
Evaluating Subsurface Uncertainty Using Modified Geostatistical Techniques

by:

William L. Wingle

A thesis submitted to the Faculty and Board of Trustees of the Colorado School of Mines in partial fulfillment of the requirements for the Degree of Doctor of Philosophy (Geological Engineer).

Golden, Colorado

Date: _____

Signed: _____
William L. Wingle

Approved: _____
Dr. Eileen P. Poeter

Golden, Colorado

Date: _____

Dr. Roger Slatt
Professor and Department Head,
Department of Geology and
Geological Engineering

ABSTRACT

There is a great deal of uncertainty about the distribution of geologic and hydrologic properties in the subsurface and the migration routes and extent of contaminants at most hazardous waste sites. This is because site data is limited. This research develops four geostatistical techniques which facilitate the assessment of and/or the reduction in the level of uncertainty associated with describing the subsurface. First, jackknifing and Latin-Hypercube sampling are used to define the uncertainty in the experimental semivariogram. Second, directional differences in the spatial variation of a semivariogram often cannot adequately be described using anisotropy factors; the kriging process is modified to accommodate three unique, orthogonal, semivariogram models. Third, the conditional simulation process is modified to use indicator classes rather than the threshold level between indicators. Fourth, zones at a site are modeled using individual and merged model semivariograms.

Using these methods is complex; consequently, a software package, UNCERT, was developed to integrate data collection, data evaluation, site interpretation, ground water flow and contaminant transport modeling, and data and model visualization. This software user interface makes the use of these modified geostatistical methods a practical endeavor.

TABLE OF CONTENTS

	ABSTRACT	v
	TABLE OF CONTENTS	vii
	LIST OF FIGURES	xi
	LIST OF TABLES	xxi
CHAPTER 1	INTRODUCTION	1
CHAPTER 2	JACKKNIFING & LATIN-HYPERCUBE SAMPLING	3
	2.1: Introduction	3
	2.2: Semivariograms	7
	2.3: Indicator Kriging And Stochastic Simulation	7
	2.4: Jackknifing	12
	2.4.1: Additional Comments About Jackknifing	15
	2.5: Latin-Hypercube Sampling	16
	2.6: Expert Opinion	18
	2.7: Results	19
	2.8: Conclusions	24
CHAPTER 3	VARIATION OF SEMIVRIOGRAM MODELS WITH DIRECTION	27
	3.1: Introduction	27
	3.2: Previous Work	28
	3.3: Theory	29
	3.3.1 Equation and Proof	31
	3.3.2: Positive Definite Matrix Issues	36
	3.4: Modification of Algorithms	37
	3.4.1: Algorithm Constraints	37
	3.4.2: Computational Cost	38
	3.5: Examples	38
	3.5.1: Comparison With the Classic Method	38
	3.5.2: Practical Applications	48
	3.6: Conclusions	61

CHAPTER 4	CLASS VS. THRESHOLD INDICATOR SIMULATION	63
	4.1: Introduction	63
	4.2: Previous Work	66
	4.3: Methods	69
	4.3.1: Semivariogram Calculation	70
	4.3.2: Data Definition	70
	4.3.3: P_1 - P_2 Calculations	71
	4.3.4: Difference Between Prior Hard and Prior Soft Data CDF's for Class and Threshold Simulations	72
	4.3.5: Order Relation Violations	75
	4.4: Applications	77
	4.4.1: Synthetic Data Set	77
	4.4.2: Colorado School of Mines Survey Field	93
	4.5: Conclusions	107
CHAPTER 5	ZONAL KRIGING	113
	5.1: Introduction and Previous Work	113
	5.2: Methodology	115
	5.3: Examples	119
	5.3.1: Synthetic Data Set Example	119
	5.3.2: Yorkshire, England Example	121
	5.3.3: Colorado School of Mines Survey Field Example	127
	5.4: Steps to Determine if Zonal Kriging is Appropriate	146
	5.5: Conclusions	153
CHAPTER 6	UNCERT: GEOSTATISTICAL, GROUND WATER MODELING, AND VISUALIZATION SOFTWARE	155
	6.1: Introduction	155
	6.2: Previous Work	155
	6.3: Platform Support	157
	6.4: UNCERT Modules	158
	6.4.1: Mainmenu	158
	6.4.2: Plotgraph	158
	6.4.3: Histo	158
	6.4.4: Distcomp	159
	6.4.5: Vario	159
	6.4.6: Variofit	159
	6.4.7: Grid	159
	6.4.8: Contour	160

TABLE OF CONTENTS

	6.4.9: Surface	160
	6.4.10: Block	160
	6.4.11: Sisim & Sisim3d	160
	6.4.12: Modmain	160
	6.4.13: Mt3dmain	161
	6.4.14: Array	161
	6.4.15: Utilities	161
CHAPTER 7	SUMMARY AND CONCLUSIONS	163
	7.1: Summary and Conclusions	163
	7.2: Recommendations for Future Work	164
CHAPTER 8	BIBLIOGRAPHY	167
APPENDIX A	UNCERT AND UNCERT USER'S MANUAL	171
	A1: Information and Comments:	171
	A1.1: Warranty:	172
	A2: Hardware / Operating System Requirements:	172
	A3: Acquiring Software:	172
	A4: Installation:	174
	A4.1: Unpacking the Software:	174
	A4.2: Compiling UNCERT:	174
	A4.3: Setting Up User Accounts:	176
APPENDIX B	SAMPLE DATA SETS	179

LIST OF FIGURES

FIGURE 2-1.	Borehole data used to interpret the subsurface may not provide a unique solution. In this case, there are eleven data samples; six of fine-grained sediments with low hydraulic conductivity, and five of coarse-grained sediments with high hydraulic conductivity. Although data for each map is identical, the nature of the geology in each map is substantially different. This illustrates that there is uncertainty associated with the interpretation of the character of subsurface at locations that have not been sampled.....	5
FIGURE 2-2.	Contaminants will migrate in different patterns within the two geologic models presented in Figure 2.1. It is important to evaluate the probable alternative scenarios when designing a remediation plan.	6
FIGURE 2-3.	Features of a semivariogram and parameters defining the search area (after Englund and Sparks, 1988).....	8
FIGURE 2-4.	Experimental and modeled semivariograms developed from the eleven labeled data points in Figure 2.1. A great deal of uncertainty is associated with the modeled semivariogram because of the limited number of data.	9
FIGURE 2-5.	These experimental semivariograms based on 315 data points from the models in Figure 2.1 were determined by overlaying a regular grid (25' x 25') on each model. The distribution of high and low conductivity materials in Figure 2.1a was determined to be isotropic and is described by the model semivariogram in 2.5a. In Figure 2.1b, the distribution is anisotropic and the major and minor axes of the model semivariogram ellipsoid are shown in 2.5b and 2.5c respectively.....	10
FIGURE 2-6.	Jackknifing the eleven data points indicated in Figure 2.1 allows evaluation of uncertainty associated with the semivariogram. The vertical error-bars define the 95% confidence intervals for the mean $\gamma^*(h)$ of each lag. The variance around the mean lag is represented by the horizontal error bars. Each data point represents 1 instance of a jackknifed experimental semivariogram. This experimental semivariogram is based on the assumption of an isotropic material distribution.	13
FIGURE 2-7.	Although anisotropy cannot be identified by evaluating single semivariograms of the eleven data points, anisotropic character is	

	hinted at when the same data are jackknifed along specified search directions. For a search direction of 45° , the range is likely to be less than 100 feet. In the 135° search direction, the range is likely to be greater than 150 feet, and possibly more than 500 feet. The anisotropy defined in Figure 2.5b-2.5c cannot be determined from the eleven data points, but its possibility is indicated by the data. Symbols are described in the caption of Figure 2.6.	14
FIGURE 2-8.	When a substantial amount of data are collected, the experimental semivariogram may be clearly defined. In this jackknifed simulation, there is little uncertainty in the lag means, and there would be little uncertainty in defining the model semivariogram.	16
FIGURE 2-9.	Reasonable models must be selected from the shaded region in 2.8a to represent the “flavor” of the alternative interpretations of the data. Four model semivariograms with a nugget selected from the lower quartile of possible nugget values are shown in 2.8b. The ranges of the four semivariograms are selected to represent each of the quartiles of possible ranges. Sixteen models would be used to represent the distribution of semivariogram models for the isotropic case. Symbols are described in the caption of Figure 2.6.	17
FIGURE 2-10.	These two simulations were generated assuming isotropy and using the model semivariogram developed from the extensive data set and illustrated in Figure 2.5a. The solutions are a reasonable approximation of the map in Figure 2.1a.	20
FIGURE 2-11.	These two simulations were generated assuming isotropy and using a latin-hypercube sample from the jackknifed model semivariogram ($C_0=0.0$, $C_1=0.25$, $a_1=115'$) developed from the eleven data points and illustrated in Figure 2.6. The solutions are a reasonable approximation of the map in Figure 2.1a, and are very similar to those generated in Figure 2.10. Much of the reason that the simulations in Figure 2.10 and 2.11 are similar is that the same random path through the grid was used to simulate 2.10a and 2.11a and another path was used to simulate 2.10b and 2.11b.	21
FIGURE 2-12.	These two stochastic simulations were generated assuming anisotropy using the jackknifed model semivariogram based on the eleven data points and illustrated in Figure 2.6. The latin-hypercube technique was applied and these are two simulations of a potential 256, as described in the text. Even though the geologic models presented in Figure 2.1 are different, use of jackknifing and Latin Hypercube sampling can produce both configurations from limited data. These solutions are a reasonable approximation of the map in Figure 2.1b. Unfortunately, the method will not indicate whether these simulations or the simulations in Figures 2.10 and 2.11 are the most likely because the data are not sufficient to draw such a conclusion. ..	22

FIGURE 2-13.	These two simulations were generated assuming anisotropy using the extensive model semivariogram based on the extensive data set and illustrated in Figures 2.5b-2.5c. The solutions are a reasonable approximation of the map in Figure 2.1b, and are very similar to those generated in Figure 2.12, indicating that extensive data are more important to determining the character of the semivariogram than they are to conditioning the simulation.....	23
FIGURE 3-1.	When directional semivariograms are used, distance alone does not determine the most influential neighboring points. In this example, all points in the minor model axis direction (b) that are separated by less than x_2 (158 m) have smaller $\gamma(h)$'s than points separated by x_1 (109 m) on the major-axis (a). The same is true for x_3 and x_2 respectively.	30
FIGURE 3-2.	Directional semivariogram analysis components.....	32
FIGURE 3-3.	Example results confirming directional semivariograms can exactly mimic anisotropy factors: a) sample data set, b) SK map using anisotropic factors, c) SK map using directional semivariograms.	40
FIGURE 3-4.	Geometric steps for calculating directional semivariogram model defined in Figure 3.1. The major axis is aligned North-South, and the minor axis is aligned East-West. Note, the 45° angle is transformed ($\rightarrow$) based on the anisotropy of the ellipsoid.....	47
FIGURE 3-5.	Semivariogram models used for synthetic directional semivariogram data set. Despite the general rule of thumb that the practical Gaussian range to a spherical range (a) is the $\text{SQRT}(3)$ multiplied by the range (a), the Gaussian (range (a) $\times$ $\text{SQRT}(3)$) model, because it mimicked the general nature of the other models more closely.....	50
FIGURE 3-6.	Results of directional semivariogram models using different assumptions about major and minor semivariogram models.....	51
FIGURE 3-7.	Differences between original SK models (Figure 3.3a-b), and directional semivariogram models (Figures 3.6a-c).....	52
FIGURE 3-8.	Distribution of differences between original SK models (Figure 3.3a-b), and directional semivariogram models (Figures 3.6a-c).....	53
FIGURE 3-9.	Experimental and model semivariograms for RMA bedrock residuals (2nd order trend removed): a) anisotropy factor model optimized to minimize MSE based on Johnson (1995), b) optimized minor-axis fit with Gaussian model (note MSE reduced by 82%), c) minor-axis Gaussian model fit with elevated nugget to reduce kriging matrix instability, d) anisotropy factor model optimized to minimize MSE, but also honor nugget defined in b).....	55

FIGURE 3-10.	Location of sample wells at RMA (a), SK map of bedrock elevation residuals (b), and estimation variance using an anisotropy factor, spherical-spherical semivariogram model I (c) (Johnson, 1995). 57
FIGURE 3-11.	RMA SK map of bedrock elevation residuals (a), and estimation variance using robust Gaussian factor semivariogram models (b), and difference between robust Gaussian (b) and original (Figure 3.10c) estimation variance maps (c)..... 58
FIGURE 3-12.	Distribution of differences between alternative estimation variance maps: (a) the difference between the robust Gaussian (III) and the anisotropy factor, spherical-spherical semivariogram model (I); (b) the difference between the robust Gaussian (III) and the anisotropy factor, spherical-spherical semivariogram model with nugget (I). The positive, average difference in (a) indicates the Gaussian model has a higher average estimation variance. The negative, average difference in (b) indicates the Gaussian model has a lower average estimation variance.... 59
FIGURE 3-13.	RMA SK map of bedrock elevation residuals (a), and estimation variance using the anisotropic factor spherical-spherical semivariogram model with a valid nugget (IV) (b), and difference between the robust Gaussian (III) (Figure 3.10c) and estimation variance maps (b). 60
FIGURE 3-14.	Q-Q plot of bedrock elevation residuals where the original Spherical model using anisotropy factors (I) is compared versus 1) the original Gaussian model (II), 2) the robust Gaussian model, and 3) the original Spherical model adjusted with a nugget. The plot suggests that the general nature of all the models are similar. 61
FIGURE 4-1.	Spatial distribution of several indicators. Defining semivariograms based on indicator classes is more intuitive, because the range reflects the average size of the indicator bodies. The class semivariogram model ranges are: silt = 112m, silty-sand = 106m, sand = 60m, and gravel = 41m. For thresholds, semivariogram model ranges are: silt vs. all others = 112m, silt and silty-sand vs. sand and gravel = 68m, and gravel vs. all others = 41m. 66
FIGURE 4-2.	Ordinary Kriging (and most other estimation methods) tends to average or smooth data to achieve a best linear unbiased estimate (BLUE) of reality. Indicator Kriging with conditional simulation provides a means for modeling the variability observed in nature, while still honoring the field data. Conditional simulation does not produce a best estimate of reality, but it yields models with characteristics similar to reality. When multiple realizations are made and averaged, values will approximate the smoothed, BLUE. 67
FIGURE 4-3.	Class and threshold indicator kriging generate the cumulative density function (CDF) in different ways. Threshold CDF's are determined directly from the probability that the specified grid location is less

	than each threshold level (a). The final CDF term should be less than 1.0 with the remaining probability attributed to the final indicator. Calculating class CDF's requires two steps. First the probability of occurrence of each class is calculated (b). The PDF is converted into a CDF by summing the individual PDF terms (c). Ideally the probabilities will sum to 1.0. For both the threshold and class approaches, a random number between 0.0 and 1.0, is generated to determine the estimated indicator for the cell. From the random number (e.g. 0.82), a horizontal line is drawn across to the CDF curve, an a vertical line is dropped from the intersection, to identify the indicator estimate (5).	68
FIGURE 4-4.	This illustration shows the step wise manner in which a grid is kriged using Indicator Kriging in conjunction with stochastic simulation. Grid cells containing sample data (hard data and some types of soft data) are defined prior to kriging. Once these points are defined, the remaining cells are evaluated. To krig an unestimated cell, a random location is selected, evaluated and redefined as a hard data point, then the next undefined cell is randomly selected. This cell selection and estimation process is continued until all grid cells have been visited and defined.	69
FIGURE 4-5.	Graphical method for calculating p_1 and p_2 values for a specific threshold (After Alabert, 1987). Data from CSM Survey Field.	73
FIGURE 4-6.	Graphical method for calculating p_1 and p_2 values for a specific class. Data from CSM Survey Field.	74
FIGURE 4-7.	For both the class and threshold approach, there are two basic types of order relation violations (ORV). a) An individual CDF probability is less than the CDF of a smaller threshold (the CDF is decreasing); this is equivalent to a class having a negative probability of occurrence. This type of ORV is resolved for thresholds by averaging the two CDF's so that they are equal; for classes, a 0.0 probability of occurrence is assigned to the PDF. b) When cumulative probabilities are greater than 1.0, the value is truncated to 1.0 for the threshold approach, while for the class approach, the probability of each class is proportionally rescaled, so that the CDF will sum to 1.0.	76
FIGURE 4-8.	Synthetic data set distribution.	78
FIGURE 4-9.	Reordering indicators in conditional simulation changes the results for an individual simulation grid cell, because the CDF changes along with the indicator ordering, even though the individual components of the PDF do not. Here the indicators from Figure 4.3 have been reordered. The same random number is used, but now, instead of indicator #5 being selected, indicator #4 is selected.	79

FIGURE 4-10.	Realization (#32) pairs for the a) original indicator ordering, b) reversed ordering, and c) arbitrary reordering for thresholds (left) and classes (right). In these realization pairs there are significant differences between the class and threshold results: a) there is more silty-sand in the class realization at location (270, 10); b) sand bisects the silt in the threshold realization at location (100, 100); this not present in the class realization; c) more sand is in the class realization at location (270, 10). The differences between the realization pairs in a, b, and c are expected, because reordering the indicators changes the CDF. 80
FIGURE 4-11.	Realization (#100) pairs for a) original indicator ordering, b) reversed ordering, and c) arbitrary reordering for thresholds (left) and classes (right). These threshold and class realization pairs are similar. The differences between the realization pairs in a, b, and c are expected, because reordering the indicators changes the CDF..... 81
FIGURE 4-12.	Maximum probability of occurrence for each indicator. At hard data locations, the indicator type is known; the probability is 0% for any other indicator type, or 100% for the specified indicator type... 83
FIGURE 4-13.	Maximum probability of occurrence of any indicator. At known data points, the maximum probability of being a particular indicator is 100%. The minimum probability is 33% (100% / # of indicators); at these locations each indicator is equally likely to occur. These maps are useful for evaluating the spatial distribution of uncertainty. 84
FIGURE 4-14.	Frequency of maximum probabilities throughout the grid. At known data points, the maximum probability is 100%. The minimum probability is 33% (100% / # of indicators); at locations where each indicator is equally likely to occur. 85
FIGURE 4-15 a,b.	Difference between a) threshold and class maximum probability maps, and b) threshold and threshold (reversed) probability maps. These maps show areas, where the different series return different averaged results. Differences are largest, although of opposite sign, in red and blue areas. Differences are smallest in green areas. The histograms are useful for identifying the magnitude and distribution of the differences. 86
FIGURE 4-15 c,d.	Difference between a) threshold and threshold (reordered) maximum probability maps, and b) class and class (reversed) probability maps. These maps show areas, where the different series return different averaged results. The histograms are useful for identifying the magnitude and distribution of the differences..... 87
FIGURE 4-15 e,f.	Difference between a) class and class (reordered) maximum probability maps, and b) threshold (reversed) and class (reversed) probability maps. These maps show areas, where the different series return

	different averaged results. The histograms are useful for identifying the magnitude and distribution of the differences.....	88
FIGURE 4-15 g.	Difference between threshold (reordered) and class (reordered) maximum probability maps. These maps show areas, where the different series return different averaged results. The histograms are useful for identifying the magnitude and distribution of the differences. ..	90
FIGURE 4-16.	Test of random number generator: a) scatter of 10,000 sequential random numbers, b) frequency distribution of 10,000 random numbers (equal distribution would put 1% in each class), and c) frequency distribution of 100,000 random numbers (equal distribution would put 1% in each class).	91
FIGURE 4-17.	Coarse grid realization (#32) pair for the original indicator ordering. These grids, are slightly different, due to an ORV occurring at (490m, 30m). Because of the ORV, the CDF's varied between the two methods at this cell, and a random number in each realization selected different material types for the cell. From this point forward, the prior sample data, and prior evaluated cells varied.....	92
FIGURE 4-18.	CSM Survey Field location map.	95
FIGURE 4-19.	CSM Survey Field site map. Dots represent borehole locations. Solid lines identify location of tomography surveys.	96
FIGURE 4-20.	CSM Survey Field borehole and tomography data converted to indicators. The dotted layer delineates the extents of the model grid.	97
FIGURE 4-21.	Individual threshold and class realization #1.	102
FIGURE 4-22.	Individual threshold and class realization #37.	103
FIGURE 4-23 a,b,c.	Maximum probability of occurrence for threshold and class indicators #1, #2, and #3. At known hard data points, probability is 0% for any other indicator type, or 100% for the specified indicator type...	104
FIGURE 4-23 d,e,f.	Maximum probability of occurrence for threshold and class indicators #4, #5, and #6. At known hard data points, probability is 0% for any other indicator type, or 100% for the specified indicator type...	105
FIGURE 4-23 g,h.	Maximum probability of occurrence for threshold and class indicators #7 and #8. At known hard data points, probability is 0% for any other indicator type, or 100% for the specified indicator type.....	106
FIGURE 4-24 a,b,c,d.	Difference between the class and threshold, individual indicators (#1-4) maximum probability of occurrence maps.	107
FIGURE 4-24 e,f,g,h.	Difference between the class and threshold, individual indicators (#5-8) maximum probability of occurrence maps.	108
FIGURE 4-25 a,b.	Maximum probability of occurrence of any indicator. At known data points, probability is 100%. The minimum probability is 12.5% (100% / # of indicators); at these locations each indicator is equally	

	likely to occur (no cells had this minimum probability). These maps are useful for identifying the spatial distribution of uncertainty.	109
FIGURE 4-26 a,b.	Histograms indicate the maximum probability of occurrence of any indicator. At known data points, probability is 100%. The minimum probability is 12.5% (100% / # of indicators); at these locations each indicator is equally likely to occur (no cells had this minimum probability). Class realizations have slightly higher mean indicating lower overall uncertainty.	110
FIGURE 4-27.	a) Spatial distribution of the difference between the class and threshold maximum probability maps; b) histogram of the same information. The positive mean difference indicates the class realizations have a lower level of uncertainty.	111
FIGURE 5-1.	Spatial statistics may vary across a site, such that a single semivariogram model may not be appropriate for the entire site.	114
FIGURE 5-2.	Basic steps involved in the standard kriging algorithm with additional steps needed to implement Zonal Kriging indicated in italics. ..	115
FIGURE 5-3.	Different methods for describing zone contacts: a) sharp, b) gradational, and c) fuzzy.	117
FIGURE 5-4.	The search for nearest neighbors varies with zone boundary type: a) sharp, b) gradational, and c) fuzzy.	118
FIGURE 5-5.	Different forms of ordinary and zonal kriging. (a) Sample data, (b) a traditional Simple Kriged map, (c) one possible zone map from a conditional indicator simulation, (d) sharp transition, (e) gradational transition, (f) fuzzy transition.	120
FIGURE 5-6.	Model definition information for the Yorkshire cross-section: a) actual cross-section sampled with 2m x 1m cells; the locations of b) 7 and c) 10 "well samples;" d) zone definition.	122
FIGURE 5-7.	Exhaustive experimental and model indicator semivariograms for SH/SH-SS threshold (full cross-section) of the Yorkshire data set.	123
FIGURE 5-8.	Exhaustive experimental and Model indicator semivariograms for SH-SS/SS threshold (full cross-section) of the Yorkshire data set.	124
FIGURE 5-9.	Sub-sampled experimental and Model indicator semivariograms for SH/SH-SS threshold (well data) of the Yorkshire data set.	125
FIGURE 5-10.	Sub-sampled experimental and Model indicator semivariograms for SH-SS/SS threshold (well data) of the Yorkshire data set.	126
FIGURE 5-11.	Realizations from single-zone simulation series.	127
FIGURE 5-12.	Realizations from two-zone simulation series.	128
FIGURE 5-13.	Impact of altering the semivariogram nugget independently in the top and bottom zones of a section: a) original simulation section; b) reduced nugget in top zone; c) reduced nugget in bottom zone.	128

LIST OF FIGURES

FIGURE 5-14.	CSM Survey Field location map.	129
FIGURE 5-15.	CSM Survey Field site map. Dots represent borehole locations. Solid lines identify location of tomography surveys.	131
FIGURE 5-16.	CSM Survey Field borehole and tomography data converted to indicators.	132
FIGURE 5-17.	Tomographic cross-section at CSM Survey Field (viewing North-West). Dashed line is approximate location of Fountain / Lyons Formations contact.	133
FIGURE 5-18.	CSM Survey Field hard and soft data distributions (converted to indicators) for the Fountain (a) and Lyons (b) Formations.	134
FIGURE 5-19.	Distribution of hard and soft (Type-A only) data for full data set and for Fountain and Lyons Formations regions.	135
FIGURE 5-20.	CSM Survey Field zone definition.	137
FIGURE 5-21.	Graphical method for calculating p_1 and p_2 values for a specific threshold (After Alabert, 1987). Data from CSM Survey Field.	140
FIGURE 5-22.	Graphical method for calculating p_1 and p_2 values for a specific class. Data from CSM Survey Field. Note the similarities between the <i>Class</i> > 10000 and the <i>Class</i> < 9250 diagrams, and the diagrams in Figure 5.21. They are fundamentally identical. This is always the case for the first and last class and threshold p_1 - p_2 calculations.	141
FIGURE 5-23.	Single-zone and two-zone realization pair #19.	143
FIGURE 5-24.	Distribution of indicators a) in the single-zone realization #19 (Figure 5.23a) and b-d) in two-zone realization #19 (Figure 5.23b). The single-zone realization poorly reproduces original data distribution (Figure 5.19a), whereas the two-zone realization reasonably reproduces the full and individual formation distributions (Figure 5.19a-c).	144
FIGURE 5-25.	Single-zone and two-zone realization pair #27.	145
FIGURE 5-26.	Distribution of indicators a) in the single-zone realization #27 (Figure 5.25a) and b-d) in the two-zone realization #27 (Figure 5.25b). The single-zone realization poorly reproduces original data distribution (Figure 5.19a), whereas the two-zone realization reasonably reproduces the full and individual formation distributions (Figure 5.19a-c).	146
FIGURE 5-27.	Single-zone and two-zone realization pair #44.	147
FIGURE 5-28.	Distribution of indicators a) in the single-zone realization #44 (Figure 5.27a) and b-d) in the two-zone realization #44 (Figure 5.27b). The single-zone realization poorly reproduces original data distribution (Figure 5.19a), whereas the two-zone realization reasonably	

	reproduces the full and individual formation distributions (Figure 5.19a-c).	148
FIGURE 5-29.	Maximum probability of occurrence of indicator #1 for single-zone and two-zone realizations. In the single-zone model, the uncertainty is fairly consistent throughout the site except near hard and soft data locations. There are significant differences between the zones in the two-zone map. The Lyons Formation has a much higher percentage of Indicator #1, and this is reflected in the maximum probability map.	149
FIGURE 5-30.	Maximum probability of occurrence of indicator #6 for single-zone and two-zone realizations. Note, in the single-zone model, the uncertainty is fairly consistent throughout the site except near hard and soft data locations. There are significant differences between the zones in the two-zone map. The Lyons Formation has a much lower percentage of Indicator #6, and this is reflected in the maximum probability map.	150
FIGURE 5-31.	Maximum probability of occurrence of any indicator. In the single-zone model, uncertainty is fairly uniform across the site except near hard and soft data locations. In the two-zone model, the Lyons Formation exhibits significantly lower uncertainty. These maps are useful for identifying the spatial distribution of uncertainty.	151
FIGURE 5-32.	Distribution of the maximum probability of occurrence of any indicator. The two-zone model (b) has fewer low probability cells and more mid-probability cells than the single-zone model (a). This implies the two-zone model has less uncertainty, and is therefore a better solution. The two-zone histogram also has a bi-modal distribution (b), and the populations may be separated by formation (c and d). Most of the model improvement comes from improved definition of the Lyons Formation.	152
FIGURE 6-1.	Detailed flow chart of uncertainty analysis software package.	156

LIST OF TABLES

TABLE 3.1.	Alternative semivariogram models for RMA residual bedrock surface. Range, sill, and nugget terms are in feet.	56
TABLE 4.1.	Kriging weight and nearest neighbors for both class and threshold realization #32. The indicators at each point, and the kriging weights are identical. 94	
TABLE 4.2.	Indicator and associated material type.	96
TABLE 4.3.	The indicator classification is based on sonic velocity measurements, and are matched to approximate hydraulic conductivity's.	98
TABLE 4.4.	Threshold p_1 - p_2 estimates.	98
TABLE 4.5.	Class p_1 - p_2 estimates.	99
TABLE 4.6.	Threshold prior hard and prior soft (Type-A) data distributions. The large difference in threshold 3.5 propagates through threshold 6.5.	99
TABLE 4.7.	Class prior hard and prior soft (Type-A) data distributions.	100
TABLE 4.8.	Threshold semivariogram models.	100
TABLE 4.9.	Class semivariogram models.	101
TABLE 5.1 a,b.	CSM Survey Field threshold (single-zone) and class (two-zone: Fountain and Lyons Formation) semivariogram models.	136
TABLE 5.1 c.	CSM Survey Field threshold (single-zone) and class (two-zone: Fountain and Lyons Formation) semivariogram models.	137
TABLE 5.2.	CSM Survey Field threshold (single-zone) and class (two-zone: Fountain and Lyons Formation) hard and soft data (Type-A only) sample data distributions.	138

When designing a remediation plan for a hazardous waste site where the ground water is contaminated, there are several questions of concern about the contaminant: where is it, where is it going, how long will it take to get there, and what can be done to contain or remove it. To answer these questions, one critical question is, what are the subsurface hydrologic flow conditions. Unfortunately, as important as this question is, a precise answer is difficult to obtain. This is largely because we can only sample a small volume of the site; on the order of one 1/100,000th of the site. Exploratory drilling is expensive and can create new migration routes between contaminated and uncontaminated aquifers or zones, outcrops are generally very limited, and the distribution of the materials that control the hydrologic conditions vary widely. Because of the complexity of the hydrogeologic flow system, and the scarcity of data, there is usually substantial uncertainty in the subsurface description.

To describe some of this uncertainty, this research project develops several geostatistical techniques with the purpose of better defining or reducing uncertainty. A software package is also developed to aid modelers with the data analysis, geostatistics and ground water flow and contaminant transport modeling. The geostatistical techniques developed here are:

- Jackknifing the semivariogram and Latin-Hypercube sampling. These methods are useful for defining the uncertainty associated with the semivariogram model definition and applying that uncertainty in conditional indicator simulation.
- Directional semivariogram models. With traditional kriging techniques, the model semivariogram is defined and oriented in the direction with the longest spatial continuity, thus the longest model range. The spatial correlation, not oriented parallel to the principal axis, is defined by anisotropy factors describing the minor perpendicular axes. This approach is computationally efficient, but it is limiting. The method developed in this research allows the modeler to describe and kriges each orthogonal axis independently.

- Class discrete indicator simulation. Traditional discrete indicator simulation techniques are based on the cumulative probability that a cell is less than a cut-off level or threshold. When non-continuous, discrete data are evaluated, this approach can be non-intuitive. A method where the probability that a discrete indicator class occurs at a cell location is developed here. For the semivariogram analysis, this method is more intuitive; for the simulation process, sensitivities due to indicator ordering are easier to test; and though order relation violations are more common, the remedy is mathematically more appropriate.
- Zonal Kriging. One of the basic assumptions in kriging is the assumption of stationarity (Journel and Huijbregts, 1978). This implies that the spatial variation across the site is approximately constant. For many sites this may be reasonable, but for others, this assumption will lead to significant errors. The zonal kriging method developed in this research project allows the model to be divided into unique and transitional regions.

A collection of program modules was developed to make these techniques practically useful for ground water modelers (as well as researchers from other disciplines). The software package is called UNCERT, for its task is to facilitate uncertainty assessment of ground water problems. It is composed of a number of individual modules: array, block, contour, distcomp, grid, histo, modmain, mt3dmain, sisim, surface, vario, and variofit. These cover a variety of statistical, geostatistical, ground water flow and contaminant transport models, and visualization applications. These run in any UNIX, X-windows/motif environment. All the major applications and tools utilize a user friendly, graphical user interface. Help manuals are also available for each package on-line using HTML (Hypertext Markup Language).

Each of these methods or tools is presented in an individual chapter. These chapters can be read as “stand-alone” documents, though they are all related to geostatistics and reducing uncertainty. Chapter 2 describes Jackknifing and Latin-Hypercube Sampling; Chapter 3, directional semivariogram analysis; Chapter 4, class versus threshold based indicator simulation; and Chapter 5, Zonal Kriging. In the final chapter, Chapter 6, there is a brief description of the UNCERT software package which contains the software described in Chapters 2 through 5, and many other statistical, geostatistical, visualization, and ground water modeling tools. A more complete description of the UNCERT package can be found on the tape (along with the source code) in Appendix A, or on the World Wide Web at <http://uncert.mines.edu/>.

JACKKNIFING & LATIN-HYPERCUBE SAMPLING

Uncertainty is associated with interpretation of the subsurface, and stochastic simulation techniques are incapable of accounting for all the uncertainty, if only a single deterministic semivariogram model is utilized. Jackknifing the sample data bounds the limits of model semivariograms, but typically indicates that a large number of simulations must be conducted to consider the full distribution of possible semivariograms. Latin-Hypercube sampling, particularly when combined with expert opinion reduces the number of simulations that must be created and evaluated. For small data sets, where there is significant uncertainty, this process provides for a more complete assessment of the potential variability of the subsurface and of flow paths for contaminants, given the available data. Such assessment can be used to guide the data collection program and decision making process.

2.1: Introduction

Hydrogeologists recognize that heterogeneity of hydraulic parameters has a major influence on groundwater flow and contaminant migration. Inaccurate description of the subsurface when modeling contaminant transport in groundwater systems can result in selection of inappropriate remedial actions. Identification and characterization of continuous high hydraulic conductivity units of complex geometry, which can dominate contaminant transport, is difficult because the amount of drilling that can be undertaken to characterize the site is less than desired, either due to expense, inaccessibility, or potential for creating pathways for contaminant migration. Thus, the modeler must settle for estimating the range of possible solutions, i.e. the modeler must evaluate the uncertainty in the site definition, and determine how each alternative subsurface interpretation may affect contaminant migration.

At this time, the best approach is to integrate all available data from a site into a range of possible subsurface interpretations and then consider the probability of satisfactory performance of alternative remedial actions. Multiple indicator conditional simulation (MCIS) blends indicator

kriging and stochastic simulation to statistically evaluate the range of possible subsurface geologic configurations. Knowledge of the range of possible subsurface conditions aids the modeler in defining best, worst, and most likely case scenarios as well as the probability of occurrence of particular scenarios given the available data. To date, such simulations have been carried out at the level where the kriging matrix is solved, without incorporation of the uncertainty associated with definition of the semivariogram. Such an approach is based on the assumption that the specified semivariogram models are absolutely correct, but this is often not true, particularly when one considers the limited data usually available at a typical hazardous waste site. In such a situation use of the estimation error to evaluate the accuracy of the kriging is misleading because it appears to characterize uncertainty associated with the result but ignores the uncertainty associated with selection of the semivariogram. The result of a kriging process is based on the definition of the semivariogram. By evaluating the uncertainty in the semivariogram, the greater range of uncertainty associated with the simulated results becomes apparent. Uncertainty in the simulation process can be more completely evaluated by using methods such as jackknifing, latin-hypercube sampling, and expert opinion in defining the semivariogram models to be used for stochastic simulation. These methods are discussed in this chapter.

Data collection is time consuming and expensive. Data collection can be performed more efficiently by examining data as they are collected, preparing experimental semivariograms, plotting estimation errors, and using the results to select subsequent data types and locations. Some projects have used estimation errors to identify areas of greater uncertainty which can be targeted for further data collection, thus optimizing dollars spent in site characterization. Similarly, evaluation of experimental semivariograms as data are collected can guide the data collection program.

Because data are usually limited, the results of kriging can be misleading; depending on the parameters used to define the semivariogram, the same data can yield different results. Although kriging will produce results that honor the data, the estimated values at locations between sample sites are non-unique. The simple examples shown in Figure 2.1 demonstrate this point. These hypothetical, two-dimensional models represent two distinctly different geologic settings that are indistinguishable by examination of only the well data. The sample data in Figure 2.1a and 2.1b are identical. Of the eleven well borings, six are in fine-grained sediments of relatively low hydraulic conductivity (low K) and five are in coarse-grained sediments of generally high hydraulic conductivity (high K). Ideally more data should be collected, but because of cost constraints or constraints on drilling locations, this may be the only data set that can be used. Because different geologic configurations can yield distinctly different contaminant plumes (Figure 2.2), incorrect modeling of the site, or failure to recognize the uncertainty associated with subsurface interpretation, can result in remedial action that does not accommodate conditions at the site.

It is not sufficient to utilize a program that calculates an experimental semivariogram and selects a suitable model. For good results, the modeler must evaluate the uncertainty of the data. In many, if not most cases, there is not enough data available to clearly and absolutely define the semivariogram, but by incorporating the modeler's knowledge or expert opinion about the site, uncertainty may be reduced, possibilities limited, and reasonable results may be identified.

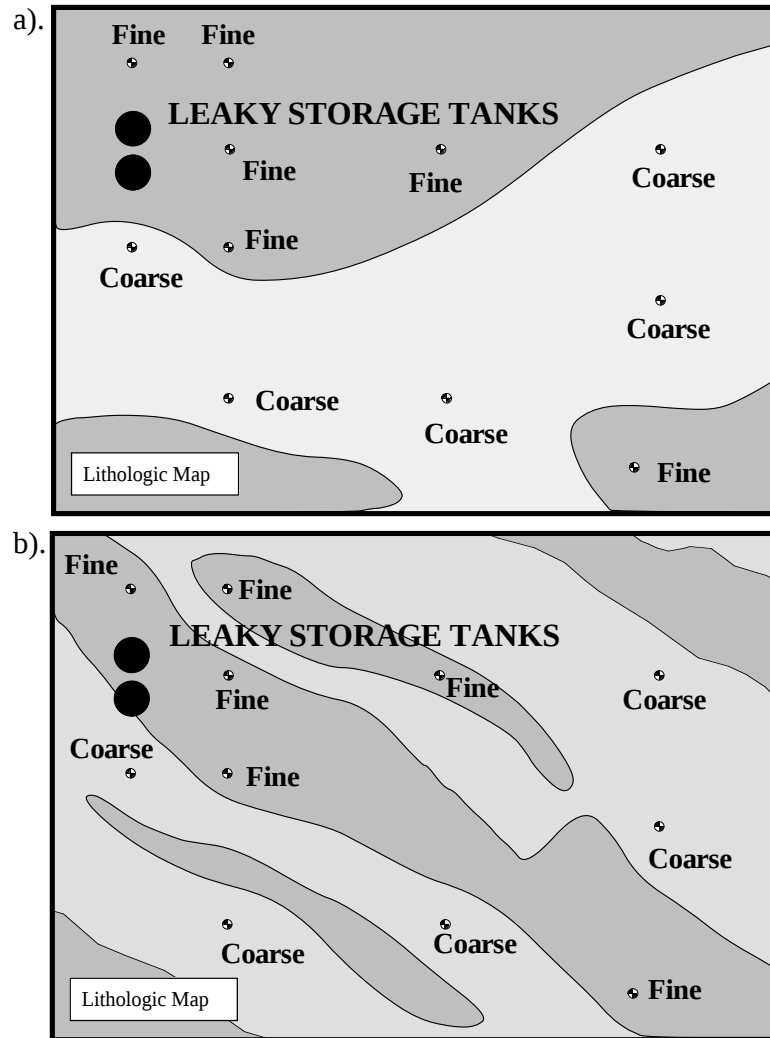


FIGURE 2-1. Borehole data used to interpret the subsurface may not provide a unique solution. In this case, there are eleven data samples; six of fine-grained sediments with low hydraulic conductivity, and five of coarse-grained sediments with high hydraulic conductivity. Although data for each map is identical, the nature of the geology in each map is substantially different. This illustrates that there is uncertainty associated with the interpretation of the character of subsurface at locations that have not been sampled.

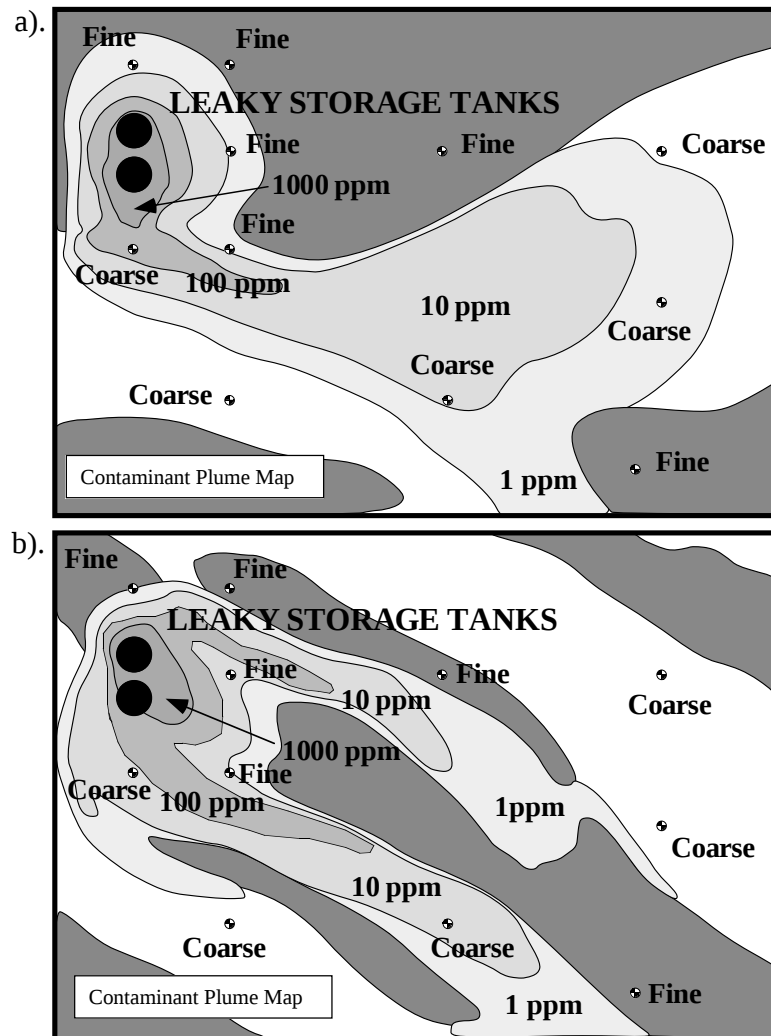


FIGURE 2-2. Contaminants will migrate in different patterns within the two geologic models presented in Figure 2.1. It is important to evaluate the probable alternative scenarios when designing a remediation plan.

2.2: Semivariograms

A semivariogram is a measure of the spatial correlation of a parameter. Samples taken close together are typically more similar than samples separated by larger distances. The semivariogram represents this change in variance with increasing separation distance. The experimental semivariogram ($\gamma^*(h)$) is defined as:

$$\gamma^*(h) = \frac{1}{2N} \sum_{i=1}^N [z(x_i) - z(x_i + h)]^2 \quad (2.1)$$

for a particular lag distance (h), where N = number of data pairs in the search area, and $z(x_i)$ and $z(x_i + h)$ are all the pairs of the N samples within the lag range, h . The search area is defined using a search direction and half angle. The search direction is measured clockwise from North (or the horizontal axis for a cross section) and defines parallel lines along which data of the given lag distance must fall in order to be used in the calculation of the semivariogram (Figure 2.3a). Often data exhibit anisotropy, consequently the experimental semivariogram is calculated in a number of directions. The major axis of the anisotropy is indicated by the search direction of the semivariogram with the longest range (range is the separation distance at which the semivariogram value reaches the population variance and is discussed later). Generally, few data will lie directly along a search direction line, therefore a tolerance angle (defined as the search half angle) is used to include data that are offset from the line (Figure 2.3a). The maximum bandwidth also excludes points that lie well to the side of the search direction. It is useful to note that any search direction accompanied by a search half angle of 90° includes all combinations of orientations of points at each spacing, thus is appropriate when evaluating data with an isotropic distribution.

The model semivariogram, $\gamma(h)$, is a function representing the experimental semivariogram. The distance at which the model semivariogram meets the data set variance is defined as the range (Figure 2.3b). The variance of the sample at a separation distance of zero is called the nugget (Figure 2.3b). This terminology arose in the mining industry where two assays from the same gold sample would sometimes yield markedly different results due to the presence of a gold nugget in one portion of the sample while another portion includes only disseminated gold. The variance of the entire data set is referred to as the sill (Figure 2.3b).

2.3: Indicator Kriging And Stochastic Simulation

One approach for generating alternative subsurface interpretations is indicator kriging combined with stochastic simulation. Indicator kriging differs from simple or ordinary kriging in that a range of parameter values are reduced to discrete indicators (integer values) by defining threshold values.

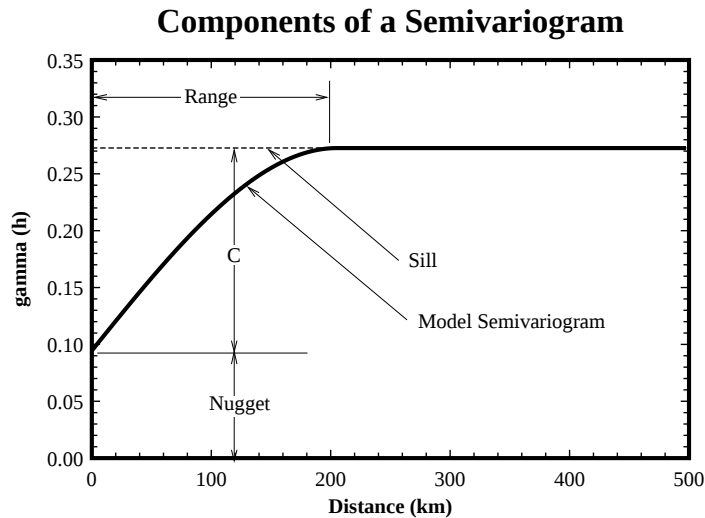
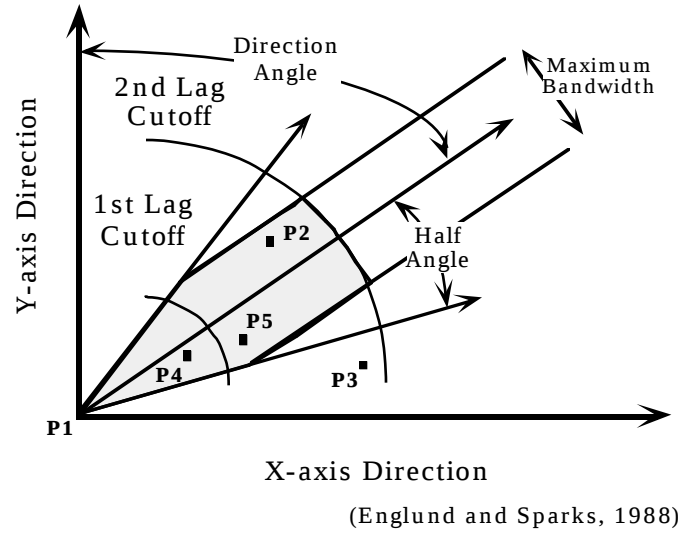


FIGURE 2-3. Features of a semivariogram and parameters defining the search area (after Englund and Sparks, 1988).

For example, materials with hydraulic conductivity less than or equal to 1×10^{-3} may be defined as indicator 1, materials with hydraulic conductivity greater than 1×10^{-3} and less than or equal to 1×10^{-1} may be defined as indicator 2, and materials with hydraulic conductivity greater than 1×10^{-1} may be defined as indicator 3. Indicator description makes it possible to kriging qualitative

parameters such as lithology which could be defined as indicator 1 for silt, indicator 2 for silty-sand, and indicator 3 for fine sand. Suffice it to say, that MCIS allows the modeler to generate multiple interpretations of the subsurface which are distinctly different, but honor all the original data and honor the nature of the model semivariogram. The modeler can use these simulations to assess the uncertainty associated with the subsurface interpretation and to evaluate the affects of the different possible geologic settings on contaminant migration. However, if it is assumed that the range of uncertainty of subsurface interpretations is completely defined by the process, then it is assumed that the model semivariogram accurately represents spatial variation at the site. This assumption is not necessarily correct.

An experimental semivariogram based on the well data from Figure 2.1 is presented as Figure 2.4.

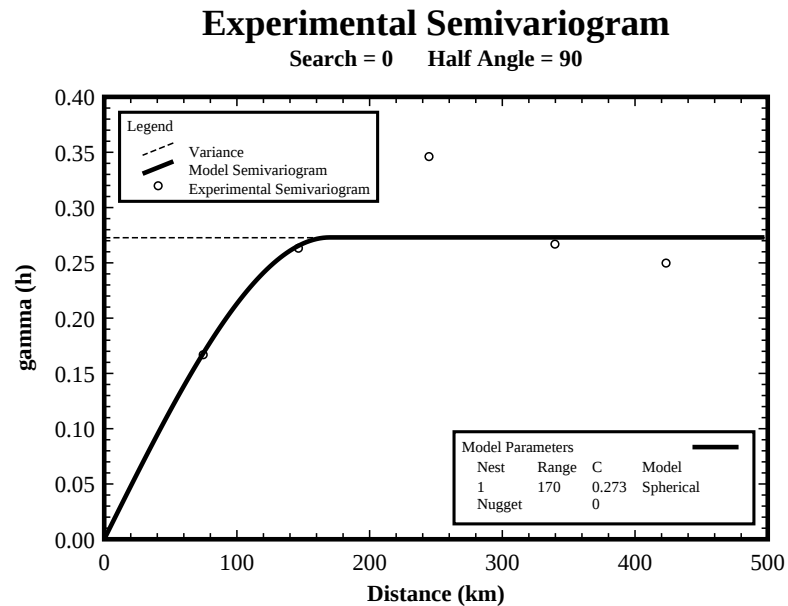


FIGURE 2-4. Experimental and modeled semivariograms developed from the eleven labeled data points in Figure 2.1. A great deal of uncertainty is associated with the modeled semivariogram because of the limited number of data.

For simplicity in demonstrating concepts, only two indicators were employed, one for low hydraulic conductivity materials and another for high hydraulic conductivity materials. Although both models in Figure 2.1 share the same experimental data, semivariograms generated using many data points selected from the two models (1750 points vs. 11 points) are substantially different (Figure 2.5). These semivariograms developed from the extensive data sets illustrate that use of only one experimental semivariogram of the raw data may lead to inaccurate simulations. The actual experimental semivariogram (Figure 2.4) is based on very few points, and arbitrary, simple

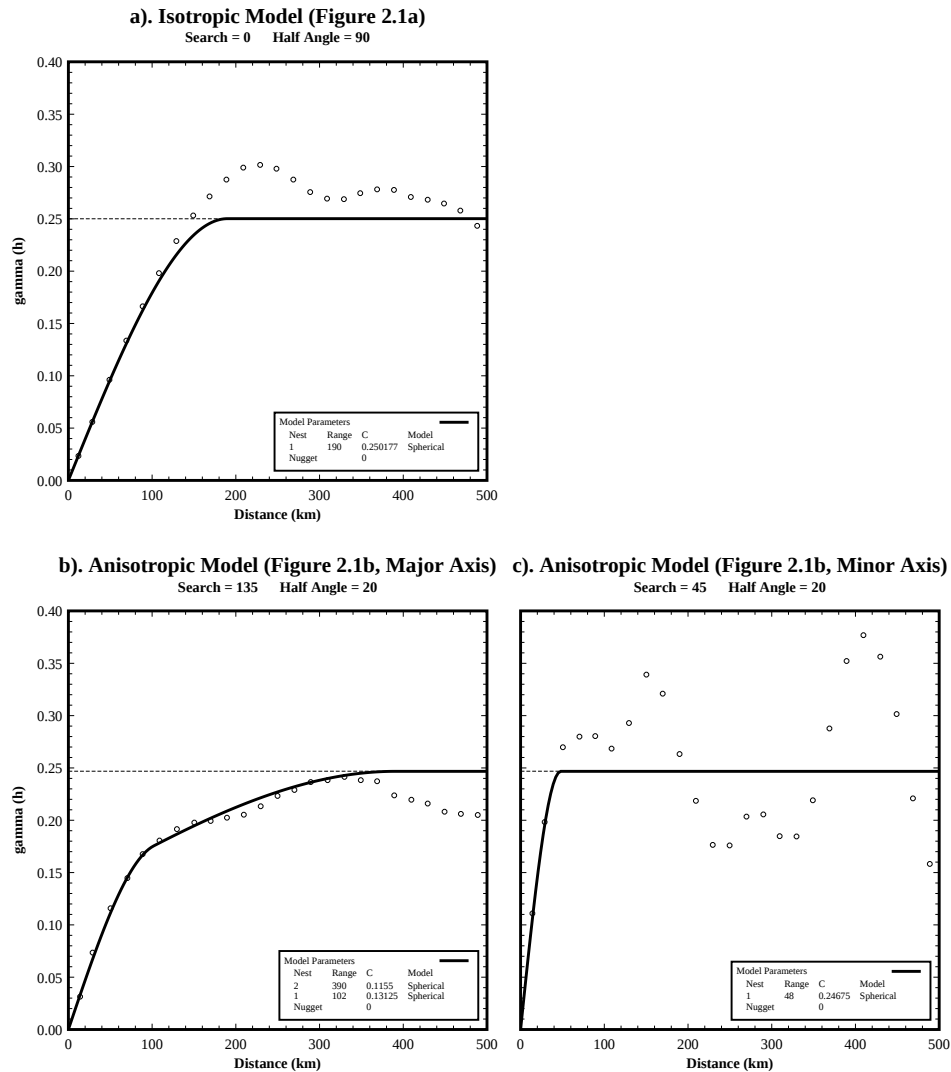


FIGURE 2-5. These experimental semivariograms based on 315 data points from the models in Figure 2.1 were determined by overlaying a regular grid (25' x 25') on each model. The distribution of high and low conductivity materials in Figure 2.1a was determined to be isotropic and is described by the model semivariogram in 2.5a. In Figure 2.1b, the distribution is anisotropic and the major and minor axes of the model semivariogram ellipsoid are shown in 2.5b and 2.5c respectively.

assumptions (in this case, a search direction of 0° with a 90° half-angle). When more restrictive searches were analyzed (i.e. searches with different directions and smaller half-angles), it was not possible to define anisotropy in the data. That is, there are not enough data to develop convincing

statistics to indicate a distinctly longer range is obtained by orienting the semivariogram in a particular direction. Based on the sample data only, using isotropic assumptions, a spherical model was defined for the experimental semivariogram shown in Figure 2.4. The model parameters are:

$$\begin{aligned}\text{Spherical Model: } & 0^\circ \text{ search direction, } 90^\circ \text{ half-angle} \\ & \text{range} = 170 \text{ feet} \\ & C_1 = 0.273 \\ & C_0 = 0.00\end{aligned}$$

where C_0 equals the nugget, and C_1 equals the portion on the data set variance, not due to the nugget. This semivariogram, developed from the 11 data points, contrasts to the semivariograms developed from the extensive data sets in Figure 2.5. An extensive data set taken from the model presented in Figure 2.1a yields a model semivariogram (Figure 2.5a) with the following characteristics:

$$\begin{aligned}\text{Spherical Model: } & 0^\circ \text{ search direction, } 90^\circ \text{ half-angle} \\ & \text{range} = 190 \text{ feet} \\ & C_1 = 0.251 \\ & C_0 = 0.00\end{aligned}$$

This model is similar to the semivariogram model determined using the field data and simple assumptions, and though they are not identical, the simulated results would be similar. Semivariograms developed from the extensive data set for the model shown in Figure 2.1b exhibit a distinctly longer range for an orientation of 135° , yielding a selected model as follows:

(Major-axis) Two-Nested Spherical Model:

$$\begin{aligned}& 135^\circ \text{ search direction, } 20^\circ \text{ half-angle} \\ & a_1 = 102 \text{ feet} \\ & a_2 = 390 \text{ feet} \\ & C_1 = 0.131 \\ & C_2 = 0.116 \\ & C_0 = 0.0.\end{aligned}$$

(Minor-axis) Spherical Model: 45° search direction, 20° half-angle

$$\begin{aligned}& a_1 = 48 \text{ feet} \\ & C_1 = 0.247 \\ & C_0 = 0.0.\end{aligned}$$

where a_i represents the range of each model nest, and C_1 and C_2 represent the non-nugget portion of the data set variance for each nest. Considering the character of the experimental semivariograms

developed from the extensive data set taken from the model in Figure 2.1b, the validity of the semivariogram model based on the limited field data and simple assumptions comes into question. These semivariograms suggest that there may be an anisotropic structure in the model. This anisotropy cannot be identified based on the field data alone, and as a result, a multiple indicator conditional simulation using the semivariogram of Figure 2.4 would not properly represent this alternative interpretation.

The purpose of this example is to illustrate that much of the uncertainty in the kriging process is directly accountable to the definition of the modeled semivariogram. The difficulty, however is that, at an actual site, sparse data often result in unsatisfactory experimental semivariograms. Two techniques, jackknifing and latin-hypercube sampling, can be used to address the uncertainty associated with the semivariograms. In some cases, it may also be reasonable to bias the results with expert opinion. Use of expert opinion in formulating semivariograms may lack statistical rigor, but may be necessary to limit the possibilities. Ground water hydrologists are hired for their expertise; exercising it, as opposed to blindly following a statistical method which we know has limitations, can improve results. Of course, hydrologists must remember to keep an open mind about the nature of the subsurface at a site and not assume the presence of trends nor assume a simple pattern (such as that of Figure 2.1a as opposed to that of Figure 2.1b) without sufficient observation.

2.4: Jackknifing

A method for directly measuring uncertainty, error, or confidence limits associated with an experimental semivariogram is not available, because for each lag, there is only a single calculable $\gamma^*(h)$ value. $\gamma^*(h)$ is calculated as the mean of squared differences for a given lag. Therefore, it is not a mean, but a variance of the data for that lag. Initially it may be thought that $\gamma^*(h)$ could be bounded by estimating the variance of the squared differences about $\gamma^*(h)$. However this is not appropriate because this is the variance about a variance which is calculated, using exactly the same data. Not only is such an approach circular and inappropriate, but as should be expected, the variance about $\gamma^*(h)$ increases with separation distance, yielding no useful information.

To circumvent this problem, a process called jackknifing is used. Jackknifing is a procedure where the experimental semivariogram is calculated with one (or more) data point(s) removed from the data set. By repeating this procedure for every point in the data set, a series of n (n = number of samples) experimental semivariograms is calculated. For each lag distance there are now n $\gamma^*(h)$ values. Using these values, confidence limits can be approximately determined, for the mean $\gamma^*(h)$ at a particular lag. When these are plotted, the error bars define the possible range of the modeled semivariogram (given a specific confidence level; 95% is used in this example). The problem with this method is that each mean value ($\gamma^*(h)$) is correlated with the other mean values ($\gamma^*(h)$) calculated at each specific lag (the same data, except for one point, is being used), therefore the

variance calculations are not strictly correct. However, this technique is not being used to select the best semivariogram model, which it cannot do. Rather it is used to guide the modeler in optimizing further data collection or identifying a likely range of reasonable model semivariograms.

A semivariogram developed by jackknifing the eleven data points from the models in Figure 2.1 (using a 0° search direction with a 90° half-angle) is presented in Figure 2.6. By examining the

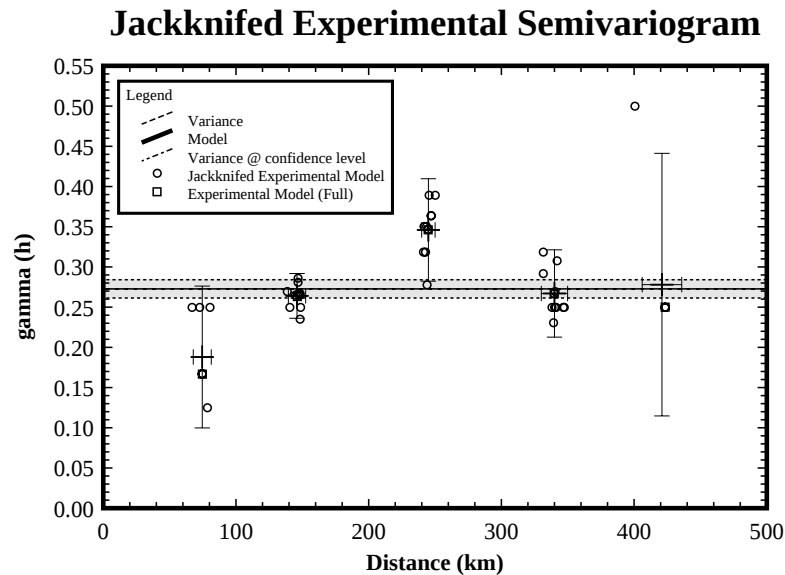


FIGURE 2-6. Jackknifing the eleven data points indicated in Figure 2.1 allows evaluation of uncertainty associated with the semivariogram. The vertical error-bars define the 95% confidence intervals for the mean $\gamma^*(h)$ of each lag. The variance around the mean lag is represented by the horizontal error bars. Each data point represents 1 instance of a jackknifed experimental semivariogram. This experimental semivariogram is based on the assumption of an isotropic material distribution.

error-bars, it can be seen that the modeled spherical range could vary from less than 70 feet to more than 155 feet, but is probably less than 220 feet (error-bars are set at 95% confidence). This compares favorably with the experimental semivariogram shown in Figure 2.5a.

