

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.

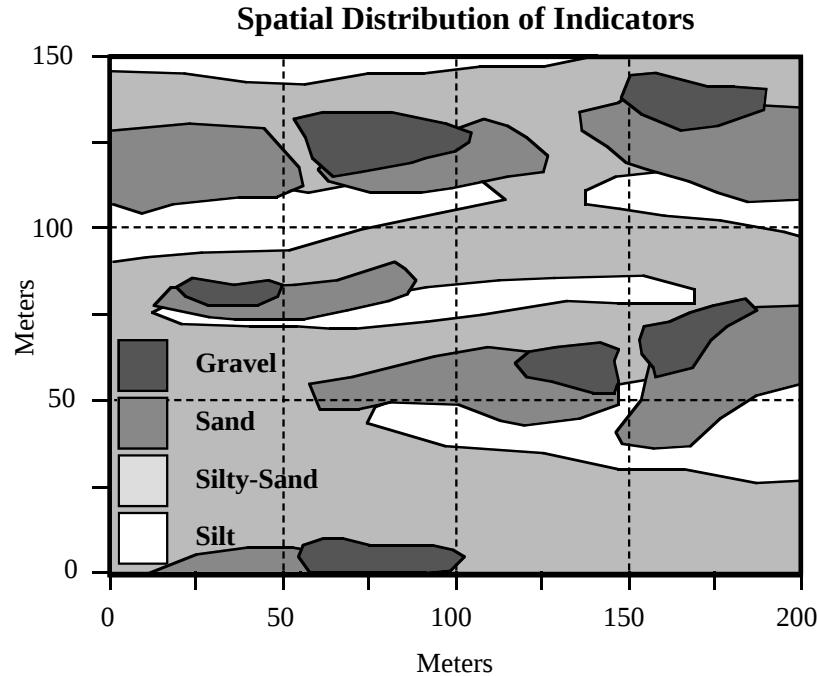


FIGURE 4-1. Spatial distribution of several indicators. Defining semivariograms based on indicator classes is more intuitive, because the range reflects the average size of the indicator bodies. The class semivariogram model ranges are: silt = 112m, silty-sand = 106m, sand = 60m, and gravel = 41m. For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

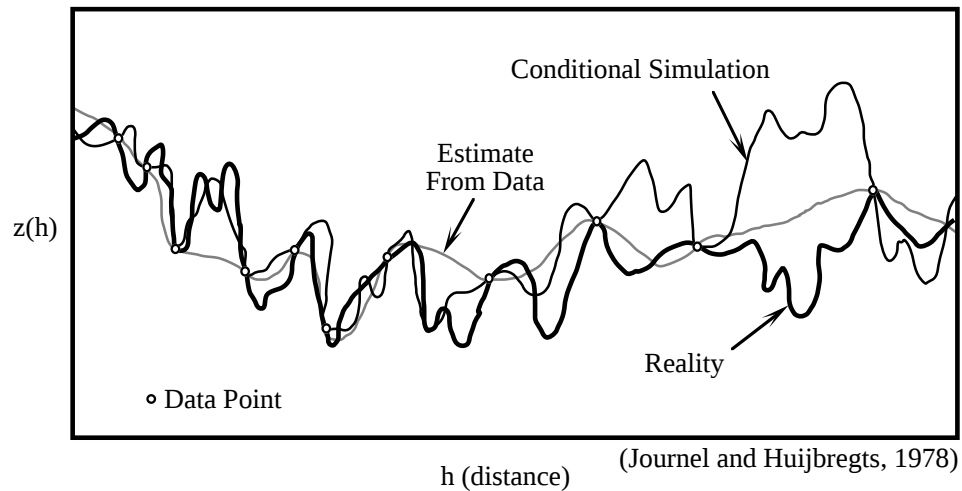


FIGURE 4-2. Ordinary Kriging (and most other estimation methods) tends to average or smooth data to achieve a best linear unbiased estimate (BLUE) of reality. Indicator Kriging with conditional simulation provides a means for modeling the variability observed in nature, while still honoring the field 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

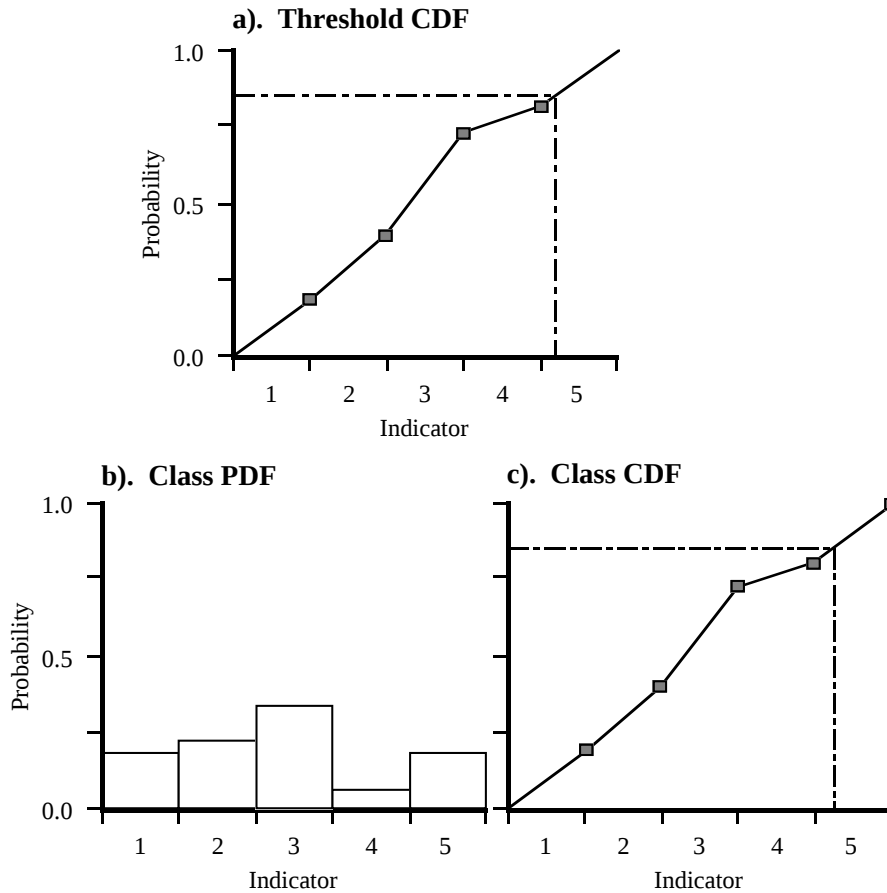


FIGURE 4-3. Class and threshold indicator kriging generate the cumulative density function (CDF) in different ways. Threshold CDF's are determined directly from the probability that the specified grid location is less than each threshold level (a). The final CDF term should be less than 1.0 with the 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

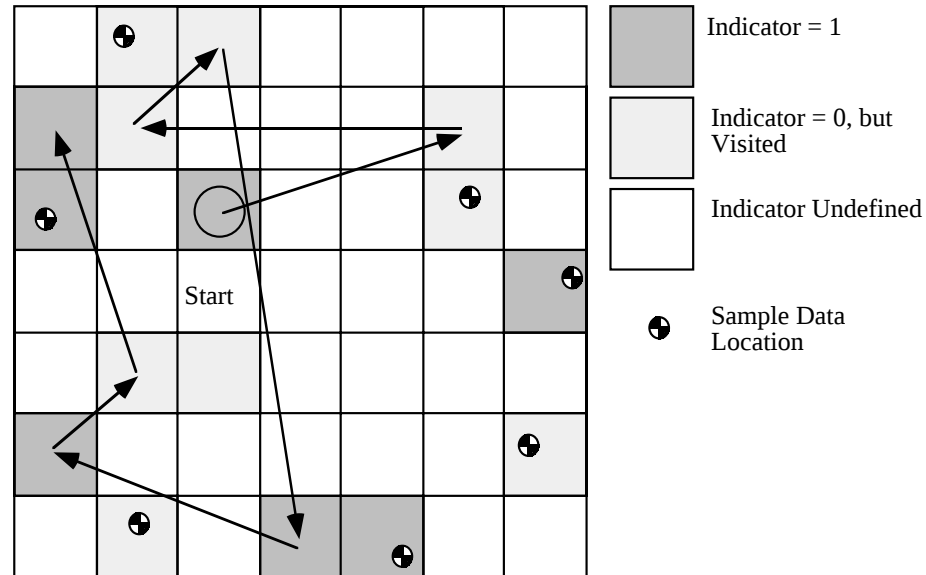


FIGURE 4-4. This illustration shows the step wise manner in which a grid is kriged using Indicator Kriging in conjunction with stochastic simulation. Grid cells containing sample data (hard data and some types of soft data) are defined prior to kriging. Once these points are defined, the remaining cells 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

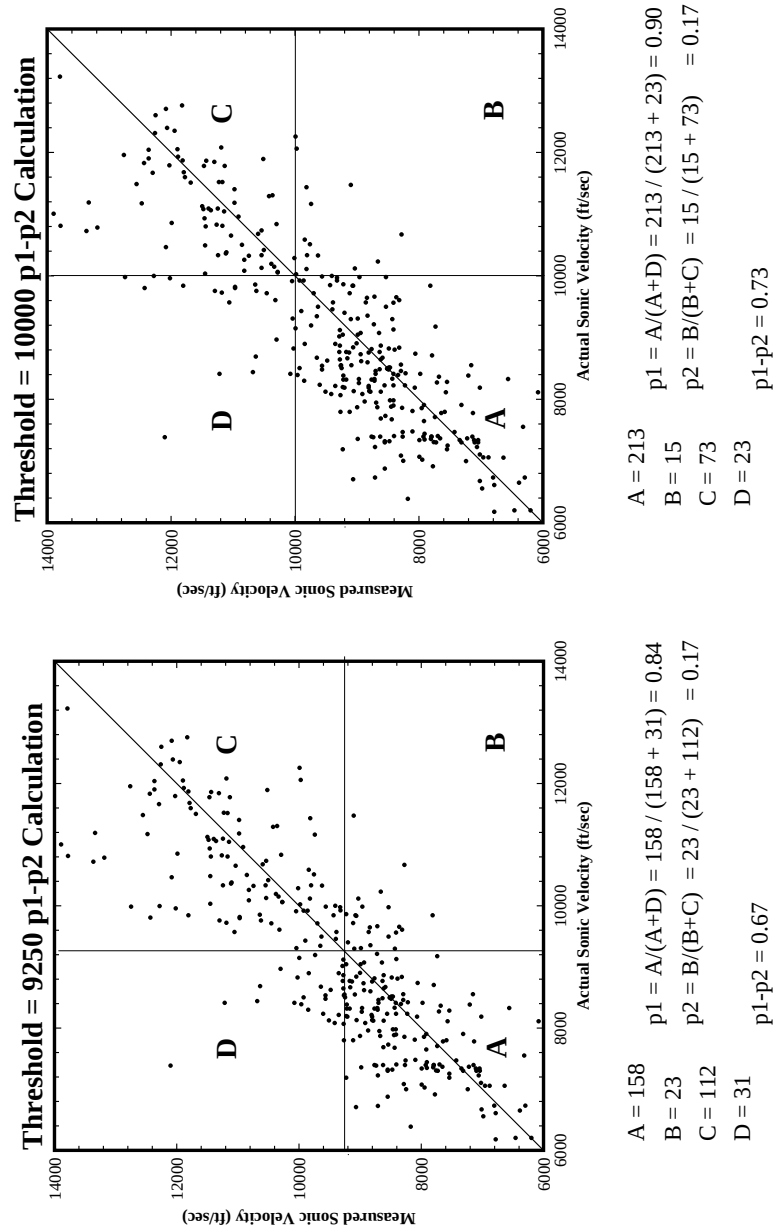


FIGURE 4-5. Graphical method for calculating p_1 and p_2 values for a specific threshold (After Alabert, 1987). Data from CSM Survey Field.

