

Examining the Performance of Kriging in the Estimation of Property Value - A Case Study

Md. Nasir Daud, PhD
National Institute Of Valuation
Valuation And Property Services Department
Malaysia

Received : October 2000

Revised : November 2000

Abstract

Location is paramount to property value. Unfortunately, the conventional methods of valuation are often implicit in their treatment of location as a value determinant. This paper investigates the practicality of a method that provides for the explicit consideration of location in valuation by using a spatial interpolation technique known as the ordinary kriging method. To evaluate the performance of this single-variable location-explicit method, a comparison was made against the results generated with the multi-variable less location explicit Multiple Regression Analysis (MRA).

Keywords: multiple regression analysis, kriging, PRESS sums, trend surface analysis (TSA), variogram

Introduction

Existing methods of valuation range from the conventional techniques such as the comparison method to statistical techniques such as regression analysis. These methods are not always based on the explicit consideration of location in their determination of housing property value; this is despite the all-time recognition that location is a critical factor to property value.

In this paper, we investigate a method that explicitly provides for the consideration of location in valuation. More specifically, the spatial interpolation technique of interest is the ordinary kriging method. This method proceeds by analysing the geographic arrangements of data and using the analysed information to interpolate value on other points of interest. For an evaluation of the method's predictive performance, the

value estimates derived by this method are compared against the actual property values. The estimation performance of this method is then compared against that of the multiple regression analysis (MRA) using the *PRESS* sums criterion. This comparison enables conclusions regarding the extent to which a single-variable location-explicit model performs against a multi-variable but location-implicit model of value prediction.

As a tool for spatial analysis, kriging is not exactly new. Its property for handling the prediction of values where points are spatially dependent between one and another is known. In the realm of property valuation, however, the application of this method has remained practically unheard of. What makes this technique interesting to investigate is the fact that it explores the patterns of spatial dependence within local areas and uses the information to interpolate value at points of interest. This

makes it a tool of different and greater precision to other surface interpolation techniques attempted before for property valuation.

To implement this case study, a simulated dataset was in use. We start with a discussion of the motivations for considering surface interpolation technique.

Spatial Interpolation for Housing Property: A Rationalisation

The rationale behind spatial interpolation is the observation that, on average, values at points closer together in space are more likely to be similar than points farther apart (Tobler, 1979). This notion of spatial 'association' is extendable to housing since house units close together tend to have similar values, concordant with the valuer's 'tone of the list' assumption applied in rating valuation (Wyatt, 1994). Further, it is common practice in dealing with the traditional comparison method that, *ceteris paribus*, nearer houses take precedence over houses farther away in the choice as comparables.

Valuers have for quite some time recognised the potential in exploiting the relationship between spatial autocorrelation and patterns in house values. Indeed, if it can be shown that meaningful associations exist between the positions in space of houses and the way the value of these houses relates to one another, then it is plausible that the former can be used as the basis for the prediction of the latter. The work of Byrne *et al* (1973) represents probably the earliest known attempt at exploring this possibility. In that study conducted in 1973, the investigators had used the trend surface analysis (TSA), a variant of the surface interpolation techniques, to determine housing prices in St. Albans on the basis of their locations.

It is well acknowledged here that interpolated surfaces are continuous while property values are discrete. However, if we can exploit a spatial interpolation method to create continuous surfaces from a sample of discrete points, we can derive the value estimate for each point from the interpolated surface. Indeed, if this interpolation method performs satisfactorily, we have indeed found another approach to mass valuation, which will be potentially useful for rating applications.

Nonetheless, underlying the above approach is the pretense that we can predict a house value solely on the basis of its locational information. It is often the case that property values vary drastically even between neighbouring housing units, such as when a vacant detached residential plot is situated next to a fully built plot. In such a situation, spatial interpolation techniques would be of limited value. Still, means can be found to mitigate such an effect, both through the choice of appropriate interpolation methods and the careful stratification of the data to achieve improved homogeneity.

The Kriging Method of Spatial Interpolation

Kriging has developed from the practice in earth science. For this, it has also been referred to as the geo-statistical method of interpolation. A number of methods are available within the kriging family, such as block kriging, co-kriging and probabilistic kriging, but in this study we utilise the ordinary kriging method.

The method of kriging was developed in the late 1960's by G. Matheron who was inspired by the contributions of D.G. Krige (Christensen, 1991). It was originally developed for use in the mining industry but has become increasingly popular in many fields of science and industry where there is

a need for evaluating data using the consideration of spatial or temporal correlations (Wackernagel, 1995).

The methods developed by Matheron produce optimal results in the sense that the interpolation weights are chosen to optimise the interpolation functions in order to provide a 'best linear unbiased estimate' (BLUE) of the value of a variable at a given point (Burrough *et al*, 1998). It is linear because it bases estimates on linear combinations of available data, unbiased because it aims to make the mean residual error equal to zero, and best because it aims to minimise the variance of the errors (Isaaks *et al*, 1989).

Kriging utilises the theory of regionalised variable which is founded on the notion that spatially distributed data behaves more like random variables and should therefore be treated stochastically (Oliver *et al*, 1990). The theory assumes that the spatial variation in the data can be described as the sum of three components, namely the deterministic structural component corresponding to either a constant mean or a trend, the stochastic spatially correlated component, and the spatially independent residual term (Burrough *et al*, 1998). In its simplest application, the theory assumes a constant local mean and a stationary variance of the differences between places separated by a given distance and direction; this constitutes the intrinsic hypothesis (Lam, 1983).

The variance of the differences, denote λ is the semi-variance. Formally

$$\text{Var}[z(x)-z(x+h)] = E[(z(x)-z(x+h))^2] = 2\gamma(h) \quad (1)$$

where $z(x)$ is the value of some attribute at position x , and $z(x+h)$ is the value at position $(x+h)$. This semi-variance value depends on the separation distance between the points, h ; the actual positions are not relevant.