The jackknifed semivariogram does not include the range exhibited in Figure 2.5b where the range of the nested structures, 390 feet, is much greater than 220 feet. This discrepancy occurs because the jackknifed experimental semivariogram in Figure 2.6 is evaluated using all points separated by a given lag distance regardless of their orientation (isotropic conditions were assumed). When the same search windows used to develop the semivariograms of Figures 2.5a and 2.5b are used to develop the jackknifed semivariogram from the limited data set, there is a hint of the character of Figures 2.5b and 2.5c (Figure 2.7). A semivariogram developed using a search direction of 135°

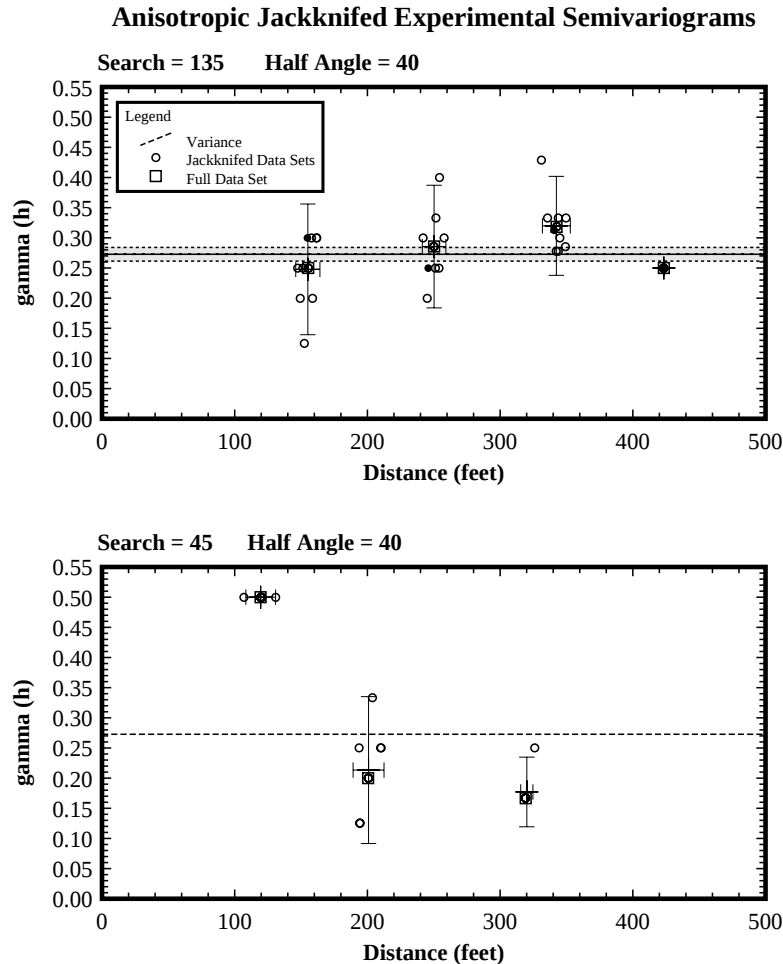


FIGURE 2-7. Although anisotropy cannot be identified by evaluating single semivariograms of the eleven data points, anisotropic character is hinted at when the same data are jackknifed along specified search directions. For a search direction of 45° , the range is likely to be less than 100 feet. In the 135° search direction, the range is likely to be greater than 150 feet, and possibly more than 500 feet. The anisotropy defined in Figure 2.5b-2.5c cannot be determined from the eleven data points, but its possibility is indicated by the data. Symbols are described in the caption of Figure 2.6.

and a 40° half-angle (the initial search used 20° half-angle, but too few pairs were found to be useful) is presented in Figure 2.7a. The range of this semivariogram cannot be determined from the data, but it is likely to be less than 500 feet (use of the extensive data set suggests the range is on the order of 390 feet). A semivariogram developed using the perpendicular search direction of 45° with a 40° half-angle (the initial search used 20° half-angle, but again too few pairs were found to be

useful) indicates that the range in this direction is probably less than 120 feet (use of the extensive data set suggests the range is approximately 48 feet). It would be difficult to justify these last two experimental semivariograms, or to identify them without first knowing the exhaustive data set, but the fact that even limited data contain a hint of the underlying structure is important.

The jackknifed experimental semivariogram (Figure 2.6) also suggests that eleven data points are not enough to correctly define the model semivariogram. The data are not even sufficient to determine if the drilling pattern is tight enough to be within the range of the local variance, as indicated by the fact that the upper limit of the uncertainty bars associated with the smallest sample separation falls above the total (population) variance (the sill). This suggests that further drilling (data collection) is required. Given an increasing number of samples, the jackknifed lag variances will decline (Figure 2.8), and ideally, a jackknifed semivariogram will appear more like that shown in Figure 2.8c. Unfortunately, uncertain semivariograms are the norm rather than the exception as indicated by the work of Shafer and Varljen (1990), and the erratic nature of published indicator semivariograms of lithology. The lack of variation in the experimental jackknifed semivariogram illustrated in Figure 2.8c allows the modeler to clearly define the model semivariogram. If the experimental jackknifed semivariogram of lithology at a site had the character of Figure 2.8c, it could be argued that, too much money was expended collecting data; the semivariogram could have been modeled adequately with fewer data (Figure 2.8b). In this case, if jackknifed semivariograms had been calculated while data were being collected, the characterization program could have been terminated sooner or redirected to focus on collecting data to reduce uncertainty in poorly characterized areas of the site (as indicated by areas of high kriging estimation error), as opposed to collecting data that would further define the semivariogram, thus saving time and money.

2.4.1: Additional Comments About Jackknifing

Several other concepts should be considered when using jackknifing in a semivariogram analysis. First, because data points are being removed from the data set to calculate the experimental semivariogram, the variance, and therefore the sill, will generally increase slightly. Second, when a single experimental semivariogram based on all the data is calculated, the results may appear to be easily modeled. However it is difficult to differentiate an experimental semivariogram that represents the true nature of the site, from one that is the product of a fortunate lag selection. Jackknifing provides error-bars which give the modeler insight on the level of confidence which can be attributed to the modeled semivariogram. Finally, jackknifing should not be considered for every data set where the experimental semivariogram is poorly behaved. Jackknifing computationally is very expensive. For N data samples, $N + 1$ semivariograms must be calculated. As N increases by one, the computational effort to calculate a single semivariogram doubles. As N gets into the hundreds, particularly thousands, the time to compute the uncertainty for a single experimental semivariogram could take days, weeks, or even longer.

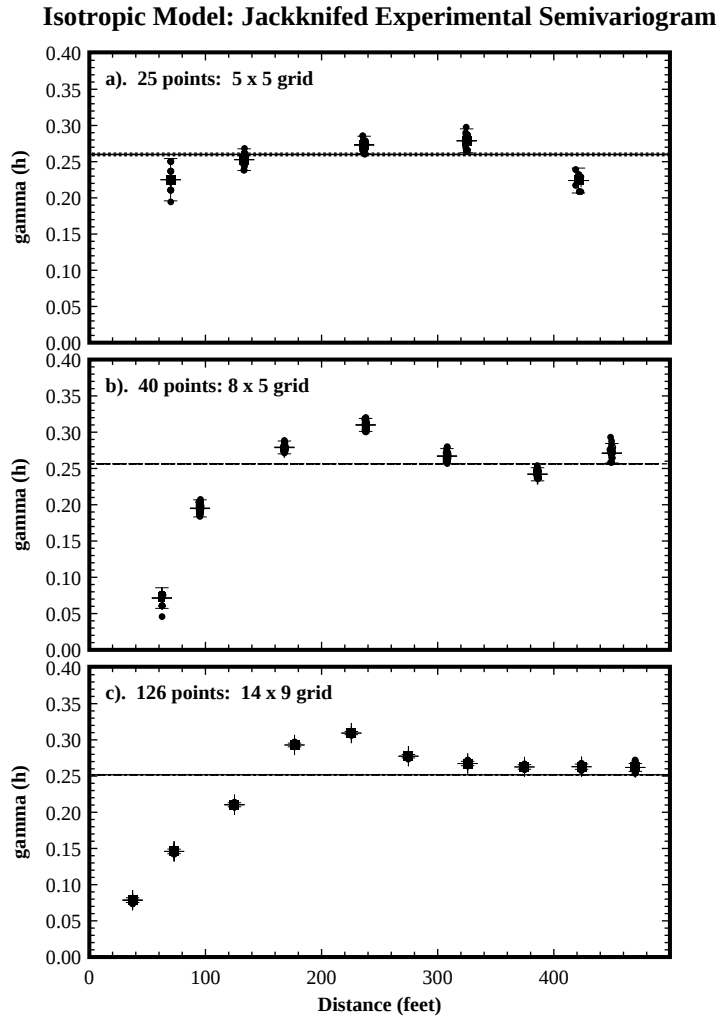


FIGURE 2-8. When a substantial amount of data are collected, the experimental semivariogram may be clearly defined. In this jackknifed simulation, there is little uncertainty in the lag means, and there would be little uncertainty in defining the model semivariogram.

2.5: Latin-Hypercube Sampling

Once the statistical distribution of experimental semivariograms has been calculated, semivariograms can be fit through the zone defined by the error-bars. The objective is not to make

a single best estimate of the character of the subsurface (i.e. a single semivariogram), rather the objective is to select model semivariograms representative of the range of possible conditions at the site. This range of semivariograms is used with the original data to conduct indicator kriging and stochastic simulation to generate multiple interpretations of the subsurface. One approach is to use Monte-Carlo techniques and randomly select, for example, 100 model semivariograms that fall within the range of reasonable solutions (Figure 2.9a). This might appear reasonable, but, for

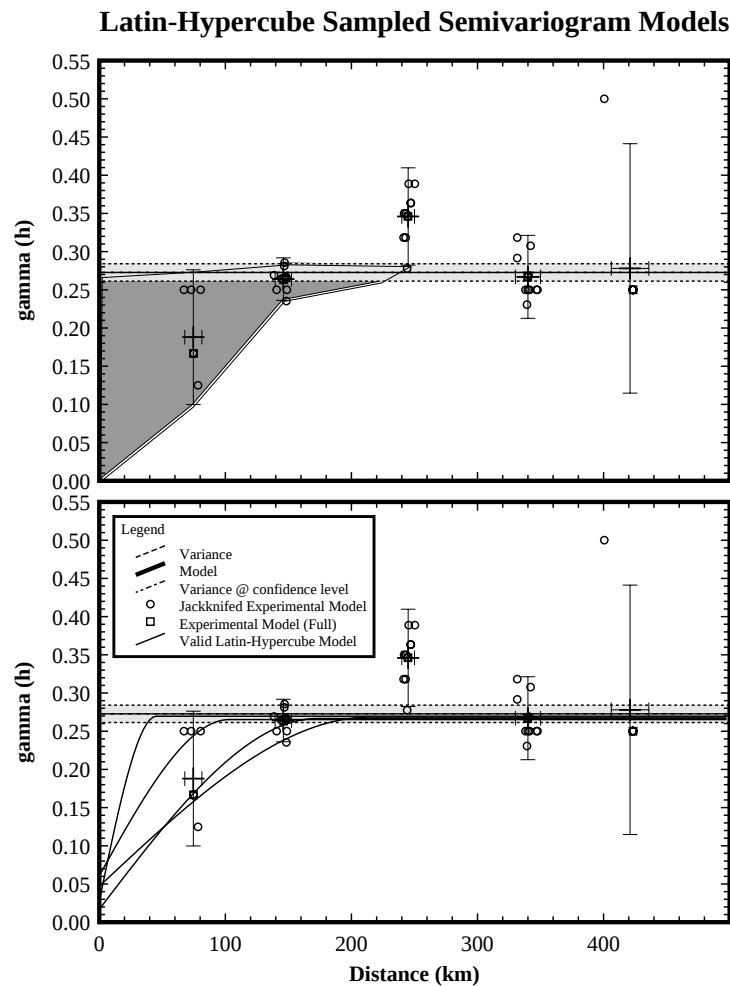


FIGURE 2-9. Reasonable models must be selected from the shaded region in 2.8a to represent the “flavor” of the alternative interpretations of the data. Four model semivariograms with a nugget selected from the lower quartile of possible nugget values are shown in 2.8b. The ranges of the four semivariograms are selected to represent each of the quartiles of possible ranges. Sixteen models would be used to represent the distribution of semivariogram models for the isotropic case. Symbols are described in the caption of Figure 2.6.

example, expert opinion of conditions at the site may indicate that models generated with nuggets approximately equal to the sill or ranges near zero are unreasonable or unrealistic, even though the jackknifed experimental semivariogram in Figure 2.6 indicates such semivariogram models of the site are possible interpretations.

An alternative approach to random selection of a large number of possible semivariogram models is to use latin-hypercube sampling. This reduces the number of simulations required to insure that the "flavor" of all alternatives is addressed. For this example, one might suggest the nugget must fall within one of four equiprobable regions, and the range also must fall within one of four equiprobable regions. The actual nugget, or range within each region is then randomly calculated (Figure 2.9b). This allows sixteen model semivariograms to be calculated for an isotropic model. For an anisotropic model, the direction and magnitude of the anisotropy can be restricted similarly. This, however requires many more simulations. If the anisotropy factor between the major and minor axis is evaluated at four ratios (e.g. 1.0, 0.5, 0.25, and 0.125 or some other ratios as determined from jackknifing the data to obtain a semivariogram in the direction of the minor axis of anisotropy), the number of semivariograms is increased to 64. If the search directions, 0° to 180° , are divided into four directions (0° , 45° , 90° , and 135°), the number of semivariograms is increased to 256.

This approach can yield a daunting number of simulations, many of which will bear little resemblance to one another if the data set is small. Such a situation results in the obvious conclusion that some data sets provide so little information about a site that more data should be collected before further assessment is undertaken. If the data are more abundant, the range of possible models will be constrained, and the simulated models may represent a modest range of possible subsurface interpretations. If the jackknifed semivariogram has small error-bars, as in Figures 2.8b and 2.8c, the entire process of using a variety of semivariograms for simulation of one site can be omitted because the process is not likely to indicate a larger uncertainty associated with the interpretation of such well characterized sites.

Recall that the objective of this approach is not to make a single best estimate of the subsurface interpretation, but to evaluate the possible range of subsurface character based on available data. From a purely mathematical approach this may be computationally intractable, however incorporation of expert opinion into the process makes it possible to limit the reasonable alternatives.

2.6: Expert Opinion

Thus far, only mathematical techniques for describing the subsurface have been discussed, and only field data from wells at the site have been used for interpretation of the subsurface configuration. Two points are important to consider; 1) these mathematical techniques do not necessarily honor geologic laws, and 2) hydrogeologists often know more about the site than the borehole data suggest.

The process of stochastic simulation uses probabilities to estimate a value at a grid location. Unfortunately, these probabilities are based on measured values near that location and, consequently, geologically impossible configurations can be simulated. For example, the “law of original horizontality” and the “principal of stratigraphic superposition” are readily broken. Eventually techniques that incorporate these concepts into stochastic simulation will be developed. Until that time, such simulations must be identified, deemed unreasonable, and discarded.

Although creation of such geologic fallacies cannot be prevented with the current simulation process, the simulations can be improved by incorporation of geologic knowledge from analog sites. An expert can infer more information about the site than is evident in the borehole data. For example, an expert may know that sand lenses in the area tend to be between 10 and 25 feet thick. The borehole data at the site may be too sparse to determine this range of thickness, but knowledge from analog sites in the area may render it reasonable to assign a range of 10 to 25 feet to the vertical modeled semivariogram. Although such action is not based on data from the site, knowledge of analogs adds information to (decreases uncertainty associated with) the simulation process. If the site is made of horizontally bedded alluvial deposits, there is no reason to run simulations which assume the material distribution is isotropic. In such settings, units are generally continuous for greater distances horizontally than vertically. The modeler may be able to confirm the presence of layered anisotropy by demonstrating that semivariograms with different search directions and limited half-angle and bandwidths have the potential to have different ranges. Even if the indications are sketchy, due to scarcity of data, the modeler can limit the simulations to produce only reasonable interpretations given the local geology. Similarly, anisotropy may be present in lateral directions and geologic knowledge of directional trends of lenses or channels may be used to limit the number of orientations considered for semivariograms which will, in turn, limit the number of simulations that must be undertaken.

There is little reason to evaluate solutions that are mathematically possible, but geologically improbable. Discarding geologically improbable solutions adds “bias” to the results that may have to be defended later. However omission of the bias means that we do not use all the information available to us. When expert opinion is used wisely, the bias is likely appropriate, and will speed the site evaluation, thus limiting exploration and analysis costs.

2.7: Results

Four examples are presented to illustrate the process of multiple indicator conditional simulation using latin-hypercube sampling of a jackknifed experimental semivariogram. The differences in these simulations demonstrate the variability of subsurface interpretation that is obtained using the limited data given in the example in Figure 2.1.

These simulations were created using the MCIS code ISIM3D. The map area was modeled in two-dimensions using a 50x35 grid, with ten foot square grid cells.

Simulations resulting from use of the modeled semivariogram using the extensive data set (Figure 2.5a) are presented in Figure 2.10. These simulations differ significantly from the model because

Isotropic: Extensive Data Set Model Semivariogram

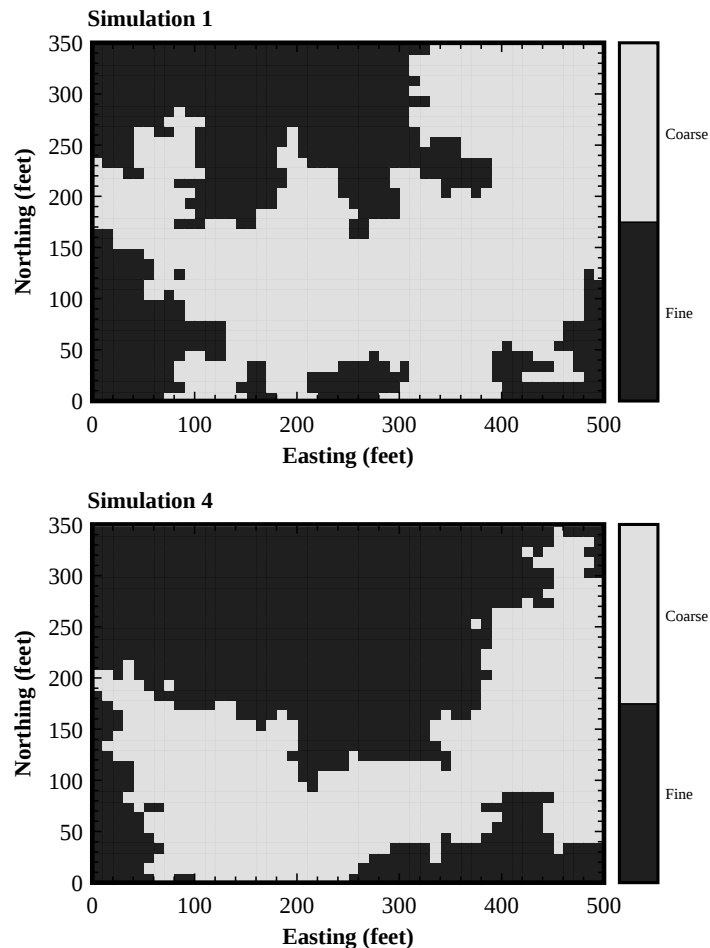


FIGURE 2-10. These two simulations were generated assuming isotropy and using the model semivariogram developed from the extensive data set and illustrated in Figure 2.5a. The solutions are a reasonable approximation of the map in Figure 2.1a.

only the 11 data points were used to condition the simulation. Simulations presented in Figure 2.11 are based on a model semivariogram ($a_1 = 115'$, $C_1 = 0.25$, $C_0 = 0.0$) sampled from the jackknifed experimental semivariogram shown in Figure 2.6. Although neither simulation (Figure 2.11a or 2.11b) is identical to the model in Figure 2.1a, they are reasonable approximations considering the limited data. The simulation in Figure 2.11b is particularly close to the model of Figure 2.1a. The

Isotropic: Jackknifed Bore-Hole Data Set Model Semivariogram

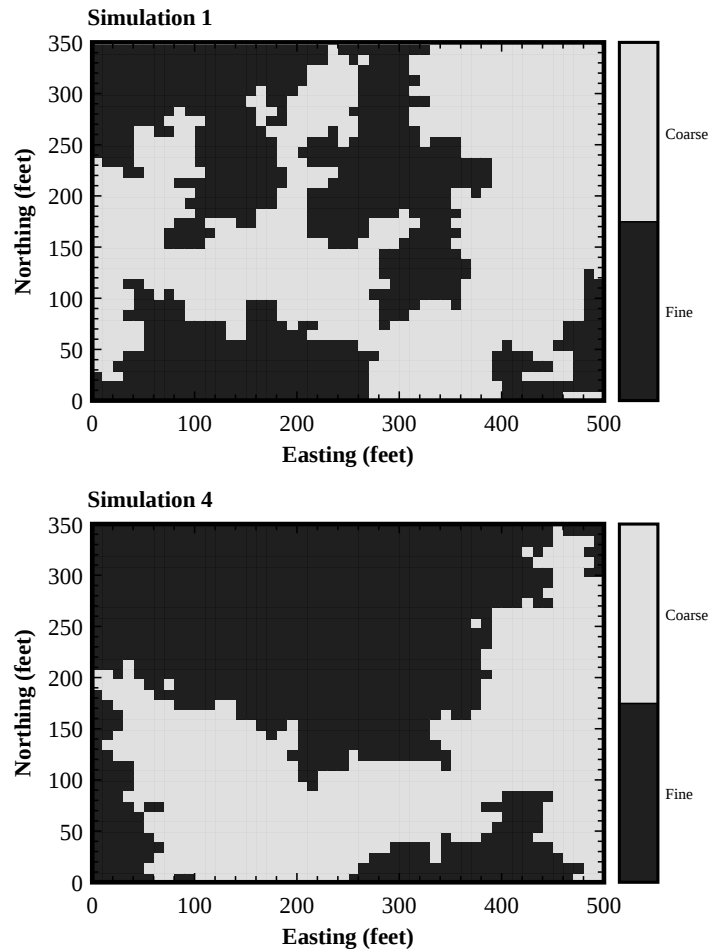


FIGURE 2-11. These two simulations were generated assuming isotropy and using a latin-hypercube sample from the jackknifed model semivariogram ($C_0=0.0$, $C_1=0.25$, $a_1=115'$) developed from the eleven data points and illustrated in Figure 2.6. The solutions are a reasonable approximation of the map in Figure 2.1a, and are very similar to those generated in Figure 2.10. Much of the reason that the simulations in Figure 2.10 and 2.11 are similar is that the same random path through the grid was used to simulate 2.10a and 2.11a and another path was used to simulate 2.10b and 2.11b.

appearance of the resulting simulation is rather insensitive to the choice of range (compare Figure 2.11 with Figure 2.10 which was generated using a range of 190' vs. 170'). Both the experimental semivariograms (Figure 2.5a and Figure 2.6) were developed based on an assumption of isotropic material distribution. The simulations in Figure 2.10a and Figure 2.11a are also similar because

the same random path (same random number seed) was used to generate all of the '(a)' simulations in Figures 2.10-2.13. A different path was used to generate the '(b)' simulations. These isotropic

Anisotropic: Jackknifed Bore-Hole Data Set Model Semivariogram

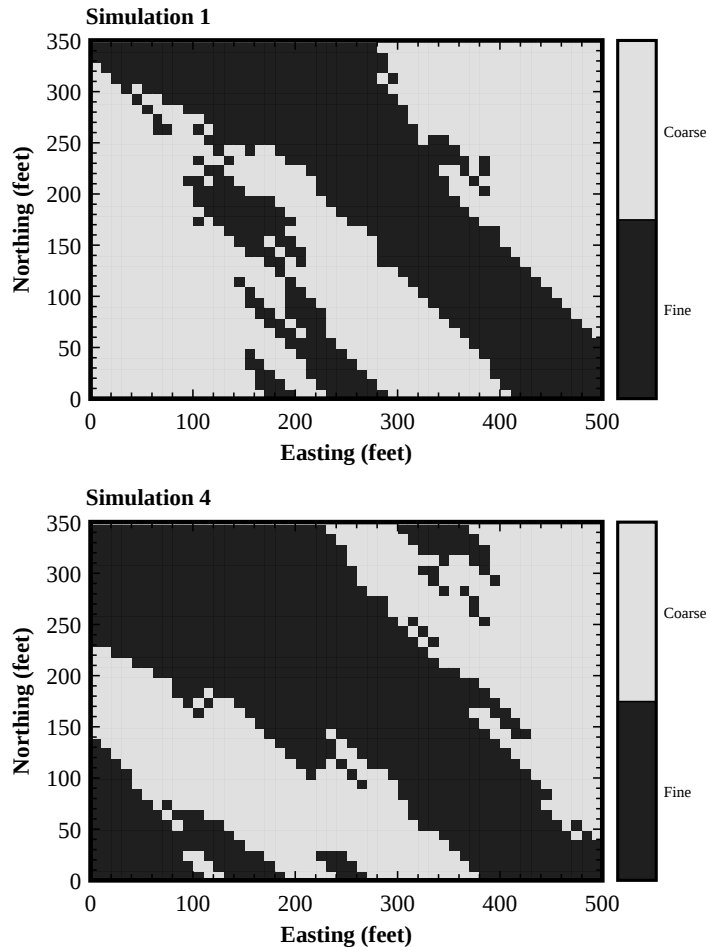


FIGURE 2-12. These two stochastic simulations were generated assuming anisotropy using the jackknifed model semivariogram based on the eleven data points and illustrated in Figure 2.6. The latin-hypercube technique was applied and these are two simulations of a potential 256, as described in the text. Even though the geologic models presented in Figure 2.1 are different, use of jackknifing and Latin Hypercube sampling can produce both configurations from limited data. These solutions are a reasonable approximation of the map in Figure 2.1b. Unfortunately, the method will not indicate whether these simulations or the simulations in Figures 2.10 and 2.11 are the most likely because the data are not sufficient to draw such a conclusion.

Anisotropic: Extensive Two-Nested Data Set Model Semivariogram

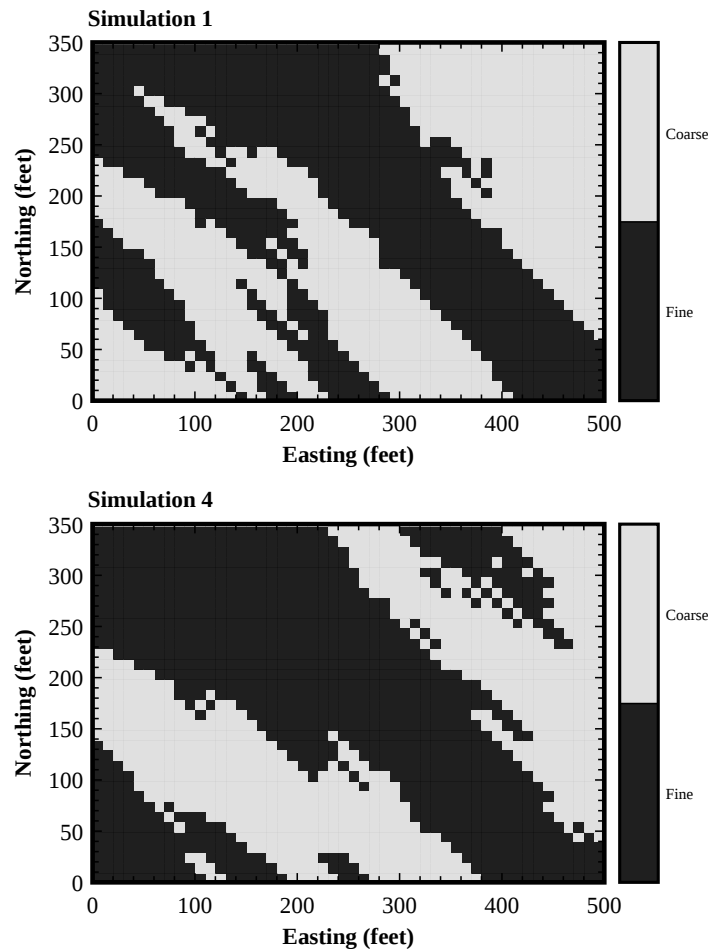


FIGURE 2-13. These two simulations were generated assuming anisotropy using the extensive model semivariogram based on the extensive data set and illustrated in Figures 2.5b-2.5c. The solutions are a reasonable approximation of the map in Figure 2.1b, and are very similar to those generated in Figure 2.12, indicating that extensive data are more important to determining the character of the semivariogram than they are to conditioning the simulation.

simulations bear little resemblance to the model in Figure 2.1b which is a viable interpretation of the data from the 11 field measurements. This inability to represent the full range of possible interpretations is not unexpected.

If expert opinion indicated that the site would be expected to exhibit the locally observed NW-SE trend of high and low hydraulic conductivity deposits, then the simulations presented in Figures 2.10 and 2.11 could be assumed to be less probable. They would be superseded by the probability of occurrence of anisotropic representations of the site. If such expert opinion were not available the two alternative configurations would have to be considered equally likely to occur. The two simulations in Figure 2.12 were generated in the latin-hypercube sampling process, using one of the semivariograms that would fall in the shaded area in Figure 2.9a with a range between 120 and 180 feet (third quartile estimate of range), a nugget between 0.0 and 0.061 (first quartile estimate of the nugget), an anisotropy factor of (minor to major axis) 0.125, a major axis orientation of 135° , and using different random paths through the grid. Although they are not identical to the model in Figure 2.1b, they mimic its nature. When using the range, sill and nugget terms identified by the semivariogram developed from the extensive data set (Figures 2.5b and 2.5c), the simulation results (Figure 2.13) are not significantly different from the simulation results (Figure 2.12) obtained using the jackknifed semivariogram (Figure 2.7), indicating that an extensive field sampling would not improve the character of these simulations but might improve the certainty of occurrence of units with a 135° orientation. That is, more data will improve the certainty of the semivariogram having a given orientation whereas the jackknife approach only indicates the possibility of units having that orientation. Of course, a larger data set improves conditioning of the simulations.

The simulations presented in Figures 2.10-2.13 demonstrate that correct definition of anisotropy is important in order to capture the character of the site. Similar results in paired simulations also suggest that the differences in model ranges are less important than the assumption of isotropy.

2.8: Conclusions

A great deal of uncertainty is associated with interpretation of the subsurface, and simulation techniques are incapable of accounting for all the uncertainty if only a single deterministic semivariogram model is utilized. Typically there are not enough data available at hazardous waste sites to adequately define a single model semivariogram in a rigorous statistical basis.

By jackknifing the data to determine a reasonable range of model semivariograms, and using latin-hypercube sampling and incorporating expert opinion to limit the required simulations, the uncertainty associated with the subsurface interpretation can be more completely assessed utilizing a reasonable amount of simulations. Unfortunately the uncertainty may be so great that little can be concluded about site. However, this is important information because it indicates that more data must be collected before conclusions are made about the site. Given one data sample, one can begin to make interpretations of the site, but quantifying the uncertainty associated with those interpretations is important.

The method presented herein is useful when a significant amount of uncertainty is associated with the experimental semivariogram. If the uncertainty is small, the process only adds unnecessary work. For small data sets, where there is significant uncertainty, this process may be the only way

to correctly assess the potential variability of the subsurface, and evaluate potential flow paths for contaminants.

Although the application considered herein pertains to indicator conditional simulation, evaluation of the uncertainty associated with a semivariogram is important whenever a semivariogram is used. Jackknifing is a practical tool for relatively small data sets, but for large data sets, the computational intensity of the jackknifing process may make the process unmanageable.

VARIATION OF SEMIVARIOGRAM MODELS WITH DIRECTION

When developing semivariogram models, it is often difficult to fit a model semivariogram to both the principle-axis and the minor-axes using traditional methods with anisotropy ratios. Currently a single model (possibly nested) is modified with anisotropy factors; these represent the relative range of the semivariogram for all three orthogonal axes. This technique is restrictive, and this discussion presents a method for relieving these restrictions by defining different semivariogram models, independently for each axis. These will be referred to as directional semivariograms. The process increases the kriging processing time by 80% to 200%, but the method offers the modeler greater flexibility, and simulations or estimations that are more representative of the site, because the spatial variation of the data can be more precisely defined.

3.1: Introduction

Semivariogram modeling is the foundation for much geostatistical analysis, and can also be the most difficult and time consuming portion of the analysis. In part, this is due to the computationally intensive calculations, but it is also due to the difficulty in defining semivariogram models which reasonably honor the experimental semivariograms in the principle and minor search directions. With the current techniques that use anisotropy factors (Englund and Sparks, 1988; Journel and Huijbregts, 1978; Deutsch and Journel, 1992), often it is not possible to model all the orthogonal experimental semivariograms exactly. Consequently compromises are required for the definition of one, or even all of the models. If the compromises are not too substantial, then this approach is acceptable, because, generally the kriged results are relatively insensitive to minor changes in the semivariogram. Though this insensitivity offers some comfort, it is not particularly satisfying.

This chapter describes a procedure, which allows the modeler to define a unique semivariogram model for each orthogonal axis of the experimental semivariogram. The algorithm uses components of each model to determine $\gamma(h)$ values between the axes. Anisotropy factors are not used; rather the modeler specifies the number of nests, sill and range components, and model

structure types independently for each axis. The only requirements are 1) the nugget must be the same for all models, and 2) the total sill must be the same at infinity. These two requirements are not particularly restrictive. Requiring the nugget to be the same is reasonable, because at zero distance, direction is irrelevant. The requirement that the total sill components are equal ensures that the kriging matrix is non-singular. If different sills are desired, then this requirement is met by defining an arbitrarily large range for the final nest to make up the balance of the sill component. The error induced by the final nest has no effect on the area of interest.

This technique allows the modeler to honor the results of the experimental semivariogram analysis, thus it is easier to model the data set and the results are more accurate. However, the calculation of $\gamma(h)$ is substantially more complex than traditional methods, therefore the method requires computational effort. The additional effort is comparable to the computational effort required for the search procedure and matrix solution portions of the kriging algorithm so, overall, the task is only increased by about 80% to 200% (based on observed differences in computation time for example data sets). This is acceptable, because the semivariogram model preparation is simplified, and the simulations or estimates should more closely honor the spatial statistics of the site.

3.2: Previous Work

Many techniques have been developed to estimate values of a variable at locations between sample points. These techniques are all based on the assumption that properties at unsampled locations are related to the properties at nearby points where samples have been taken. Some techniques are inaccurate due to assumptions related to the spatial variation and the relative importance of nearby data. For example, the inverse-distance method states that surrounding data (n = number of samples) have less importance with increased distance:

$$g_i = \frac{\sum_{i=1}^n \left(\frac{x_i}{d_i^p} \right)}{\sum_{i=1}^n \left(\frac{1}{d_i^p} \right)} \quad (3.1)$$

The rate (p) at which increasing distance (d) reduces the influence of a neighboring sample value (x_i) is subject to debate. Various factors for p have been suggested; 1 (linear), 2, $1/x$, based on the modelers previous experience with the technique, and its performance at similar sites. Consequently the results are subjective.

Kriging eliminates much of this subjectivity by utilizing the semivariogram as the spatial weighting function. The variance of the data and the rate of change in variance with direction and distance can be defined with the equation:

$$\gamma(h) = \frac{1}{2N} \sum_{i=1}^N (x_i - (x_i + h))^2 \quad (3.2)$$

where $\gamma(h)$ describes spatial variance of all data pairs separated by a distance h . N is the number of pairs separated by the distance h , and x_i and x_i+h are the values at two points in the pair. Experimental semivariograms are complex, and for practical reasons are represented with one or more functions selected from a limited number of model types (e.g. spherical, exponential, Gaussian). These models are used because they guarantee the matrices in the kriging solution will be positive definite (i.e., the matrix is not singular). Even with these constraints, the semivariogram is a powerful mathematical tool for describing how a variable varies in space at a particular site.

Although semivariograms could be defined for an infinite number of directions, for practical reasons, variation is only defined along the principle orthogonal axes (X, Y, Z), creating an ellipsoid. As defined here, the X-axis is equivalent to the direction with the longest range (the principle axis), and the Y and Z-axes (orthogonal to X), have shorter, though not necessarily equal ranges. Although modeling and solving a more complicated system is theoretically possible, it would be extremely expensive computationally. To further simplify the solution, the semivariogram models for the Y and Z-axes have traditionally been described using anisotropy factors related to the X-axis. This simplification is used extensively in current kriging models (Deutsch and Journel, 1992; Gómez-Hernández and Srivastava, 1990), because it is computationally efficient, however it compromises accuracy and requires more time of the modeler when the same semivariogram model does not fit the experimental semivariogram in all directions. The technique presented here allows the modeler to specify unique semivariogram models for each axis.

3.3: Theory

Two steps of the kriging process are modified to incorporate directional semivariograms into the kriging algorithm: 1) the search for nearest neighbors, and 2) the calculation of the covariance components of the kriging matrix.

The first step in estimating the value for a grid location is to find the influential neighboring points. For isotropic situations the closest sample points are the best estimators. For anisotropic situations, the best estimators are those points with the smallest spatial variance calculated from the model semivariogram ($\gamma(h)$). Using anisotropy factors, the sample point locations are transformed to equivalent isotropic space, using a simple transformation and rotation, based on the orientation of the principle model axis, and the anisotropy factors of the minor-axes. Once transformed, the estimation variance is solely a function of the distance between the grid location and the sample point, therefore $\gamma(h)$ doesn't have to be calculated. Conventional techniques (Deutsch and Journel, 1992; Gómez-Hernández and Srivastava, 1990), use Pythagoras' Theorem to find the closest points. When directional semivariogram models are used, direction as well as distance is important. When

different model equations, sills, and structures are used for the orthogonal axes, a simple transformation and rotation is not possible, because the magnitude of $\gamma(h)$ is not solely related to distance (Figure 3.1). For this reason, $\gamma(h)$ must be calculated for each sample in the search

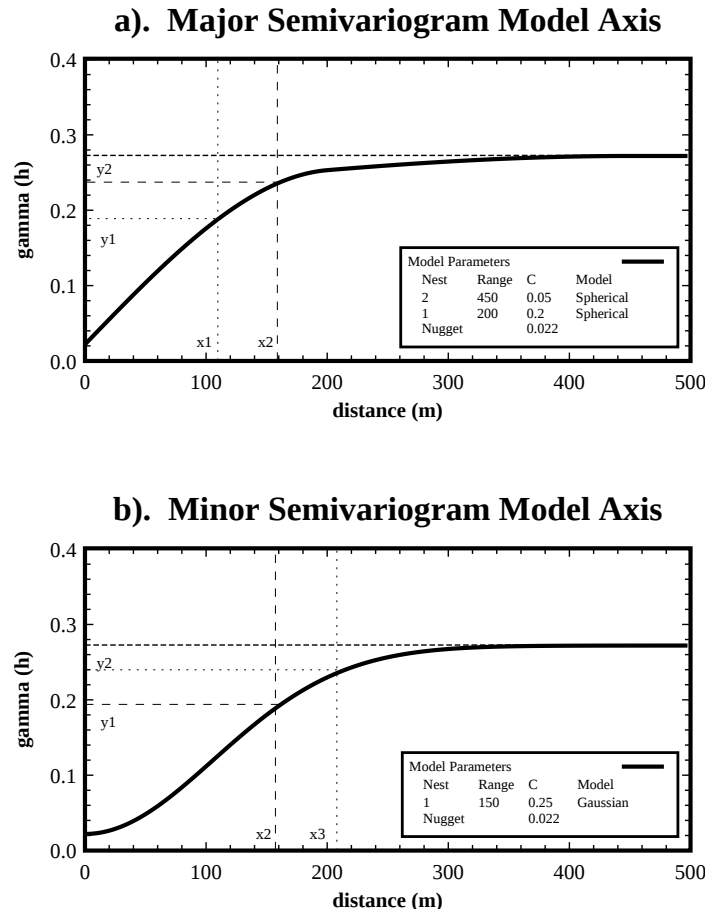


FIGURE 3-1. When directional semivariograms are used, distance alone does not determine the most influential neighboring points. In this example, all points in the minor model axis direction (b) that are separated by less than x_2 (158 m) have smaller $\gamma(h)$'s than points separated by x_1 (109 m) on the major-axis (a). The same is true for x_3 and x_2 respectively.

neighborhood, and those points with the smallest $\gamma(h)$ values are the best estimators.

Once the most influential neighboring data points have been selected, the kriging matrix is solved as usual, with the exception of the $\gamma(h)$ calculation. Again, for directional semivariograms, it is not possible to transform points into isotropic space, therefore component of the individual axes must be resolved. Whether $\gamma(h)$ is being calculated to determine the most influential neighbors or

individual components of the kriging matrix, the same technique is used as described in the following section.

3.3.1 Equation and Proof

Calculating $\gamma(h)$ to determine the nearest neighbors for a grid location, or to define individual components of the kriging matrix, requires the equation for an ellipsoid:

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} = 1 \quad (3.3)$$

Using this equation, it is possible to separate the components of each semivariogram model for any vector (Figure 3.2). One point is translated to the axis origin, and the second point is positioned at $|x|$, $|y|$, and $|z|$, along the separation vector (h). Here a , b , and c , represent the maximum practical ranges of the semivariograms model along the X-, Y-, and Z-axes respectively. In this section only, when the actual ranges are used the a_{actual} , b_{actual} , and c_{actual} subscripts will be used. The practical range refers to the distance where the semivariogram model meets the variance. For the Exponential and Gaussian models, this is defined as 95% of the variance. The practical ranges for different models are defined (Journel and Huijbregts, 1978):

Model Type	Practical Range
Spherical	range
Exponential	3 x range
Gaussian	sqrt(3) x range
Linear	range

If the unadjusted range and not the practical range is used, the axis defined with the model using the longest practical range will be under-weighted. The equations for determining each component $\gamma(h)_{X,Y,Z}$ and the resultant $\gamma(h)$ are derived below. The components of each axis for each structure of the nested semivariogram model can be related through the aspect factors:

$$f = a/b \quad (3.4a)$$

$$g = a/c \quad (3.4b)$$

$$p = b/c \quad (3.4c)$$

Rearranging Equation 3.3, the ellipsoid factors a^2 , b^2 , and c^2 for the search vector are solved:

$$\frac{x^2}{a^2} + \frac{f^2 y^2}{a^2} + \frac{g^2 z^2}{a^2} = 1$$

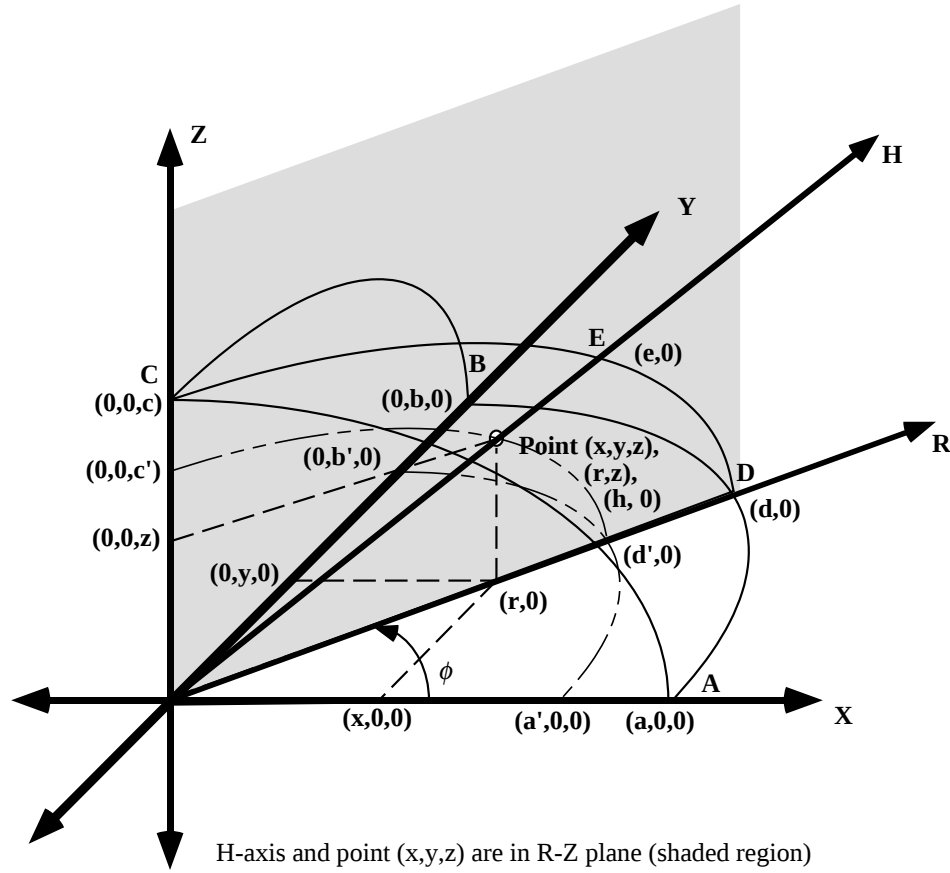


FIGURE 3-2. Directional semivariogram analysis components.

$$a^2 = x^2 + f^2 y^2 + g^2 z^2 \quad (3.5a)$$

$$\frac{x^2}{f^2 b^2} + \frac{y^2}{b^2} + \frac{p^2 z^2}{b^2} = 1$$

$$b^2 = \frac{x^2}{f^2} + y^2 + p^2 z^2 \quad (3.5b)$$

$$\frac{x^2}{g^2 c^2} + \frac{y^2}{p^2 c^2} + \frac{z^2}{c^2} = 1$$

$$c^2 = \frac{x^2}{g^2} + \frac{y^2}{p^2} + z^2 \quad (3.5c)$$

Where a' , b' , and c' represent the X, Y, and Z-axis intercepts for an ellipsoid passing through an arbitrary point, (x, y, z) along the same vector, where the aspect ratios defined by a , b , and c remain true, the following relationships are also true:

$$\frac{a'}{b'} = \frac{a}{b} = f \quad (3.5d)$$

$$\frac{b'}{c'} = \frac{b}{c} = p \quad (3.5e)$$

$$\frac{a'}{c'} = \frac{a}{c} = g \quad (3.5f)$$

An additional axis, R , is also required. R is defined by the intersection of the X-Y plane, and the vertical plane passing through the point (x, y, z) . To determine the semivariogram components, the point r , which lies on the R -axis, vertically below the point (x, y, z) is defined:

$$r^2 = x^2 + y^2 \quad (3.5g)$$

Two additional points of interest are where the semivariogram model ellipsoid and the ellipsoid passing through (x, y, z) cross the R -axis; these are d and d' respectively. Defining these two ellipsoids, with the aspect ratios described above, the components of each semivariogram model can be derived. One new aspect ratio is needed between the R - and Z -axes.

$$q = d/c \quad (3.6)$$

The distances a , b , c , and d represent the practical model range and a' , b' , c' , and d' represent the practical component range along each axis for the point (x, y, z) . The parameter d represents the semivariogram model along the R -axis and is a combination of models a and b . Once a , b , x , and y are known, then d can be determined. For a circle, the angle f is described as:

$$\phi = \tan^{-1}\left(\frac{y}{x}\right) \quad (3.7)$$

Usually, the model semivariogram ellipse (X-Y plane) will not be a circle, therefore the anisotropy must be removed to determine the component angle f . This is the product of y and the aspect ratio of the ellipse (f in the X-Y plane, major/minor dimension):

$$\phi = \tan^{-1}\left(\frac{yf}{x}\right) \quad (3.8)$$

The components of a and b can then be described by dividing $(90^\circ - \phi)$ by 90° , and multiplying by b and a , plus a :

$$a_{\phi \text{ comp}} = \frac{90^\circ - \phi}{90^\circ} a \quad (3.9)$$

$$b_{\phi \text{ comp}} = \frac{\phi}{90^\circ} b \quad (3.10)$$

The components can then be summed to calculate d' :

$$\begin{aligned} d &= a_{\phi \text{ comp}} + b_{\phi \text{ comp}} \\ d &= \frac{90^\circ - \phi}{90^\circ} a + \frac{\phi}{90^\circ} b \\ d &= a + \frac{\phi}{90^\circ} b - \frac{\phi}{90^\circ} b \\ d &= \frac{\phi}{90^\circ} (b - a) + a \end{aligned} \quad (3.11)$$

By expanding f and solving, using radians, the equation may be rewritten:

$$d = \frac{\tan^{-1}\left(\frac{yf}{x}\right)}{\frac{\pi}{2}} (b - a) + a \quad (3.12)$$

d' can be determined by proportion:

$$d' = d \frac{a}{a'} = d \frac{b}{b'} \quad (3.13)$$

Given distances a' , b' , c' , and d' , it is possible to solve directly for $\gamma(a')$, $\gamma(b')$, and $\gamma(c')$. To solve for $g(d'_{\text{actual}})$, the argument is used for d is repeated, The components of $\gamma(a'_{\text{actual}})$ and $\gamma(b'_{\text{actual}})$ can

then be described by dividing $(90^\circ - f)$ by 90° , and multiplying by $\gamma(b'_{\text{actual}})$ and $\gamma(a'_{\text{actual}})$, plus $\gamma(a'_{\text{actual}})$:

$$\gamma(a'_{\text{actual}})_{\phi \text{ comp}} = \frac{90^\circ - \phi}{90^\circ} \gamma(a'_{\text{actual}}) \quad (3.14)$$

$$\gamma(b'_{\text{actual}})_{\phi \text{ comp}} = \frac{\phi}{90^\circ} \gamma(b'_{\text{actual}}) \quad (3.15)$$

The components can then be summed to calculate $\gamma(d'_{\text{actual}})$:

$$\begin{aligned} \gamma(d'_{\text{actual}}) &= \gamma(a'_{\text{actual}})_{\phi \text{ comp}} + \gamma(b'_{\text{actual}})_{\phi \text{ comp}} \\ \gamma(d'_{\text{actual}}) &= \frac{90^\circ - \phi}{90^\circ} \gamma(a'_{\text{actual}}) + \frac{\phi}{90^\circ} \gamma(b'_{\text{actual}}) \\ \gamma(d'_{\text{actual}}) &= \gamma(a'_{\text{actual}}) + \frac{\phi}{90^\circ} \gamma(b'_{\text{actual}}) - \frac{\phi}{90^\circ} \gamma(a'_{\text{actual}}) \\ \gamma(d'_{\text{actual}}) &= \frac{\phi}{90^\circ} (\gamma(b'_{\text{actual}}) - \gamma(a'_{\text{actual}})) + \gamma(a'_{\text{actual}}) \end{aligned} \quad (3.16)$$

By expanding f and solving, the equation may be rewritten:

$$\gamma(d'_{\text{actual}}) = \frac{\tan^{-1}\left(\frac{yf}{x}\right)}{\frac{\pi}{2}} (\gamma(b'_{\text{actual}}) - \gamma(a'_{\text{actual}})) + \gamma(a'_{\text{actual}}) \quad (3.17)$$

To solve for $\gamma(e'_{\text{actual}})$, where e' is the distance from the origin to the point (x, y, z) , steps similar to those used to generate d and $\gamma(d'_{\text{actual}})$ are required. Allowing $\gamma(d'_{\text{actual}})$ to be equivalent to $\gamma(a'_{\text{actual}})$, and $\gamma(c'_{\text{actual}})$ equivalent to $\gamma(b'_{\text{actual}})$, this yields:

$$\gamma(e'_{\text{actual}}) = \frac{\tan^{-1}\left(\frac{zq}{r}\right)}{\frac{\pi}{2}} (\gamma(c'_{\text{actual}}) - \gamma(d'_{\text{actual}})) + \gamma(d'_{\text{actual}}) \quad (3.18)$$

These calculation must be evaluated for each nest of the model structure except the nugget ($\gamma(h)_0$). The nugget, having zero distance, by definition is the same for all axes. This also implies that the number of structures in every direction must be equal. This restriction can be negated by giving undesired nests a zero variance component and the same range as the previous structure. The final $\gamma(e'_{\text{actual}})$ estimate is the summation on the nugget and the nested structure components:

$$\gamma(e'_{\text{actual}}) = \sum_{i=0}^s \gamma(e'_{\text{actual}})_i ; \text{where } s = \text{number of structures} \quad (3.19)$$

3.3.2: Positive Definite Matrix Issues

The models selected for the semivariogram, must yield a positive definite kriging matrix (Journel and Huijbregts, 1978). If the matrix is not positive definite, there may be no solution or there may be several different solutions (Isaaks and Srivastava, 1989), and the kriging variance may be negative (Journel and Huijbregts, 1978). The various model types used here (spherical, exponential, Gaussian, and logarithmic) have proven to be positive definite both individually and in combination as nested structures (Journel and Huijbregts, 1978). Although the equations are merged in a different manner for directional semivariograms than for traditional kriging, it is assumed that the matrix remains positive definite. In practice, several indicators used to determine whether the matrix is not positive definite, are 1) matrices that are singular, 2) have large positive or negative kriging weights (much larger or smaller than ± 1.0), and 3) the occurrence of negative estimation variances. Proving that the equations are positive definite is a difficult task (Christakos, 1984; Isaaks and Srivastava, 1989), but in summary, for a symmetric ($n \times n$) matrix to be positive definite, it must satisfy any one of the following conditions (Burden and Faires, 1985; Isaaks and Srivastava, 1989; Strang, 1988):

- i) $x^t A x > 0$ for all non-zero vectors x .
- ii) All the eigenvalues (λ_i) of A are greater than 0.
- iii) All the upper left submatrices A_k have positive determinants.
- iv) All the pivots (d_i , without row exchange) are greater than 0.

for every n -dimensional column vector $x \neq 0$, where A is the kriging matrix, A_k is a submatrix of A , x is any vector (appropriately dimensioned), and x^t is the transpose of x .

3.3.2.1: Problems With the Positive Definite Assumption

Some problems with large positive and negative weights and negative kriging variances were encountered when Gaussian models were used with the directional kriging method. Many of the problematic matrices were confirmed to be positive definite based on tests i) and ii). The problems were attributed to the unstable nature of the Gaussian model with small nuggets (Ababou, Bagtzoglou, et al., 1994; Posa, 1989). Ababou, et al (1994), state that this is a common problem, particularly with Gaussian models that have small nugget values. The kriging matrix becomes more unstable and approaches singularity at small h values. This tendency can be estimated using the kriging matrix (A) conditioning number $\kappa(A)$:

$$\kappa(A) = |\text{MAX eigenvalue}| / |\text{MIN eigenvalue}|$$

The Gaussian model, is one of the most problematic (2 to 14 times worse than hole-exponential models (which are one of the best) (Ababou, et al., 1994)) and most unstable, and tends to have minimum eigenvalues near 0.0. Also, because the model is relatively flat at small h values (unlike all other models), the problem prevails at larger h values than for other models (Posa, 1989). Because of this instability, it is sometimes better to select a model which does not physically fit the data as well as another model, but is more robust (Posa, 1989).

3.4: Modification of Algorithms

In this project, the GSLIB ktb3dm (Deutsch and Journel, 1992), and SISIM3D (Gómez-Hernández and Srivastava, 1990; McKenna, 1994) algorithms were modified to build the kriging matrix using both anisotropic semivariogram models and directional semivariogram models.

3.4.1: Algorithm Constraints

Although directional semivariogram models relax many of the constraints in defining spatial variation, there are several limitations. Some of these limitations arise from the theory, and some from the implementation. The limitations are:

- The number of semivariogram models for each axis must be equal. This is a minor limitation, because extra models can be added as needed with a zero variance component, and a range equal to the final desired range. If this is not done, there may be ambiguity in how the semivariogram models are evaluated.
- The sill for all axes must be equal. Again this is a minor limitation. If the variance in one direction is smaller than another in the grid area, the remaining variance component may be added to the final nest, while the range for the final nest is set to a range much greater than the size of the simulated area (or the search distance for that matter). This constraint is required to ensure positive definite matrix solutions.
- Gaussian models may be used with small or zero nugget values, but the modeler must be aware that the results can be unstable. The algorithm presented here tests for large ($\pm$) weights and warns the modeler. The algorithm can remove data points (the point associated with the largest absolute kriging weight) from the kriging matrix until the results stabilize, or until there are too few points to estimate the grid location, however estimates resulting from such elimination should be considered highly suspect (See Rocky Mountain Arsenal example described below), and one may wish to compromise and use a larger nugget (often the problem with Gaussian models), or a different model type all together.

3.4.2: Computational Cost

Although use of directional semivariograms is computationally intensive, the increased computation, is significant, but not excessive. In several test cases, the computation time increased between 80% and 200%. The increased computation occurs mainly in the search algorithm that identifies the most influential neighbors for each grid location and in the additional overhead for calculating each covariance value of the kriging matrix.

The search algorithm in the traditional technique includes two main steps: 1) transformation of the sample data to isotropic space, and 2) calculation of the distance between each point and each grid location being estimated. In addition to these steps, the directional semivariogram technique requires that $\gamma(h)$ be calculated for each sample point, relative to the position of the grid position being estimated. To determine the neighboring points with the smallest spatial variances, traditional techniques calculate only the relative distance between sample points in isotropic space, and the point being estimated. This is adequate, because the spatial variance is only a function of distance. When using directional semivariogram models, transformed isotropic distances are not sufficient to rank sample points; direction is also important, thus $\gamma(h)$ for the separation between each sample and the grid location must be ranked. Calculating $\gamma(h)$ in the search phase, adds most of the increased computational effort.

The calculation of $\gamma(h)$ for the kriging matrix also requires additional effort. Although this step is generally less expensive in computation time than the search step because it is applied only to the selected nearest neighbor points and not to all points within the search neighborhood.

To solve the kriging problem, the kriging matrix must also be solved using either Gauss elimination or a more efficient LU decomposition (Alabert, 1987). These calculations are unaffected by the method used to define the semivariogram models, but because this is a computationally intensive task, the increased cost due to directional semivariogram modeling is less severe.

3.5: Examples

Several example models and data sets are used to demonstrate the applicability and validity of using directional semivariogram models.

3.5.1: Comparison With the Classic Method

In addition to the mathematical proof above, it is also important to demonstrate that the algorithm and the software are correct. Two approaches are pursued to evaluate the algorithm. First, conditions modeled using anisotropy factors with traditional methods are duplicated using directional semivariogram models that mimic the anisotropy factors. Then, model results using directional semivariogram models, that cannot be described with anisotropy factors, are compared with manual calculations.

3.5.1.1: Anisotropic Case: Directional Components Equivalent to Anisotropy Factors

To demonstrate that the directional semivariogram model technique produces the same results as the conventional technique using anisotropy factors, a small synthetic data set with eleven data points was created (Figure 3.3a). Given equivalent model input, the results are identical. The semivariogram models for each case are: The principal axis is oriented to the Northwest. The map

Method	Axis	Model Type	Range	Sill	Nugget	Y-Anisotropy
Anisotropic	All	Spherical	100	0.14	0.02	0.4
		Spherical	250	0.11		0.75
Directional	X	Spherical	100	0.14	0.02	NA
		Spherical	250	0.11		NA
	Y	Spherical	40	0.14	0.02	NA
		Spherical	187.5	0.11		NA

in Figure 3.3b is the traditional simple kriged map using a single semivariogram model with anisotropy factors. Figure 3.3c was produced using directional semivariograms. When the two maps are subtracted from one another, the difference is zero at every grid location, indicating that the directional semivariogram method is able to correctly reproduce the simple case where anisotropic conditions exist and the perpendicular semivariogram models are related by anisotropy factors.

