

University of Nebraska at Omaha ILL

ILLiad TN: 8917

Borrower: EAU

Lending String: VMW,*NBU,VNS,WEA,SSC

Patron: Jackson, Monica - au fac

Journal Title: Geographical analysis.

Volume: 37 **Issue:** 4

Month/Year: October 2005**Pages:** 371-382

Article Author:

Article Title: M. Jackson and L. Waller; Exploring Goodness-of-Fit and Spatial Correlation Using Components of Tango's Index of Spatial Clustering

Imprint: [Columbus] Ohio State University Press.

ILL Number: 23457334



Call #: G1 .G334

Location: WEB

ARIEL

Charge

Maxcost: \$25IFM

Shipping Address:

American University Library

ILL Borrowing

4400 Massachusetts Ave. NW

Washington, DC 20016-8046

Fax:

Ariel: 147.9.82.70 LVIS

ARIEL

9-26-06

NIC

12X

Exploring Goodness-of-Fit and Spatial Correlation Using Components of Tango's Index of Spatial Clustering

Monica C. Jackson, Lance A. Waller

Department of Biostatistics, Emory University, Atlanta, GA

The ability to detect anomalies such as spatial clustering in data sets plays an important role in spatial data analysis, leading to interest in test statistics identifying patterns exhibiting significant levels of clustering. Toward this end, Tango (1995) proposed a statistic (and its associated distribution under a null hypothesis of no clustering) assessing overall patterns of spatial clustering in a set of observed regional counts. Rogerson (1999) observed that Tango's index may be decomposed into the summation of two distinct statistics, the first mirroring standard tests of goodness-of-fit (GOF), and the second an index of spatial association (SA) similar to Moran's I. In this article, we investigate the effectiveness of Rogerson's expression of Tango's statistic in separating GOF from SA in data sets containing clusters. We simulate data under the null hypothesis of no clustering as well as two alternative hypotheses. The first alternative hypothesis induces a poor fit from the null hypothesis while maintaining independent observations and the second alternative hypothesis induces spatial dependence while maintaining fit. Using Rogerson's decomposition and leukemia incidence data from upstate New York, we show graphically that one is unable to statistically distinguish poor fit from autocorrelation.

Introduction

One purpose of spatial analysis in public health is to detect local clusters or anomalies in observed patterns of disease. Historically, the spatial analysis literature tends to fall into one of two groups regarding spatial clustering. The statistical literature often assumes independent regional counts and seeks to identify spatial trends (STs) in local disease risk, with particular interest in identifying areas inconsistent with patterns of disease risk defined by local values of relevant risk factors (e.g., age). In contrast, the geography and spatial econometric literature often builds inference based on global and local indices of spatial association (SA). In this

Correspondence: Lance A. Waller, Department of Biostatistics, Emory University 1518 Clifton RD, Rm 326 Atlanta, GA 30322
e-mail: lwaller@sph.emory.edu

Submitted: May 25, 2004. Revised version accepted: June 20, 2005.

article, we explore the use of Rogerson's (1999) expression of Tango's (1995) index of spatial clustering as a hybrid between these two types of approaches.

Rogerson (1999) noticed that Pearson's χ^2 statistic focuses on a summary of fit across individual regions rather than patterns of fit between regions. Thus extreme values of the statistic may not imply clustering, just lack of fit. In addition, Rogerson notes that traditional tests of SA (e.g., Moran's I) detect proximity of similar values but not proximity of local deviations (lack of fit) from a hypothesis of constant risk.

