

Tesis: Evaluación de Métodos No Paramétricos Aplicados a la Elasticidad-Crecimiento de la Pobreza

**Maestría en Economía
Facultad de Ciencias Económicas
Universidad Nacional de La Plata**

Tesista: Lic. Epele, Nicolás

Director: Dr. Walter Sosa Escudero

Evaluación de Métodos No Paramétricos Aplicados a la Elasticidad-Crecimiento de la Pobreza

Epele, Nicolás¹

Facultad de Ciencias Económicas,

UNLP, La Plata, Argentina

En este artículo se evalúan métodos no paramétricos de estimación de densidades univariadas con el fin de medir el efecto del crecimiento económico sobre la tasa de pobreza, que es el indicador de pobreza de mayor difusión en opinión pública. La literatura de crecimiento pro-pobre ha calculado la elasticidad-crecimiento de la pobreza a partir de ajustar la distribución del ingreso por medio de funciones de densidad conocidas. Sin embargo, cuando se ignoran las formas funcionales subyacentes, las aproximaciones por funciones teóricas pueden no captar cualidades de la densidad poblacional, que sí pueden deducirse por medio de métodos no paramétricos.

Debido a que la elasticidad-crecimiento de la tasa de pobreza es sensible a dichas características, en este trabajo se comparan los métodos no paramétricos de kernel, k-ésimo vecino cercano, kernel adaptativo y logspline en la estimación de dicha elasticidad en base a muestras aleatorias de distribuciones teóricas habitualmente empleadas por la literatura de distribución del ingreso. La comparación se realiza a partir del error cuadrático medio integrado y como resultado de este ejercicio se obtiene que los métodos no paramétricos ganan mucho en precisión cuando se calcula la elasticidad para tasas de pobreza superiores al 10%. En base a esto se estima la elasticidad-crecimiento de la tasa de pobreza para Gran Buenos Aires en 1974, 1986, 1998 y 2005, observándose una caída en la sensibilidad de dicho indicador al crecimiento económico a lo largo de los años.

1- Introducción

El objetivo central de este artículo es evaluar la bondad en el ajuste de los métodos no paramétricos de estimación de densidades univariadas más populares para medir el efecto del crecimiento económico sobre la tasa de pobreza. A pesar de sus inconvenientes para representar la pobreza, el interés por este indicador, definido como la proporción de la población con ingresos inferiores a uno de referencia, reside en que su fácil interpretación lo ha convertido en el índice de pobreza de mayor difusión en la opinión pública.

Conocer la capacidad del crecimiento para reducir la pobreza ha sido uno de los interrogantes que ha concentrado la atención de los investigadores en los últimos años, cuyo esfuerzo se ha centrado en calcular la elasticidad-crecimiento de algún índice de pobreza, entendida como la variación porcentual de dicho indicador ante un aumento en un punto porcentual del ingreso medio.

Con este objetivo, la literatura de crecimiento pro-pobre ha recurrido a diversas metodologías. En primer lugar puede citarse el análisis que intenta explicar la variación de la pobreza entre dos puntos del tiempo por medio de regresiones lineales sin hacer importantes especificaciones funcionales sobre la distribución del ingreso, que tiene al trabajo de Dollar y Kraay [7] como uno de sus principales representantes. Sobre un panel de 953 observaciones, correspondiente a 137 países desarrollados y en vías de desarrollo del período 1950-1999, estos autores estiman empíricamente la relación

¹ Becario del CONICET. e-mail: e_nico77@yahoo.com. Este trabajo ha sido dirigido por el Dr. Walter Sosa Escudero.

entre el ingreso medio de los pobres y el de la población total, hallando que el cambio de este último afecta al primero en igual medida que al ingreso de los restantes sectores de la sociedad. Sin embargo, este tipo análisis no permite trabajar con una línea de pobreza fija, por lo que impone la necesidad de adoptar el criterio de pobreza relativa, considerándose en este caso como pobres al primer quintil de la población. Como refiere Ravallion [22], este criterio implica que los países ricos tengan una línea de pobreza mayor que los países en vías de desarrollo, lo que suaviza el efecto del crecimiento sobre la pobreza. Por otra parte, realizar regresiones lineales sin consideraciones sobre las formas funcionales subyacentes allana la posibilidad de encontrar una relación más compleja entre crecimiento y pobreza, sugerida por Bourguignon [3].

En segundo lugar pueden mencionarse los análisis de regresiones que parten de suponer alguna forma teórica de la función de densidad de los ingresos. Entre estos se encuentra el trabajo de López y Servén [19] que, sobre la misma base de datos de Dollar y Kraay [7], estudia la relación pobreza-crecimiento-desigualdad a partir de suponer una distribución lognormal de los ingresos, lo que les permite incorporar el criterio de pobreza absoluta y calcular la tasa de pobreza para cada país. Con este propósito los autores dividen su trabajo en dos etapas. En la primer etapa testean la bondad en el ajuste del modelo lognormal regresando por mínimos cuadrados ordinarios el vector de quintiles empíricos contra el de quintiles teóricos construido bajo dicha especificación, para 3.176 observaciones correspondientes a 130 países. Como resultado, obtienen una relación de identidad entre ambos vectores con un R^2 superior a 96%.² Dado esto, en la segunda etapa realizan regresiones lineales de variaciones de la tasa de pobreza teórica contra variaciones del ingreso medio y la desigualdad. Si bien la estimación de una elasticidad-crecimiento de la pobreza promedio de distintos países y períodos tiene poca relevancia informativa, esta metodología les permite deducir, entre otras conclusiones, que la elasticidad-crecimiento de la pobreza depende de la forma funcional de la densidad de los ingresos. El test de lognormalidad utilizado por López y Servén [19] presenta algunas debilidades, ya que los quintiles empíricos no pueden diferir mucho de los teóricos y en consecuencia difícilmente pueda rechazarse la hipótesis nula de la especificación lognormal. En la misma línea de análisis, Bourguignon [2] concluye que la hipótesis de lognormalidad de los ingresos no es tan satisfactoria. Sobre una muestra de 114 observaciones correspondientes a 50 países en vías de desarrollo de las décadas del 80 y el 90, este autor regresa las variaciones de la tasa de pobreza y de la brecha de pobreza contra las del ingreso medio multiplicada por las respectivas elasticidades teóricas bajo la especificación lognormal junto a una serie de variables de control, obteniendo un R^2 de 69% y 25% para cada caso.