3.5.1.2: Anisotropic Case - Manual Solution

To demonstrate that the method produces the answers we intuitively expect, several $g(h)$ values are calculated manually for several points at various orientations with one model set (three orthogonal directional semivariogram models). A second calculation will be made for the multi-nested model defined in Figure 3.1. The first model set is defined as:

Axis	Model Type	Range	Sill	Nugget
X	Spherical	125	5	1
Y	Gaussian	75	5	1
Z	Spherical	30	5	1

for the points:

Sample	X	Y	Z
I	87	0	0
II	43	26	0
III	43	26	-11

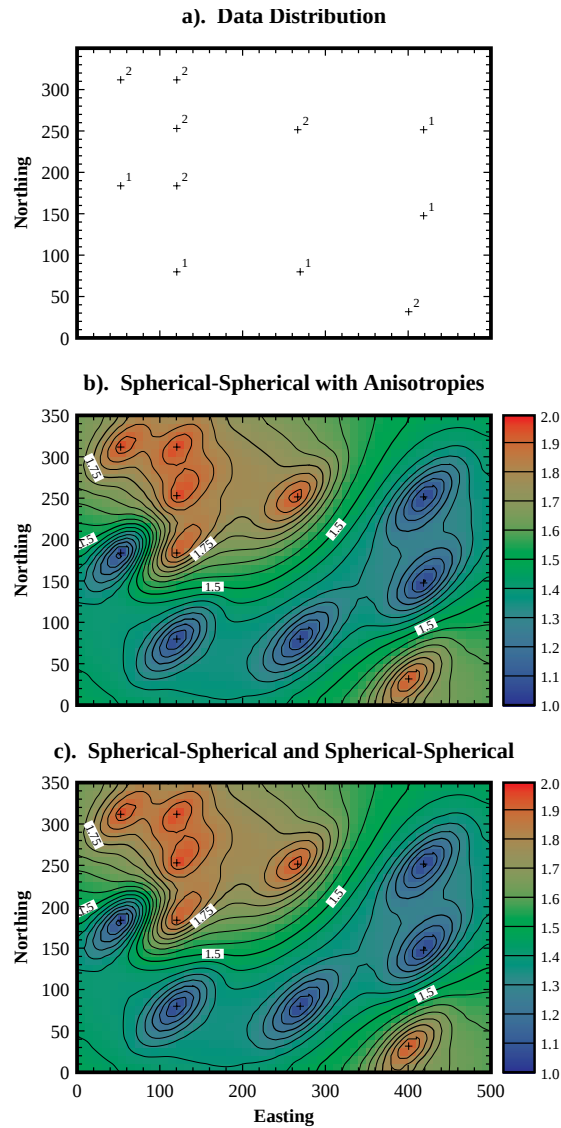


FIGURE 3-3. Example results confirming directional semivariograms can exactly mimic anisotropy factors: a) sample data set, b) SK map using anisotropic factors, c) SK map using directional semivariograms.

the solutions are calculated:

- I). For the point (87, 0, 0), the solution is simple. The point lies along a principal axis, and in this case has a zero Y and Z component. $\gamma(h)$ can therefore be calculated directly, using the standard spherical equation for the X direction:

$$\gamma(h) = \sum_{i=0}^s C_i \text{ ; where } s = \text{number of structures} \quad (3.19)$$

$$\gamma(h)_1 = C_1 \left[1.5 \frac{h}{r} - 0.5 \frac{h^3}{r^3} \right] \quad (3.20)$$

$$\gamma(h)_1 = 5 \left[1.5 \frac{87}{125} - 0.5 \frac{87^3}{125^3} \right] = 4.377$$

$$\gamma(h)_0 = \text{nugget} = 1.0$$

$$\gamma(h) = \gamma(h)_1 + \gamma(h)_0 = 4.377 + 1.0 = 5.377$$

where h is the separation distance, and r is the model range (for this equation only).

- II). For the point (43, 26, 0), the first step is to define the nugget; $\gamma(h)_0 = 1.0$. Next the X and Y directional components must be calculated (the Z axis has a zero component). The X and Y intercepts of the ellipse that passes through (43, 26) and has an X/Y aspect ratio of 125/75 (a/b) (remember the actual Gaussain range must be multiplied by the to determine the practical ellipsoid range). The intercepts, a' and b' are determined using the standard equation for an ellipse (Equation 3.3):

$$a^2 = (43)^2 + \left(\frac{125}{75\sqrt{3}} \right)^2 (26)^2 \quad (\text{using Equation 3.5a})$$

$$a' = 49.74$$

$$b^2 = \frac{(43)^2}{\left(\frac{125}{75\sqrt{3}} \right)^2} + (26)^2 \quad (\text{using Equation 3.5b})$$

$$b' = 50.97$$

Given the X and Y intercepts, $\gamma(h)$ for each ellipsoid axis is calculated:

$$\gamma_x(h)_1 = 5 \left[1.5 \frac{49.74}{125} - 0.5 \frac{49.74^3}{125^3} \right] = 2.827$$

$$\gamma_y(h)_1 = C_1 \left(1 - e^{-(h)^2 / (r)^2} \right) \quad (3.21)$$

$$\gamma_y(h)_1 = 5 \left(1 - e^{-(50.97)^2 / (75)^2} \right) = 1.849$$

Note that the actual and not practical range is used to calculate $\gamma_y(h)_1$. Once the maximum contributions for each axis have been determined, the component contribution of each must be determined. This is done by determining the effective angle of the vector (43, 26) in the X-Y plane. The effective angle is:

$$\phi = \tan^{-1} \left(\frac{26 \frac{125}{75\sqrt{3}}}{43} \right) = 30.19^\circ$$

Given this angle, the components of $\gamma_x(h)$ and $\gamma_y(h)$ can be determined. Intuitively the X-axis component can be defined as:

$$2.827 \frac{(90^\circ - 30.19^\circ)}{90^\circ} = 1.879$$

and the Y-axis component is:

$$1.849 \left(\frac{30.19^\circ}{90^\circ} \right) = 0.620$$

Adding the two components together yields a directional $\gamma(h)$ of 2.499. This yields the same result as if Equation 3.18 were used:

$$\gamma_d(h)_1 = \frac{\tan^{-1}\left(\frac{26 \frac{125}{75\sqrt{3}}}{43}\right)}{\frac{\pi}{2}}(1.849 - 2.827) + 2.827 = 2.449$$

When the $\gamma(h)_i$ components are summed, the total estimate for $\gamma(h)$ is $2.449 + 1.0$ which equals 3.449.

- III). The same approach may be used for the last point (43, 26, -11), $\gamma(h)_0 = 1.0$, and the X, Y, and Z directional components must be calculated. The first step is to calculate the X, Y, and Z intercepts for the ellipsoid that passes through (43, 26, -11) and has an X/Y aspect ratio of 125/75, a X/Z aspect ratio of 125/30, and a Y/Z aspect ratio of 75/30. The intercepts, a' , b' , and c' are determined using the standard equation for an ellipsoid (Equation 3.3):

$$a'^2 = (43)^2 + \left(\frac{125}{75\sqrt{3}}\right)^2 (26)^2 + \left(\frac{125}{30}\right)^2 (-11)^2$$

(using Equation 3.5a)

$$a' = 67.64$$

$$b'^2 = \frac{(43)^2}{\left(\frac{125}{75\sqrt{3}}\right)^2} + (26)^2 + \left(\frac{75\sqrt{3}}{30}\right)^2 (-11)^2$$

(using Equation 3.5b)

$$b' = 70.29$$

$$c'^2 = \frac{(43)^2}{\left(\frac{125}{30}\right)^2} + \frac{(26)^2}{\left(\frac{75\sqrt{3}}{30}\right)^2} + (-11)^2$$

(using Equation 3.5c)

$$c' = 16.23$$

Given the X, Y, and Z intercepts, $\gamma(h)_1$ for each ellipsoid axis is calculated:

$$\gamma_x(h)_1 = 5 \left[1.5 \frac{67.64}{125} - 0.5 \frac{67.64^3}{125^3} \right] = 3.662$$

$$\gamma_y(h)_1 = 5 \left(1 - e^{-(70.29)^2 / (75)^2} \right) = 2.923$$

$$\gamma_z(h)_1 = 5 \left[1.5 \frac{16.23}{30} - 0.5 \frac{16.23^3}{30^3} \right] = 3.662$$

Using the same methods as described for point II, $\gamma_{d'}(h)_1$ can be determined:

$$\gamma_{d'}(h)_1 = \frac{\tan^{-1} \left(\frac{26 \frac{125}{75\sqrt{3}}}{43} \right)}{\frac{\pi}{2}} (2.923 - 3.662) + 3.662 = 3.414$$

Now that the X-Y axis contributions have been merged, the Z-axis component is incorporated. This requires that d' and r be calculated. d' is calculated by merging equations 3.12 and 3.13:

$$d' = \frac{\tan^{-1} \left(\frac{yf}{x} \right)}{\frac{\pi}{2}} (b' - a') + a' \quad (3.22)$$

$$d' = \frac{\tan^{-1} \left(26 \frac{125}{75\sqrt{3}} \right)}{\frac{\pi}{2}} (70.29 - 67.64) + 67.64 = 68.53$$

$$r = \sqrt{(43)^2 + (26)^2} = 50.25$$

Similar steps are used in the vertical R-Z plane through (43, 26, -11) as were undertaken in the X-Y plane,. For purposes of calculating the vector length h , only absolute values for each coordinate are used, and c' and d' are substituted for c and d . The angle f from R to Z is:

$$\phi = \tan^{-1} \left(\frac{11 \frac{68.53}{16.23}}{50.25} \right) = 42.75^\circ$$

Given this angle, the components of $\gamma_d(h)_1$ and $\gamma_z(h)_1$ can be determined.

$$3.414 \frac{(90^\circ - 42.75^\circ)}{90^\circ} = 1.792$$

and the Z-axis component is:

$$3.662 \frac{42.75^\circ}{90^\circ} = 1.739$$

Adding the two components together yields a directional $\gamma(h)_1$ of 3.531. This yields the same result as if Equation 3.23 were used:

$$\gamma_e(h)_1 = \frac{\tan^{-1} \left(\frac{11 \frac{68.53}{16.23}}{50.25} \right)}{\frac{\pi}{2}} (3.662 - 3.414) + 3.414 = 3.531$$

When the $\gamma(h)_i$ components are summed, the total estimate for $\gamma(h)$ is $3.531 + 1.0$ which equals 4.531.

For the final example, the one and two-nested structure models shown in Figure 3.1 are used. This example demonstrates that the modeler is not required to specify the same number of model structures in all directions. Although the algorithm requires the number of structures to be equal, the algorithm can internally add extra structures as needed without affecting the model description. The calculations will be made for two points separated by 200m at a 45° angle ($x = 141.1$, $y = 141.1$). The models are defined:

Direction	Range	C	Model	Nest
North-South	450	0.05	Spherical	2
	200	0.20	Spherical	1
		0.022	Nugget	0

Direction	Range	C	Model	Nest
East-West	150	0.00	Gaussian	2
	150	0.25	Gaussian	1
		0.22	Nugget	0

Note that the number of structures are the same in both directions, but the East-West models second nest has a zero sill (C) component. As described in the earlier examples, the nugget is a constant with direction, therefore $\gamma(h)_0 = 0.022$. The remaining $\gamma(h)_i$ values are calculated as follows (geometric interpretations are shown in Figure 3.4):

$$a_1^2 = (141.4)^2 + \left(\frac{150\sqrt{3}}{200} \right)^2 (141.4)^2$$

$$a'_1 = 231.9$$

$$b_1^2 = \frac{(141.4)^2}{\left(\frac{150\sqrt{3}}{200} \right)^2} + (141.4)^2$$

$$b'_1 = 178.4$$

$$\gamma_x(h)_1 = 0.25 \left(1 - e^{-\left(\frac{150\sqrt{3}}{200} \right)^2 / (231.9)^2} \right) = 0.1789$$

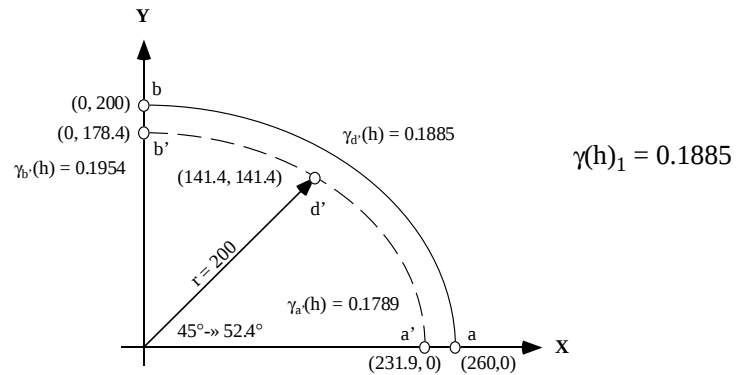
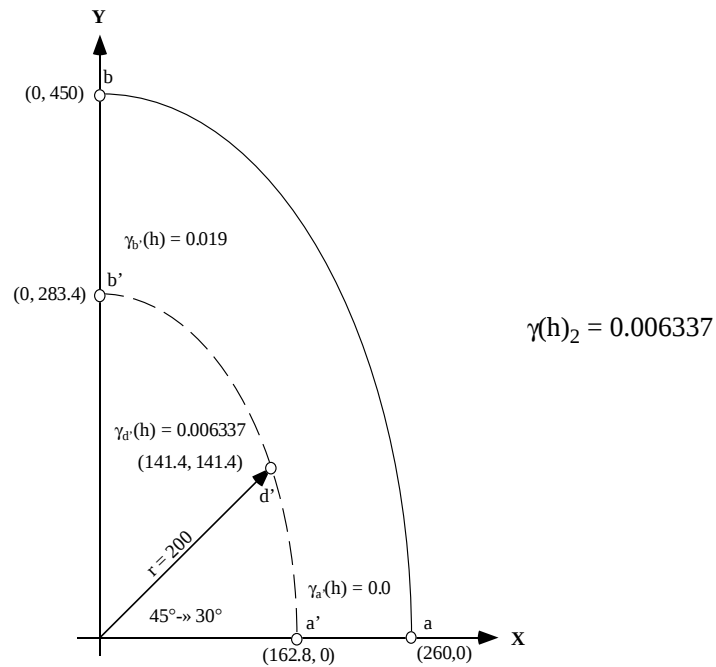
$$\gamma_y(h)_1 = 0.20 \left[1.5 \frac{200}{178.4} - 0.5 \frac{200^3}{178.4^3} \right] = 0.1954$$

$$\gamma_d(h)_1 = \frac{\tan^{-1} \left(\frac{141.4 \frac{150\sqrt{3}}{200}}{141.4} \right)}{\frac{\pi}{2}} (0.1954 - 0.1789) + 0.1789 = 0.1885$$

For the second nest, there is no East-West component. For the algorithm to work correctly, an additional East-West structure must be defined (the number of structures for all axes must be equal). To satisfy the algorithm and the specified Gaussian model, a zero sill component is used, and the range is set equal to the previous nest. This manipulation satisfies the algorithm, and leaves the model definition unchanged:

STEP 1: Determine Nugget

$$\gamma(h)_0 = 0.022$$

STEP 2: Determine C_1 Geometry (Nest 1)**STEP 3: Determine C_2 Geometry (Nest 2)****STEP 4: Sum Components**

$$\gamma(h) = 0.2168$$

FIGURE 3-4. Geometric steps for calculating directional semivariogram model defined in Figure 3.1. The major axis is aligned North-South, and the minor axis is aligned East-West. Note, the 45° angle is transformed ($\rightarrow$) based on the anisotropy of the ellipsoid.

$$a_2^2 = (141.4)^2 + \left(\frac{150\sqrt{3}}{450} \right)^2 (141.4)^2$$

$$a_1' = 162.8$$

$$b_2^2 = \frac{(141.4)^2}{\left(\frac{150\sqrt{3}}{450} \right)^2} + (141.4)^2$$

$$b_1' = 283.4$$

$$\gamma_x(h)_2 = 0.00 \left(1 - e^{-(150\sqrt{3})^2 / (162.8)^2} \right) = 0.0000$$

$$\gamma_y(h)_2 = 0.05 \left[1.5 \frac{450}{283.4} - 0.5 \frac{450^3}{283.4^3} \right] = 0.01900$$

$$\gamma_{d'}(h)_2 = \frac{\tan^{-1} \left(\frac{141.4 \frac{150\sqrt{3}}{450}}{141.4} \right)}{\frac{\pi}{2}} (0.1900 - 0.0000) + 0.0000 = 0.006337$$

The final step is to sum the $\gamma_{d'}(h)_i$ components:

$$\begin{aligned} \gamma(h) &= \sum_{i=0}^2 \gamma_{d'}(h)_i = \gamma_{d'}(h)_0 + \gamma_{d'}(h)_1 + \gamma_{d'}(h)_2 \\ &= 0.022 + 0.1885 + 0.006337 \\ &= 0.2168 \end{aligned}$$

3.5.2: Practical Applications

A synthetic and a field data set are used to demonstrate the effectiveness and usefulness of the technique. For the synthetic case, the same data set that was used in section 3.4.1.1 is utilized, though different assumptions about the X and Y semivariogram models are made. The field data set is residual bedrock elevation data from the Rocky Mountain Arsenal, Commerce City, Colorado.

3.5.2.1: Synthetic Directional Semivariogram Demonstration Set

To demonstrate that directional semivariogram models can have a significant impact on model results, the data set in Figure 3.3a is used, but in this case anisotropy factors are not used, rather directional semivariogram models are defined. Since this is a synthetic data set, none of the following models can be argued to be the best representation of site conditions, any more than the other models, but the exercise demonstrates that directional semivariograms offer great flexibility in adjusting the estimations to match perceived or measured site conditions. Three different site scenarios were calculated based on the following directional semivariogram models (Figure 3.5):

Scenarios	Axis	Nest	Model Type	Range	Sill	Nugget
I	X	1	Spherical	100	0.14	0.02
		2	Spherical	250	0.11	
	Y	1	Gaussian	23.1	0.14	0.02
		2	Exponential	62.5	0.11	
II	X	1	Spherical	100	0.14	0.02
		2	Spherical	250	0.11	
	Y	1	Gaussian	23.1	0.07	0.02
		2	Exponential	62.5	0.18	
III	X	1	Spherical	100	0.14	0.02
		2	Spherical	250	0.11	
	Y	1	Gaussian	62.5	0.25	0.02

The ranges of the exponential and Gaussian models are significantly different from the spherical model ranges used for the Y-axis (40m and 187.5m for two nests) in Section 3.4.1.1. The range of a exponential and Gaussian model must be multiplied by the following factors to yield the equivalent spherical range (Deutsch and Journel, 1992):

Model Type	Practical Range (a)
Exponential	3a
Gaussian	$a(\sqrt{3})$

Despite these rules of thumb, the range for the scenario III Gaussian model was set to 1/3 of the two-nested Spherical model's range. This configuration more closely resembles the original and alternate Y-axis models (Figure 3.5, the climbing limbs of the models are more similar, even if the full Gaussian range is somewhat reduced). These models, are oriented with their major axes to the Northeast. In Figure 3.6a, the structures for the minor axis were substituted with Gaussian and Exponential models. In Figure 3.6b, the sill terms for the first structure in the minor axis (Y) was lowered to 0.7, and the sill for the second structure was raised to 0.18. Finally, in Figure 3.6c, the minor axis was substituted with a single Gaussian model (sill = 0.25, a second structure with a 0.0 sill component is assumed by the algorithm). The estimations (Figure 3.3a,b, and Figure 3.6a-c) show the same general NE-SW trend, but vary in detail. The differences are easiest to see near the

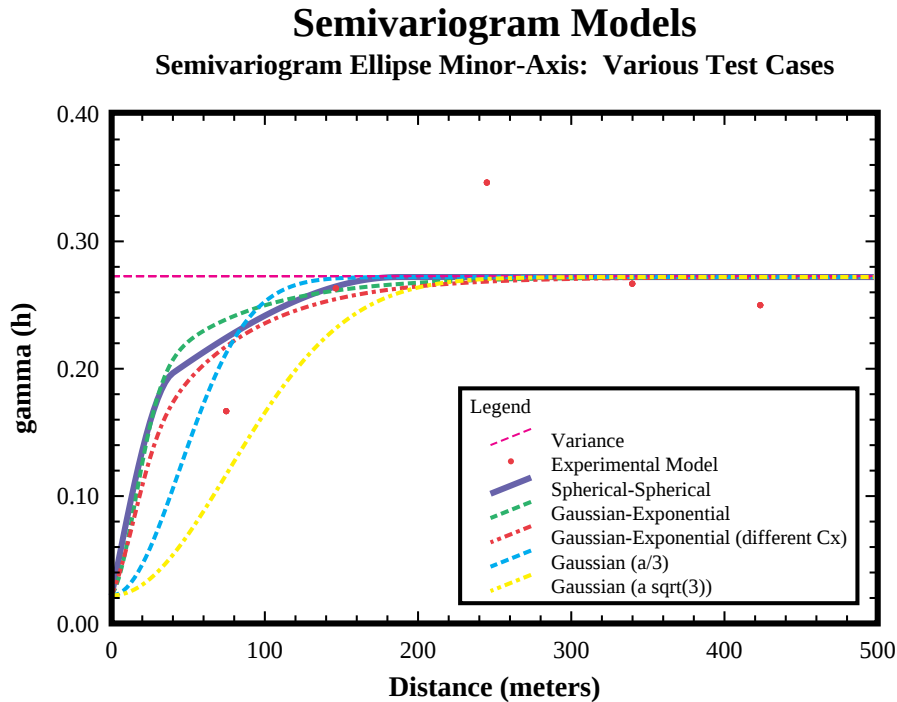


FIGURE 3-5. Semivariogram models used for synthetic directional semivariogram data set. Despite the general rule of thumb that the practical Gaussian range to a spherical range (a) is the $\text{SQRT}(3)$ multiplied by the range (a), the Gaussian (range (a) $\times$ $\text{SQRT}(3)$) model, because it mimicked the general nature of the other models more closely.

peaks at (100,300: in red), the valley depressions near (425, 210: in blue), and the slope transition at (70, 130). Although Figure 3.3b, and Figures 3.6a through 3.6c appear similar to each other, the mean absolute differences are as much as 7%, and differences between individual cells are up to 37% (Figures 3.7a, 3.7c, and 3.8), when compared to the kriged mapped using anisotropy factors (Figure 3.3b). These scenarios demonstrate how the use of directional semivariogram model descriptions impacts the resulting maps, relative to a scenario which utilizes a compromise semivariogram model with anisotropy factors.

3.5.2.2: Rocky Mountain Arsenal Demonstration Data Set

To demonstrate the effectiveness, and some of the difficulties, of directional kriging, a data set of bedrock surface elevations (actually residuals from a second-order trend-surface) from the Rocky Mountain Arsenal (RMA), Commerce City, Colorado is used. With this data set, use of correct directional semivariogram models reduced the average estimation variance over the map area, even though an artificially large nugget was used. Because of problems with the Gaussian

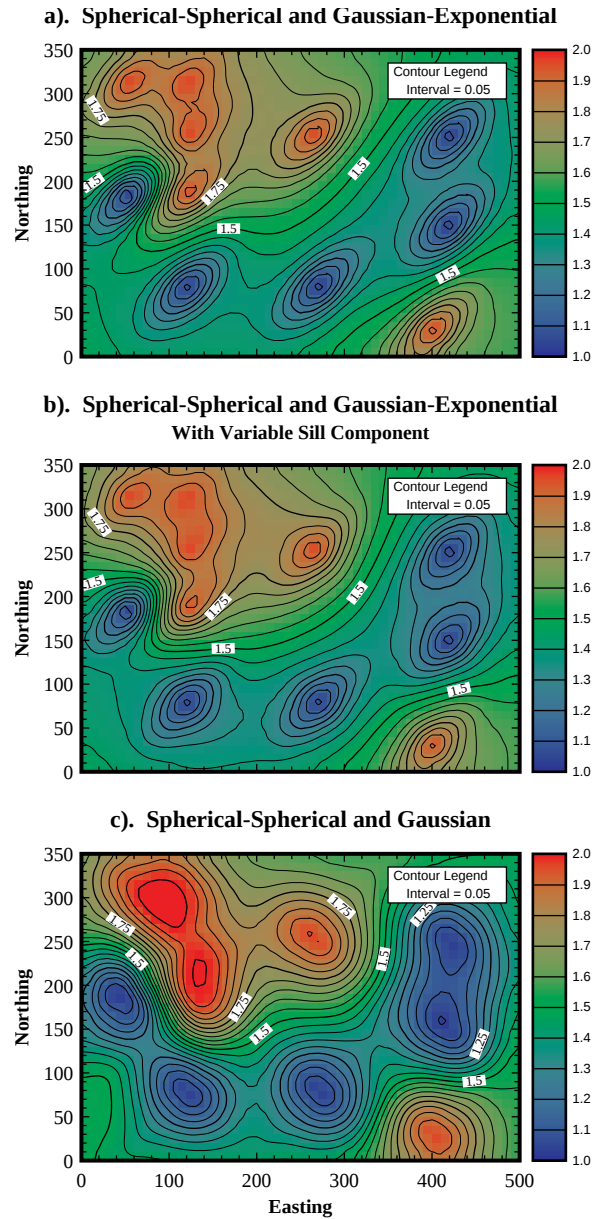


FIGURE 3-6. Results of directional semivariogram models using different assumptions about major and minor semivariogram models.

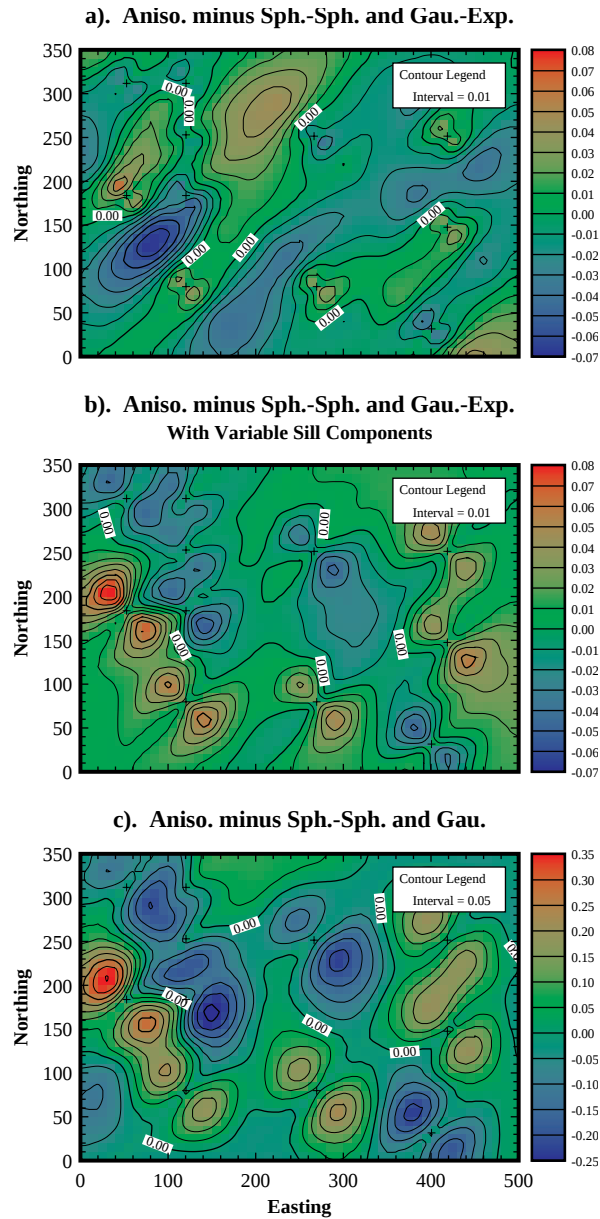


FIGURE 3-7. Differences between original SK models (Figure 3.3a-b), and directional semivariogram models (Figures 3.6a-c).

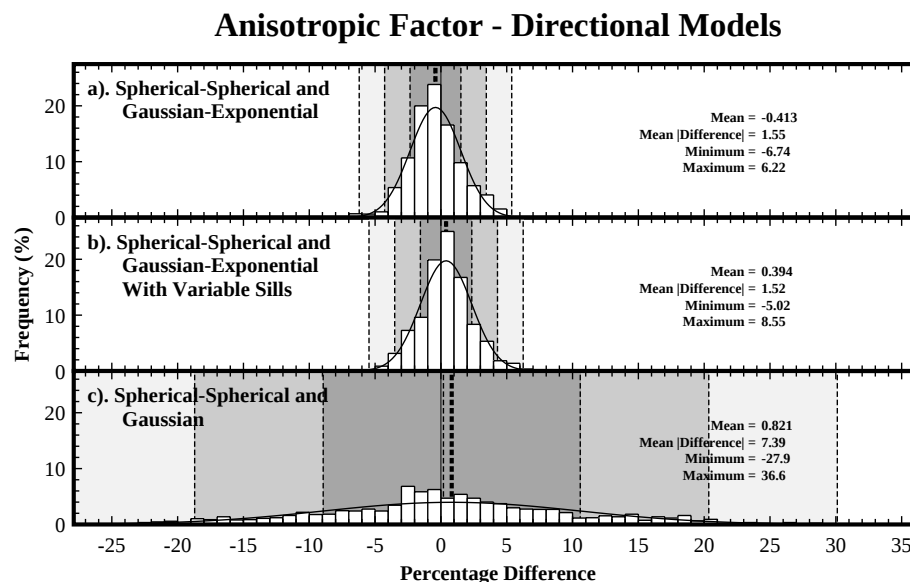


FIGURE 3-8. Distribution of differences between original SK models (Figure 3.3a-b), and directional semivariogram models (Figures 3.6a-c).

semivariogram model, the nugget term was increased by 260% to stabilize the kriging matrix (Gaussian models can cause singular matrix problems with small nuggets (Ababou, et al., 1994; Posa, 1989)).

3.5.2.2.1: Background

Johnson (1995) had trouble evaluating this site due to constraint related to the semivariogram model definition. She recognized directional differences in spatial statistics, but anisotropy factors would not allow her to model them correctly. As a result, Johnson compromised with a two-nested spherical model. It is important that the RMA be modeled accurately, because, summarizing Johnson (1995), there are many serious environmental concerns:

The RMA was established in 1942 for the production of chemical and incendiary munitions. From 1947 to 1982, herbicides and pesticides were also produced (Environmental Science and Engineering, 1987). During this time chemical agents, such as levinstein mustard (H), phosgene, napalm, isopropylmethyl fluorophosphonate (Sarin or GB), and dichlorodiphenyltrichloroethane (DDT) were produced (Harding Lawson Associates, 1992). Problems arose at the site because liquid wastes were disposed of in lined and unlined evaporation basins, and waste was initially held in settling ponds or transported by sewer or drainage ditch to the basin (Kuznear and Trautmann, 1980). By the 1950's the effects of ground water contamination were noted; there was high waterfowl mortality and extreme crop loss (Harding, Lawson, et al., 1992). By 1974,

disopropylmethylphosphonate (DIMP) and dicyclopentadiene (DCPD) contamination was detected off site (Environmental Science and Engineering, 1987).

Johnson (1995) investigated potential transport routes for contaminants from the RMA. To accomplish this, Johnson (1995) identified and simulated (using conditional indicator simulation) paleo-river channels, coarse and fine sediment distribution, and the ground water surface. The paleo-river channels are of interest because they provide potential pathways for ground water and contaminant movement. To identify these paleo-river channels Johnson (1995) simulated the bedrock surface using boring data from 842 wells. This bedrock surface was identified as an ancient erosional surface which dips slightly to the Northwest towards the Platte River (Harding, et al., 1992).

In Johnson's (1995) work, the regional dip was removed from the bedrock data using a second-order trend-surface. Using the residual data, Johnson performed semivariogram analyses and conditional simulation. A problem arose during the semivariogram analysis; the experimental semivariograms in the minor and major search directions couldn't be modeled well using a single model semivariogram with anisotropy factors. As a result, compromises were made in selecting semivariogram models (Figure 3.9a) with the hope that, by honoring the short lag data, errors would be acceptably small.

3.5.2.2.2: Directional Semivariogram Kriging

The directional semivariogram kriging technique was used to separate the directional components in semivariogram models. The full series of simulations presented by Johnson (1995) is not repeated here, but the new estimates of the bedrock surface honor the spatial distribution of the data better than the estimates made by Johnson (1995). This is accomplished by using simple kriging and evaluating the estimation variance. The estimation variance is a function of the data locations, and the differences between ordinary kriging and indicator kriging, do not effect the estimation variance. It is important to note that the estimation variance only provides a comparison of alternative data configurations; it is independent of the data values (Deutsch and Journel, 1992).

Four semivariogram models were evaluated: (I) one is similar to Johnson's (1995) two-nested spherical-spherical model with anisotropy factors, but an improved model with a lower mean square error (MSE) is used (Figure 3.9a); (II) another is an accurate directional spherical-spherical / Gaussian model (Figure 3.9b); (III) a second directional model based on II, but with a much larger nugget to accommodate difficulties with the Gaussian model is used (Figure 3.9c), and finally (IV) another two-nested spherical-spherical model with anisotropy factors, but an appropriate nugget is used (Figure 3.9d). The semivariogram models are summarized in Table 3.1.

Model I fits the major-axis (East-West) well, but its spherical-spherical structure is not able to represent the inflection in the early portion of the minor axis (North-South) experimental semivariogram. This model assumes a zero nugget. When this model is used with Simple Kriging on the site data (Figure 3.10a), using a 50 by 50, two-dimensional grid, the smallest estimation variance results of all the semivariogram models are obtained. The kriged surface and estimation variance are shown in Figures 3.10b-c. It is thought that this model underrates the estimation

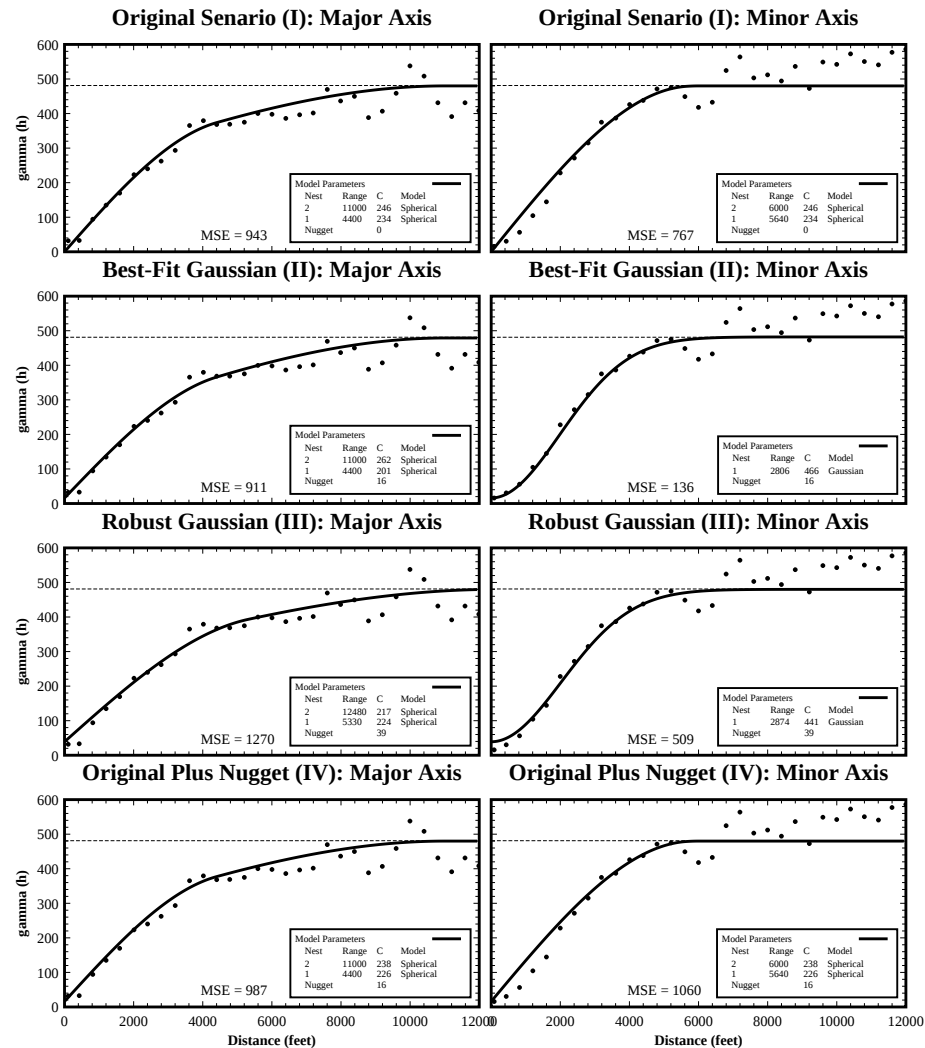


FIGURE 3-9. Experimental and model semivariograms for RMA bedrock residuals (2nd order trend removed): a) anisotropy factor model optimized to minimize MSE based on Johnson (1995), b) optimized minor-axis fit with Gaussian model (note MSE reduced by 82%), c) minor-axis Gaussian model fit with elevated nugget to reduce kriging matrix instability, d) anisotropy factor model optimized to minimize MSE, but also honor nugget defined in b).

variance due to the zero nugget. It is clear that the nugget has a $\gamma(h)$ value of approximately 16 (Figure 3.10b). This incorrect assumption is corrected with model IV.

Model II (Figure 3.9) has the best fit of the four models evaluated, based on MSE measurements of the experimental semivariograms. The fit is particularly good for the minor axis. The MSE for this

Model	Axis	Model Type	Range	Sill	Nugget	Y-Aniso	MSE
I	X/Y	Spherical	4400	234	0	1.833	943/767
		Spherical	11000	246		0.780	
II	X	Spherical	4400	201	16	NA	911
		Spherical	11000	262		NA	
III	Y	Gaussian	2806	466	16	NA	136
	X	Spherical	5330	224	39	NA	1270
		Spherical	12480	217		NA	
	Y	Gaussian	2874	441		NA	509
IV	X/Y	Spherical	4400	226	39	1.833	987/1060
		Spherical	11000	238	16	0.780	

TABLE 3.1. Alternative semivariogram models for RMA residual bedrock surface. Range, sill, and nugget terms are in feet.

axis model is only 12% to 25% of all the other models evaluated. From this model, it was concluded that the nugget has a $\gamma(h)$ value of 16.0. Due to theoretical problems with Gaussian models and the small nugget (less than 10% of the variance) associated with these data, this model has no acceptable solution. Many individual grid cell kriging matrices are singular, or have huge kriging weights (weights greater than ± 1.05 were considered unacceptable; weights greater than ± 200 were found). Regrettably, this behavior is inherent with the Gaussian model, but increasing the nugget increases the stability of every matrix solution (Ababou, et al., 1994; Posa, 1989).

Model III (Figure 3.9c) was developed in an attempt to stabilize the solution, without completely compromising the model results, the nugget was increased until there are no singular matrices or individual kriging weights greater than 1.05 (this allows for some negative kriging weights). To attain this, the nugget was increased to 39.0 (a 244% increase); this is still only 8% of the data set variance. The kriged bedrock surface and estimation variance are shown in Figure 3.11a-b. The average estimation variance is significantly larger for this model than for model I. The difference between the estimation variances (Model III - Model I) are shown in Figures 3.11c and 3.12a. The estimation variance for model III, on average, is 12.7% larger than the estimation variance for model I, but this is not a reasonable reflection of model quality, because the results of model I do not account for the variance due to the nugget.

Model IV (Figure 3.9c) is a modification of model I and accounts for the nugget (although not exaggerated as is necessary for the Gaussian model (III)). The kriged surface and estimation variances are shown in Figure 3.13a-b. The difference between the estimation variances (Model III - Model IV) are shown in Figures 3.13c and 3.12b. Now that the nugget is included, it is reasonable to compare the results of using the traditional anisotropy factor model, to those obtained by using the directional semivariogram model approach. Even though model III has increased the nugget by 244% to stabilize the Gaussian model, the mean difference in the estimation variance between models III and IV is -5.20%. This implies model III's (the directional models) results are better, or at least less uncertain, than the results from model IV. A Q-Q (quantile-quantile) plot is also shown in Figure 3.14 comparing the original (I) estimated residuals vs. each of the other models. It shows

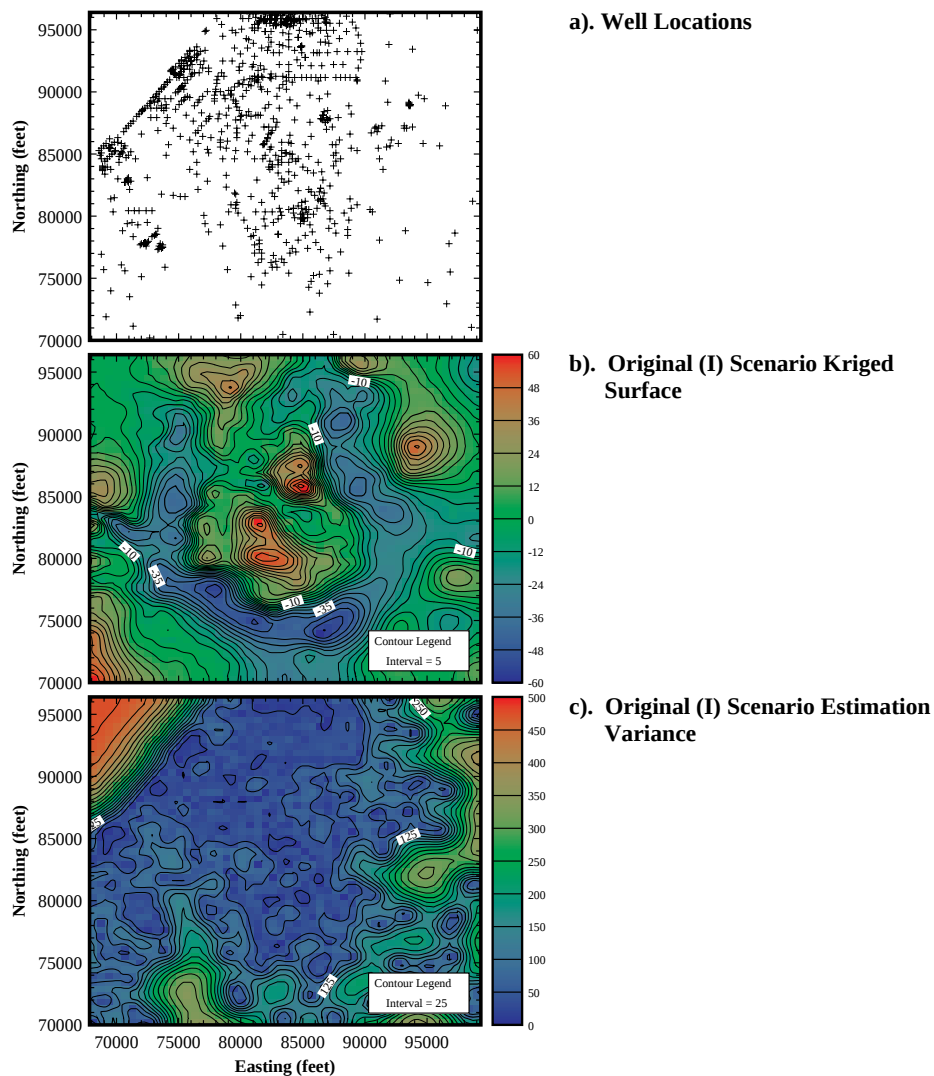


FIGURE 3-10. Location of sample wells at RMA (a), SK map of bedrock elevation residuals (b), and estimation variance using an anisotropy factor, spherical-spherical semivariogram model I (c) (Johnson, 1995).

that the results are similar in all models. Each model generates a similar number of sample values in each of 100 quantiles, but by fine tuning the semivariogram models the estimation variance can be reduced without making any dramatic changes in the overall model statistics.

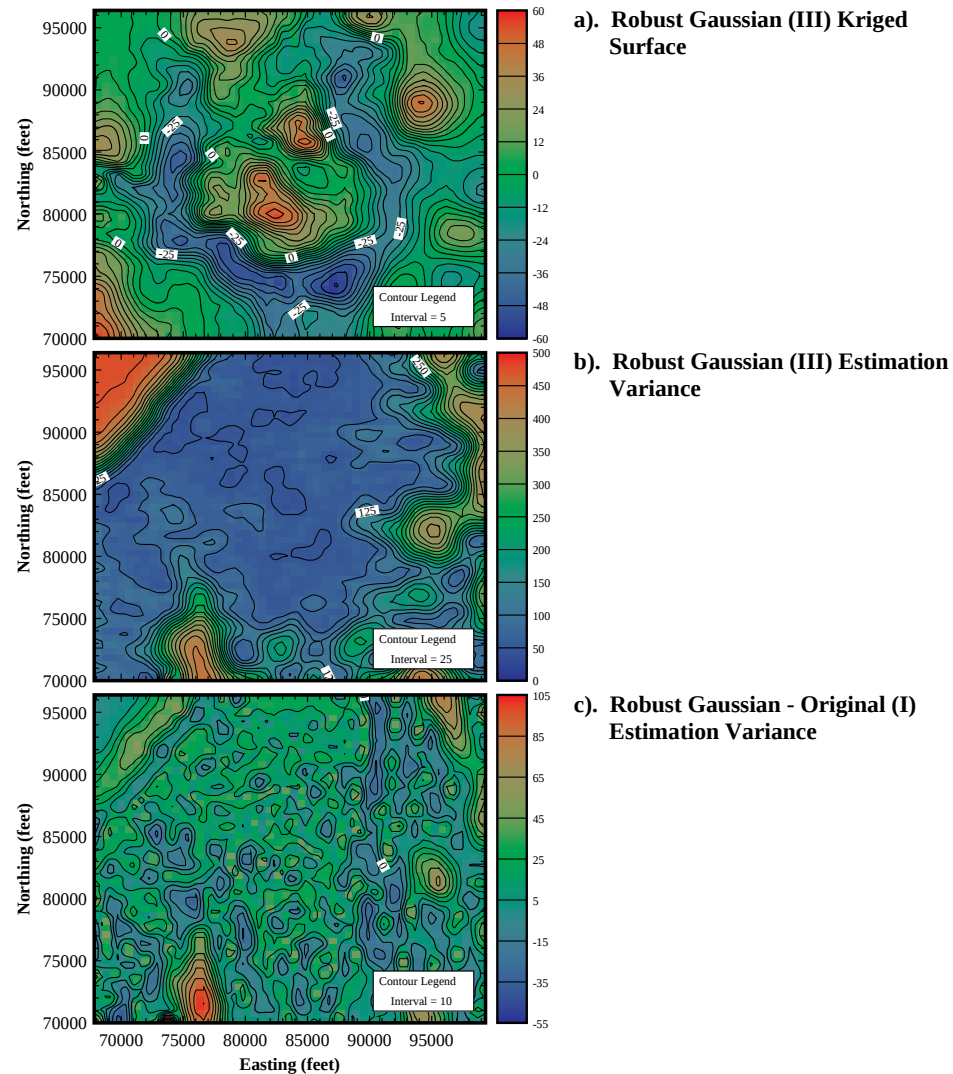


FIGURE 3-11. RMA SK map of bedrock elevation residuals (a), and estimation variance using robust Gaussian factor semivariogram models (b), and difference between robust Gaussian (b) and original (Figure 3.10c) estimation variance maps (c).

In this example, if the nugget is accounted for, the directional semivariograms yield a better result, even when the nugget was artificially exaggerated only for the directional model to prevent problems associated with use of the Gaussian model. In some cases though, unstable models (such as Gaussian) may make the use of directional models undesirable, even when they would, at first, appear justified. As Posa (1989) argues, and his conclusion is supported here, it is sometimes better

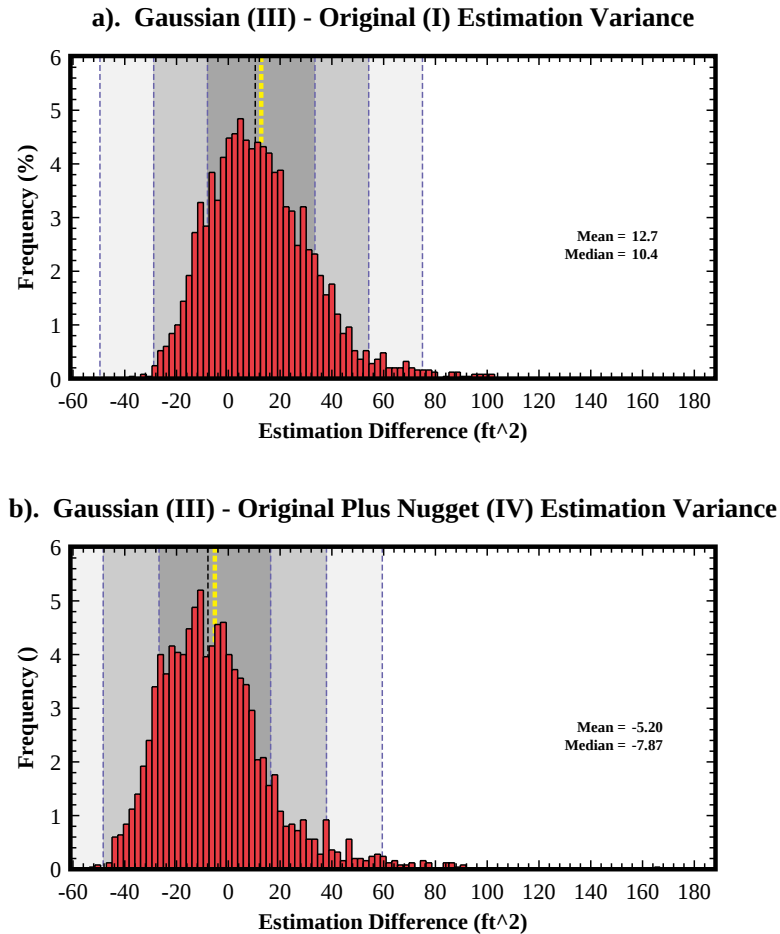


FIGURE 3-12. Distribution of differences between alternative estimation variance maps: (a) the difference between the robust Gaussian (III) and the anisotropy factor, spherical-spherical semivariogram model (I); (b) the difference between the robust Gaussian (III) and the anisotropy factor, spherical-spherical semivariogram model with nugget (I). The positive, average difference in (a) indicates the Gaussian model has a higher average estimation variance. The negative, average difference in (b) indicates the Gaussian model has a lower average estimation variance.

to use a semivariogram model which is not as physically correct, but which is numerically more robust (i.e. a spherical model).

The main problem with implementing directional semivariograms, in this case, was related to the instability in the kriging matrix resulting from theoretical problems associated with using a small nugget and a Gaussian semivariogram model. This, however is a general problem for all kriging methods, and should not reflect adversely on the directional semivariogram method.

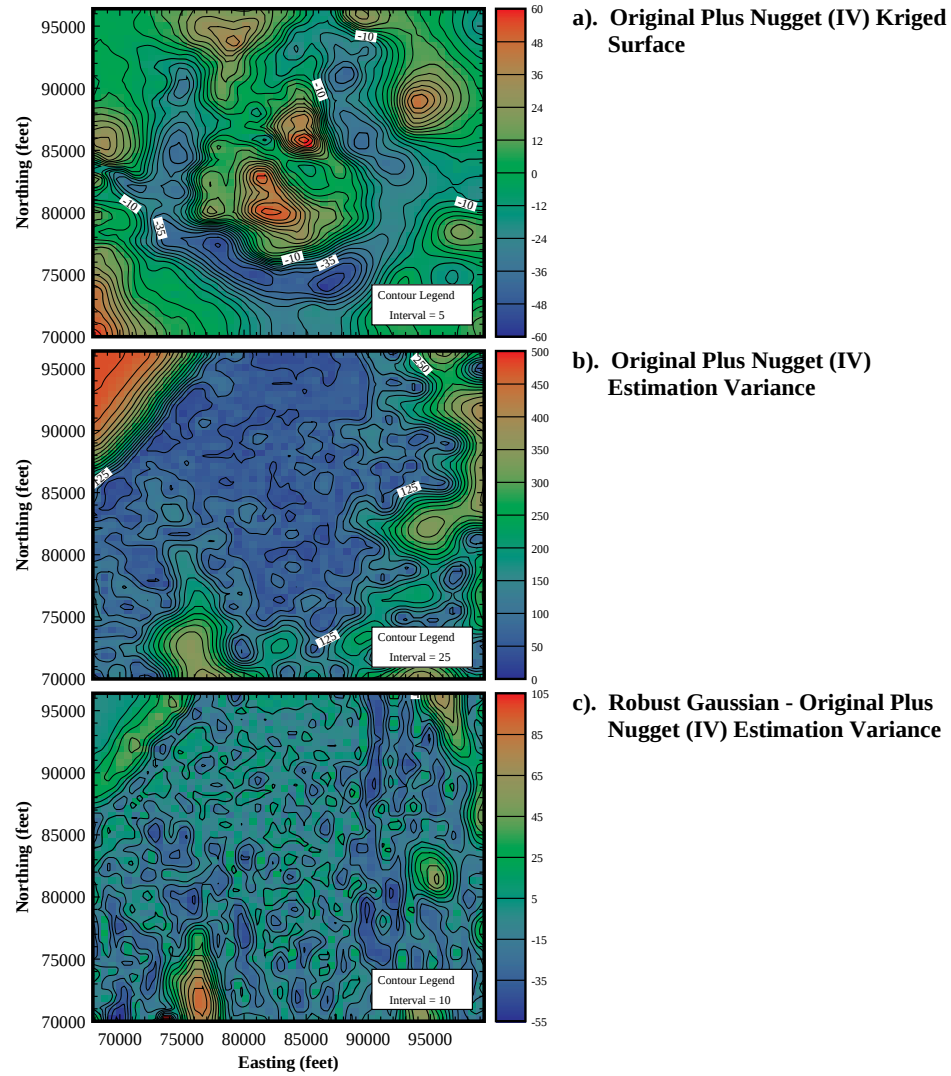


FIGURE 3-13. RMA SK map of bedrock elevation residuals (a), and estimation variance using the anisotropic factor spherical-spherical semivariogram model with a valid nugget (IV) (b), and difference between the robust Gaussian (III) (Figure 3.10c) and estimation variance maps (b).

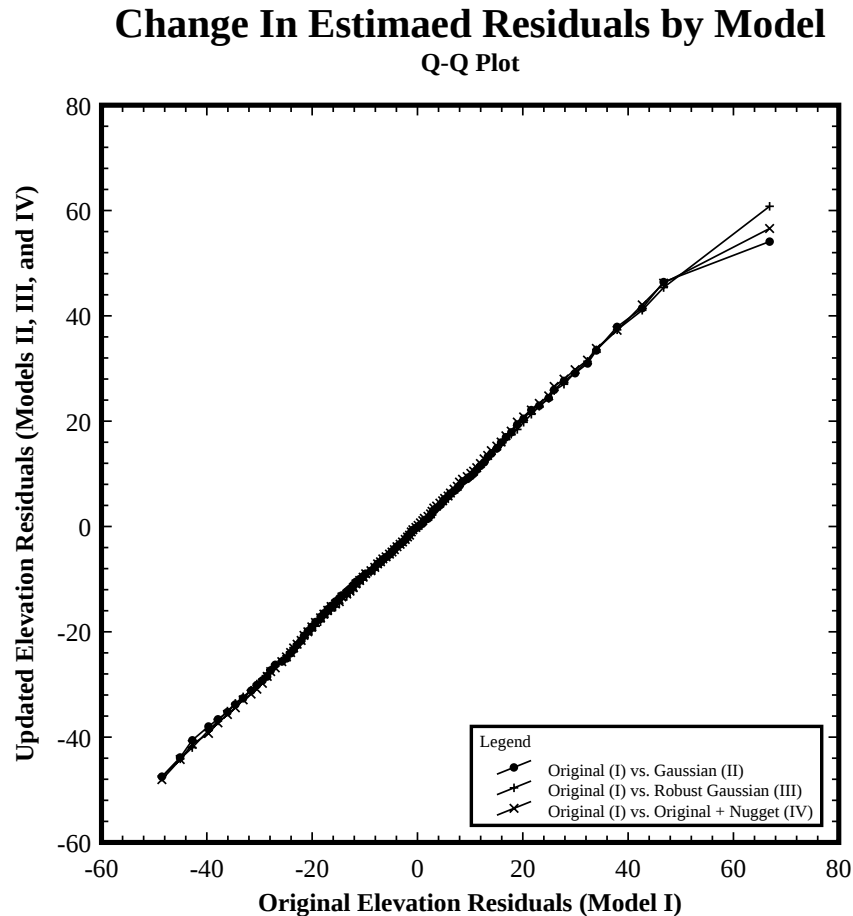


FIGURE 3-14. Q-Q plot of bedrock elevation residuals where the original Spherical model using anisotropy factors (I) is compared versus 1) the original Gaussian model (II), 2) the robust Gaussian model, and 3) the original Spherical model adjusted with a nugget. The plot suggests that the general nature of all the models are similar.

3.6: Conclusions

This chapter demonstrates that better definition of the experimental semivariogram, yields results which better honor the spatial statistics of the sample data. This is illustrated by reduced estimation variance when factors other than model definition are removed. This is accomplished by defining unique model semivariograms along each of the three principle axes of the semivariogram model ellipsoid. In addition to improving the results, the procedure also makes it easier to model

experimental semivariograms, because one need not compromise when selecting model types and sills for each axis. There is an increase in computational effort which increases total processing time in this study (observed times increased 80% to 200%), but this cost is relatively minor when compared to the total time the modeler spends developing semivariogram models. Overall, use of directional semivariogram modeling requires some additional computational time, but modeler effort is reduced, and most important, a significant increase in accuracy may be attained.

CLASS VS. THRESHOLD INDICATOR SIMULATION

With traditional discrete multiple indicator conditional simulation, semivariogram models are based on the spatial variance of data above and below selected thresholds (cut-offs). The spatial distribution of a threshold is difficult to conceptualize. Also, in some cases, ordering of the indicators may influence the results, and changing the arbitrary order, to test sensitivity of the results to the order, involves a substantial effort. If the conditional simulations instead are based on the indicators themselves, rather than the thresholds separating the indicators, then the spatial statistics are more intuitive, and reordering the indicators is a trivial endeavor. When class indicators are used, the indicator order can be switched at any time without recalculating the semivariograms. If thresholds are used, and the ordering is changed, all the semivariograms must be recalculated. Despite the significant difference in methods, the model results are nearly identical.

4.1: Introduction

In traditional Multiple Indicator Conditional Simulation (MICS), the kriged model results are based on semivariograms describing the spatial distribution of the cut-off's between indicators. The affect of the order of the indicators on the resulting realizations is rarely evaluated even though the numerical order is arbitrary. For traditional simulation, the estimated indicator at a location is based on the probability that the location is below each threshold or cut-off (the number of thresholds equals the number of indicators minus one). A more intuitive approach is based on calculating the probability of occurrence of each individual indicator. This chapter presents a technique which uses semivariogram models based on individual indicators (classes), as opposed to the traditional threshold semivariograms which are based on all the indicators below a cut-off versus all the indicators above the cut-off.

These differences can be described mathematically as follows. Where the data set has been differentiated into a finite number of indicators, it is possible to define a random function ($Z(x)$)

whose outcomes will have values in the range $z_{\min}$ to $z_{\max}$. From the definition of the indicators, K thresholds can be defined ($K + 1$ equals the number of indicators) where:

$$z_1 < z_2 < \dots < z_K \quad (4.1)$$

The random variable $Z(x)$ can then be transformed into an indicator random variable $I(x; z_k)$ by:

$$I(x : z_k) = \begin{cases} 1, & \text{if } Z(x) \leq z_k \\ 0, & \text{if } Z(x) > z_k \end{cases} \quad k = 1, \dots, K \quad (4.2)$$

The first moment of the indicator transform yields :

$$\begin{aligned} E\{I(x : z_k)\} &= 1 \times P\{Z(x) \leq z_k\} + 0 \times P\{Z(x) > z_k\} \\ &= P\{Z(x) \leq z_k\} \end{aligned} \quad (4.3)$$

where $E\{I(x; z_k)\}$ is the expectation of $I(x; z_k)$, and $P\{Z(x) \leq z_k\}$ and $P\{Z(x) > z_k\}$ are the probabilities $Z(x)$ is less than or greater than the threshold z_k . This equation is equivalent to the univariate cumulative distribution function (CDF) of $Z(x)$. For classes, similar equations can be defined. Classes (c_i) are equivalent to the indicators defined using thresholds in equation (4.1); they can also be defined by:

$$c_i = \begin{cases} 1, & \text{if } Z(x) \leq z_1 \\ 2, & \text{if } z_1 < Z(x) \leq z_2 \\ \dots & \\ K, & \text{if } z_{K-1} < Z(x) \leq z_K \\ K+1, & \text{if } Z(x) > z_K \end{cases} \quad (4.4)$$

Once the classes are defined, the random variable $Z(x)$ can then be transformed into an indicator random variable $I(x; z_k)$ by:

$$I(x : c_i) = \begin{cases} 1, & \text{if } Z(x) = c_i \\ 0, & \text{if } Z(x) \neq c_i \end{cases} \quad i = 1, \dots, K+1 \quad (4.5)$$

and the first moment of the indicator transform yields:

$$\begin{aligned} E\{I(x : c_i)\} &= 1 \times P\{Z(x) = c_i\} + 0 \times P\{Z(x) \neq c_i\} \\ &= P\{Z(x) = c_i\} \end{aligned} \quad (4.6)$$

Here, instead of this defining the univariate CDF, the univariate probability distribution function (PDF) is defined. By summing the PDF components though, it can be converted into the univariate CDF defined by equation (4.3).