Kriging is similar to inverse distance weighted averaging in that it uses a linear contribution of weights for calculating value estimates, but differs in that the weights are derived according to the variance minimisation and unbiasedness criteria that it self-provides. Kriging explores the nature of spatial autocorrelation in the data and produces a semi-variogram to reveal the spatial variation underlying the data. The semi-variogram conveys useful information concerning the size, orientation and shape of the neighbourhood from which the sample points are drawn. Kriging also provides a measure of the error or uncertainty of the estimated surface (Lam, 1983).

Kriging analyses the pattern of spatial variation present in the data. The character of this variation is captured in functions such as the autocovariogram and (semi)-variogram, which provide the information for optimising interpolation weights and search radii. Unlike most other interpolation methods, kriging involves an interactive investigation of the spatial behaviour of the phenomenon of interest represented by the z values prior to the selection of the best interpolation strategy for generating output surface.

a. The rationale of Kriging for the current case study

In a sense, this study performs a revisit of the surface interpolation problem in the context of property value prediction. As indicated earlier, Byrne *et al* (1973) attempted the use of trend surface analysis (TSA) to achieve the prediction of property value. In that study, the authors found the method of limited value in investigating local detail such as required by the problem. We argue here that TSA is by nature a global interpolator and is therefore

not the most appropriate tool for dealing with short-range or local influences. Kriging is different in the sense that it investigates local patterns of spatial dependence and uses the analysed information to interpolate the values of interest. Due to this, it is thought that kriging would make for a more appropriate tool on the problem of such nature as the one at hand. This motivates the experimentation with kriging for the current investigation.

Most methods of interpolation neither provide the means for determining the number of sample points, the size, shape and orientation of the sample neighbourhood to use, nor look beyond the simple function of distance for the estimate of interpolation weights; the errors of estimates are not given (uncertainties associated with interpolated values). Kriging provides all of the above.

It has long been recognised that property prices tend to be similar for properties nearer to one another. This is explained in terms of a multitude of factors, but spatial separation has been known to play a significant contribution. This is so when dealing with point observations because units closer together will have similar values, concordant with the valuers' tone of the list assumption applied in rating valuation (Wyatt, 1994). Estimates of the dependent variable are made on the basis of location rather than by reference to independent variables (Shaw *et al*, 1985).

In terms of predictive performance, kriging has emerged superior to most other interpolation methods empirically. Isaaks *et al* (1989), in

conducting the comparison using the dataset from geological activity, shows that the ordinary kriging estimates not only lead to lower standard deviation of errors, but also that "the estimates are also very good according to many other criteria" such as the mean absolute error and mean squared error. In another study, by Burrough *et al* (1998), kriging performs favourably against other methods of interpolation. These results, although specific to the context of the individual studies, indicate kriging as capable of improving the quality of prediction.

Kriging is unique compared to other estimation procedures in that it does not limit the weights to between 0 and 1. Rather, it extends the weights' boundaries to include negative values as well as values greater than unity. As a result, it allows the possibility of estimates that are not necessarily constrained to the minimum-maximum range as defined by sample values. This allows estimates that lie beyond the minimum and maximum of sample values, which is useful because in reality, there is also the likelihood that the true values being estimated lie beyond the extremes of the available samples. Procedures that restrict the weights to within 0 and 1 can only attain estimates that lie between the minimum and maximum sample values.

The Methodology of Ordinary Kriging

For a more comprehensive treatment of the ordinary kriging methodology, the reader is referred to Isaaks *et al* (1989). The steps involved in ordinary kriging are as follows:

1. Compute the experimental variogram and deduce from the output whether it is feasible to interpolate the data
2. If feasible, use a suitable variogram model to generate value surface on a regular grid
3. Use the surface to interpolate values at unvisited sites

a. Computing the experimental variogram

The experimental variogram is the first step towards a description of the regionalised variation. It provides useful information for interpolation, optimising sampling and determining spatial patterns. The variogram reveals the nature of the variance-covariance structure given by the actual data and provides an insight into the pattern of spatial continuity present in the dataset.

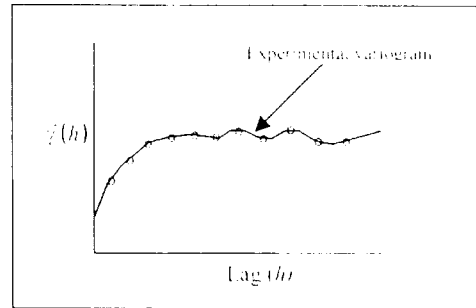
If the conditions specified by the intrinsic hypothesis are fulfilled, the semi-variance $\hat{\gamma}(h)$ can be estimated from sample data:

$$\hat{\gamma}(h) = \frac{1}{2n} \sum_{i=1}^n \{z(x_i) - z(x+h)\}^2$$

where n is the number of pairs of sample points of observations of the values of attributes z separated by distance h . A plot of $\hat{\gamma}(h)$ against h is known as the *experimental variogram* (Isaaks *et al*, 1989).

With h representing the lag, a typical experimental variogram takes the appearance of a curve that shows a steep rise from lag 0 but which changes to a more gradual one at larger h to eventually reach some kind of a plateau. The shape is illustrated in Figure 1.

Figure 1: A Typical Shape for Experimental Variogram



Since the experimental variogram is a plot of the variances of difference against lag distances, it provides an indication about the nature of spatial dependence present in the data. The shape derives from the fact that the spatial dependence between data points is greater when two points are at a shorter distance of one another than when they are further apart. The rising part of the curve describes how inter-site differences are spatially dependent: the closer the sites together, the more similar their z values are, as indicated by the low semi-variance values.

b. Modelling the experimental variogram

The next step is to decide on the suitable model to fit the data, taking into account the configuration of the experimental variogram, the sill and the nugget values, which are further explained below. A curve from mathematical models is fitted to experimentally derived semi-variances in order to describe the way in which semi-variance changes with the lag (Burrough, 1986). This curve displays several important features. First, at large lag values, it levels to what is known as a *sill*, implying that at these lag values, no

spatial dependence exists between the data points because all estimates of variances of difference are invariant with distance. Second, the curve rises from a low value to the sill, reaching it at a value h , known as the *range*. This describes the range of spatial distances at which inter-site differences are spatially dependent. Within this range, the closer together the sites, the more similar the values at the sites are. The *range* gives an idea about the size to consider for a search window. If the distance separating an unvisited site from a sampled point is greater than the range, then the latter can make no useful contribution to interpolation – it is too far away.