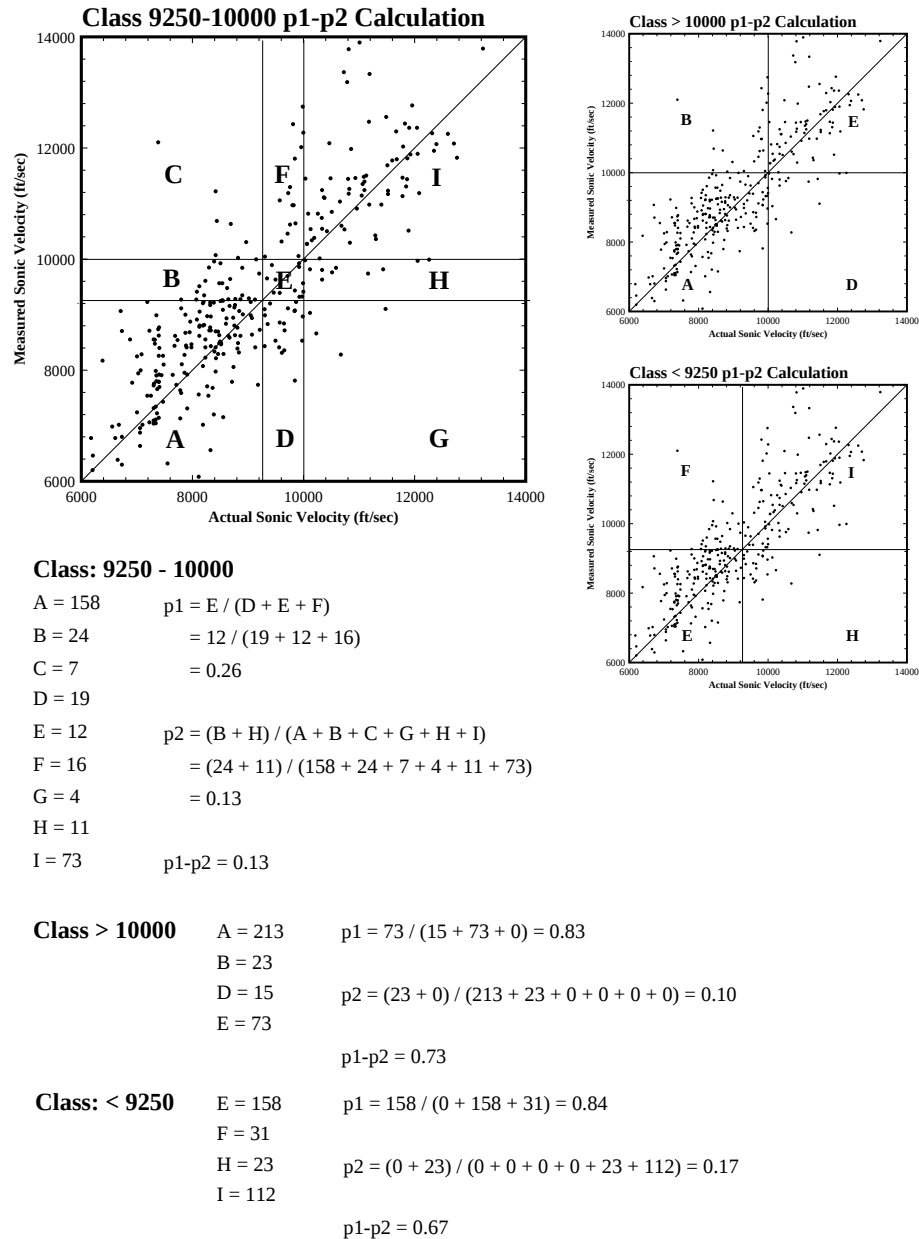


FIGURE 4-6. Graphical method for calculating p_1 and p_2 values for a specific class. Data from CSM Survey Field.

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

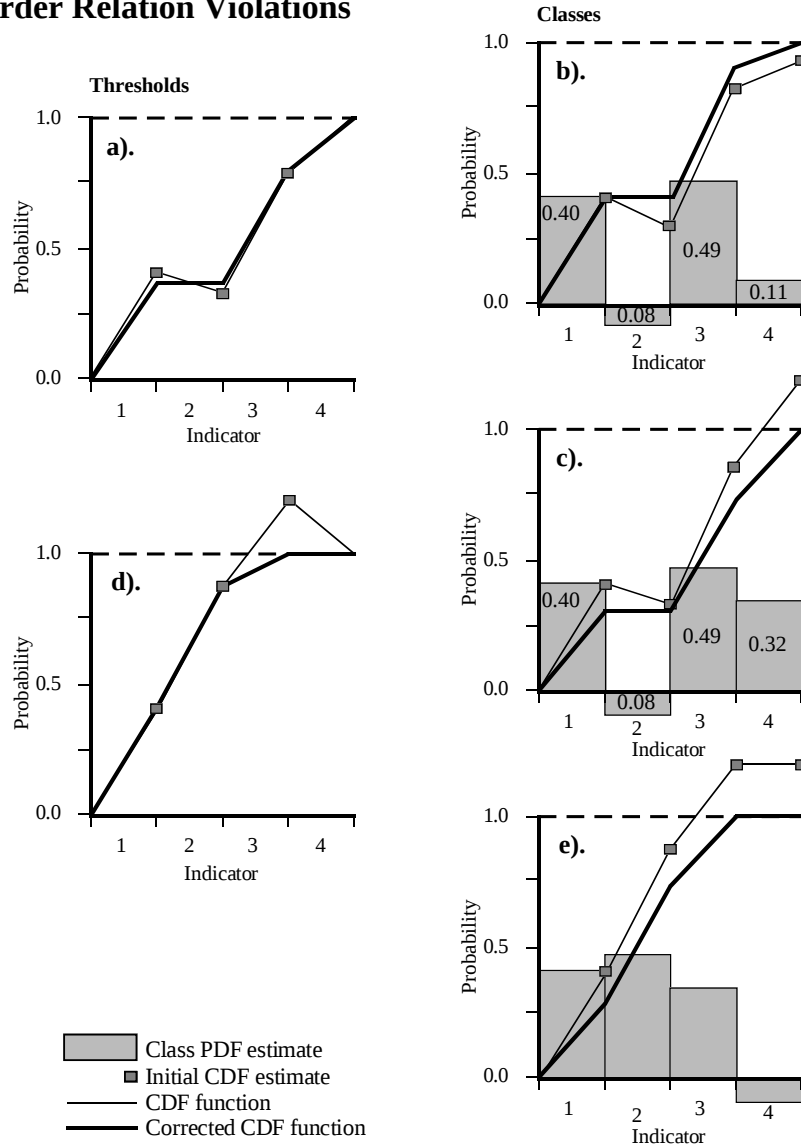


FIGURE 4-7. For both the class and threshold approach, there are two basic types of order relation violations (ORV). a) An individual CDF probability is less than the CDF of a smaller threshold (the CDF is decreasing); this is equivalent to a class having a negative probability of occurrence. This type 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 ($PDF_{new} = PDF_{old} / CDF_{final\ value}$). 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 ($PDF_{new} = PDF_{old} / CDF_{final\ value}$; 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:

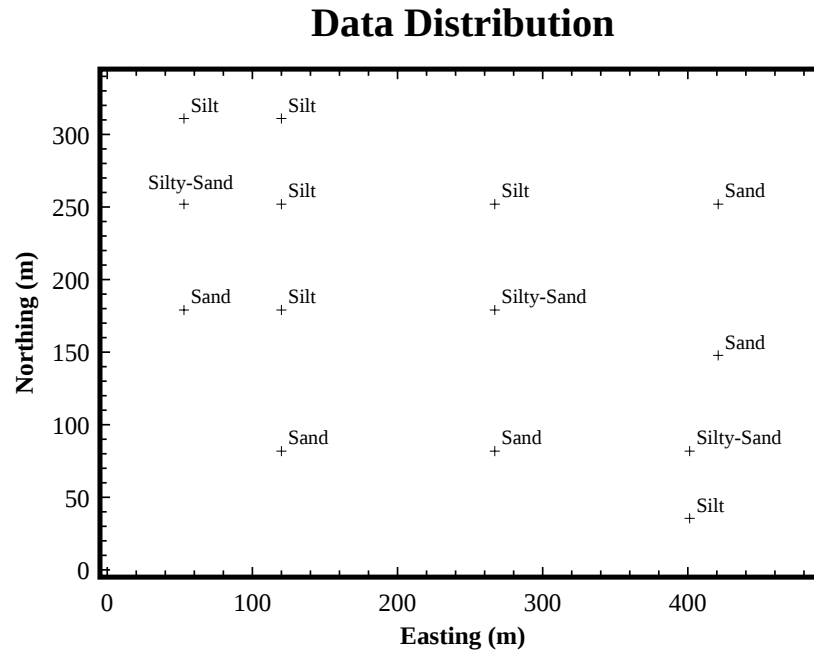


FIGURE 4-8. Synthetic data set distribution.

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).

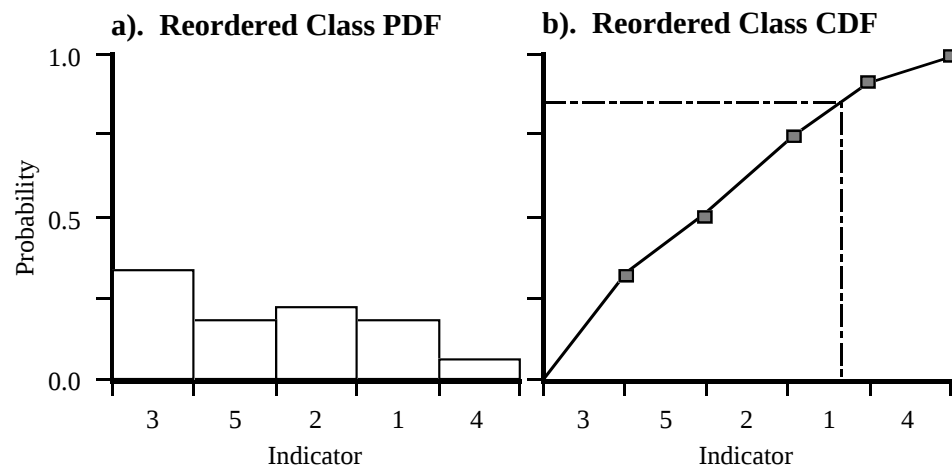


FIGURE 4-9. Reordering indicators in conditional simulation changes the results for an individual simulation grid cell, because the CDF changes along with the indicator ordering, even though the individual components of the PDF do not. Here the indicators from Figure 4.3 have been reordered. The same 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

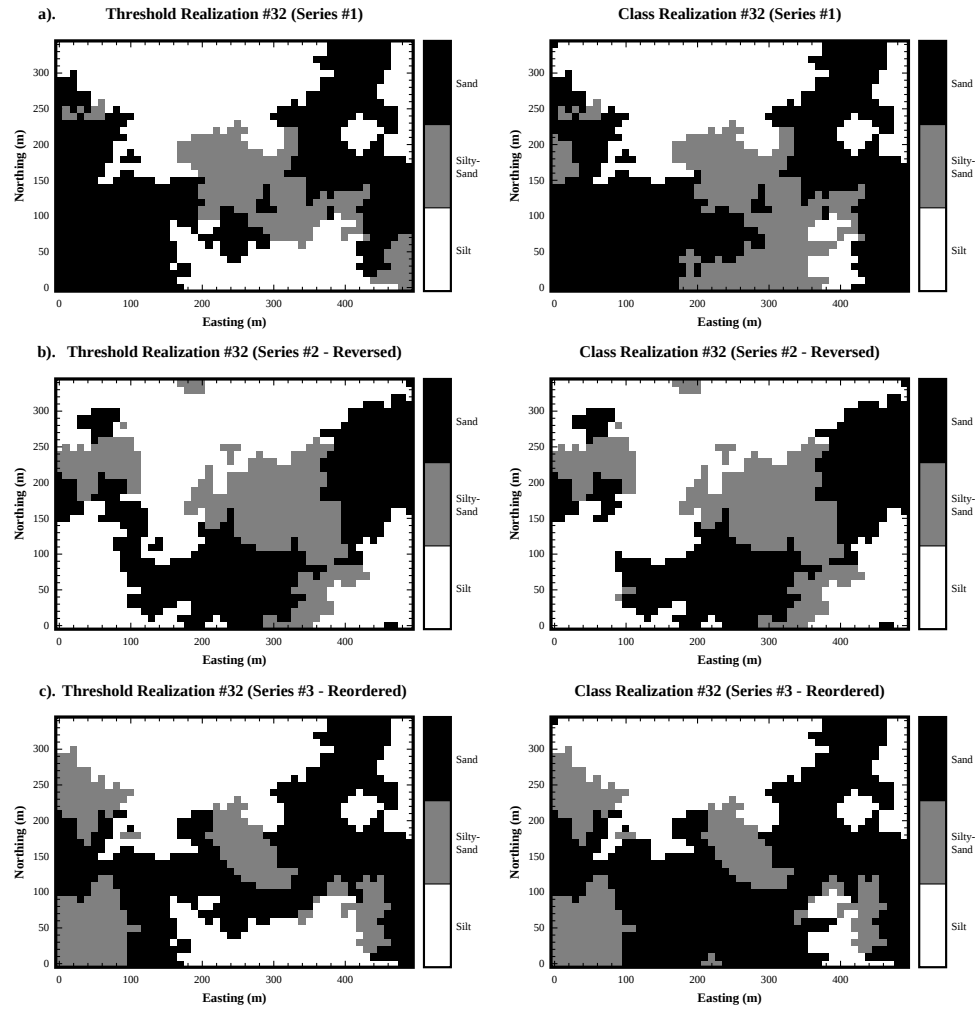


FIGURE 4-10. Realization (#32) pairs for the a) original indicator ordering, b) reversed ordering, and c) arbitrary reordering for thresholds (left) and classes (right). In these realization pairs there are significant differences between the class and threshold results: a) there is more silty-sand in 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,