Tanto el trabajo de López y Servén [19] como el de Bourguignon [2] requieren definir una línea de pobreza común para todos los países de la muestra, lo que redundo no sólo una gran heterogeneidad en los niveles de pobreza obtenidos, sino que los resultados quedan sesgados por las observaciones en que dicha línea de pobreza es relevante. Asimismo, cuando se trabaja con información en tiempo discreto, las especificaciones funcionales pueden generar grandes errores de aproximación, que crecen con la distancia entre los puntos de tiempo considerados. Por ello, es necesario avanzar en dirección del estudio específico por países.

Los artículos de Datt y Ravallion [4] y Kakwani [16] de principios de la década del 90, realizados para Brasil e India, y Costa de Marfil respectivamente, son pioneros del análisis específico por países. La metodología seguida por estos artículos consiste

² Este R^2 es alcanzado cuando López y Servén [19] trabajan con la muestra completa. Los autores consiguen aún mejores resultados cuando trabajan sólo con las observaciones de gastos o sólo de los ingresos netos, alcanzando un R^2 de 98%.

en estimar la elasticidad-crecimiento de varios indicadores de pobreza empleando una aproximación paramétrica de la curva de Lorenz. No obstante, como muestran ambos estudios, la elasticidad de la tasa de pobreza depende de la estimación puntual de la función de densidad en la línea de pobreza y por tanto de las características dicha función, como sesgos y multimodalidades, que en muchos casos sólo pueden deducirse por medio de métodos no paramétricos.

Por ello, en la Sección 2 se exponen los métodos no paramétricos de kernel, k-ésimo vecino cercano, kernel adaptativo y logspline y se los compara en su bondad en el ajuste. Con este propósito se realizan distintas simulaciones de distribuciones teóricas habitualmente empleadas por la literatura de distribución del ingreso y se calcula el error cuadrático medio integrado (*MISE*) de las estimaciones por cada método, tanto de la estimación de la función de densidad como de la elasticidad-crecimiento de la pobreza.

En base a los resultados obtenidos, en la Sección 3 se calcula la elasticidad-crecimiento de la pobreza de Gran Buenos Aires (GBA) para distintos períodos, sobre información surgida de la Encuesta de Permanente de Hogares (EPH) de Argentina elaborada por el Instituto de Estadísticas y Censos (INDEC). Finalmente en la Sección 4 se presentan las conclusiones.

2- Estimación no paramétrica de la elasticidad-crecimiento de la pobreza

La variación de los ingresos puede descomponerse, por un lado, en un cambio proporcional de los mismos para todas las personas por igual, es decir, del ingreso medio de la población sin modificar la forma funcional de la curva de Lorenz, y por otro lado, en transferencias entre ellas, lo que varía la distribución del ingreso conservando el ingreso promedio. Dada una línea de pobreza fija, ambos movimientos afectan a la pobreza de la siguiente manera:

$$d\rho = \frac{\partial \rho}{\partial \bar{x}} d\bar{x} + \frac{\partial \rho}{\partial \theta} d\theta$$

siendo ρ un indicador de pobreza, $\bar{x}$ el ingreso medio y θ algún índice de desigualdad.

La elasticidad-crecimiento de la pobreza mide la primera de las variaciones, que en la literatura de crecimiento pro-pobre se la conoce como “efecto crecimiento puro”. Si se identifica a los ingresos con la variable x , la línea de pobreza puede definirse como el valor z de la misma. En este caso, la tasa de pobreza, H , es igual a la función de distribución acumulada de los ingresos evaluada en z , $F(z)$. Si se aumenta infinitesimalmente el ingreso de todos los individuos en ε por ciento, se obtiene una nueva función de distribución acumulada $\tilde{F}$ con un nuevo ingreso medio $\tilde{\bar{x}}$, de manera que la elasticidad-crecimiento de H se anota como:

$$\begin{aligned} \eta_H &= \lim_{\varepsilon \rightarrow 0} \frac{(\tilde{F}(z) - F(z)) / F(z)}{(\tilde{\bar{x}} - \bar{x}) / \bar{x}} \\ &= \lim_{\varepsilon \rightarrow 0} \frac{(\tilde{H} - H) / H}{(\tilde{\bar{x}} - \bar{x}) / \bar{x}} \end{aligned}$$

Dado que $\tilde{F}(x) = F(x - \varepsilon x)$, en el Apéndice se muestra que η_H puede describirse del siguiente modo:

$$\eta_H = -\frac{zf(z)}{F(z)} = -\frac{zf(z)}{H} \quad (1)$$

donde f representa la función de densidad de los ingresos. Por tanto, la estimación de η_H se reduce a aproximar H y $f(z)$. Kakwani [17] prueba que H puede aproximarse no paramétricamente por:

$$H = \frac{1}{n} \sum_{i=1}^n 1(X_i \leq z) \quad (2)$$

siendo $1(\cdot)$ una función indicadora que vale uno cuando $x \leq z$ y cero en el resto de los casos. En tanto que el valor de la función de densidad evaluada en z puede estimarse por diferentes métodos no paramétricos. Por ello, en esta a continuación se exponen y evalúan los métodos no paramétricos de kernel, k-ésimo vecino cercano, kernel adaptativo y logspline.