The *range* of the variogram therefore provides information about the size of the search window to consider. In effect, it defines the radius of distance from the point under investigation within which sample data points should lie to be considered influential to the estimation. These distances can vary as a result of anisotropy, which modifies the shape of the search neighbourhood from a circle to an ellipse.

Another feature of the fitted model is that it does not necessarily pass through the origin but cuts at a positive value of $\hat{\gamma}(h)$ despite the theoretical assertion that the semi-variance should be zero at lag 0. This situation arises because the positive value estimates the residual, spatially uncorrelated noise g^n . Also known as the nugget, g^n represents the variance of measurement errors combined with that from spatial variation at distances much shorter than the sample spacing, which cannot be resolved (Isaaks *et al*, 1989).

A number of variogram models are possible but the more common models in use are the spherical, exponential and gaussian models. Basic variogram models can be divided into two broad groups, known generally as transition and non-transition models. Transition model is 'bounded' in the sense that its variogram reaches towards an upper bound in the sill. Non-transition model, on the other hand, is unbounded since its variogram rises continuously as a function of lag distance h . Some transition models reach their sill asymptotically; for such models, the range is arbitrarily defined to be the distance at which 95% of the sill is reached (Isaaks *et al*, 1989).

Briefly, the spherical model has a linear behaviour at small separation distances near the origin but one that flattens out at larger distances to reach the sill at a . The tangent at the sill crosses the sill at about two-thirds of the range a (Isaaks *et al*, 1989). This model is normally used where there is a clear range and sill (Burrough *et al*, 1998). The exponential model is linear at very short distances near the origin but rises more steeply compared to the spherical model and flattens more gradually. The model is more appropriate where the approach to the range is more gradual. The Gaussian model has a parabolic shape near the origin and is often used to model extremely continuous phenomena. Like the exponential model, the Gaussian model reaches its sill asymptotically and the range is defined as the distance at which the variogram value is 95% of the sill. It is the only transition model whose shape has an inflexion point.

c. Fitting a model

Variogram fitting is an interactive process requiring considerable judgement and skill (Burrough *et al*, 1998). Model fitting is in order to produce the values of the parameters a , c_o and c_f . These values can be obtained by least-squares or maximum likelihood procedure.

When the nugget variance dominates the local variation and the experimental variogram shows no tendency to diminish as $h \rightarrow 0$, the interpretation is that the data are so noisy that interpolation is not sensible. In such a situation, the best estimate of $z(x)$ is the overall mean computed from all sample points in the region of interest without taking spatial dependence into consideration.

A noisy variogram, in which the experimentally derived semi-variances are scattered, suggests that too few examples have been used to compute it. As a rule of thumb, at least 50 - 100 data points are necessary to achieve a stable variogram although smooth surfaces require fewer points than those with irregular variation. Smoother variograms can also be obtained by increasing the size of the search window.

The presence of a hole effect in the experimental variogram (a dip in the semi-variances at distances greater than the range) may indicate a pseudo-periodic pattern due to long range variation over a study area that is too small to encompass the total range of variation (Burrough *et al*, 1998). If the range is large, then long-range variation dominates: if it is too small, then the major variation occurs over short distances.

d. Dealing with the directional issues

Anisotropy in the experimental variogram suggests a directional effect in value pattern, but directional differences can also occur if there are insufficient samples to get robust estimates in all directions. In many cases where samples are spaced irregularly, a circular search radius is used to define a zone whose mid-point is h from its centre. All data points falling within the circle are used to estimate the contribution of $(z_i - z_j)^2$ from all pairs. If directional effect is absent, the resulting variogram is isotropic, i.e., it results from averaging over all directions (Burrough *et al*, 1998). However, variograms can also be computed in specific directions β , in which case they are known as anisotropic variograms. If different ranges and sills are obtained for different variograms, they may indicate spatial variation that varies with direction.

Kriging the Property Value: A Case Study

The implementation of this kriging case study was performed with GS+, a commercial geo-statistical package available from Gammadesign. This software provides the functionality necessary for performing the various tasks required in kriging analysis. Further, it generates its output in ASCII files which can be read into other GIS application software such as ArcView.

As a tool for kriging, GS+ is rather versatile (Robertson, 1998). It copes well with interactive needs of the user. Its interactivity allows models to be refined, or the parameters to be adjusted on the fly. This is very useful particularly when the

need is to consider several alternative models instead of just one. Speed is also its plus point. Finally, the 3-d mapping it provides along with the zoom and rotation capabilities allow the user full control over the display.

a. Preparing the Data

The dataset for this case study comes from simulated house prices in Newcastle upon Tyne, UK. It embraces a total of 37,812 housing properties located within the eight sub-areas of Benwell, Byker, Fenham, Gosforth, Heaton, Jesmond, Kenton, Longbenton and Walker. Each house is represented by its seed point in the digital map, and this has been extracted from the Landline data, which provides the original positional information.

The data was then split into two smaller sub-samples in the proportion of 80% to 20%. The larger sub-sample, consisting of 30,250 data points, is to be used in the modelling. The smaller sub-sample consisting of the remaining 7,562 data points is to be retained as an independent 'test' sample for the purposes of tests on the models derived with the first sub-sample.

In undertaking this case study, three effects are of particular interest to investigate in terms of their influences on the performance of kriging models: first, the effect of sample data density; second, the effect of sample stratification by house type; and third, the effect of variable normalisation.

In the real world, house price data are not as abundant as the simulated data suggests since such data does not regularly become available.

House price data arises when the property undergoes market transactions, but for any particular property, this does not occur with any regularity. Further, property value is not static and changes over time. Since value has its validity period as dictated by the market to which the property relates, not all price data is relevant for a particular time-period of interest. This introduces a further limitation to the availability of value data. To study the influence of data availability on kriging's predictive performance, tests at three levels of data density corresponding to 5,000, 10,000 and 15,000 house units will be performed.

