

XXIX REUNIÓN DE ESTUDIOS REGIONALES

Santander, 27-28 de noviembre de 2003

DISEÑO DE UNIDADES TERRITORIALES REPRESENTATIVAS DEL FENÓMENO A ESTUDIAR: UNA PROPUESTA METODOLÓGICA.

Juan Carlos Duque, Raúl Ramos, Jordi Suriñach

jduque@eco.ub.es, rramos@ub.edu, jsurinach@ub.edu

Grup d'Anàlisi Quantitativa Regional - Universitat de Barcelona

Espai de Recerca en Economia. Avda. Diagonal 690. 08034 Barcelona

Resumen:

Uno de los principales inconvenientes al realizar estudios que impliquen la utilización de información agregada geográficamente consiste en la sensibilidad de los resultados a la forma como se han configurado las unidades territoriales, unidades que en muchas ocasiones no se relacionan con los objetivos del estudio. Dicho problema ha sido estudiado en la literatura como el Problema de la Unidad Espacial Modificable (PUEM).

En este trabajo se presenta un modelo de optimización y un heurístico para el diseño de unidades territoriales que permiten superar algunas de las deficiencias de las formulaciones propuestas en la literatura. Para ilustrar el potencial de dicha metodología, se presenta una aplicación para la ciudad de Barcelona.

Palabras clave: Zonificación, Unidad espacial modificable, Optimización, Restricción de continuidad.

Códigos JEL: R22, R12, C61

1. Introducción*

La utilización de información estadística a nivel regional y urbano cada vez está siendo más utilizada para llevar a cabo análisis de carácter económico¹. Sin embargo, una de las principales cuestiones a resolver a la hora de realizar dicho tipo de estudios consiste en el diseño de unidades territoriales apropiadas, es decir, que estén directamente relacionadas con el fenómeno que se estudia. En este sentido, hay que destacar que frecuentemente estas unidades no coinciden con las oficialmente establecidas, motivo que no impide que se lleve a cabo el estudio en cuestión a pesar de las consecuencias adversas que se pueden producir. De hecho, es conocido que la sensibilidad de los resultados a la agregación de datos geográficos puede traer consecuencias no deseables en el análisis (Openshaw, 1984), lo cual implica que dicha agregación no puede hacerse arbitrariamente. En este contexto, son numerosas las metodologías que se han desarrollado para el diseño de unidades territoriales que cumplan con unos requerimientos preestablecidos y que minimicen el peligro de obtener resultados erróneos, o poco representativos, como consecuencia del diseño de unidades territoriales inapropiadas. Sin embargo, dichas metodologías han sido formuladas para solucionar problemas de regionalización muy específicos, lo cual implica que en determinados ámbitos pueden resultar demasiado restrictivas o inapropiadas para cumplir con los requerimientos deseados.

Así pues, el objetivo de este trabajo consiste en aportar una nueva herramienta de regionalización que incorpore algunas posibilidades que no ofrecen las metodologías propuestas en la literatura.

El resto del trabajo se organiza en cuatro apartados. En primer lugar, en el segundo apartado se realiza una breve revisión de la literatura existente sobre diferentes metodologías para el diseño de unidades territoriales; a continuación, se presenta un modelo de optimización lineal para el diseño automatizado de regiones y la posibilidad de aplicar heurísticos para la solución de problemas de gran dimensión; en el cuarto apartado se presenta una aplicación del modelo en un entorno real y, por último, se presentan las principales conclusiones del estudio.

* Los autores agradecen al doctor Ernest Pons por sus aportes a lo largo de la elaboración de este trabajo.

¹ Así lo confirma el estudio bibliométrico realizado por Suriñach *et al* (2003) para el período 1991-2000, tomando como base las 10 revistas internacionales de mejor prestigio en el campo de la ciencia regional y urbana.

2. Revisión de la literatura

Gran parte de las metodologías utilizadas en el campo de la regionalización utilizan técnicas enmarcadas dentro del análisis cluster. En este contexto, el problema de agregación de datos espaciales puede ser visto como un caso particular de cluster en el cual se debe asegurar la continuidad geográfica entre los elementos agrupados. Este caso especial de análisis cluster es llamado generalmente análisis cluster con restricción de continuidad espacial o simplemente regionalización². Dentro de este tipo de análisis, se pueden destacar tres estrategias metodológicas: La agrupación en dos etapas, la inclusión de las variables geográficas como variables de clasificación y la utilización de instrumentos adicionales para controlar la restricción de continuidad geográfica.

La agrupación en dos etapas consiste en aplicar, en una primera etapa, un método clásico de clustering utilizando únicamente los datos no-geográficos³, y posteriormente, en una segunda etapa se utiliza la información geográfica⁴ para controlar las condiciones de continuidad. Así pues, si a un cluster pertenecen áreas que son geográficamente discontinuas, entonces cada una de estas áreas serán asumidas como regiones diferentes (Ohsumi, 1984). Entre las ventajas asociadas a dicha metodología Openshaw y Wymer (1994) resaltan el hecho de que en la primera etapa se garantiza la homogeneidad de las zonas creadas utilizando únicamente las variables de interés, esto es, sin incorporar la distorsión que puedan generar las variables geográficas. Por otro lado, este método puede resultar útil como un medio para obtener evidencias de una dependencia espacial entre los elementos. Sin embargo, para los objetivos de la regionalización, puede convertirse en un inconveniente el hecho de que el número de grupos no dependa del investigador sino del grado de dependencia espacial⁵.

La segunda estrategia consiste en incluir las variables geográficas y no-geográficas como variables de clasificación (Perruchet, 1983). En esta estrategia, a diferencia de la anterior, las variables geográficas forman parte activa en el cálculo de las similitudes o disimilitudes entre los objetos que serán agrupados utilizando algoritmos convencionales de clasificación. Los principales inconvenientes asociados a esta estrategia de regionalización son la dificultad que implica el tratamiento simultáneo de variables expresadas en diferentes unidades de

² Un resumen detallado de estas metodologías de agrupación se encuentra en Gordon (1999) y para los casos especiales de cluster restringido en Gordon (1996) y Murtagh (1985).

³ Aquellos atributos propios de los elementos o de sus relaciones, por ejemplo, población, ingreso medio, número de hospitales, movilidad laboral entre municipios, etc.

⁴ Aquella información utilizada para controlar la restricción de continuidad geográfica de los elementos agrupados, por ejemplo, las coordenadas geográficas, las distancias entre elementos, entre otras.

⁵ A mayor (menor) dependencia espacial habrá una tendencia hacia la creación de pocas (muchas) regiones.

medida y la asignación objetiva de los pesos para las variables que intervienen en el proceso de clasificación, cuidando que el peso dado a las variables geográficas asegure la formación de regiones continuas (Wise, Haining y Ma, 1997). Los principales inconvenientes derivados de la inclusión de variables geográficas como variables de clasificación son: la necesidad de que el componente geográfico tenga el peso suficiente para asegurar la continuidad geográfica y, dado que los elementos a agrupar son de tipo área, la solución puede variar en función del método utilizado para localizar el punto que representa cada una de ellas, sobretodo en aquellos casos donde las áreas son considerablemente grandes (Horn, 1995).

Por último, y quizá la más utilizada de las estrategias para la resolución de problemas de agrupación territorial, consiste en controlar la restricción de continuidades geográficas utilizando instrumentos adicionales como la matriz de contactos. Dicha matriz puede ser graficada como una red, donde los arcos muestran las relaciones de vecindad entre los nodos que representan los elementos a agrupar⁶, y, a partir de esta representación, el siguiente paso consiste en realizar las modificaciones apropiadas a los algoritmos clásicos de clasificación, jerárquicos o de particionamiento, con el objetivo de asegurar el cumplimiento de la restricción de continuidad.