2.a- Métodos no paramétricos de estimación de densidades univariadas

La estimación no paramétrica de funciones de densidad más conocida es el histograma. Si bien éste nació como una aproximación visual de la información estadística, constituyó la primer estimación no paramétrica de una función de densidad. Sin embargo, su construcción le impone limitaciones para captar las características de la misma, ya que la forma funcional obtenida puede depender fuertemente tanto de la elección de los rangos y su punto de partida, como de las discontinuidades que presenta, que no son propias de la distribución subyacente.

La necesidad de resolver estas dificultades motivaron la aparición de nuevos estimadores de la función de densidad, siendo el más popular de ellos el de kernel:

$$\hat{f}_1(x) = \frac{1}{nh} \sum_{i=1}^n K\left(\frac{x - X_i}{h}\right) \quad (3)$$

Este estimador corrige la primera de las limitaciones del histograma al ponderar, por medio de la función de kernel K , todas las observaciones de la muestra para aproximar f en cada punto del dominio. En tanto que las discontinuidades que exhibe el histograma se resuelve eligiendo K continua. Básicamente se requiere que K sea una función de densidad simétrica con media cero.³ Sin embargo, el estimador $\hat{f}_1$ (así como sus derivados, que se presentan a continuación) tiene inconvenientes cuando la variable de análisis es acotada, como es el caso de los ingresos que sólo toman valores positivos, que pueden corregirse empleando alguna variante de la función K , como la presentada en la Sección 2.b.

No obstante, La forma funcional de $\hat{f}_1$ depende fundamentalmente de la elección del parámetro h , conocido como “ancho de banda”. Cuando éste toma valores chicos la distribución resultante presenta una gran variabilidad, ya que para la estimación de f evaluada en x se tendrán en cuenta sólo aquellas observaciones

³ Algunos ejemplos de funciones $K(\psi)$ son:

- Triangular: $(1 - |\psi|) \mathbf{1}(|\psi| \leq 1)$
- Cuadrático: $\frac{15}{16} (1 - \psi^2) \mathbf{1}(|\psi| \leq 1)$
- Gaussiana: $\frac{1}{\sqrt{2\pi}} e^{-\psi^2}$
- Epanechnikov: $\frac{3}{4} (1 - \psi^2) \mathbf{1}(|\psi| \leq 1)$

La función de Epanechnikov resulta de elegir K que minimice el *MISE* (Pagan y Ullah [21]). En las estimaciones realizadas en este artículo se emplea la función gaussiana, con la variación presentada en la Sección 2.b.

cercanas a dicho punto. En tanto que si h es grande la función obtenida será suave, debido a que la estimación de f para cada punto del dominio dará una alta ponderación a todas las observaciones de la muestra. La necesidad de contar con una regla automática para la elección de h motivó su obtención por algún criterio de optimización. El más difundido es que minimiza el *MISE*, del que se tiene que $h \propto n^{-1/5}$ (Silverman [26]).⁴

Del estimador de kernel se derivan otros en los que h no es constante, lo que les permite una mayor convergencia a la función estimada (Pagan y Ullah [21]). La motivación para ello es que en las regiones del dominio en que la muestra presenta una importante concentración de observaciones se espera la función de densidad subyacente esté cambiando de valor rápidamente, es decir, que $f^{(2)}(x)$ es grande. Esto requiere que en dicha región se dé una alta ponderación a las observaciones próximas a x , y por tanto, que el valor de h sea chico. En tanto que en las regiones con pocas observaciones se supone que la función de densidad es plana. Allí las ponderaciones de observaciones alejadas de x pueden ser mayores y por tanto h puede ser grande. El primero de estos estimadores considerado es el del k -ésimo vecino cercano:

$$\hat{f}_2(x) = \frac{1}{2nd_k(x)} \sum_{i=1}^n K\left(\frac{x - X_i}{2nd_k(x)}\right) \quad (4)$$

siendo $d_k(x)$ la distancia a la k -ésima observación más cercana. En la práctica el valor elegido de k suele ser $\lceil \sqrt{n} \rceil$, de manera que este parámetro se ajusta al tamaño de la muestra. Como puede observarse la distancia $d_k(x)$ depende del punto en que se esté haciendo la estimación, lo que hace que la función resultante pueda no ser una función de densidad, ya que no necesariamente integra en la unidad.

Este inconveniente se corrige calculando anchos de banda que sólo dependen de la ubicación de las observaciones muestrales, como en el estimador de kernel adaptativo:

⁴ Existe una extensa discusión acerca del criterio de optimización elegido para la obtención del ancho de banda. Los criterios comúnmente aceptados son dos: El primero de ellos consiste en elegir h tal que minimice el error cuadrático integrado (*ISE*) determinado por:

$$ISE(h) = \int \{f(x) - \hat{f}_1(x)\}^2 dx \quad (5)$$

Dado que el *ISE* es una variable aleatoria, ya que depende de la muestra aleatoria a la que se aplica, tiene una distribución de la que puede obtenerse su valor medio *MISE*, que es el segundo criterio de optimización:

$$MISE(h) = E_x \int \{f(x) - \hat{f}_1(x)\}^2 dx \quad (6)$$

Si bien hay evidencias de que prácticamente no existen desemejanzas en los resultados de uno y otro criterio (Grund, Hall y Marron [10]), la diferencia fundamental entre ambos criterios surge asintóticamente, debido a que el cálculo del h óptimo empleando (5) depende en mayor medida de f de lo que lo hace (6) (Jones [13]).