More specifically, Rogerson (1999) considered Tango's (1995) index, denoted T , and defined as

$$T = \sum_{i=1}^N \sum_{j=1}^N w_{ij} \left(\frac{y_i}{y_+} - \frac{n_i}{n_+} \right) \left(\frac{y_j}{y_+} - \frac{n_j}{n_+} \right) \quad (1)$$

where w_{ij} represents a spatial weight between region i and region j , N is the total number of regions, y_i denotes the number of cases observed in region i , n_i is the population size at location i ,

$$y_+ = y_1 + y_2 + \cdots + y_N \quad (2)$$

and

$$n_+ = n_1 + n_2 + \cdots + n_N \quad (3)$$

Verbally, T represents a spatially weighted sum of the product of the observed minus the expected proportion of cases in regions i and j , where the "expected" proportion corresponds to the population proportion, that is, the proportion of cases expected when all individuals are subject to the same probability of disease. We note a slight difference between Tango's use of the proportion of all cases and the more common epidemiologic proportion of cases among individuals at risk in a given region. We also note that Kulldorff, Tango, and Park (2003) define a test similar to Tango's index known as Tango's Excess Events Test which differs from equation (1) only by a constant.

Rogerson (1999) partitioned Tango's index, T , into the sum of (1) a goodness-of-fit component (GOF) and (2) a SA component, yielding

$$T = \overbrace{\sum_{i=1}^N w_{ii} \left(\frac{y_i}{y_+} - \frac{n_i}{n_+} \right)^2}^{\text{GOF}} + \overbrace{\sum_{i=1}^N \sum_{\{j:j \neq i\}} w_{ij} \left(\frac{y_i}{y_+} - \frac{n_i}{n_+} \right) \left(\frac{y_j}{y_+} - \frac{n_j}{n_+} \right)}^{\text{SA}} \quad (4)$$

where all variables are as defined in equation (1). Note that $\{j:j \neq i\}$ refers to the set of all j such that j does not equal i . In this setting, Rogerson and Tango define $w_{ii} = 1$, in contrast to the usual convention in spatial analysis of $w_{ii} = 0$. In this form, the first part of the right-hand side of equation (4) refers to a GOF component similar to the numerator of Pearson's χ^2 (i.e., a sum of squared "observed" minus "expected" terms). The second part of equation (4) refers to a SA component similar to the numerator of Moran's I when applied to the disease proportions and

where we replace the overall average proportion with the local expected proportion under an assumption of constant disease risk. For both components, we compare local “observed” values (the proportion of all cases in region i) to “expected” values (the proportion of the population in region i). Equating the expectation of the observed and population proportions provides us with the null hypothesis:

$$H_0 : \frac{E(y_i)}{y_+} = \frac{n_i}{n_+} \quad \text{for } i = 1, \dots, N \quad (5)$$

conditioning on the total number of observed cases, y_+ , and reflecting our assumption of constant risk (i.e., if every individual is subject to the same risk of disease, the expected proportion of cases in region i should correspond to the proportion of people in region i).

In this article, we take things a step further and explore the behavior of the two components under two specific alternative hypotheses with the same New York leukemia data used by Rogerson. One hypothesis addresses the GOF component and the other the SA component of equation (4). In fairness, our study perhaps asks more of the decomposition than was originally intended. Rogerson’s motivation for the decomposition was to see if the hybrid nature of Tango’s index provided improvements in statistical power over approaches focusing solely on either fit or spatial autocorrelation (e.g., the Pearson and Moran statistics, respectively). That is, Rogerson (1999) addresses the gains in *sensitivity* (detecting pattern when pattern is indeed present) in Tango’s index over the GOF or SA components alone, and we extend the work to address *specificity* (correctly identifying the specific type of alternative) of the two components.

The outline of this article is as follows: In the next section, we briefly review the leukemia data set. In the third section, we review the null distribution of Tango’s index. In the penultimate section, we specify our two alternative hypotheses. In the final section, we discuss results and draw conclusions.

Data

For this study, we use the New York leukemia data set collected by the New York Department of Health’s cancer registry, introduced by Turnbull et al. (1990) (see also Waller et al. 1992, 1994) and used by Rogerson (1999). The data include a total of 592 reported cases of leukemia among 1,057,673 people at risk during the time frame 1978–1982 in eight counties (Broome, Cayuga, Chenango, Cortland, Madison, Onondaga, Tioga, and Tompkins) of central New York. The counties are divided into 790 subregions (U.S. census tracts in Boome County, U.S. census block groups in the other seven counties). The data contain population counts and leukemia incidence counts for each subregion and are available in Waller et al. (1994) and online at <http://lib.stat.cmu.edu/datasets/csb/>.