Because the equations to define the class or threshold expectation are fundamentally the same, the class method generates realizations that are equally accurate to threshold realizations, but it has two main advantages. First, it is easier to conceptually relate the model semivariograms to the spatial distribution of the materials. When class semivariograms are calculated, the range reflects the average size of the indicator bodies (Figure 4.1):

Class	Horizontal Range
Silt	112
Silty-Sand	106
Sand	60
Gravel	41

where as, the threshold semivariograms represent the distribution of indicators above or below a threshold:

Threshold	Horizontal Range
Silt vs. (Silty-Sand, Sand, & Gravel)	112
(Silt & Silty-Sand) vs. (Sand & Gravel)	68
(Silt, Silty-Sand, & Sand) Vs. Gravel	41

and these can be difficult to conceptually relate back to the original data in complex geologic settings. It is important to note, that the first and last class and threshold semivariograms will always be identical (they are based on equivalent indicator sets (0's and 1's)). The intermediate semivariograms, though may vary substantially. The intuitive sense for the threshold semivariogram range also tends to decrease with an increasing number of indicators. Class semivariogram ranges though, still reflect the average size of the indicator body. The second advantage to using classes is that sensitivity to indicator ordering can be evaluated without developing additional semivariogram models. If thresholds are used, the full suite of threshold semivariogram models must be recalculated for each reordering. The class approach does have several disadvantages: 1) more order relation violations occur (discussed later), 2) it is computationally more expensive (one additional kriging matrix must be solved per grid cell), and 3) it requires one additional semivariogram model (the number of class semivariograms equals the number for thresholds, plus one). The last two items are only an issue, if ordering sensitivity is not a concern. If sensitivities are a concern, preparation for the threshold method requires far more human effort and computer time to develop the additional semivariogram models.

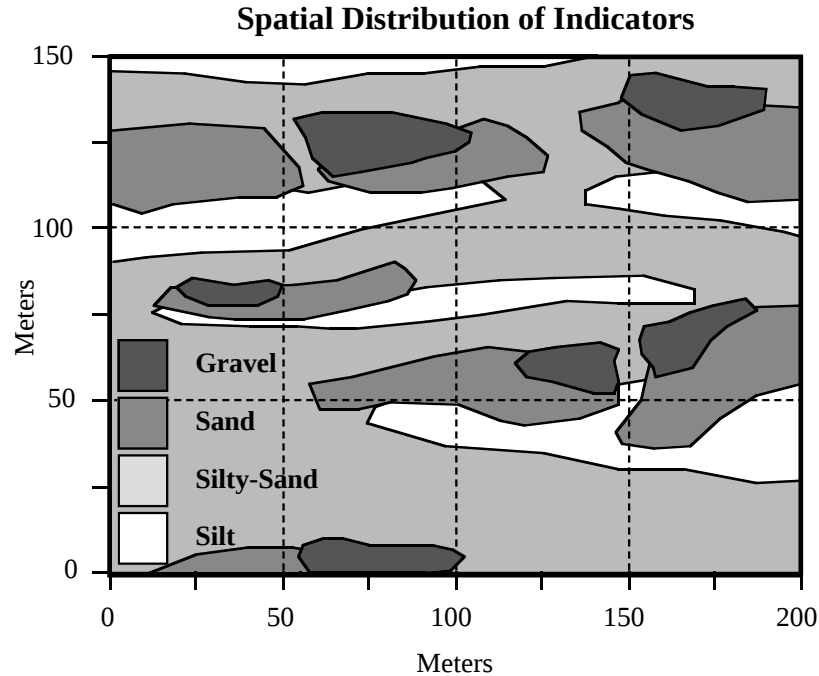


FIGURE 4-1. Spatial distribution of several indicators. Defining semivariograms based on indicator classes is more intuitive, because the range reflects the average size of the indicator bodies. The class semivariogram model ranges are: silt = 112m, silty-sand = 106m, sand = 60m, and gravel = 41m. For thresholds, semivariogram model ranges are: silt vs. all others = 112m, silt and silty-sand vs. sand and gravel = 68m, and gravel vs. all others = 41m.

4.2: Previous Work

The "best-estimate" of the conditions at a site may not necessarily be a realistic interpretation of actual site conditions. By their nature, estimation techniques such as Ordinary Kriging, are averaging algorithms which smooth much of the true site variability (Figure 4.2). To address this issue and develop a technique that would both honor the data and their spatial statistics, conditional (constrained by field data) simulation techniques were developed. These techniques use a probabilistic (Monte-Carlo) approach to estimate site conditions. When estimating a value for a particular location, the probability that the value is less than each threshold is determined, a random number is generated, and an indicator value is assigned based on that random number. As a result a single "realization" will retain much of the variability exhibited by the field data, but a single "realization" may be a poor representation of actual site conditions. When using simulation techniques, many models must be calculated; each preserves the nature of the spatial data, but each

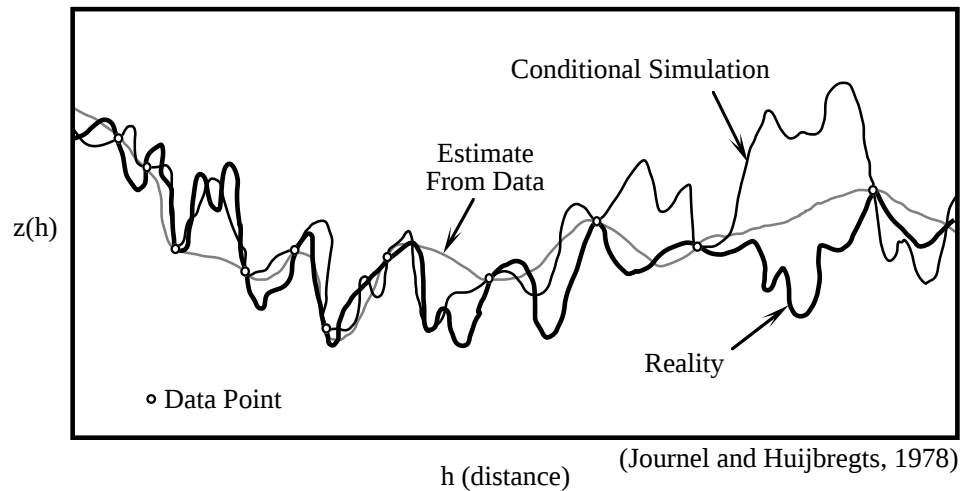


FIGURE 4-2. Ordinary Kriging (and most other estimation methods) tends to average or smooth data to achieve a best linear unbiased estimate (BLUE) of reality. Indicator Kriging with conditional simulation provides a means for modeling the variability observed in nature, while still honoring the field data. Conditional simulation does not produce a best estimate of reality, but it yields models with characteristics similar to reality. When multiple realizations are made and averaged, values will approximate the smoothed, BLUE.

has a random component. When grouped together (assuming an adequate number of simulations are run), the average result is, in theory, the same as a "best-estimate" kriged map.

Several different varieties of conditional simulation are commonly used. Some are based on continuous data (e.g. contaminant concentrations; Deutsch and Journel, 1992), and others on discrete data (e.g. geologic units; Deutsch and Journel, 1992). A simple example to distinguish the two methods is to consider two points representing two different indicators (1 and 3). If an estimate for a point mid-way between the indicators is desired, the results can be quite different depending on which method is used. If a continuous simulator is used, the result would be the average, or indicator 2. If the indicators represent concentrations (indicators #1 = 1 ppm, #2 = 10 ppm, and #3 = 50 ppm) the result is reasonable. If the indicators represent geologic units (indicators #1 = clay, #2 = sand, #3 = basalt), the averaged solution does not have a physical basis (sand is not intermediate to clay and basalt). Discrete simulation should be used for the latter case, and the result would be either indicator 1 or 3. If continuous data are used, either simulation approach can be applied, but only discrete simulation is considered in this chapter.

Indicator simulation requires that one or more semivariogram models be calculated; one for each threshold (sometimes a single median semivariogram model based on the median sample value is applied to all thresholds). To estimate a value for a particular location, the distance, direction, and value of the neighboring samples is used to determine the probability that the estimate will be less than each threshold. This process generates a cumulative density function (CDF). Once the CDF is

calculated, a random number is generated, and for that realization, a specific estimate is selected. Class and threshold indicator kriging generate the CDF in different ways as shown in Figure 4.3. A

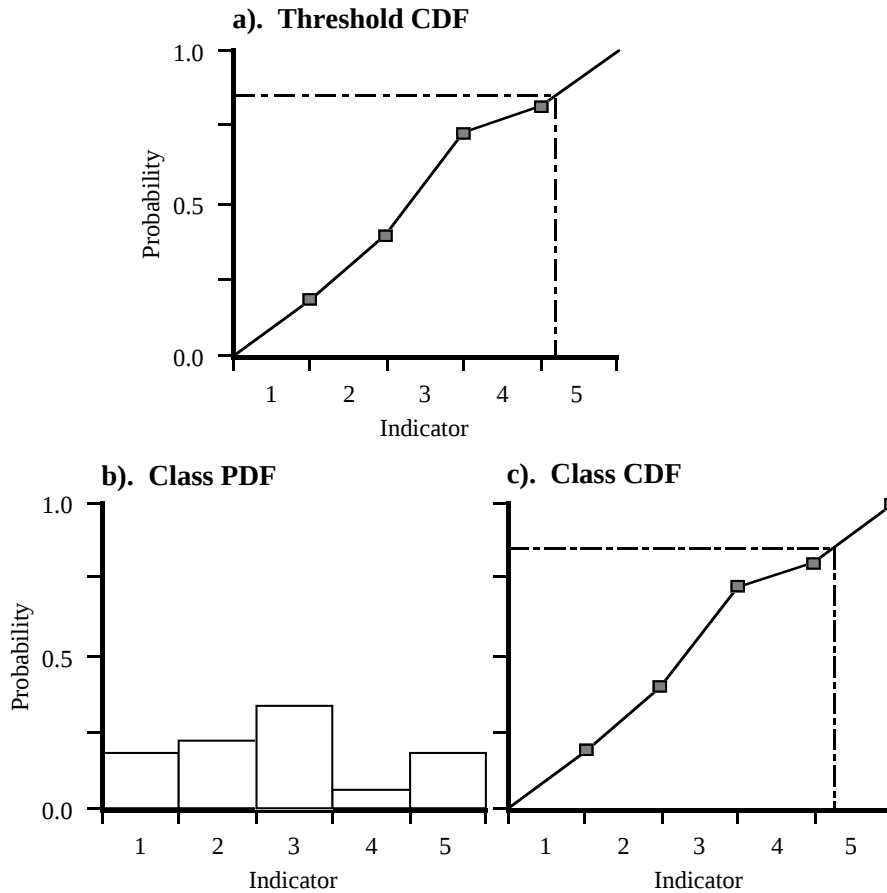


FIGURE 4-3. Class and threshold indicator kriging generate the cumulative density function (CDF) in different ways. Threshold CDF's are determined directly from the probability that the specified grid location is less than each threshold level (a). The final CDF term should be less than 1.0 with the remaining probability attributed to the final indicator. Calculating class CDF's requires two steps. First the probability of occurrence of each class is calculated (b). The PDF is converted into a CDF by summing the individual PDF terms (c). Ideally the probabilities will sum to 1.0. For both the threshold and class approaches, a random number between 0.0 and 1.0, is generated to determine the estimated indicator for the cell. From the random number (e.g. 0.82), a horizontal line is drawn across to the CDF curve, and a vertical line is dropped from the intersection, to identify the indicator estimate (5).

random path is followed through the grid (Figure 4.4), because, unlike ordinary kriging methods, previously estimated values are treated as hard data samples and influence subsequent estimates. Finally, to perform a full analysis of a site, many realizations (the resultant map from one

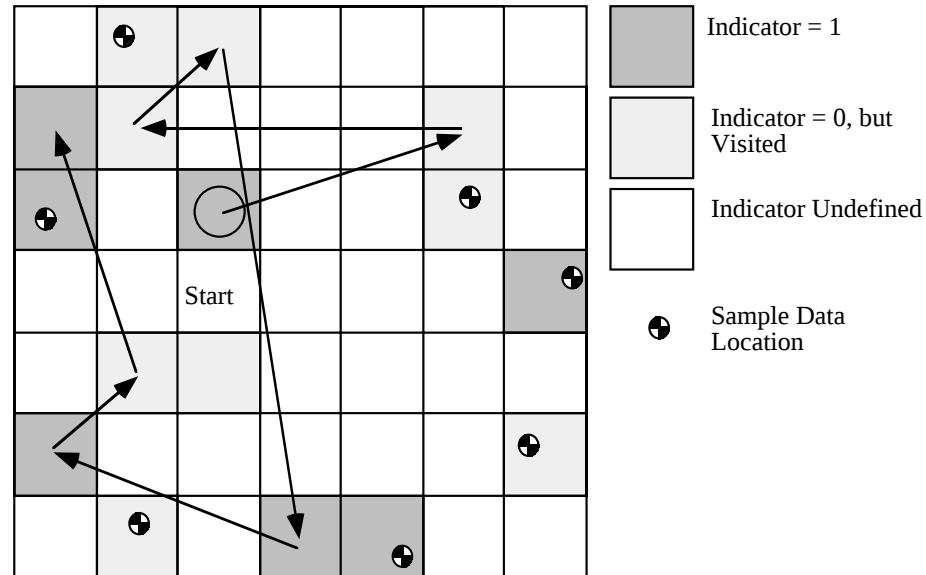


FIGURE 4-4. This illustration shows the step wise manner in which a grid is kriged using Indicator Kriging in conjunction with stochastic simulation. Grid cells containing sample data (hard data and some types of soft data) are defined prior to kriging. Once these points are defined, the remaining cells are evaluated. To krig an unestimated cell, a random location is selected, evaluated and redefined as a hard data point, then the next undefined cell is randomly selected. This cell selection and estimation process is continued until all grid cells have been visited and defined.

simulation) must be generated, because each realization represents only one possible interpretation; not the best or most likely interpretation.

4.3: Methods

Two steps of the simulation process are modified in order to use classes rather than thresholds for simulation. First, indicator semivariograms are calculated based on the individual indicators rather than thresholds, and second, the indicator kriging algorithm defines the kriging matrix based on the probability an indicator occurs, as opposed to the probability that the location being estimated is below a given threshold. These changes were incorporated into an existing computer program, SISIM3D .

4.3.1: Semivariogram Calculation

To calculate a traditional indicator (“threshold”) semivariogram, an individual threshold or cut-off is selected (Journel and Huijbregts, 1978). All values below the cutoff are assigned a 1, and values above the cutoff are assigned a 0. When using “class” semivariograms, data locations with sample values that equal the indicator value being simulated are set to 1, the remaining values are set to 0. The class approach differs from threshold approach in that both a low and a high cut-off are defined.

4.3.2: Data Definition

The hard and soft data labeling conventions are defined differently for class and threshold simulations. For both approaches, each data point is transformed into an indicator mask composed of 0’s and 1’s (some soft data may have an associated probability distribution reflecting a weight between 0 and 1, for a particular class or threshold level). Using traditional methods, the mask is set to 0 if the data value is less than the specified indicator threshold, and the mask is set to 1, if the data value is greater than the specified indicator threshold. For example, hard data with the indicator order basalt, clay, silt, sand, gravel, and cobbles, would have the following traditional indicator masks:

```
Basalt   = 11111
Clay     = 01111
Silt     = 00111
Sand     = 00011
Gravel   = 00001
Cobbles  = 00000
```

There is one less mask (5) than there are indicators (6). For class semivariograms, the mask indicates whether the data point is (1) or is not (0), the specified indicator. For the same example given above, the masks would be:

```
Basalt   = 100000
Clay     = 010000
Silt     = 001000
Sand     = 000100
Gravel   = 000010
Cobbles  = 000001
```

Using the class method, the number of masks equals the number of indicators.

Soft data are those associated with non-negligible uncertainty. Three different soft data types are summarized below:

- Type-A: Imprecise data. These data are classed as a given indicator with associated misclassification probabilities described in the next section.

- Type-B:Interval bound data. The value at these locations is known to fall within a given range (i.e., the probability of occurrence is zero outside of the interval), but the probability distribution within that range is unknown.
- Type-C:Prior CDF data. A probability density function (PDF) is known for these data. The data could be one of several indicators; the most likely indicator is defined by the PDF. The PDF could be defined from an analogous site or expert opinion.

Several masking examples for both class and threshold indicators are given below:

<u>Class</u>	<u>Threshold</u>	<u>Comments</u>
• Type-A = 001000	00111	The quality of this information is described with a p_1 - p_2 term (see next section).
• Type-B = 001110	00111	The datum is known to represent one of several indicators. There is no information available though to describe which indicator is most likely. The PDF is built by kriging the surrounding data.
• Type-C = 001110	00111	The datum is known to represent one of several indicators, and there is a PDF available to describe the probability of occurrence for each indicator (e.g., 0%, 0%, 20%, 50%, 30%, 0%).

For Type-B data, the threshold method requires that additional information be stored defining the top of the interval. These notation methods are not strict theoretical requirements, but are conventions for this particular algorithm.

4.3.3: P_1 - P_2 Calculations

When describing Type-A (imprecise) data, the probability that the data correctly, or incorrectly, reflect the value being classified is defined by the misclassification probabilities, p_1 and p_2 . They are defined as:

p_1 : Given that the actual value is less than the threshold (or in the class), p_1 is the probability that the measured value is less than the threshold (or in the class) (correctly classified).

p_2 : Given that the actual value is NOT less than the threshold (or not in the class), p_2 is the probability that the measured value is less than the threshold (or in the class) (incorrectly classified).

These values are determined by comparing the soft data to co-located hard data with a training set. After p_1 and p_2 have been determined, the misclassification probabilities can be used for the same type of soft data, at locations where hard data are not present.

Using indicator thresholds, p_1 and p_2 are determined by measuring the ability of soft information to correctly classify the hard training set data above and below a specified threshold level. This is

shown graphically, for two thresholds, in Figure 4.5. The misclassification probabilities are defined as:

$$p_1 = A / (A + D) \quad (4.7)$$

$$p_2 = B / (B + C) \quad (4.8)$$

In region A, points are correctly classified as being below the specified threshold. In region C, they are correctly defined as being above the threshold. In regions B and D, the soft data incorrectly classify the sample. Ideally p_1 is greater than p_2 . For hard data $p_1 = 1.0$ and $p_2 = 0.0$. If the soft data are not correlated with the hard data $p_1 = p_2$ (NOTE: p_1 and p_2 are not expected to sum to 1.0). The difference between p_1 and p_2 indicates how useful the soft data are for classifying the samples. When using indicator classes, rather than thresholds, the implications of p_1 and p_2 are the same, but calculating p_1 and p_2 is more complex and the misclassification probabilities tend to increase as the number of classes increases. A graphical representation for calculating p_1 and p_2 is shown for three classes in Figure 4.6. The misclassification probabilities are defined as:

$$p_1 = E / (D + E + F) \quad (4.9)$$

$$p_2 = (B + H) / (A + B + C + G + H + I) \quad (4.10)$$

In region E, points are correctly classified as being included in the specified class. In regions A, C, G, and I, they are correctly defined as being outside of the class. In regions B, D, F, and H, the soft data incorrectly classify the sample.

Due to the nature of the p_1 - p_2 classification scheme, the results for the class and threshold p_1 - p_2 terms are identical for the first and last indicators (Figures 4.5 and 4.6: Threshold₉₂₅₀ p_1 - p_2 = Class_{<9250} = 0.67, and Threshold₁₀₀₀₀ p_1 - p_2 = Class_{>10000} = 0.73). This is because, in the class method, the upper or lower bound is missing, and the equations reduce to that used for thresholds. For other classes and thresholds, the p_1 - p_2 will vary significantly. If a single threshold is used, there are only four possible classifications (A, B, C, D). Basically the soft sample values need only be on the correct side of the cut-off for the threshold to correctly identify the hard data sample. Using classes, the soft data have both high and low cut-offs, therefore the soft data precisely identify a location as being, or not being, a member of a hard data class (Figure 4.6). This is a more restrictive constraint and as a result, the interior class p_1 - p_2 values are lower than those for threshold simulation. The quality of the soft data has not changed, it is just defined differently. What has changed is the ability of the algorithm to describe the imprecision. Other approaches have been proposed, but are not implemented here.

4.3.4: Difference Between Prior Hard and Prior Soft Data CDF's for Class and Threshold Simulations

An additional and important difference between class and threshold simulation is the definition and treatment of the difference in the hard data and soft data prior probability distributions. After the kriging matrix has been solved, the CDF is estimated for each class or threshold. The CDF

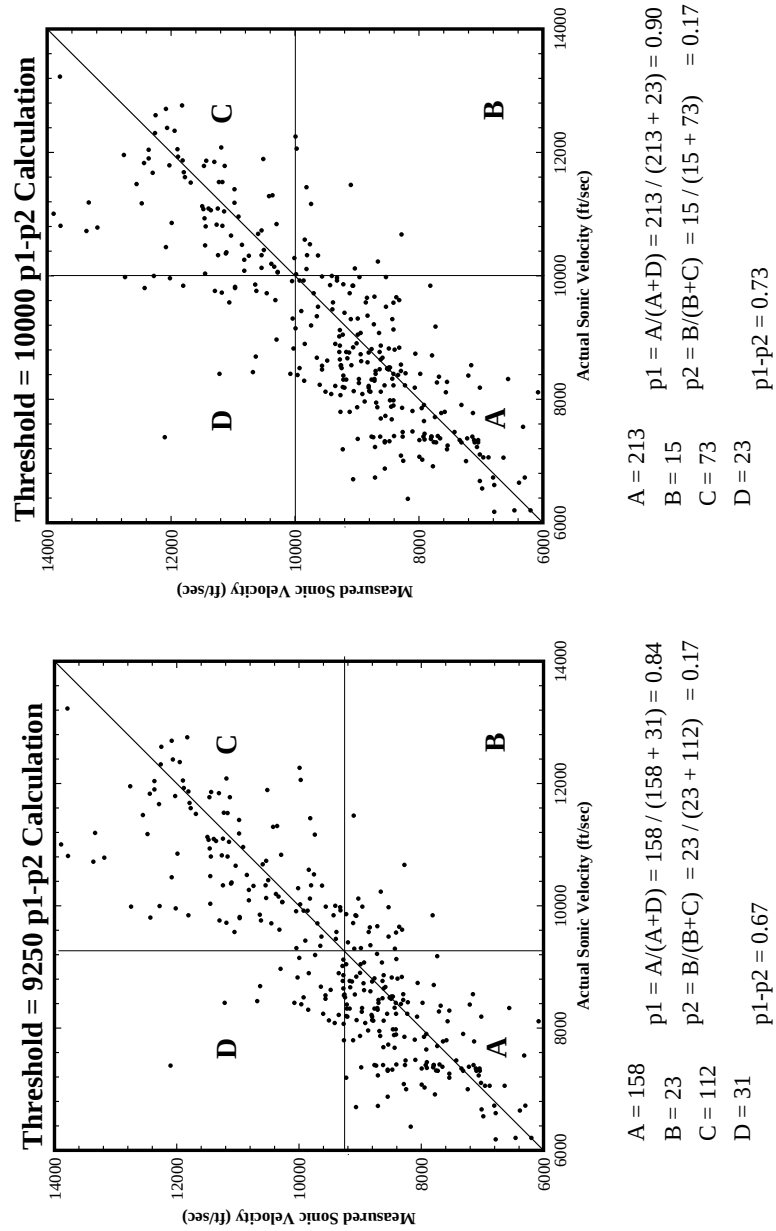


FIGURE 4-5. Graphical method for calculating p_1 and p_2 values for a specific threshold (After Alabert, 1987). Data from CSM Survey Field.

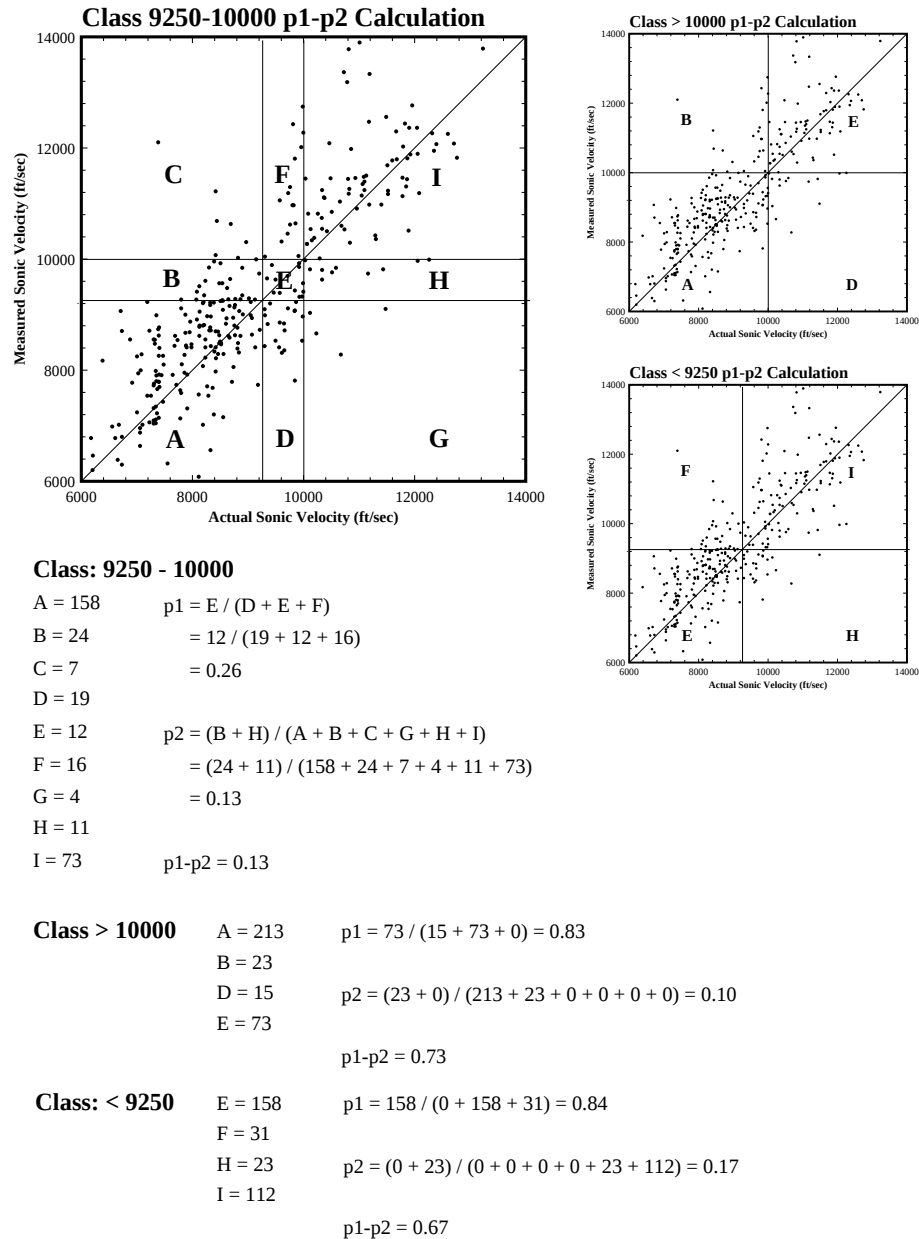


FIGURE 4-6. Graphical method for calculating p_1 and p_2 values for a specific class. Data from CSM Survey Field.

estimate is calculated from the hard and the soft data points in the search neighborhood, the relative frequency of each indicator in the prior hard and soft data, and by the soft data scaled by the p_1 - p_2 term discussed above. Often, hard and soft data collection techniques suggest different percentages of each indicator occurring at the site. If the simulator uses thresholds, the correction term is based on :

$$| (\text{percentage of hard data} < \text{threshold}) - (\text{percentage of soft data} < \text{threshold}) |$$

If classes are used, the correction term is based on:

$$| (\text{percentage of hard data} = \text{class}) - (\text{percentage of soft data} = \text{class}) |$$

The difference is subtle but important. For the threshold approach, if the probability of a single threshold varies significantly between the hard and soft data, particularly if it is the first threshold, the importance of the remaining thresholds can be under-valued. Reordering the indicators could alleviate some of this problem. For the class approach, the relative occurrence of each indicator is directly compared, therefore when one class has very different prior hard and prior soft probabilities, these will not seriously affect other class estimates. This is because, indicators are directly compared, and errors are not cumulative.

4.3.5: Order Relation Violations

As with traditional threshold simulation, the class CDF for a particular grid location may not be monotonically increasing and may not sum to 1.0. These are order relation violations (ORV's). They can be caused by use of inconsistent semivariogram models for the different thresholds or classes, or by use of different prior probabilities and p_1 - p_2 weights applied to soft data. Using the algorithms described here, thresholds and classes manage ORV's in slightly different manners. This is, in part, due to theoretical differences in how the CDF's are generated, but it is also due to technical difficulties in equating the threshold CDF and the class PDF.

The method for managing threshold ORV's in SISIM3D was not modified, but the methods used were not appropriate for classes. Therefore a new set of tools for managing class ORV's was developed. The differences between the two methods are described below and are diagrammed in Figure 4.7.

For the threshold method, one type of ORV occurs when the CDF declines from one threshold to the next (Figure 4.7a). A CDF is a cumulative probability, so a declining CDF is physically impossible. It is not possible to determine which threshold causes the problem, therefore to remedy the situation, the average of the two probabilities is assigned to both thresholds (note, the indicator associated with the declining CDF term, has zero probability of occurrence). For classes, the equivalent problem is an individual class having a negative probability of occurrence (Figure 4.7b: indicator #2), which is also physically impossible. In this case though, it is reasonable to simply assign that class a zero probability of occurrence. There is no obvious reason to distribute the error to another unrelated indicator.

Another type of ORV occurs when the CDF sums to a value greater than, or less than 1.0. For thresholds, the last threshold CDF term is often less than 1.0 (Figure 4.7a), and it is assumed that

Order Relation Violations

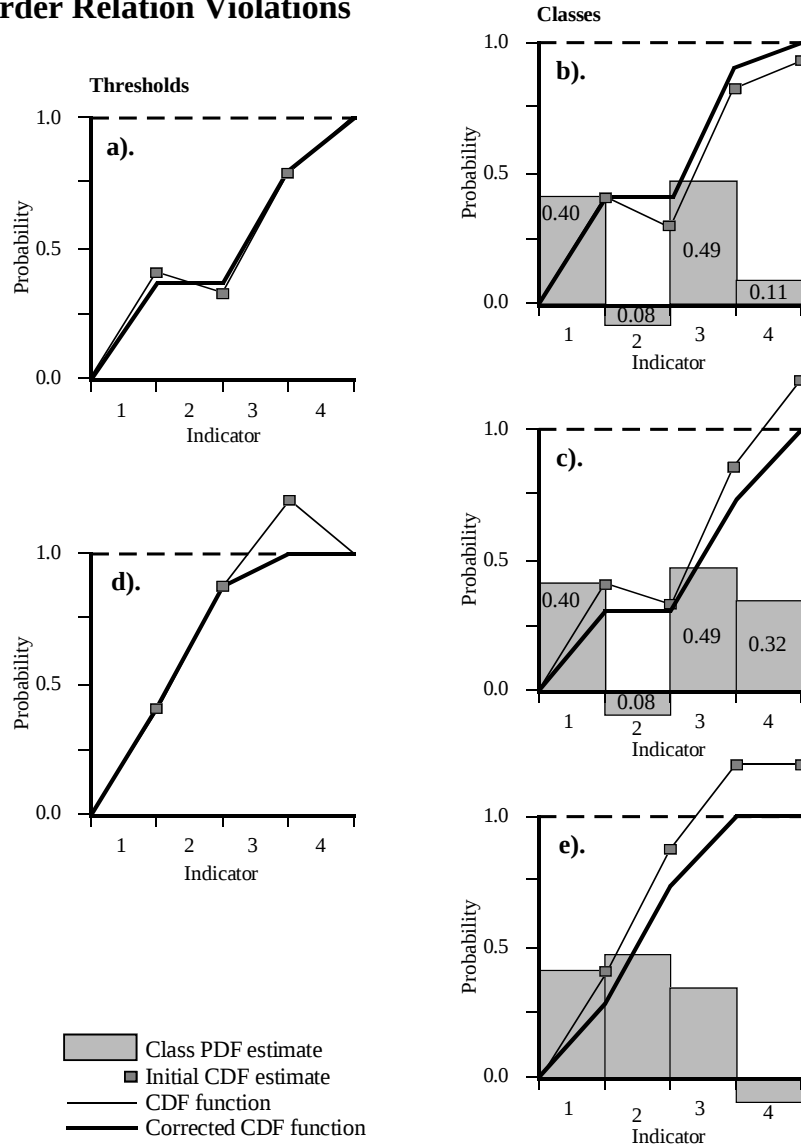


FIGURE 4-7. For both the class and threshold approach, there are two basic types of order relation violations (ORV). a) An individual CDF probability is less than the CDF of a smaller threshold (the CDF is decreasing); this is equivalent to a class having a negative probability of occurrence. This type of ORV is resolved for thresholds by averaging the two CDF's so that they are equal; for classes, a 0.0 probability of occurrence is assigned to the PDF. b) When cumulative probabilities are greater than 1.0, the value is truncated to 1.0 for the threshold approach, while for the class approach, the probability of each class is proportionally rescaled, so that the CDF will sum to 1.0.

the balance of the CDF is described by the final indicator. This is generally the case, but as shown in an equivalent class example, it is possible for the final indicator to account for significantly more (or less) than the remaining portion of the CDF (Figure 4.7c). With classes, the overestimate in the CDF can be proportionally absorbed by each of the PDF components ($\text{PDF}_{\text{new}} = \text{PDF}_{\text{old}} / \text{CDF}_{\text{final value}}$). In this example though, the threshold method would not have recognized that an ORV occurred. It is also possible that the threshold CDF will exceed 1.0 (Figure 4.7d). Currently thresholds manage this problem by truncating the CDF to 1.0 for the affected threshold (and all following thresholds). This solution is not very satisfying, in part, because it implies the offending threshold level is fully responsible for the error, even though the CDF is a cumulative probability (i.e., an earlier threshold could be the root cause of the problem). It is also unsatisfying because it biases results to the lower order indicators. Classes again, manage this situation by distributing the error over all the PDF components ($\text{PDF}_{\text{new}} = \text{PDF}_{\text{old}} / \text{CDF}_{\text{final value}}$; Figure 4.7e).

As implemented, the techniques for managing class ORV's are less biased than the threshold method. This is fortunate, since the class method also produces more ORV's. It is felt though, that some of these extra ORV's occur when the threshold CDF does not sum to 1.0, and a mistaken assumption is made that the remaining indicator exactly contributes the remaining portion of the CDF (compare Figures 4.7a vs. 4.7c).

4.4: Applications

Two data sets are used to demonstrate that class indicator simulations generate statistically identical realizations as threshold indicator simulations. The first is a simple synthetic data set with fourteen hard data points. The two series of solutions yield similar results, but are not exactly the same, because of the differences in how ORV's are managed. The second data set is from the Colorado School of Mines Survey Field in Golden, Colorado and includes hard data, as well as Type-A, B, and C soft data. The use of classes rather than thresholds is not meant to improve results, rather it is intended to render the process more intuitive, and to facilitate testing the sensitivity of simulations to the order of the indicators.

4.4.1: Synthetic Data Set

The synthetic data set is composed of fourteen samples distributed in two-dimensional space, representing one of three indicators (silt, silty-sand, and sand) (Figure 4.8). A single isotropic median semivariogram model is assumed for each indicator threshold or class, because with this small data set, it was not possible to generate useful experimental semivariograms for each threshold or class. Use of a median indicator semivariogram model under these conditions is a reasonable and recommended approach. This also ensures that the differences between the resulting simulations is a function of the algorithm and not due to differences in the semivariogram models. The median semivariogram model is:

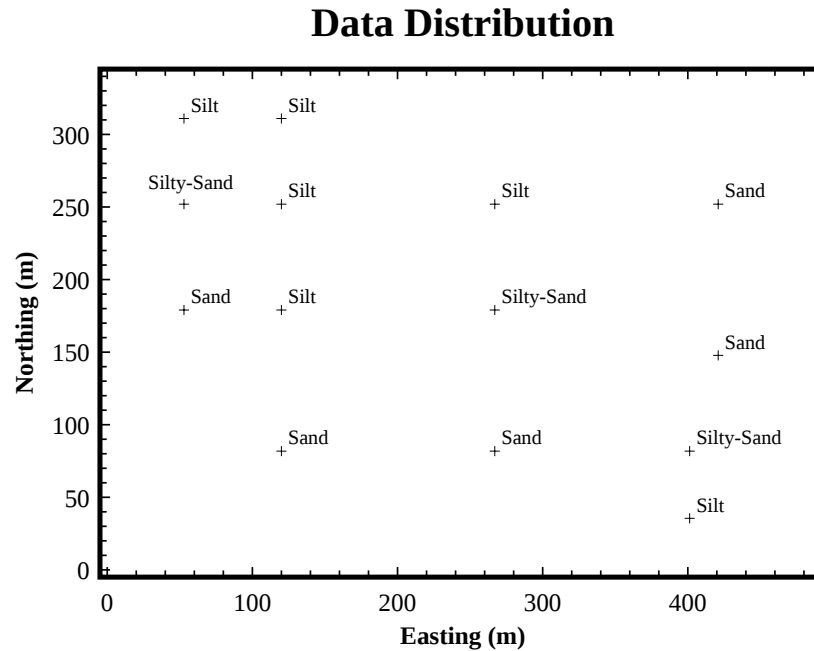


FIGURE 4-8. Synthetic data set distribution.

Model Type = Spherical
 Range = 162m
 Sill = 0.244
 Nugget = 0.0

A regularly spaced, two-dimensional, 50 by 35 grid of 10 m by 10 m grid cells is used to create six series of 200 realizations each (this defines the final grid; a coarse pre-grid 20m by 20m was first calculated. This is managed within SISIM3D and is fairly transparent to the modeler). Three series are generated for the threshold approach and for the class approach with the same reordering of indicators. For the first simulation series, the indicators are defined as:

Silt = 0
 Silty-Sand = 1
 Sand = 2

For the second series, the indicator order was reversed:

Silt = 2
 Silty-Sand = 1
 Sand = 0

and because the order is arbitrary, the indicators in the last series were defined as:

Silt = 0
 Silty-Sand = 2
 Sand = 1

If a median indicator semivariogram model was not being used, the threshold semivariogram models would have to be recalculated, but the class semivariogram models would remain unchanged. The averaged results of the simulation series are expected to be nearly identical in these cases, because a median indicator semivariogram is used, but in a field application, different semivariogram models would be used for each class and threshold, and the results of the class simulations are likely to vary from the threshold results. For individual simulations, changing the indicator order will change results, because the new ordering also changes the CDF. The indicator components of the PDF are unaltered, but with the new order, a different CDF is built (Figure 4.9).

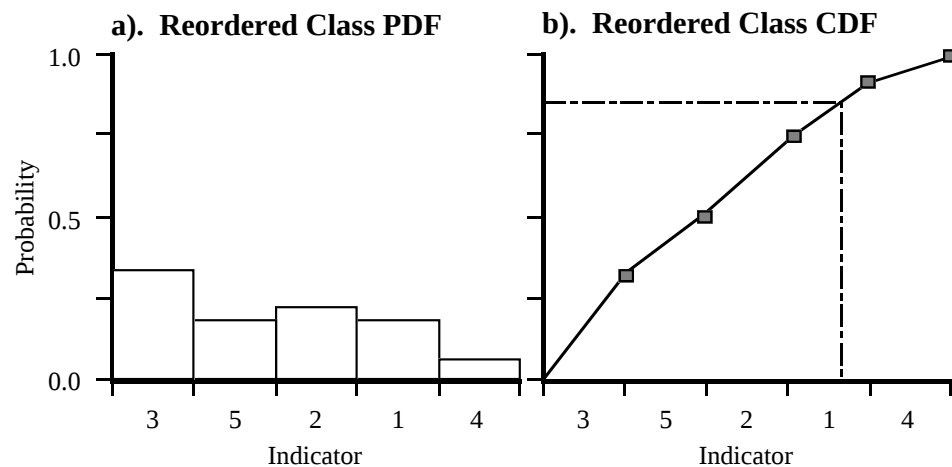


FIGURE 4-9. Reordering indicators in conditional simulation changes the results for an individual simulation grid cell, because the CDF changes along with the indicator ordering, even though the individual components of the PDF do not. Here the indicators from Figure 4.3 have been reordered. The same random number is used, but now, instead of indicator #5 being selected, indicator #4 is selected.

As a result, when the same "random" number is used are used to select from the CDF, a different indicator is selected (Figure 4.9).

4.4.1.1: Initial Indicator Ordering

The class and threshold simulations generate visually similar, but not identical realizations for the original indicator ordering scheme (silt = 0, silty-sand = 1, sand = 2). Cell by cell comparison of the realizations reveals that differences do occur, but these differences are caused by differences in the way order relations violations are resolved in the two methods. Despite these differences, the results are sufficiently similar, that both sets of results are considered reasonable and acceptable.

Several realizations from each simulation series are shown in Figures 4.10 and 4.11. For each

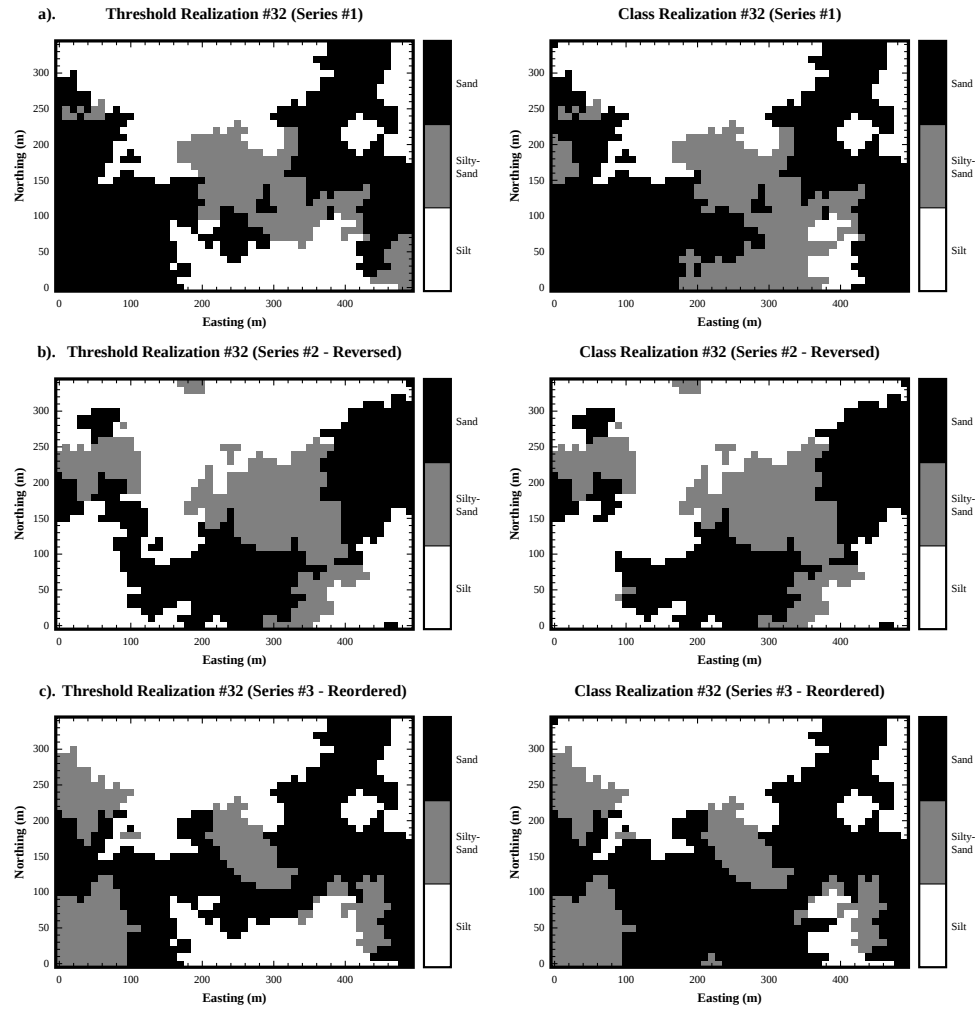


FIGURE 4-10. Realization (#32) pairs for the a) original indicator ordering, b) reversed ordering, and c) arbitrary reordering for thresholds (left) and classes (right). In these realization pairs there are significant differences between the class and threshold results: a) there is more silty-sand in the class realization at location (270, 10); b) sand bisects the silt in the threshold realization at location (100, 100); this not present in the class realization; c) more sand is in the class realization at location (270, 10). The differences between the realization pairs in a, b, and c are expected, because reordering the indicators changes the CDF.

paired realization set (realizations using the same random number seed) there are clear similarities in the results, but there are also significant differences (e.g. the Southern portion of Realization #32,

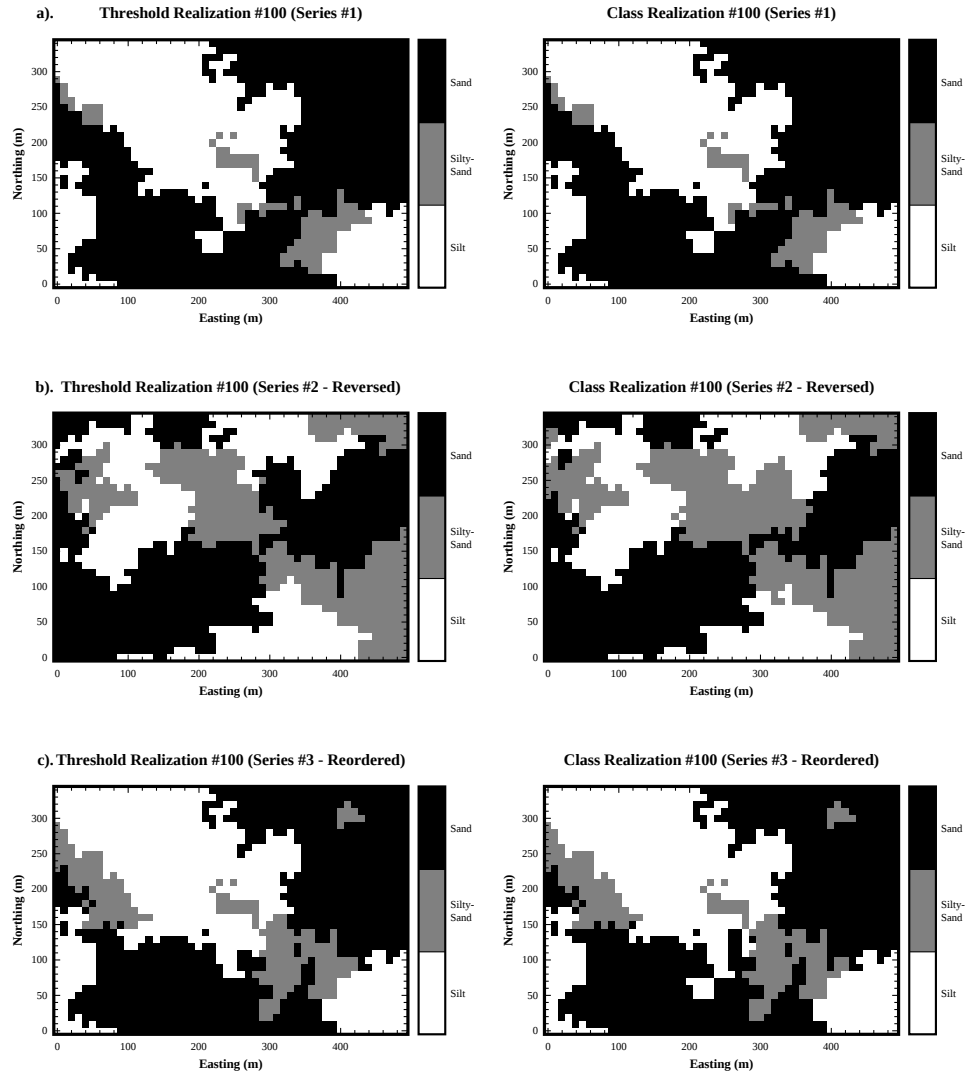


FIGURE 4-11. Realization (#100) pairs for a) original indicator ordering, b) reversed ordering, and c) arbitrary reordering for thresholds (left) and classes (right). These threshold and class realization pairs are similar. The differences between the realization pairs in a, b, and c are expected, because reordering the indicators changes the CDF.

Figure 4.10). The class realization has more silty-sand near location (270, 10) than the threshold simulation (Figure 4.10a). In other model pairs, there are only minor differences (e.g. Realization #100, Figure 4.11). The similarities occur because the same "random" path is used to generate each

model grid pair (Figure 4.4), and the same "random" number is used to select from the CDF. For most cells, the CDF is sufficiently similar that the same estimate is made for each cell. Because there are differences in how negative probabilities are generated and resolved, the CDF's are slightly different in some instances. When this occurs, a random number can be generated which results in a different estimate for the cell. From that point on, the results of the two realizations will diverge, because previously estimated cells are treated as hard data values in subsequent calculations for, as yet, unestimated cells in the simulation. Depending on the location and timing of the divergent estimate, results may be quite similar or different. When the indicators are reordered, the results change substantially (Figures 4.10b,c and 4.11b,c vs. 4.10a and 4.11a). For example, in the second set of Figures 4.11a-c, the amount and distribution of the silty-sand varies substantially. This occurs because the order of the CDF has been changed, yet the same random numbers are used. Individual realizations are not expected to be similar when the indicator order is changed; the averaged results of many realizations though, should be the same.

It is useful to compare differences between the uncertainties associated with the realization series instead of comparing individual simulations. In Figures 4.12a-f, the 200 realizations for each simulation series have been summed and averaged for each indicator, showing the probability that a particular indicator will occur at each location (red indicates areas where the indicator always occurs, and blue indicates where the indicator never occurs). If these maps are summed (Figures 4.12a + c + e, or 4.12b + d + f) every cell will equal 100 percent. The maps and histograms in Figures 4.13a-f and 4.14a-f present the distribution of the maximum probability of any indicator occurring for each approach for each cell, providing an overall measure of uncertainty. With three indicators, the minimum value is 33% (blue: all indicators equally likely to occupy cell) and the maximum value is 100% (red: at locations with hard data point). Ideally these maps would be nearly identical, signifying that, although different estimation techniques are used, the same net result is obtained. However, in this case the threshold orderings, always have a slightly higher mean probability of occurrence (Figures 4.14a, c, and e mean values versus 14b, d, and e mean values), which implies the threshold results are slightly better than the class results. The differences though are small, and as will be shown in section 4.4.2.2, threshold results are not always associated with greater certainty. In this example, class and threshold realizations vary by as much as 12% in some areas (Figure 4.15a: the largest differences are indicated in red and blue (+ and - errors), with green areas yielding nearly identical results). This is because the variation of uncertainty caused by simply reordering the indicators is of a similar magnitude for threshold simulations. This is demonstrated in the next section and illustrated in Figures 4.15b and 4.15c.

4.4.1.2: Reverse and Arbitrary Indicator Ordering

Although the initial comparison of class and threshold simulations indicate small, local areas that are significantly dissimilar, the variations are on the same scale (are simply reordered (Figures 4.15b and 4.15c). If the differences which occur from an arbitrary reordering of the indicators are no larger than those that result from using classes, it is concluded that the class and threshold techniques are essentially the same.

When the order of the indicators is simply reversed (silt = 0 2, silty-sand = 1, sand = 2 0), the differences in the threshold simulations are relatively minor, again about rily reordered (silt = 0,

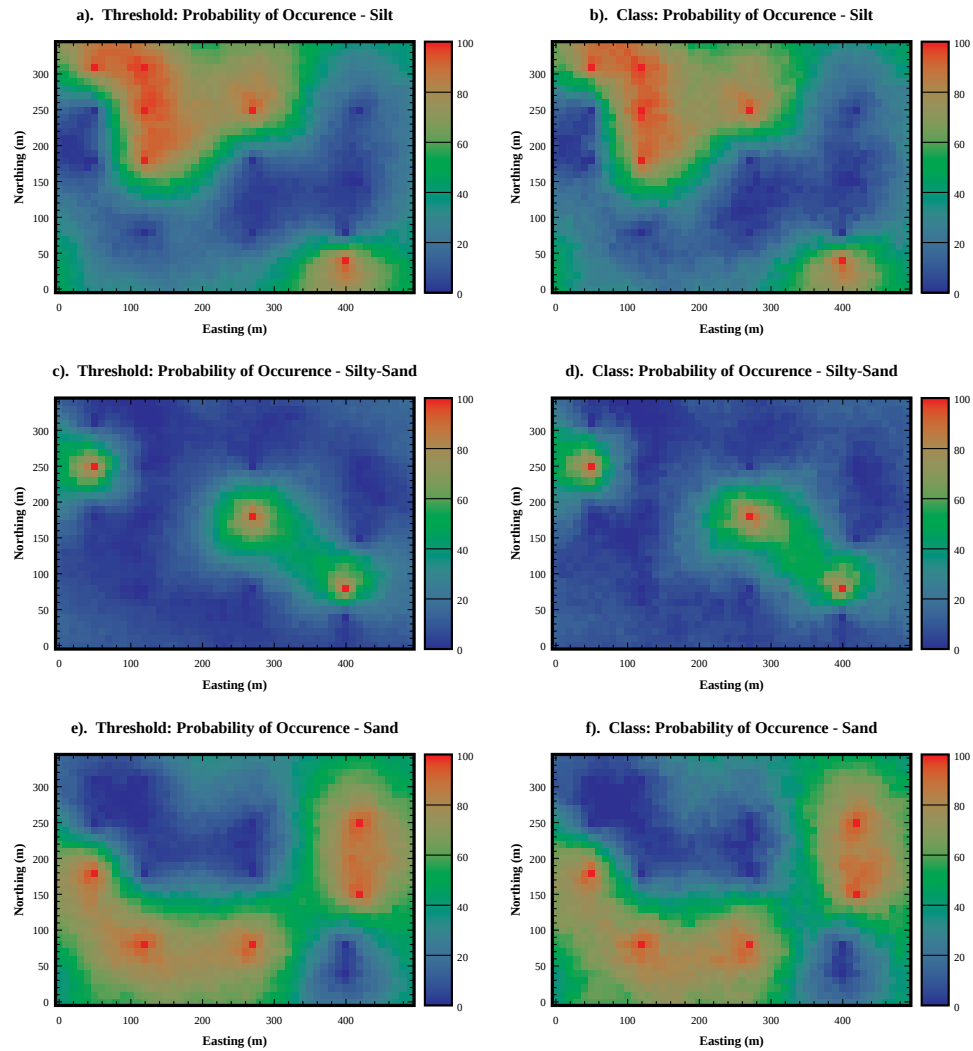


FIGURE 4-12. Maximum probability of occurrence for each indicator. At hard data locations, the indicator type is known; the probability is 0% for any other indicator type, or 100% for the specified indicator type.

silty-sand = 1 2, sand = 2 1, Figure 4.15c). Given this level of variability in realization results, due only to the order of the indicators, similar variations due to class simulation indicate that the approach is as acceptable as the threshold approach. Additional realizations (1000's) are being computed to determine if these differences are due to the size of the simulation series (200). These results though, are not yet available.

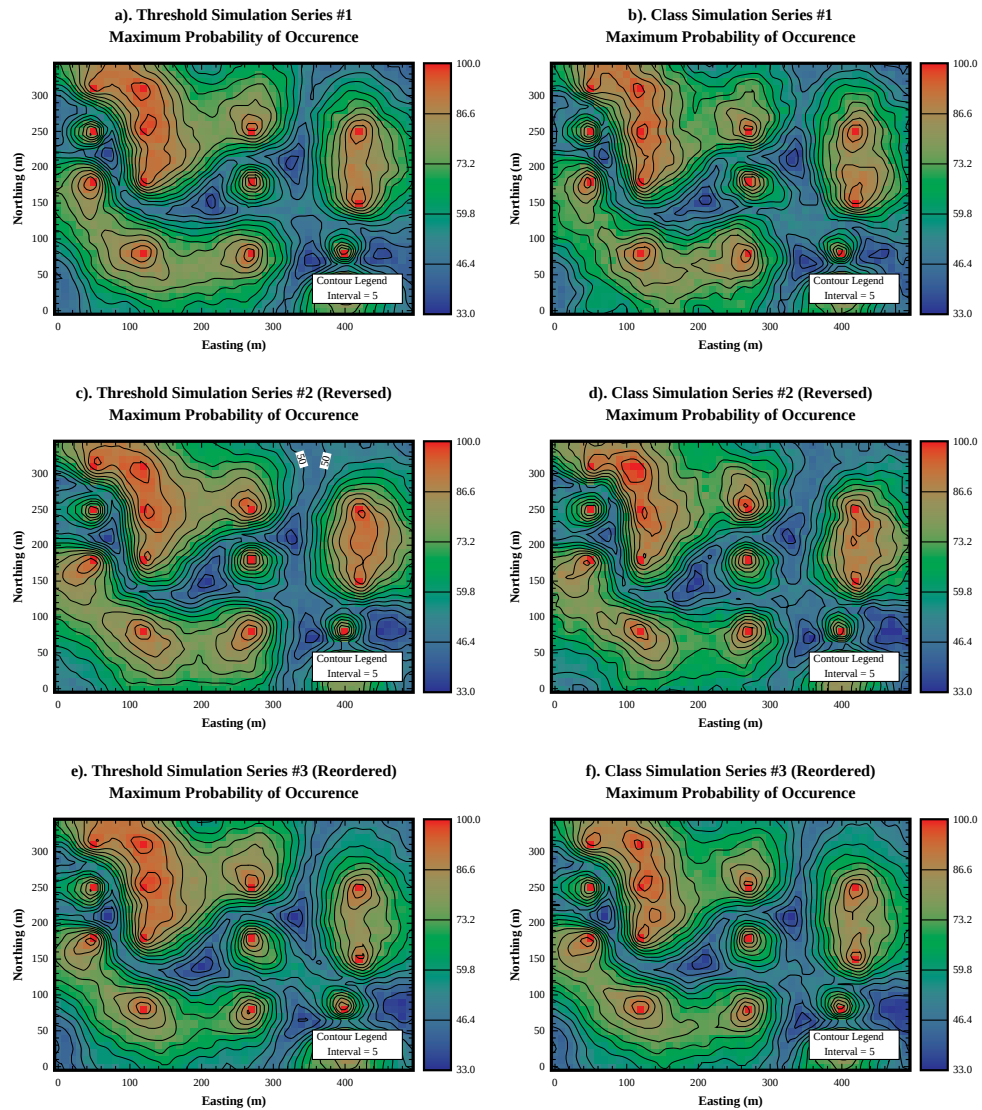


FIGURE 4-13. Maximum probability of occurrence of any indicator. At known data points, the maximum probability of being a particular indicator is 100%. The minimum probability is 33% (100% / # of indicators); at these locations each indicator is equally likely to occur. These maps are useful for evaluating the spatial distribution of uncertainty.