Por ello, el ancho de banda con el que se trabaja en la Sección 3 es aquel que minimiza el *AMISE*, que es la aproximación de *MISE* cuando $n \rightarrow \infty$:

$$h_{AMISE} = \left[\frac{\int K^2(x) dx}{\left(\int f^{(2)}(x) dx \right) \left(\int x^2 K(x) dx \right)^2} \right]^{1/5} n^{-1/5}$$

La dificultad práctica que surge en este caso es aproximar $(f^{(2)})^2$, para lo que se han desarrollado varios métodos de optimización que aproximan su valor (Jones, Marron y Sheater [15]).

$$\hat{f}_3(x) = \frac{1}{n} \sum_{i=1}^n \frac{1}{h_{ni}} K\left(\frac{x - X_i}{h_{ni}}\right) \quad (7)$$

El ancho de banda $h_{ni} = h\delta_{ni}$, siendo $\delta_{ni} = \left(\frac{G}{\dot{f}(X_i)}\right)^\gamma$, $\dot{f}$ una estimación de f con

el ancho de banda fijo h y G la media geométrica de $\{\dot{f}(X_i)\}_{i=1}^n$ (Hall y Marron [11]). γ es un parámetro suavizador que suele tomárselo igual a 0,5, aunque no hay ningún motivo claro para ello (así como para emplear la media geométrica), de modo que el estimador de kernel es un caso particular del de kernel adaptativo para el que γ es igual a cero.

El último estimador considerado es el de logspline, propuesto por Kooperberg y Stone [18], que resuelve las limitaciones del histograma siguiendo una lógica diferente de (3), (4) y (7). Dada una secuencia de números reales $-\infty < t_1 < t_2 < \dots < t_T < \infty$, llamados nodos, que forman una partición del dominio $(-\infty, t_1], [t_1, t_2], \dots, [t_{T-1}, t_T], [t_T, \infty)$, el estimador de logspline se obtiene ajustando una función s , que es una combinación lineal de polinomios cúbicos $B_1, \dots, B_{T-1}$ definidos bajo ciertos criterios para los intervalos mencionados, junto a la función constante 1.⁵ La expresión del logspline es:

$$\hat{f}_4(x) = \exp\{\theta_1 B_1(x) + \dots + \theta_{T-1} B_{T-1}(x) - c(\theta)\} \quad (8)$$

siendo:

$$c(\theta) = \ln \int \exp\{\theta_1 B_1(x) + \dots + \theta_{T-1} B_{T-1}(x)\}$$

la constante que normaliza $\hat{f}_4$. El valor de los parámetros $\theta_1, \dots, \theta_T$ se estima por máxima verosimilitud, siendo:

$$l(\theta) = \sum_{i=1}^n \ln(\hat{f}_4(\theta; X_i)) \quad (9)$$

el logaritmo de la función de verosimilitud.

La ubicación de los nodos no es determinante de la forma funcional de $\hat{f}_4$, y por ello se sigue una regla automática. En tanto que sí lo es su cantidad, que equivale a determinar el ancho de banda en el estimador de kernel: cuando hay muchos nodos la función estimada presenta una gran variabilidad y cuando son pocos, la función resultante es suave. Kooperberg y Stone [18] muestran que la cantidad de nodos debe estar relacionada con la estructura de la muestra analizada, como son las modas, asimetría y sesgos, y desarrollan un algoritmo basado en el criterio de información Akaike para su determinación. Este consiste en partir de una cantidad T' de nodos e ir eliminando nodos hasta minimizar la función:

$$AIC_{am} = -2\hat{l}_m + \alpha(T' - m - 1)$$

siendo $\hat{l}_m$ la función (9) evaluada en su máximo valor para el modelo (8). α es un parámetro de penalización, que de aquí en adelante se toma igual a $\ln(n)$, propuesto por Schwartz [24].

⁵ Kooperberg y Stone [18] proponen que los polinomios cúbicos B_i cumplan con que en el intervalo $(-\infty, t_1]$ B_1 sea lineal con pendiente negativa y $B_2, \dots, B_{T-2}$ sean constantes, en tanto que en el intervalo $[t_T, \infty)$ B_{T-1} sea lineal con pendiente positiva y $B_2, \dots, B_{T-2}$ constantes.

2.b- Elección de la función kernel

Como se dijo anteriormente, cuando la variable de análisis es acotada los métodos de kernel, k-ésimo vecino cercano y kernel adaptativo generan un sesgo en la estimación de la función de densidad en los bordes, conocido como “efecto límite”, que en el caso de los ingresos corresponde a los valores cercanos a cero. Esto se debe a que (3), (4) y (7) permiten dar una probabilidad de ocurrencia positiva a puntos fuera del tramo relevante, sesgando la estimación de la función de densidad, y por tanto de la elasticidad-crecimiento de la pobreza en los casos en que la línea de pobreza se encuentra próxima al límite.

La solución más elemental a este inconveniente es truncar $\hat{f}$ al tramo de valores relevantes de x y renormalizarla de modo que siga integrando en la unidad, es decir, que continúe siendo una función de densidad. No obstante, esta estrategia no corrige el sesgo original de la estimación en la cola inferior de la distribución. Por ello, se han propuesto diferentes soluciones para aliviar esta dificultad, descritas en el trabajo de Jones [14].