Null distribution of Tango's index

Following Tango (1995) we define $\mathbf{W}$ to be an $N \times N$ weight matrix with elements

$$w_{ij} = e^{-d_{ij}/k} \quad (6)$$

and

$$w_{ii} = 1$$

Notice that the value of k affects the spatial scale of the test. We use values of $k = 1, 7, 15$, and 20 km for all simulations in this article, as these reflect a wide range of spatial neighborhoods for the New York data as specified in the results below. Also, at $k = 7$ all regions have $w_{ij(j \neq i)} > 0.05$ for at least one other region.

We may reexpress the null hypothesis defined in equation (5) as

$$H_0 : E[y_i] = n_i \lambda \quad (7)$$

where the y_i are independent Poisson random variables, n_i is the population size in region i , and λ is the overall rate (noting that we set $\lambda = y_+/n_+$, the overall observed rate, again conditioning on the total number of counts, y_+).

Note that H_0 involves both a fit component ($E[y_i] = n_i \lambda$) and an independence component. Assunção and Reis (1999) discuss the importance of defining the null hypothesis, noting that different tests of spatial pattern assume different null (and alternative) hypotheses, for example, some tests include independence in the null hypothesis (as we do here), and some do not. In our case, we are interested in deviations from H_0 that maintain fit (i.e., $E[y_i] = n_i \lambda$) but no longer maintain independence (SA only) or deviations from H_0 that do not fit but maintain independence (GOF only). For data following a Gaussian (normal) distribution, it is straightforward to treat the fit and independence components of H_0 separately, but (as we shall see) inducing such patterns in Poisson data is much more involved.

To observe Rogerson's result graphically using the simulated data, first we decompose Tango's index into the two separate parts: GOF and SA as defined in equation (4) for each of 1000 simulations. We then plot a graph of the pair (GOF, SA) for each simulated data set. Plotting points (GOF, SA) reveals a joint (bivariate) distribution of the two components of T . As inference for T is univariate and based on the sum of GOF+SA, the 95% critical value for T , denoted t^* , corresponds to the line

$$SA + GOF = t^* \quad (8)$$

also displayed on the plot. Using a simulation-based value of t^* , by definition, 5% of the (GOF, SA) pairs lie above and to the right of this line. We could also consider bivariate inference based on the joint distribution and return to this point in the discussion section. Figure 1 illustrates the (GOF, SA) values, under H_0 , for $k = 1, 7, 15$, and 20 . We note:

