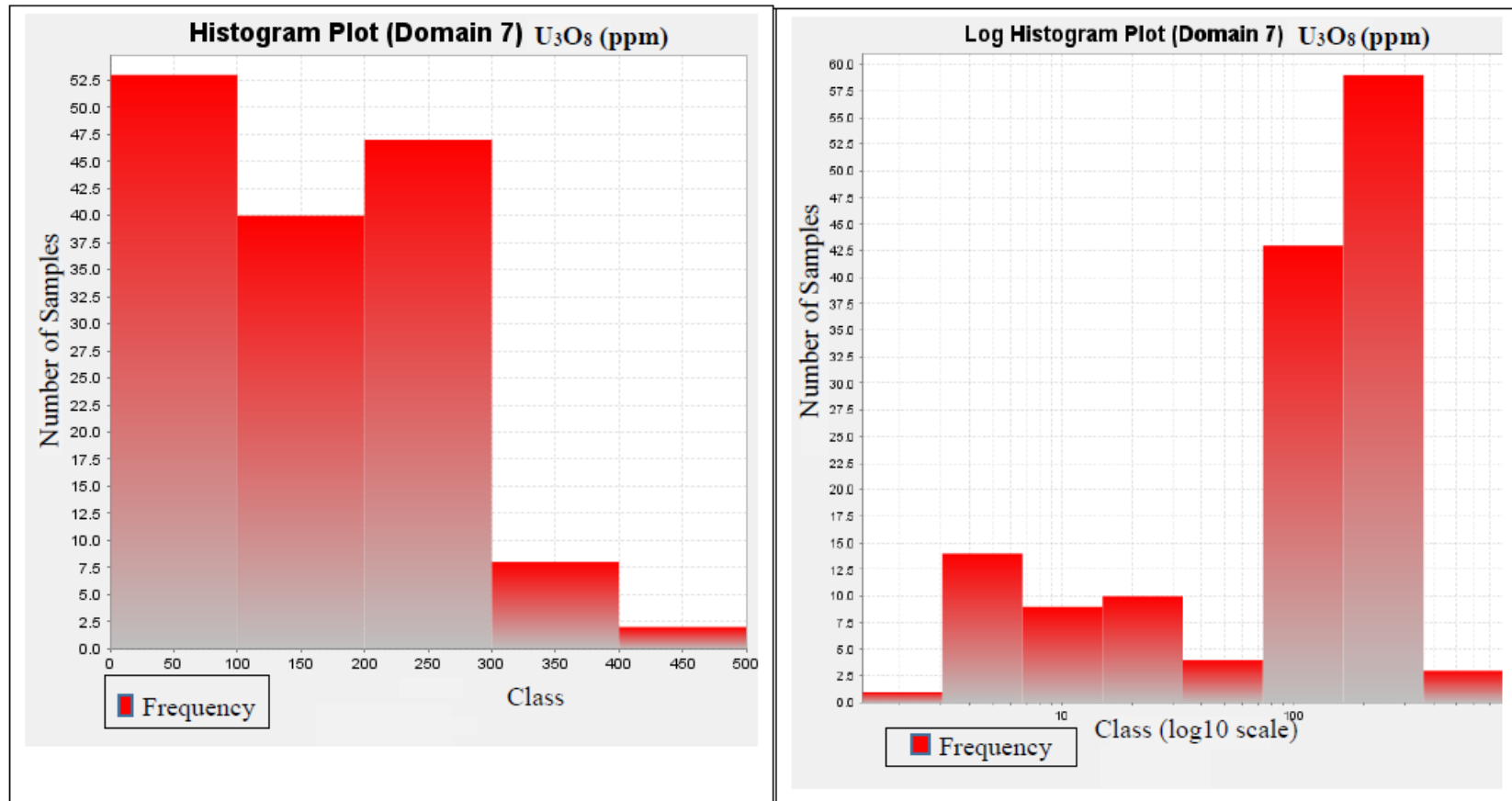


Figure B13: Domain 7 U_3O_8 (ppm) histogram (left), and the histogram of the $\log_{10}(U_3O_8)$ (right) of 1.5m composites.