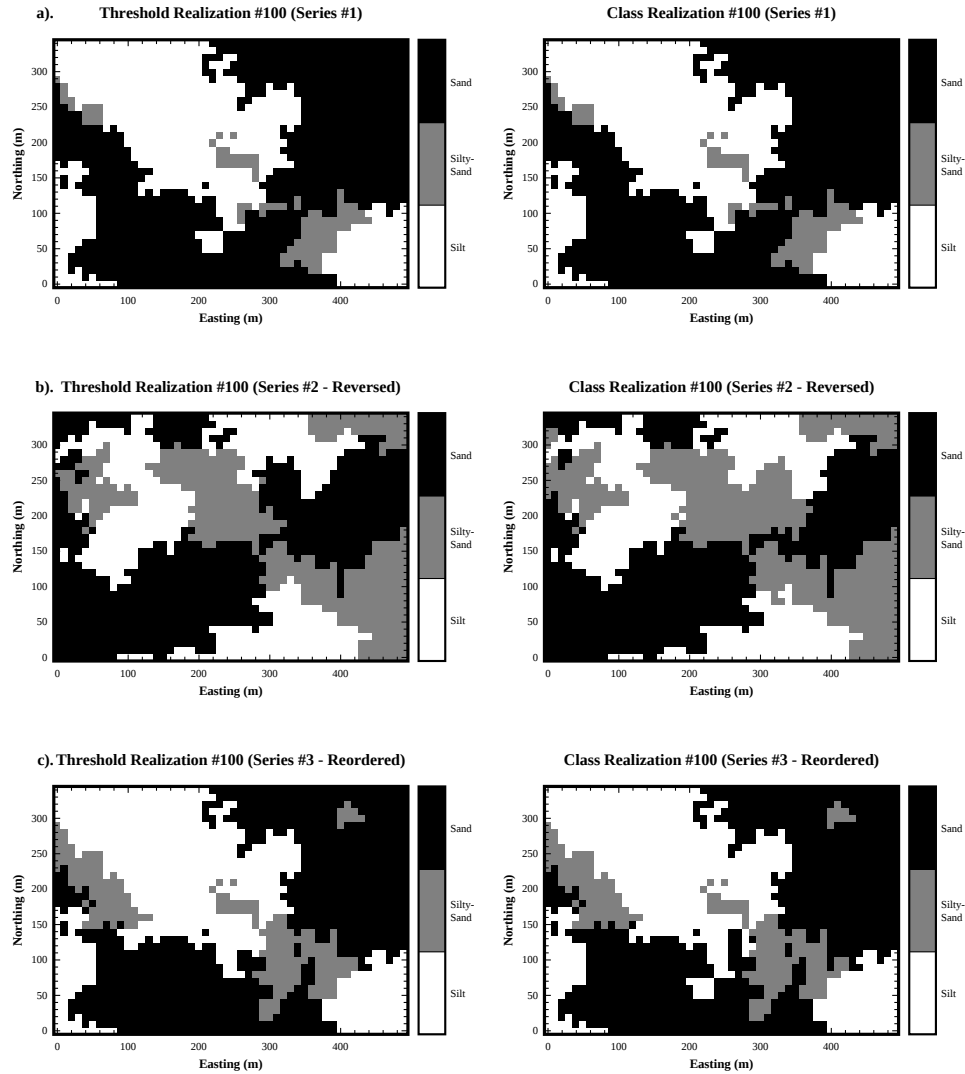


FIGURE 4-11. Realization (#100) pairs for a) original indicator ordering, b) reversed ordering, and c) arbitrary reordering for thresholds (left) and classes (right). These threshold and class realization pairs are similar. The differences between the realization pairs in a, b, and c are expected, because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

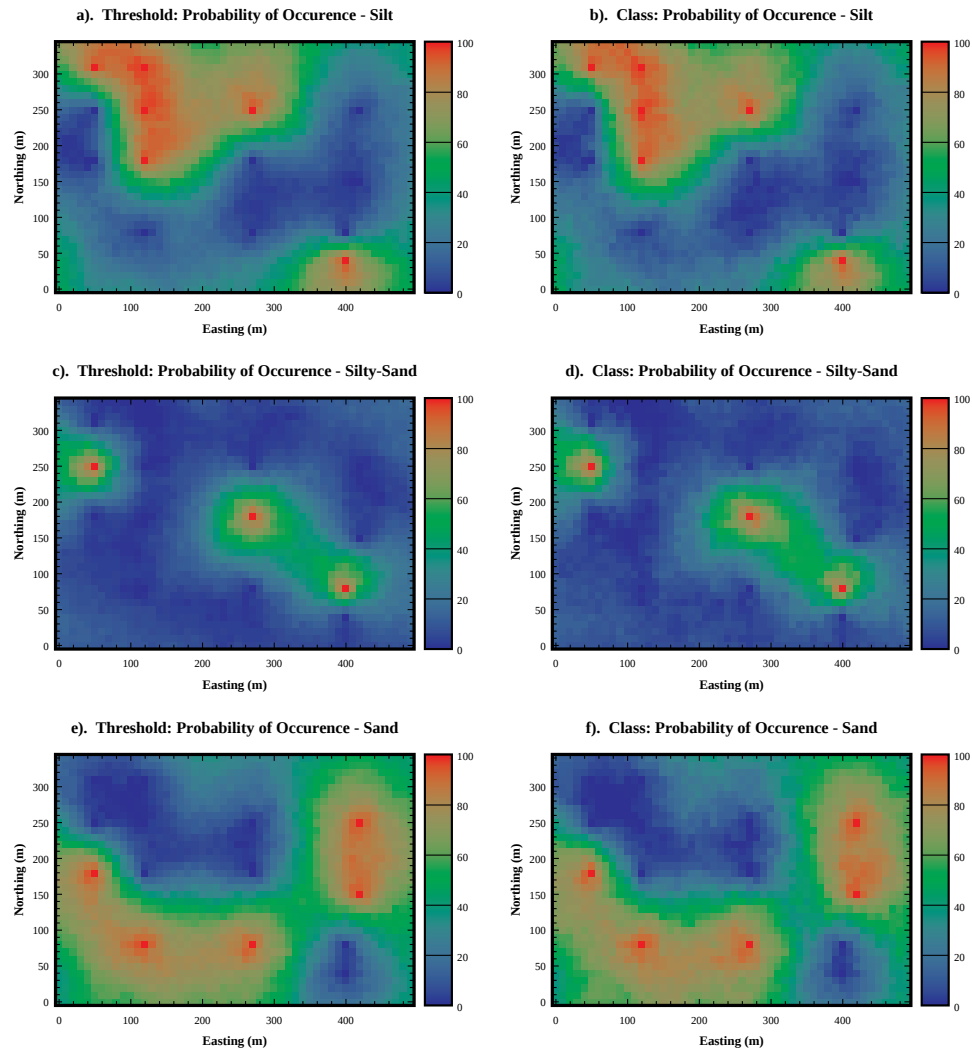


FIGURE 4-12. Maximum probability of occurrence for each indicator. At hard data locations, the indicator type is known; the probability is 0% for any other indicator type, or 100% for the specified indicator type.

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.

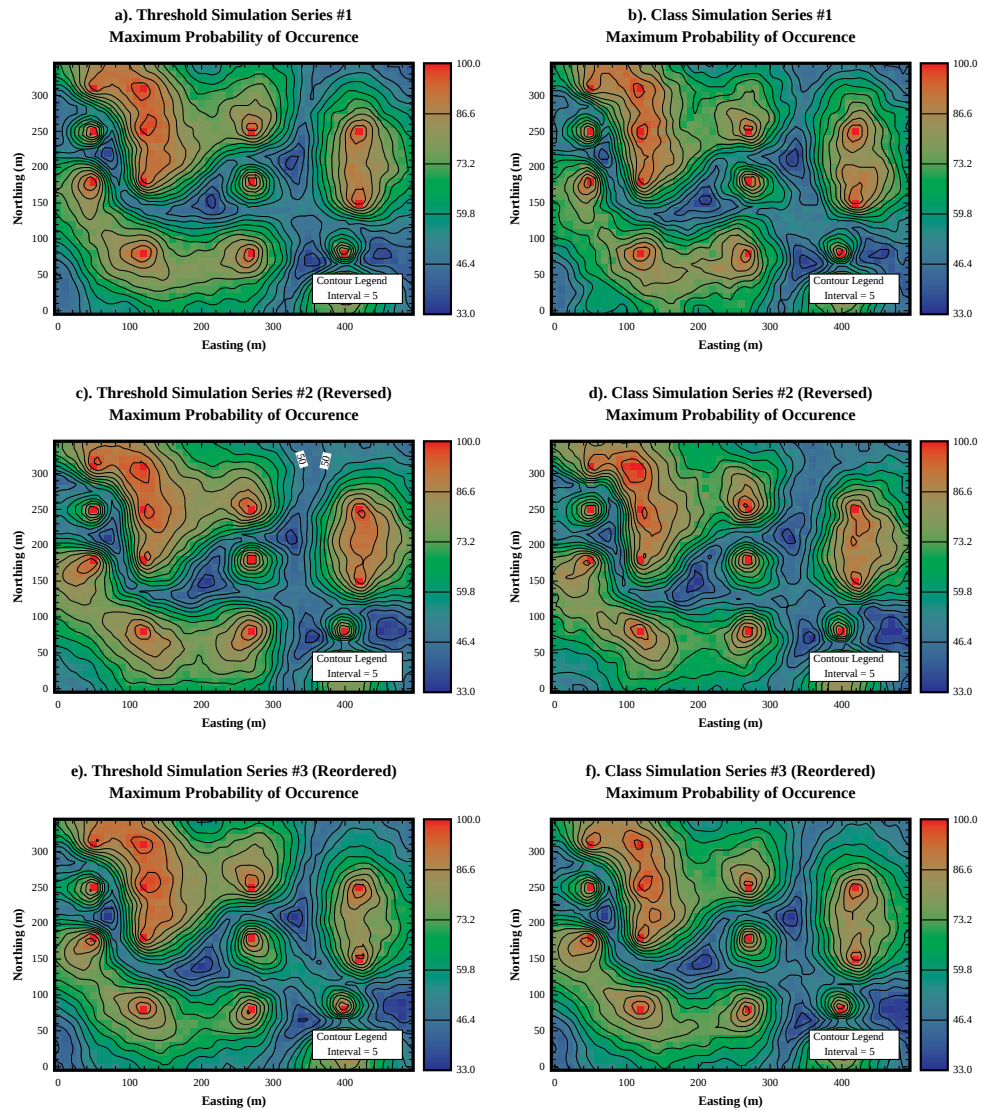


FIGURE 4-13. Maximum probability of occurrence of any indicator. At known data points, the maximum probability of being a particular indicator is 100%. The minimum probability is 33% (100% / # of indicators); at these locations each indicator is equally likely to occur. These maps are useful for evaluating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