- (1) A consistent range of the GOF component, and

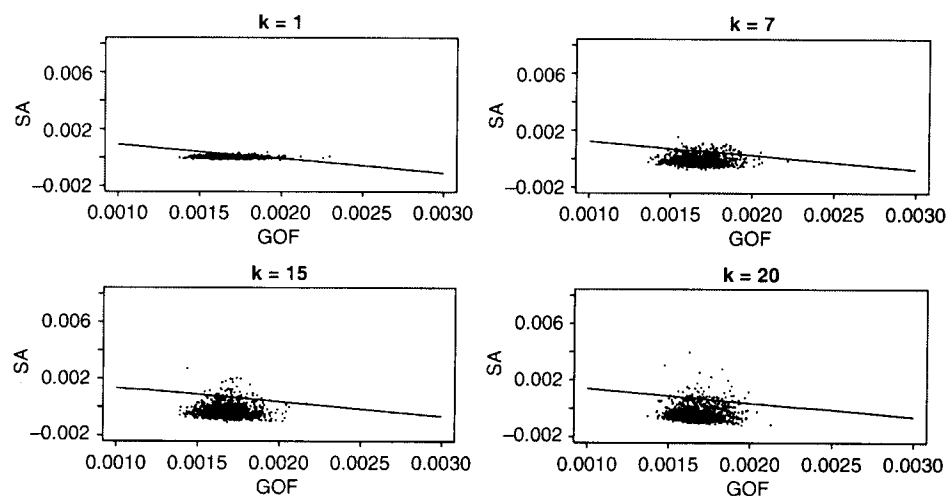


Figure 1. Plots of the goodness of fit (GOF) and spatial association (SA) components of Tango's index for the New York leukemia data for 500 data sets simulated under the null hypothesis. The line represents the 95% critical value of the statistic (5% of simulated values fall above and to the right of the line) under the null hypothesis for $k = 1, 7, 15$, and 20 .

- (2) An expanding range and skewed distribution for the SA component with increasing neighborhood (increasing k).

Two alternative hypotheses

We now define two alternative hypotheses which separately address the GOF and SA components of equation (4). We note that these two alternatives do not represent an exhaustive catalog of such alternatives, but represent an attempt to define families of clusters based on primarily first- (mean) or second-order (correlation) properties of the probability model defining the observed data. First-order clustering assumes the risk of disease may change from location to location but that individual cases remain independent of one another. Second-order clustering assumes that the risk of disease in a location is increased by the presence of other cases occurring nearby. (In reality, a combination could occur for an infectious disease with environmental influences.)

Simulating clustered counts based on only first- or only second-order properties in regional data with heterogeneous population sizes is a delicate task for several reasons. First and foremost, it is important to understand the potential interplay between the spatial scale of any clustering and the level of aggregation in the data. If cluster occurs at a scale smaller than the level of aggregation (and if we assume clusters do not cross regional boundaries), then, at the aggregate level, clustering would only appear to be first order, as neighboring counts would not be related by a shared cluster.

For the purposes of this article, we assume a constant mean disease risk for individuals within each region (ignoring any first-order variations within regions),

and limit attention to second-order effects at a scale larger than that of the regions. As noted, this does not represent a full spectrum of possible alternatives, but serves to illustrate several interesting points.

ST alternative hypothesis

We define a ST hypothesis relating to the GOF component of equation (4) as

$$H_{ST} : E[y_i] = n_i \lambda e^{x_i \beta}, \quad y_1, \dots, y_n \text{ independent} \quad (9)$$

where, n_i is the population at site i , λ is the total disease rate, and x_i is computed as

$$x_i = \frac{1}{\text{distance from cluster center to region } i\text{'s centroid}} \quad (10)$$

(Assume for simplicity, that the cluster center does not correspond exactly to a tract centroid.) Note that if $\beta = 0$, H_{ST} corresponds to a hypothesis of constant risk (λ). Essentially, H_{ST} creates a spatial pattern in the mean number observed even though we maintain independent counts, hence the deviation from the null hypothesis is due to a misspecified expected value (the null omits the $\exp(x_i \beta)$ term). Note that there is a spatial pattern in the means defined by equations (9) and (10). We avoid the term “spatial dependence” or “SA” here as we maintain (statistical) independence. The literature is not consistent on limiting the phrase “spatial dependence” to only spatial autocorrelation, but we prefer to make this distinction here as it is central to interpreting our results. We note that typical measures of SA such as Moran’s I or the SA component above are based on measures of statistical (auto)correlation. However, it is well known that measures of SA based on null expectations (here, $n_i \lambda$ rather than $n_i \lambda^* \exp(x_i \beta)$) are sensitive to spatially patterned trends in the alternative expectation, a feature confounding the separation between our two alternatives and a result we observe below. Briefly, the trend in expectations creates a measurable impact in the SA component as the null expectations define population proportions n_i/n_+ resulting in deviations from ST (regions with poor fit tend to be near other regions with poor fit). The sensitivity of Moran’s I (and related indicators) to multiple spatial alternatives is part of its appeal; however this results in reduced specificity for defining what sort of alternative is actually present in the data, a feature discussed in the regression setting by Anselin and Rey (1991).