Entre otras, las técnicas de “filo de navaja” corrigen satisfactoriamente el “efecto límite” por medio de una solución sencilla: reemplazan la función de kernel K de (3), (4) y (7) por una combinación lineal de la misma.⁶ En las secciones siguientes la función de kernel empleada es la sugerida por Schucany y Sommers [25], que resulta de una combinación lineal de la función gaussiana:

$$K(x) = \frac{(a_2(p) - a_1(p)x)\phi(x)}{a_0(p)a_2(p) - a_1^2(p)} \quad (10)$$

donde $a_l(p) = \int_{-1}^p u^l \phi(u) du$ y $p=xh$. Como puede observarse, la expresión (10) se comporta como el kernel gaussiano para valores grandes de x .

2.c- Evaluación de los métodos no paramétricos

La función de densidad del ingreso ha sido aproximada por una gran variedad de modelos capaces de captar sus cualidades más elementales: sólo da probabilidad positiva a los valores positivos de la variable y presenta asimetría con sesgo hacia la derecha. Para evaluar la bondad en el ajuste de los métodos no paramétricos en estimar la elasticidad-crecimiento de la pobreza, en esta sección se practican una serie de simulaciones a partir de algunas de las densidades habitualmente empleadas por la literatura de distribución del ingreso.

La primera de las funciones consideradas es la lognormal, introducida en el contexto de distribución del ingreso por Gibrat en 1931. La forma funcional de la lognormal es:

$$f_{LN}(x) = \frac{1}{x\sqrt{2\pi}} \exp\left\{-\frac{1}{2} \frac{(\ln(x) - \mu)^2}{\sigma^2}\right\} \quad 0 < x$$

siendo μ y σ la media y desvío estándar del logaritmo natural de x . La razón de su popularidad se debe a que sus parámetros son de fácil interpretación y se relacionan directamente con las medidas de desigualdad.

⁶ El estimador de kernel tiene una tasa de convergencia en la reducción del sesgo del orden de $n^{-4/5}$. En el caso de las variables acotadas, la tasa de convergencia en puntos cercanos a los bordes se reduce a $n^{-2/3}$. La técnica de “filo de la navaja” permite recuperar la tasa de convergencia de $n^{-4/5}$ en dichos puntos (Jones [14]).

Salem y Mount [23] prueban que la distribución gamma ajusta mejor a la densidad del ingreso familiar en EE.UU. entre 1960 y 1970 que la aproximación lognormal. Dado que ésta es empleada en el contexto de la distribución del ingreso frecuentemente, la gamma es la segunda función de densidad sobre la que se hacen las simulaciones:

$$f_G(x) = \frac{x^{\alpha-1} e^{-x/\beta}}{\Gamma(\alpha)\beta^\alpha} \quad 0 < x ; \alpha, \beta > 0$$

donde $\Gamma(\alpha) = \int e^u u^{\alpha-1} du$. A α y β se los denomina parámetros de “escala” y “forma” respectivamente.

Si bien la lognormal y la gamma son de fácil estimación, Chakravarty y Majumder [8] muestran que su ajuste es peor que las aproximaciones por distribuciones de tres parámetros, debido a que estas tienen una mayor flexibilidad funcional. Por esto, la última distribución considerada es la propuesta por Dagum en 1977:

$$f_D(x) = apb \left(\frac{x}{b} \right)^{ap-1} \left(1 + \left(\frac{x}{b} \right)^a \right)^{-p-1} \quad 0 < x ; a, p, b > 0$$

siendo b un parámetro de “escala” y a y p de “forma”.

Las simulaciones consisten en estimar por medio de los métodos no paramétricos las tres distribuciones teóricas, para distintos valores sus parámetros, a partir de 30 muestras aleatorias de 4.000 observaciones cada una y evaluar su bondad en el ajuste en base al *MISE*. El *MISE* es una norma de funciones definida por:

$$MISE = E_x \int_a^b \{f(x) - \hat{f}(x)\}^2 dx$$

en la que la integral hace que esta medida se ocupe de la diferencia entre f y $\hat{f}$ en todo el intervalo $[a,b]$ y no sólo de problemas locales, en tanto que el hecho de que el término entre llaves esté elevado al cuadrado hace que el integrando sea no negativo y cercano a cero cuando $\hat{f}(x)$ está próxima a $f(x)$.

Con este propósito, los parámetros dados para cada especificación son los que surgen de su estimación por máxima verosimilitud de la función de densidad del ingreso per cápita familiar por adulto equivalente (ipcf-ae) de GBA del tercer trimestre de 1974, 1986, 1998 y 2005, cuyos resultados se presentan en la siguiente tabla.⁷

⁷ Para 1974, 1986 y 1998 la información de la EPH corresponde al mes de septiembre, obtenida por el método de ondas. Por su parte, la información de 2005 pertenece la EPH-Continua del último trimestre.

Tabla 1: Estimación de los parámetros poblacionales por máxima verosimilitud del la densidad del ipcf-ae para las distribuciones lognormal, gamma y dagum Gran Buenos Aires. 1974, 1986, 1998 y 2005

Distribución	Parámetro	1974	1986	1998	2005
Lognormal	μ	6,1	5,8	5,4	5,3
	σ	0,6	0,7	0,8	0,8
Gamma	α	3,0	2,1	1,6	1,6
	β	181,3	207,0	198,9	168,9
Dagum	a	3,3	2,5	2,3	2,3
	b	524,2	321,3	263,3	238,1
	p	0,8	1,1	0,8	0,8
Observaciones		6.899	8.678	8.207	5.742

Fuente: Elaboración propia en base a EPH-INDEC

La selección de estos años se debe a que corresponden a períodos de estabilidad de la economía Argentina en los que la distribución del ingreso ha variado de modo tal de presentar tasas de pobreza muy disímiles, como se observa en la Tabla 2. Esto permite evaluar la bondad en el ajuste de los estimadores en diferentes situaciones.