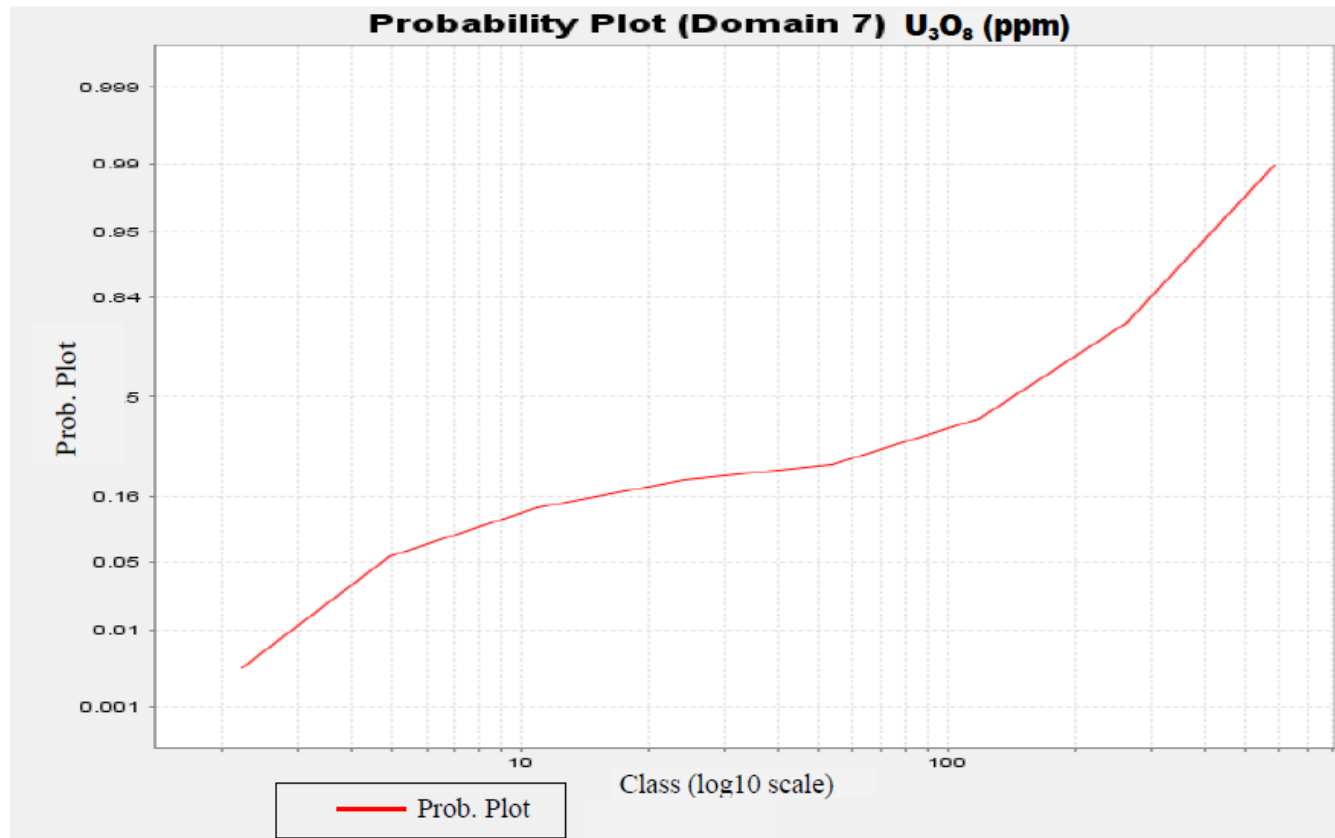


Figure B14: Domain 7- cumulative log-probability plot of 1.5m composites

Appendix C: KNA Results

Table C1: SUPERVISOR: Well-informed blocks - Block size optimisation

Block size Xm x Ym x Zm	SLOR: Slope of regression	KE: Kriging efficiency	No. of negative weights	Proportion of negative weights
10 x 10 x 10	0.994	0.908	43	0.439
10 x 10 x 15	0.993	0.911	39	0.398
10 x 10 x 20	0.991	0.910	34	0.347
20 x 20 x 10	0.993	0.921	40	0.408
20 x 20 x 15	0.992	0.922	34	0.347
20 x 20 x 20	0.990	0.919	33	0.337
30 x 30 x 10	0.992	0.926	33	0.337
30 x 30 x 15	0.991	0.925	31	0.316
30 x 30 x 20	0.989	0.922	28	0.286
40 x 40 x 10	0.990	0.922	28	0.286
40 x 40 x 15	0.989	0.921	27	0.276
40 x 40 x 20	0.987	0.918	26	0.265
50 x 50 x 10	0.988	0.913	20	0.204
50 x 50 x 15	0.987	0.913	20	0.204
50 x 50 x 20	0.985	0.910	21	0.214
60 x 60 x 10	0.985	0.901	21	0.214
60 x 60 x 15	0.984	0.901	21	0.214
60 x 60 x 20	0.982	0.900	20	0.204
70 x 70 x 10	0.982	0.888	21	0.214
70 x 70 x 15	0.981	0.888	21	0.214
70 x 70 x 20	0.979	0.887	21	0.214
80 x 80 x 10	0.979	0.874	21	0.214
80 x 80 x 15	0.978	0.874	21	0.214
80 x 80 x 20	0.976	0.873	21	0.214
90 x 90 x 10	0.976	0.859	21	0.214
90 x 90 x 15	0.974	0.860	22	0.224
90 x 90 x 20	0.972	0.860	22	0.224
100 x 100 x 10	0.972	0.845	22	0.224

Table C1 continued: SUPERVISOR: Well-informed blocks - Block size optimisation

Block size Xm x Ym x Zm	Slope of regression	Kriging efficiency	No. of negative weight	Proportion of negative weights
100 x 100 x 15	0.970	0.846	22	0.224
100 x 100 x 20	0.968	0.846	22	0.224