For our simulations, we assign β to quadruple disease risk in the highest exposed subregion (block group or track). We denote the maximum exposure value M , and we choose β such that

$$e^{M\beta} = 4 \quad (11)$$

Thus,

$$\beta = \frac{\ln 4}{M} \quad (12)$$

As a side note, in each simulation, we place the cluster center at a new point, effectively averaging results across all possible cluster centers. The statistical power of tests often depends on the location of the cluster with respect to the edges of the study area and (often more importantly) due to variations in local population density (local sample size). Gangnon and Clayton (2001, 2004) and Rudd (2001) illustrate such geographic variations in power. While important in practice, we choose to average over cluster locations in order to focus attention on the sorts of alternatives under consideration.

SA alternative hypothesis

The SA alternative hypothesis is defined as

$$H_{SA} : E[y_i] = n_i\lambda, \text{ where the } y_i \text{ are correlated} \quad (13)$$

As noted above, simulating correlated Poisson counts without adding some lack of fit is not straightforward as the mean and variance of a Poisson distribution are not orthogonal parameters as they are in the normal distribution. In addition, we cannot simply rearrange Poisson random variates among regions, due to the dependence of each region's mean (and variance) on local sample size. An anonymous referee outlined an approach for inducing SA without impacting the expected counts via reorganizing the percentiles of observations such that high values are clustered together and low values are clustered together, and, after clarification and some experimentation, we implement a version of the suggested method below. We refer to the approach as "spatial rank redistribution" and compare results with those from the ST alternative above.

We simulate data under the spatial rank redistribution alternative as follows:

- *Step 1:* Generate $u_1, \dots, u_N$ as independent, uniformly distributed random variables between 0 and 1.
- *Step 2:* Order the N observed uniform random variates from lowest to highest and denote the ordered values as $u_{(1)}, \dots, u_{(N)}$.
- *Step 3:* Generate $(g_1, \dots, g_N)^T \sim \text{MVN}(\mathbf{0}, \Sigma)$ where Σ is the spatial variance covariance matrix defined with elements $\sigma^2 = \exp(-\tau * d_{ij})$. We fix τ at 0.1 and 1.0, for our examples. We compute d_{ij} as the Euclidean distance between regions i and j . This yields a set of spatially correlated, normally distributed random variates $g_1, \dots, g_N$.
- *Step 4:* Define $r_1, \dots, r_N$ as the rank of $g_1, \dots, g_N$. (e.g., if $g_c = \min(g_i)$ then $r_c = 1$).
- *Step 5:* Using the data in Step 2, assign $u_{(1)}$ to location j where $r_j = 1$, and so on for each $u_{(i)}$. At the conclusion of this step, the uniformly distributed variates will have been redistributed to new locations so they have ranks matching those of the multivariate normal data in Step 3.

- *Step 6:* Assign the number of cases in location j as the u_{ij} th percentile of a $\text{Poisson}(\lambda n_j)$ distribution where λ is the overall rate and n_j is the population at risk in location j . Moving from uniform random variates to Poisson values via percentiles represents the very general (but sometimes slow) “inversion” method for generating (pseudo)random numbers (Bratley, Fox, and Schrage 1987, pp. 182–183).

By using percentiles for each region’s associated Poisson distribution, we maintain local fit, but the spatial pattern in percentiles induces spatial similarity. That is, the result of this algorithm is a set of Poisson random variates with the same (heterogeneous) expectations as under H_0 whose observed values are spatially structured (higher percentiles are near high percentiles, and low percentiles are near low percentiles). However, the approach is not completely satisfactory in that the total number of cases assigned in this algorithm is no longer fixed to be 592. Uniform “thinning” of the sample would not be appropriate due to the heterogeneous means associated with each region. This said, the approach does allow us to explore sensitivity to SA patterns unobstructed by structural changes in fit.

Finally, we note that no alternative is “right,” rather each represents a different set of deviations from the null hypothesis, essentially a different family of “clusters,” and we may expect different performance versus each family of alternatives.