Variability of results for different ordering of indicators using class indicator simulation is equally consistent. Result from two reordering schemes are shown in Figures 4.15d (silt = 0 2, silty-sand =

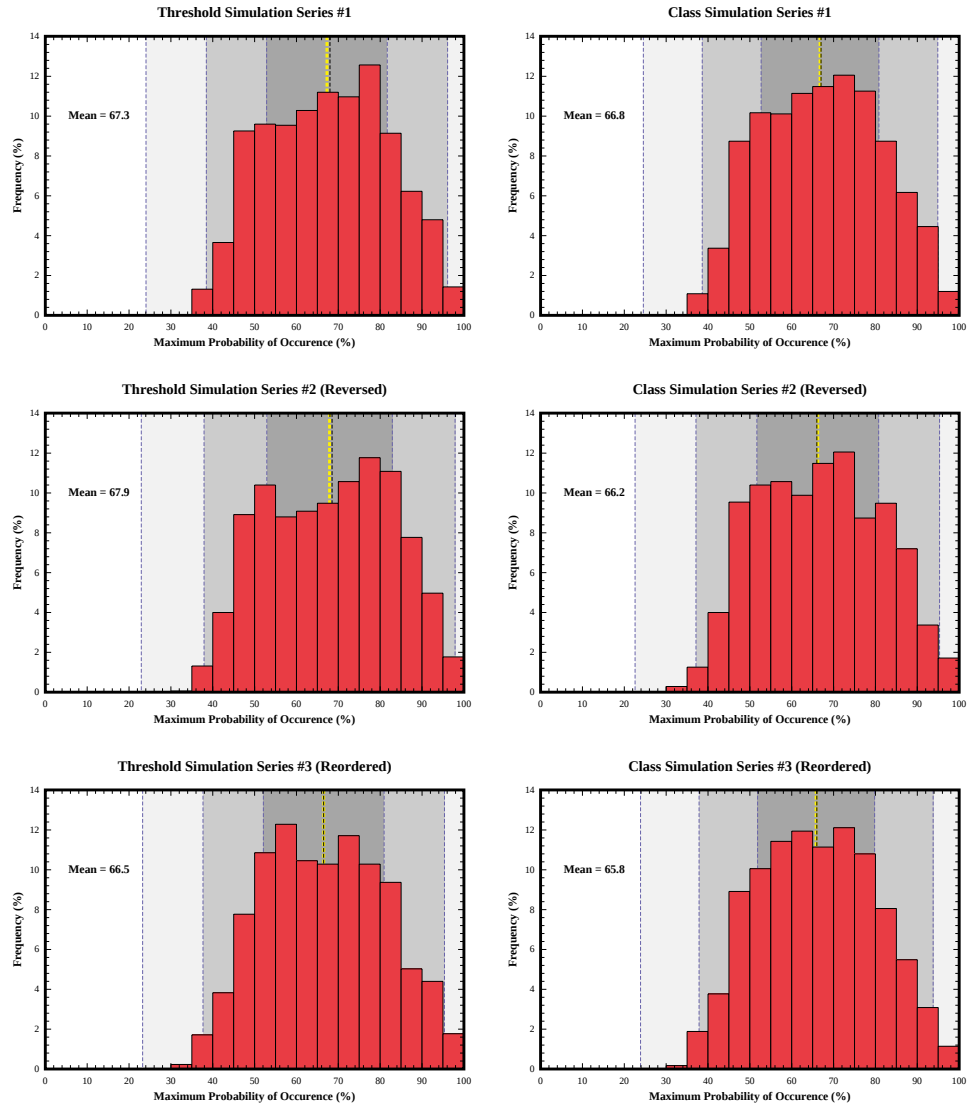


FIGURE 4-14. Frequency of maximum probabilities throughout the grid. At known data points, the maximum probability is 100%. The minimum probability is 33% ($100\% / \#$ of indicators); at locations where each indicator is equally likely to occur.

1, sand = 2 0) and Figures 4.15e (silt = 0, silty-sand = 1 2, sand = 2 1). As with the threshold realization series, the differences in the class realization series appear random and are limited to approximately

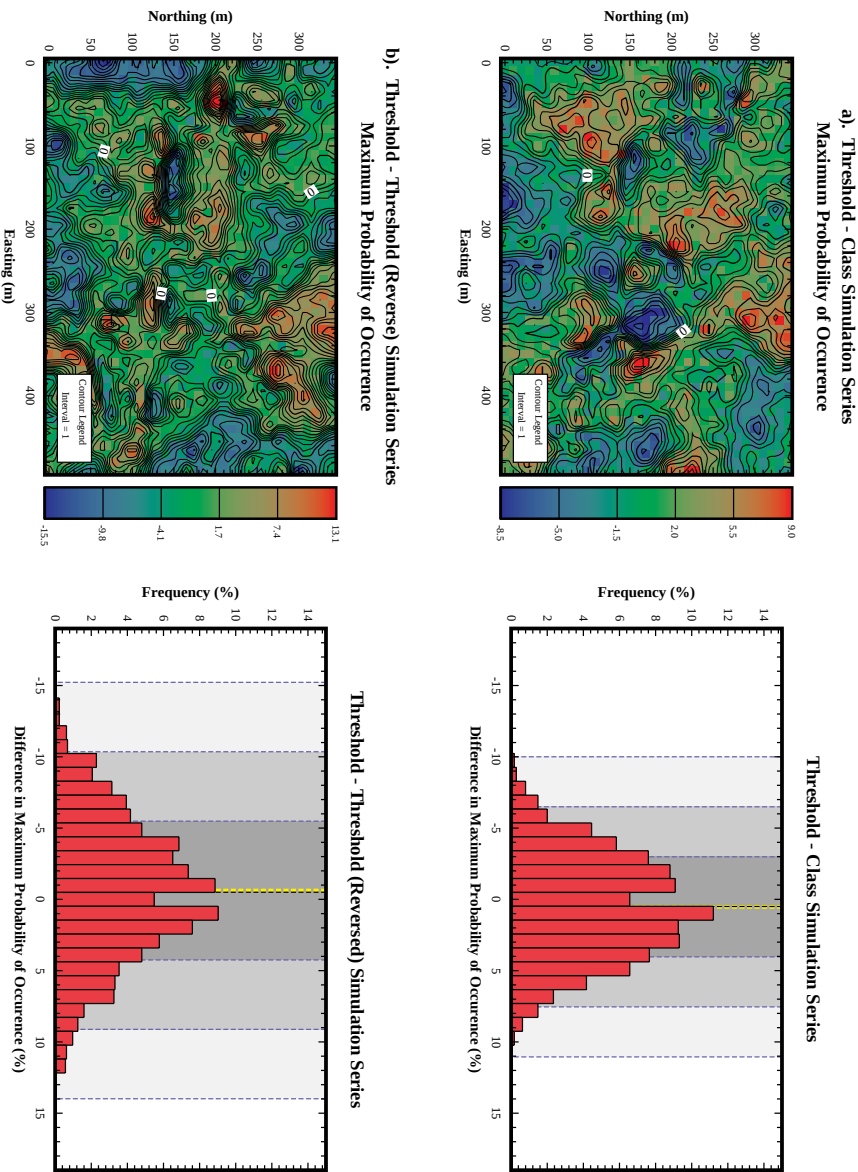


FIGURE 4-15 a,b. Difference between a) threshold and class maximum probability maps, and b) threshold and threshold (reversed) probability maps. These maps show areas, where the different series return different averaged results. Differences are largest, although of opposite sign, in red and blue areas. Differences are smallest in green areas. The histograms are useful for identifying the magnitude and distribution of the differences.

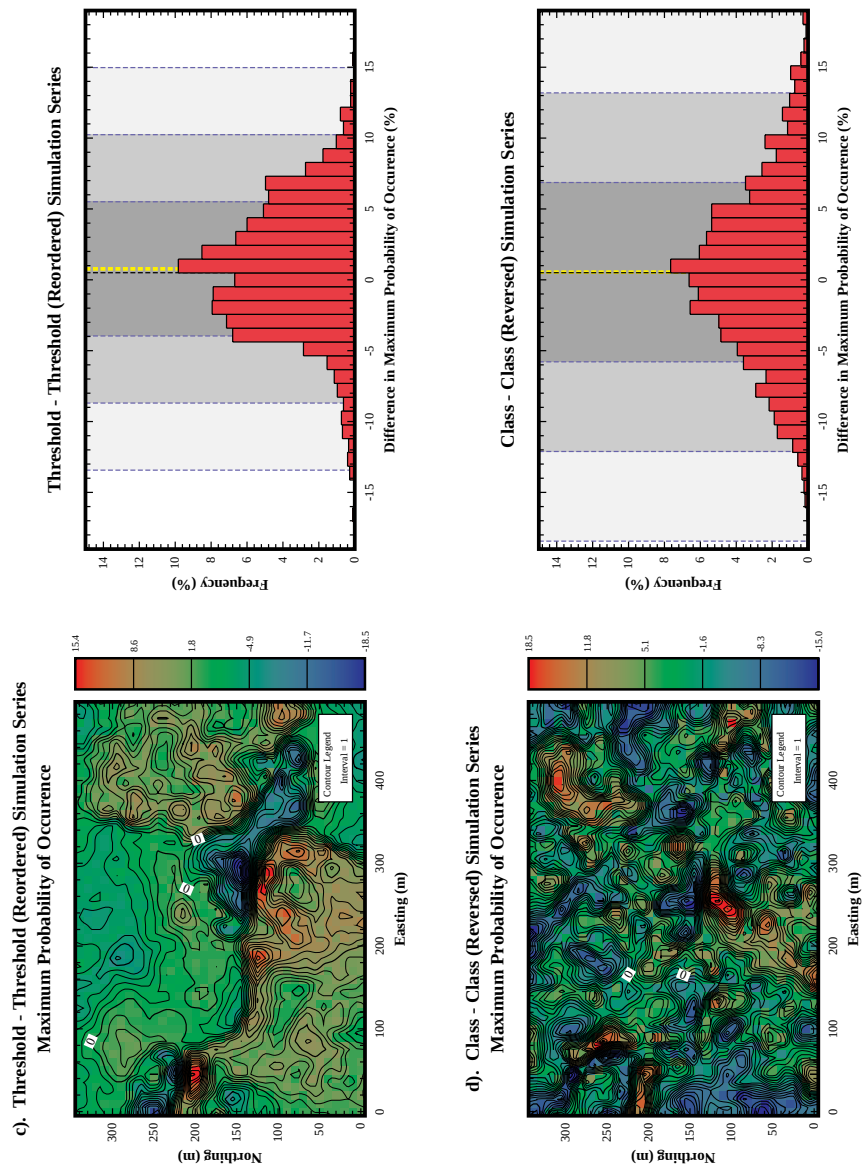


FIGURE 4-15 c,d. Difference between a) threshold and threshold (reordered) maximum probability maps, and b) class and class (reversed) probability maps. These maps show areas, where the different series return different averaged results. The histograms are useful for identifying the magnitude and distribution of the differences.

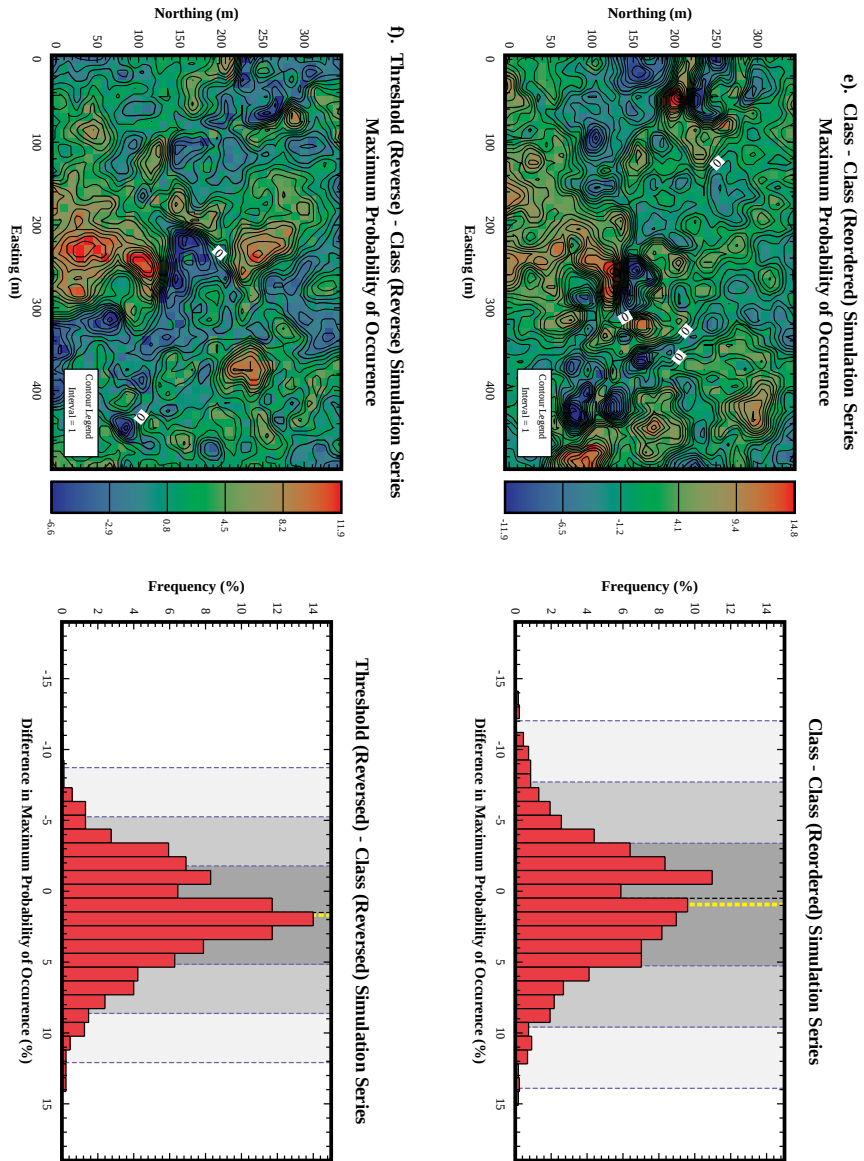


FIGURE 4-15 e, f. Difference between a) class and class (reordered) maximum probability maps, and b) threshold (reversed) and class (reversed) probability maps. These maps show areas, where the different series return different averaged results. The histograms are useful for identifying the magnitude and distribution of the differences.

For each ordering scheme, the differences between the threshold and class simulations (Figures 4.15a, 4.15f, and 4.15g) are less than the differences between different ordering schemes when the same method was used. Given these similarities, and knowing that differences result in differences in managing ORV's, it is concluded that class and threshold simulation generate equivalent results.

4.4.1.2.1: Simulation Differences Due to Indicator Order

One of the initial motivations for using classes was to eliminate the differences in simulations resulting due to indicator ordering. As seen in the examples above, the class simulations have a similar problem. The probability for each class indicator is calculated independently, therefore order should not make a difference, yet it does. If the differences are not due to computer round-off, there should be differences in the kriging matrices or the kriging weights, however cells with different results were identified and compared, and this was not the case. It is possible that some of the differences due to indicator ordering are associated with the random number generator but this is difficult to demonstrate or prove. The random number generator used in this program was evaluated for a group of 10,000 and 100,000 random numbers (Figure 4.16a,b,c), and no serious, or obvious bias was found, but all numbers are not equally sampled. These differences could explain some of the differences in the simulation results, because when the indicators are reordered, preferences to different "random" number ranges could cause a bias. It is thought that the source of the problem, is the management of the ORV's.

4.4.1.2.2: Simulation Differences Due to Order Relation Violations

Realization #32 (Figure 4.17 (these are the coarse pre-grids for the realizations in Figure 4.10a)) is used to demonstrate the difficulties that arise, due to ORV's. The first violation occurred at cell ((25, 2) (490m, 30m)) during the simulation of the coarse grid (realizations are calculated in two passes; a coarse grid is simulated first, then it is used to condition the fine grid). Calculations for this cell were based on 37 values (including original sample points and prior estimated grid cells) (Table 4.1). Because the same semivariogram models were used for all class and threshold levels, the class and threshold kriging matrices were identical, therefore the kriging weights were identical. The following uncorrected CDF (threshold) and PDF (class) values were calculated:

CDF		PDF	
Threshold	F(Z≤thr.)	Threshold	F(Z=class)
0.5	0.563	1.0	0.563
1.5	1.089	2.0	0.526
		3.0	0.000

Both the threshold ($1.089 > 1.0$) and class ($0.563 + 0.526 = 1.089 > 1.0$) probabilities needed to be rescaled to 1.0. The threshold method truncates CDF values greater than 1.0 to 1.0, and the class

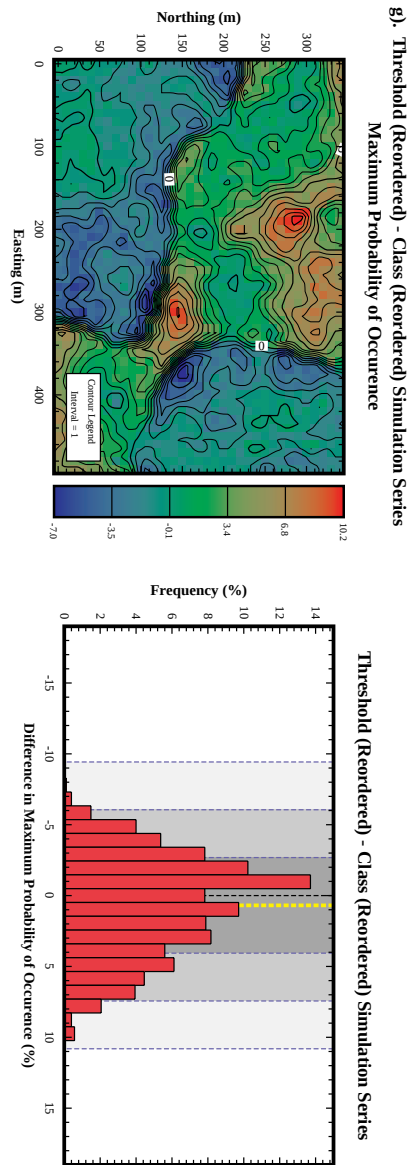


FIGURE 4.15 g. Difference between threshold (reordered) and class (reordered) maximum probability maps. These maps show areas, where the different series return different averaged results. The histograms are useful for identifying the magnitude and distribution of the differences.

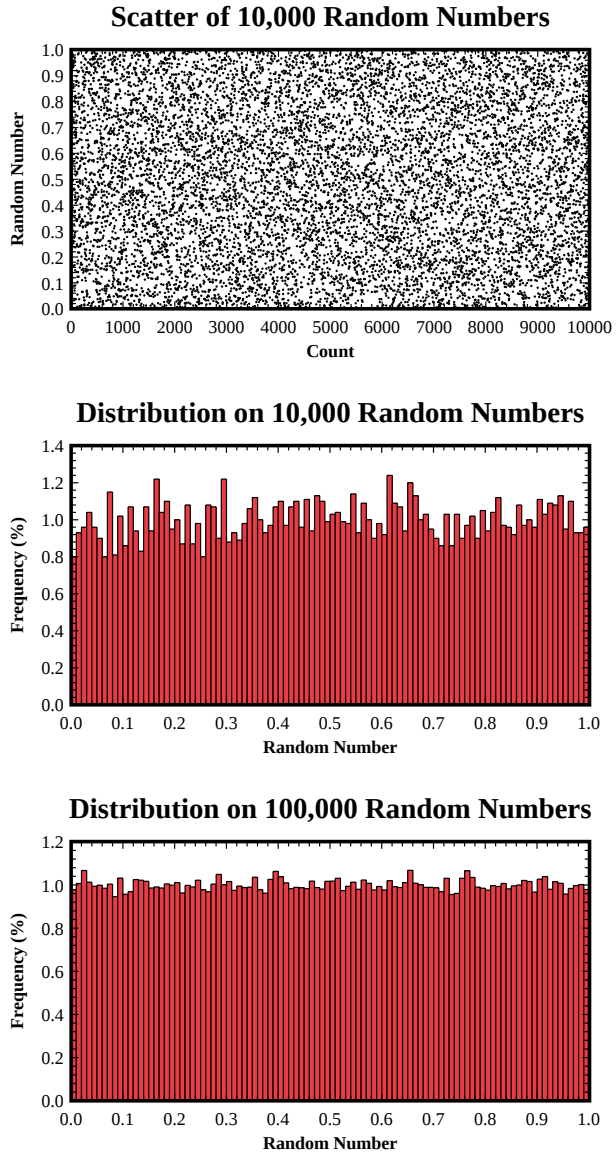


FIGURE 4-16. Test of random number generator: a) scatter of 10,000 sequential random numbers, b) frequency distribution of 10,000 random numbers (equal distribution would put 1% in each class), and c) frequency distribution of 100,000 random numbers (equal distribution would put 1% in each class).

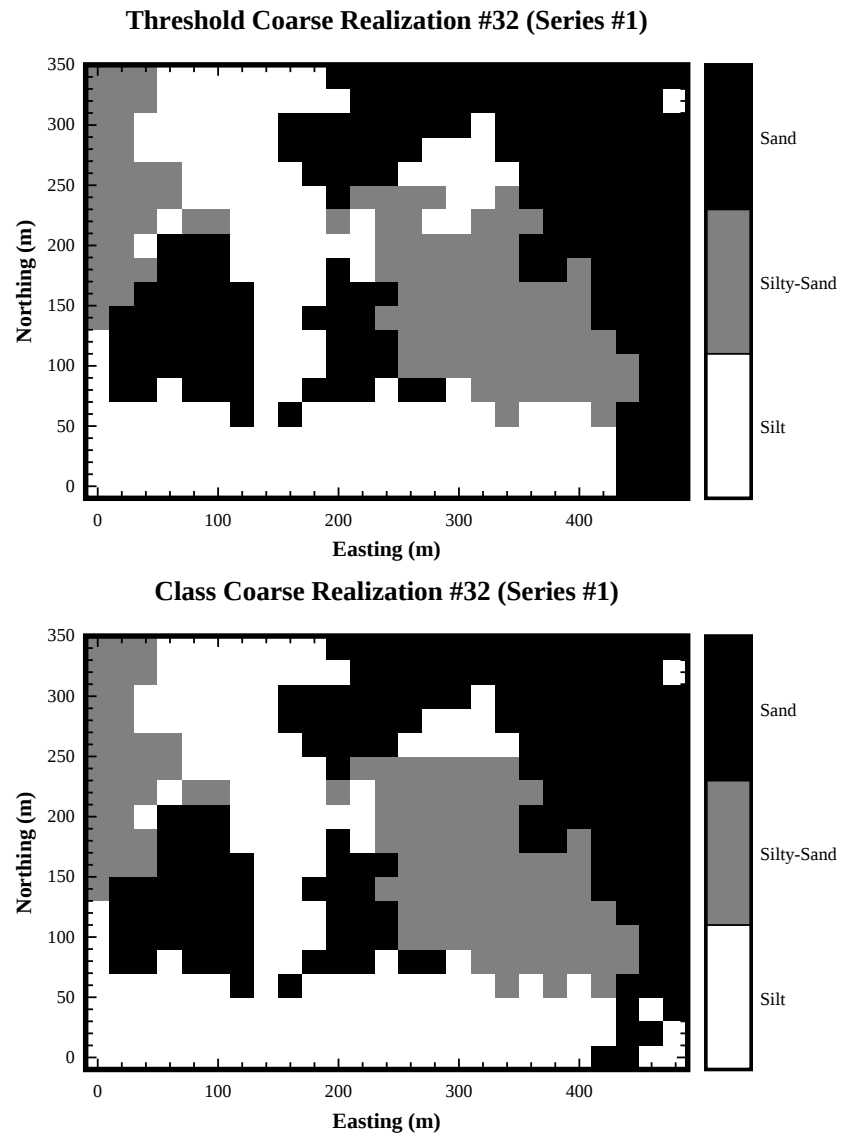


FIGURE 4-17. Coarse grid realization (#32) pair for the original indicator ordering. These grids, are slightly different, due to an ORV occurring at (490m, 30m). Because of the ORV, the CDF's varied between the two methods at this cell, and a random number in each realization selected different material types for the cell. From this point forward, the prior sample data, and prior evaluated cells varied.

method scales all the PDF terms so they will sum to 1.0. The corrected distributions are:

CDF		PDF	
Threshold	$F(Z \leq \text{thr.})$	Threshold	$F(Z = \text{class})$
0.5	0.563	1.0	0.517
1.5	1.000	2.0	0.483
		3.0	1.000

and the final class PDF is converted to a CDF:

CDF		CDF	
Threshold	$F(Z \leq \text{thr.})$	Threshold	$F(Z = \text{class})$
0.5	0.563	1.0	0.517
1.5	1.000	2.0	1.000
		3.0	1.000

The probabilities are no longer the same, and in this realization, a random number of 0.547 was generated to select the indicator class. As a result, for the threshold realization, the cell was defined as silt (indicator #1; $0.547 < 0.563$). For the class realization, the cell was defined as silty-sand (indicator #2; $0.547 > 0.517$).

Although the methods initially agreed exactly on the probability of occurrence for indicators #1 and #2, the different procedures for correcting the ORV's, resulted in different CDF's. As a consequence, the results for this grid cell pair, and those following, diverged. With different estimates for this cell, future class and threshold calculations using this cell as a conditioning point would generate different CDF's even without further ORV's.

4.4.2: Colorado School of Mines Survey Field

At the CSM survey field, located on the west edge of Golden, Colorado (Figure 4.18), hard and soft data were collected to investigate the use of soft data for reducing uncertainty associated with groundwater flow models. The site data include core and chip samples; borehole geophysical logs; and eight (Figure 4.19), two-dimensional, cross-hole, tomographic sections. A sub-region of the data set is used to demonstrate class vs. threshold indicator simulation.

This data set is used to compare the class and threshold simulation techniques, using, not only hard data values, but also three different types of soft data. Because of differences in the semivariogram models, and management of soft data, the simulations generated using threshold and class approaches were not identical. In this case, the class realizations have a slightly higher average certainty level, though uncertainty at individual cells can be substantially higher or lower than the threshold models in the same cell. Some of the differences may be due the random number generator, but most likely they result from differences in managing ORV's.

Coarse Grid Position			Indicator ID		Kriging
X	Y	Z	Threshold	Class	Weight
21	3	1	111	100	0.328
21	5	1	011	010	0.437
20	2	1	111	100	0.117
18	2	1	111	100	0.128
23	4	1	001	001	-0.051
21	7	1	011	010	0.048
17	7	1	011	010	0.104
22	8	1	001	001	-0.015
24	6	1	001	001	-0.018
24	7	1	001	001	-0.014
17	8	1	011	010	-0.007
15	3	1	111	100	0.044
15	2	1	111	100	-0.030
14	5	1	001	001	0.010
14	8	1	011	010	-0.023
23	11	1	001	001	-0.045
14	9	1	011	010	-0.033
24	11	1	001	001	-0.011
22	12	1	001	001	-0.004
14	10	1	011	010	0.000
24	12	1	001	001	0.041
11	3	1	111	100	-0.012
11	7	1	001	001	-0.005
11	1	1	111	100	-0.002
11	8	1	001	001	0.003
22	14	1	001	001	0.004
11	9	1	001	001	0.017
16	14	1	111	100	-0.001
14	14	1	111	100	0.004
7	5	1	001	001	0.001
7	10	1	111	100	-0.006
5	2	1	111	100	0.000
7	14	1	111	100	-0.001
4	10	1	001	001	0.000
3	2	1	111	100	-0.002
7	17	1	111	100	-0.002
1	2	1	111	100	-0.002
Silt			111	100	
Silty-Sand			011	010	
Sand			001	001	

TABLE 4.1. Kriging weight and nearest neighbors for both class and threshold realization #32. The indicators at each point, and the kriging weights are identical.

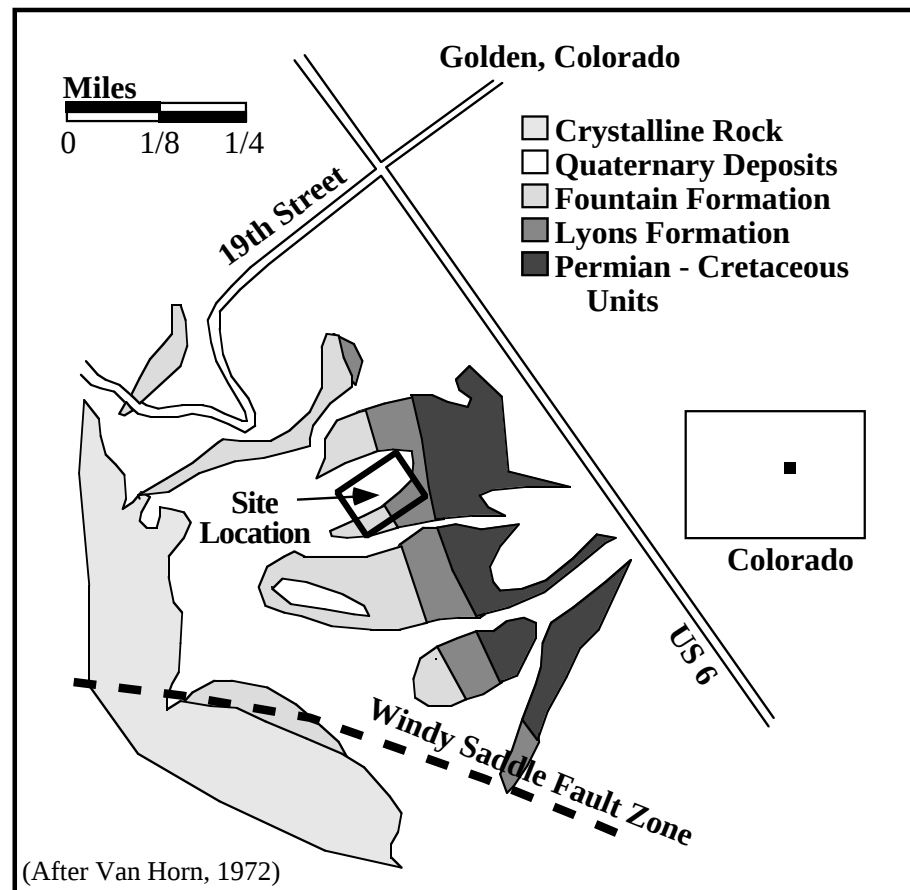


FIGURE 4-18. CSM Survey Field location map.

4.4.2.1: Model Definition and Simplifying Assumptions

The model and grid dimensions are based on the same data and indicator classes as described by McKenna and Poeter (1994) for the CSM survey field (Figure 4.20). However, only a small sub-grid was used for this demonstration. The grid was dimensioned 80 columns equally spaced between 2,045 feet - 2,450 feet in the X direction, 60 rows equally spaced between 4,228 feet - 4,533 feet in the Y direction, and 2 layers, each two feet thick between 5,917 feet - 5,921 feet in the Z direction. This same grid was used for both the threshold and class simulation.

The eight indicator classes are based on seismic velocities of different materials at the site (Table 4.2). The indicators were selected after thorough analysis which concluded that the seismic

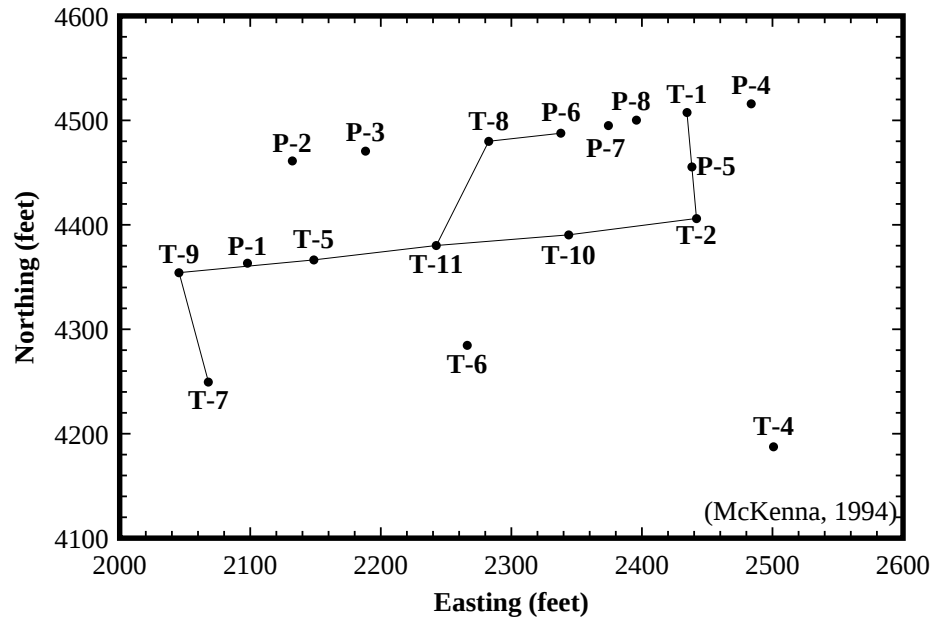


FIGURE 4-19. CSM Survey Field site map. Dots represent borehole locations. Solid lines identify location of tomography surveys.

Indicator	Material Description
1	Conglomerate (Lyons Formation)
2	Fine to coarse sandstone with conglomerate lenses
3	Fluvial sandstone (Lyons Formation)
4	Very fine to very coarse sandstone
5	No core recovered. Moderately consolidated with low-moderate clay
6	Two materials: 1) silty sandstone, and 2) poorly sorted sandstone with siltstone and conglomerate lenses
7	No core recovered. Poorly consolidated, low clay material
8	No core recovered. Well fractured area of any material type.

TABLE 4.2. Indicator and associated material type.

velocities reflected differences in hydrologic flow properties. Several distinctly different lithologies were grouped together because they displayed similar hydraulic properties and spatial correlation's. The indicator classes are defined in Table 4.3. Initial estimates of hydraulic conductivity values were assigned to each indicator based on either material type, permeability measurements, or

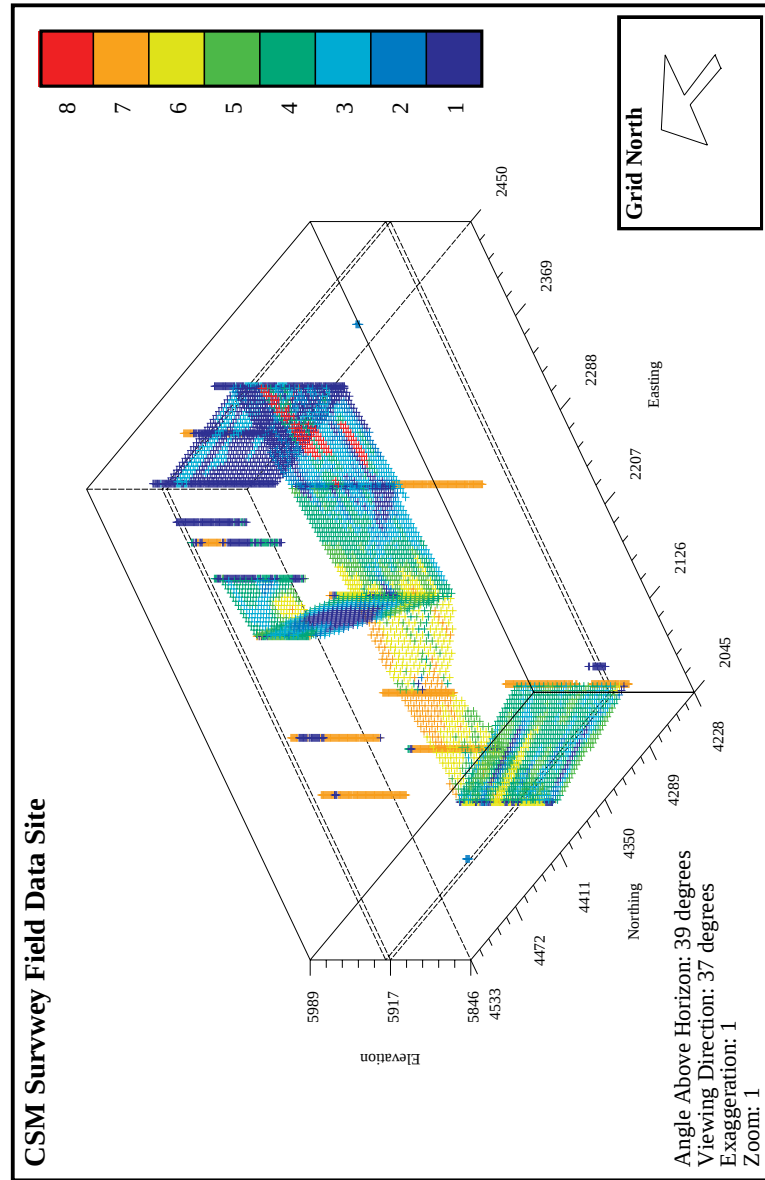


FIGURE 4-20. CSM Surveyway Field borehole and tomography data converted to indicators. The dotted layer delineates the extents of the model grid.

Indicator	Sonic Velocity (ft/sec)	Hydraulic Conductivity (ft/day)
1	> 10870	0.0011
2	10000 - 10870	0.0011
3	9050 - 10000	0.0025
4	8550 - 9050	0.0043
5	8050 - 8550	0.040
6	7250 - 8050	0.0043
7	6060 - 7250	0.40
8	< 6060	7.8

TABLE 4.3. The indicator classification is based on sonic velocity measurements, and are matched to approximate hydraulic conductivity's.

estimated material sonic velocities. Later, inverse flow modeling was used to improve the hydraulic conductivity estimates. The sonic-velocities were used as Type-A data as described in Tables 4.4

Threshold

Threshold Velocity (ft/sec)	Threshold	P ₁	P ₂	P ₁ - P ₂
6060	7	0.00	0.00	0.00
7250	6	0.56	0.04	0.52
8050	5	0.58	0.05	0.53
8550	4	0.63	0.10	0.63
9050	3	0.84	0.17	0.67
10000	2	0.90	0.17	0.73
10870	1	0.91	0.15	0.74

6060 velocity: No hard data to calibrate against. Found only in tomography cross sections.

TABLE 4.4. Threshold p_1 - p_2 estimates.

and 4.5. The p_1 - p_2 values are significantly lower for the class method, because the class method is more restrictive.

The frequency distribution of indicators was based on the same sample information (Tables 4.6 and 4.7). The soft data prior distributions are based only on hard and the Type-A data. Type-B and C data were available, but assigning them to an individual class or threshold is not possible.

In order to make realizations for the class and threshold simulation as similar as possible, the same data distributions and grid were used. It was not possible to use the same semivariogram models

Class

Velocity Range (ft/sec)	Class	P ₁	P ₂	P ₁ - P ₂
< 6060	8	0.00	0.00	0.00
6060 - 7250	7	0.56	0.00	0.56
7250 - 8050	6	0.39	0.04	0.35
8050 - 8550	5	0.25	0.12	0.13
8550 - 9050	4	0.45	0.18	0.27
9050 - 10000	3	0.26	0.13	0.13
10000 - 10870	2	0.38	0.05	0.33
> 10870	1	0.84	0.00	0.84

6060 velocity: No hard data to calibrate against. Found only in tomography cross sections.

TABLE 4.5. Class p_1 - p_2 estimates.

Threshold	Cumulative Probability		
	Hard	Soft	Difference
1.5	0.1638	0.2306	0.0668
2.5	0.2370	0.3120	0.0750
3.5	0.2706	0.5031	0.2325
4.5	0.3693	0.7306	0.3613
5.5	0.3886	0.8491	0.4605
6.5	0.5066	0.9429	0.4363
7.5	1.0000	0.9748	0.0252

TABLE 4.6. Threshold prior hard and prior soft (Type-A) data distributions. The large difference in threshold 3.5 propagates through threshold 6.5.

(Tables 4.8 and 4.9) in both sets of simulations though, because the class and threshold methods calculate the semivariogram models on different portions of the data set. The first and last semivariogram models will always be identical, but there can be significant differences in the intermediate models. For example, the maximum range for threshold 3.5 is 282 feet in the East-West direction and 174 feet in the North-South direction. For classes three and four, the respective ranges are much shorter (111, 77 feet and 117 feet respectively). It is thought that most of the differences in the simulation results are due to the differences in the semivariogram models.

Class	Individual Probability		
	Hard	Soft	Difference
1	0.1638	0.2306	0.0668
2	0.0732	0.0814	0.0082
3	0.0336	0.1911	0.1575
4	0.0987	0.2275	0.1288
5	0.0193	0.1184	0.0991
6	0.1180	0.0938	0.0242
7	0.4934	0.0319	0.4615
8	0.0000	0.0252	0.0252

TABLE 4.7. Class prior hard and prior soft (Type-A) data distributions.

Threshold	East-West		North-South		Vertical		Nugget
	Range	Sill	Range	Sill	Range	Sill	
1.5	81.0	0.118	155.6	0.118	81.0	0.0623	0.0516
2.5	93.0	0.207	126.0	0.207	54.0	0.0061	0.0
3.5	75.0	0.169	174.0	0.246	21.0	0.0468	0.0
	282.0	0.0783					
4.5	90.0	0.0831	15.0	0.149	3.0	0.0405	0.0
	204.0	0.145	99.0	0.0765	47.0	0.0358	
5.5	132.0	0.189	12.0	0.158	36.0	0.0692	0.0
			48.0	0.0308			
6.5	78.0	0.129	15.0	0.129	93.0	0.0284	0.0
7.5	61.5	0.0208	18.5	0.0208	36.9	0.0079	0.0

NOTE: Multi-nested models require two rows.

TABLE 4.8. Threshold semivariogram models.

4.4.2.2: Geostatistical Realizations and Results

A total of 100 realizations were calculated for both the class and threshold models. In this section, several realization pairs are described, the probability that any individual indicator will occur is defined, then the difference between the class and threshold realizations and the maximum probability that any indicator will occur in each cell are calculated.

Examining the paired realizations from each set, it is clear that the same site is being simulated, but there are subtle, yet distinct differences in the realizations (Figures 4.21 and 4.22). The general

Class	East-West		North-South		Vertical		Nugget
	Range	Sill	Range	Sill	Range	Sill	
1	81.0	0.118	155.6	0.118	81.0	0.0623	0.0516
2	64.5	0.0720	20.3	0.0720	41.0	0.0374	0.0
3	110.7	0.0878	43.1	0.0397	92.3	0.0878	0.0447
			116.9	0.0480			
4	76.9	0.0880	116.9	0.0880	59.0	0.0880	0.0709
5	104.6	0.0878	64.6	0.0878	27.7	0.0658	0.0
6	150.7	0.0910	31.0	0.0910	31.0	0.0910	0.0
7	61.5	0.116	15.4	0.116	49.2	0.0569	0.0
8	61.5	0.0208	18.5	0.0208	36.9	0.0079	0.0

NOTE: Multi-nested models require two rows.

TABLE 4.9. Class semivariogram models.

distribution of indicators is similar, but the threshold realizations have more scatter, and produce smaller regions of indicator #8 (examine grids near (2380, 4420): indicator #8 is red). The reason for the differences are based on three factors: 1) differences in the semivariograms, 2) differences in the p_1 - p_2 values, and 3) differences in applying the prior hard minus prior soft prior probabilities. In the individual realization pairs, the indicators in the threshold realizations tend to be slightly less continuous. It appears this is caused by the large differences between the hard and soft prior distributions and the ORV's.

In addition to the greater randomness in the threshold realizations, there is also larger uncertainty associated with the spatial distribution of indicators. Several figures were prepared to illustrate the differences in the results of the threshold and class simulations and to show that smaller uncertainties are associated with the class simulations. Figure 4.23 illustrates the probability a particular indicator will occur in each cell for both the threshold (left) and class (right) simulation series. The highest certainty levels coincide with the hard and soft data locations (red = 100%, blue = 0%). Figure 4.24 shows the difference between probability of occurrence for the class and threshold simulation series for each indicator. The largest differences (red: class probability >> threshold probability; green: class probability $\cong$ threshold probability; blue: class probability << threshold probability) occur where the model has the most data. By examining the kriging matrix results, and CDF development in these areas for several cells, the differences are largely due to differences in how the prior data probabilities modify the CDF and how ORV's are handled between the class and threshold techniques. The differences in uncertainty also tend to be small away from the control data, because both methods are very uncertain as to what is occurring in those areas (a small number minus a small number equals a small number). In areas of the model grid with little or no data, the averaged results are more similar (green: differences $\cong$ 0.0), but there is more uncertainty.

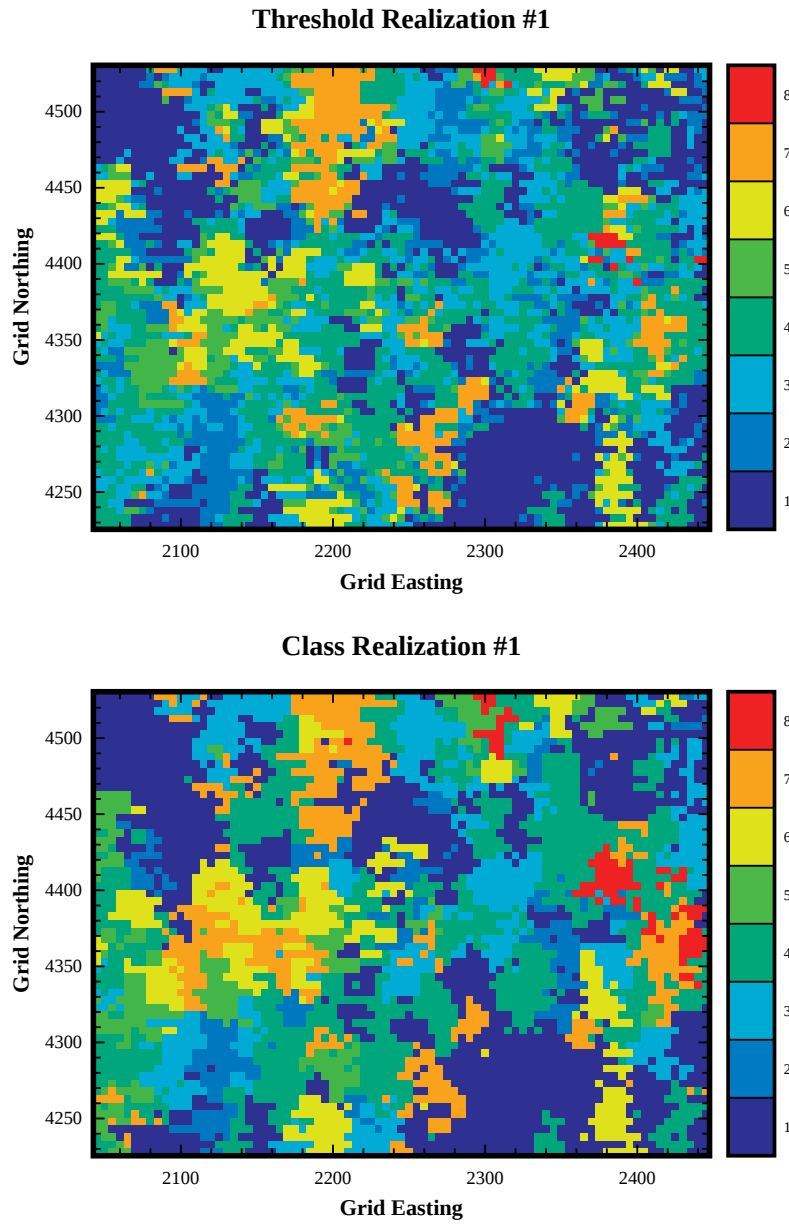


FIGURE 4-21. Individual threshold and class realization #1.

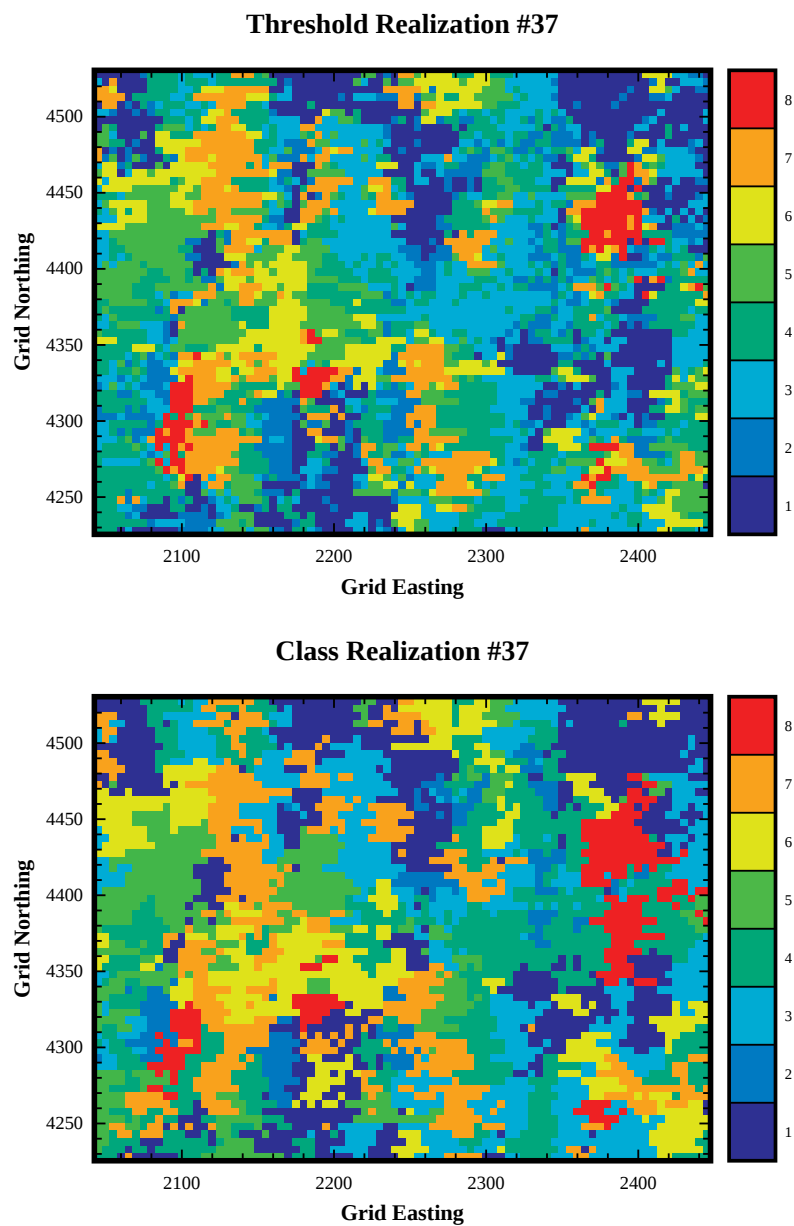


FIGURE 4-22. Individual threshold and class realization #37.

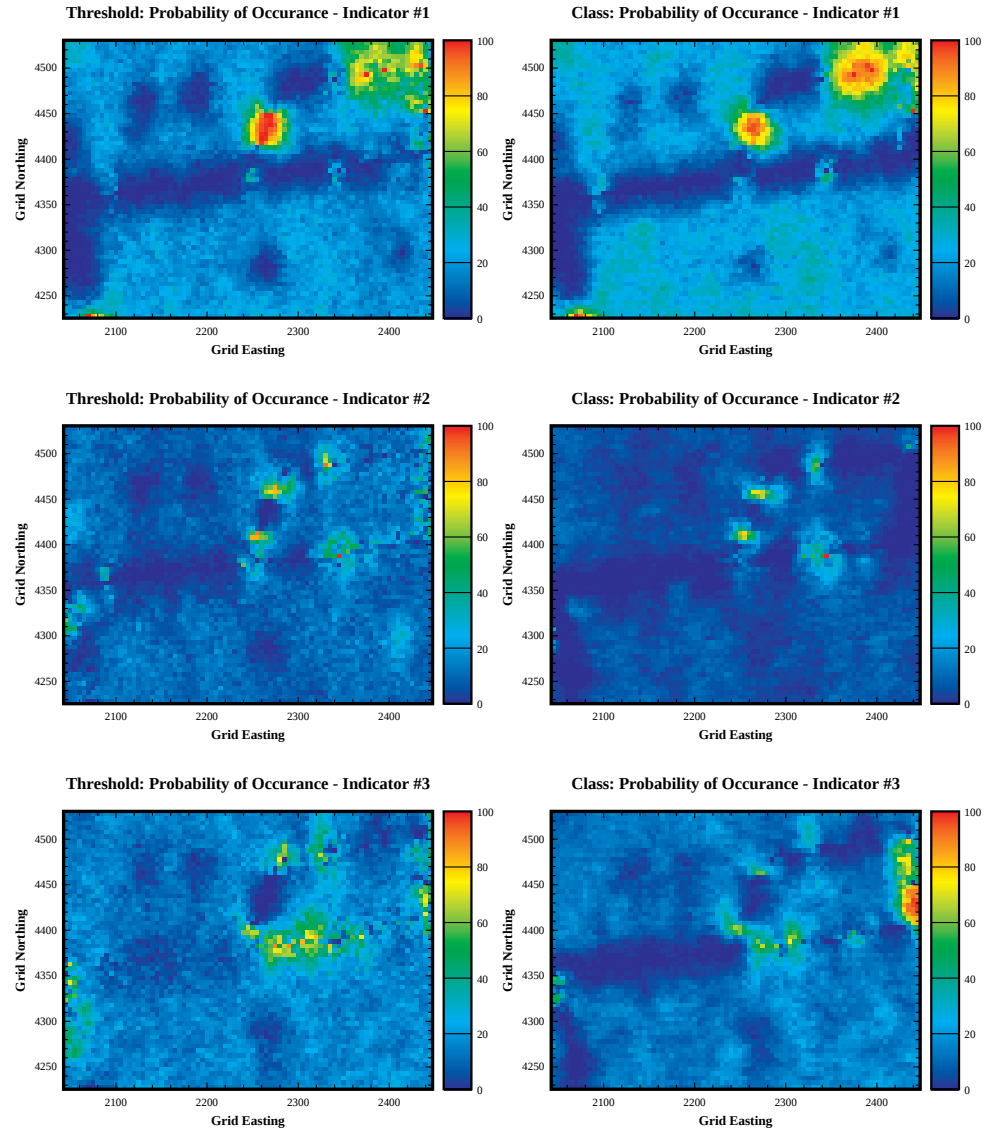


FIGURE 4-23 a,b,c. Maximum probability of occurrence for threshold and class indicators #1, #2, and #3. At known hard data points, probability is 0% for any other indicator type, or 100% for the specified indicator type.

The maximum probability that any indicator will occur in each cell is shown in Figure 4.25 (red implies more certainty, up to 100% at hard data locations; blue less, as little as 12.5% (100% / # indicators)) with associated histograms presented in Figure 4.26. From the histograms, it can be

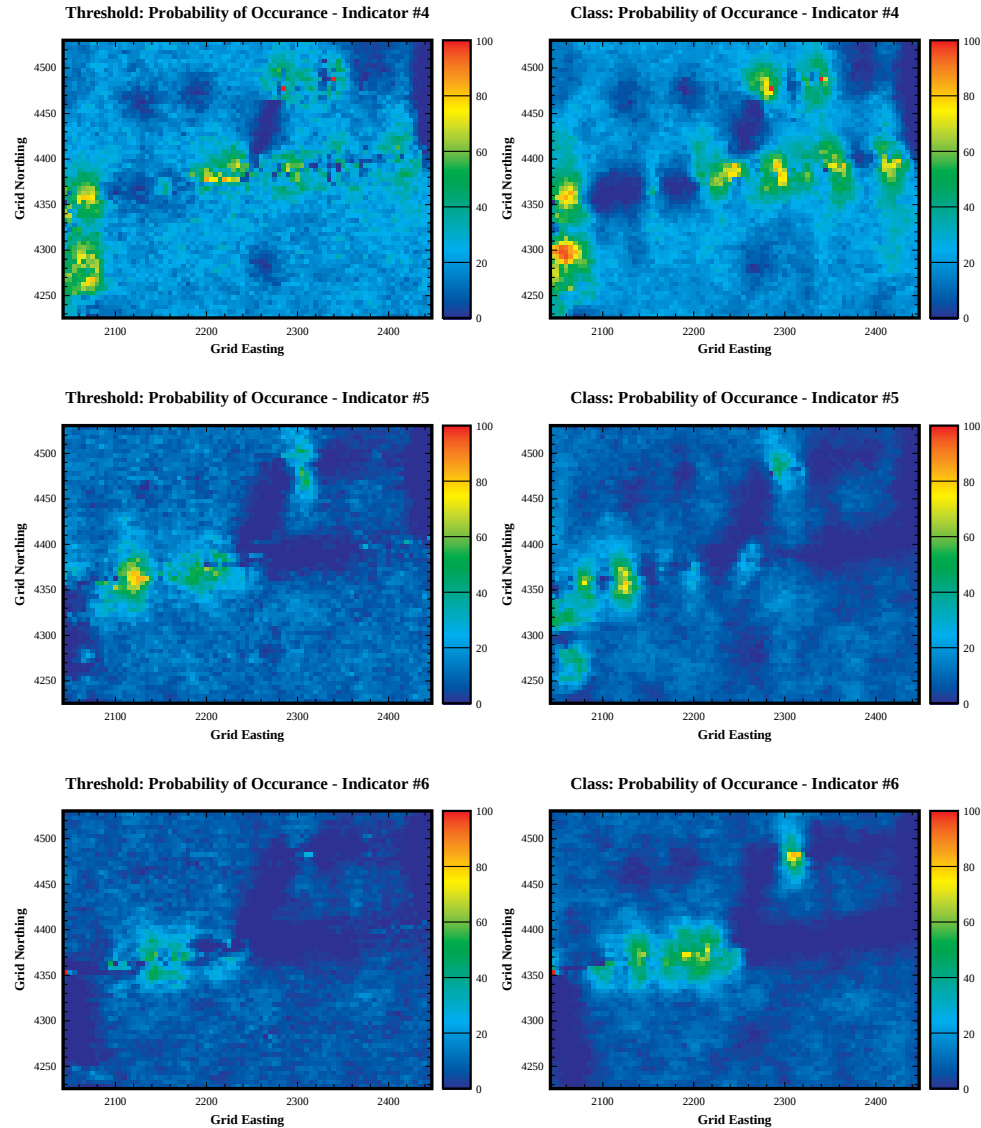


FIGURE 4-23 d,e,f. Maximum probability of occurrence for threshold and class indicators #4, #5, and #6. At known hard data points, probability is 0% for any other indicator type, or 100% for the specified indicator type.

seen that there is slightly less uncertainty in the class realizations (class mean maximum probability (certainty level) is 37% compared the 33% for thresholds) and the differences in uncertainty are presented in Figure 4.27a. Positive differences (green to red) show areas where the class

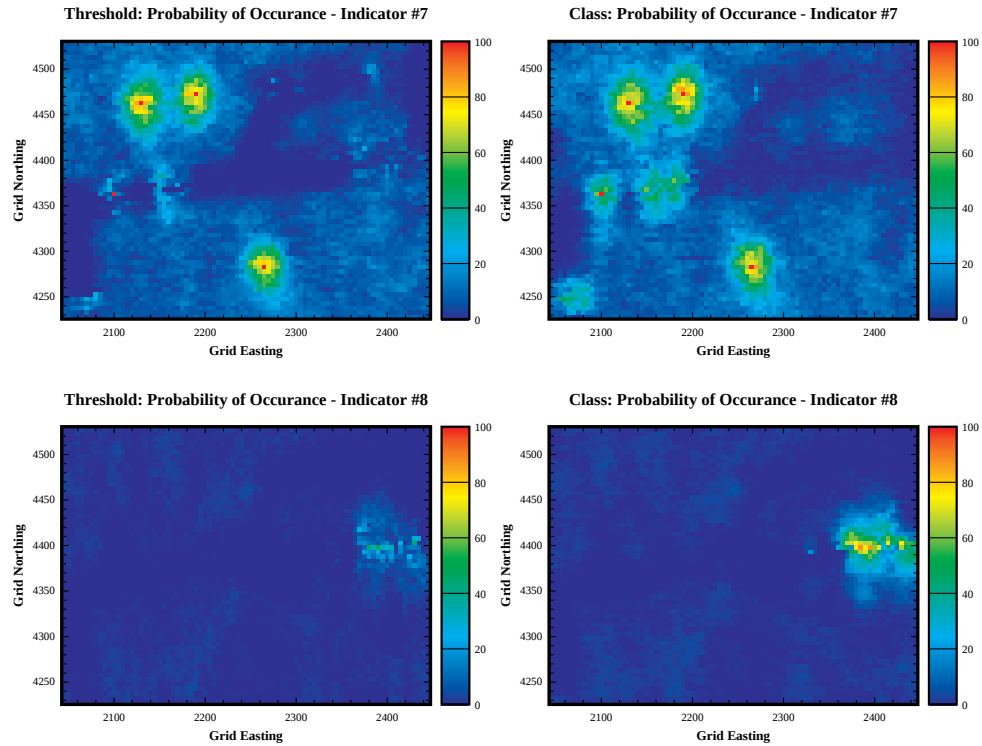


FIGURE 4-23 g,h. Maximum probability of occurrence for threshold and class indicators #7 and #8. At known hard data points, probability is 0% for any other indicator type, or 100% for the specified indicator type.

realizations are less uncertain than the threshold realizations; negative differences (green to blue) show areas where the class realizations are more uncertain than the threshold realizations. These maps are useful for defining the uncertainty in the model, but are not useful for analyzing the distribution of a particular indicator. Again the largest differences are near areas of hard and soft data (red: class certainty $\gg$ threshold certainty; green: class certainty $\cong$ threshold certainty; blue: class certainty $\ll$ threshold certainty), and this relates to the problems associated with the difference in how ORV's are managed. For this site, the class realizations show less uncertainty than the threshold realizations. On average the mean uncertainty reduction is 3.9% with a standard deviation of 9.3% (Figure 4.27b). Based on this model alone it is premature to suggest that the class method may help reduce model uncertainty. In the synthetic models described in section 4.4.1.1, the threshold realizations had slightly smaller uncertainties than the class realizations.