Table C2: SUPERVISOR: Well-informed 50m x 50m x 15m blocks - Number of samples optimisation.

Number of samples	SLOR: Slope of regression	KE: Kriging efficiency	No. of negative weights	Proportion of negative weights
2	0.710	0.631	0	0.000
4	0.758	0.726	0	0.000
6	0.781	0.765	0	0.000
8	0.833	0.845	0	0.000
10	0.846	0.862	0	0.000
12	0.855	0.873	0	0.000
14	0.862	0.880	0	0.000
16	0.865	0.882	0	0.000
18	0.867	0.884	0	0.000
20	0.886	0.888	0	0.000
22	0.888	0.888	0	0.000
24	0.889	0.888	0	0.000
26	0.896	0.890	0	0.000
28	0.897	0.890	0	0.000
30	0.898	0.890	1	0.033
32	0.913	0.895	2	0.063
34	0.915	0.895	3	0.088
36	0.916	0.896	3	0.083
38	0.916	0.896	3	0.079
40	0.916	0.896	4	0.100

Table C2 continued: SUPERVISOR: Well-informed 50m x 50m x 15m blocks

Number of samples optimisation

Number of samples	SLOR: Slope of regression	KE: Kriging efficiency	No. of negative weights	Proportion of negative weights
42	0.916	0.896	6	0.143
44	0.926	0.897	9	0.205
46	0.927	0.898	9	0.196
48	0.928	0.898	9	0.188
50	0.932	0.898	10	0.200
52	0.932	0.898	10	0.192
54	0.933	0.898	10	0.185
56	0.969	0.910	9	0.161
58	0.971	0.910	9	0.155
60	0.971	0.910	9	0.150
62	0.971	0.910	11	0.177
64	0.971	0.910	13	0.203
66	0.971	0.910	15	0.227
68	0.978	0.912	14	0.206
70	0.978	0.912	14	0.200
72	0.979	0.912	14	0.194
74	0.979	0.912	16	0.216
76	0.979	0.912	17	0.224
78	0.979	0.912	18	0.231
80	0.980	0.912	20	0.250
82	0.981	0.912	20	0.244
84	0.981	0.912	20	0.238
86	0.982	0.912	21	0.244
88	0.983	0.192	21	0.239
90	0.983	0.192	21	0.233
92	0.985	0.192	21	0.228
94	0.986	0.192	21	0.223
96	0.986	0.192	21	0.219
98	0.987	0.913	20	0.204