En la Tabla 4 se presenta el *MISE* de las estimaciones no paramétricas de f_{LN} , f_G y f_D para distintos valores de sus parámetros, calculado en los intervalos [0,200], [0,100] y [100,200]. Asimismo, esta tabla incluye el *MISE* de las estimaciones por modelos paramétricos realizadas empleando máxima verosimilitud. Esta tabla se corresponde con las Figuras 1.1, 2.1, 3.1 y 4.1 en las que se presentan las diferencias entre cada distribución teórica y el promedio de las 30 estimaciones realizadas por cada muestras. En estos gráficos se destacan la línea de indigencia y pobreza del ipcf-ae, que en Argentina es de \$70 y \$165 (a precios de 1999).

Como muestra dicha tabla, las mejores aproximaciones se alcanzan utilizando los modelos paramétricos siempre que la especificación funcional elegida para ello sea igual a la de la densidad a estimar. Sin embargo, cuando la forma funcional del modelo paramétrico y la distribución a estimar no coinciden, las mejores estimaciones se obtienen empleando los estimadores no paramétricos en la gran mayoría de los casos. (Las únicas excepciones se producen en las estimaciones de las distribuciones lognormal y gamma con los parámetros de 1974, en las que el estimador de k-ésimo vecino cercano genera un mayor *MISE* que el de la aproximación paramétrica suponiendo una especificación f_D tanto en [0,200] como en [0,100].) Esto se debe a que el k-ésimo vecino cercano tiene dificultades en estimar la densidad en los valores próximos a cero, a pesar de la corrección sugerida en la Sección 2.b.

De los métodos no paramétricos, el estimador de logspline es el que realiza el mejor ajuste en el tramo [0,200] prácticamente en todas las ocasiones, seguido por el estimador de kernel, kernel adaptativo y k-ésimo vecino cercano en ese orden. Sin embargo, cuando se observan los *MISE* en [100,200], el estimador de logspline conserva su superioridad estimativa, aunque su ventaja sobre los restantes métodos disminuye y ya no puede establecerse un orden claro de los restantes estimadores. Ejemplo de ello es que en dicho intervalo el kernel adaptativo se muestra como la peor opción para aproximar las distribuciones con parámetros de 1974, es decir, para distribuciones con una baja tasa de pobreza. En tanto que para los años 1998 y 2005, el kernel adaptativo se comporta mejor que los métodos de kernel y k-ésimo vecino cercano.

A medida que se pasa de 1974 a 2005 todos los estimadores mejoran su ajuste promedio en los tres intervalos, como se observa en las Figuras 1.1 a 4.1. Esto se debe a las dificultades de estimación de las densidades de estos métodos en los primeros cuantiles, que se evidencia gráficamente. A la vez, se aprecia cómo crece la dispersión de dicha estimación año a año, siendo el logspline el estimador de mayor dispersión en el intervalo [0,100], y sólo es superado por el k-ésimo vecino cercano en [100,200].

Por su parte, la Tabla 5 muestra los *MISE* de la estimación de la elasticidad-crecimiento de la pobreza en los intervalos [50,200], [50,125] y [125,200] para cada caso, siendo las Figuras 1.2, 2.2, 3.2 y 4.2 sus correspondientes representaciones gráficas. Como puede apreciarse, las estimaciones no paramétricas de η_H de las tres densidades con parámetros de 1974 y 1986 son peores que las paramétricas bajo las especificaciones f_G y f_D en varias oportunidades, aún cuando dicha especificación no coincide con la función de densidad a estimar. Aunque no existe un orden definido para los estimadores no paramétricos en los años 1974 y 1986, el de k-ésimo vecino cercano es el mejor estimador de la elasticidad. En las Figuras 1.2 y 2.2 puede verse que en varios puntos del intervalo [50,200] la diferencia entre la estimación no paramétrica de dicha elasticidad y su verdadero valor para los años 1974 y 1986 es mayor a 1 punto. Esto se debe a que, dadas las características de las distribuciones bajo los parámetros de 1974 y 1986, tanto $f(z)$ como H , coeficientes del numerador y denominador respectivamente de η_H en la ecuación (1), toman valores chicos y por tanto pequeños errores en su estimación se traducen en grandes errores en la aproximación de la elasticidad-crecimiento de la pobreza. En tanto que, para las distribuciones f_{LN} , f_G y f_D con parámetros de 1998 y 2005, la mejor estimación de la elasticidad-crecimiento de la pobreza se consigue no paramétricamente, siendo el método de mejor ajuste el de logspline. Para estos años el error en la estimación es inferior al 0,25 puntos para todos los estimadores no paramétricos.

Los estimadores no paramétricos mejoran la estimación de η_H a medida que se la calcula para líneas de pobreza mayores, o en el mismo sentido para mayores tasas de pobreza. A diferencia de la estimación de la función de densidad, las Figuras 1.2, a 4.2 muestran que la dispersión en la estimación de η_H se reduce a medida que aumenta el valor de la línea de pobreza considerada.

3- Elasticidad-crecimiento de la pobreza en GBA

En esta sección se presentan las estimaciones no paramétrica de la elasticidad-crecimiento de la tasa de indigencia y pobreza de GBA. La elección de esta región se debe a que es la de mayor población de Argentina para la que se dispone de la serie de encuestas de hogares más extensa del país. En la Tabla 2 se muestra como el *ipcf-ae* medio cae de 1974 a 2005, con el consecuente aumento de las tasas de indigencia y pobreza.