Los algoritmos jerárquicos suelen ser utilizados en aquellos casos donde el investigador está interesado en conocer la clasificación óptima en función del número de grupos a conformar, es decir para diferentes niveles de escalamiento. Esta solución es generalmente esquematizada con dendogramas⁷, los cuales, para el caso de regionalización, pueden presentar reveses o rupturas de monotonidad, lo cual implica que la distancia que separa dos objetos puede ser mayor que la distancia que separa la unión de estos dos objetos con un tercero (Gordon 1996). Uno de los principales inconvenientes asociados a dicha metodología, además de los requerimientos de una alta capacidad de memoria para almacenar las soluciones (Wise, Haining y Ma, 1997), es la alta posibilidad de obtener óptimos locales debido a que una vez dos elementos están agrupados, estos no volverán a ser evaluados por separado para resolver niveles de escalamientos más bajos (Semple y Green, 1984). Como principal ventaja se puede resaltar el no tener que especificar particiones iniciales como punto de inicio del algoritmo (Macmillan y Pierce, 1994).

En cambio, cuando el investigador está interesado en realizar un proceso de regionalización dado un escalamiento predefinido, resulta ser muy apropiada la utilización de

⁶ Para una descripción más detallada sobre los diferentes métodos para la construcción de éste tipo de gráficos, ver Gordon (1996, 1999).

⁷ Representación gráfica de las soluciones obtenidas en un proceso de cluster jerárquico. (ver Gordon, 1996).

algoritmos de particionamiento. Para solucionar problemas de particionamiento, destacan dos metodologías: la programación matemática y los algoritmos iterativos.

Con respecto a los métodos de programación matemática, Macmillan y Pierce (1994) definen el problema de regionalización como un problema de optimización combinatoria con el cual se determinará una agregación territorial óptima dado un número predeterminado de grupos a conformar. La solución propuesta por estos autores para asegurar la continuidad geográfica de los elementos consiste en exponenciar la matriz de contactos K_{ij} (matriz binaria que toma el valor de 1 si i y j son continuos; y 0 en caso contrario), teniendo en cuenta que para la formación de un grupo con n elementos continuos es necesario que la $(n-1)$ ésima potencia de la matriz de contactos no contenga elementos nulos. Esta solución implica que el espacio factible definido por las restricciones sea no convexo lo cual no asegura la obtención de un óptimo global.

Dado que la solución de este tipo de problemas por métodos de optimización convencionales resulta extremadamente complejo⁸, se han desarrollado otras metodologías en el campo de la regionalización que han resultado muy efectivas en aquellos casos donde el número de elementos a agrupar es muy grande. Dentro de este tipo de soluciones destacan distintos algoritmos como el “*automatic zoning problem*” (AZP) (Openshaw, 1977) que tienen como objetivo dividir un territorio en distintas áreas buscando además que dichas unidades territoriales minimicen, en lo posible, los efectos asociados al “problema de la unidad espacial modificable” (PUEM)^{9,10}.

Los modelos expuestos hasta ahora, se caracterizan por ser métodos supervisados, lo cual implica que el investigador conoce *a priori* la estructura de los datos que esta utilizando y que describen un fenómeno que se desea estudiar. Pero también existen métodos no controlados, útiles en casos donde el investigador desea analizar una gran cantidad de datos sobre los cuales no existe información de los factores que afectan el sistema. En estos casos surge la posibilidad de realizar un análisis no paramétrico de los datos que permita encontrar patrones y relaciones en un gran volumen de información. Una de las principales aplicaciones

⁸ Openshaw (1984) calculó que para agregar 1000 objetos en 20 grupos existen 101,260 soluciones diferentes. Para más información sobre problemas de tipo combinatorio, ver Aarts y Lenstra (1997).

⁹ En inglés “*modifiable areal unit problem*” (MAUP).

¹⁰ Openshaw planteó el problema de la unidad espacial modificable (PUEM), definiéndolo como una fuente potencial de error que puede afectar los resultados de aquellos estudios basados en información agregada geográfica, ya que estos pueden variar en función de la configuración de dicha agregación. El PUEM agrupa dos problemas específicos con respecto al análisis de datos espaciales: El problema de escalamiento, relacionado con el número de zonas que se desean establecer para el estudio, y el problema de agregación, relacionado con la configuración de pequeñas zonas dentro de zonas más grandes. Para más información ver Openshaw (1977), Openshaw y Taylor (1981), y en un contexto econométrico ver Fotheringham y Wong (1991) y Amrhein y Flowerdew (1992).

de estos métodos dentro del campo de la regionalización son los “*Self Organization Maps*” (SOM) propuesto por Kohonen (1984). Esta metodología, desarrollada en el campo de la inteligencia artificial ha generado opiniones divergentes entre los investigadores que trabajan en este campo, principalmente por la falta de fundamentación teórica que dificulta la interpretación de los resultados (Openshaw, 1992).

3. Modelo de optimización lineal para la configuración de regiones

3.1. Descripción del modelo

Antes de su formalización matemática, se expondrán a continuación las características de dicho modelo y las hipótesis sobre las cuales fue construido.

El punto de partida de cualquier proceso de regionalización es determinar el territorio que se desea regionalizar, el cual debe estar formado por un número finito (n) de áreas geográficas de menor tamaño que forman un continuo geográfico $A = \{a_1, a_2, a_3, \dots, a_n\}$. Este conjunto geográfico puede ser resumido en un grafo compuesto por n nodos representando cada una de las áreas y cuyos arcos representan las relaciones de continuidad entre éstas.

Para tal representación existen diferentes métodos de los cuales se ha seleccionado el más general, el “*Delaunay Triangulation*” (DT) (Aurenhammer, 1991), en el cual cada arco une aquellas áreas que comparten una frontera común. Una de las principales ventajas de este método consiste en que la localización del punto que representa cada una de las áreas no afecta el resultado del grafo¹¹.

El paso a seguir es la incorporación de una serie de descriptivos sobre las relaciones entre cada una de las áreas (o nodos del grafo). Estas relaciones entre áreas son uno de los elementos más importantes dentro del proceso de regionalización propuesto en este trabajo, ya que su incorporación permite considerar interacciones entre las áreas para que su agrupación sea más coherente. Estas relaciones se incorporan en el modelo, a través de una matriz cuadrada D_{ij} ($i = 1, 2, \dots, n$ y $j = 1, 2, \dots, n$), y simétrica ($r_{ij} = r_{ji} \forall i, j$) que contiene valores de disimilaridades entre cada par de áreas i, j . La función seleccionada para el cálculo de las disimilaridades entre pares de áreas debe cumplir las siguientes propiedades:

¹¹ Otros métodos, como el “*Gabriel Graph*” (Matula y Sokal, 1980), el “*Relative Neighbourhood Graph*” (Toussaint, 1980) o el “*Minimum Spanning Tree*” (Graham y Hell, 1985), están contenidos dentro de DT y a diferencia de éste los resultados obtenidos pueden variar en función del punto donde se localice el centroide de cada área.

$$d_{ij} = d_{ji} \quad \forall i, \forall j = 1, \dots, n \quad (1)$$

$$d_{ij} \geq 0, \quad (d_{ij} = 0 \quad \text{sí } i = j) \quad \forall i, \forall j = 1, \dots, n \quad (2)$$

Lo cual implica que las funciones no tienen que ser métricas, es decir, no tienen que cumplir necesariamente con la propiedad de desigualdad triangular¹²:

$$d_{ij} \leq d_{ik} + d_{kj} \quad \forall i, \forall j, \forall k = 1, \dots, n \quad (3)$$

La posibilidad de utilizar funciones de distancia que no sean necesariamente métricas puede ser entendida como una relajación de la hipótesis utilizada en aquellos modelos de regionalización basados en la selección de áreas centroides a las cuales se asignan las áreas restantes en función de la cercanía en tiempo o espacio, funciones que al ser métricas aseguran que el proceso final de asignación a los centroides cumpla con la restricción de continuidad geográfica.

Una vez se tienen los elementos descritos anteriormente (información sobre la configuración del territorio con el cual se trabaja y sobre las relaciones entre las áreas que lo forman), se puede iniciar el proceso de regionalización, que puede ser definido como la agrupación de las n áreas $\{a_1, a_2, \dots, a_n\}$ en m conjuntos no vacíos o regiones $\{1, 2, \dots, m\}$ tal que las áreas que pertenecen a cada región conformen un continuo geográfico.