Results

Using the simulated data described above, we now explore the sensitivity and specificity of the components of Tango’s test. The plots below illustrate the (GOF, SA) values derived from data simulated under both the null distribution and under the two alternative distributions.

ST simulations

From the plots shown in Fig. 2 we assess the difficulties of separating the impact of the ST alternative between the GOF and SA components of Tango’s index. For each of the figures, the sensitivity or statistical power (probability of rejecting H_0 when H_{ST} is true) is available by calculating the proportion of simulated H_{ST} values above and to the right of the critical line. We observe that the shift of (GOF, SA) values is not solely in the direction of increased lack of fit (i.e., resulting in increased GOF values but maintaining similar SA values), indicating poor specificity in identifying the ST nature of the alternative. The increase in GOF values is accompanied by an increase in SA values, even though no spatial dependence is induced in the simulation. As noted above, this is due to the sensitivity of the SA component to spatial patterns in expectation differing from the null expectations, even with no spatial correlation the components are not specific to the correlated alternative. Hence, from the graph or observed values of the GOF and SA components we are unable to determine the GOF source of the clustering.

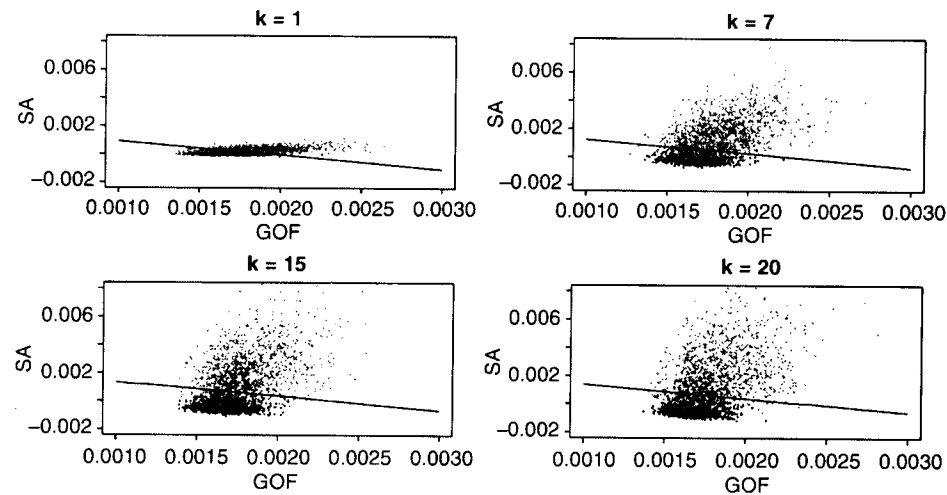


Figure 2. Plots of goodness-of-fit (GOF) and spatial association (SA) components of Tango's index (see text). Light points represent the density of (GOF, SA) pairs under the spatial trend (ST) alternative. Dark points represent the density under the null hypothesis. The line represents the 95% critical value of the statistic under the null hypothesis.

SA simulations

The plots shown in Fig. 3 show the results from the SA simulations based on spatial rank redistribution simulations. In the plots, we see (GOF, SA) pairs under the null hypothesis (dark dots) and under the SA alternative hypotheses (light dots). As with the ST alternative hypothesis, we again notice the difficulties in separating the impact on the GOF and SA components. We observe shifts in both the GOF and SA directions, rather than in the SA direction only. The reason for the shift in the GOF component (which is fairly consistent across the particular alternatives considered) is not immediately obvious, but may be due to the positive correlation of percentiles, which may generate more pockets of high percentiles than we might expect from independent observations. While these would be offset by pockets of low percentiles, the heavy upper tails of the Poisson distribution may result in a greater impact among high percentiles than low. As before, we are unable to distinguish between the ST and GOF alternatives based on the graphs alone, and again, the components are not specific to the source of spatial pattern.

Discussion

In this article, we apply Rogerson's mathematical decomposition of Tango's statistic into GOF and SA components. While the decomposition provides an excellent illustration of how Tango's statistic incorporates both lack of fit and spatial dependence, we explore whether the decomposition is useful in empirically identifying whether a detected pattern is driven more by lack of fit or by spatial autocorrelation.