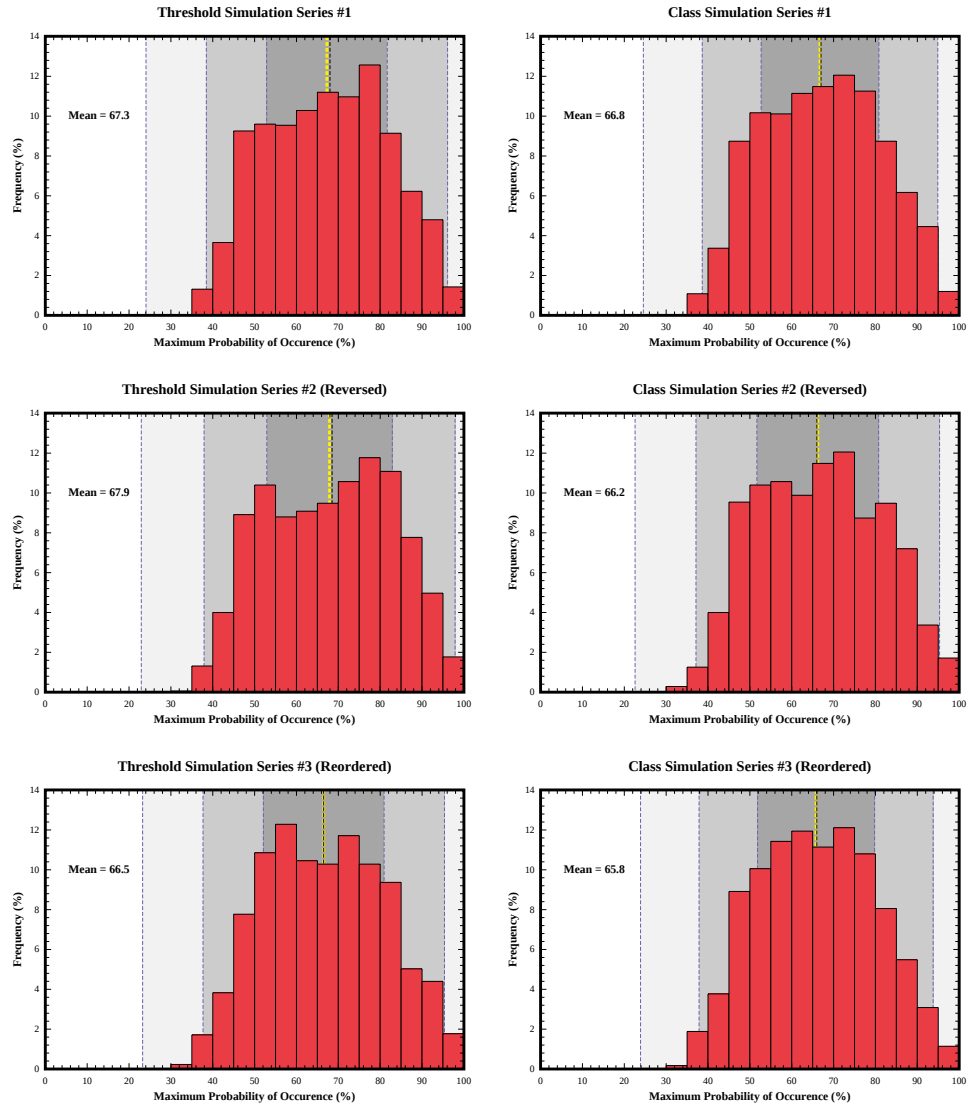


FIGURE 4-14. Frequency of maximum probabilities throughout the grid. At known data points, the maximum probability is 100%. The minimum probability is 33% ($100\% / \#$ of indicators); at locations where each indicator is equally likely to occur.

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

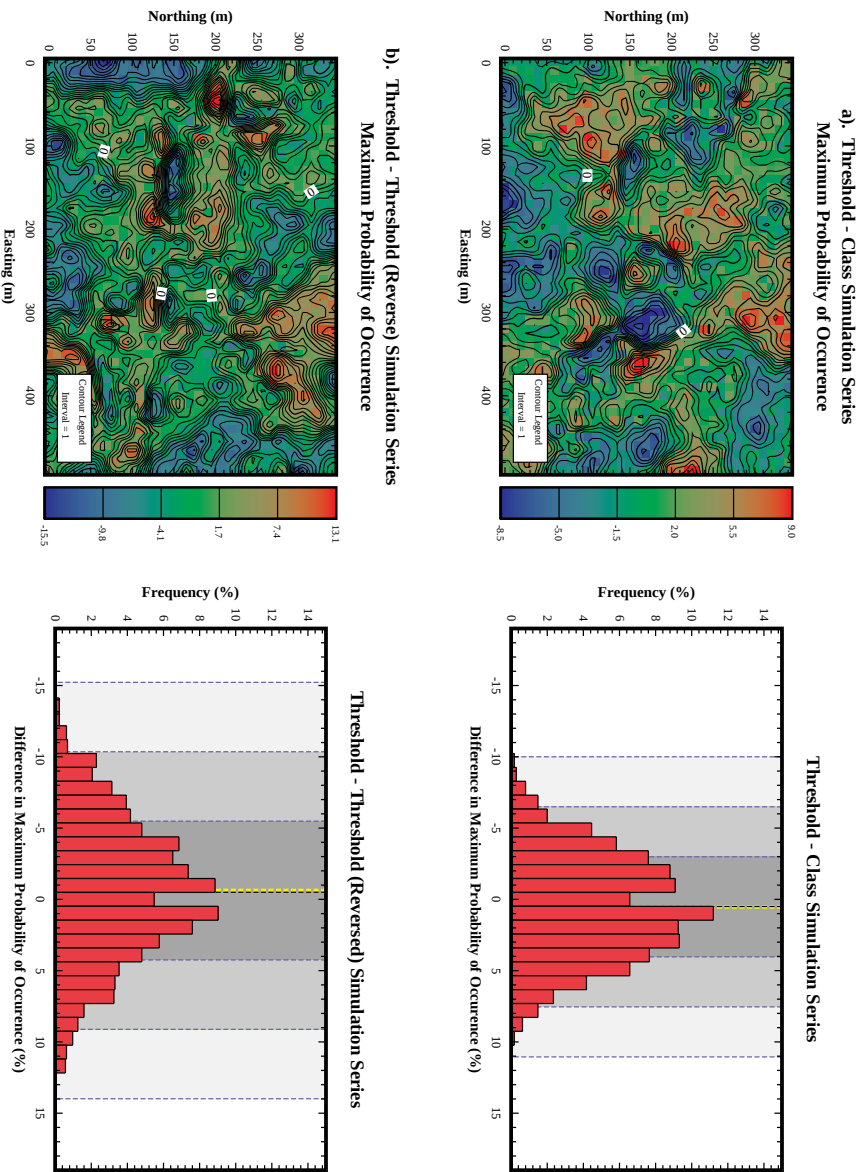


FIGURE 4-15 a,b. Difference between a) threshold and class maximum probability maps, and b) threshold and threshold (reversed) probability maps. These maps show areas, where the different series return different averaged results. Differences are largest, although of opposite sign, in red and blue areas.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

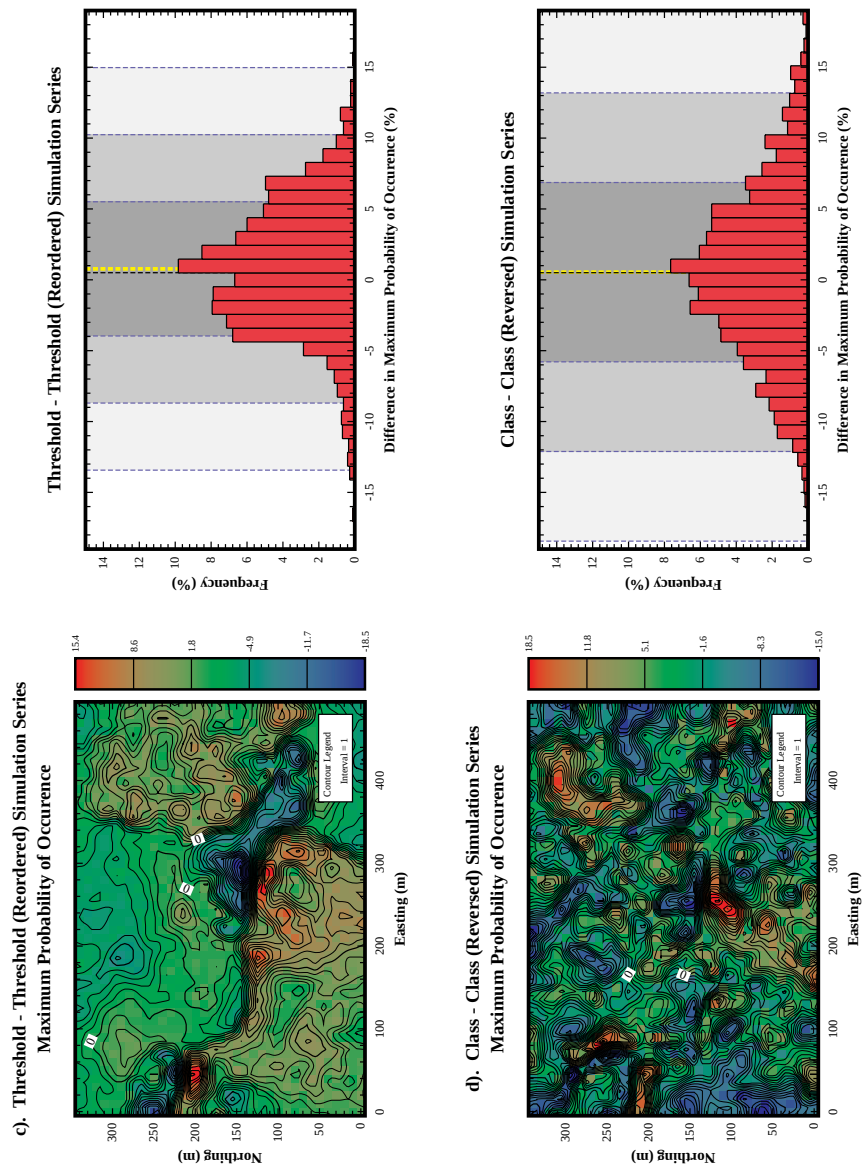


FIGURE 4-15 c,d. Difference between a) threshold and threshold (reordered) maximum probability maps, and b) class and class (reversed) probability maps. These maps show areas, where the different series return different averaged results. The histograms are useful for identifying the magnitude and distribution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