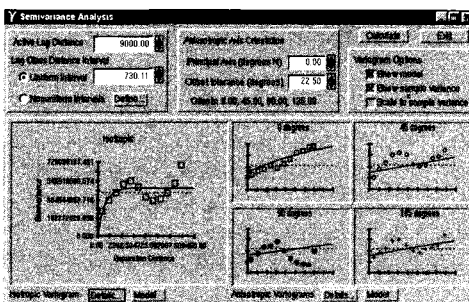
It is commonly observed that different house types have different characteristics and this heterogeneity leads to their different value classes in the market. House type is therefore a significant influence to consider and needs to be accounted for in modelling. Kriging, due to its single-variable nature, has no intrinsic means for dealing with this problem. As such, it would also be of interest to investigate if separate modelling of each property type would lead to better performance with kriging.

So far, the basis of interpolation is the value for the total property (land and building). However, it is also of interest to investigate the effect of value normalisation on kriging's predictive performance. Value normalisation in this study means the 'devaluation' of whole property value in which the value of a house is averaged over its parcel size. Underlying this approach is the supposition that normalising this way will improve homogeneity in the data and hence help improve the modelling with kriging.

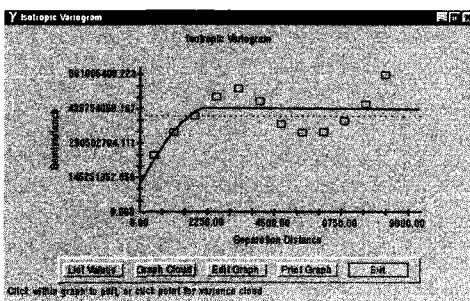
b Variogram modelling on property value

Figure 2 shows the isotropic variogram produced from 10,000 house values in this study. Plotting at uniform lag distance intervals of 730 metres, the plot shows a typical rising trend, hitting the peak at an approximate lag distance of 3,400 metres followed by a declining trend to the trough at about 5,400 metres and a rising trend again thereafter. The whole shape takes the appearance of a reflected S-curve. The semi-variance analysis for the data taken at lower and higher densities of 5,000 and 15,000 points respectively reveals similar variogram shapes, as shown in Figure 3.

Figure 2: Isotropic Variogram for the 10,000 Data Points

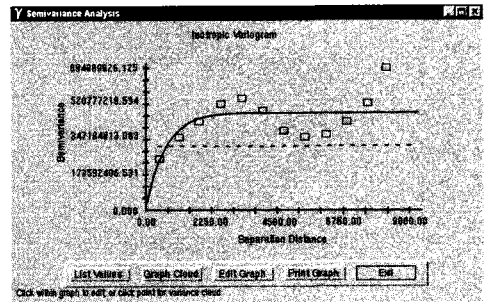


Semi-variance analysis based on 10,000 points

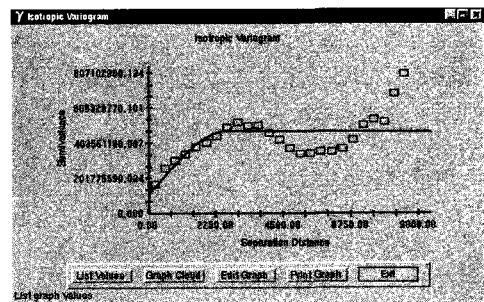


A close-up of the isotropic variogram

Figure 3: Isotropic Variogram for the 5,000 and 15,000 Data Points Respectively



Semi-variance analysis based on 5,000 points

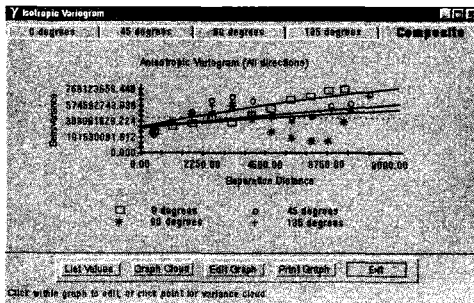


Semi-variance analysis based on 15,000 points

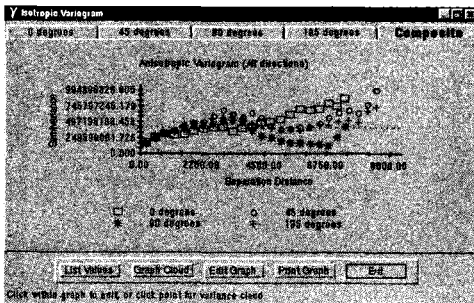
The variogram indicates spatial dependence that weakens as separation distance increases up to around 3,400 metres but this dependence seems to grow thereafter for some larger separation distances before diminishing again finally. This behaviour of spatial dependence is generally typical apart from the temporary dip, which is rather anomalous when the ideal experimental variogram would show a clear and horizontal sill after 3,400 metres. In the circumstances, it seems best to treat the data with one of the transition models, by assuming the presence of a sill that cuts a path roughly midway between the peak and the trough. Indeed the default model fitted by the GS+ is precisely of this nature.

The variogram in Figure 4 shows no clear anisotropy or directional

Figure 4: An Analysis of the Multi-Directional Effect of Spatial Dependence (anisotropy)



Variograms for the multi-directional semi-variance analysis on 5,000 data points



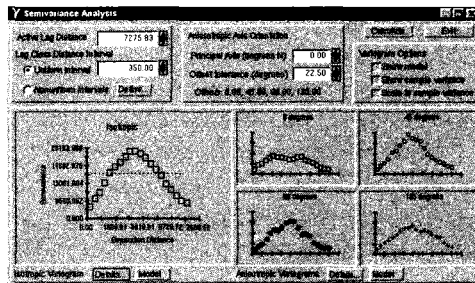
Variograms for the multi-directional semi-variance analysis on 15,000 data points

variation of spatial dependence in the data, and this negates the need to consider models of non-onnidirectional nature.

c. Modelling for property value as normalised by parcel size

Figure 5 shows the experimental variogram produced from the 5,000 data points by the normalised value variable. The semi-variogram pattern seen here defies the textbook description of an ideal variogram shape. Here, the typical rising trend at the beginning is followed by an ever declining trend to give the overall appearance of a peaked hill. This appears to indicate that at separation distances larger than about 3,300 metres, the square metre value of houses tends to become