Table C3: SUPERVISOR: Well-informed 50m x 50m x 15m blocks - Discretisation configurations optimisation

Discretisation X x Y x Z	Slope of regression	Kriging efficiency	No. of negative weight	Proportion of negative weights
2 x 2 x 3	0.971	0.910	10	0.179
5 x 5 x 3	0.968	0.910	8	0.143
10 x 10 x 3	0.967	0.910	8	0.143
2 x 2 x 5	0.971	0.911	10	0.179
5 x 5 x 5	0.967	0.910	8	0.143
10 x 10 x 5	0.967	0.910	8	0.143

Table C4: SUPERVISOR: Poorly-informed blocks - Block size optimisation

Block size Xm x Ym x Zm	SLOR: Slope of regression	KE: Kriging efficiency	No. of negative weights	Proportion of negative weights
50 x 50 x 10	1.021	0.791	33	0.423
50 x 50 x 15	1.019	0.796	32	0.410
50 x 50 x 20	1.017	0.801	32	0.410
100 x 100 x 10	0.995	0.883	15	0.192
100 x 100 x 15	0.993	0.886	15	0.192
100 x 100 x 20	0.990	0.888	15	0.192
150 x 150 x 10	0.938	0.872	8	0.103
150 x 150 x 15	0.936	0.875	8	0.103
150 x 150 x 20	0.933	0.877	8	0.103
200 x 200 x 10	0.842	0.766	3	0.038
200 x 200 x 15	0.840	0.771	3	0.038
200 x 200 x 20	0.838	0.774	3	0.038
250 x 250 x 10	0.730	0.594	5	0.064
250 x 250 x 15	0.728	0.597	5	0.064
250 x 250 x 20	0.725	0.598	5	0.064
300 x 300 x 10	0.617	0.352	10	0.128
300 x 300 x 15	0.615	0.353	10	0.128
300 x 300 x 20	0.613	0.351	10	0.128

Table C5: SUPERVISOR: Poorly-informed 200m x 150m x 15m blocks – Number of samples optimisation

Number of samples	SLOR: Slope of regression	KE: Kriging efficiency	No. of negative weights	Proportion of negative weights
2	0.355	-0.946	0	0.000
4	0.451	-0.243	0	0.000
6	0.473	-0.125	0	0.000
8	0.483	-0.075	0	0.000
10	0.564	0.249	0	0.000
12	0.582	0.308	0	0.000
14	0.590	0.336	0	0.000
16	0.598	0.359	0	0.000
18	0.641	0.477	0	0.000
20	0.693	0.600	0	0.000
22	0.704	0.622	0	0.000
24	0.709	0.634	0	0.000
26	0.713	0.644	0	0.000
28	0.717	0.654	0	0.000
30	0.724	0.666	0	0.000
32	0.725	0.668	0	0.000
34	0.729	0.675	0	0.000
36	0.730	0.676	0	0.000
38	0.731	0.677	0	0.000
40	0.732	0.678	0	0.000
42	0.733	0.679	0	0.000
44	0.750	0.691	1	0.023
46	0.752	0.692	1	0.022
48	0.753	0.693	1	0.021
50	0.812	0.770	3	0.060
52	0.817	0.776	3	0.058
54	0.820	0.778	3	0.056

Table C5 continued: SUPERVISOR: Poorly-informed 200m x 150m x 15m blocks –Number of samples optimisation.