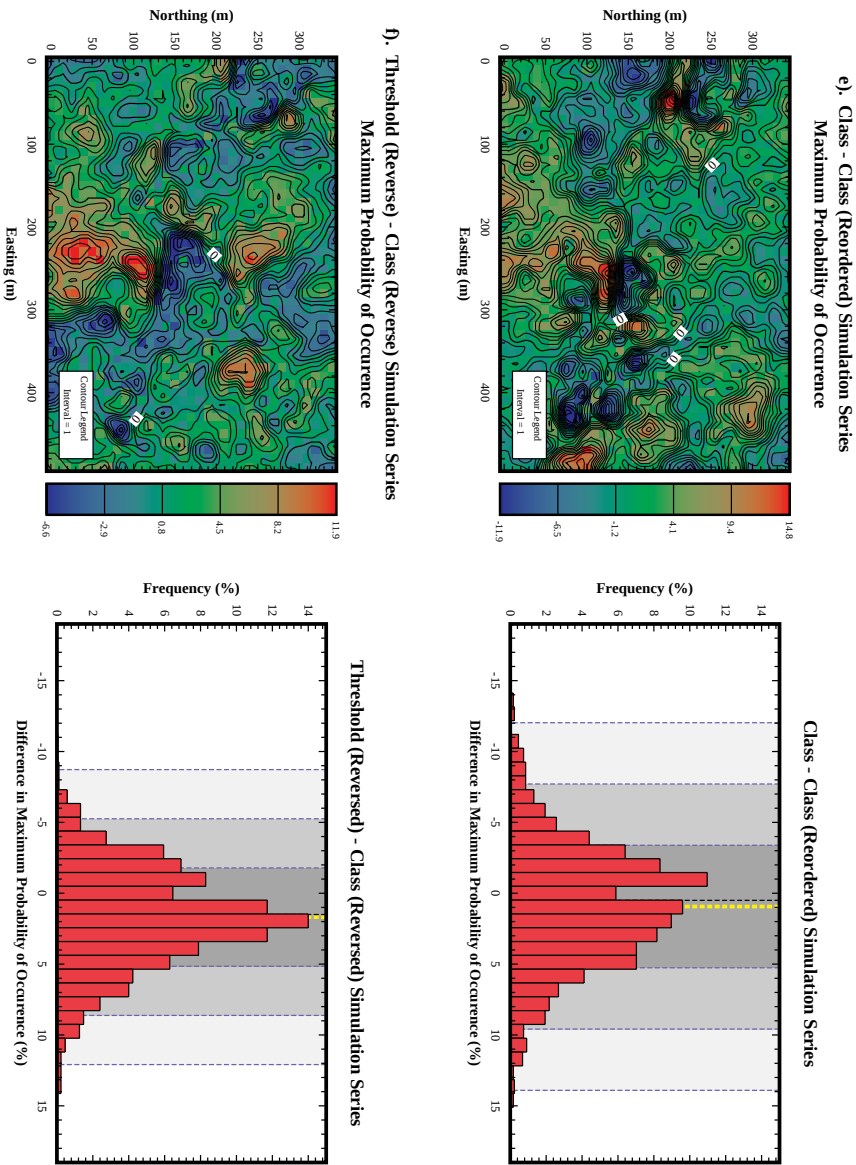


FIGURE 4-15 e, f. Difference between a) class and class (reordered) maximum probability maps, and b) threshold (reversed) and class (reversed) probability maps. These maps show areas, where the different series return different averaged results. The histograms are useful for identifying the magnitude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

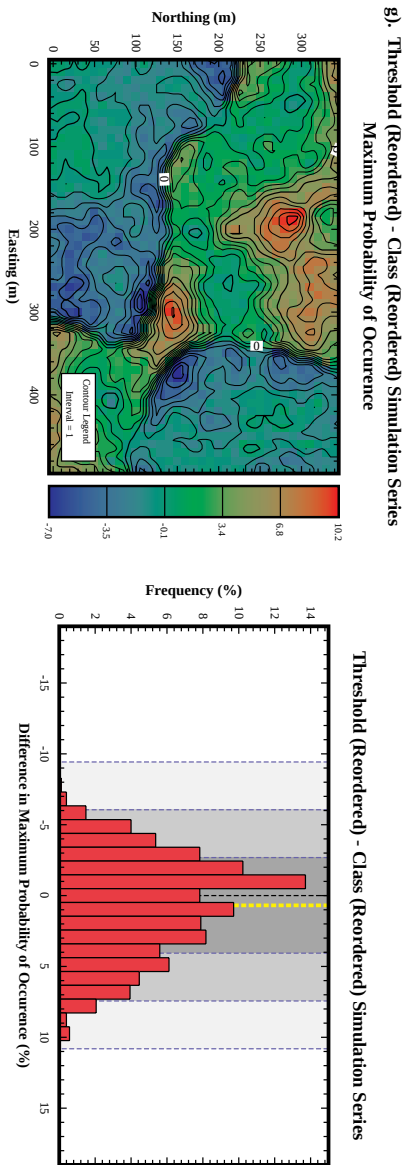


FIGURE 4.15 g. Difference between threshold (reordered) and class (reordered) maximum probability maps. These maps show areas, where the different series return different averaged results. The histograms are useful for identifying the magnitude and distribution of the differences.

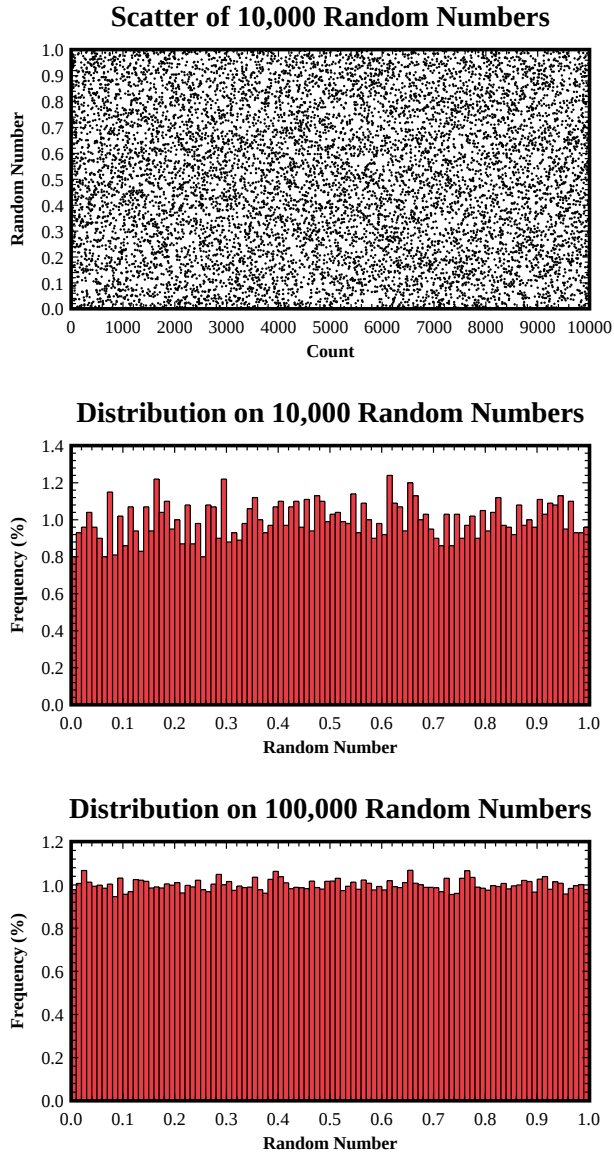


FIGURE 4-16. Test of random number generator: a) scatter of 10,000 sequential random numbers, b) frequency distribution of 10,000 random numbers (equal distribution would put 1% in each class), and c) frequency distribution of 100,000 random numbers (equal distribution would put 1% in each class).

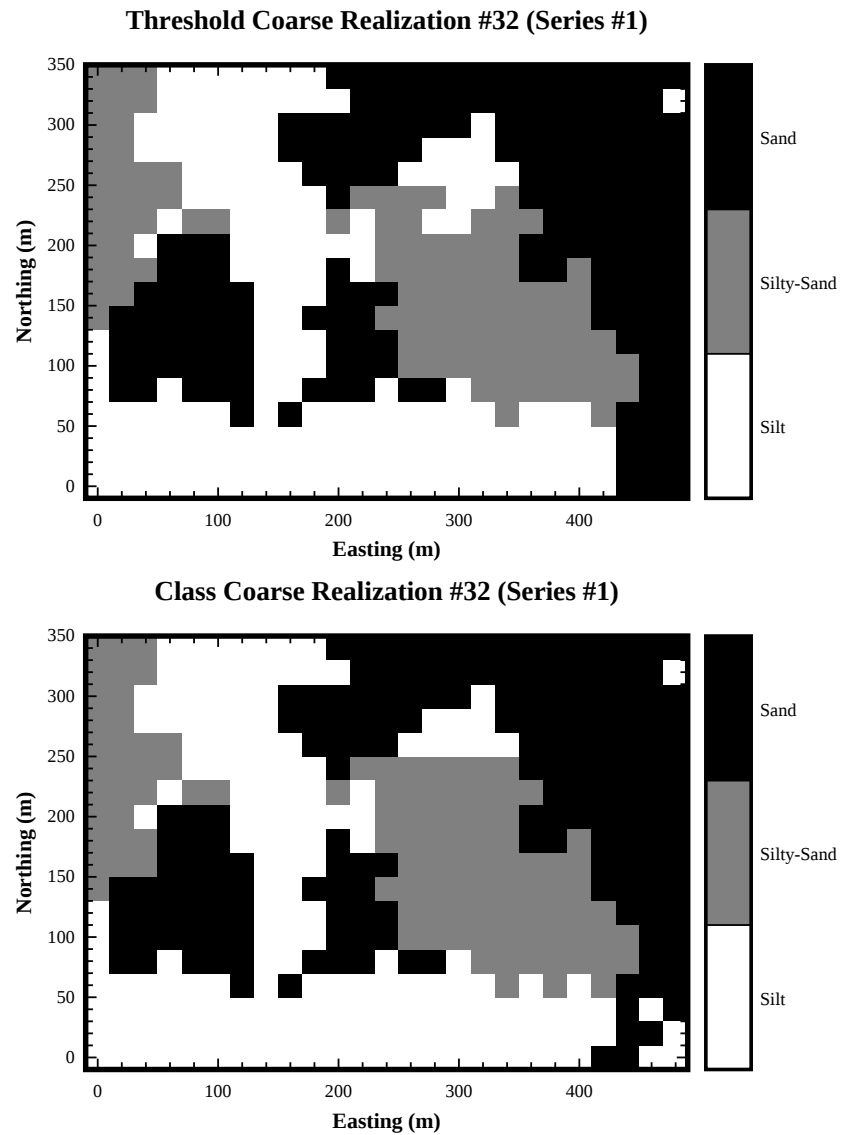


FIGURE 4-17. Coarse grid realization (#32) pair for the original indicator ordering. These grids, are slightly different, due to an ORV occurring at (490m, 30m). Because of the ORV, the CDF's varied between the two methods at this cell, and a random number in each realization selected different material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

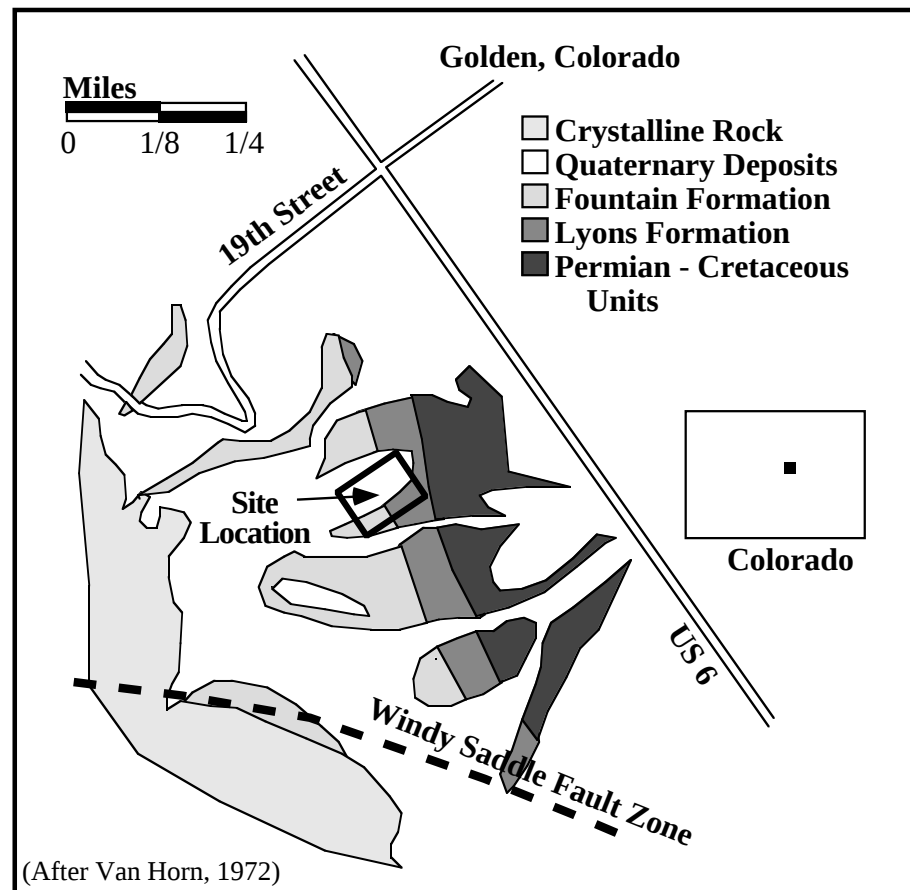


FIGURE 4-18. CSM Survey Field location map.