Para definir dichas regiones será necesario seleccionar $n-m$ arcos dentro del conjunto de arcos que conforman el grafo de continuidades. Estos $n-m$ arcos serán suficientes, bajo el cumplimiento de ciertas condiciones, para formar cada una de las m regiones tal que cada región esté internamente conectada pero completamente desconectada de las demás regiones. Este sistema de configuración de regiones implica que el mínimo número de áreas en cada región será dos (un arco uniendo dos áreas), esto significa que $m \leq [n/2]$. Esta condición se hace menos restrictiva a medida que el número de áreas dentro del territorio sea mayor.

La ubicación de los arcos dentro de cada región no tiene influencia en el resultado final, por ejemplo la región formada por las áreas conectadas con los arcos 1-2, 2-3 y 2-4 puede ser configurada con los arcos 1-3, 2-4 y 3-4, esto se debe a que la función de los arcos, al no tener valores asociados, es únicamente la de asegurar la continuidad geográfica.

¹² Para más información, ver Gower y Legendre (1986).

Esta estrategia de arcos resulta especialmente útil para identificar configuraciones regionales con variedad de formas, por ejemplo franjas litorales, pre-litorales e interior, o regiones en forma de coronas, entre otras.

El objetivo de agrupar n áreas en m regiones tal que las áreas que pertenecen a cada región formen un continuo geográfico, y que cada una de las regiones agrupe áreas similares, hace necesario definir un criterio de partición que permita detectar cual de las posibles configuraciones de n áreas en m regiones es la mas adecuada. Para alcanzar dicho objetivo se puede calcular el grado de heterogeneidad de las áreas asignadas a una región o, como segunda alternativa, se puede calcular el grado de aislamiento de las áreas dentro de una región, con respecto a las áreas localizadas fuera de dicha región. La medida de heterogeneidad seleccionada en este trabajo como criterio de agrupación se basa en el sumatorio de los elementos de la diagonal superior de la matriz de distancias (métricas o no métricas) entre los elementos que están dentro de una región. Conservando la notación utilizada por Gordon (1999), esta medida de heterogeneidad para la clase (o región) r (C_r) se expresa como:

$$H(C_r) \equiv \sum_{\{i,j \in C_r | i < j\}} d_{ij} \quad (4)$$

Así pues, el objetivo para la construcción de r clases homogéneas puede ser planteado como la minimización del sumatorio de las medidas de heterogeneidad de cada clase r :

$$P(H, \Sigma) \equiv \sum_{r=1}^c H(C_r) \quad (5)$$

o también se puede seguir una estrategia MIN-MAX para que el valor de la clase más heterogénea sea lo más bajo posible, lo cual asegura que las demás clases tendrán una heterogeneidad menor o igual:

$$P(H, Max) \equiv \max_{\{r=1, \dots, c\}} H(C_r) \quad (6)$$

Un inconveniente asociado a la segunda estrategia es que una vez se logra minimizar el valor asociado a la clase más heterogénea, la configuración del resto de clases no será revisada, descartando la posibilidad de hacer modificaciones que mejoren los valores de heterogeneidad de dichas clases.

Así pues, la estrategia adoptada en el modelo a proponer consistirá en la minimización del sumatorio de las medidas de heterogeneidad de cada clase ($P(H, \Sigma)$). Los objetivos de minimización de heterogeneidades, $H(C_r)$, dentro de las clases y de maximización de aislamientos entre clases, $I(C_r)$, no son objetivos inconexos, para el objetivo seleccionado existe un equivalente en términos de aislamientos:

$$P(H, \Sigma) \equiv P(I, \Sigma) \equiv \sum_{r=1}^c I(C_r) \text{ con } I(C_r) \equiv \sum_{i \in C_r} \sum_{j \notin C_r} d_{ij} \quad (7)$$

3.2. Modelo matemático.

A continuación se presentará la formalización matemática que incorpora los elementos introducidos anteriormente en tres bloques: parámetros, variables de decisión y la función objetivo y restricciones del modelo.

Parámetros :

- i, I índice y conjunto de áreas, $i = \{1, \dots, n\}$;
- k, K índice y conjunto de regiones, $k = \{1, \dots, m\}$;
- w_{ij} $\begin{cases} 1, \text{ si } i \text{ y } j \text{ son continuas (comparten una frontera), con } i < j, \\ 0, \text{ en caso contrario;} \end{cases}$
- M $\text{Max} \left(\sum_{j=1}^n w_{1j}, \dots, \sum_{j=1}^n w_{nj} \right)$;
- N_i $\{j | w_{ij} = 1\}$;
- $D_{i,j}$ Medida de distancia entre áreas i y j , con $i < j$; la propiedad métrica no es importante

Variables de Decisión :

- X_{ijk} $\begin{cases} 1, \text{ si las áreas } i \text{ y } j | j \in N_i \text{ pertenecen a la misma región } k, \text{ con } i < j, \\ 0, \text{ en caso contrario;} \end{cases}$
- Y_{ik} $\begin{cases} 1, \text{ si la área } i \text{ pertenece a la región } k, \\ 0, \text{ en caso contrario;} \end{cases}$
- T_{ij} $\begin{cases} 1, \text{ la medida de distancia entre } i \text{ y } j \text{ es considerada si las dos áreas pertenecen a} \\ \text{la misma región } k, \text{ } i < j, \\ 0, \text{ en caso contrario;} \end{cases}$

Función Objetivo :
$$\text{Min} \sum_{i=1}^n \sum_{j=1}^n D_{ij} \cdot T_{ij}$$

Sujeto a :

$$T_{ij} \geq Y_{ik} + Y_{jk} - 1, \quad \forall i, \forall j = 1, \dots, n ; \forall k = 1, \dots, m \quad (8)$$

$$\sum_{i=1}^n Y_{ik} \geq 2, \quad \forall k = 1, \dots, m \quad (9)$$

$$\sum_{k=1}^m Y_{ik} = 1, \quad \forall i = 1, \dots, n \quad (10)$$

$$\sum_{j \in N_i} X_{ijk} \leq Y_{ik} \cdot M, \quad \forall i = 1, \dots, n ; \forall k = 1, \dots, m \quad (11)$$

$$\sum_{j \in N_i} X_{jik} \leq Y_{ik} \cdot M, \quad \forall i = 1, \dots, n ; \forall k = 1, \dots, m \quad (12)$$

$$\sum_{i=1}^n \sum_{j \in N_i} X_{ijk} = \sum_{i=1}^n Y_{ik} - 1, \quad \forall k = 1, \dots, m \quad (13)$$

$$\sum_{i,j \in C} X_{ijk} \leq |C| - 1, \quad \forall \text{subconjunto no vacío de } C \subseteq \{3, \dots, (n-2m+1)\}; \forall k = 1, \dots, m \quad (14)$$

$$X_{ijk} \in \{1,0\}; Y_{ik} \in \{1,0\}; T_{ij} \geq 0, \quad \forall i, \forall j = 1, \dots, n ; \forall k = 1, \dots, m \quad (15)$$

Tal como se comentó anteriormente, la función objetivo adoptada en el modelo propuesto busca la minimización del sumatorio de los elementos de la diagonal superior de la matriz de distancias entre los elementos que pertenecen a un mismo grupo (minimización de la heterogeneidad total). La restricción (8) permite seleccionar los elementos de la matriz de distancia que corresponden a las áreas que pertenecen a un mismo grupo, en otras palabras, se eliminan de la matriz de distancia todos aquellos elementos que relacionan áreas que pertenecen a diferentes grupos. La restricción (9) impone que el mínimo número de áreas que conforman una región es dos. Como se explico anteriormente, esta restricción se hace menos fuerte a medida que aumenta el número de áreas. La restricción (10) impone que cada área debe ser asignada a una sola región. Las restricciones (11) y (12) imponen que sólo en el momento en el que una área i está asignada a una región k será posible asignar arcos de entrada o salida hacia las vecindades de dicha área ($j \in N_i$). Para no restringir más de lo necesario las posibilidades de configuraciones territoriales, el número de arcos de entrada y/o salida desde una área puede ser mayor que uno. La restricción (13) impone que el número de arcos necesarios para asegurar la continuidad de las áreas asignadas a una región debe ser

igual al número de áreas asignadas a dicha región menos uno. Sin embargo esta restricción por sí sola no asegura que la solución final genere regiones continuas, pues se pueden presentar casos como el que se ilustra en la figura 1, donde la región A, compuesta por las áreas 1, 2, 3, 6 y 7, satisface la restricción (13), pues para 5 áreas asignadas a dicha región fueron utilizados 4 arcos de unión, pero la existencia del ciclo de arcos 1-2, 1-3, 2-3 genera el rompimiento de la continuidad geográfica de la región. Por lo tanto será necesario controlar, además del número de arcos, la existencia de ciclos, punto que se tratará más detenidamente a continuación y del cual se deriva la restricción (14).