Number of samples	SLOR: Slope of regression	KE: Kriging efficiency	No. of negative weights	Proportion of negative weights
56	0.840	0.794	3	0.054
58	0.842	0.796	3	0.052
60	0.843	0.797	3	0.050
62	0.869	0.811	1	0.016
64	0.877	0.813	1	0.016
66	0.878	0.813	1	0.015
68	0.880	0.814	1	0.015
70	0.880	0.814	1	0.014
72	0.882	0.816	1	0.014
74	0.883	0.816	1	0.014
76	0.883	0.816	1	0.013
78	0.883	0.816	3	0.013

Table C6: SUPERVISOR: Poorly-informed 200m x 150m x 15m blocks – Discretisation configuration optimisation.

Discretisation X x Y x Z	SLOR: Slope of regression	KE: Kriging efficiency	No. of negative weight	Proportion of negative weights
2 x 2 x 3	0.896	0.773	1	0.0161
3 x 4 x 3	0.862	0.815	1	0.0161
5 x 5 x 3	0.856	0.821	1	0.0161
6 x 8 x 3	0.855	0.821	1	0.0161
10 x 10 x 3	0.850	0.820	1	0.0161
2 x 2 x 5	0.895	0.776	1	0.0161
3 x 4 x 5	0.862	0.816	1	0.0161
5 x 5 x 5	0.855	0.821	1	0.0161
6 x 8 x 5	0.852	0.821	1	0.0161
10 x 10 x 5	0.850	0.820	1	0.0161

Table C7: SURPAC: Well-informed 50m x 50m x 15m blocks – Number of samples optimisation

Min no. of samples	Max no. of samples	SLOR: Slope of regression	KE: Kriging efficiency	Kriging variance	No. of negative weights	Proportion of negative weights	No. of samples
2	4	0.758	0.726	0.164	0	0.000	4
4	6	0.781	0.765	0.140	0	0.000	6
6	8	0.833	0.845	0.093	0	0.000	8
8	10	0.846	0.862	0.082	0	0.000	10
10	12	0.855	0.873	0.076	0	0.000	12
12	14	0.862	0.880	0.072	0	0.000	14
14	16	0.865	0.882	0.072	0	0.000	16
16	18	0.867	0.884	0.070	0	0.000	18
18	20	0.886	0.888	0.067	0	0.000	20
20	22	0.888	0.888	0.067	0	0.000	22
22	24	0.889	0.888	0.067	0	0.000	24
24	26	0.896	0.890	0.066	0	0.000	26
26	28	0.897	0.890	0.066	0	0.000	28
28	30	0.898	0.890	0.066	1	0.033	30
30	32	0.913	0.895	0.063	2	0.063	32
32	34	0.915	0.895	0.063	3	0.088	34
34	36	0.916	0.896	0.062	3	0.083	36
36	38	0.916	0.896	0.062	3	0.079	38
38	40	0.916	0.896	0.062	4	0.100	40
40	42	0.917	0.896	0.062	5	0.119	42
42	44	0.926	0.897	0.061	9	0.205	44
44	46	0.927	0.898	0.061	9	0.196	46
46	48	0.928	0.898	0.061	9	0.188	48
48	50	0.932	0.898	0.061	9	0.180	50
50	52	0.932	0.898	0.061	10	0.192	52
52	54	0.933	0.898	0.061	10	0.185	54
54	56	0.969	0.910	0.054	9	0.161	56
56	58	0.971	0.910	0.054	9	0.155	58
58	60	0.971	0.910	0.054	9	0.150	60

Table C7 continued: SURPAC: Well-informed 50m x 50m x 15m blocks –
Number of samples optimisation