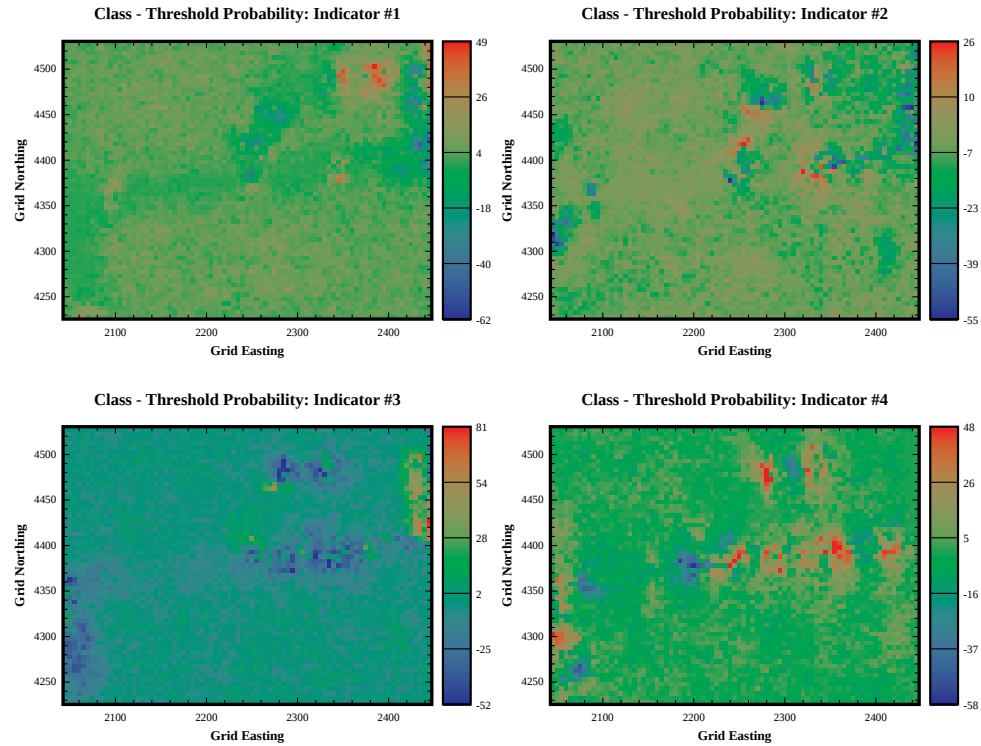


FIGURE 4-24 a,b,c,d. Difference between the class and threshold, individual indicators (#1-4) maximum probability of occurrence maps.

4.5: Conclusions

It cannot be argued that class simulation is a numerically better technique than threshold simulation, but overall it is not worse. Class simulation has some significant advantages over threshold simulation:

- Class simulation is more intuitive. The range of a class semivariogram is easier to understand conceptually than the range of a threshold semivariogram.
- Testing simulation sensitivities due to indicator ordering is easy to perform. Recalculation of semivariogram models is not required. In contrast, if thresholds are used, a new semivariogram model must be calculated for each threshold for each reordering, adding significant work for the modeler.

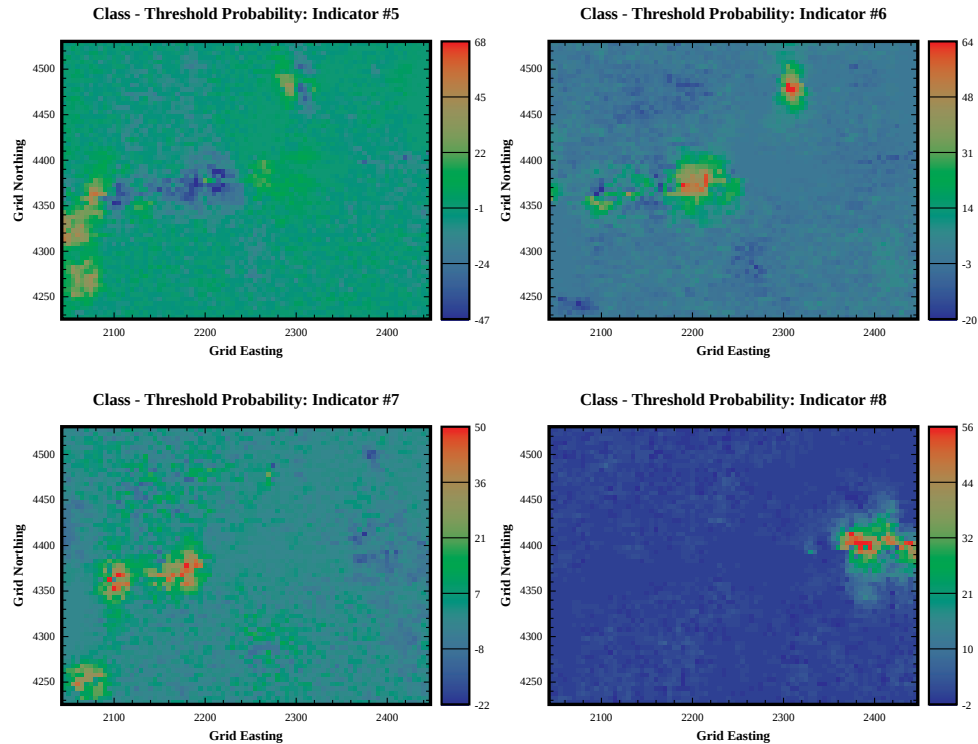


FIGURE 4-24 e,f,g,h. Difference between the class and threshold, individual indicators (#5-8) maximum probability of occurrence maps.

Several other advantages of using the classes approach were revealed during this analysis:

- ORV's, are more common with the class approach (a negative), because the last CDF value is calculated rather than implied. However, ORV's are more logical when the CDF is greater than 1.0. It can be argued that the increase in ORV's occurs because class simulation better identifies problem cells, and correctly adjusts the weights.
- Hard and soft data prior probabilities differences tend to be smaller. Using thresholds, a large difference in an early class propagates through remaining indicators, making the differences artificially large.
- Theoretically, though not proven here, the indicator ordering should not affect the simulation results. Eventually though, it may be possible to relate the class semivariograms to geological sequences ; geostatistics has not yet been able to consistently observe geologic laws.

There are some disadvantages to using classes to:

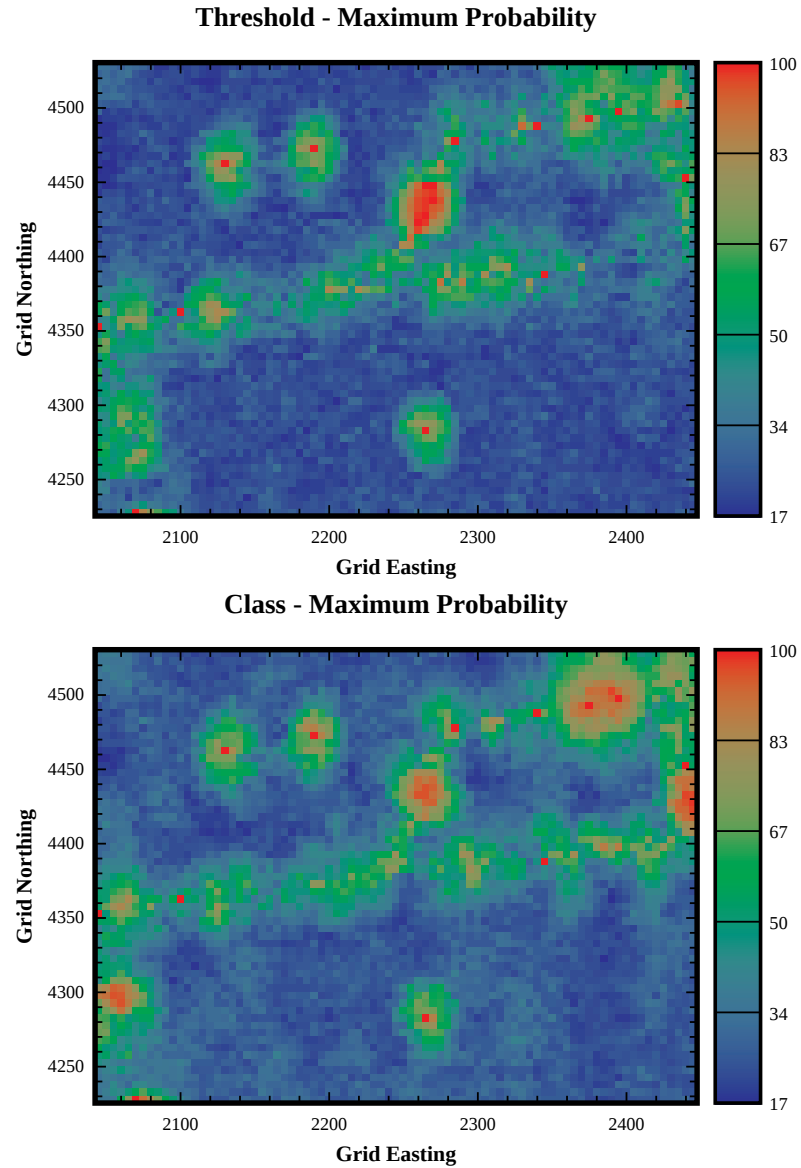


FIGURE 4-25 a,b. Maximum probability of occurrence of any indicator. At known data points, probability is 100%. The minimum probability is 12.5% ($100\% / \#$ of indicators); at these locations each indicator is equally likely to occur (no cells had this minimum probability). These maps are useful for identifying the spatial distribution of uncertainty.

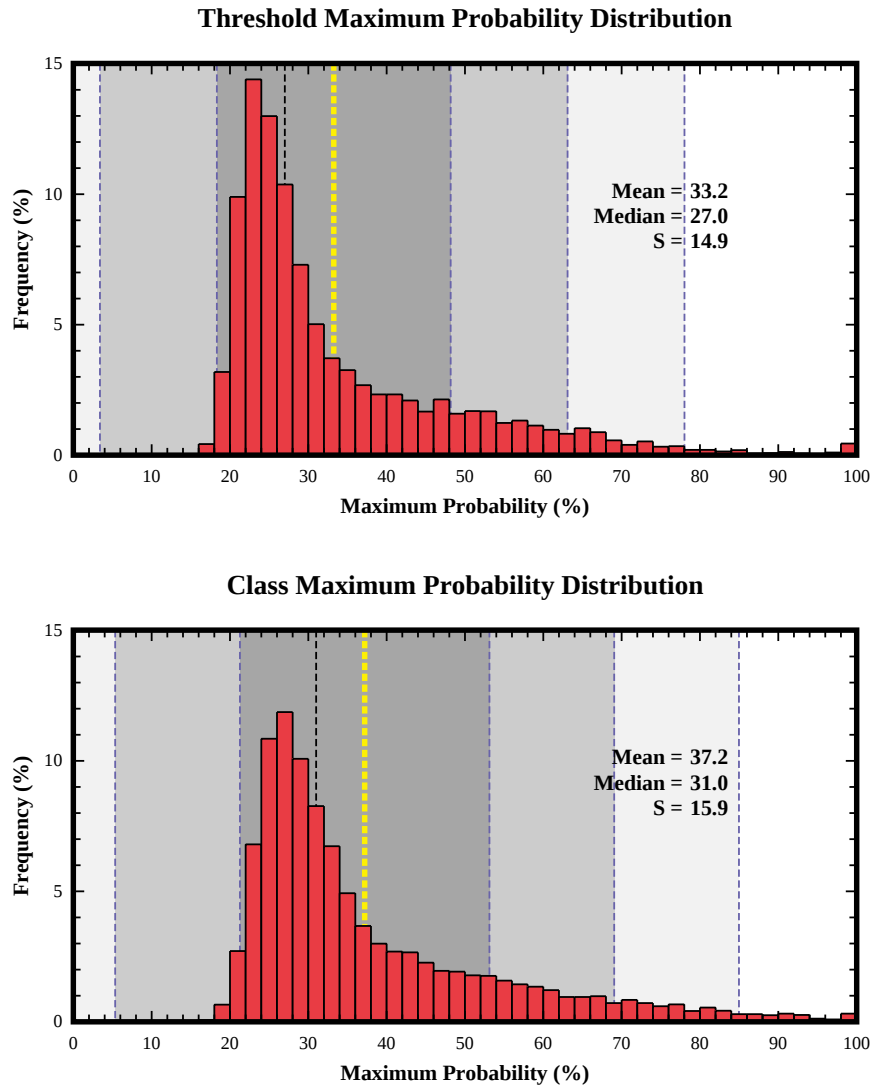
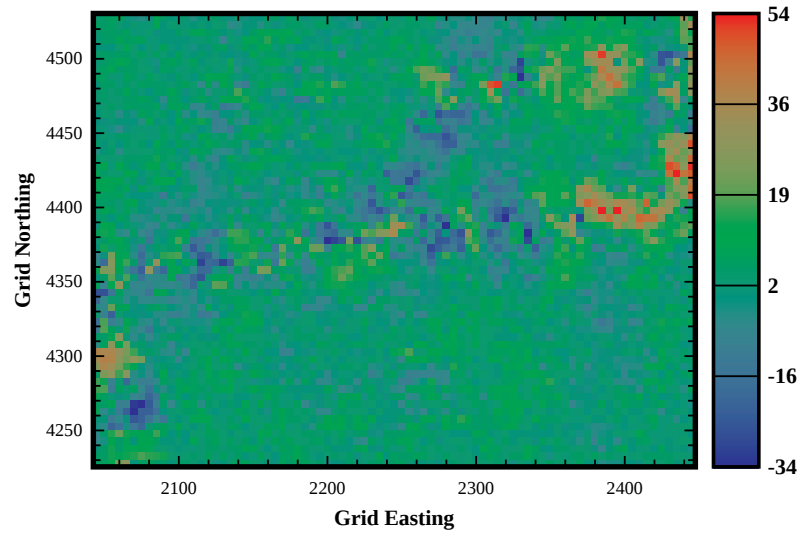


FIGURE 4-26 a,b. Histograms indicate the maximum probability of occurrence of any indicator. At known data points, probability is 100%. The minimum probability is 12.5% ($100\% / \#$ of indicators); at these locations each indicator is equally likely to occur (no cells had this minimum probability). Class realizations have slightly higher mean indicating lower overall uncertainty.

- Class simulation depends on poorer p_1 - p_2 values for Type-A soft data. It is not that the data quality has changed, but the method in handling the data has changed.

Class - Threshold: Maximum Probability Difference



Class - Threshold: Maximum Probability Difference

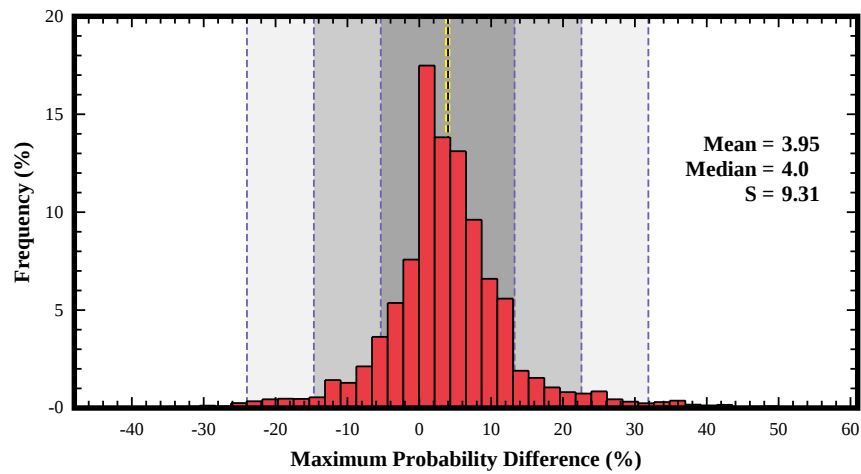


FIGURE 4-27. a) Spatial distribution of the difference between the class and threshold maximum probability maps; b) histogram of the same information. The positive mean difference indicates the class realizations have a lower level of uncertainty.

- Class simulation requires one additional semivariogram model. This requires more effort on the modelers part if sensitivity to indicator order is not being evaluated.

- Class simulation is computationally more expensive. An additional kriging matrix must be solved for every grid cell.

The last two disadvantages, are insignificant if even one indicator reordering is done to test model sensitivity to the indicator order. The modeler's effort to develop new threshold semivariograms for the new ordering outweighs the initial setup effort for the class approach. Class simulation is a useful technique, if only because it is more intuitive.

Kriging and conditional indicator simulation are valuable tools for evaluating uncertainty in the subsurface, but are limited by the assumption of stationarity (i.e. the mean and spatial variance are constant across the site). If the regions are distinct and unrelated, then zonal kriging can be accomplished manually by modeling each region and merging the results into a single model. However, merging results is expensive in human resources and computer processing time, and merged results cannot represent gradational transitions. A technique called zonal kriging was developed and is presented in this chapter so that different spatial equations can be applied to separate regions of a site. This zonal kriging algorithm is applied to a synthetic data set, to data from an extensively sampled outcrop in Yorkshire, England, and to a subsurface site at the Colorado School of Mines (CSM) survey field, in Golden, Colorado. Estimation of the synthetic data set demonstrates the advantages and shortcomings of the technique, and conditional multiple indicator simulations of both the Yorkshire outcrop and the CSM survey field data sets illustrate the improvement attained through use of zonal kriging.

5.1: Introduction and Previous Work

Significant variation of spatial statistics across a site can violate the basic assumptions of stationarity and this can lead to strongly biased estimates. Depending on the magnitude of the deviation from stationarity and the importance of the results, two approaches are often taken. One assumes the problem can be controlled with the local stationarity of the neighboring data samples, and a spatial model which reflects the mean behavior of the entire site. The other divides the area into an appropriate number of zones, describes the spatial statistics for each zone (Figure 5.1), estimates each zone, and merges the results. One problem with the second method is that the boundary between zones is often abrupt. The second method may be appropriate where the contact is a fault or an unconformity, but the results are unsatisfactory for sites with gradational transitions.

One alternative approach, that can be applied in cases with gradational boundaries is to transform all the points in the data set to match the spatial statistics of the cell currently being estimated in

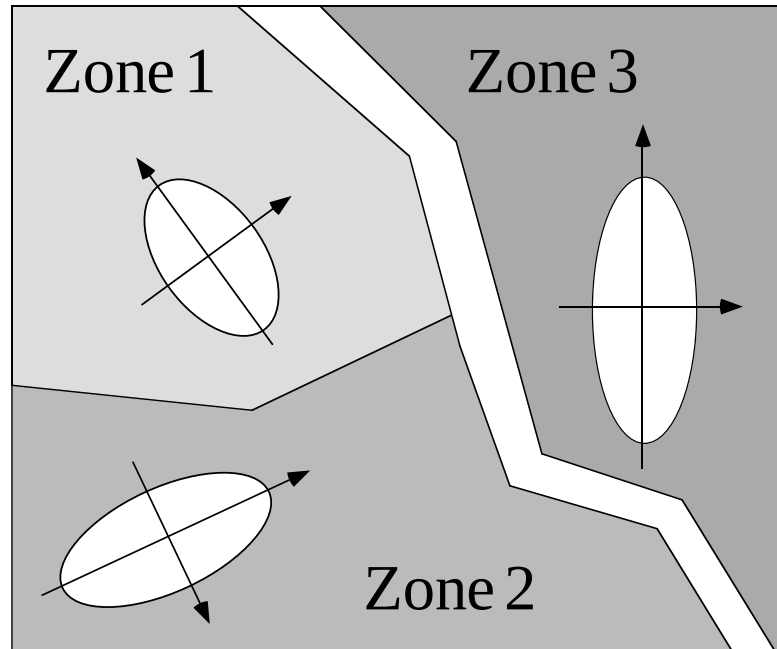


FIGURE 5-1. Spatial statistics may vary across a site, such that a single semivariogram model may not be appropriate for the entire site.

which case all the site data are considered, whether they are from the zone of interest or not. This method eliminates the problem of sharp zone boundaries, and addresses the possible gradation of properties between zones, and eliminates the need to manually merge individual zones into a single model. However, it does not accommodate sharp boundaries, nor does it recognize that some points from neighboring zones may have no bearing on the estimated value, even though they are in the transformed search neighborhood.

The approach presented here has the advantage of both of the zoning techniques described above and adds utilities to define inter-zone relationships. Such relationships describe how data, located in one zone are treated when sampled for a cell calculation located in another zone. This technique is applied using Simple (SK) and Ordinary Kriging (OK), and Multiple Conditional Indicator Simulation (MCIS). MCIS is used in two ways in this chapter. It is first used to define the zonation boundaries. MCIS can generate multiple, unique, realizations of zone boundaries, which honor the statistics of the data. This is a useful technique when the data are limited. This makes the simulation a two step process; first the zone boundaries are defined using discrete MCIS, then the interior of each zone is estimated using SK or OK. The second use of MCIS, is to populate the zones (predefined with some other method) with indicator based parameter estimates. MCIS can be used to generate zone boundaries; and MCIS, SK, or OK can be used to estimate parameter variations within the zones.

5.2: Methodology

Existing kriging algorithms (ktb3dm ; and SISIM3D) were modified to implement zonal kriging. Both codes have the required mathematical tools, but the calculation sequence was reordered, additional input describing zones and their transitional character was defined, and the search algorithm was modified. The standard kriging algorithm and modifications for zonal kriging (*italicized steps*) are shown in a flowchart in Figure 5.2. The key new aspects that have been added

Standard Kriging Algorithm with Zonal Kriging Modifications

(*Additional Zonal Kriging tasks in italics*).

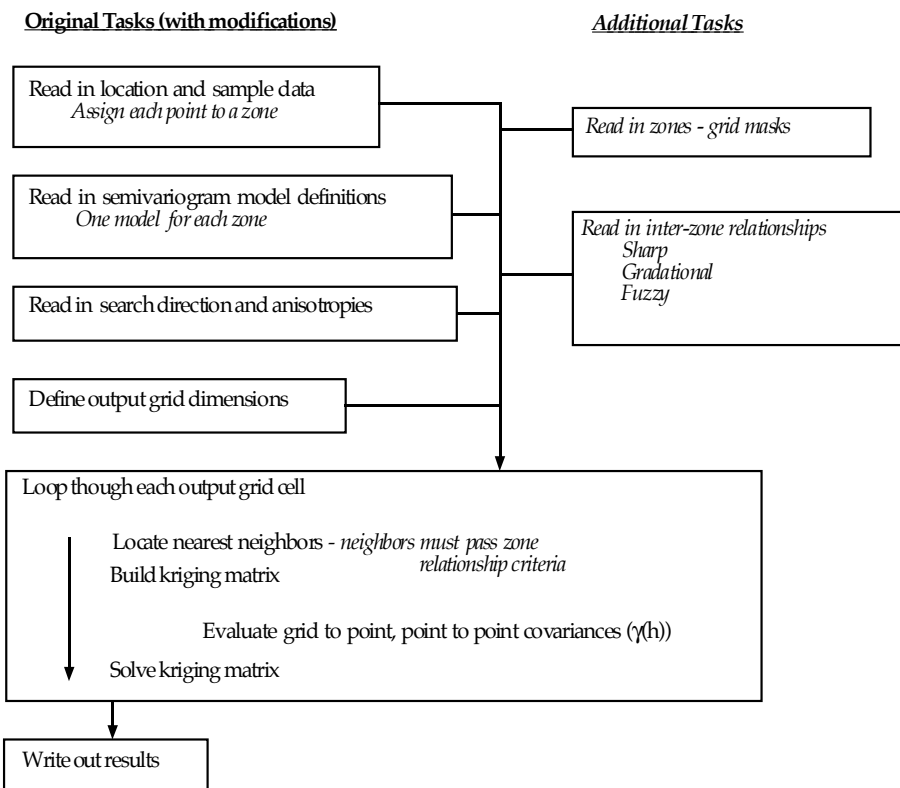


FIGURE 5-2. Basic steps involved in the standard kriging algorithm with additional steps needed to implement Zonal Kriging indicated in italics.

to the previous algorithms are:

- 1) Defining zones: Defining zones is the most arbitrary portion of the process. Choosing the location of boundaries between zones is subjective, particularly when the data are

sparse. In the synthetic example described below, a realization from a MCIS is used to define the zonation. Repeating the process using zones from a series of realizations, addresses much of the uncertainty associated with the location of the boundaries (although this is not done in this chapter). In the Yorkshire, England example MCIS is not used to define zones because there are sufficient data to determine zones by stratigraphic interpretation. At the CSM survey field, the zonal boundary is defined along the contact between two geologic formations.

- 2) Defining inter-zone relationships. Inter-zone relationships may be *Sharp* (the zones are completely unrelated), *Gradational* (one zone merges infinitely into the other), or *Fuzzy* (the zones are gradational over a limited distance and then are distinct). If the inter-zone relationship is *Fuzzy* (this term does not refer to fuzzy logic), the width of the boundary must also be defined. The rationale behind each type of transition is as follows:

Sharp: In many cases, two units are in contact with one another (Figure 5.3a), but otherwise are unrelated. Examples are faults and geologic unconformities. In this situation, it is not appropriate to use data from one zone to estimate the spatial distribution in the other.

Gradational: In some environments, units grade into one another (Figure 5.3b). This is typical of coastal deposits where beach sands grade into marine clays and shales. In each environment, the depositional systems are very different, but the change is gradational. Selecting this option requires the assumption that the region being estimated lies fully within the gradational region.

Fuzzy: Fuzzy inter-zone relationships are similar to the Gradational relationship, however the transitional distance is limited in extent (Figure 5.3c). Beyond a defined distance, data from the other zone is no longer correlated to the location being considered. The width of this boundary is subjective, as it is not necessarily related to the range of the semivariogram of either zone. Defining the width of the zone is left to the modeler, and is based on their experience and knowledge of site conditions. The gradational method, is a special case of the Fuzzy method with an infinite boundary width.

Once the user has defined the zones and the inter-zone relationships, the algorithm:

- 3) assigns each data point to a zone.
- 4) creates a mask for each zone, describing how each cell within the grid will be treated.

Once these prerequisite details are defined, each grid cell is evaluated. When estimating a grid cell value, the modified programs use the properties associated with the zone in which that cell lies. This includes search criteria and semivariogram model information. The algorithm then:

- 5) finds the nearest neighboring data points: Neighboring points may be selected in several ways depending on the zone inter-relationship. If the points are in the same zone they are treated normally. If the boundary is *sharp*, no points from across the

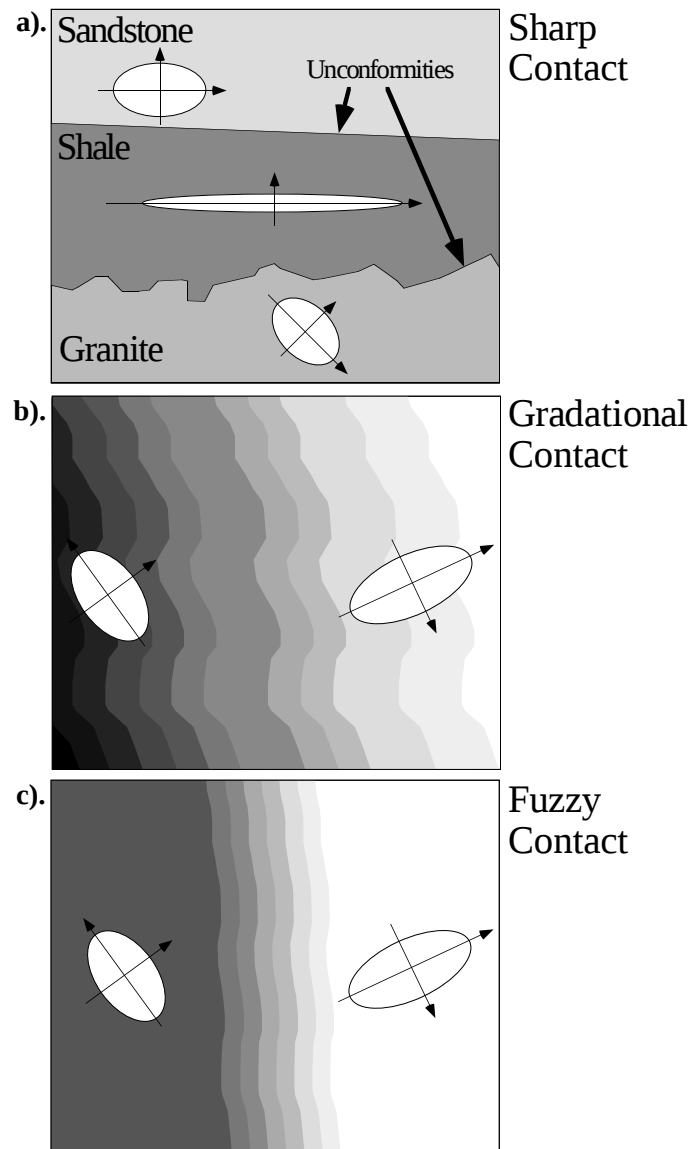


FIGURE 5-3. Different methods for describing zone contacts: a) sharp, b) gradational, and c) fuzzy.

boundary will be used (Figure 5.4a). If the boundary is *gradational*, the nearest points will be used regardless of the zone they belong to (Figure 5.4b). If the transition is fuzzy, points can be selected from the neighboring zone, but only to a limited distance (Figure 5.4c), and all other points are ignored.

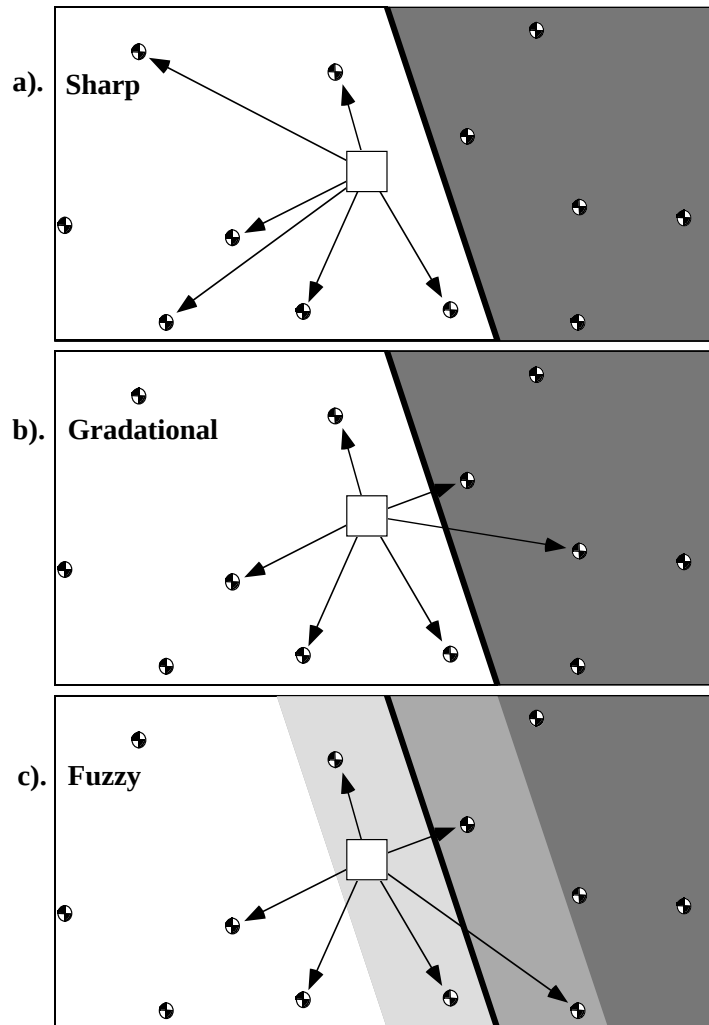


FIGURE 5-4. The search for nearest neighbors varies with zone boundary type: a) sharp, b) gradational, and c) fuzzy.

- 6) generates and solves the kriging matrix: Once the nearest neighbor points have been selected, the kriging algorithm proceeds, evaluating all points using the semivariogram model from the zone of the cell being estimated. This is regardless of which zone each point is from, or how deep the point is into the neighboring zone. The information describing the semivariogram from the neighboring zone is ignored because merging the models may lead to a matrix that is not positive definite, a basic requirement for kriging.

5.3: Examples

Three examples are used to demonstrate the utility of zonal kriging. The first is based on a small sample of eleven synthetic data points. This example demonstrates 1) the differences between SK with and without zoning; 2) how different inter-zone relationships can effect results; and 3) some of the shortcomings of zonal kriging. The second example applies zonal kriging and indicator simulation to an extensively sampled fluvio-deltaic outcrop in Yorkshire, England. This outcrop exhibits two zones with different spatial characteristics, and was "sampled" on seven vertical lines representing bore-holes, thus allowing the comparison of simulations based on a small sample of points, with the full, known section. This example demonstrates that dividing the cross-section into two zones yields more realistic realizations as compared with modeling the site using a single zone. The final example uses a data set from the Colorado School of Mines survey field, Golden, Colorado. Zonal kriging is applied to field data in combination with techniques described in earlier chapters (directional semivariograms, class indicators).

5.3.1: Synthetic Data Set Example

A simple, synthetic, two-dimensional data set of porosity's (%) with eleven data points (Figure 5.5a) was evaluated using SK (Figure 5.5b). If the spatial statistics of the site are relatively consistent, this may be a good interpretation. However, if the data reflect material properties from three distinctly different areas, where the spatial statistics are substantially different, another interpretation is needed (three zones is excessive given the amount of data, but is useful for this demonstration). Given a map of the material zones (Figure 5.5c is one possible zonation realization created using MCIS), the site can be modeled using one of several assumptions: 1) the zones are completely unrelated; 2) there is an infinite gradation between zones; or 3) there is a short distance over which the zones are gradational. For these examples, the following isotropic semivariogram models were used:

	Single Zone	Zone 1	Zone 2	Zone 3
Model Type	Spherical	Spherical	Spherical	Spherical
Range	150 m	150 m	175 m	200 m
C₁	0.24	0.24	0.22	0.22
Nugget	0.01	0.01	0.03	0.03

Although these models are not dramatically different, some interesting results are obtained.

Simple kriging with sharp, non-gradational contacts yields the map shown in Figure 5.5d. In this model, Zone 1 (Figure 5.5c - blue) and Zone 3 (red) are gradational, but Zone 2 (green) is completely unrelated. This yields sharp transitions between Zone 2 and the other two zones. Within each zone though, the surfaces are smooth as would be expected with SK. The sharp transitions match the zone definition shown in Figure 5.5c.

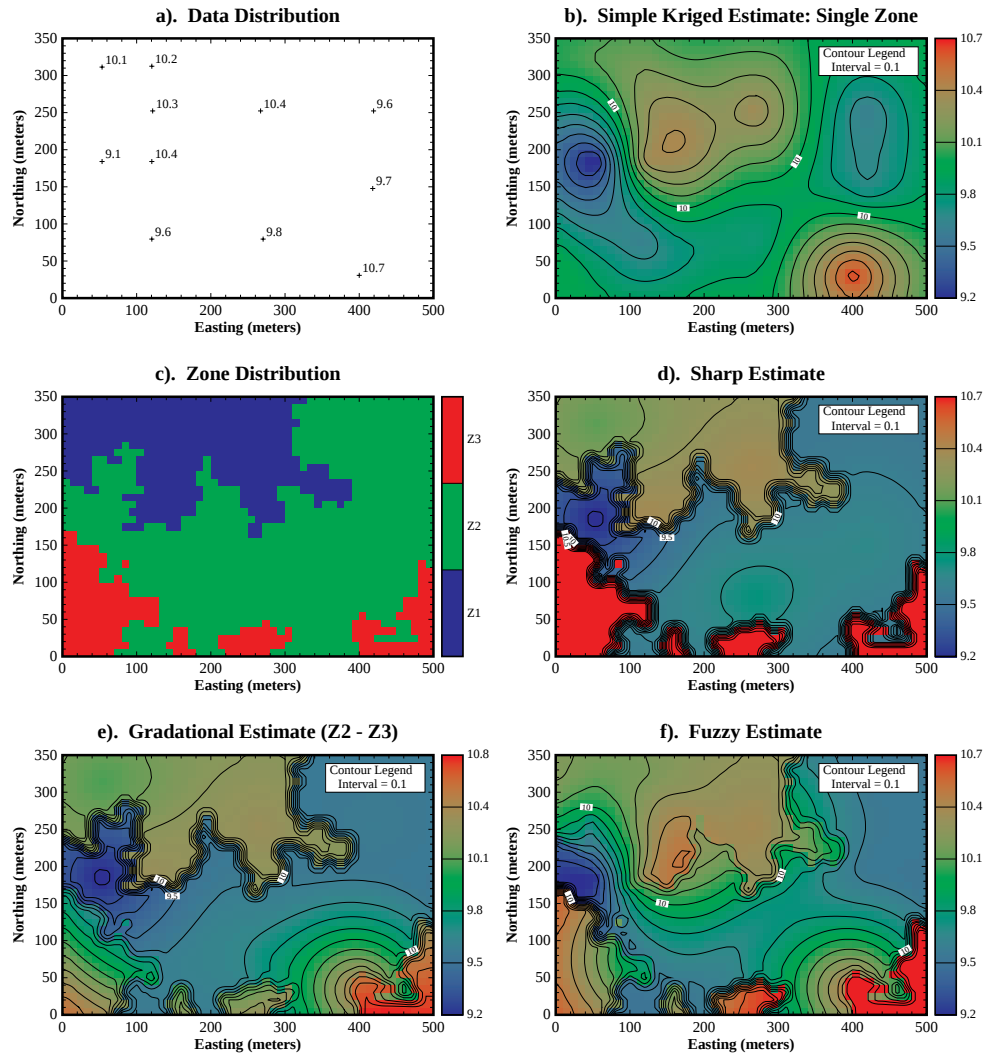


FIGURE 5-5. Different forms of ordinary and zonal kriging. (a) Sample data, (b) a traditional Simple Kriged map, (c) one possible zone map from a conditional indicator simulation, (d) sharp transition, (e) gradational transition, (f) fuzzy transition.

A gradational contact between adjacent units produces a different map. In this example (Figure 5.5e), Zones 1 and 3 are gradational, and the contact between Zones 2 and 3 is also defined as gradational (the contact between Zone 1 and 2 is sharp). The contour lines in Zone 1 are unchanged from Figure 5.5d, but there is substantial smoothing between Zones 2 and 3, although it is not complete, and the boundaries are somewhat abrupt. This incomplete smoothing occurs because the

data control for neighboring cells in different zones is coupled with different semivariogram models, and model ranges near the sample spacing of the data. These rough transitions will disappear with finer sampling.

The final option is explored by defining transitional (fuzzy) contacts of finite thickness. In this example, Zone 1 was defined to have a fuzzy boundary of 20 meters with Zone 2, and Zone 3 was defined to have a fuzzy boundary of 40 meters with Zone 2. Although there is substantial smoothing, each zone maintains much of its own character (Figure 5.5f), and the map is still substantially different than the map generated using single zone SK (Figure 5.5b). In some of the fuzzy boundary zones, particularly near the southern map border (Zone 2 vs. Zone 3), the contacts are still quite abrupt. This is, as noted for the gradational method, due to lack of data within the model range.

Which of the above models best represents the synthetic site is not an issue for this discussion. What is important, is that the modeler can control zonal differences and inter-relationships as appropriate for the site under evaluation, thus providing additional flexibility. Neither SK, nor OK could generate anything other than the first map (Figure 5.5b) or the second map if results were merged manually (Figure 5.5d).

As shown in Figures 5.5e and 5.5f, gradational and fuzzy boundaries can be abrupt. This is a numerical phenomenon, not a physical feature. The problem occurs when the set of nearest neighboring points are substantially and consistently different. In this model, this occurs because sample spacings are close to the range of the zonal semivariogram models. If the site was modeled with 1) more data points, or 2) with longer semivariogram model ranges, the contacts would be less abrupt.

5.3.2: Yorkshire, England Example

An outcrop cross-section from a fluvio-deltaic deposit near Yorkshire, England was sampled on a 17 x 20 cm grid spacing. For practical reasons (computation time) the data were upscaled to 2m x 1m grid blocks. The full 600m x 30m cross-section is shown in Figure 5.6a, and is composed of three materials: shale (SH), shaley-sandstone (SH-SS), and sandstone (SS).

To demonstrate the advantages of zonal kriging, the cross-section was "sampled" with seven vertical lines representing wells (Figure 5.6b) (1m vertical samples). Two distinct zones exist: a lower zone with high continuity in the SS and SH-SS units, and an upper zone dominated by SH, with small lenses of SS and SH-SS. To confirm that the spatial statistics were indeed different, the original section was divided into two zones (Figure 5.6d) based on a fence diagram of the bore-hole data. The facies frequencies and semivariograms for the entire cross-section were compared with those from the two zones. Whether examining the data from the exhaustive data set of the full

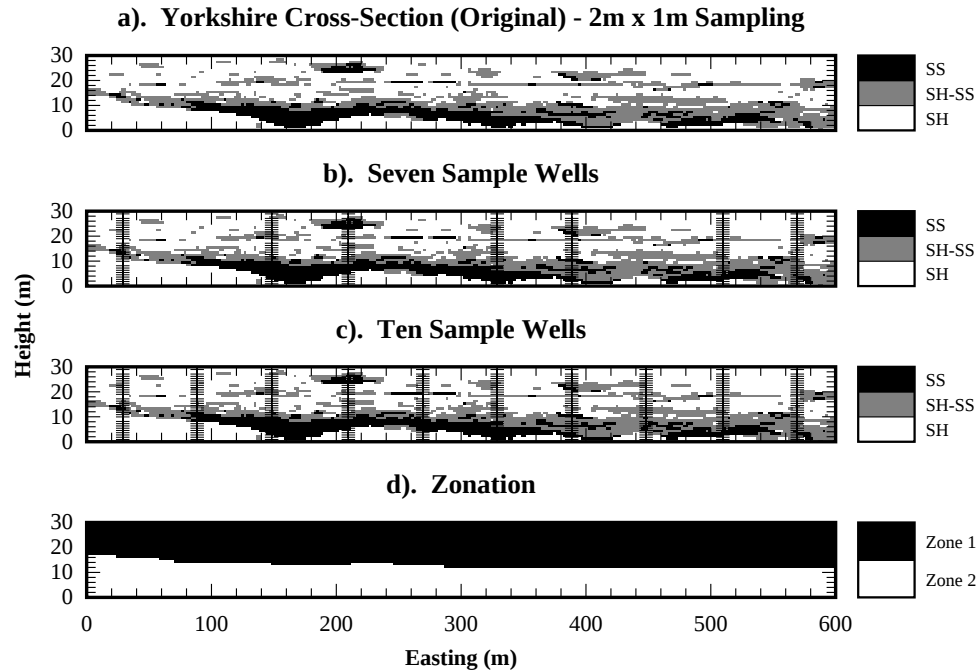


FIGURE 5-6. Model definition information for the Yorkshire cross-section: a) actual cross-section sampled with 2m x 1m cells; the locations of b) 7 and c) 10 “well samples;” d) zone definition.

cross-section, or the well samples, the results are similar. About 61% of the samples are SH, 24% SH-SS, and 15% SS. When the individual zones are separated, the distributions are very different:

	SH (%)	SH-SS (%)	SS (%)
Cross-Section (All)	63.4	23.6	13.0
Wells (All)	59.0	24.8	16.2
Cross-Section (Top)	81.2	16.4	2.4
Cross-Section (Bottom)	40.2	32.9	26.9
Wells (Top)	80.7	16.8	3.5
Wells (Bottom)	37.2	32.6	30.2

The top zone contains more than twice as much SH as the bottom zone and almost no SS, whereas the materials are more evenly distributed in the bottom zone. When semivariograms are calculated for the entire section, and for the top and bottom zones, using both the full cross-section and the well data, the zonation is again apparent. The horizontal and vertical indicator semivariograms are shown in Figures 5.7 through 5.10. The spatial statistics of the top zone are clearly different than

Exhaustive Yorkshire Cross-Section Indicator Semivariograms: SH/SH-SS Threshold

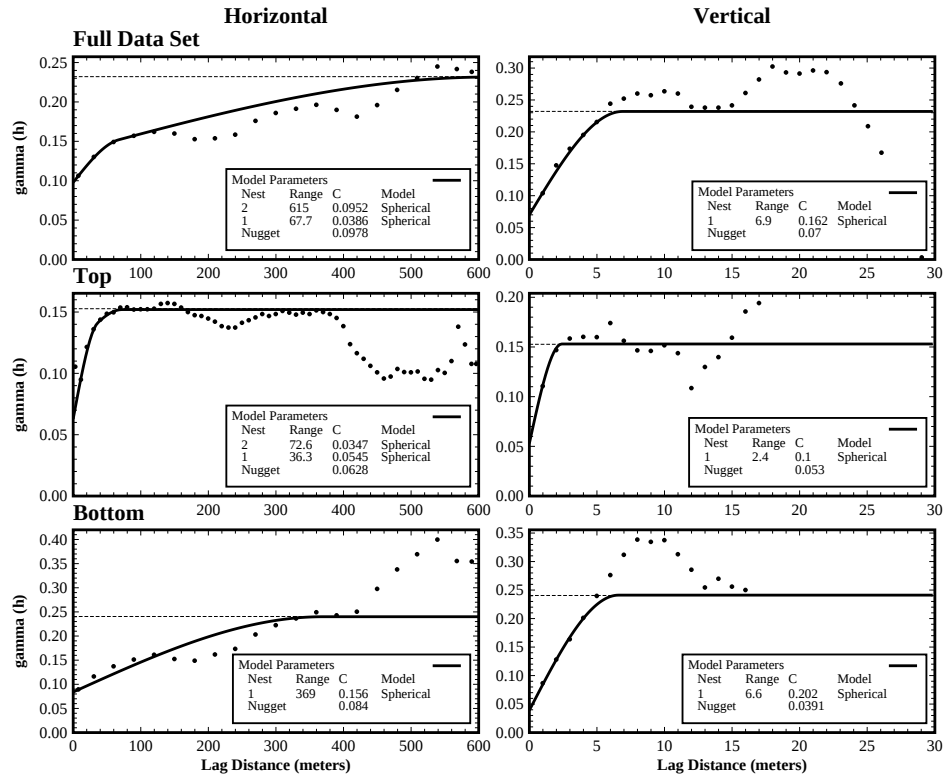


FIGURE 5-7. Exhaustive experimental and model indicator semivariograms for SH/SH-SS threshold (full cross-section) of the Yorkshire data set.

those of the bottom zone, and those of the entire cross-section. The horizontal range for the SH/SH-SS threshold in the top zone is only about 15% to 25% that of the bottom zone (Figures 5.7 and 5.9). For the SH-SS/SS threshold, the horizontal range in zone 1 is about 35% of that in zone 2 (Figures 5.8 and 5.10).

The assumption of stationarity is not applicable to this cross-section, thus modeling this site using a single set of semivariograms is statistically inappropriate. Two ensembles of 100 realizations each, demonstrate that zonal kriging yields more accurate results than a single semivariogram model. The first ensemble is based on the assumption that stationarity is valid, and only one semivariogram model set is required (the full well data set semivariogram models, Figures 5.9 and 5.10). The second set is based on the observation that stationarity is violated between zones, but is valid within each zone, thus a sharp transition is assumed, and the "Top" and "Bottom" semivariogram models (Figures 5.9 and 5.10) are used. The ranges of both horizontal and vertical semivariograms, in the top zone are much shorter than the ranges for the full section and the bottom zone. The model grid

Exhaustive Yorkshire Cross-Section Indicator Semivariograms: SH-SS/SS Threshold

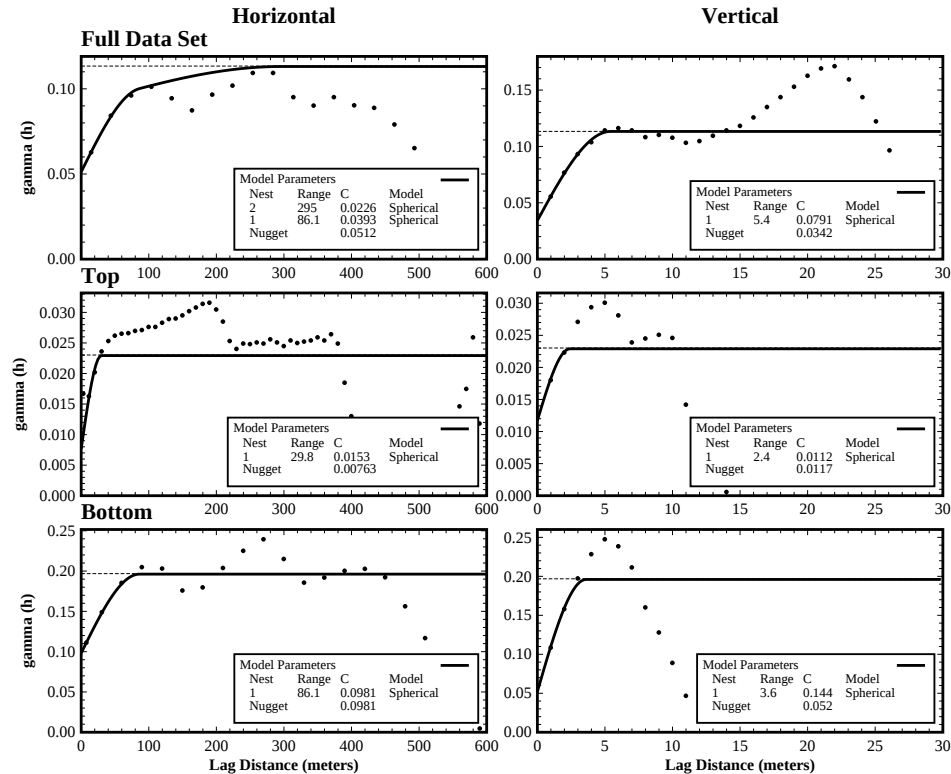


FIGURE 5-8. Exhaustive experimental and Model indicator semivariograms for SH-SS/SS threshold (full cross-section) of the Yorkshire data set.

matches the original cross-section grid exactly; with cells 2m x 1m in 300 columns and 30 layers. Results of several realizations using one and two zones respectively, are shown in Figures 5.11 and 5.12. Differences between the two sets are subtle, but discernible in the top zone. There are differences in the bottom zone, but they are minor. To evaluate the results, realization pairs were compared in order (Realization #1 [1 Zone] vs. #1 [2 Zones], #2 [1 Zone] vs. #2 [2 Zones], etc.) because matched realizations use the same "random" search path to evaluate the grid, and start using the same random seed. Therefore differences are only attributed to the use of zoning, and the realization with the smallest number of misclassifications, is defined to be the more accurate model. The number of misclassifications was calculated by subtracting the actual grid value from the estimated grid value at each grid location; any cell with a non-zero error was misclassified.

In the four realizations shown (using either method), there is a largely continuous SH-SS/SS unit in the bottom zone, which is present in the actual section, though there are random fluctuations in the realizations. The random fluctuations in the realizations are due to the large nugget terms (up to

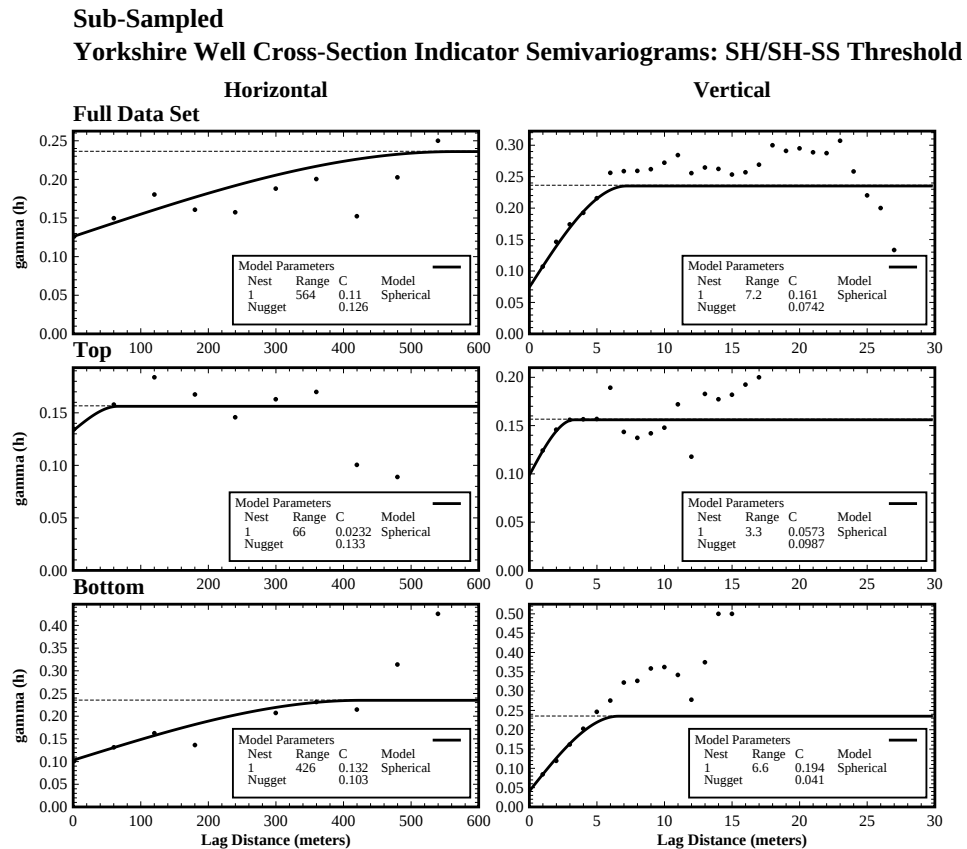


FIGURE 5-9. Sub-sampled experimental and Model indicator semivariograms for SH/SH-SS threshold (well data) of the Yorkshire data set.

64% of variance) in the model semivariograms. The difference between the lower portions of the two series of simulations is minor, because the global semivariogram models are fairly similar to the semivariogram models for the lower zone.

There are greater differences between the one and two zone simulations in the top zone. Both exhibit a great deal of randomness due to large nuggets, but the results of using two zones create more solid lenses without as many small isolated cells (compare the single and two-zone versions of realization #66, this was one of the most accurate models, and #22, one of the least accurate). The single zone model tends to create long, thin layers with considerable scatter, as compared to the more connected, but less laterally continuous units in the upper section of the two zone realization.

One-hundred pairs of realizations (one zone and two zone) were generated using the same random number sequence, and random path through the model grid. These pairs were compared to the actual cross-section. Realizations generated with the two-zone approach had fewer (10 to 328; both

**Sub-Sampled
Yorkshire Well Cross-Section Indicator Semivariograms: SH/SH-SS Threshold**

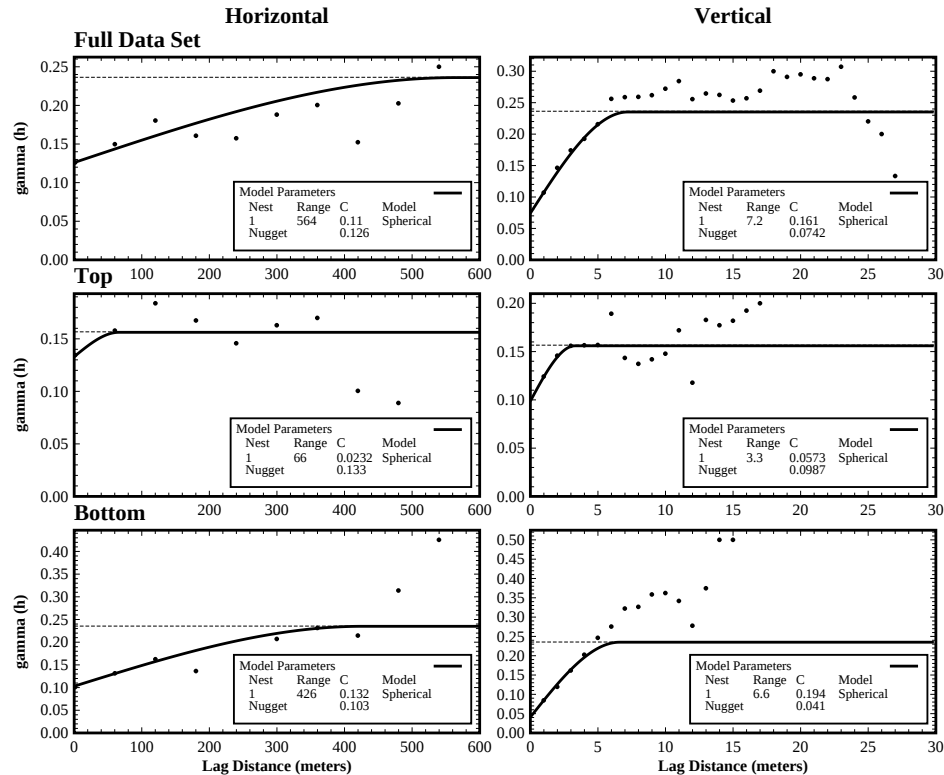


FIGURE 5-10. Sub-sampled experimental and Model indicator semivariograms for SH-SS/SS threshold (well data) of the Yorkshire data set.

methods consistently, correctly, identified about 5700 of the 9000 cells) misclassifications in 80 of the 100 pairs. When ten equally spaced wells (Figure 5.6c) were used with the same semivariogram models, instead of the original seven wells (Figure 5.6b), realizations generated with the two zone approach had fewer misclassifications in 91 of the 100 pairs (new semivariogram models may have improved results even more). The reduced number of misclassifications indicate that modeling the site with two zones yields better results. However more work is required to set up the model (additional semivariogram calculations, data entry, zone definition).

In addition to improved accuracy, the modeler has more flexibility in fine tuning the model solutions for each zone. The results can be dramatic in one zone without affecting the other. A sample two-zone realization is shown in Figure 5.13a. Realizations resulting from independently decreasing the nuggets in the top and bottom zones are shown in Figures 5.13b and 5.13c respectively. In both cases, the continuity is increased and the random scatter is reduced, without

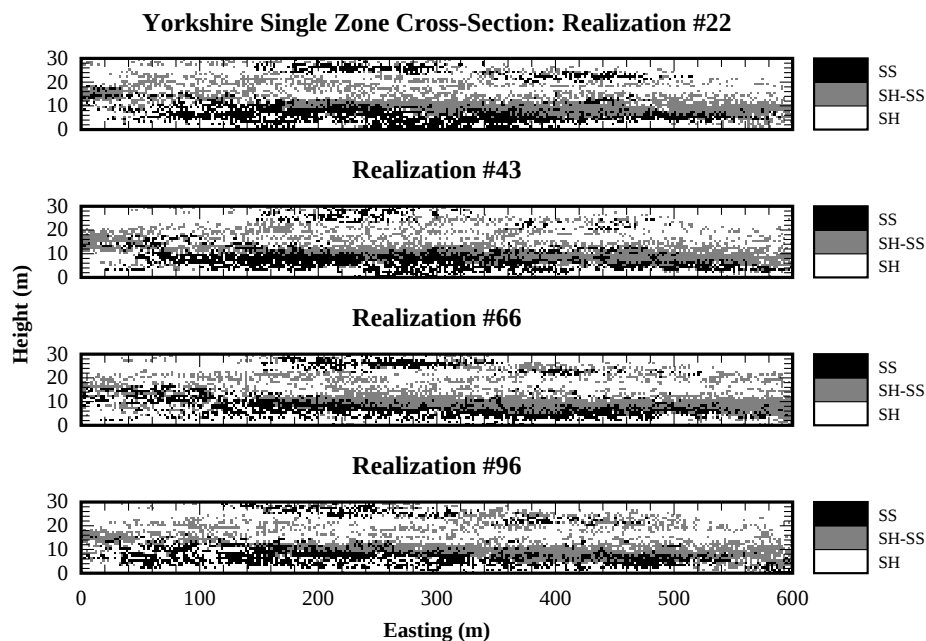


FIGURE 5-11. Realizations from single-zone simulation series.

affecting the alternate zone. If these same changes were made on a single zone realization, they would effect the entire grid, and improvement in one portion of the grid would have to be balanced against the degradation of the other zone.