El problema de los ciclos ha sido tratado bajo el nombre de “*subtours*” en modelos de transporte como el “*Vehicle Routing Problem*” (VRP)¹³ donde la configuración de una ruta, o “*tour*”, para un vehículo no puede contener “*subtours*”, y para controlar esta condición, el VRP incorpora la siguiente restricción:

$$\sum_{j,i \in S} X_{ijk} \leq |S| - 1, \forall \text{ subconjunto no vacío de } S \subseteq \{2, \dots, n\}; k=1, \dots, m. \quad (16)$$

FIGURA 1.

El inconveniente de este planteamiento es que el número de restricciones crece exponencialmente con n y m . Por este motivo, aunque la propuesta es teóricamente correcta, a nivel práctico ha sido necesario implementar otras restricciones que permitan resolver el problema de una manera más eficiente. Dichas propuestas, además de no asegurar la eliminación de “*subtours*” en problemas de ciertas dimensiones, son apropiadas para el problema específico de VRP pero no son adaptables al problema de regionalización por incompatibilidades como: la necesidad de determinar a priori un nodo “*depot*” que será el punto de inicio y final de todos los “*tours*”, y la necesidad de que cada nodo tenga un arco de entrada y uno de salida para asegurar un recorrido secuencial.

Sin embargo, la restricción del VRP propuesta a nivel teórico puede ser adaptada en el ámbito geográfico si se tiene en cuenta que el previo conocimiento de la configuración territorial permite reducir considerablemente el número de elementos que pertenecen al conjunto $|S|$. Por ejemplo, en la configuración territorial presentada a la izquierda del cuadro 1 se pueden identificar claramente las distintas combinaciones de arcos que generan ciclos, con $i < j$, la

¹³ Propuesto originalmente por Dantzing y Ramser (1959). Un resumen sobre modelos derivados del planteamiento original se encuentra en Laporte y Osman (1995).

combinación de arcos 1-2, 1-3, 2-3 (o 2-3, 2-4, 3-4) daría lugar a un ciclo en el que se verían involucradas tres áreas, la 1, 2 y 3 (o 2, 3, 4), mientras que la combinación de arcos 1-2, 1-3, 3-4, 2-4 generaría un ciclo entre las cuatro áreas. Por otro lado, en una configuración territorial como la que aparece a la derecha del cuadro 1, no hay ninguna combinación de arcos i,j (con $i < j$), que pueda generar un ciclo. Por tanto, a nivel territorial, no todos los subconjuntos de S pueden presentar ciclos, puesto que el conjunto de potenciales arcos i,j está limitado a aquellos pares i,j en los que la matriz de contactos $w_{ij} = 1$. Dicho conjunto de potenciales arcos i,j son los que están recogidos en N_j .

CUADRO 1.

¿Existe algún patrón especial que permita detectar los potenciales ciclos en una configuración territorial específica? Sí, sólo basta con identificar aquellas combinaciones de arcos donde el número de arcos utilizados coincida con el número de áreas conectadas con dichos arcos, para el ejemplo anterior, los tres arcos 1-2, 1-3, 2-3 (o 2-3, 2-4, 3-4) conectan tres áreas, 1,2,3 (o 2,3,4), por lo tanto, 3 arcos y 3 áreas implica la existencia de un ciclo. Lo mismo ocurre con la combinación de arcos 1-2, 1-3, 3-4, 2-4 con la cual se conectan cuatro áreas (1,2,3,4), y por lo tanto, 4 arcos y 4 áreas darían lugar a un ciclo de 4 elementos.

La siguiente pregunta a formular será: Para una configuración territorial compuesta por n áreas donde se busca diseñar m regiones, *¿Cuál puede ser el máximo número de áreas involucrados en un ciclo?* Puesto que el planteamiento del modelo, en su restricción (9), exige que el mínimo número de áreas dentro de una región puede ser 2, el caso extremo en la agregación de n áreas en m regiones es aquel en el que $(m-1)$ regiones estén formadas cada una por dos áreas, las cuales no tendrían la posibilidad de que existan ciclos, pues para cada una de estas regiones solo se asignaría un arco (restricción (13)). Por tanto, al crear $m-1$ regiones con 2 áreas cada una, se tendría una región formada por $n-2(m-1)$ áreas, a las cuales serán asignados $(n-2(m-1))-1$ arcos, que sería el máximo numero de arcos con los cuales se puede crear un ciclo. Simplificando la expresión, se tiene:

$$n-2m+1 \quad (17)$$

Por otro lado, el mínimo número de áreas a partir de las cuales se deben evaluar las posibilidades de ciclos es tres, pues para un número inferior de áreas no es posible que exista este problema.

Así, la restricción (14) surge como una modificación de la definición del conjunto $|S|$ propuesto en el VRP, lo cual trae como ventaja adicional una importante reducción del número de restricciones a satisfacer, evitando que el número de restricciones crezca exponencialmente con n y m . Esto hace posible la utilización de software comercial para la solución de problemas de regionalización con un mayor número de áreas y regiones.

Por último, en la restricción (15) del modelo, sólo es necesario fijar como variables binarias a X_{ijk} y a Y_{ik} . Aunque la variable T_{ij} se ha definido como positiva, no como binaria, siempre tomará valores de 0 o 1 por la combinación de la restricción (8) y el objetivo de minimización del modelo¹⁴.

3.3. Utilización de heurísticos para la configuración de regiones: “*Regionalization Algorithm with Selective Search*” (RASS)

Una de las características que más interesa a la hora de aplicar los modelos matemáticos a problemas reales, los cuales involucran generalmente una gran cantidad de áreas, es el tiempo que tarda dicho modelo en encontrar la solución óptima. Con el objetivo de mejorar este tiempo, en la literatura se ha propuesto la utilización de diferentes heurísticos.

Los heurísticos de interés teniendo en cuenta las características del modelo formulado son aquellos que buscan la partición del territorio en un número preestablecido de regiones y dónde no se asigne un papel diferenciado a las áreas a ser agrupadas. Entre los heurísticos más utilizados en la literatura se encuentra el “Automatic Zoning Procedure (AZP)”. Este heurístico, propuesto por Openshaw (1977), está basado en un procedimiento iterativo de prueba y error. Consiste en la optimización de una función objetivo $F(Z)$, donde Z es la asignación de cada una de las N zonas a una de las M regiones, cada zona debe ser asignada a un sola región y cada región debe tener como mínimo una zona. Entre las principales ventajas atribuidas a este heurístico destaca la posibilidad de utilizar cualquier función objetivo sensible a la agregación de zonas. Esta característica resulta de gran utilidad para tener una aproximación a los efectos límite (la deferencia entre maximizar y minimizar su valor) de la agregación de

¹⁴ La posibilidad de definir una variable que toma valores 0 o 1 como positiva, y no como binaria, representa una ventaja para la utilización del algoritmo de bifurcación y acotamiento, pues el número de sub-problemas se reduce significativamente. Para más información sobre este algoritmo, ver Hiriart, Oettli y Store (1983).

cualquier estadístico o modelo. También fue muy útil para demostrar la existencia del PUEM. Por otro lado, los principales inconvenientes atribuidos a este heurístico son su procedimiento de búsqueda local (restringida a la región seleccionada) y la fuerte dependencia de los resultados al punto de partida seleccionado (paso 1). Por último, la estrategia de descartar la posibilidad de mover una zona que genere un descenso en la función objetivo puede hacer que el heurístico quede atrapado con relativa facilidad en óptimos locales¹⁵.