Figure 5: The Variogram for the Normalised Value Based on the 5,000 Data Points

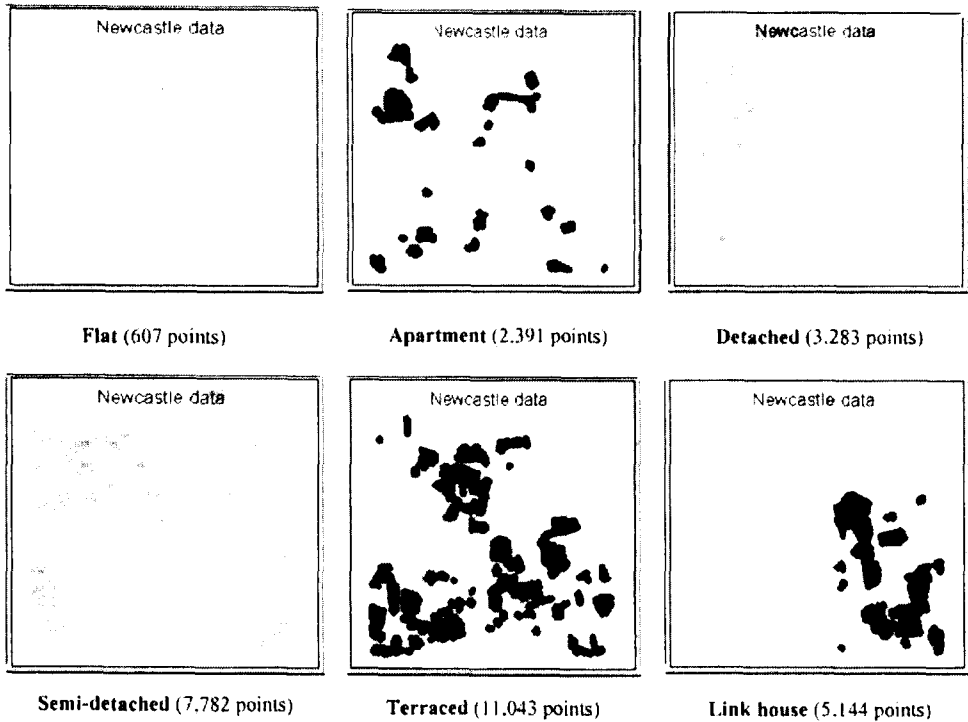


more similar again, in fact more similar the further the houses are apart. This behaviour is neither intuitive nor easy to explain in the present data set. The nature of spatial dependence in the data as described by the experimental variogram appears too anomalous to be treated with kriging. On this basis, the planned extension to this investigation involving the normalised value is abandoned and is not pursued further in this case study.

d. Modelling by property type

Modelling by property type requires separate semi-variance analysis on the data points on different property types. The implications are two-fold here: each property type is a much smaller sample size than the original work dataset, and each property type is an uneven distribution of data points spatially. Figure 6 shows the spatial distribution and sample size for each house type.

It is clear from Figure 7 that the variograms are not as smooth as for the whole work dataset. The combined effects of reduced sample size and the scatteredness of data points could have contributed to this situation. The variogram for flat

Figure 6: The Spatial Distribution of Houses by Type

units, which has the lowest number of points, is particularly jagged. It is also noticed that variogram smoothness in this data generally improves for the larger sample house types although the semi-detached variogram provides an exception; this is certainly true in the case of the link house.

The variogram for the apartment data is devoid of a sill but instead shows a continuous linear rise over the spatial extent considered. For such spatial structure, theory recommends the use of a linear model (Burrough *et al*, 1998).

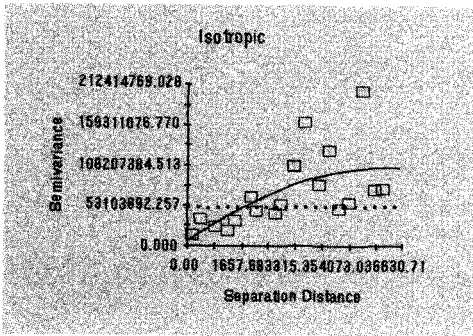
e. Value interpolation by Kriging

Given the information from their respective variograms, kriging interpolations are performed on the independent 7,562 point test dataset.

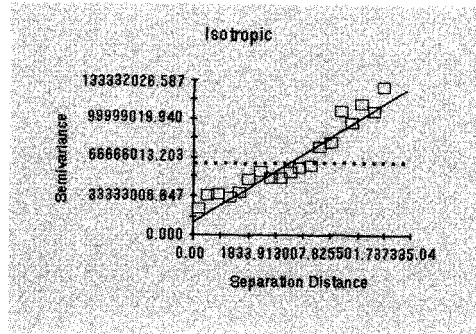
This allows the house value represented by each point to be estimated by the kriging method. Since the spatial structure in the normalised values does not provide a clear case of spatial dependence for modelling with kriging and also, since the patterns of spatial dependence for individual property types are too erratic, the interpolations based on the individual types are abandoned.

A series of kriging interpolation are performed based on the three levels of sample density used for the weight calculations. The interpolated values are then compared against their corresponding original final values in the dataset and the *PRESS* statistics calculated. Although kriging does produce estimates of errors, these estimates are not looked at because the *PRESS* statistics present a more

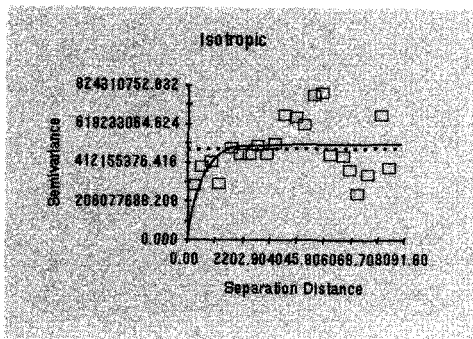
Figure 7: The Isotropic Variograms by House Type