5.3.3: Colorado School of Mines Survey Field Example

Unlike the Yorkshire, England example, field geology is rarely known in such detail. Usually only a limited amount of data are available at a site, typically far less than would be desired. At the CSM survey field, located on the West edge of Golden, Colorado (Figure 5.14), hard and soft data were collected to investigate the use of soft data for reducing uncertainty associated with ground water flow models. The site contains core and chip data from eighteen boreholes; borehole geophysical logs; and eight (Figure 5.15 and 5.16), two-dimensional, cross-hole, tomographic sections. Even though the site crosses two geologic formations, the site was initially simulated as a single zone. This was done because a zonal simulation tool was not available. In this study, zonal kriging is used in combination with directional semivariograms and class indicator simulation to model the CSM survey field.

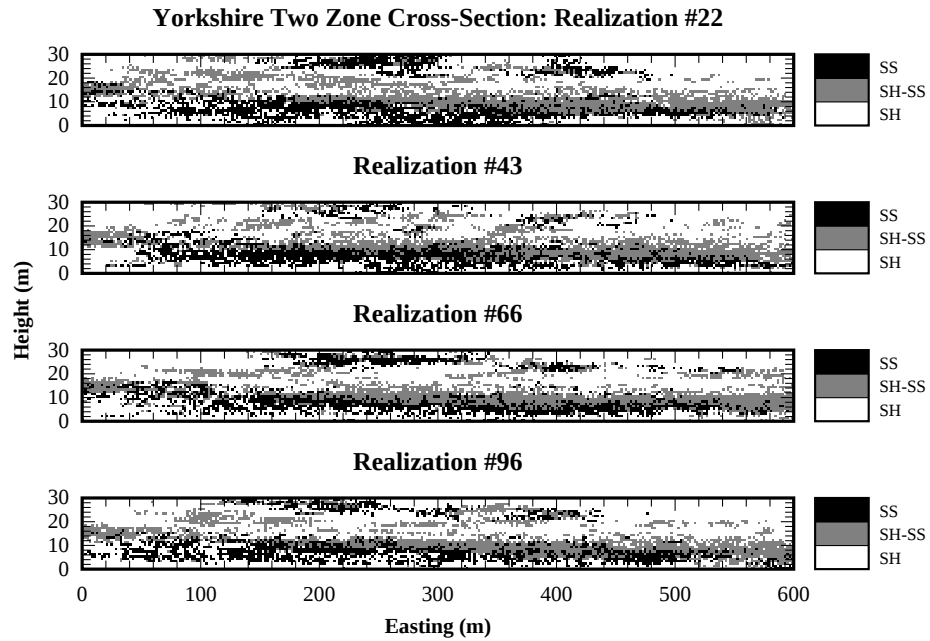


FIGURE 5-12. Realizations from two-zone simulation series.

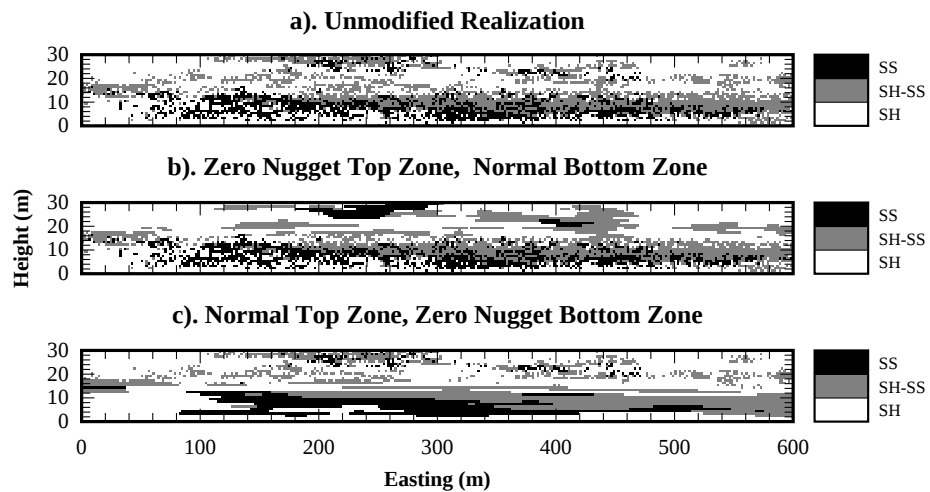


FIGURE 5-13. Impact of altering the semivariogram nugget independently in the top and bottom zones of a section: a) original simulation section; b) reduced nugget in top zone; c) reduced nugget in bottom zone.

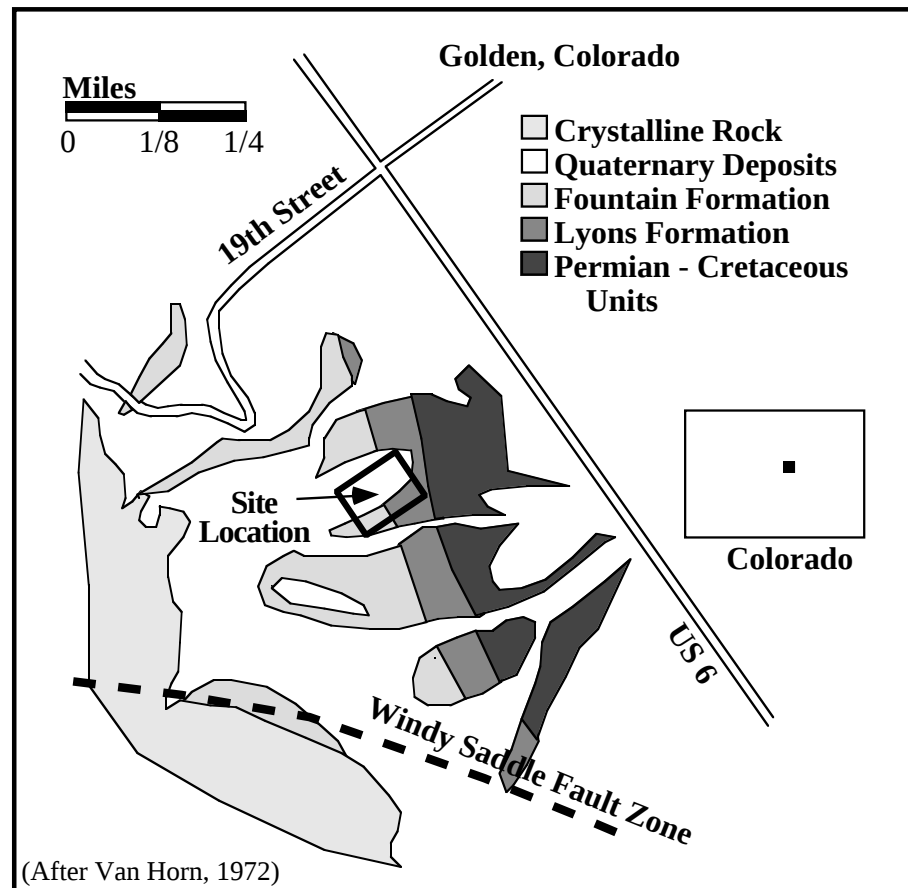


FIGURE 5-14. CSM Survey Field location map.

5.3.3.1: Evidence That Zonation is Required

Zonal kriging is needed because 1) the data populations and 2) the spatial statistics of the data from the two formations are different. The difference in spatial statistics, applied to kriging, is the most important reason for using zonal kriging. The data are divided into two data sets along the formation contact between the Fountain and Lyons Formations. The contact dips approximately 40° ENE. The contact is clearly defined on the seismic tomogram cross-sections (Figure 5.17). The full data set is shown in Figure 5.16 with all data converted to one of eight indicator values. The indicator values represent discrete sonic velocity ranges; hydraulic conductivity's (K) were

estimated from field and laboratory tests, or inferred from lithologic character and sonic velocity. Later optimal values of hydraulic conductivity were estimated using inverse flow modeling:

Indicator	Sonic Velocity (ft/sec)	Initial K Estimates (ft/day)	Optimized K Estimates (ft/day)	
			Hard Only	Hard/Soft
1	> 10870	0.0011	0.010	0.0020
2	10000 - 10870	0.0011	0.0063	0.00079
3	9050 - 10000	0.0025	0.63	0.050
4	8550 - 9050	0.0043	0.079	1.6x10 ⁻⁵
5	8050 - 8550	0.040	0.025	2.5x10 ⁻⁶
6	7250 - 8050	0.0043	NO	NO
7	6060 - 7250	0.40	0.016	0.0040
8	< 6060	7.8	NO	NO

NO = Not Optimized

Data sets, for the Fountain and Lyons Formations are shown in Figure 5.18. Examination of the full data set does not readily reveal sub-populations (Figure 5.19), but independent examination of data from the two formations reveals their differences. Indicators #4 through #7 have a high frequency in the Fountain Formation. In the Lyons Formation, indicators #1 through #4 are more frequent. Indicators #5 and #6 are poorly represented, and indicator #8 occurs exclusively in the Lyons Formation. The difference in the distributions can be explained by the materials the indicator represent:

Indicator	Material Description
1	Conglomerate (Lyons Formation)
2	Fine to coarse sandstone with conglomerate lenses
3	Fluvial sandstone (Lyons Formation)
4	Very fine to very coarse sandstone
5	No core recovered. Moderately consolidated with low-moderate clay
6	Two materials: 1) silty sandstone, and 2) poorly sorted sandstone with siltstone and conglomerate lenses
7	No core recovered. Poorly consolidated, low clay material
8	No core recovered. Well fractured area of any material type.

The spatial statistics of the indicators also differ between the two formations. Semivariograms were calculated for three principle axes (North-South, East-West, and Vertical) for each indicator threshold. The semivariogram model results are summarized in Table 5.1a-c. For all the indicators (except #8; which does not occur in Fountain Formation) there were significant differences between

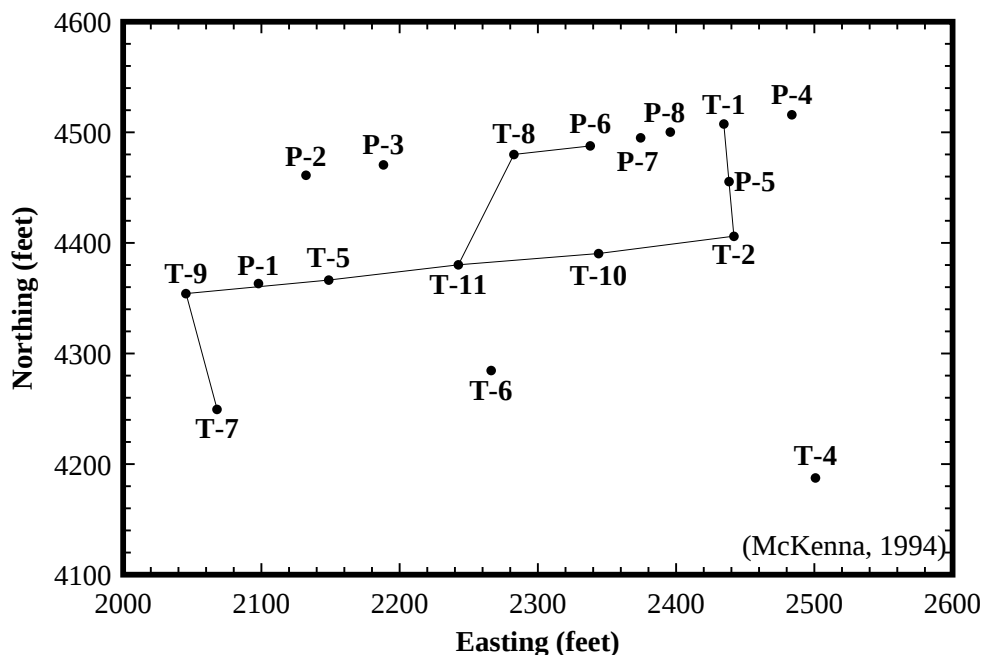


FIGURE 5-15. CSM Survey Field site map. Dots represent borehole locations. Solid lines identify location of tomography surveys.

the full data set, the Fountain Formation, and the Lyons Formation semivariogram models. For class 4, Fountain and Lyons Formation semivariograms, in the North-South direction the ranges are 72 ft. and 18 ft. respectively. In the East-West direction they are; 27 ft. and 48 ft. respectively. Not only are the ranges for the same indicator substantially different, but the principle anisotropy directions are 90° apart. Consequently it was concluded that two distinct zones are present at the site.

5.3.3.2: Grid and Model Definition

Single and two zone simulations were conducted. The regular model grid was defined as 80 columns representing 1600 feet in the X direction, 60 rows representing 1200 feet in the Y direction, and 72 layers representing 144 feet in the Z direction. The search ellipsoids for locating data were identical for both zones. The differences between the models were defined by 1) the semivariogram models used (Table 5.1), 2) the zone definitions (Figure 5.20), and 3) the quality of the soft data (Table 5.2). The first two differences have already been defined. The final difference though is less obvious, because the same data are used, however the calculation of the misclassification probabilities, p_1 and p_2 for Type-A soft data differ for the threshold and class simulations. In this example, the single zone model is simulated using thresholds, and as such, the

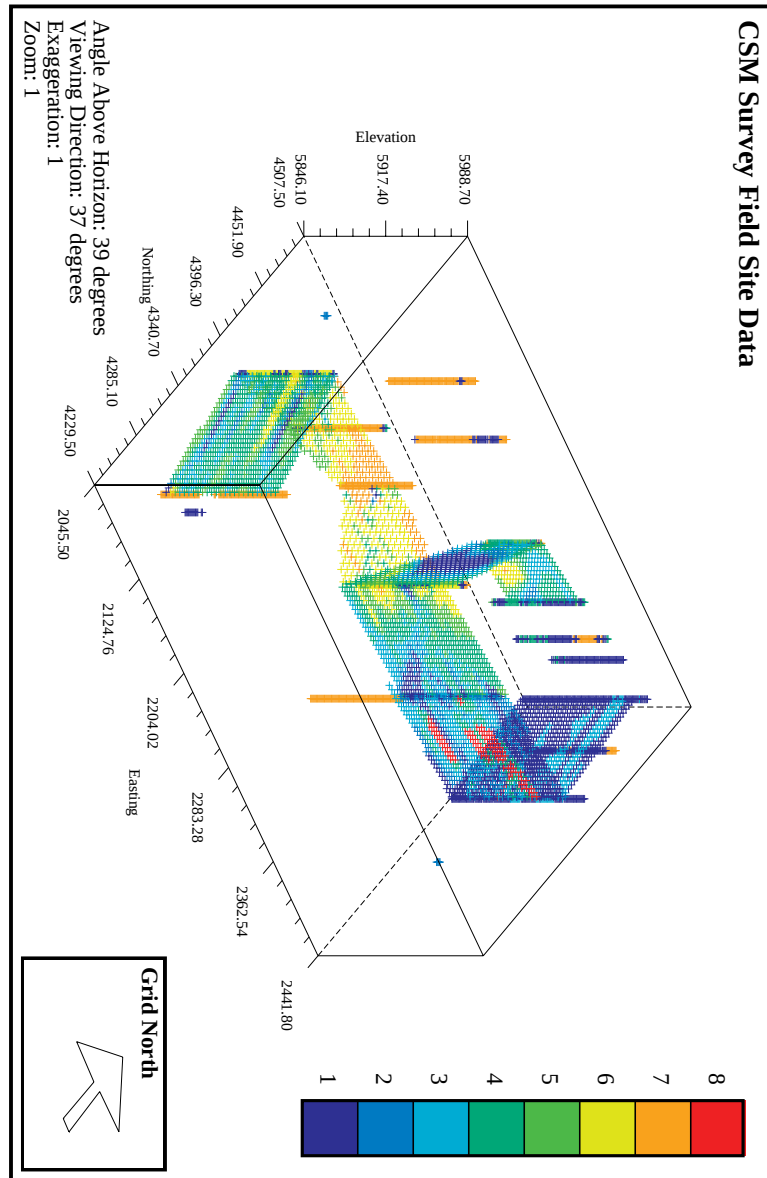


FIGURE 5-16. CSM Survey Field borehole and tomography data converted to indicators.

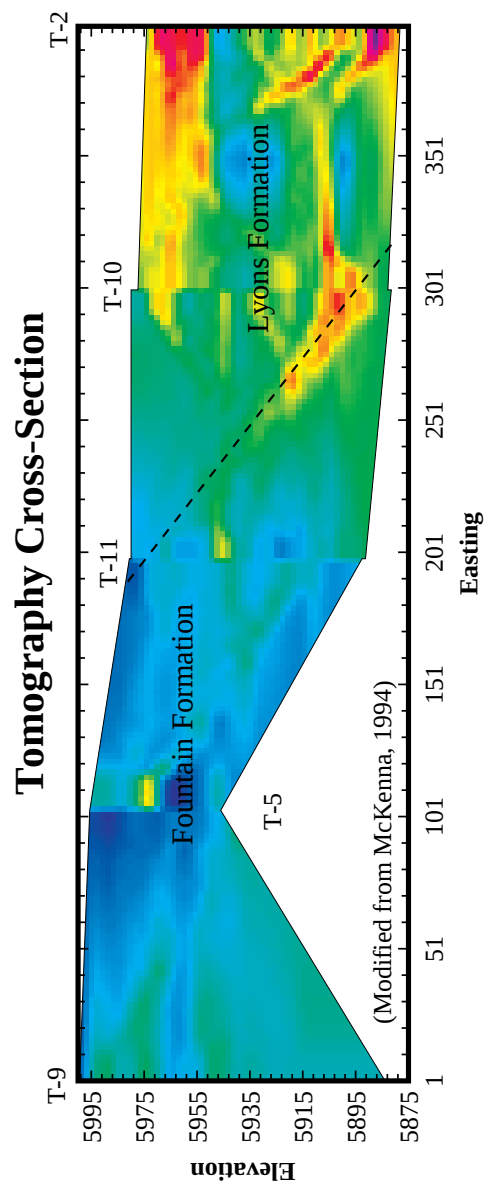


FIGURE 5-17. Tomographic cross-section at CSM Survey Field (viewing North-West). Dashed line is approximate location of Fountain / Lyons Formations contact.

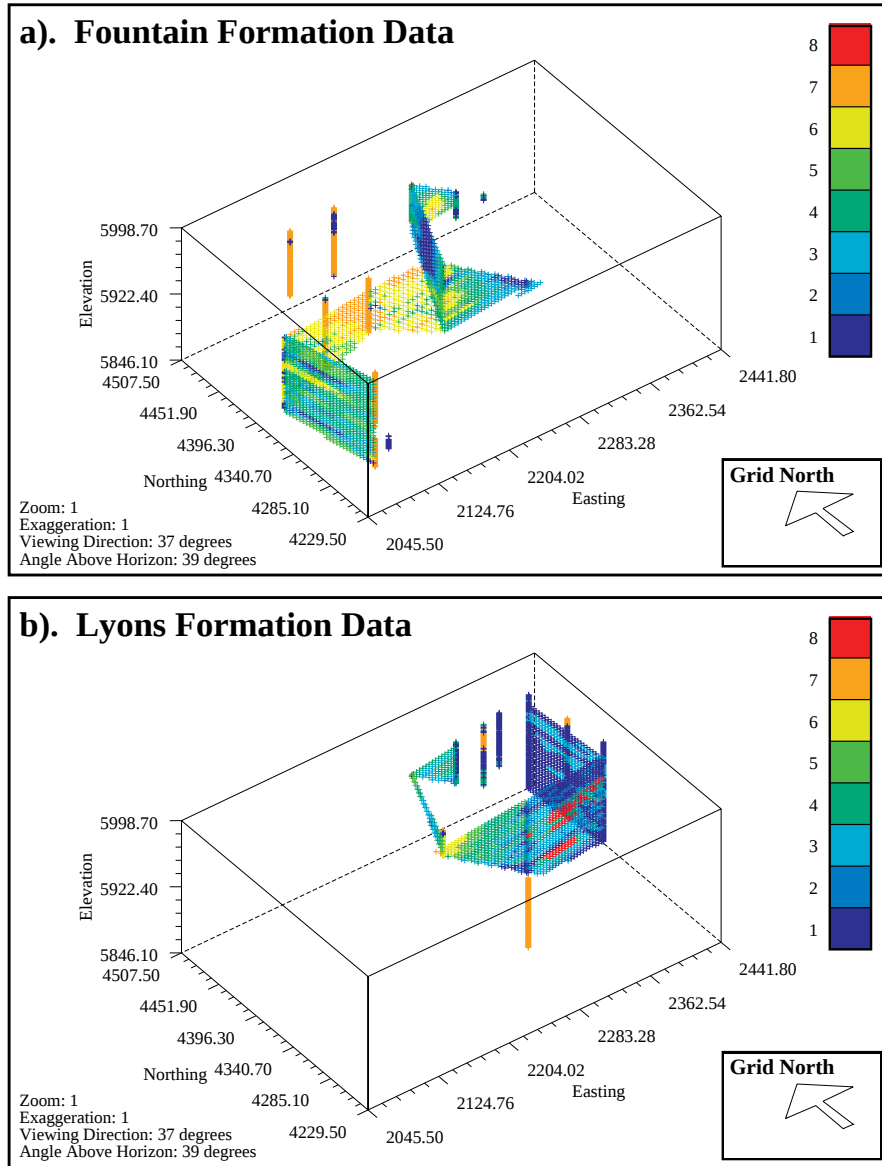


FIGURE 5-18. CSM Survey Field hard and soft data distributions (converted to indicators) for the Fountain (a) and Lyons (b) Formations.

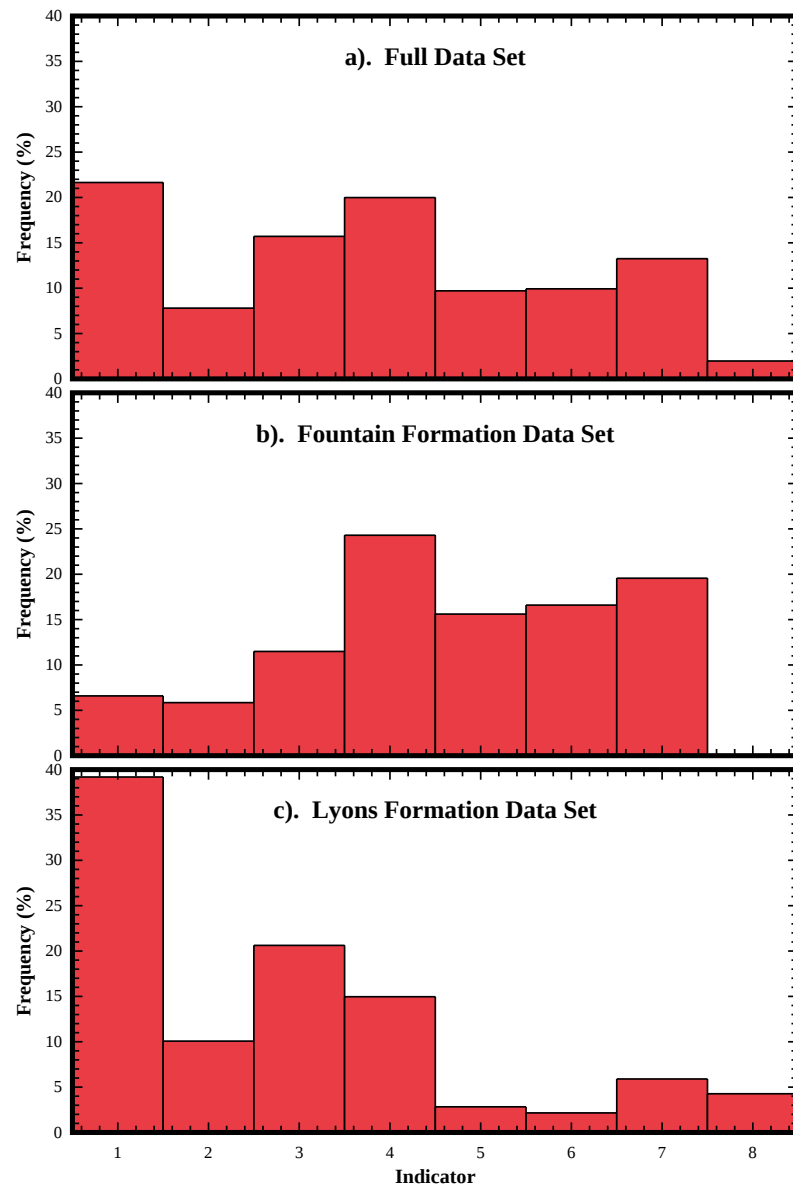


FIGURE 5-19. Distribution of hard and soft (Type-A only) data for full data set and for Fountain and Lyons Formations regions.

a). Single Zone - Threshold - All Models Spherical

Threshold	East-West		North-South		Vertical		Nugget
	Range	Sill	Range	Sill	Range	Sill	
1.5	81.0	0.118	155.6	0.118	81.0	0.0623	0.0516
2.5	93.0	0.207	126.0	0.207	54.0	0.0061	0.0
3.5	75.0	0.169	174.0	0.246	21.0	0.0468	0.0
	282.0	0.0783					
4.5	90.0	0.0831	15.0	0.149	3.0	0.0405	0.0
	204.0	0.145	99.0	0.0765	47.0	0.0358	
5.5	132.0	0.189	12.0	0.158	36.0	0.0692	0.0
			48.0	0.0308			
6.5	78.0	0.129	15.0	0.129	93.0	0.0284	0.0
7.5	61.5	0.0208	18.5	0.0208	36.9	0.0079	0.0

NOTE: Multi-nested models require two rows.

b). Fountain Formation - Zone 1/2 - All Models Spherical

Class	East-West		North-South		Vertical		Nugget
	Range	Sill	Range	Sill	Range	Sill	
1	18.0	0.0353	59.9	0.0353	69.0	0.0353	0.0265
2	15.0	0.0256	30.0	0.0256	21.0	0.0139	0.0297
					48.0	0.0115	
3	18.0	0.0215	15.0	0.0256	30.0	0.0256	0.0297
	60.0	0.0229					
4	27.0	0.111	72.0	0.111	72.0	0.111	0.0253
5	15.0	0.0500	57.0	0.100	27.0	0.100	0.0317
	156.0	0.500					
6	9.0	0.103	36.0	0.137	78.0	0.137	0.000
	60.0	0.0353					
7	66.0	0.137	27.0	0.137	44.0	0.0786	0.0214
8	74.0	0.142	125.0	0.142	74.0	0.0799	0.0100

NOTE: Multi-nested models require two rows.

TABLE 5.1 a,b. CSM Survey Field threshold (single-zone) and class (two-zone: Fountain and Lyons Formation) semivariogram models.

c). Lyons Formation - Zone 2/2 - All Models Spherical

Class	East-West		North-South		Vertical		Nugget
	Range	Sill	Range	Sill	Range	Sill	
1	75.0	0.160	114.0	0.160	75.0	0.160	0.0776
2	54.0	0.0902	12.0	0.0902	42.0	0.0583	0.000
3	84.0	0.148	84.0	0.148	69.0	0.148	0.0141
4	48.1	0.126	18.0	0.126	75.1	0.126	0.000
5	27.0	0.0275	18.0	0.0275	45.0	0.0275	0.000
6	15.0	0.0211	15.0	0.0211	5.0	0.0211	0.000
7	12.0	0.0468	12.0	0.0468	63.0	0.0468	0.00856
8	66.0	0.0410	15.0	0.0410	63.0	0.0410	0.000

NOTE: Multi-nested models require two rows.

TABLE 5.1 c. CSM Survey Field threshold (single-zone) and class (two-zone: Fountain and Lyons Formation) semivariogram models.

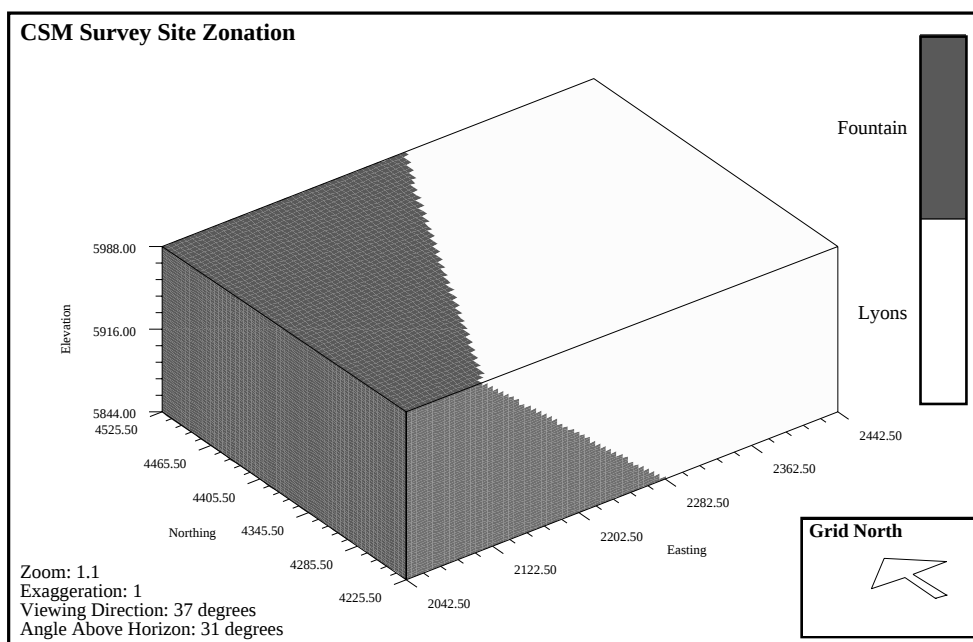


FIGURE 5-20. CSM Survey Field zone definition.

Threshold	Cumulative Probability		
	Hard	Soft	Difference
1.5	0.1638	0.2306	0.0668
2.5	0.2370	0.3120	0.0750
3.5	0.2706	0.5031	0.2325
4.5	0.3693	0.7306	0.3613
5.5	0.3886	0.8491	0.4605
6.5	0.5066	0.9429	0.4363
7.5	1.0000	0.9748	0.0252

Class	Individual Probability		
	Hard	Soft	Difference
1	0.0018	0.0843	0.1317
2	0.0148	0.0710	0.0524
3	0.0000	0.1479	0.1400
4	0.1107	0.2810	0.1552
5	0.0314	0.1919	0.1502
6	0.1697	0.1649	0.0137
7	0.6716	0.0589	0.6159
8	0.0000	0.0000	0.0000

Class	Individual Probability		
	Hard	Soft	Difference
1	0.3628	0.3936	0.0308
2	0.1451	0.0888	0.0563
3	0.0748	0.2415	0.1667
4	0.0839	0.1673	0.0834
5	0.0045	0.0347	0.0302
6	0.0544	0.0128	0.0416
7	0.2744	0.0012	0.2732
8	0.0000	0.0541	0.0541

TABLE 5.2. CSM Survey Field threshold (single-zone) and class (two-zone: Fountain and Lyons Formation) hard and soft data (Type-A only) sample data distributions.

p_1 - p_2 values are calculated based on a value being above or below a particular threshold (threshold #3, Figure 5.21). The p_1 - p_2 values used in the single-zone, threshold simulations are:

Threshold Velocity (ft/sec)	Threshold	P_1	P_2	$P_1 - P_2$
6060	7	0.00	0.00	0.00
7250	6	0.56	0.04	0.52
8050	5	0.58	0.05	0.53
8550	4	0.63	0.10	0.63
9050	3	0.84	0.17	0.67
10000	2	0.90	0.17	0.73
10870	1	0.91	0.15	0.74

In the two-zone case, class simulation is performed and the p_1 - p_2 values are calculated based on a value being between two thresholds (class #3, Figure 5.22). Because classes are more restrictive, the probability of misclassification is higher and the p_1 - p_2 values will be lower. This suggests that the soft data (Type-A) are less useful at reducing the uncertainty in the two-zone model. As will be shown, the two-zone simulations produce smaller uncertainties. The p_1 - p_2 values used in the two-zone, class simulations are:

Velocity Range (ft/sec)	Class	P_1	P_2	$P_1 - P_2$
< 6060	8	0.00	0.00	0.00
6060 - 7250	7	0.56	0.00	0.56
7250 - 8050	6	0.39	0.04	0.35
8050 - 8550	5	0.25	0.12	0.13
8550 - 9050	4	0.45	0.18	0.27
9050 - 10000	3	0.26	0.13	0.13
10000 - 10870	2	0.38	0.05	0.33
> 10870	1	0.84	0.00	0.84

Note, the less than 6060 ft/sec velocity measurements are only observed in the tomography cross sections. Hard data are not available to for calibration.

In addition to the differences in the p_1 - p_2 values, there are significant differences in the prior hard data CDF's and soft data CDF's. These differences also affect the realization calculations and are summarized in Table 5.2.

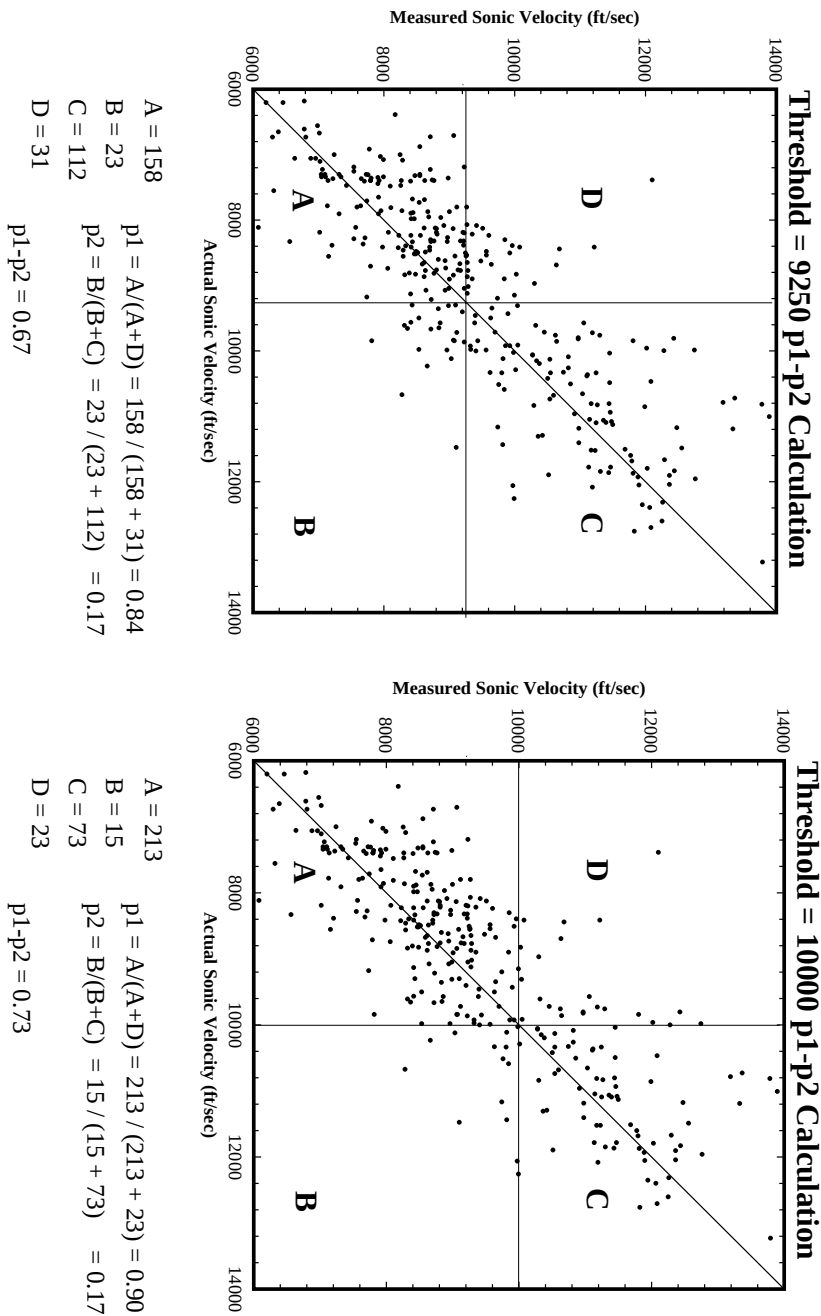


FIGURE 5-21. Graphical method for calculating p_1 and p_2 values for a specific threshold (After Alabert, 1987). Data from CSM Survey Field.

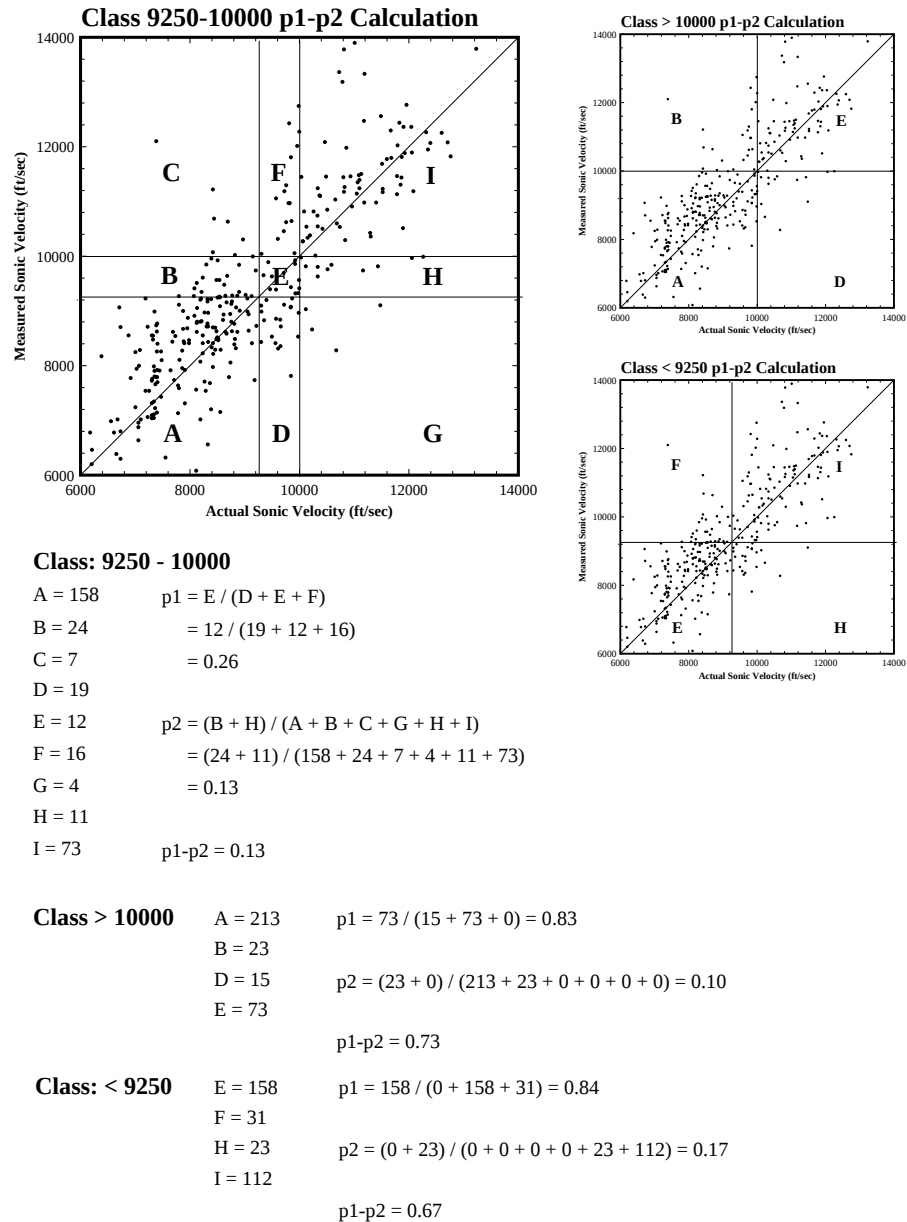


FIGURE 5-22. Graphical method for calculating p_1 and p_2 values for a specific class. Data from CSM Survey Field. Note the similarities between the *Class > 10000* and the *Class < 9250* diagrams, and the diagrams in Figure 5.21. They are fundamentally identical. This is always the case for the first and last class and threshold p_1 - p_2 calculations.

5.3.3.3: *Realizations and Indicator Populations*

To demonstrate the differences in simulations with and without zones, fifty conditional simulations were calculated for each model assumption (single and two-zone). When evaluating results, it is reasonable to expect that the calculated data distribution should approximately reproduce the original data sample population. Several sets of realizations shown here (Figures 5.23 through 5.28) demonstrate that modeling with two zones yields significantly different results than modeling with a single zone. With a single zone, the character of individual indicators is fairly uniform throughout the site. The Fountain Formation should primarily exhibit indicators 4-7, but the realization indicator populations poorly reproduced the initial data population when a single zone was used (Figures 5.24a, 5.26a, and 5.28a vs. Figure 5.19a). Using two-zones, the Lyons Formation is dominated by indicators 1-4 and 7, and the indicator populations for each two-zone model reasonably reproduce the field data distributions (Figures 5.24b-d, 5.26b-d, and 5.28b-d vs. Figure 5.19a-c). This is not proof, but it strongly suggests that the two-zone realizations are better approximations than the single-zone realizations.

5.3.3.4: *Minimizing Uncertainty*

Another approach to comparing the single and two-zone realizations, is to evaluate which series produces the smallest uncertainty. This is done visually and graphically in two steps. First, the probability that a particular indicator will occur at any specific location is calculated; and second, based on these calculations, the maximum probability any individual indicator will occur in a particular cell is calculated. The first series of maps are useful for evaluating where a particular indicator is likely to occur, and the second series is useful for identifying areas of the site where the modeler can be relatively certain or uncertain about the model results.

The maps in Figures 5.29 and 5.30 were developed by combining the results of 50 realizations into individual indicator probability maps. The difference resulting from using a single zone and two zone model was most pronounced with indicators #1 and #6. From the maps, it can be seen that the single zone realizations do not identify a transition of indicator frequency across the site; the uncertainties are fairly uniform except immediately near conditioning data. In the two-zone realizations, there are distinct differences. There is a high frequency of indicator #1 in the Lyons Formation (indicator #1 identifies two conglomerate Lyons Formation facies), consequently most cells have a relatively high probability of being indicator #1. Indicator #6, is rare in the Lyons Formation, thus its relative probability of occurrence to the Fountain Formation is low.

Ideally zonal kriging will yield more definitive results. The maps in Figure 5.31 indicate the maximum probability any particular indicator will occur in each cell. Visually, for the single-zone simulations, the only areas of low uncertainty are near the hard and soft conditioning data; in the two-zone model though, the improved definition of indicator #1, has reduced much of the uncertainty (green areas indicate much lower uncertainty than blue areas) in the Lyons Formation (right). Comparing the histograms from the single and two-zone simulations (Figures 5.32a and 5.32b), the two-zone model have fewer low probability cells, and substantially more mid-probability and high probability cells. The two-zone model exhibits a bi-modal distribution of uncertainty, due to separate populations from the two formation zones (Figures 5.32c and 5.32d).

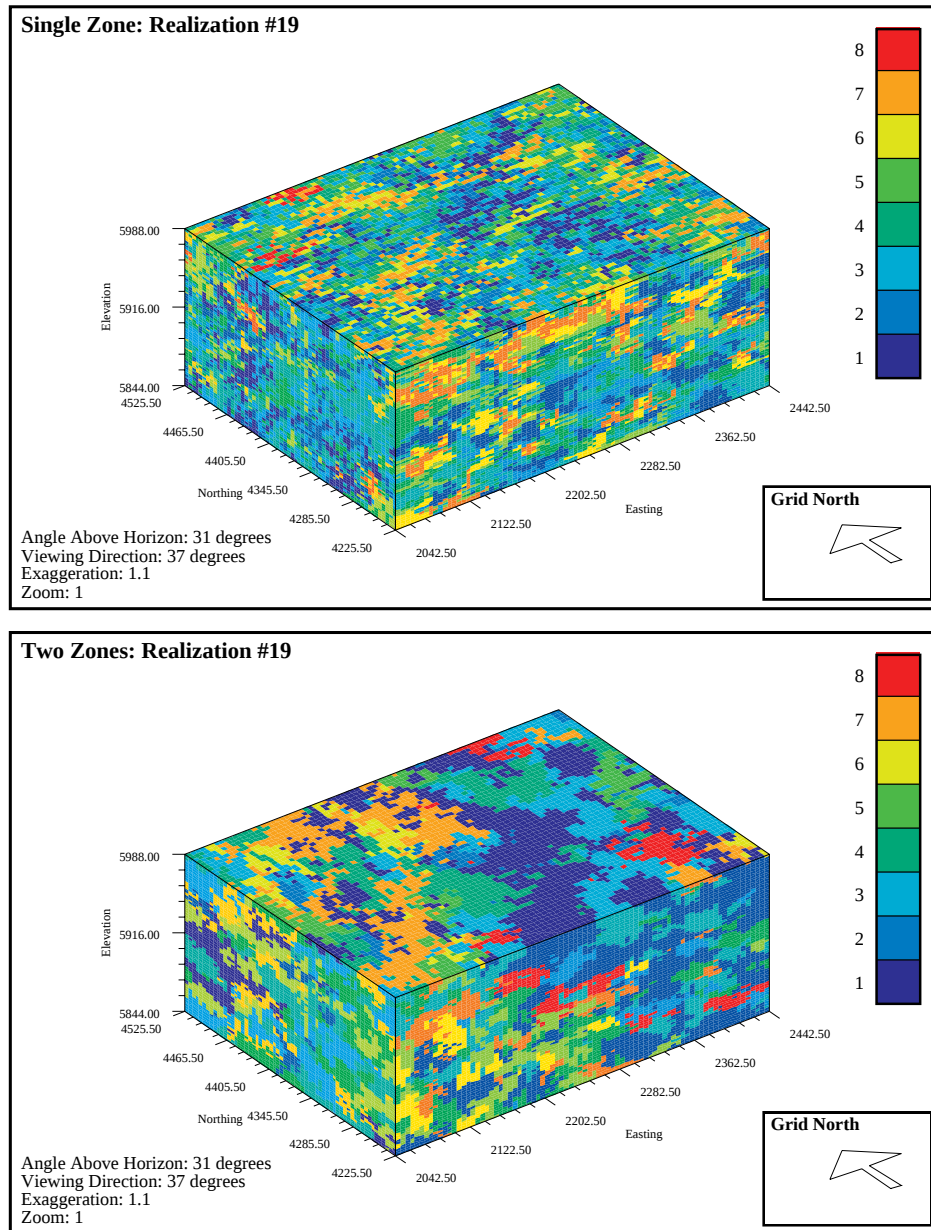


FIGURE 5-23. Single-zone and two-zone realization pair #19.

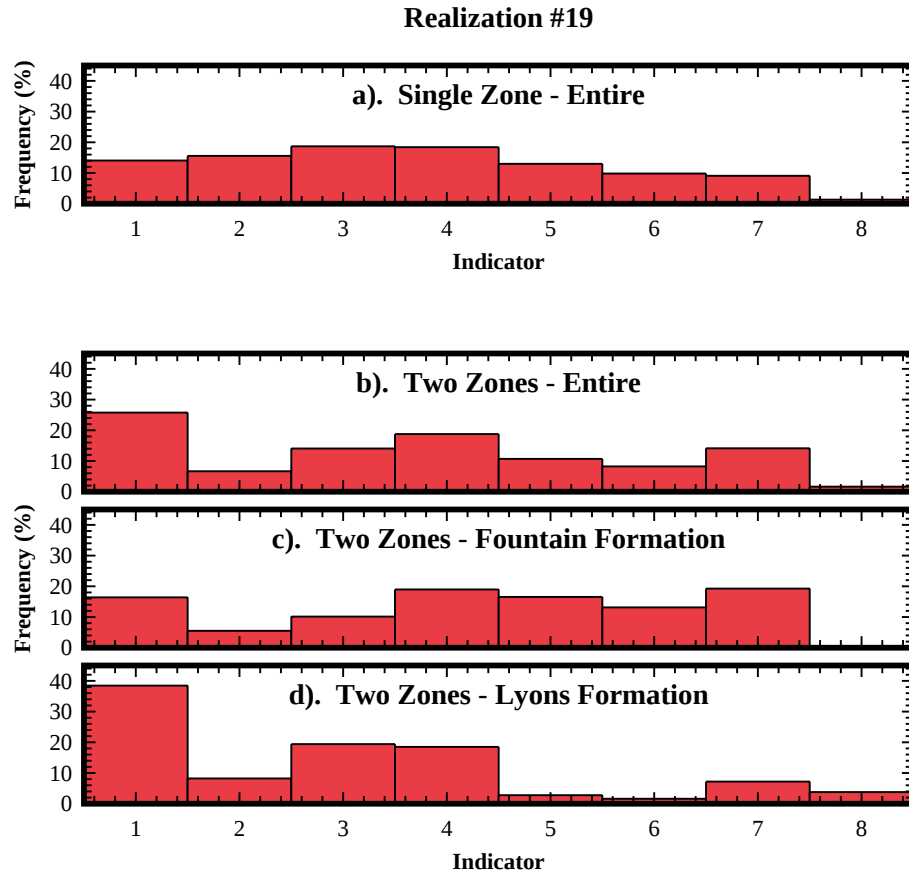


FIGURE 5-24. Distribution of indicators a) in the single-zone realization #19 (Figure 5.23a) and b-d) in two-zone realization #19 (Figure 5.23b). The single-zone realization poorly reproduces original data distribution (Figure 5.19a), whereas the two-zone realization reasonably reproduces the full and individual formation distributions (Figure 5.19a-c).

From these histograms (Figure 5.32) we can see that most of the model uncertainty reduction is due to the improvement in the definition of the Lyons Formation. The Fountain Formation uncertainty is slightly less than occurs when the entire site is modeled with a single-zone. These histograms imply that there is less uncertainty in the two-zone model which suggests the realizations have improved. Without further exploration or groundwater flow modeling, however, it is not possible to confirm this conclusion.

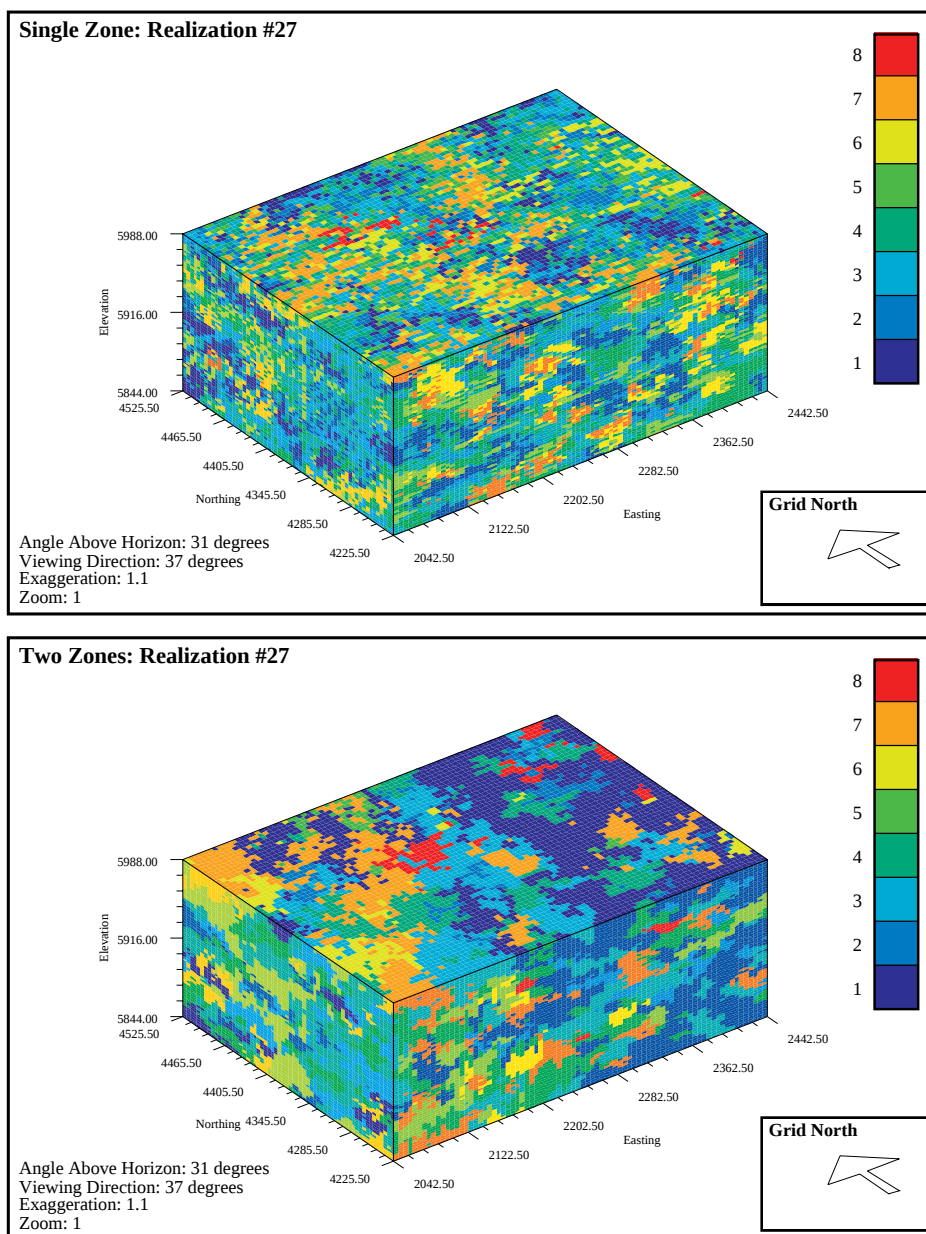


FIGURE 5-25. Single-zone and two-zone realization pair #27.

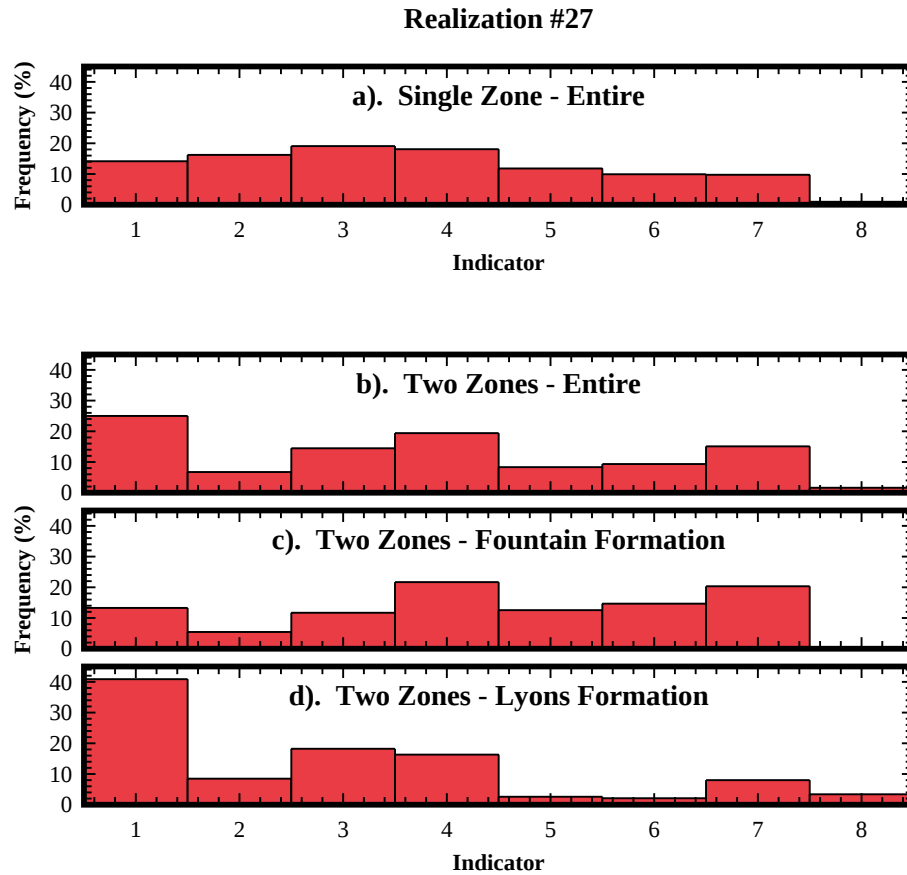


FIGURE 5-26. Distribution of indicators a) in the single-zone realization #27 (Figure 5.25a) and b-d) in the two-zone realization #27 (Figure 5.25b). The single-zone realization poorly reproduces original data distribution (Figure 5.19a), whereas the two-zone realization reasonably reproduces the full and individual formation distributions (Figure 5.19a-c).

5.4: Steps to Determine if Zonal Kriging is Appropriate

Before using zonal kriging, it is important to determine whether the statistics of the data suggest that zonation is appropriate. Several items to consider are whether:

- a visual display of the field geologic data suggests distinct zones.
- the full data set exhibits a bi/multi-modal population frequency distribution.
- there are statistical differences between the populations in each suspected zone.

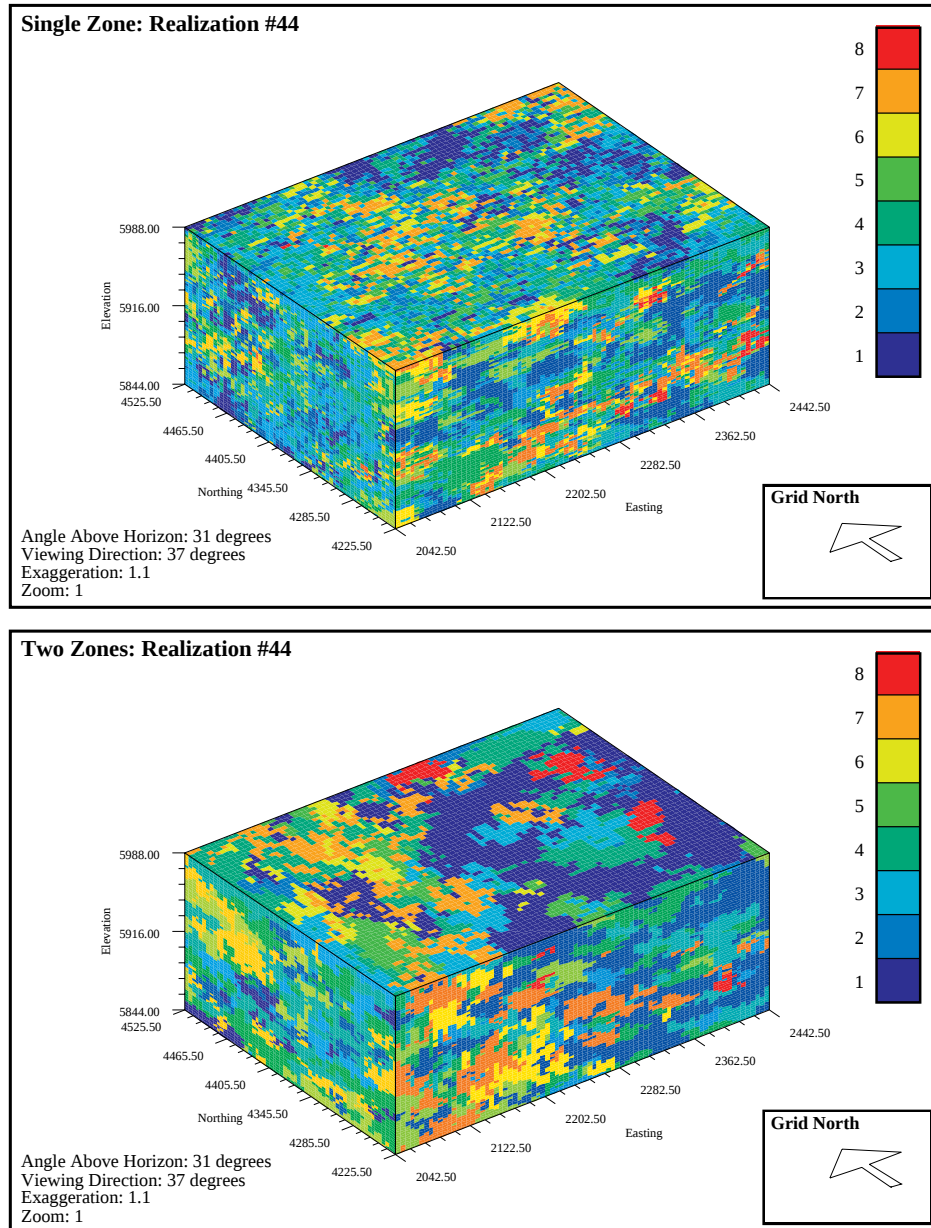


FIGURE 5-27. Single-zone and two-zone realization pair #44.