Así pues, con el objetivo de implementar una estrategia de regionalización que evite en lo posible quedar atrapado en óptimos locales, el algoritmo RASS (*Regionalization Algorithm with Selective Search*) tiene como principal característica la capacidad de dirigir el proceso de búsqueda de forma selectiva, basándose en la información disponible sobre las relaciones entre áreas. El RASS incorpora dentro de su algoritmo el modelo de optimización presentado anteriormente en este mismo apartado y para conseguir mejoras locales de la función objetivo, consta de los siguientes pasos:

Paso 1: Partir de una solución factible de m regiones que agrupen n áreas.

Paso 2: Seleccionar de las m regiones el continuo geográfico mas heterogéneo compuesto por r regiones (con $r \geq 2$).

$$H(C_m) \equiv \sum_{\{i,j \in C_k | i < j\}} d_{ij} \rightarrow \text{Max} \left(\sum_{m \in M_i} H(C_m) \right) \quad (18)$$

Dónde M_i es el conjunto formado por las diferentes alternativas de selección de r regiones contiguas del las m disponibles.

Paso 3: Aplicar a las áreas contenidas dentro de las r regiones seleccionadas el modelo de optimización lineal propuesto para crear r^* regiones.

Paso 4: Seleccionar la región que entra (e): De las $(m-r)$ regiones que quedaron fuera y que además estén en las vecindades del área conformada por r^* (f), seleccionar aquella región más similar a alguna de las r regiones que están dentro (d).

¹⁵ Este último se ha intentado solucionar en posteriores propuestas como el "Simulated Annealing Variant of AZP (AZP-SA)" propuesto por Openshaw, el "ANNEAL redistricting problem" propuesto por Macmillan y Pierce (1994), el "Tabu Search Algorithm" (AZP-TABU) adaptado por Openshaw para problemas de regionalización, y el heurístico basado en árboles de recorrido para las agrupaciones territoriales, como MIDAS, propuesto por Maravalle y Simeone (1995).

$$I(C_{d,f}) \equiv \sum_{i \in C_f} \sum_{j \in C_d | j > i} d_{ij} \rightarrow \text{Min}(I(C_{d,f})) \quad (19)$$

Dónde d es el conjunto formado por las r^* regiones que están dentro, y donde f es un subconjunto de las regiones que están fuera, formado por aquellas regiones que se encuentran en las vecindades de d . Cada una de las $(m-r)$ regiones que quedaron fuera en el paso 2 sólo podrán ser seleccionadas como regiones entrantes una vez por ciclo (pasos 2 a 8).

Paso 5: Selección de la región que sale (s): Sale la región más diferente a la región que fue seleccionada para entrar (e) en el paso 4.

$$I(C_{d,e}) \equiv \sum_{i \in C_e} \sum_{j \in C_d | j > i} d_{ij} \rightarrow \text{Max}(I(C_{d,e})) \quad (20)$$

Paso 6: Incorporar al conjunto de r regiones la región seleccionada para entrar y extraer la región seleccionada para salir $d=(d+e-s)$. A esta nueva configuración de r regiones aplicar el modelo de optimización propuesto para crear r^* regiones.

Paso 7: Repetir 4 a 6 hasta que las $(m-r)$ regiones que quedaron afuera en el paso 2 hayan pasado por d .

Paso 8: Calcular la función objetivo.

Paso 9: Repetir los pasos 2 al 8 hasta que no se encuentre una mejora significativa de la función objetivo entre el último ciclo y el anterior.

Algunas características a destacar en el algoritmo *RASS* son que la aplicación de optimización directa a un grupo de regiones, en los pasos 3 y 6, permite lograr mejoras en la función objetivo que pueden generar cambios importantes en las configuraciones regionales, por la reasignación de un gran número de áreas entre las regiones. Además, los criterios utilizados en el paso 2 para la selección de las r regiones, junto con los criterios de entrada y salida de regiones de los pasos 4 y 5, buscan mantener dentro del modelo de optimización, paso 3, aquellas regiones que tengan el mayor potencial de mejorar la función objetivo tras su reconfiguración. De esta manera, se busca que la región que entra sea aquella que presenta la mayor posibilidad de contener áreas que óptimamente pertenezcan a otra región. Esta potencial reasignación de áreas se identifica asumiendo que dos regiones con áreas intercambiadas (en uno o ambos sentidos), aumenta la similitud entre dichas regiones. Por último, una vez identificada el área

entrante, el criterio de salida (para mantener un número apropiado de áreas dentro del modelo de optimización) establece que la región que sale es la que, en los términos de disimilaridad utilizados, está mas alejada del área entrante, y que por tanto dicha región tiene menor posibilidad de intercambiar áreas con la región entrante. Otro aspecto a destacar es que *RASS* contiene, además de las condiciones de entrada y salida comentadas en el punto anterior y de la utilización de un método de optimización directo, dos condiciones cuyo objetivo es el de escapar de óptimos locales. En primer lugar, en cada ciclo todas las regiones que están afuera deben entrar una vez: para asegurar que todas las áreas sean evaluadas dentro del modelo de optimización, y, en segundo lugar, el retorno al paso 1 del algoritmo cada vez que termina un ciclo: para que la solución final no dependa de la selección inicial de las *r* regiones que entrarán dentro del modelo de optimización y para que la condición 1 no mantenga separadas dos regiones mas tiempo de lo necesario.

Por último, el hecho de aplicar el modelo de optimización a una parte del territorio no implica que una mejora local pueda empeorar la solución global, así, al final de cada ciclo, el valor de la función objetivo siempre será menor o igual al valor de la función objetivo al principio del ciclo.

4. Aplicación del modelo¹⁶

Este apartado persigue un doble objetivo: en primer lugar, conocer la potencialidad del modelo propuesto y del algoritmo *RASS* en procesos de regionalización y, en segundo lugar, analizar los posibles efectos del nivel de escalamiento sobre las propiedades de las estimaciones mínimo cuadrático ponderadas (MCP) ¹⁷, ampliamente utilizadas en el análisis econométrico aplicado a nivel territorial.

4.1 Potencialidad del modelo y del algoritmo RASS en procesos de regionalización

Para analizar la potencialidad de la metodología propuesta, en este apartado se presenta una aplicación del modelo a partir de la utilización de datos simulados de tal forma que se conociera *a priori* la configuración regional óptima a la que debería llegar el modelo. El territorio

¹⁶ El software utilizado para correr el modelo fue el Extended LINGO/PC 6.0 en un ordenador Pentium 4 a 2.40C GHz y 256 Mb de RAM. El heurístico fue implementado en Gauss 3.2 y los procesos de optimización en Extended LINGO/PC 6.0.

¹⁷ La utilización de este tipo de estimaciones surge como un instrumento para controlar la heteroscedasticidad generada al trabajar con datos agregados. Para mas información sobre este tema ver Artís *et al* (2001).

seleccionado es la ciudad de Barcelona dividida en 38 áreas (“*Zones Estadístiques Grans*”). De esta configuración territorial interesan, para la aplicación del modelo, las relaciones de continuidad entre las 38 áreas a partir de las cuales se crea la matriz de contactos de primer orden. La matriz de relaciones fue creada de tal forma que la solución óptima agrupase las 38 áreas en 10 regiones con diferentes formas y tamaños (entre 2 y 6 áreas por región)¹⁸. El procedimiento seguido para la elaboración de la matriz de relaciones se resume a continuación.