**Tabla 2: *ipcf-ae* medio, tasa de indigencia y tasa de pobreza ⁸
En pesos (a precios de 1999) y puntos porcentuales
Gran Buenos Aires. 1974, 1986, 1998 y 2005**

Año	1974	1986	1998	2005
ipcf-ae medio	539,2	444,5	325,1	274,9
Tasa de indigencia (z=70)	0,7%	1,5%	7,4%	10,2%
Tasa de pobreza (z=165)	4,5%	12,3%	32,1%	39,1%

Fuente: Elaboración propia en base a EPH-INDEC

⁸ El *ipcf-ae* medio corresponde al promedio muestral sin ponderaciones de la EPH de cada año, y las tasas de indigencia y pobreza están calculadas siguiendo la ecuación (2) para las líneas de pobreza.

Por su parte, la Tabla 3 indica la estimación de la elasticidad-crecimiento de la tasa de indigencia y de pobreza por los distintos métodos no paramétricos, presentada gráficamente en las Figuras de 5.1 a 5.4.

Tabla 3: Elasticidad crecimiento de la tasa de indigencia y de la tasa de pobreza sobre el ipcf-ae
En puntos porcentuales
Gran Buenos Aires. 1974, 1986, 1998 y 2005

Método	$\eta_H (z=70)$				$\eta_H (z=165)$			
	1974	1986	1998	2005	1974	1986	1998	2005
Kernel	1,677	2,735	1,740	1,735	3,042	2,417	1,463	1,171
K-ésimo vecino cercano	2,168	2,803	1,748	1,724	3,109	2,308	1,468	1,176
Kernel adaptativo	1,314	2,292	1,760	1,749	2,919	2,321	1,476	1,166
Logspline	1,355	2,963	1,734	1,591	3,187	2,279	1,503	1,134
Promedio	1,628	2,698	1,745	1,700	3,064	2,331	1,477	1,162

Fuente: Elaboración propia en base a EPH-INDEC

A excepción de 1974, la elasticidad-crecimiento de la pobreza es mayor para la línea de pobreza de \$70 que para la de \$165; cabe recordar la subestimación de η_H por los cuatro métodos para las distribuciones teóricas con parámetros de 1974 presentada en la sección anterior. Con la misma excepción, η_H disminuye notablemente a medida que aumentan la indigencia y la pobreza. (Para la línea de pobreza de \$165, en 1974 el crecimiento económico en un 1% es capaz de reducir la pobreza en 3 puntos porcentuales, en tanto que en 2005 sólo lo hace en 1,1%.) Sin embargo, no hay que perder de vista que este efecto del crecimiento sobre la pobreza se realiza sobre niveles del indicador de pobreza cada vez más altos, como los mostrados en la Tabla 2.

En las Figuras 5.1 a 5.4 se observa un notorio cambio en la pendiente de la curva de η_H , lo que está asociado al valor de $F(z)$ sobre la que se calcula la elasticidad-crecimiento de la pobreza. Adicionalmente la estimación no paramétrica de la elasticidad-crecimiento de la pobreza presenta algunos máximos locales, que se deben a la concentración de la población en dichos niveles del ipcf-ae. Si la línea de pobreza estuviera casualmente ubicada en dichos máximos, un pequeño crecimiento de la economía generaría una importante caída porcentual en la tasa de pobreza para dicho período.

4- Conclusiones

En este artículo se evaluaron los métodos no paramétricos de kernel, k-ésimo vecino cercano y logspline en cuanto a su bondad en el ajuste para estimar puntualmente la función de densidad de la distribución del ingreso y la elasticidad-crecimiento de la pobreza y se los comparó con algunas aproximaciones paramétricas. Se pudo comprobar que si bien los estimadores no paramétricos generan mejores estimaciones de la densidad que los paramétricos en la gran mayoría de los casos, esta supremacía no se refleja en los casos en que la tasa de pobreza es chica. Para tasas de pobreza superiores a 10% estos métodos dieron una aproximación de dicha elasticidad con un error menor a 0,25 en promedio, que se redujo notablemente a medida que se la calculó para tasas mayores. También se encontró que, si bien el método de logspline generó el menor *MISE* en la estimación puntual de la función de densidad en prácticamente todos los casos estudiados, su ventaja no ser trasladó necesariamente a la aproximación de η_H .

Todos los estimadores no paramétricos mejoraron en su estimación puntual de la función de densidad y de la elasticidad de la pobreza a medida que se las calculó

para líneas de pobreza mayores, o bien para tasas de pobreza más grandes. Sin embargo, esta mejora se produjo con una mayor dispersión para la primera y menor para la última.

En el caso particular de GBA, se observó que la sensibilidad de la tasa de pobreza al crecimiento económico en la línea de pobreza oficial se redujo al pasar de 1974 a 2005. Aunque cabe destacar que la elasticidad-crecimiento de la pobreza se la calculó año a año para tasas de pobreza cada vez mayores, lo que indica que un aumento del 1% del ingreso medio hubiera implicado una caída en niveles de H cada vez mayor. Los métodos no paramétricos permitieron mostrar que hay ciertos valores del ipcf-ae para los que la elasticidad es mayor que los de su entorno, debido a que son ingresos en los que hay mayor concentración local de la población.