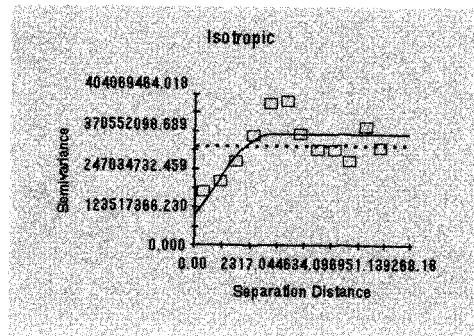
Flat



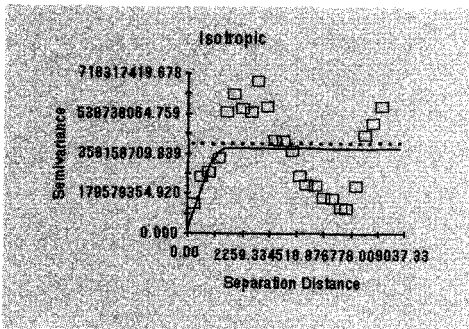
Apartment



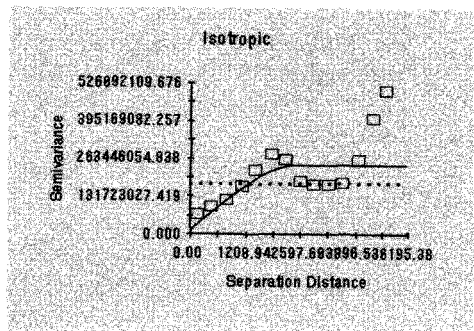
Detached



Semi-detached



Terraced



Link

desirable criterion for the comparison of performance due to their tests on independent samples.

The *PRESS* statistic is defined as: where,

$$PRESS = \sum_{i=1}^n (P_i - A_i)^2$$

N = the number of observations in the test sample

P_i = predicted house value

A_i = the observed house price

The aim is to achieve minimum value for *PRESS*. The model with the smallest *PRESS* statistic prevails as the best model for predictive performance.

The results on four of the interpolation runs are of interest and their *PRESS* sum calculations are presented in Table 1.

Table 1: The *PRESS* Statistics Obtained for the Different Levels of Data Density Used in the Kriging of Independent Test Data Points

Data density and model fitted	<i>PRESS</i> Sum	<i>PRESS</i> Maximum
Spherical model on 5,000 data points	915,481,005,731	25,913,695,301
Exponential model on 5,000 data points	864,920,477,807	26,800,453,327
Spherical model on 10,000 data points	834,042,753,822	25,876,385,828
Spherical model on 15,000 data points	793,489,424,232	26,114,944,609

There is a progressive reduction in the *PRESS* sum with greater number of data points used. Since lower *PRESS* sums are associated with higher predictive performance, the above results suggest that overall, the use of a greater number of sample points has led to improved predictive power of kriging in this investigation. At the individual property level, however, the position is not so clear. The statistics on *PRESS* maximum show that increasing the sample size does not necessarily reduce the gap between the actual and the predicted values in the property with the largest value difference.

Comparing the Predictive Performances of Kriging and MRA

For the comparison of predictive performance between kriging and MRA, *PRESS* statistics are again used. They are computed from the predictions made by both the methods on house values in the common (7,562 points) test dataset. The *PRESS* statistics for the kriging models are already available as a result of the comparison performed on the effect of different sample sizes used on predictive performance. It is now necessary to obtain similar statistics for the MRA models.

For the MRA, two models have emerged as the best in terms of predictive performance in this research. The models are *EnterA* and *StepA* (Appendix A). To arrive at the

prediction of these models on the test data, their equation forms are applied on the appropriate variables in the dataset concerned. Once the predicted values for each model have been obtained, *PRESS* statistics are calculated in the same way as before.

Table 2 presents the *PRESS* statistics for both the kriging and MRA models. The *PRESS* sums indicate that the kriging models have not outperformed the MRA models in terms of predictive performance. In fact, the *PRESS* sums of the latter are more than four times lower than that achieved by the best model from kriging. The means and standard deviations of the *PRESS* suggest that the variability in the gap between predicted and actual values is much greater in kriging than in the MRA. Given the fact that the MRA models have been derived with a larger sample of 30,250 data points and bearing in mind the finding that predictive improvements have been achieved with successive increases in sample size used, it is interesting how much further kriging models would have improved if this investigation has had the opportunity of modelling them with the larger sample.

However, kriging results do have their interesting aspects too in this investigation. As Table 3 shows, the minimum predicted values obtained by the kriging models are

Table 2: PRESS Statistics Obtained from Modelling the Independent Test Dataset: Kriging V. MRA

Model	PRESS Sum	PRESS Minimum	PRESS Maximum	PRESS Mean	PRESS Std. Dev
KRIGING					
<i>15000, spherical</i>	793,489,424,232	0.29	26,114,944,609	104,931,159	423,214,465
<i>10000, spherical</i>	834,042,753,822	0.38	25,876,385,828	110,293,937	430,836,079
<i>5000, exponential</i>	864,920,477,807	2.69	26,800,453,327	114,377,212	441,553,144
<i>5000, spherical</i>	915,481,005,731	0.10	25,913,695,301	121,063,344	435,826,102
MRA					
<i>Stepwise (StepA)</i>	181,942,677,523	0.23	3,714,051,752	24,060,127	82,654,039
<i>EnterA (EnterA)</i>	181,972,992,980	1.93	3,719,835,025	24,064,136	82,697,784

much closer to the actual minimum property value in the sample compared to those of the MRA models. Further, the kriging models do not produce negative predictions whereas the MRA models do. If this is taken in isolation, it means that kriging is more realistic than MRA since their estimates are more 'acceptable' in terms of the common perceptions in the real estate community. The krigings' predicted

means are in line with the mean of the actual value. So are the standard deviations; in fact the krigings' standard deviations are smaller than that of the actual value. Unfortunately, the ceiling values of prediction in kriging are much lower than the maximum of the actual value. This has meant that very poor predictions have been made on the properties with the very largest actual values.