- a) Agrupación de las n áreas en m regiones continuas geográficamente, asignando cada área $i = \{1, \dots, n\}$ a una región $k = \{1, \dots, m\}$. A partir de esta agrupación se crea el conjunto $R_k \{i | i \in k\}$.
- b) Asignar un valor a cada una de las áreas $i = \{1, \dots, n\}$ en función de la región a la cual hayan sido asignadas. Este valor está compuesto de una constante a la cual se suma un término de perturbación generado a partir de una distribución uniforme entre 0 y 1. El valor de la constante es diferente entre regiones, pues entre ellas existe un diferencial D lo suficientemente grande para obtener valores medios significativamente diferentes entre regiones y muy representativos de las áreas que contiene cada región. La expresión a partir de la cual se genera el vector de valores para cada área es:

$$A_{i \in R_k} = C + (D * k) + \varepsilon \quad \forall i = 1..n; \forall k = 1..m; \varepsilon \sim U[0,1] \quad (21)$$

- c) Relaciones entre áreas: Una vez se han determinado los valores de cada una de las áreas, el siguiente paso consiste en calcular las relaciones entre áreas aplicando una función de distancia, que no debe ser necesariamente métrica. La función de distancia a utilizar en el ejemplo será la distancia euclídea ponderada, calculada entre los elementos del vector A_i después de centrarlo.

$$d_{ij} = \sqrt{\left(\frac{A_i^c}{S} - \frac{A_j^c}{S} \right)^2}, \quad \forall i, j = 1, \dots, n \mid i < j \quad (22)$$

Donde S es la desviación estándar del vector A_i , y A_i^c corresponde la vector A_i centrado:

¹⁸ Por falta de espacio no se presenta aquí la matriz de relaciones entre las regiones. Aquellas personas interesadas pueden solicitarla directamente a los autores.

$$A_i^c = A_i - \left(\sum_{i=1}^n A_i / n \right), \forall i = 1, \dots, n \quad (23)$$

En la parte izquierda del cuadro 2 se presenta la configuración regional óptima a la cual se debería acercar la solución heurística, en el centro de dicho cuadro se presenta la partición inicial, paso 1 del RASS, formada por regiones que se alejan considerablemente de la configuración regional óptima y en el extremo derecho se presenta el resultado obtenido, después de 6 ciclos^{19,20,21}, el cual coincide con la agrupación óptima que se había determinado *a priori*.

CUADRO 2.

En la figura 2, se muestra el comportamiento de la función objetivo. Tal y como se puede observar es en los primeros ciclos donde se logran las mejoras más importantes. También se confirma como en cada ciclo del heurístico se consigue mejorar, o en el peor de los casos mantener estable, el valor de la función objetivo con respecto a su valor en el ciclo anterior.

FIGURA 2.

Cabe destacar que el número de regiones dentro del modelo de optimización se fijó en 4 (es decir $r = 4$), con esto, el número promedio de áreas con las cuales se ejecutó cada modelo de optimización fue de 15 áreas, las suficientes para que los tiempos de ejecución fueran apropiados, un promedio de 2.43 minutos por modelo. En la figura 3 se presentan los tiempos de ejecución de cada uno de los modelos de optimización ejecutados dentro del RASS. Como se puede observar, el tiempo de ejecución de los modelos de optimización tiende a ser mayor al comienzo de cada ciclo y es especialmente elevado la primera vez que se ejecuta, pero también la que más aporta a la minimización de la función objetivo. Esto se debe a que el primer modelo de cada ciclo se ejecuta tomando las $r=4$ regiones mas heterogéneas, lo cual implica que la reasignación de las áreas contenidas en estas r regiones puede ser elevada, para este ejemplo en particular, el primer modelo reasigna el 37% de dichas áreas (o un 18.4%

¹⁹ Entendiendo por “ciclo” la ejecución de los pasos 2 a 8 dentro del RASS.

²⁰ Por falta de espacio, no se presentan las diferentes configuraciones regionales que se crearon a lo largo de la aplicación del RASS. Aquellas personas interesadas pueden solicitarlas directamente a los autores.

²¹ En la presentación de los resultados se excluye el ciclo 6 pues es igual al ciclo 5 lo cual representa una variación nula de la función objetivo y por lo tanto la indicación de que el heurístico ha terminado (Paso 9).

teniendo en cuenta las 38 áreas) y consigue una disminución de la función objetivo de 13.18 unidades, un 54.6% de la disminución obtenida en el primer ciclo (o un 39.6% de la disminución total)²².

FIGURA 3.

Los resultados obtenidos apuntan a que *RASS*, gracias a la incorporación de una rutina de optimización directa dentro de su algoritmo, tiene gran capacidad para alcanzar óptimos globales en un problema de regionalización. Sin embargo, es necesario tener en cuenta que la relación entre regiones a configurar (m) y áreas (n) debe ser establecida de tal forma que el número de regiones que entran al modelo de optimización (r) sea como mínimo de 2 y que dichas regiones contengan un total de áreas que se adapten a las capacidades computacionales del modelo. Para cumplir con dicho requerimiento se ha calculado que la relación m/n que se considera conveniente es $(m/n) \geq 14\%$ ²³.

4.2 Efectos de la agregación territorial en el análisis econométrico aplicado

¿Cómo se ve afectada la capacidad explicativa de una variable a medida que se utilizan distintos niveles de agrupación territorial? Aproximaciones similares a este tipo de análisis han sido realizadas por Fotheringham y Wong (1991) para modelos de regresión múltiple, lineales y logits, donde se concluye que los efectos de agregación y escalamiento son impredecibles y con efectos mucho mayores que en los análisis univariados y bivariados, y por Amrhein y Flowerdew (1992) para modelos Poisson, donde no se encuentran efectos de agregación.

En la presente aplicación se buscará explicar el Índice de Capacidad Económica Familiar (*ICEF*) en función de las variables sociodemográficas: porcentaje de personas con estudios superiores (*Estudios*), porcentaje de personas propietarias de viviendas de habitación

²² Este mismo problema fue solucionado partiendo de diferentes particiones iniciales y los resultados fueron los mismo. También se observó que el número de ciclos necesarios para alcanzar el óptimo global varían en función de la cercanía existente entre la partición inicial y el óptimo global.

²³ Por ejemplo, si se tiene un territorio formado por 8,000 áreas, el número de regiones a conformar puede ser mayor o igual a 1,120 regiones, lo cual implica que cada región contendrá 7 áreas en promedio. Esto asegura que r puede tomar valores mayores o iguales a 2 sin que el número de áreas dentro del modelo impliquen tiempos de ejecución excesivamente largos. En un caso donde la relación entre el número de regiones y el número de áreas sea muy baja, la estrategia a seguir puede ser diseñar agrupamientos anidados, lo cual implicaría la aplicación del *RASS* más de una vez.

(*Propiedad*), porcentaje de personas con cargos directivos en empresas o administraciones públicas (*Laboral*) y el porcentaje de la población mayor de 65 años (*Mayor65*)²⁴.

$$ICEF_{MCP} = \beta_0 + \beta_1 \cdot \text{Estudios} + \beta_2 \cdot \text{Propiedad} + \beta_3 \cdot \text{Laboral} + \beta_4 \cdot \text{Mayor65}$$

El mayor nivel de desagregación del cual se partirá serán las 38 “Zonas Estadísticas Grandes” (ZEG) que corresponde a una de las divisiones estadísticas oficialmente establecidas para la ciudad de Barcelona. La agrupación de dichas áreas forman los 10 distritos en los cuales se encuentra dividida la ciudad. La primera parte de la aplicación consiste agrupar las 38 ZEG en 10 regiones utilizando como criterio de agrupación la homogeneidad de las regiones en función de las variables exógenas del modelo. Para esto se aplican dos metodologías de agregación a saber: análisis cluster k-medias, lo cual implica que no se impondrá la restricción de continuidad geográfica en las agrupaciones, y el algoritmo RASS que asegura la continuidad geográfica de las agrupaciones. Las 10 agrupaciones obtenidas con ambas metodologías podrán ser comparadas con los 10 distritos oficialmente establecidos.