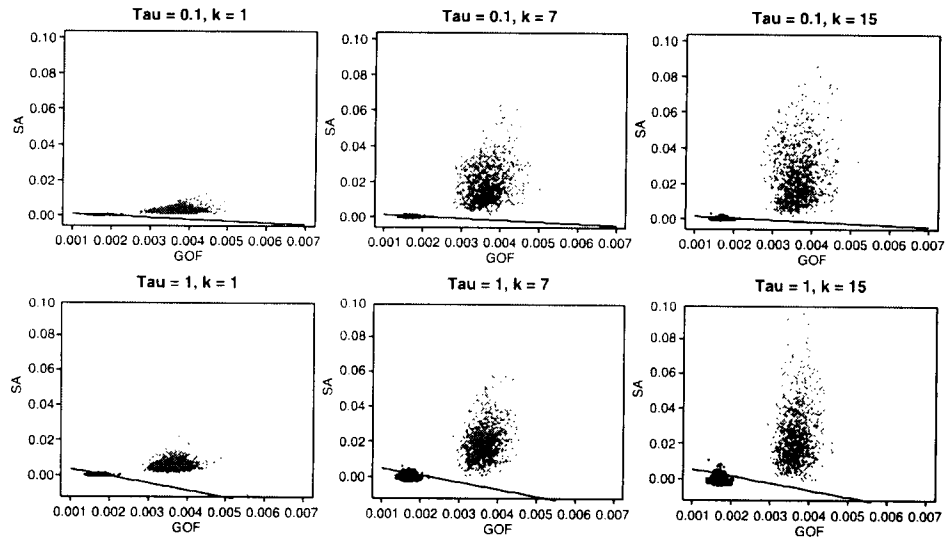


Figure 3. Plots of goodness-of-fit (GOF) and spatial association (SA) components of Tango's index for the New York leukemia data. Light points represent the density of the (GOF, SA) pairs under the spatial rank redistribution alternative (see text). The dark points represent the density under the null hypotheses. The line represents the 95% critical value of the test statistic under the null hypothesis.

If it were possible to graphically separate the GOF and SA components we would expect an upward shift for the SA alternative hypothesis away from the null. Likewise, we would expect a shift to the right from the null for the ST alternative hypotheses. Neither result was achieved. Instead both ST and SA alternative hypotheses yielded a diagonal shift from the null. Thus, even though mathematically Rogerson's decomposition splits Tango's index into two distinct parts with attractive interpretation, we must take care not to overinterpret their relative values, as they often do not provide information sufficient to categorize an observed pattern as driven primarily by GOF or SA. Our results serve to illustrate in a practical setting the theoretical findings of Bartlett (1964) regarding the mathematical intractability of separating ST (first-order effects) and spatial autocorrelation (second-order effects) from a single observation of spatial data. These results are also found in Bailey and Gatrell (1995, p. 33), Diggle (2003), and Waller and Gotway (2004, p. 150), who all note that first-order (expectation, here, GOF) and second-order (correlation, here, SA) effects do not uniquely identify a point pattern.

There are many directions for future research. The spatial rank redistribution algorithm requires additional consideration, for example, a more detailed exploration of the observed lack of fit and identifying optimal choices of the user-defined w_{ij} s in Tango's index to maximize power for particular spatial covariances in the latent Gaussian random variates. In addition, there is much work to be done in defining other algorithms to induce correlation in Poisson random variables

according to appropriate alternative hypotheses modeling clustering deviations of interest. Such simulation methods may build from generalized linear mixed models (Breslow and Clayton 1993), conditional autoregressive and simultaneous autoregressive correlation models (see Kaluzny et al. 1998, p. 128), or automodels (Ferrandiz et al., 1995). Furthermore, it may be possible to consider the bivariate distribution of (GOF, SA), although the plots above do not provide much optimism for such an approach.