Table 3: Descriptive Statistics for the Predicted and Actual Values

	Range	Minimum	Maximum	Mean	Std. Dev.
KRIGING					
<i>15000, spherical</i>	134,014.74	9,656.66	143,671.40	43,143.26	17,565.00
<i>10000, spherical</i>	173,495.59	8,361.57	181,857.16	43,209.54	18,506.91
<i>5000, exponential</i>	149,443.39	9,010.08	158,453.47	43,063.52	17,429.49
<i>5000, spherical</i>	131,490.60	9,441.27	140,931.87	43,079.91	16,699.61
MRA					
<i>Stepwise (StepA)</i>	192,053.77	- 7,641.80	184,411.98	43,291.57	19,700.88
<i>EnterA (EnterA)</i>	192,055.74	- 7,691.19	184,364.55	43,291.72	19,701.67
ACTUAL VALUE	238,271.00	7,084.00	245,355.00	43,346.78	20,533.08

Conclusions

This study shows that kriging is a poorer predictor of property values compared to MRA. However, this has to be set against the fact that a simulated dataset has been worked with and that this dataset has been

geared directly towards modelling with MRA. Given that this is the case, the results are perhaps not too surprising. It would be interesting if we can make similar comparisons based on real data where both methods are on the same level of advantage or disadvantage.

Despite its comparatively inferior predictive performance, the potential of kriging for mass appraisal of housing properties is perhaps not to be dismissed altogether. The experimental variograms arrived at in this study show that the patterns of spatial dependence exist in property values, suggesting that the exploitation of spatial correlation for value determination does have its basis. It is just that reliance on this pattern of spatial dependence alone may not be good enough to arrive at estimates that compare favourably with MRA, or perhaps that the kriging models are in need of further refinements. Further investigations are necessary. As the bottom line, kriging should be useful where the concern is with the investigation on locational factor in isolation in valuation as opposed to the investigation on property value in this study, which involves a multiplicity of factors.

The attraction of kriging comes from the fact that it utilises the information about localised spatial variation to estimate values at local positions. To arrive at this information, however, a large number of sample points and their fair distribution over the study area are important. It is this that probably makes the technique rather workable for housing properties, where the volume of data and their spatial omnipresence are relatively more favourable compared to most other types of properties. For properties that do not have such advantages (industrial properties, for example, are clustered around certain locations only), the practicality of a similar exercise remains to be tested and requires a separate study. For the moment we can only presume that the reliability of the estimates will be lower due to the greater presence of regions where no sample data points are available to draw information from for interpolation.

One issue remains particularly outstanding from this study: the lack of intuitive appeal that some of the kriging results provide. It is difficult to reconcile the fact that the

variogram is better behaved for the data that combines all the property types than for the data that has been made more homogeneous by dealing with only a particular property type. Probable explanations lie in the reduced sample size and the unevenness of spatial spread that occurs in the individual property dataset, but these are just possibilities. Could there be other more valid explanations?

It is argued that a major issue with kriging lies in the fitting of appropriate models to interpolate. This is because the fitting involves examining the variogram plots and choosing a model that is considered the best fit, a procedure that can entail arbitrary decisions on the part of the user. Collins (1996), for example, remarks that kriging has been criticised due to the subjective nature of variogram fitting – a central component of kriging. Nonetheless, arbitrariness is not something the user can avoid completely in dealing with problems of this nature. For that matter, not even the MRA can claim to be entirely free from arbitrary decisions, particularly in the choice of variables and of equation forms to use. On this score, kriging cannot be said to be any less desirable than the regression technique.

On the basis of the above initial evaluation, this study recommends that kriging should be investigated further before decisions are made about its utility for valuation. This constitutes yet another benefit the consideration of geography contributes towards the practice in valuation.

References

- Burrough, P. A. and McDonnell, R. A. (1998), *Principles of Geographical Information Systems*, Oxford University Press
- Burrough, P. A. (1986), *Principles of Geographic Information Systems for Land Resources Assessments*, Clarendon Press, New York

- Byrne, P. J. and Bowcock, P. (1973), *The Structure of Residential Property Values in St. Albans*, College of Estate Management, University of Reading
- Christensen, R. (1991), *Linear Models for Multivariate, Time Series, and Spatial Data*, Springer - Verlag, New York Inc.
- Collins Jr., F. C. (1996), A Comparison of Spatial Interpolation Techniques in Temperature Estimation, The Third International Conference/Workshop on Integrating GIS and Environmental Modelling, January 1996, Santa Fe
- Isaaks, E. H. and Srivastava, R. M. (1989), *An Introduction to Applied Geostatistics*, Oxford University Press, New York
- Lam, N. S. (1983), Spatial Interpolation Methods: A Review, *The American Cartographer*, 10(2):129-149
- Oliver, M. A. and Weber, R. (1990), Kriging: A Method of Interpolation, *International Journal of Geographic Information System*, 4(3):313-332
- Robertson, G. P. (1998), GS+: Geostatistics for the Environmental Sciences, Gamma Design Software, Planiwell, Michigan, USA
- Shaw, G. and Wheeler, D. (1985), *Statistical Techniques in Geographical Analysis*, John Wiley & Sons, Chichester
- Tobler, W. R. (1979), 'Smooth Pycnophylactic Interpolation for Geographical Regions', *Journal of the American Statistical Association* 74(367):519-536
- Wackernagel, H. (1995), *Multivariate Geostatistics*, Springer, Berlin
- Wyatt, P. J. (1994), Property Valuation Using a GIS: An Investigation of the Potential Impact that a GIS-based Property Information System Will Have on Property Valuation, with Particular Reference to the Analysis of the Spatial Element of Retail Property Value, Ph.D. thesis, University of Brighton, U.K