Para conocer la homogeneidad de cada una de las regiones se han tomado los valores de las variables para las 38 ZEG y se han calculado las desviaciones estándar entre ZEG que fueron asignadas a cada grupo. En el cuadro 3 aparecen los promedios de las desviaciones estándar intragrupos. El mayor grado de homogeneidad intragrupo se logra al agrupar las variables sin la restricción de continuidad (cluster). Cuando se impone la restricción de continuidad geográfica de los grupos el algoritmo RASS consigue formar unas regiones más homogéneas que las obtenidas con la división administrativa distrital.

CUADRO 3.

Así pues, con la aplicación del algoritmo RASS ha sido posible configurar regiones de estudio con un menor grado de dispersión intragrupo, lo cual asegura que las características obtenidas para cada grupo pueden ser más representativas en comparación con las obtenidas a partir de la división distrital oficialmente establecida. Sin embargo, es claro que un mayor nivel

²⁴ . Índice de Capacidad Económica Familiar (centrado), ICEF, 1996. Fuente: Departament d'Estadística. Ajuntament de Barcelona.
 · Nivel de instrucción, 1996, población clasificada de 10 y mas años. *Nivel de instrucción = Superior*. Fuente: Estadística de Població 1996. Institut d'Estadística de Catalunya.
 · Número de viviendas principales, 1991. *Régimen de tenencia = Propiedad*. Fuente: Cens de Població i Habitatge, 1991. Instituto Nacional de Estadística.
 · Profesión de la población ocupada, 1996. *Ocupación = Personal directivo de empresas y administraciones públicas*. Fuente: Estadística de Població 1996. Institut d'Estadística de Catalunya.
 · Edades por grandes grupos de edad de la población, 1996. *Población = 65 años y más*. Fuente: Padró Municipal d'Habitants 1996. Departament d'Estadística. Ajuntament de Barcelona.

de homogeneidad intragrupo puede ser alcanzado si se relaja la restricción de continuidad geográfica de las agrupaciones.

La segunda parte de la aplicación se plantea la siguiente pregunta: ¿Cómo puede afectar el nivel de escalamiento la bondad de las estimaciones obtenidas por mínimos cuadrados ponderados²⁵? En el cuadro 4 se presentan las desviaciones estándar y, con asteriscos, los *p_value* asociados al estadístico *t* de significación individual de las variables que han sido seleccionadas para explicar el ICEF. En la primera regresión se han utilizado los datos desagregados, para las 38 ZEG. A este nivel de desagregación todas las variables aparecen como significativas al 1%. El segundo bloque de regresiones fueron elaboradas a partir de las agrupaciones obtenidas al aplicar análisis cluster. Se puede observar como a medida que disminuye el nivel de escalamiento disminuye el potencial explicativo de las variables. Resultados similares se obtienen con las agrupaciones formadas por el *RASS*, pero los resultados parecen indicar que la velocidad con la cual las variables exógenas pierden poder explicativo es más lenta al aplicar esta metodología. Con respecto a las desviaciones estándar, se observa un importante incremento de su valor a medida que el nivel de escalamiento disminuye. A nivel distrital, y en comparación con los resultados obtenidos con el *RASS*, las variables Estudios y Propiedad conservan mejor su potencial explicativo, mientras que las variables Laboral y >65 pierden mayor potencial explicativo.

CUADRO 4.

En general, la bondad del ajuste es elevada para todos los modelos y, como era de esperar, aumenta a medida que disminuye el nivel de escalamiento. Sin embargo se puede resaltar que al hacer un análisis transversal de los resultados, los mayores R^2 siempre son los obtenidos con el *RASS*.

Así pues, parece que los resultados de las estimaciones obtenidas con las agrupaciones definidas a partir de la aplicación del *RASS* se ven menos afectados por los problemas derivados de la agrupación de territorios.

²⁵ Véase nota 18.

5. Conclusiones

En el presente trabajo se ha propuesto una nueva metodología para realizar procesos de regionalización a partir de la agregación de unidades territoriales (áreas) teniendo en cuenta no sólo sus características, sino también, las relaciones que se establecen entre ellas. Para ello, se ha formulado un modelo de optimización lineal el cual, a partir de una matriz de continuidades y una matriz de relaciones, encuentra la agrupación óptima de las áreas en un número predeterminado de regiones. El criterio utilizado para definir la calidad de las agrupaciones es la medida de heterogeneidad dentro de cada una de las regiones configuradas, medida que debe ser minimizada con el fin de obtener regiones que contengan áreas muy similares. La posibilidad de plantear el problema de regionalización como un modelo lineal permite asegurar, por sus propiedades matemáticas, que la región factible es convexa y que, por tanto, es posible encontrar una solución óptima. Otra ventaja de este tipo de formulación es que su implementación es posible en una gran variedad de software comercial.

La metodología propuesta tiene, además, una gran capacidad para identificar configuraciones territoriales complejas en cuanto a su forma, así pues, se consigue formular un modelo que cumpla con la restricción de continuidad entre las áreas que pertenecen a una región, pero sin que la metodología utilizada para cumplir con esta restricción condicione de alguna manera la forma que pueden adoptar dichas regiones. También se ha presentado un algoritmo, con el nombre de *RASS* (“*Regionalization Algorithm with Selective Search*”), que permite aumentar la capacidad de cómputo del modelo de optimización aprovechando los beneficios de aplicar la optimización directa en una porción del territorio que varía en cada ejecución gracias a una estrategia de búsqueda selectiva que define el esquema de entrada y salida de regiones. Estas características dotan al algoritmo de una buena capacidad de escapar de óptimos locales.

Los resultados obtenidos con el *RASS* son alentadores, en el sentido de que el óptimo local es encontrado independientemente de la partición inicial, sin embargo, serán necesarias múltiples ejecuciones del algoritmo para tener resultados concluyentes.

Por último, el análisis de los posibles efectos del nivel de escalamiento sobre las propiedades de las estimaciones mínimo cuadrático ponderadas (MCP), ampliamente utilizadas en el análisis econométrico aplicado a nivel territorial, ha puesto de manifiesto que los resultados de las estimaciones obtenidas con las agrupaciones definidas a partir de la

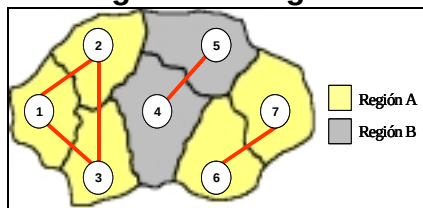
aplicación del RASS parecen verse menos afectados por los problemas derivados de la agrupación de territorios.