Acknowledgements

This work was supported by a grant from the National Institute of Environmental Health Sciences (NIEHS) R01-ES07750. The opinions expressed are ours and do not necessarily represent those of NIEHS or of NIH. We also thank three anonymous referees for insightful constructive criticism of an earlier draft of the manuscript.

References

- Anselin, L., and S. Rey. (1991). "Properties of Tests for Spatial Dependence in Linear Regression Models." *Geographical Analysis* 23, 112–31.
- Assunção, R. M., and E. A. Reis. (1999). "A New Proposal to Adjust Moran's *I* for Population Density." *Statistics in Medicine* 18, 2147–62.
- Bartlett, M. S. (1964). "The Spectral Analysis of Two-Dimensional Point Processes." *Biometrika* 51(3–4), 299–310.
- Bratley, P., B. L. Fox, and L. E. Schrage. (1987). *A Guide to Simulation*, 2nd ed. New York: Springer.
- Breslow, N. E., and D. G. Clayton. (1993). "Approximate Inference in Generalized Linear Mixed Models." *Journal of the American Statistical Association* 88(421), 9–25.
- Bailey, T. C., and A. C. Gatrell. (1995). *Interactive Spatial Data Analysis*. New York: Wiley.
- Diggle, P. J. (2003). *Statistical Analysis of Spatial Point Patterns*, 2nd ed. London: Arnold Publishers.
- Ferrandiz, J., A. Lopez, A. Llopis, M. Morales, and L. Tejerizo. (1995). "Spatial Interaction Between Neighbouring Counties: Cancer Mortality Data in Valencia (Spain)." *Statistics in Medicine* 20, 2977–87.
- Gangnon, R. E., and M. K. Clayton. (2001). "A Weighted Average Likelihood Ratio Test for Spatial Clustering." *Statistics in Medicine* 20, 2977–87.
- Gangnon, R. E., and M. K. Clayton. (2004). "Likelihood-Based Tests for Localized Spatial Clustering of Disease." *Environmetrics* 15, 797–810.
- Kaluzny, S. P., S. C. Vega, T. P. Cardoso, and A. A. Shelly. (1998). *S+ SpatialStats*. New York: Springer-Verlag.
- Kulldorff, M., T. Tango, and P. Park. (2003). "Power Comparisons for Disease Clustering Tests." *Computational Statistics and Data Analysis* 42, 665–84.
- Rogerson, P. (1999). "The Detection of Clusters Using a Spatial Version of the Chi-Square Goodness-of-Fit Statistic." *Geographical Analysis* 31(1), 128–47.
- Rudd, R. A. (2001). "The Geography of Power: A Cautionary Example for Tests of Spatial Patterns of Disease." Unpublished M.S. Thesis, Department of Biostatistics, Rollins School of Public Health, Emory University.

- Tango, T. (1995). "A Class of Tests for Detecting General and Focused Clustering of Rare Diseases." *Statistics in Medicine* 14, 2323–34.
- Turnbull, B. W., E. J. Iwano, W. S. Burnett, H. L. Howe, and L. C. Clark. (1990). "Monitoring for Clusters of Disease: Application to Leukemia Incidence in Upstate New York." *American Journal of Epidemiology* 132(Suppl), S136–43.
- Waller, L. A., and C. A. Gotway. (2004). *Applied Spatial Statistics for Public Health Data*. New York: Wiley.
- Waller, L. A., B. W. Turnbull, L. C. Clark, and P. Nasca. (1994). "Spatial Pattern Analyses to Detect Rare Disease Clusters." In *Case Studies in Biometry*, 3–23, edited by N. Lange, L. Ryan, L. Billard, D. Brillinger, L. Conquest, and J. Greenhouse. New York: Wiley.
- Waller, L. A., B. W. Turnbull, L. C. Clark, and P. Nasca. (1992). "Chronic Disease Surveillance and Testing of Clustering of Disease and Exposure: Application to Leukemia Incidence and TCE-Contaminated Dumpsites in Upstate New York." *Environmetrics* 3, 281–300.