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

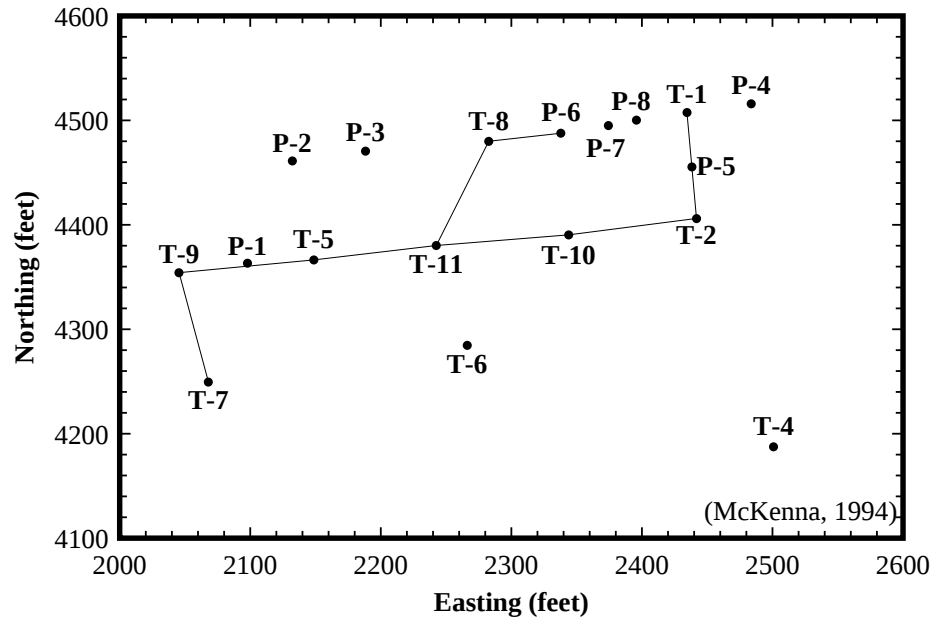


FIGURE 4-19. CSM Survey Field site map. Dots represent borehole locations. Solid lines identify location of tomography surveys.

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

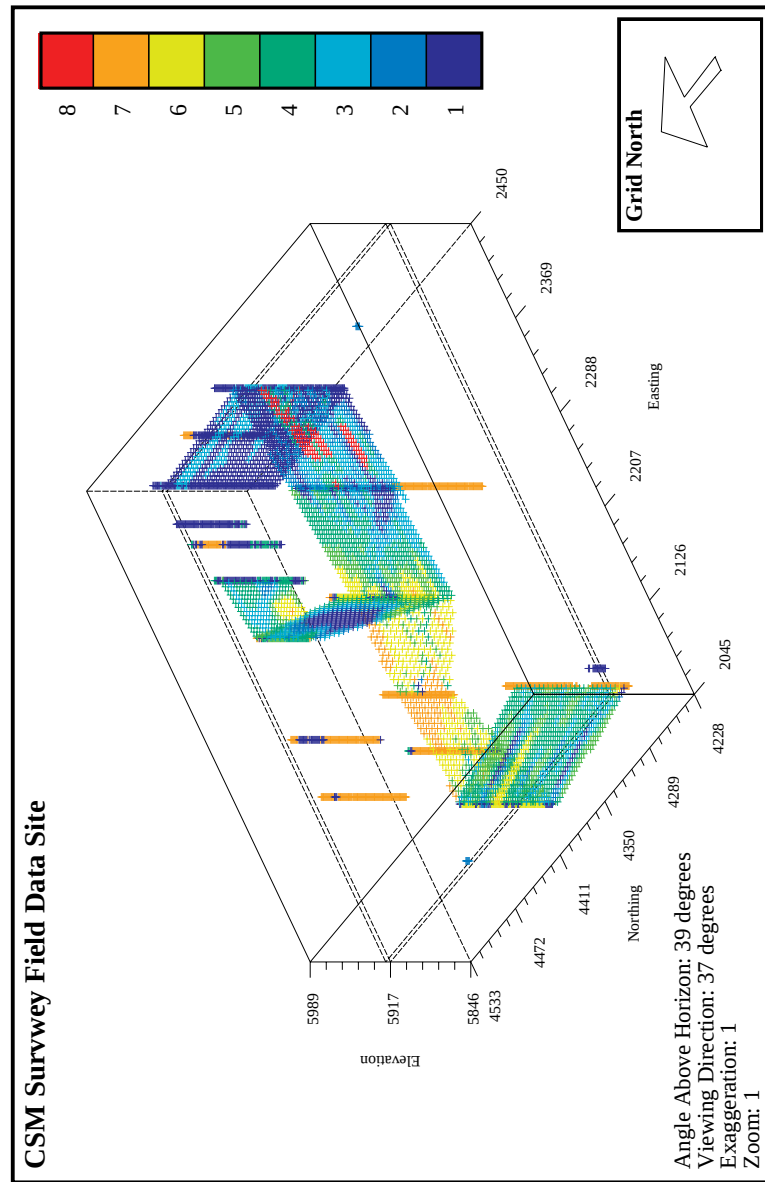


FIGURE 4-20. CSM Surveyway Field borehole and tomography data converted to indicators. The dotted layer delineates the extents of the model grid.

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.

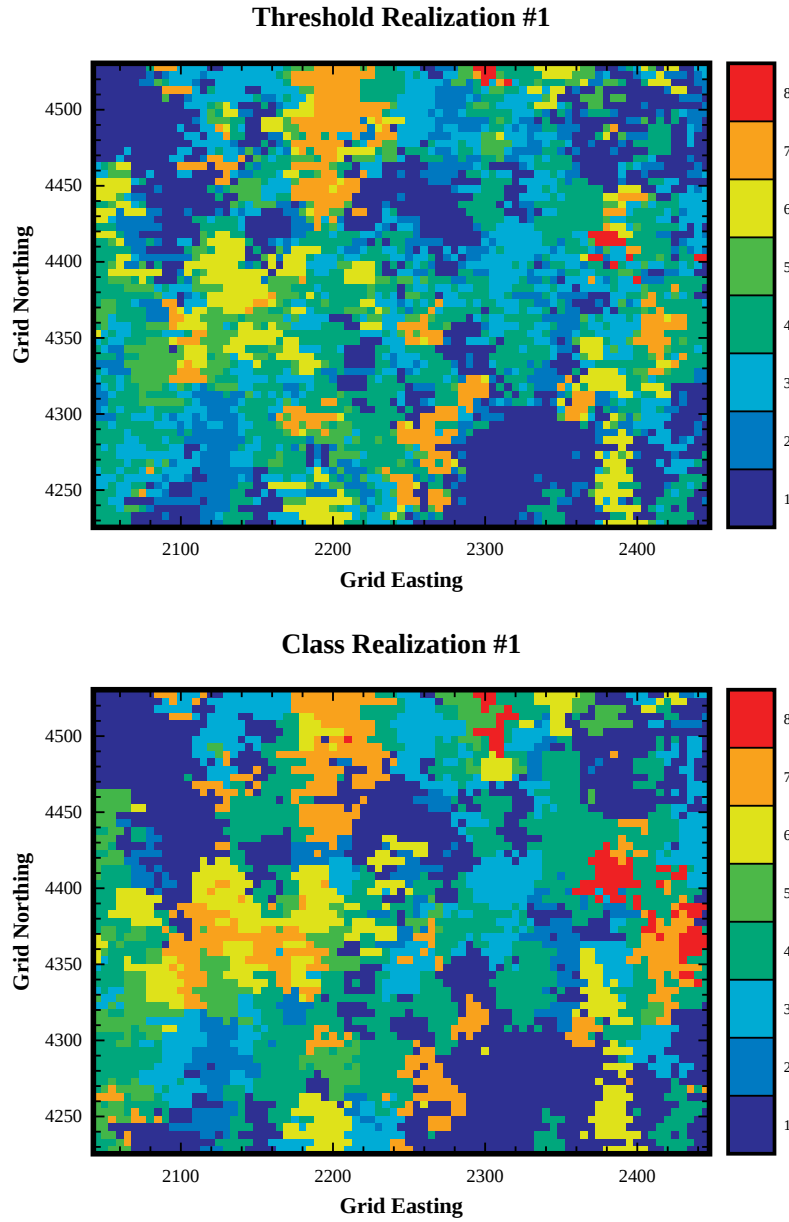


FIGURE 4-21. Individual threshold and class realization #1.

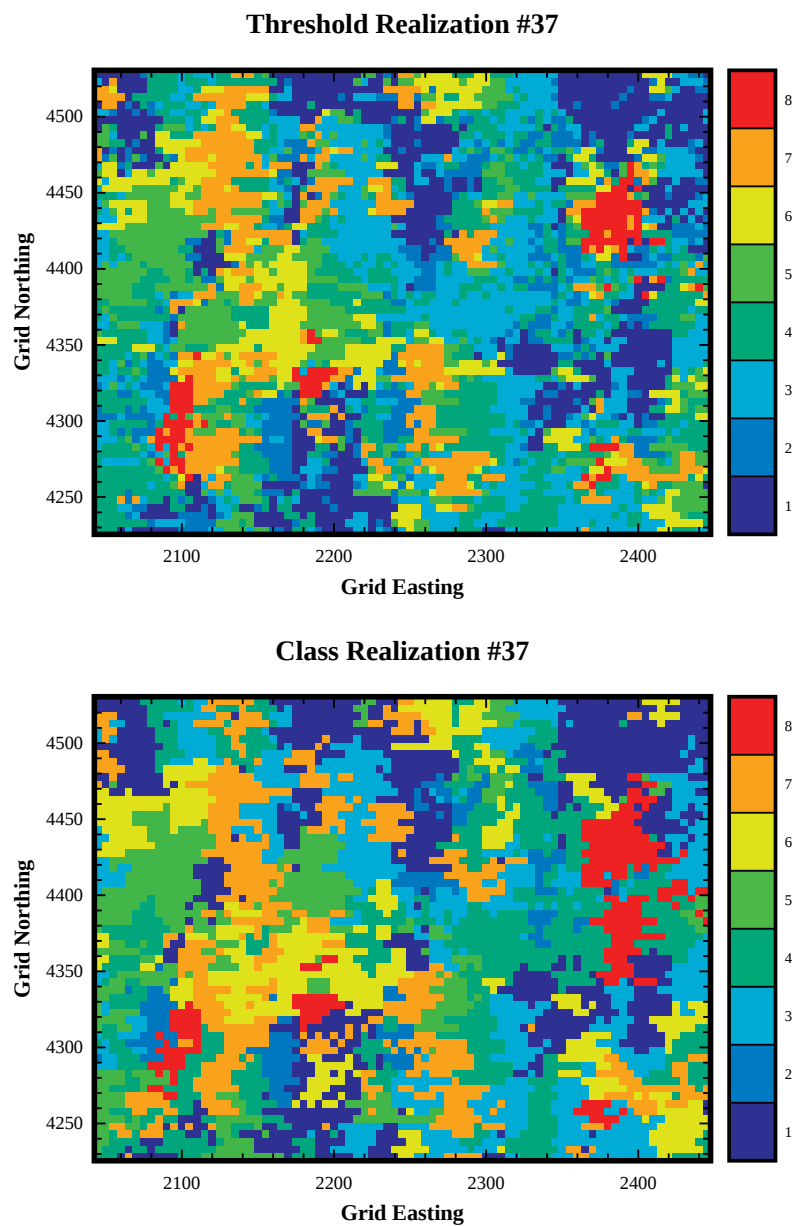


FIGURE 4-22. Individual threshold and class realization #37.