Min no. of samples	Max no. of samples	SLOR: Slope of regression	KE: Kriging efficiency	Kriging variance	No. of negative weights	Proportion of negative weights	No. of samples
60	62	0.971	0.910	0.054	11	0.177	62
62	64	0.971	0.910	0.054	13	0.203	64
64	66	0.971	0.910	0.054	15	0.227	66
66	68	0.978	0.912	0.053	14	0.206	68
68	70	0.978	0.912	0.053	14	0.200	70
70	72	0.979	0.912	0.053	14	0.194	72
72	74	0.979	0.912	0.053	16	0.216	74
74	76	0.979	0.912	0.053	17	0.224	76
76	78	0.979	0.912	0.053	18	0.231	78
78	80	0.980	0.912	0.053	20	0.250	80
80	82	0.981	0.912	0.053	20	0.244	82
82	84	0.981	0.912	0.053	20	0.238	84
84	86	0.982	0.912	0.053	21	0.244	86
86	88	0.983	0.912	0.053	21	0.239	88
88	90	0.983	0.912	0.053	21	0.233	90
90	92	0.985	0.912	0.052	22	0.239	92
92	94	0.986	0.912	0.052	21	0.223	94
94	96	0.986	0.912	0.052	21	0.219	96
96	98	0.987	0.913	0.052	20	0.204	98
98	100	0.987	0.913	0.052	20	0.204	98

Table C8: SURPAC: Poorly-informed 200m x 150m x 15m blocks – Number of samples optimisation.

Min no. of samples	Max no. of samples	SLOR: Slope of regression	KE: Kriging efficiency	Kriging variance	No. of negative weights	Proportion of negative weights	No. of samples
2	4	0.451	-0.243	0.342	0	0.000	4
4	6	0.473	-0.125	0.310	0	0.000	6
6	8	0.483	-0.075	0.296	0	0.000	8
8	10	0.564	0.249	0.207	0	0.000	10
10	12	0.582	0.308	0.191	0	0.000	12
12	14	0.590	0.336	0.183	0	0.000	14
14	16	0.598	0.359	0.176	0	0.000	16
16	18	0.643	0.483	0.142	0	0.000	18
18	20	0.697	0.609	0.108	0	0.000	20
20	22	0.705	0.627	0.103	0	0.000	22
22	24	0.711	0.640	0.099	0	0.000	24
24	26	0.715	0.648	0.097	0	0.000	26
26	28	0.719	0.659	0.094	0	0.000	28
28	30	0.725	0.667	0.092	0	0.000	30
30	32	0.727	0.673	0.090	0	0.000	32
32	34	0.728	0.674	0.090	0	0.000	34
34	36	0.730	0.677	0.089	0	0.000	36
36	38	0.734	0.683	0.087	0	0.000	38
38	40	0.735	0.684	0.087	0	0.000	40
40	42	0.735	0.684	0.087	0	0.000	42
42	44	0.753	0.696	0.084	0	0.000	44
44	46	0.754	0.697	0.083	0	0.000	46
46	48	0.755	0.698	0.083	0	0.000	48
48	50	0.815	0.775	0.062	3	0.060	50
50	52	0.820	0.780	0.061	3	0.058	52
52	54	0.822	0.782	0.060	3	0.056	54
54	56	0.843	0.799	0.055	3	0.054	56

Table C8 continued: SURPAC: Poorly-informed 200m x 150m x 15m blocks –
Number of samples optimisation.

Min no. of samples	Max no. of samples	SLOR: Slope of regression	KE: Kriging efficiency	Kriging variance	No. of negative weights	Proportion of negative weights	No. of samples
56	58	0.846	0.801	0.055	3	0.052	58
58	60	0.847	0.802	0.055	3	0.050	60
60	62	0.873	0.816	0.051	0	0.000	62
62	64	0.880	0.818	0.050	0	0.000	64
64	66	0.881	0.818	0.050	0	0.000	66
66	68	0.883	0.819	0.050	0	0.000	68
68	70	0.883	0.819	0.050	0	0.000	70
70	72	0.885	0.820	0.050	0	0.000	72
72	74	0.886	0.820	0.049	0	0.000	74
74	76	0.886	0.820	0.049	0	0.000	76
76	78	0.886	0.820	0.049	0	0.000	78