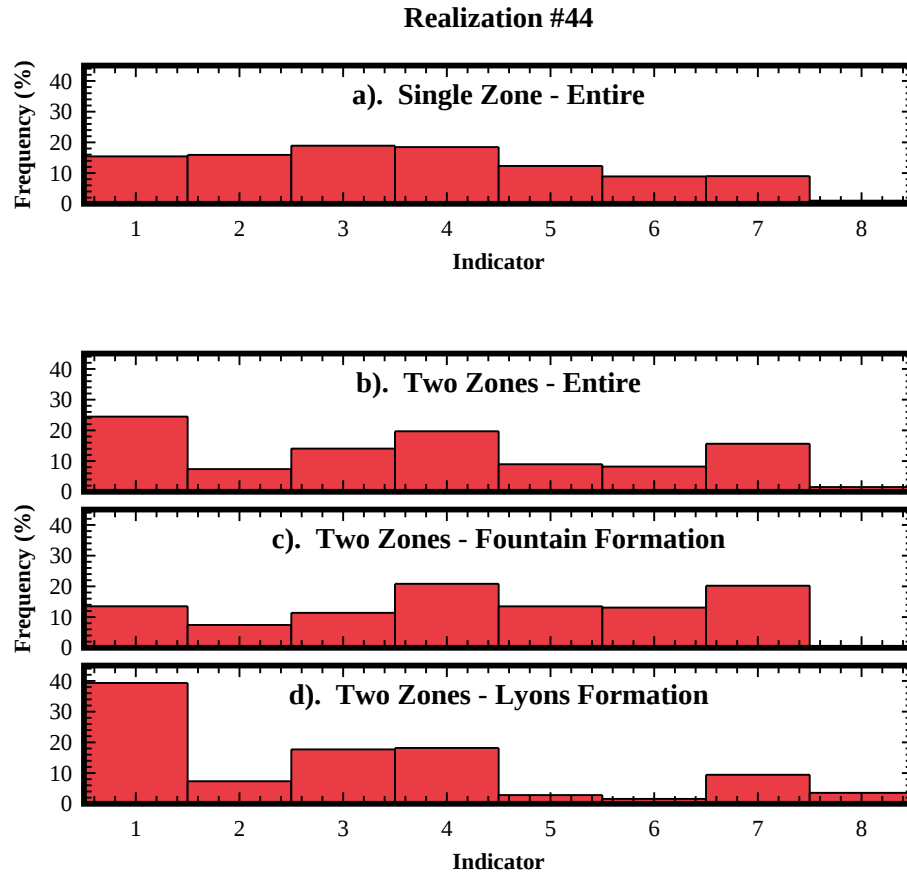


FIGURE 5-28. Distribution of indicators a) in the single-zone realization #44 (Figure 5.27a) and b-d) in the two-zone realization #44 (Figure 5.27b). The single-zone realization poorly reproduces original data distribution (Figure 5.19a), whereas the two-zone realization reasonably reproduces the full and individual formation distributions (Figure 5.19a-c).

- the frequency distribution between the populations in each of the suspected zones varies significantly.
- the spatial statistics (semivariogram models) vary significantly between zones.

If these conditions exist, consider using zonal kriging.

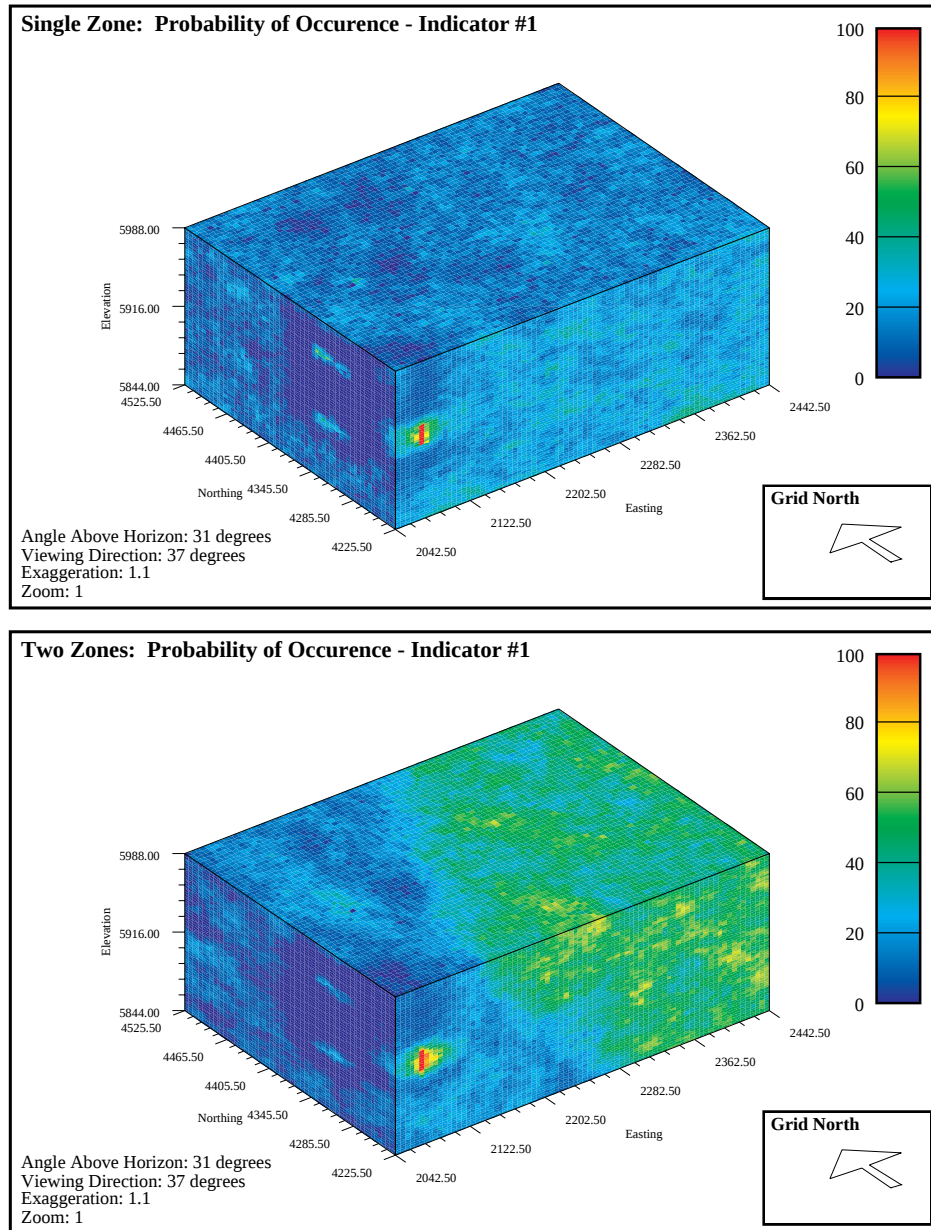


FIGURE 5-29. Maximum probability of occurrence of indicator #1 for single-zone and two-zone realizations. In the single-zone model, the uncertainty is fairly consistent throughout the site except near hard and soft data locations. There are significant differences between the zones in the two-zone map. The Lyons Formation has a much higher percentage of Indicator #1, and this is reflected in the maximum probability map.

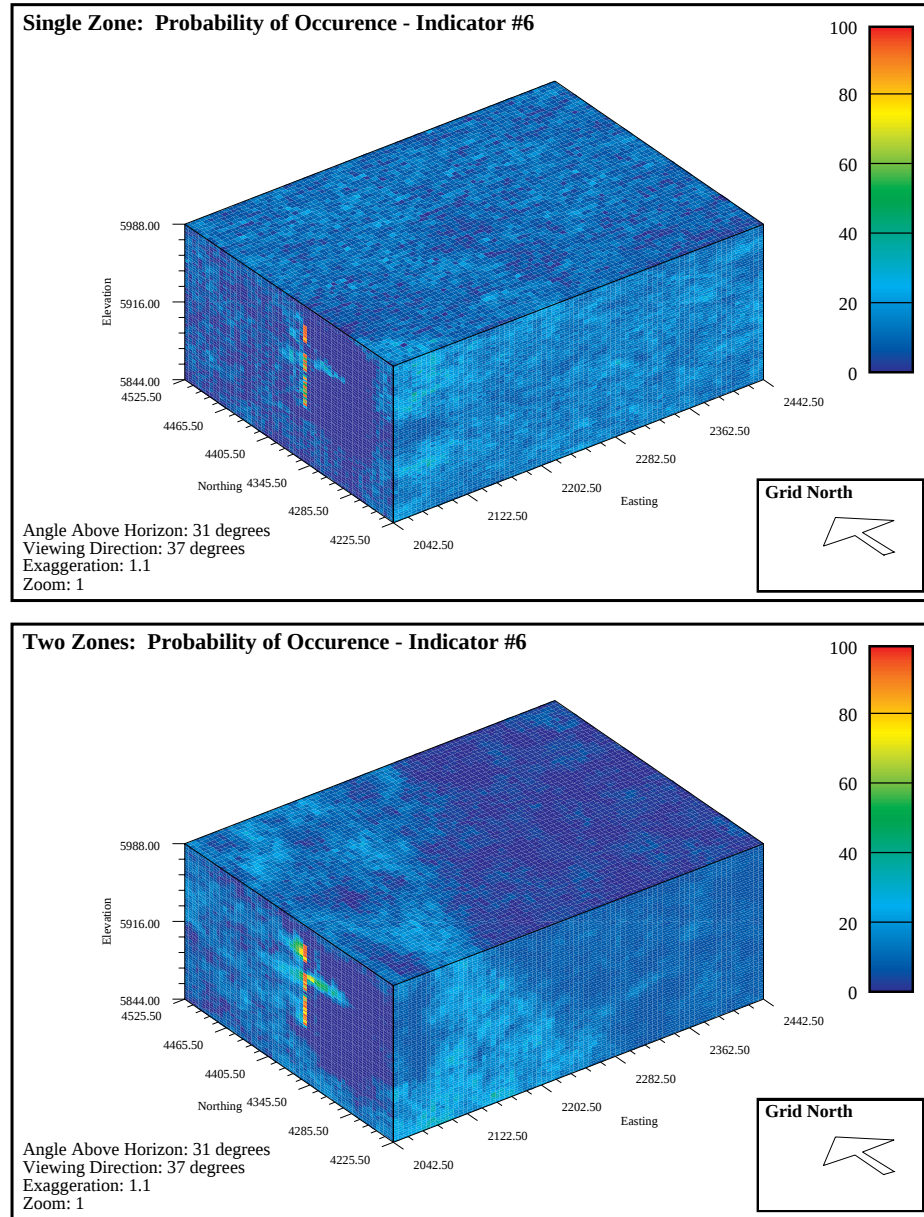


FIGURE 5-30. Maximum probability of occurrence of indicator #6 for single-zone and two-zone realizations. Note, in the single-zone model, the uncertainty is fairly consistent throughout the site except near hard and soft data locations. There are significant differences between the zones in the two-zone map. The Lyons Formation has a much lower percentage of Indicator #6, and this is reflected in the maximum probability map.

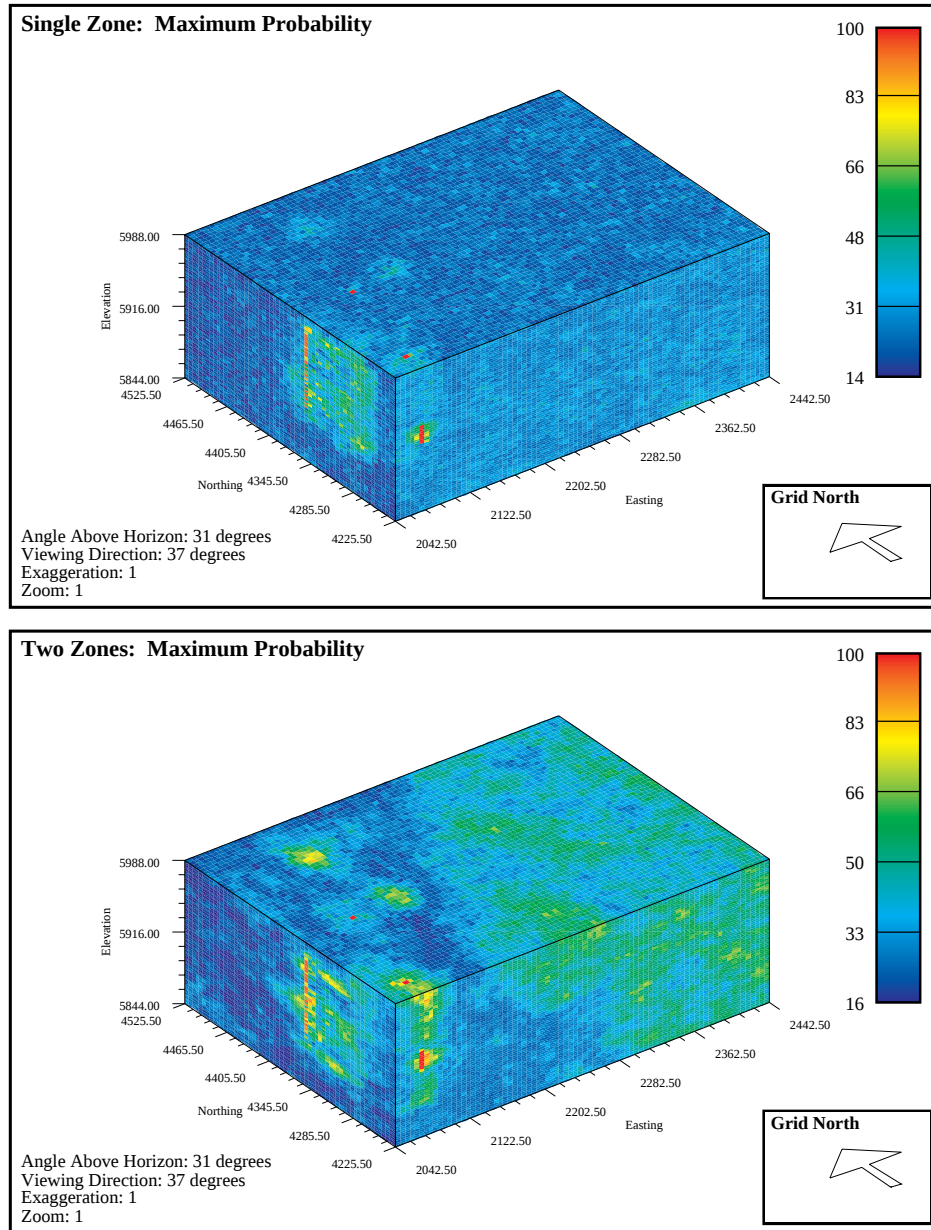


FIGURE 5-31. Maximum probability of occurrence of any indicator. In the single-zone model, uncertainty is fairly uniform across the site except near hard and soft data locations. In the two-zone model, the Lyons Formation exhibits significantly lower uncertainty. These maps are useful for identifying the spatial distribution of uncertainty.

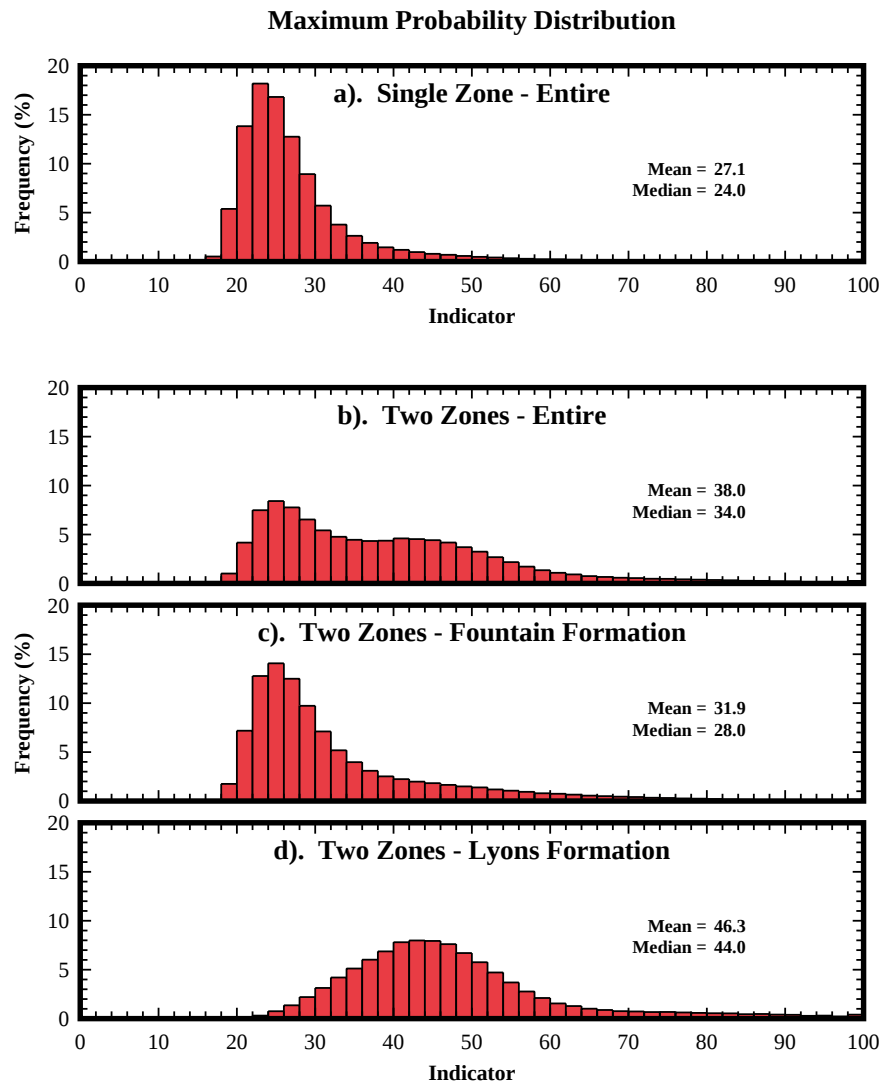


FIGURE 5-32. Distribution of the maximum probability of occurrence of any indicator. The two-zone model (b) has fewer low probability cells and more mid-probability cells than the single-zone model (a). This implies the two-zone model has less uncertainty, and is therefore a better solution. The two-zone histogram also has a bi-modal distribution (b), and the populations may be separated by formation (c and d). Most of the model improvement comes from improved definition of the Lyons Formation.

5.5: Conclusions

Through a series of examples, it has been shown that zonal kriging can yield significantly different results than those obtained using SK or MCIS alone. At sites where the assumption of stationarity is not valid, correctly applied zonal kriging produces realizations that more accurately represent site conditions with greater certainty. The technique requires additional data processing to define the model, and unusual boundary effects may occur when sample data are sparse or are located at spacings near the range of the semivariogram models. These shortcomings, however are offset by the increased certainty, improved accuracy, and modeling flexibility.

UNCERT: GEOSTATISTICAL, GROUND WATER MODELING, AND VISUALIZATION SOFTWARE

UNCERT is an uncertainty analysis, geostatistical, ground water modeling, and visualization software package. It was developed for evaluating the uncertainty associated with the characterization and prediction of subsurface geology, hydraulic properties, and the migration of hazardous contaminants in groundwater flow systems. The package is well suited for evaluating hazardous waste sites and evaluating remediation methods, but it also includes general modules which are usable by researchers from a wide range of disciplines.

6.1: Introduction

UNCERT is a collection of software program modules designed to work together to aid ground water modelers through the data analysis, site modeling, and site evaluation processes. A flow chart of the basic UNCERT processes is shown in Figure 6.1. Many of the tools are also applicable to scientists and engineers from many other fields such as mining, oil exploration, meteorology, and criminology.

6.2: Previous Work

Much of the programming in UNCERT is original work, some of which has been described in previous chapters, but some of the codes are taken from or based on previous work. Software code from other sources, was either 1) public domain, 2) offered with unrestricted use such that copyright notice remained intact and was referenced, 3) transferred to me with permission from the author or the authors agent, or 4) a previous code/algorithm was used as a reference, and original code was written. These resources are described in Appendix G of the UNCERT User's Manual (Appendix A, CD-ROM) and are summarized below:

Project Design: Indicator Kriging and Stochastic Simulation Analysis Flowchart

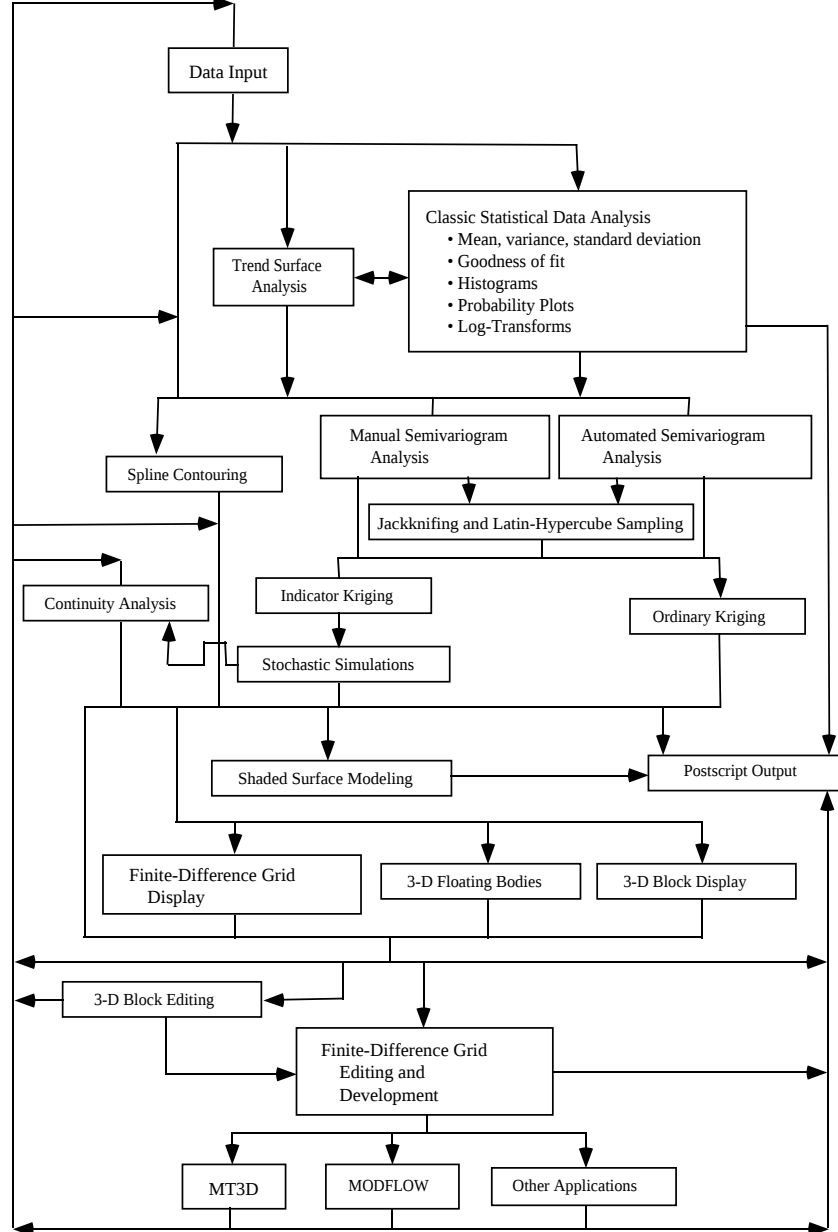


FIGURE 6-1. Detailed flow chart of uncertainty analysis software package.

- Three-dimensional rotations and transformations for 2D and 3D visualization. The algorithms are based on work by Foley, et. al. (1984).
- MODFLOW (McDonald and Harbaugh, 1988).
- MT3D(Zheng, 1990).
- Simple and ordinary kriging algorithm ktb3d (Deutsch and Journel, 1992).
- Irregularly spaced data semivariogram algorithm gamv3 (Deutsch and Journel, 1992).
- Regularly gridded data semivariogram algorithm gam3 (Deutsch and Journel, 1992).
- Indicator conditional simulation program sisim3d (Gómez-Hernández and Srivastava, 1990; McKenna, 1994).
- Contouring and spline algorithm (Wessel and Smith, 1991).
- Rotated text (Richardson, 1993).
- Text editor/viewer program editor (Heller, 1991).
- HTML help browser (NCSA, 1993; Punin, 1994).
- Linear algebra routines from LINPACK (Dongarra, Bunch, et al., 1984).

6.3: Platform Support

In selecting a development platform for UNCERT, two main issues were of concern; portability and processing power. To balance these two issues, UNIX computers, ANSI-C compilers, the Postscript printer language, and the X-windows/Motif graphical user interface was selected. UNIX computers were selected because they 1) had the processing power required by many of the tasks in UNCERT, 2) supported true multi-tasking, and 3) at the beginning of the project, offered one of the best, high resolution graphical work environments. ANSI-C was selected for its 1) computational efficiency, 2) structured programming capability, and 3) portability between platforms. Some programs and code segments were left in FORTRAN, because translating these sections would serve little purpose, cost time, and potentially introduce software errors. Postscript was selected for printer output, because it is 1) a non-hardware dependent printer language and 2) a standard in the UNIX environment. X-windows (developed at MIT) and Motif (developed by the OSF (Open Software Foundation)) were selected as the window manager interface because they have become an industry standard on UNIX computer systems. As a standard, software developed on one system, is easily ported between platforms.

UNCERT has currently been tested on, and is running on eight different 32-bit UNIX platforms (Data General, Dec, IBM RISC-6000, HP, Linux, SGI, Solaris, and SUN) using native ANSI-C and FORTRAN compilers, and gcc (a public domain ANSI-C compiler, (GNU, 1995)) and f2c (a public domain FORTRAN77 to ANSI-C converter (Feldman, Gay, et al., 1990)).

6.4: *UNCERT Modules*

Below are descriptions of each of the modules in UNCERT and a brief discussion of the mathematical techniques used. For a complete description of each module, refer to the UNCERT User's Manual (Appendix A: Attached Tape). This is an HTML (Hyper-Text Markup Language, NCSA (1993)) document, viewable on a wide range on freeware and commercial World Wide Web (WWW) browsers available on Microsoft-Windows, UNIX, and Macintosh platforms. Current versions of UNCERT and the User's Manual may also be viewed or downloaded from the WWW, from:

<http://uncert.mines.edu/>

or by using anonymous ftp from:

<ftp://uncert.mines.edu/pub/uncert/manual/>

Listed below is a brief description and summary of the features of each UNCERT module.

6.4.1: Mainmenu

The mainmenu module is a simple user interface to execute the different modules in the UNCERT software package. It is designed to be a user friendly interface, so that user's can progress through the software to evaluate their field data, and to model the site of concern.

Currently mainmenu is a very simple interface which allows the user only to execute the different software modules within UNCERT. As it stands now, mainmenu is used mainly as a convenience in executing software which the user may not be familiar with, and allows the user to minimize working in the UNIX command line environment. It is a simple attempt to bring the entire UNCERT package together into a unified, windows based environment. It is not recommended that the user try to use this interface exclusively. A great deal of functionality in the software would be lost, trying to do so.

6.4.2: Plotgraph

The plotgraph application is used for plotting two-dimensional X-Y graphs. The application allows the user to plot lines, points with various symbols, and calculate regression lines (up to a tenth order polynomial). The data can be plotted using normal, semi-log, and log-log axes.

6.4.3: Histo

The histo application is used to calculate and display univariant statistics for different data sets. For a sample population, histo can be used to calculate basic statistics such as the mean, standard deviation, and variance. It may also be used to display the behavior of several different populations at once using stacked histograms, cumulative distribution plots, probability plots, and box and whisker plots.

6.4.4: Distcomp

The distcomp application calculates all of the statistical information calculated by histo, but focuses on how sample populations vary between different data sets. For a multiple sample population distcomp can be used to calculate basic statistics such as the mean, and variance. It may also be used to display the behavior of several different populations at once using stacked histograms, cumulative distribution plots, probability plots, P-P (probability vs. probability) plots, and Q-Q (quantile vs. quantile) plots.

6.4.5: Vario

The vario application is used for calculating one- and two-dimensional experimental semivariograms for scattered and regularly gridded data. The package is not limited to the classic semivariogram but will also calculate covariances, madograms, rodograms, cross-semivariograms, etc. Jackknifing the sample data set is also an option. Three types of soft indicator data can be used with hard data to calculate spatial continuity. The application displays the measure of covariance ($\gamma(h)$) versus lag.

6.4.6: Variofit

The variofit application is used to fit model semivariograms to experimental and jackknifed experimental semivariograms (generated by vario). This can be done manually or automatically using least-squares regression or latin-hypercube sampling techniques. Ergodic variations of the model semivariogram from simulation series may also be evaluated.

6.4.7: Grid

The Grid module interpolates parameter values at locations where there are no physical data. This is done using various interpolation algorithms (inverse-distance, kriging, trend-surface analysis) based on irregularly spaced data. Sometimes it is of interest to estimate what is occurring between data locations. For other applications, for convenience, or for clarity, irregularly spaced data must be interpolated onto a regular grid. For example contour, surface, and block require that the data being viewed be gridded with a rectangular pattern. These programs then allow the user to visually view the interpolated estimate of the field data. Grid is used to interpolate values at locations of convenience based on field data.

Within grid there are several gridding algorithms; inverse-distance, simple and ordinary kriging, and trend-surface analysis. Inverse-distance is a relatively simple method which estimates the value of a location based on the distance and value of surrounding sample data points. Kriging does much the same thing as inverse-distance, except kriging also considers spatial statistics describing how the field data vary spatially. Kriging is often referred to as the best unbiased estimator for evaluating a value at a given location. Trend-surface analysis is a least-squares regression technique which assumes the data values are a function of a “regional” trend with minor “local” variations. The calculated trend-surface attempts to describe the “regional” component.

6.4.8: Contour

The contour application contours and performs gradient analysis for two- and three-dimensional, regularly gridded data. Only two-dimensional views are possible however. Three-dimensional data sets can be viewed along X-Y, X-Z, and Y-Z planes. Profile lines along the contoured surface can also be plotted.

6.4.9: Surface

Surface is a 2-1/2 dimensional visualization program for viewing regularly gridded data as a color contoured, gradient, or shaded relief surface. It is included in the UNCERT software as a tool to view two-dimensional grids as three dimensional surfaces. This is referred to as a 2-1/2 dimensional surface because for each X-Y grid location, there is only one Z value. In a true 3D model (see block) each X-Y location may have multiple Z values. This package is used to view gridded surface data generated from grid, or for examining layers or cross-sections from sisim, MODFLOW (McDonald and Harbaugh, 1984), and MT3D output files. Surface may also be used to display any regularly gridded data from other sources (some data file format manipulation maybe required); DEM's (Digital Elevation Model's) are an example.

6.4.10: Block

Block is a 3-dimensional visualization program for viewing regularly gridded data or scattered data points and lines (both cannot be viewed at the same time). It is included in the UNCERT software as a tool to view the values in three-dimensional grids as three dimensional blocks. This package is used to view gridded block data generated from grid or for examining output from sisim, modmain, and mt3dmain.

6.4.11: Sisim & Sisim3d

Sisim is a graphical user interface (GUI) for sisim3d, an indicator kriging and conditional stochastic simulation program for discrete data (non-continuous data: e.g. clay, sand, gravel) developed at Stanford University by Gómez-Hernández and Srivastava (ISIM3D, 1990) and modified at the Colorado School of Mines by McKenna (1994) to utilize soft data. Up to eight indicators can be modeled in a single simulation. In its basic form sisim3d can be awkward to use, particularly when many simulations are required based on varying semivariogram models. This interface assists the user in handling data files, input parameters, coordinating multiple simulations, tasking jobs to other computers, calculating simulation statistics, and visualizing results.

6.4.12: Modmain

Modmain is a graphical user interface for MODFLOW, the MODular three-dimensional, finite difference FLOW model developed by the United States Geological Survey (McDonald and Harbaugh, 1988). MODFLOW is a program designed to model ground water flow and heads (pressure and elevation) in confined and unconfined aquifer systems. In its basic form,

MODFLOW can be difficult, or awkward to use. The modmain program module is designed to simplify data entry, model editing, and analysis of results.

6.4.13: Mt3dmain

Mt3dmain is a graphical user interface for MT3D, a modular three-dimensional transport program (Zheng, 1990). MT3D is a program designed to model contaminant transport based on a pre-solved ground-water flow model (MODFLOW is often used to solve the ground water flow equations. MT3D uses the solution aquifer heads to base the transport results). In its basic form, MT3D can be difficult, or awkward to use. The mt3dmain program module is designed to simplify data entry, model editing, and analysis of results.

6.4.14: Array

The Array module is used to manipulate mathematically one, two, or a series of block, sisim, contour, or surface 2D or 3D grid files. Depending on the options selected, operations include addition, subtraction, multiplication, division, averaging, minimum, maximum, probability value within a range, reclassification, and basic statistics. These are basic grid tools similar to those used in Geographical Information System (GIS) software. This tool can be useful for data preparation, or for data and result analysis. For example, by reclassifying a contaminant plume map to a cost of remediation map, estimates can be made about site clean up costs.

6.4.15: Utilities

In addition to the main modules, there are also several utility modules: calc, lpr_ps, ps_merge, editor, and xhelp. These are very simple user-aid utilities. Calc is a simple RPN scientific calculator. Lpr_ps is used to print ASCII text files with variable margins, line numbers, and variable font sizes. Ps_merge is used to combine, translate, and scale two UNCERT Postscript files into a single Postscript file. Editor is a simple text editor. It is convenient for viewing wide (132 column) and extremely large files. Xhelp is a simple HTML viewer, though it does not display any graphics figures.

SUMMARY AND CONCLUSIONS

This research presents several new geostatistical methods for modeling subsurface site conditions and a geostatistical ground water modeling and visualization software package. The overall goals of developing these methods and tools are to better define site uncertainty, reduce site uncertainty, or simplify the process of modeling site uncertainty. These methods assist hydrogeologists in defining uncertainty in the ground water flow and contaminant transport modeling process, so that risks can be more accurately accessed and appropriate remediation methods can be better designed. Many of these methods and tools are also applicable to other scientific disciplines.

7.1: Summary and Conclusions

Jackknifing the semivariogram, with small data sets (10's to 100's of samples), can be useful in describing the uncertainty associated with the definition of the model semivariogram. It can be combined with Latin-Hypercube Sampling and conditional indicator simulation to model overall site uncertainty, but its most useful feature may be facilitating quantitative evaluation of when enough data has been collected at the site to sufficiently describe the site spatial variation.

Directional semivariograms more accurately model the spatial variation of the sample data. Instead of defining a single semivariogram model for the principle axis of site variability, and forcing the orthogonal axes of variation to use the same model adjusted with anisotropy factors, each orthogonal axis may be described and modeled independently. This simplifies the modeling process for the modeler, because it is not necessary to compromise by selecting one model that acceptably depicts all orientations. This advantage is partially offset, by the overall computational effort in the kriging process, because overall processing time approximately doubles. This is a reasonable tradeoff, because 1) modeler time is usually more valuable than computer time, and 2) most importantly, results are more accurate.

Class conditional indicator simulation, versus threshold conditional indicator simulation, doesn't appear to reduce model uncertainty nor improve model results, yet, it is a, valuable tool, because; 1)

it is a more intuitive approach to defining and modeling indicator semivariograms; 2) it facilitates testing of model sensitivity to indicator ordering; and 3) it can be argued that, although the technique generates more order relation violations, these may more accurately represent the true number of problem matrices in the solution. In some situations, the threshold method fails to identify problems associated with the matrix solution.

Zonal kriging addresses a common limitation of the geostatistical method at many sites. The sample data at many sites, do not honor the assumption of second-order stationarity, that is, the spatial variation of the data, as described by the experimental semivariogram, varies across the site. The zonal kriging method developed in this research automates some of the methods previously performed manually to address this situation, and adds several new tools to cope with transitions between zones. Again, by allowing the modeler to more accurately describe site conditions, this technique generates more accurate site estimates.

The UNCERT software package incorporates all the methods described above in addition to other statistical and geostatistical methods, ground water flow and contaminant transport models, and visualization tools. This package simplifies the data handling, and the use of complex tools, thus aiding hydrogeologists in site evaluation and remediation design. It is also useful to scientists from other scientific disciplines.

7.2: *Recommendations for Future Work*

There are several aspects of this research that could be further developed. Some relate to how the methods were tested and evaluated, while others relate to limitations of the methods themselves.

First, the methods developed here are tested and evaluated under the premise that if uncertainty is reduced, ground water flow and contaminant transport model results will be more accurately match site conditions. This is a reasonable supposition, but it may not be true. Intuitively, a model with less uncertainty better describes true site conditions. The model results from these techniques should be tested using flow and transport models at controlled sites, to test this premise. The questions to ask are: 1) do these methods produce more consistent and more accurate results, 2) do they help reduce uncertainty in defining contaminant migration pathways, and 3) when inverse parameter estimation methods are implemented, are the sensitivities for the estimated parameters reduced?

Two other issues that should be further researched relate specifically to the class and threshold conditional simulation methods. It is argued that, although the results between the methods are not identical, the methods produce substantially the same results. This conclusion is based on two observations. One is that the differences between the methods are approximately equivalent to the differences using the same method, with a different indicator ordering. Because ordering is arbitrary, theoretically it should not affect model results, however, there were small, local, differences for the limited number of simulations calculated (50 to 200). It would be useful to test and confirm that the differences due to indicator ordering, or class vs. threshold techniques, become smaller as the number of simulations is increased, possibly to 1000 or more. This was beyond the

scope of this research due to limited computing facilities. The second observation related to the model differences, involves the significant differences in methods for resolving the order relation violations in the class and threshold simulations. Some of these differences can probably be eliminated, but they cannot be completely eliminated due to the nature of the two approaches.

Finally the procedure for handling gradational and fuzzy transitions between zones is simplistic, but functional. When a cell is estimated, points from neighboring zones may be used in the kriging calculation. The equation used to describe the spatial variance between the sample point and the cell being estimated, is the semivariogram model for the zone in which the cell is located. This disregards 1) the relative amount of separation distance in each zone, and 2) the possibility that a sharp bounding zone separates the point and the cell. These conditions were ignored for computational efficiency, but most importantly, due to concerns that the kriging matrix could no longer be guaranteed to be positive definite. This last issue however was not explored, but could be investigated in future work.

-
- Ababou, R., A.C. Bagtzoglou and E.F. Wood, 1994, "On the Condition Number of Covariance matrices in Kriging, Estimation, and Simulation of Random Fields." *Mathematical Geology*, Vol. 26, No. 1, pp. 99-133.
- Alabert, F.G., 1987, *Stochastic Imaging and Spatial Distributions Using Hard and Soft Information*. Master's Thesis, Department of Applied Earth Sciences. Stanford, Stanford University.
- Burden, R.L. and J.D. Faires, 1985, *Numerical Analysis*. Boston, Prindle, Weber, and Schmidt.
- Carle, S.F. and G.E. Fogg, 1996, "Transition Probability-Based Indicator Geostatistics." *Mathematical Geology*, Vol. 28, No. 4, pp. 453-476.
- Christakos, G., 1984, "On the Problem of Permissible Covariance and Variogram Models." *Water Resources Research*, Vol. 20, No. 2, pp. 251-265.
- Dagdelen, K. and A.K. Turner, 1996, *Importance of Stationarity for Geostatistical Assessment of Environmental Contamination. Geostatistics for Environmental and Geotechnical Applications*, ASTM STP 1283. . R. M. Srivastava, S. Rouhani, M. V. Cromer and A. I. Johnson, Eds., Philadelphia, American Society For Testing and Materials.
- Davis, B.M., 1987, "Uses and Abuses of Cross-Validation in Geostatistics." *Mathematical Geology*, Vol. 19, No. 3, pp. 241-248.
- Deutsch, C.V. and A.G. Journel, 1992, *GSLIB: Geostatistical Software Library and User's Guide*. New York, Oxford Press.
- Dongarra, J., J. Bunch, C. Moler and P. Stewart, 1984, *LINPACK*.
- Englund, E. and A. Sparks, 1988, *GEO-EAS*. U.S. Environmental Protection Agency, Environmental Monitoring Systems Laboratory, EPA/600/4-88/033.
- Environmental Science and Engineering, 1987, *Rocky Mountain Arsenal Water Quantity/Quality Survey Final Initial Screening Program Report. Part I of III*, Prepared for the Rocky Mountain Arsenal.

- Feldman, S.I., D.M. Gay, M.W. Maimone and N.L. Schryer, 1990, A Fortran-to-C Converter, Computing Science Technical Report No. 149, Bell Communications Research, and Carnegie Mellon University.
- Fogg, G.E., 1986, Stochastic Analysis of Aquifer Interconnectedness, with a test case in the Wilcox Group, East Texas. Ph.D. dissertation. University Microfilms International, Ann Arbor, MI., University of Texas at Austin.
- Fogg, G.E., 1989, Stochastic Analysis of Aquifer Interconnectedness: Wilcox Group, Trawick Area, East Texas. Bureau of Economic Geology (Texas), University of Texas at Austin, Report of Investigations No. 189.
- Foley, J.D. and A.V. Dam, 1984, Fundamentals of Interactive Computer Graphics. Reading, MA, Addison-Wesley Publishing Co.
- GNU, 1995, gcc. Boston, MA, Free Software Foundation , Inc.
- Gómez-Hernández, J.J. and R.M. Srivastava, 1990, "ISIM3D: An ANSI-C Three Dimensional Multiple Indicator Conditional Simulation Program." Computers in Geoscience, Vol. 16, No. 4, pp. 395-440.
- Harding, Lawson and Associates, 1992, Groundwater Monitoring Program, Final Annual Groundwater Monitoring Report for 1991, Vol. I of II, Prepared for the Rocky Mountain Arsenal.
- Heller, D., 1991, Motif Programming Manual For OSF/Motif Version 1.1 (Motif Edition). Sebastopol, California, O'Reilly & Associates, Inc.
- Isaaks, E.H. and R.M. Srivastava, 1989, An Introduction to Applied Geostatistics. New York, Oxford University Press.
- Johnson, A.D., 1995, A Geostatistical Evaluation of the Unconfined Aquifer Underlying the Rocky Mountain Arsenal, Commerce City, Colorado. Master's Engineering Report ER-4537, Department of Geology and Geological Engineering. Golden, Colorado School of Mines.
- Johnson, N.M. and S.J. Dreiss, 1989, "Hydrostratigraphic Interpretation Using Indicator Geostatistics." Water Resources Research, Vol. 25, No. 12, pp. 2501-2510.
- Journal, A.G., 1983, "Nonparametric Estimation of Spatial Distributions." Mathematical Geology, Vol. 15, No. 3, pp. 445-468.
- Journal, A.G. and F.G. Alabert, 1988, Focusing on Spatial Connectivity of Extreme-Valued Attributes: Stochastic Indicator Models of Reservoir Heterogeneities. Society of Petroleum Engineers 63rd technical Conference, Richardson, Texas, SPE Paper 18324, pp. 621-632.
- Journal, A.G. and C.J. Huijbregts, 1978, Mining Geostatistics. London, Academic Press.
- Journal, A.G. and E.H. Isaaks, 1984, "Conditional Indicator Simulation: Application to a Saskatchewan Uranium Deposit." Mathematical Geology, Vol. 17, No. 7, pp. 685-718.

-
- Kushnir, G. and J.M. Yarus, 1992, Modeling Anisotropy in Computer Mapping of Geologic Data. Computer Modeling of Geologic Surfaces and Volumes, AAPG Computer Applications in Geology 1, Tulsa, The American Association of Petroleum Geologists, pp. 75-92.
- McDonald, M.G. and A.W. Harbaugh, 1988, A Modular Three-Dimensional Finite-Difference Ground-Water Flow Model. U.S. Geological Survey, Techniques of Water-Resources Investigation, Book 6.
- McGill, R.E., 1984, An Introduction to Risk Analysis, 2nd Edition. Tulsa, Oklahoma, PennWell Publishing.
- McKay, M.D. and R.J. Beckman, 1979, "A Comparison of Three Methods for Selecting Values of Input Variables in the Analysis of Output From a Computer Code." Technometrics, Vol. 21, No. 2, pp. 239-245.
- McKenna, S.A. and E.P. Poeter, 1994, Simulating Geological Uncertainty with Imprecise Data for Groundwater Flow and Advective Transport Modeling. Simulating Geological Uncertainty with Imprecise Data for groundwater Flow and Advective Transport Modeling. AAPG Computer Applications in Geology. J. M. Yarus and R. L. Chambers, Eds., Tulsa, Association of Petroleum Geologists, pp. 241-248.
- McKenna, S.A., 1994, Utilization of Soft Data for Uncertainty Reduction in Groundwater Flow and transport Modeling. Ph.D. Dissertation T-4291, Department of Geology and Geological Engineering. Golden, Colorado School of Mines.
- Myers, D.E. and A.G. Journel, 1990, "Variograms and Zonal Anisotropies and Noninvertible Kriging Systems." Mathematical Geology, Vol. 22, No. 7, pp. 779-785.
- NCSA, 1993, libhtmlw.a. Champaign, Illinois, National Center for Supercomputing Applications, University of Illinois.
- Olea, R.A., 1974, "Optimal Contour Mapping Using Universal Kriging." Journal of Geophysical Research, Vol. 79, No. 5, pp. 695-702.
- Orr, E.D. and A.R. Dutton, 1983, "An Application of Geostatistics to Determine Regional Ground-water Flow in the San Andreas Formation, Texas and New Mexico." Ground Water, Vol. 21, No. 5, pp. 619-624.
- Posa, D., 1989, "Conditioning of the Stationary Kriging Matrices for Some Well-Known Covariance Models." Mathematical Geology, Vol. 21, No. 7, pp. 755-765.
- Poeter, E.P. and D.R. Gaylord, 1990, "Influence of Aquifer Heterogeneity on Contaminant Transport at the Hanford Site." Ground Water, Vol. 28, No. 6, pp. 900-909.
- Press, W.H., S.A. Teukolsky, W.T. Vetterling and B.P. Flannery, 1992, Numerical Recipes in C, The Art of Scientific Computing. Second Edition, Cambridge University Press.
- Punin, J.R., 1994, ASHE (A Simple HTML Editor) - xhtml 1.3. Troy, New York, Rennseler Polytechnic Institute.

- Rautman, C.A. and A.H. Treadway, 1991, "Geologic Uncertainty in a Regulatory Environment: An Example From the Potential Yucca Mountain Nuclear Waste Repository Site." *Environmental Geology Water Sciences*, Vol. 18, No. 3, pp. 171-184.
- Ravenne, C. and H. Beucher, 1988, Recent Developments of Description of Sedimentary Bodies in a Fluvio Deltaic Reservoir and Their 3D Conditional Simulations, Richardson, Texas, Society of Petroleum Engineers 18310, , pp. 463-476.
- Ravenne, C., R. Eschard, A. Galli, Y. Mathieu, L. Montadert and J.L. Rudkiewicz, 1987, Heterogeneities and Geometry of Sedimentary Bodies in a Fluvio-Deltaic Reservoir. SPE 16753, SPE Annual Technical Conference and Exhibition, Dallas, Texas, , pp. 115-124.
- Richardson, A., 1993, xvertext, mppa3@uk.ac.sussex.syms.
- Shafer, J.M. and M.D. Varljen, 1990, "Approximation of Confidence Limits on Sample Semivariograms From Single Realizations of Spatially Correlated Random Fields." *Water Resources Research*, Vol. 26, No. 8, pp. 1787-1802.
- Strang, G., 1988, *Linear Algebra and its Applications*. Fort Worth, Saunders College Publishing.
- Van Horn, R., 1972, *Geologic Map of the Morrison Quadrangle, Jefferson County, Colorado*, United States Geological Survey, 1:24,000, 7.5 minute.
- Wessel, P. and W.H.F. Smith, 1991, *GMT - The Generic Mapping Tools*, School of Ocean & Earth Science Technology, University of Hawaii.
- Wingle, W.L. and E.P. Poeter, 1992, Evaluation of Uncertainty Associated with Contaminant Migration in Ground Water - A Technically Feasible Approach. *Proceedings of the Fifth NGWA Conference on Solving Ground Water Flow Problems with Models*, , , pp. 687-695.
- Zheng, C., 1990, MT3D, A Modular Three-Dimensional Transport Model. U.S. Environmental Protection Agency and Papadopoulos and Associates, Inc.
- Zhu, H. and A.G. Journel, 1992, *Formatting and Integrating Soft Data: Stochastic Imaging via the Markov-Bayes Algorithm*. *Geostatistics Tróia '92*, Tróia, Kluwer Academic Publishers, *Quantitative Geology and Geostatistics*, Vol. 1, pp. 1-12.

UNCERT AND UNCERT USER'S MANUAL

The UNCERT software package and User's Manual (Version 1.20) are contained on the CD-ROM at the back of the dissertation. UNCERT is still a growing software package though, and I recommend that if you want to use it, you download the most recent version from the ftp site described below. The instructions below are designed for someone downloading the software from the internet using anonymous ftp. Special instructions for retrieving the software from the CD-ROM will be described in *italics*.

A1: Information and Comments:

This is a freeware software package, so there is little technical support. However, we are trying to make this package as useful as possible, and if you have questions or comments contact Bill Wingle at:

e-mail wwingle@mines.edu

or

Department of Geology and Geological Engineering
Colorado School of Mines
Golden, Colorado 80401
(303) 273-3905
(303) 273-3859 (FAX)

If you find software bugs, or better yet, fix software bugs, please contact us so that we can improve future releases.

A1.1: Warranty:

The UNCERT package, the program modules within, and the user's manual are distributed in the hope that they will be useful, but WITHOUT ANY WARRANTY. No author or distributor accepts any responsibility to anyone for the consequences of using them or for whether they serve any particular purpose or work at all, unless stated so in writing by the authors. No author or distributor accepts responsibility for the quality of data generated, nor the damage to existing data. Everyone is granted permission to copy, modify, and redistribute the UNCERT package, but only under the condition that the copyright notice in the software remain intact. The software is provided "as is" without express or implied warranty.

A2: *Hardware / Operating System Requirements:*

The UNCERT software package was written in ANSI C and FORTRAN (very little) using X-windows and motif under the UNIX operating system. Currently the software has been tested on IBM-RISC 6000, Dec Alpha, Data General (so I'm told), HP, Linux, Silicon Graphics (SGI), Sun OS, Sun Solaris, and SCO UNIX workstations with 8-bit color graphics cards. The software was written to be easily ported, and where possible follows ANSI-C standards. To use this software you must have:

1. UNIX operating system computer.
2. ANSI C compiler and a FORTRAN compiler (A FORTRAN to C preprocessor is available upon request). These must be 32-bit compilers. The new 64-bit compilers break the software.
3. X-windows Release 4+ and motif Release 1.1+ window manager.
4. X-windows and motif development packages (used at compile time to build X-windows/motif interfaces).
5. 8-bit color graphics card (allows 256 colors simultaneously).
6. 16 MB's of RAM.
7. 35 MB's of free Hard Disk space.

A3: *Acquiring Software:*

The software may be acquired on the internet, by using one of the following anonymous ftp or http addresses:

`ftp://uncert.mines.edu/`
`http://uncert.mines.edu/`

For anonymous ftp, the UNCERT software is stored in the directory:

`/pub/uncert`

In this directory, releases with executables are available for several UNIX platforms, but only the file:

uncert.ver_#.##.tar.Z

is guaranteed to be current. I have limited access to most platforms, and all versions may not be current. Check the dates on the files. The file “uncert.ver_#.##.tar.Z” contains the full UNCERT release, but no executable files. You will have to compile UNCERT yourself if you retrieve this file. The UNCERT User’s Manual is located in:

/pub/uncert/manual/

There are also several useful files in:

/pub/misc

that may be useful. These include public domain and shareware programs available from other locations on the internet. These versions may not be the most recent, but they may save you time trying to locate them elsewhere. These files include f2c (a FORTRAN to C preprocessor), gcc (an ANSI C compiler. You need a C compiler to build it), gs (a Postscript previewer), xv (a GIF/JPEG viewer), and gzip (a good compression utility). There are other files too.

A typical anonymous ftp session might look like:

```
your prompt > ftp uncert.mines.edu
user name: anonymous
password: (your e-mail address, e.g., wwingle@mines.edu)
ftp > binary
ftp > cd /pub/uncert
ftp > get uncert.ver_1.20.tar.Z
ftp > quit
```

If you want to recover UNCERT from the CD-ROM, there are several steps you must follow. You may need to have root privilege to mount and unmount the CD-ROM.

- 1). *Mount the CD-ROM. On IBM RS-6000 workstations, you must have root permission. Use the following commands:*

```
prompt> su -
prompt> mkdir /cdrom
prompt> mount -v 'cdrfs' -p -r /dev/cd0 /cdrom
```

- 2). *Install UNCERT. The full UNCERT (IBM-RS6000) release is located in the /cdrom/uncert/ directory. All files are directly accessible; they are not tarred or compressed. You can now copy it to your computer. For example:*

```
prompt> cd /usr/local
prompt> cp -pR /cdrom/uncert uncert
```

You do not have to install UNCERT in /usr/local. This is a common location however.

3). Unmount CD-ROM. This will have to be done as root also.

```
prompt> umount /cdrom
```

If you are using a system other than an IBM-RS-6000 , the commands for mounting and unmounting the CD-ROM may vary. Consult your system administrator. You will also have to recompile UNCERT. Instructions are given below for compiling UNCERT on different UNIX systems.

A4: Installation:

Once you have downloaded the UNCERT software there are several steps you need to follow to install UNCERT: 1) Unpack the software, 2) compile all the UNCERT modules (This step can be skipped if you downloaded a version with executable), and 3) set up user accounts.

A4.1: Unpacking the Software:

Once you have downloaded the UNCERT package file, move it to the directory above where you want UNCERT stored (e.g. /usr/local). To unpack UNCERT type:

```
uncompress uncert.tar.Z
tar xvf uncert.tar
```

(On Linux computers use: `gzip -d uncert.tar.Z' instead of uncompress)

This will uncompress and un-tar UNCERT. During the tar process, all files will be put in their appropriate locations. If you get warning that directories cannot be created, you will need to download a script file called mk_uncert_dirs. This script also needs to be executed from the directory above where you want UNCERT stored. After running this command, execute the tar command given above again.

A4.2: Compiling UNCERT:

At this point you should read the README file in the uncert directory.

If you did not download a file with executables, or you have trouble with the executables you did download (e.g. cannot find a shared library ..., etc.), you will have to compile UNCERT. Once the files are unpacked, change directories to the uncert directory, for example:

```
prompt > cd /usr/local/uncert
```

If you are not running on an IBM RS-6000 computer, you will need to setup the Makefile's for each module. This is done using the following command (select the appropriate command based on your system):

```
prompt> set_make ibm: IBM RS6000
prompt> set_make hp: HP
prompt> set_make sun: Sun OS
prompt> set_make sol: Sun Solaris
prompt> set_make sgi: Silicon Graphics
prompt> set_make sco: SCO
prompt> set_make linux: Linux/Slackware
```

If your machine type is not listed, you will probably need to modify each Makefile in the directories:

```
?/uncert/src/*
```

This will mainly involve defining where the X-windows and motif library and include files are located. You may also have to define your C and FORTRAN compilers. Once the Makefile's are correctly defined, type:

```
prompt> build
```

This script will go into each ?/uncert/src directory and try to make each program. This may or may not work. Several things can go wrong.

1. The Makefile's do not have the right libraries specified. See if there is a Makefile specific to your machine (e.g. Makefile.ibm). If there is, copy it to "Makefile." If a correct Makefile does not exist, you may have to determine which libraries are missing.
NOTE: on some computers library order is important.
2. You do not have an ANSI C compiler, or your compiler is named something other than "cc." If you have compiler other than "cc," set the variable "CC" to your compiler name. If you don't have a ANSI C compiler, you can get gcc from our ftp site. gcc is a shareware C and C++ compiler. It may take some effort to compile. Note: our posted version is not the most recent version.
3. You do not have a FORTRAN compiler or your compiler is named something other than "xlf" (or "f77"). If you have a compiler other than "xlf" (or "f77"), set the variable "F77" to your compiler. If you don't have a FORTRAN compiler, you can get f2c from our ftp site. It is a shareware FORTRAN to C conversion program. You then compile the C. Contact me (Bill Wingle) if you have this problem. I'm still working on an instruction set.
4. The FORTRAN compiler does not recognize the -qextname compile option. Delete it. This is an IBM FORTRAN/C compile option.
5. In block, you cannot find su.h, segy.h, libcwp.a, libpar.a, or libsu.a. Remove the -DSU compile option. This is an option to compile SU (Seismic UNIX) which most users probably won't have.

Once you get the Makefile's corrected, you can type "make" in each src directory, or you can type "build" from the ~/uncert directory.

NOTE: When you start to port the code, you must do a make in the ?/uncert/src/Xs directory first. This builds the library libXs.a which most of the programs depend on. You can then move to any of the other src directories and start compiling code.

If you are compiling on a non-supported system, I doubt that you will have to make more than a few changes to get the UNCERT modules compiled. There are a couple of important notes though.

1. Many of the program directories repeat the same files. These are common tool object files that will eventually go into a single library. Unfortunately I can't keep all of the files current as I develop UNCERT, therefore, from directory to directory, files may vary slightly and be incompatible. This means that if you find a problem in a common file, you need to change each file, not copy the fixed file to the different directories.

A4.3: Setting Up User Accounts:

To run the programs correctly, each user will have to have several environment variables defined in their login file (ksh -> .profile, csh -> .cshrc, etc.). If you use ksh, they are defined as follows:

```
export UNCERT=/usr/local/uncert
export PATH=$PATH:$UNCERT/bin

export UNCERT_TMPDIR=/tmp
export UNCERT_HELP_DIR=$UNCERT/help/
export XAPPLRESDIR=$UNCERT/app-defaults/
export WWWVIEWER=xhelp
```

If you use csh, they are defined as follows:

```
setenv UNCERT /usr/local/uncert
setenv PATH $PATH:$UNCERT/bin

setenv UNCERT_TMPDIR /tmp
setenv UNCERT_HELP_DIR $UNCERT/help/
setenv XAPPLRESDIR $UNCERT/app-defaults/
setenv WWWVIEWER=xhelp
```

If you use another shell, you may have to modify the syntax slightly. On some systems XAPPLRESDIR can be replaced with (ksh):

```
export XUSERFILESEARCHPATH=$XUSERFILESEARCHPATH: \
$UNCERT/app-defaults/%N
```

On SGI's you must make this substitution. In general, if your platform supports this option, it is better than XAPPLRESDIR.

You must also define a help browser. If you do not define a browser, you will still have on line text help, but no graphics for figures. We are currently developing the xhelp package, but we

recommend you use netscape (Netscape Communications Corporation) or Mosaic (NCSA). These viewers may be downloaded from:

```
http://home.netscape.com/  
http://www.ncsa.uiuc.edu/SDG/Software/Mosaic/NCSAMosaicHome.html
```

A version a Mosaic (old) may be downloaded from our anonymous ftp site if you do not have a web browser. The ftp site and file are:

```
uncert.mines.edu  
/pub/misc/xmosaic-2.5.tar.gz
```

At this point netscape has more features, but it is a commercial application, though they have been letting educational institutions use unlicensed versions. To define a browser other than xhelp, modify the WWWVIEWER environment variable described above with one of the following commands:

```
export WWWVIEWER=/usr/local/netscape  
export WWWVIEWER=/usr/local/Mosaic  
  
setenv WWWVIEWER=/usr/local/netscape  
setenv WWWVIEWER=/usr/local/Mosaic
```

At this point the UNCERT software should be installed, compiled and ready for use. In order to set some of the environment variables it is suggested that you logout and then login before you try to run the applications.

The data sets used in this research are contained on the CD-ROM at the back of the dissertation. These files are accessible from UNIX systems (see mounting instructions in Appendix A) and from DOS/MS-Windows computers with CD-ROM drives. The files for each chapter are located in the following directories:

Chapter 2: Jackknifing and Latin-Hypercube Sampling

/data/jackknife; synthetic test case

Chapter 3: Variation of the Semivariogram Models With Direction

/data/direct/test; synthetic test case

/data/direct/rma; Rocky Mountain Arsenal test case

Chapter 4: Class vs. Threshold Indicator Simulation

/data/class/test; synthetic test case

/data/class/survey; CSM Survey Field test case

Chapter 5: Zonal Kriging

/data/zone/test; synthetic test case

/data/zone/yorkshire; Yorkshire, England test case

/data/zone/survey; CSM Survey Field test case

Chapter 6: UNCERT

/uncert; full UNCERT distribution

/uncert/html/index.html; on-line UNCERT Users's Manual

NOTE: The filenames follow UNIX naming conventions. On DOS and MS-Windows 3.11 (or older) systems, filenames may be truncated or modified.

Descriptions of each data set are contained in README files in each directory.