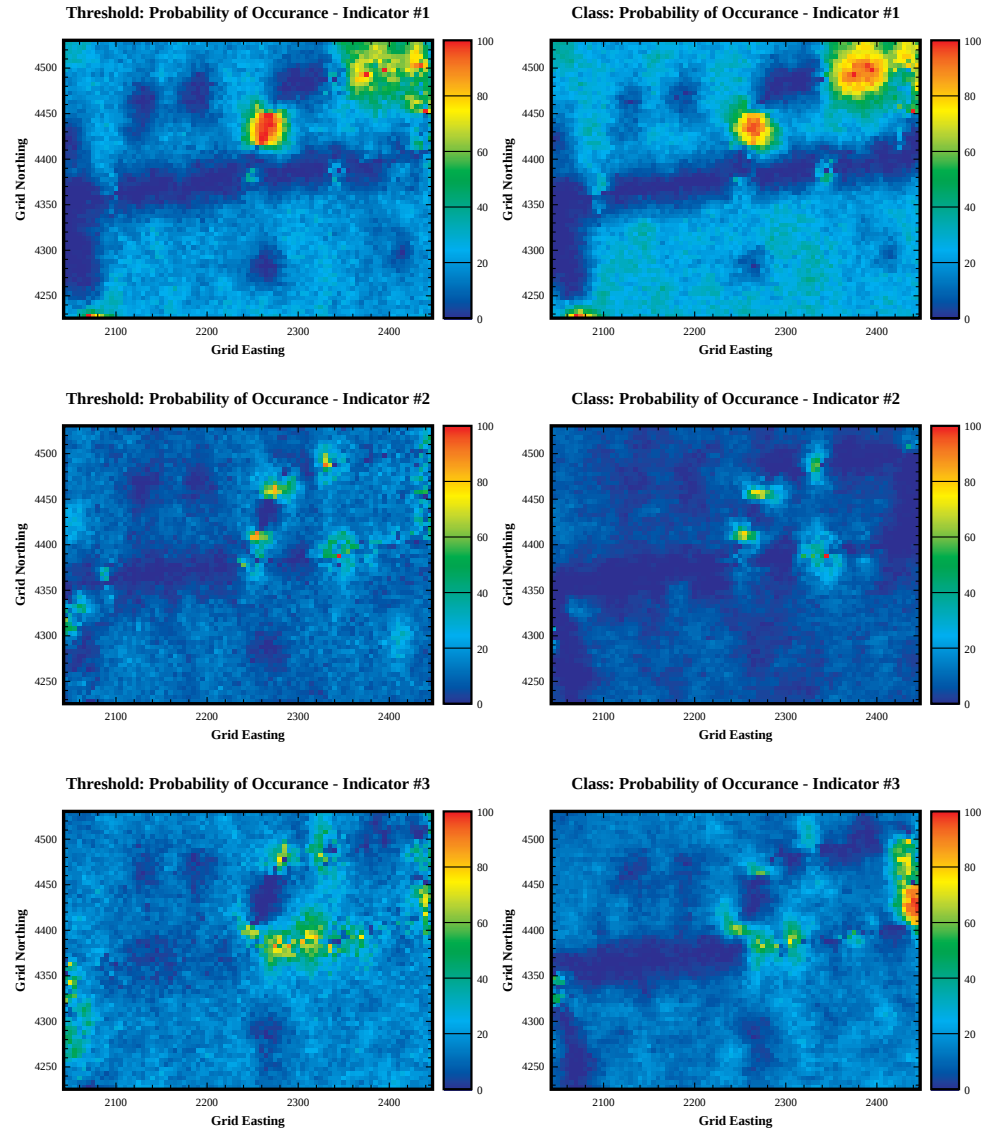


FIGURE 4-23 a,b,c. Maximum probability of occurrence for threshold and class indicators #1, #2, and #3. At known hard data points, probability is 0% for any other indicator type, or 100% for the specified indicator type.

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

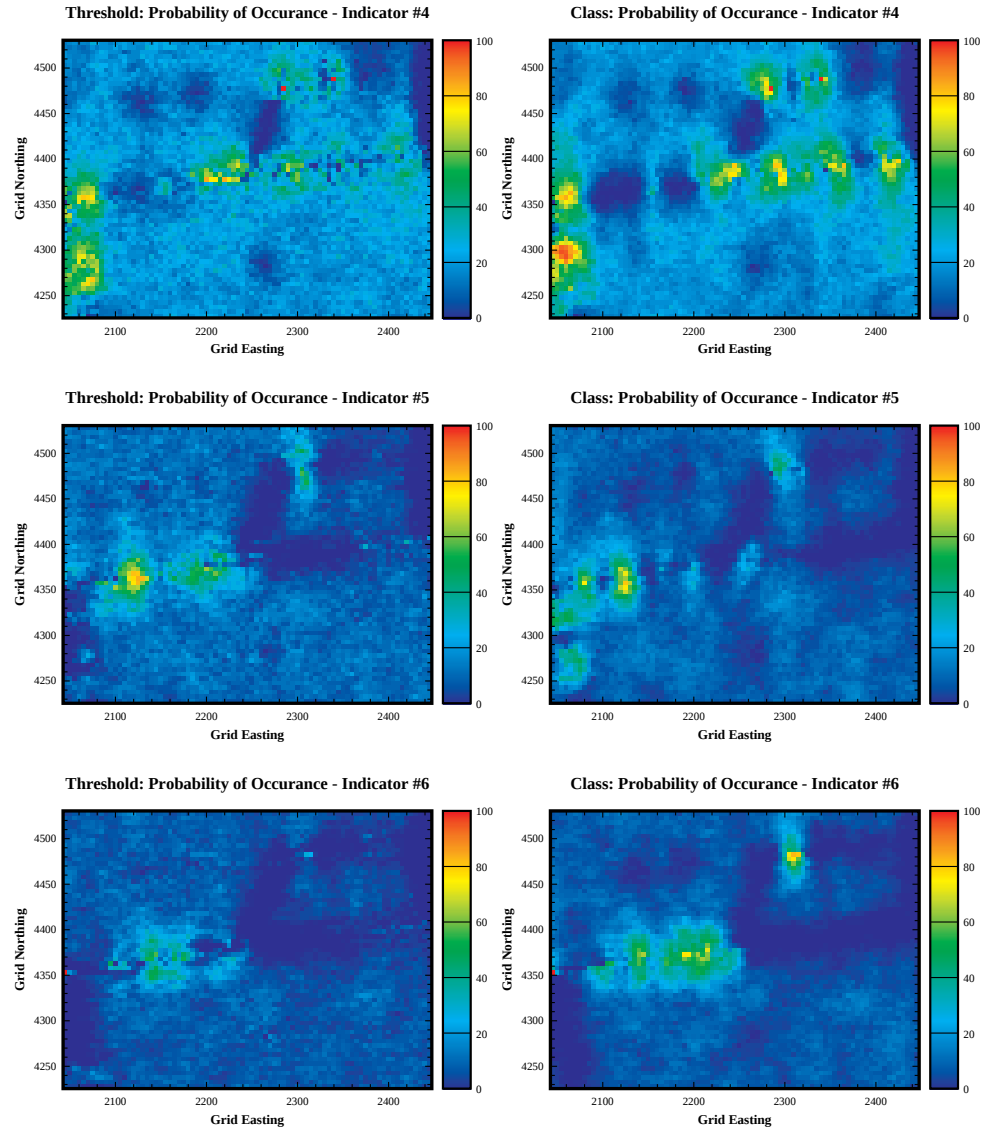


FIGURE 4-23 d,e,f. Maximum probability of occurrence for threshold and class indicators #4, #5, and #6. At known hard data points, probability is 0% for any other indicator type, or 100% for the specified indicator type.

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

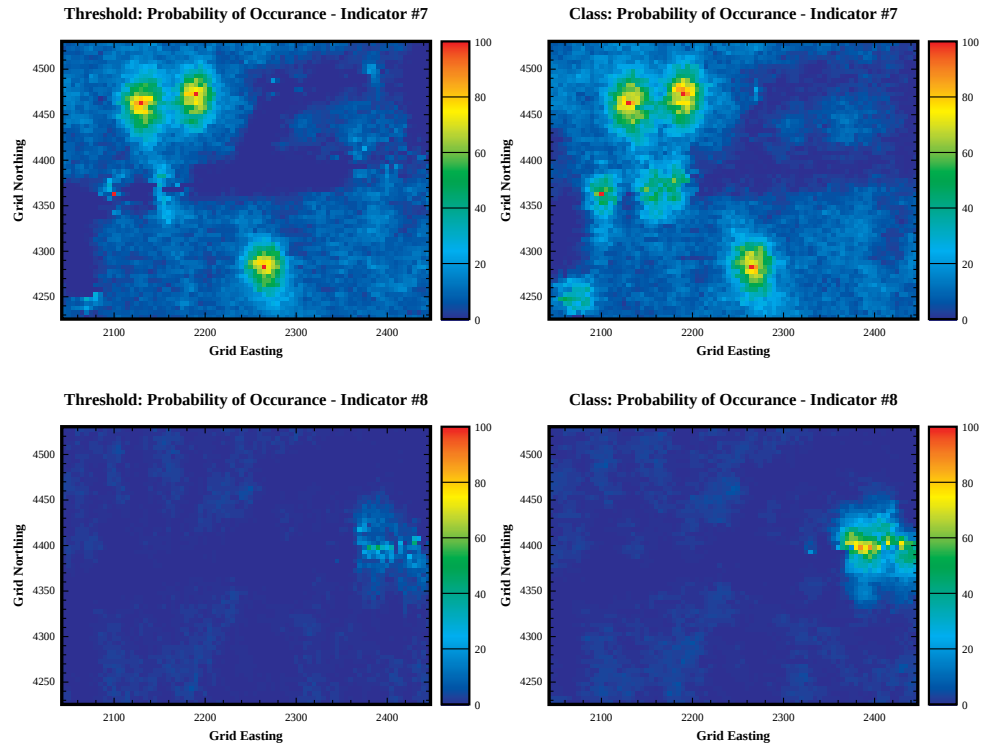


FIGURE 4-23 g,h. Maximum probability of occurrence for threshold and class indicators #7 and #8. At known hard data points, probability is 0% for any other indicator type, or 100% for the specified indicator type.

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.

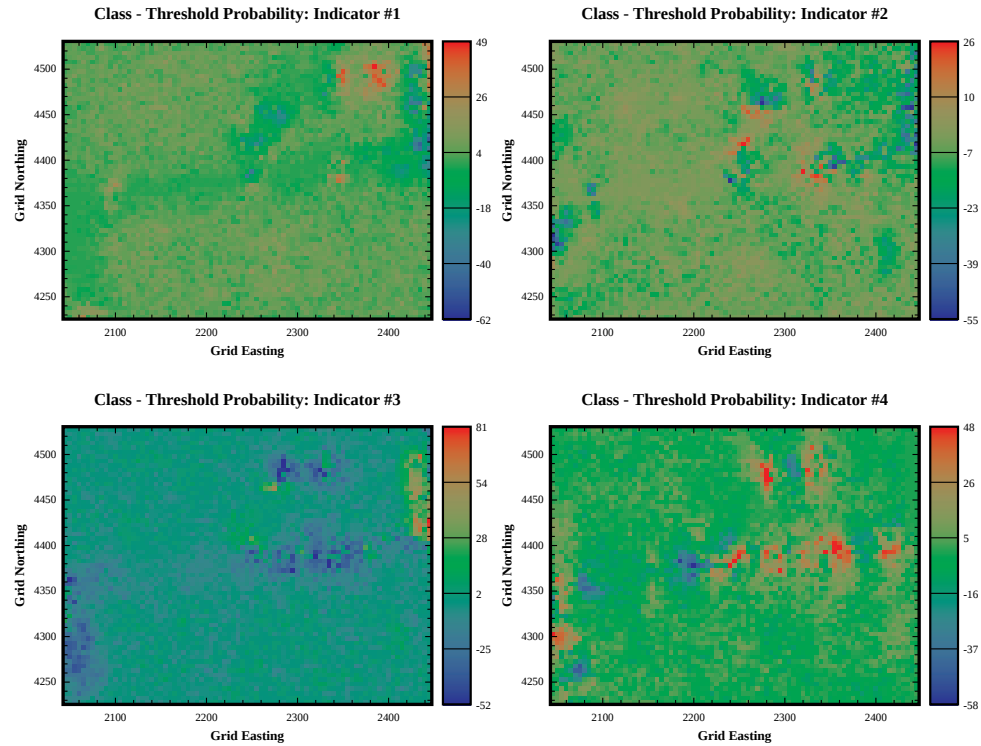


FIGURE 4-24 a,b,c,d. Difference between the class and threshold, individual indicators (#1-4) maximum probability of occurrence maps.