6. Referencias

- Aarts, E. y Lenstra, J. K. 1997. *Local search in combinatorial optimization*. Chichester, New York, [etc.]: John Wiley and Sons.
- Amrhein, C. G. y Flowerdew, R. 1992. "The effect of data aggregation on a Poisson regression model of Canadian migration", *Environment and Planning A*, **24** 1381-91.
- Artís, M., Clar, M., del Barrio, T., Guillén, M. y Suriñach, J. 2001. *Tòpics d'econometria*. Col·lecció Manuals 27. Edicions de la Universitat Oberta de Catalunya.
- Aurenhammer, F. 1991. "Voronoi diagrams - a survey of a fundamental geometric data structure", *ACM Computing Surveys*, **23** 345-405.
- Dantzing, G. B. y Ramser, J. H. 1959. "The truck dispatching problem", *Management Science*, **6** 80-91.
- Fotheringham, A. S. y Wong, D. W. S. 1991. "The modifiable areal unit problem in multivariate statistical analysis", *Environment and Planning A*, **23** 1025-44.
- Gordon, A. D. 1996. "A survey of constrained classification", *Computational Statistics & Data Analysis*, **21** 17-29.
- Gordon, A. D. 1999. *Classification*, C. Hall/CRC, (eds.) Boca Raton [etc.].
- Gower, J. C. y Legendre, P. 1986. "Metric and euclidean properties of dissimilarity coefficients", *Journal of Classification*, **3** 5-48.
- Graham, R. R. y Hell, P. 1985. "On the history of the minimum spanning tree problem", *Annals of the history of computing*, **7** 43-57.
- Hiriart, J. B., Oettli, W. y Stoer, J. 1983. *Optimization: Theory and Algorithms*, J. Baptiste y otros, (eds.) New York [etc.]: Marcel Dekker, cop.
- Horn, M. E. T. 1995. "Solution techniques for large regional partitioning problems", *Geographical Analysis*, **27** 230-48.
- Kohonen, T. 1984. *Self-Organisation and Associative Memory*. Berlin [etc.]: Springer.
- Laporte, G. y Osman, I. H. 1995. "Routing problems: A bibliography", *Annals of Operations Research*, **61** 227-62.
- Macmillan, B. y Pierce, T. 1994. "Optimization modelling in GIS framework: the problem of political redistricting", en *Spatial analysis and GIS*, S. Fotheringham y P. Rogerson, (eds.) London [etc.]: Taylor & Francis, pp. 221-46.
- Maravalle, M. y Simeone, B. 1995. "A spanning tree heuristic for regional clustering", *Communications in Statistics - Theory and Methods*, **24** 3 625-39.
- Matula, D. W. y Sokal, R. R. 1980. "Properties of Gabriel graphs relevant to geographic variation research and the clustering of points in the plane", *Geographical Analysis*, **12** 205-22.
- Murtagh, F. 1985. "A survey of Algorithms for Contiguity-constrained Clusterind and Related Problems", *The Computer Journal*, **28** 1 82-88.
- Ohsumi, N. 1984. "Practical techniques for areal clustering", en *Data analysis and informatics, Vol III*, E. Diday, M. Jambu, L. Lebart, J. Pagès y R. Tomassone, (eds.) North-Holland, Amsterdam pp. 247-58.
- Openshaw, S. 1977. "Algorithm3: a procedure to generate pseudo random aggregation of N zones into M zones where M is less than N", *Environment and Planning A*, **9** 1423-28.
- Openshaw, S. y Taylor, P. J. 1981. "The modifiable areal unit problem", en *Quantitative Geography*, N. Wrigley y R. J. Bennett, (eds.) London pp. 60-70.
- Openshaw, S. 1984. "The modifiable areal unit problem", *Concepts and Techniques in Modern Geography*, **38** GeoAbstracts, Norwich.
- Openshaw, S. 1992. "Some suggestions concerning the development of artificial intelligence tools for spatial modelling and analysis in GIS", *The Annals of Regional Science*, **26** 35-51.
- Openshaw, S. y Wymer, C. 1994. "Classifying and regionalizing census data", en *Census Users Handbook*, S. Openshaw, (eds.) Cambridge, UK: Geo Information International, pp. 239-70.
- Perruchet, C. 1983. "Constrained agglomerative hierarchical classification", *Pattern Recognition*, **16** 213-17.
- Semple, R. K. y Green, M. B. 1984. "Classification in human geography", en *Spatial statistics and models*, G. L. Gaile y C. J. Wilmott, (eds.) Reidel, Dordrecht pp. 55-79.
- Suriñach, J., Duque, J. C., Ramos, R. y Royuela, V. 2003. "Publication patterns in regional and urban analysis: have topics, techniques and applications changed during the 1990s?", *Regional Studies*, **37** 4 351-63.
- Toussaint, G. T. 1980. "The relative neighbourhood graph of a finite planar set", *Pattern Recognition*, **12** 261-68.
- Wise, S. M., Haining, R. P. y Ma, J. 1997. "Regionalization Tools for Exploratory Spatial Analysis of Health Data", en *Recent Developments in Spatial Analysis: Spatial statistics, behavioural modelling, and computational intelligence*, M. M. Fisher y A. Gentis, (eds.) Berlin [etc.]: Springer, pp. 83-100.

7. Cuadros y figuras

Figura 1. Configuración regional no factible



Cuadro 1. Configuración de áreas con y sin ciclos potenciales

Con ciclos potenciales	Sin ciclos potenciales

Cuadro 2. Partición inicial y solución obtenida por el heurístico

Óptimo preestablecido	Partición inicial	Óptimo RASS
AGRUPACIONES ÓPTIMAS	PARTICIÓN INICIAL	AGRUPACIONES ÓPTIMAS

Figura 2. Evolución de la función objetivo a lo largo del heurístico

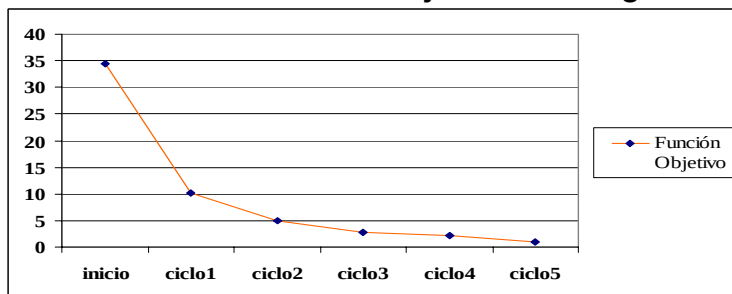
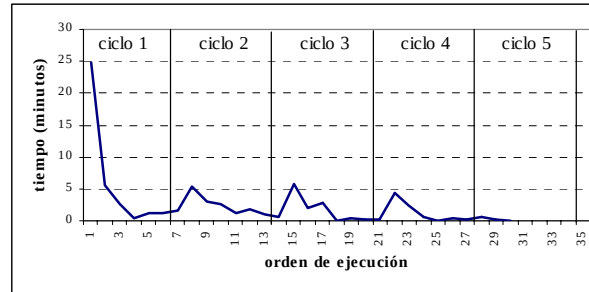


Figura 3. Tiempos de ejecución de los modelos de optimización



Cuadro 3. Promedio de las desviaciones estándar intragrupos (m=10)

Método	Estudios	Propiedad	Laboral	>65
Cluster ^(a)	2.09%	2.95%	0.89%	1.72%
Cluster ^(b)	2.61%	3.69%	1.11%	2.16%
RASS	3.19%	7.57%	1.42%	2.71%
Distritos	4.38%	8.35%	2.08%	3.00%

^(a) Asignando una desviación estándar = 0 a las regiones compuestas por una ZEG.

^(b) Calculado a partir de las regiones formadas por mas de una ZEG.

Cuadro 4. Desviaciones estándar y *p*_values por estimación MCP por método y nivel de escalamiento

Método	m=38	m=25	m=19	m=14	m=10
ZEG					
(Constante)	12.05 ***				
Estudios	26.81 ***				
Propiedad	8.54 ***				
Laboral	66.00 ***				
>65	27.56 ***				
<i>R</i> ² corregido	0.962				
<i>F</i>	0.00				
Cluster					
(Constante)	14.95 **	17.70	18.27	91.27	
Estudios	31.89 ***	37.11 ***	39.83 ***	59.31	
Propiedad	10.10 ***	11.05 ***	12.76 ***	62.21 *	
Laboral	79.64 ***	94.26 ***	91.43 ***	155.20 **	
>65	34.61 **	40.97	41.65	226.73	
<i>R</i> ² corregido	0.966	0.971	0.979	0.980	
<i>F</i>	0.00	0.00	0.00	0.00	
RASS					
(Constante)		13.88 ***	19.25	26.49	
Estudios		29.24 ***	39.19 ***	45.44 *	
Propiedad		11.04 ***	12.56 ***	19.43 **	
Laboral		76.93 ***	109.76 **	112.02 ***	
>65		35.78 **	44.40	69.47	
<i>R</i> ² corregido		0.985	0.987	0.995	
<i>F</i>		0.00	0.00	0.00	
Distritos					
(Constante)					31.70
Estudios					36.79 ***
Propiedad					20.79 **
Laboral					105.86 *
>65					72.19
<i>R</i> ² corregido					0.993
<i>F</i>					0.00

*** Significativa al 1%

** Significativa al 5%

* Significativa al 10%