Appendix A

Comparison of Regression Coefficient from the Different Regression Models

	NEWCI1 DATASET (24,033 records)			
	UNLOGGED		LOGGED	
	FULL VARIABLE			
	EnterA	StepA	EnterB	StepB
(Constant)	- 2,286.615	- 1,454.892	14,923.027	15,188.478
ACCESS	2,484.021	2,361.140	2,377.030	2,338.168
AGE	- 239.279	- 239.322	-	-
BLD_AREA	189.742	189.771	190.577	190.525
C_FIN	455.750	455.487	466.173	465.982
DF_SCH	- 1.097	- 1.087	- 0.545	- 0.567
DISTCITY	- 2.503	- 2.504	- 2.536	- 2.535
DISTROAD	0.750	-	0.220	-
DISTTOWN	- 1.359	- 1.357	- 1.442	- 1.444
DSTMETRO	- 0.302	- 0.307	-	-
GEOG1	-22,623.081	- 22,632.347	- 20,957.772	- 20,967.179
GEOG2	-19,876.360	- 19,874.560	- 19,079.748	- 19,091.350
GEOG3	- 20,161.101	- 20,179.121	- 18,599.337	-18,614.185
GEOG4	2,586.272	2,583.948	3,245.591	3,246.849
GEOG5	24,098.266	24,095.323	24,504.722	24,502.517
GEOG6	10,676.208	10,676.893	10,913.277	10,918.283
GEOG7	-	-	-	-
GEOG8	-13,572.188	- 13,577.371	- 12,851.418	-12,841.671
GEOG9	-16,538.263	- 16,544.409	- 15,696.363	- 15,698.330
LANDAREA	83.290	83.253	83.486	83.542
LEVEL_NO	- 5,909.085	- 5,912.154	- 5,600.600	- 5,594.127
LNDMETRO	-	-	- 1,122.217	- 1,122.864
LNAGE	-	-	- 4,428.472	- 4,429.468
LNNOBATH	-	-	4,821.346	4,820.199
LNROOMNO	-	-	7,629.416	7,626.449
NBOR_QUA	3,144.258	3,147.381	3,212.139	3,212.718
NO_BATH	2,913.904	2,914.328	-	-
NOISE	- 1,324.245	- 1,363.038	- 1,402.315	- 1,413.945
Q_FIN	1,005.874	1,005.625	1,004.582	1,004.176
ROOM_NO	2,431.791	2,431.219	-	-
UNITYPE1	28.357	-	- 785.530	- 781.006
UNITYPE2	3,268.569	3,262.827	3,302.410	3,308.332
UNITYPE3	2,455.196	2,455.183	681.470	682.187
UNITYPE4	-	-	-	-
UNITYPE5	874.841	874.400	- 572.934	- 579.376
UNITYPE6	- 2,990.528	- 2,992.763	- 3,596.005	- 3,615.763
ZONE1	-	-	-	-
ZONE2	6,842.258	6,840.429	6,733.303	6,729.109
ZONE3	- 2,897.712	- 2,897.610	- 2,856.022	- 2,857.522
ZONE4	- 1,479.899	- 1,495.077	- 1,055.444	-
ZONE5	8,296.337	-	7,626.575	-
ZONE6	- 3,460.108	- 3,462.749	- 4,829.184	- 4,818.577

Note: The base reference of the above models is a hypothetical property of the LINK HOUSE type located in RESIDENTIAL zone in HEATON.

Appendix A1

Comparison of Regression Coefficient from the Different Regression Models (contd.)

	NEWC11 DATASET (24,033 records)			
	UNLOGGED		LOGGED	
	DROPPED VARIABLE			
	EnterC	StepC	EnterD	StepD
(Constant)	- 8,907.997	- 8,894.750	6,741.762	6,744.574
ACCESS	2,443.673	2,442.214	2,426.390	2,425.849
AGE	- 229.271	- 229.268	-	-
BLD_AREA	134.132	134.132	136.723	136.719
C_FIN	440.913	440.545	451.437	451.023
DF_SCH	- 1.427	- 1.436	- 0.864	- 0.866
DISTCITY	- 2.772	- 2.771	- 2.787	- 2.786
DISTROAD	-	-	-	-
DISTTOWN	- 1.159	- 1.161	- 1.241	- 1.242
DSTMETRO	- 0.197	- 0.199	-	-
GEOG1	-22,462.486	-22,459.119	- 20,852.729	-20,848.820
GEOG2	-19,587.310	-19,589.577	- 18,810.346	-18,811.732
GEOG3	-19,962.375	-19,963.678	- 18,431.178	-18,431.520
GEOG4	3,115.646	3,115.653	3,712.447	3,712.134
GEOG5	24,209.355	24,209.986	24,587.238	24,588.189
GEOG6	11,387.087	11,387.021	11,579.619	11,579.535
GEOG7	-	-	-	-
GEOG8	-13,018.798	-13,015.500	- 12,375.471	-12,374.898
GEOG9	-16,882.447	-16,880.815	- 16,063.152	-16,063.079
LANDAREA	118.690	118.695	117.635	117.634
LEVEL_NO	-	-	-	-
LNDMETRO	-	-	- 992.465	- 992.343
LNAGE	-	-	- 4,219.521	- 4,219.169
LNNOBATH	-	-	4,298.640	4,298.650
LNROOMNO	-	-	5,989.403	5,988.106
NBOR_QUA	3,152.117	3,152.194	3,217.570	3,217.661
NO_BATH	2,741.783	2,742.075	-	-
NOISE	- 1,252.817	- 1,252.382	- 1,306.638	- 1,305.851
Q_FIN	1,003.710	1,003.366	1,002.005	1,001.738
ROOM_NO	1,828.453	1,827.830	-	-
UNITYPE1	3,703.790	3,703.473	2,671.978	2,671.296
UNITYPE2	7,028.773	7,027.252	6,740.750	6,738.879
UNITYPE3	2,667.579	2,667.199	1,020.444	1,020.749
UNITYPE4	-	-	-	-
UNITYPE5	- 2,079.645	- 2,081.788	- 3,250.310	- 3,250.959
UNITYPE6	- 8,526.287	- 8,531.673	- 8,850.190	- 8,851.125
ZONE1	-	-	-	-
ZONE2	7,333.813	7,331.049	7,221.978	7,221.186
ZONE3	- 2,950.443	- 2,952.494	- 2,927.396	- 2,927.868
ZONE4	- 478.218	-	- 110.863	-
ZONE5	7,747.334	-	7,150.754	-
ZONE6	- 4,821.144	- 4,816.834	- 6,084.707	- 6,082.082

Note: The base reference of the above models is a hypothetical property of the LINK HOUSE type located in RESIDENTIAL zone in HEATON.