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.

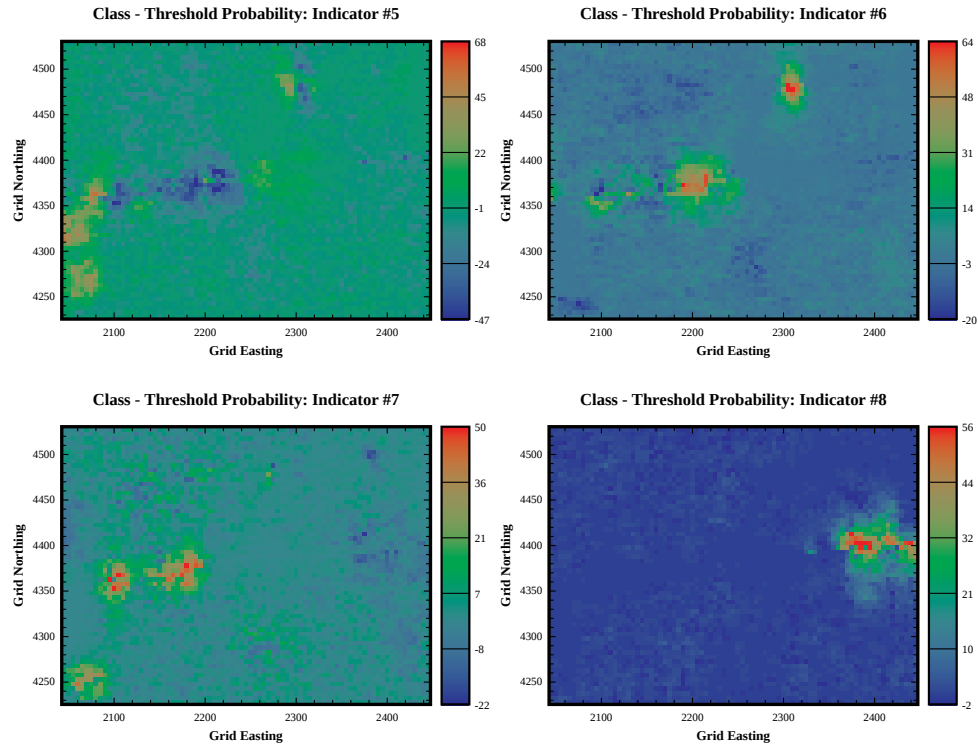


FIGURE 4-24 e,f,g,h. Difference between the class and threshold, individual indicators (#5-8) maximum probability of occurrence maps.

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:

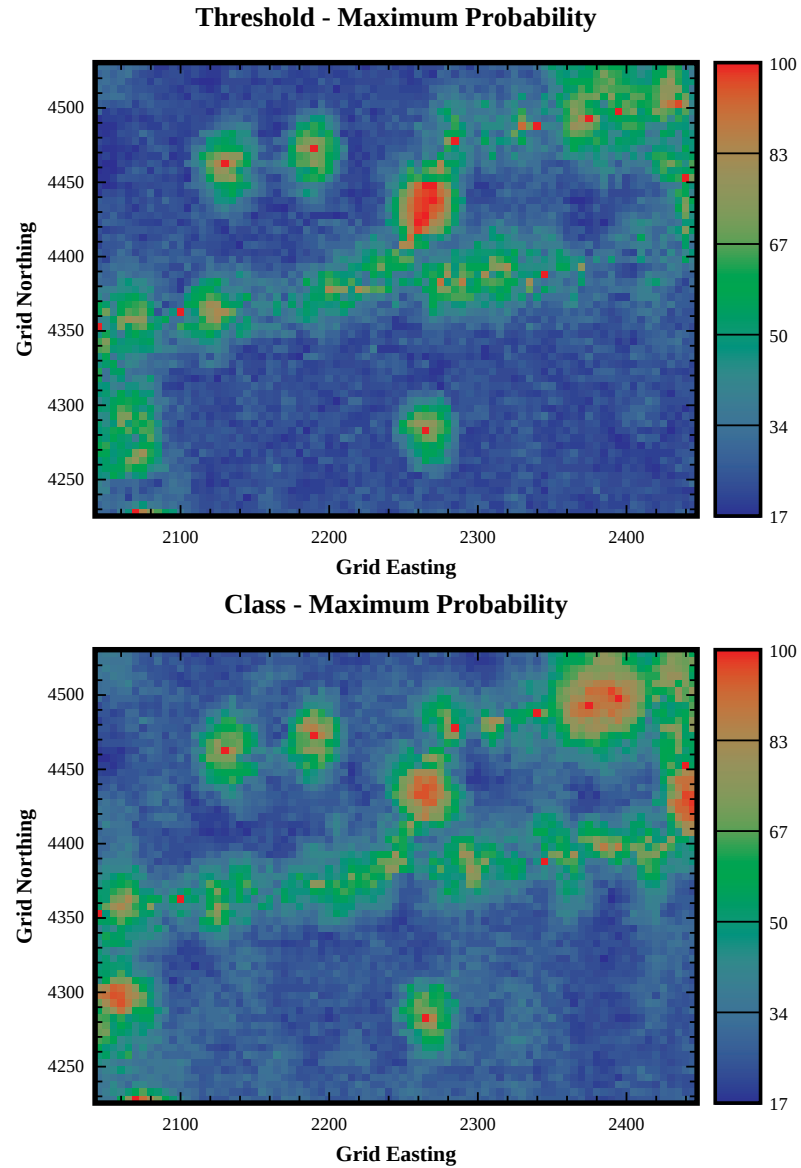


FIGURE 4-25 a,b. Maximum probability of occurrence of any indicator. At known data points, probability is 100%. The minimum probability is 12.5% ($100\% / \#$ of indicators); at these locations each indicator is equally likely to occur (no cells had this minimum probability). These maps are useful for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

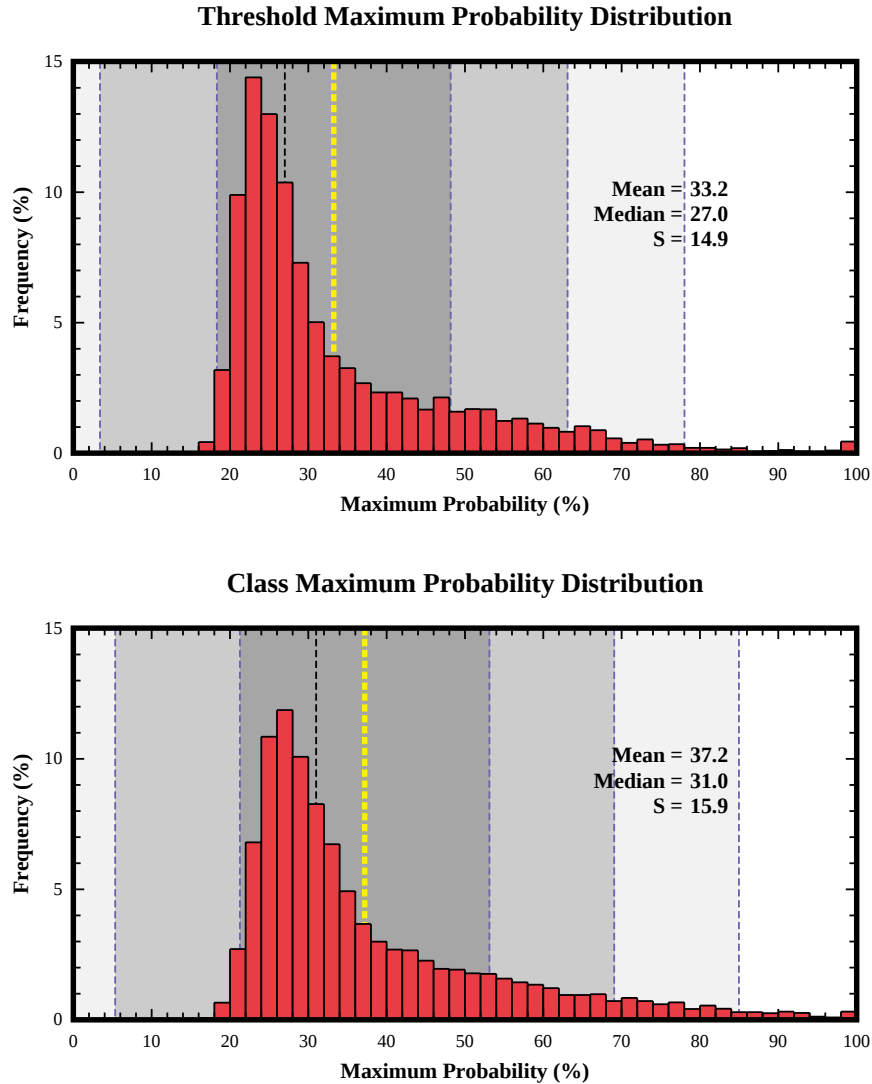
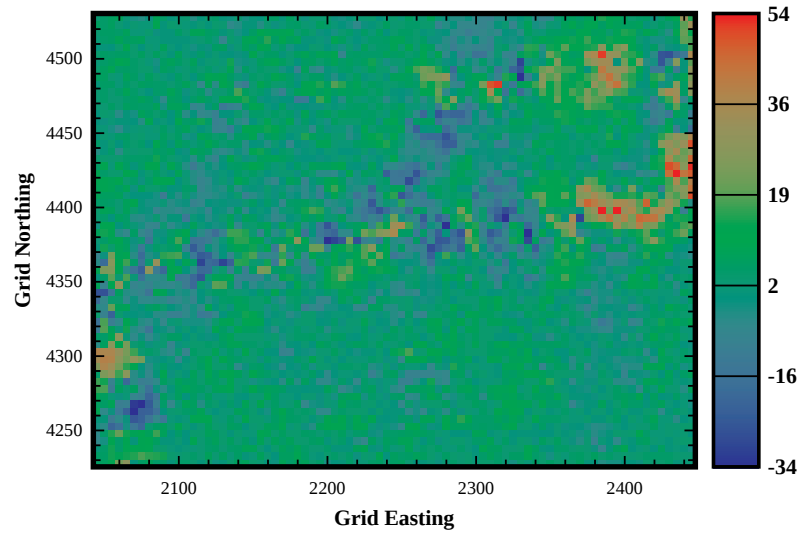


FIGURE 4-26 a,b. Histograms indicate the maximum probability of occurrence of any indicator. At known data points, probability is 100%. The minimum probability is 12.5% (100% / # of indicators); at these locations each indicator is equally likely to occur (no cells had this minimum probability). Class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

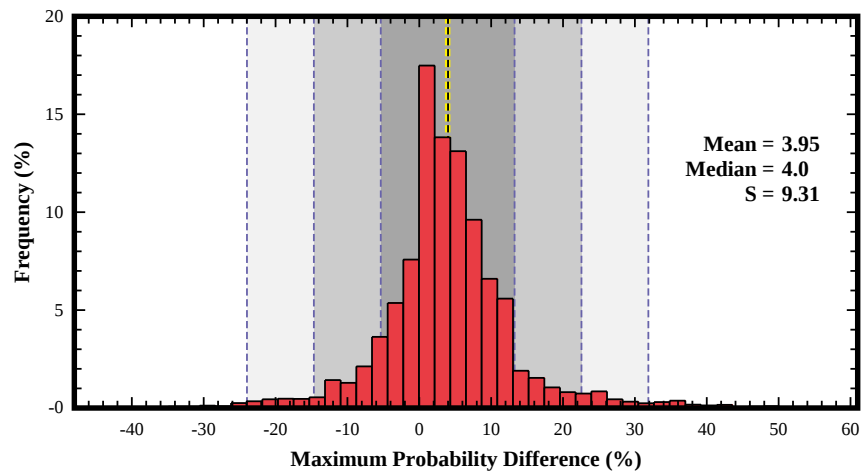


FIGURE 4-27. a) Spatial distribution of the difference between the class and threshold maximum probability maps; b) histogram of the same information. The positive mean difference indicates the class realizations have a lower level of uncertainty.

- 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.