5- Apéndice

A continuación se demuestra cómo se obtiene la expresión (1) de la elasticidad-crecimiento de la tasa de pobreza. Como se dijo en la Sección 2, a partir de un cambio en el ingreso de todos los individuos de ε por ciento se obtiene una nueva función acumulada $\tilde{F}$ tal que $\tilde{F}(x) = F(x - \varepsilon x)$, o bien, $\tilde{f}(x) = f(x - \varepsilon x)(1 - \varepsilon)$. Esto determina que el índice H tome el nuevo valor:

$$\begin{aligned}\tilde{H} = \tilde{F}(z) &= \int_0^z f(x)dx \\ &= \int_0^z f(x - \varepsilon x)(1 - \varepsilon)dx\end{aligned}$$

Definiendo $u := (1 - \varepsilon)x$, de modo que $du = (1 - \varepsilon)dx$, y haciendo el cambio de variable correspondiente se tiene que:

$$\begin{aligned}\tilde{F}(z) &= \int_0^{z - \varepsilon z} f(u)du \\ &= \int_0^z f(u)du - \int_{\varepsilon z}^z f(u)du \\ &= F(z) - \varepsilon z f(z)\end{aligned}$$

Por tanto, la diferencia $\tilde{F}(z) - F(z)$ es igual a $-\varepsilon z f(z)$. Por su parte, el nuevo ingreso medio $\bar{\tilde{x}}$ es:

$$\bar{\tilde{x}} = \int_0^\infty x \tilde{f}(x)dx = \int_0^\infty x f(x - \varepsilon x)(1 - \varepsilon)dx$$

Realizando el mismo cambio de variable:

$$\bar{\tilde{x}} = \frac{1}{1 + \varepsilon} \int_0^\infty u f(u)du$$

Y finalmente, aproximando por Taylor la razón $1/(1 - \varepsilon)$, se tiene que:

$$\begin{aligned}\bar{\tilde{x}} &= (1 + \varepsilon) \int_0^\infty u f(u)du \\ &= \bar{x} + \varepsilon \bar{x}\end{aligned}$$

de modo que $\bar{\tilde{x}} - \bar{x} = \varepsilon \bar{x}$. Reemplazando ambos resultados en la definición de elasticidad-crecimiento de la pobreza:

$$\eta_H = \lim_{\varepsilon \rightarrow 0} \frac{(\tilde{F}(z) - F(z)) / F(z)}{(\bar{\tilde{x}} - \bar{x}) / \bar{x}}$$

se llega a la expresión:

$$\eta_H = -\frac{zf'(z)}{H}$$

Bibliografía

- [1] Botargues P. y Petrecolli (1999) D. "Estimaciones paramétricas y no paramétricas de la distribución del ingreso de los ocupados del Gran Buenos Aires, 1992-1997" Económica Año XLV, Nro.1 - Enero/Junio
- [2] Bourguignon F. (2003) "The growth elasticity of poverty reduction: explaining heterogeneity across countries and time periods" Inequality and Growth: Theory and Policy Implications. Cambridge, MA: MIT Press
- [3] Bourguignon F. (2004) "The poverty-growth-inequality triangle" Indian Council for Research on International Economic Relations, New Delhi. Working Papers 125
- [4] Datt G. y Ravallion M. Growth and redistribution components of changes in poverty measures. A decomposition with applications to Brazil and India in the 1980s
- [5] Datt G. y Ravallion M. (2002) Why has economic growth been more pro-poor in some states of India than others? Journal of Development Economics 68.
- [6] Deaton A. (1997) The Analysis of household surveys Microeconomic Approach to Development Policy Baltimore: Johns Hopkins University Press
- [7] Dollar, D. y Kraay, A. (2002) Growth is good for the poor Journal of Economic Growth, vol. 7(3)
- [8] Chakravarty y Majumder (1990) "Distribution of personal income: development of a new model and its application to U.S. income data" Journal of Applied Econometrics Vol.5 No.2
- [9] Chen, S. y Ravallion, M. (2001) "Measuring pro-poor growth". Policy Research Working Paper 2666
- [10] Grund B., Hall P. y Marron J.S. (1994) "Loss and risk in smoothing parameter selection". Nonparametric Statistics. Vol.4
- Journal of the American Statistical Association. Marzo de , Vol.91, No.433.
- [11] Hall P. y Marron J.C. (1987) "On the amount of noise inherent in bandwidth selection for a kernel density estimation" The Annals of Statistics. Vol.15, No.1
- [12] Janvry, A. y Sadoulet, E. (2000) "Growth, poverty, and inequality in Latin America: a causal analysis", 1970-74 Review of Income and Wealth 46 (3)
- [13] Jones M.C. (1990) "The role of ISE and MISE in the density estimation". Statistical & Probability Letters.
- [14] Jones M.C. (1993) "Simple boundary correction for kernel density estimation" Statistics and Computing 3.
- [15] Jones M.C., Marron J.C. y Sheater S.J. (1996) "A brief survey of bandwidth selection for density estimation" Journal of American Statistics Association. 91
- [16] Kakwani N. (1993) "Poverty and economic growth with application to Cote D'Ivoire" Review of Income and Wealth No.2
- [17] Kakwani N. (1993) "Statistical inference in the measurement of poverty" Econometrica. Vol.48 No.2
- [18] Kooperberg C. y Stone C.J. (1991) "A study of logspline density estimation" Computational Statistics & Data Analysis 12
- [19] Lopez y Servén (2005) "A normal relationship?" World Bank WPS3814
- [20] McDonald J.B. y Mantrala A. (1995) "Distribution of personal income: revisited" Journal of Applied Econometrics Vol.10 No.2

- [21] Pagan A. y Ullah A. (1999) Nonparametric econometrics. Themes in modern econometrics, Cambridge University Press, New York, USA.
- [22] Ravallion, M. (2003) "The debate on globalization, poverty and inequality: why measurement matters". World Bank Policy Research Working Paper 3038
- [23] Salem y Mount (1974) "A convenient descriptive model of income distribution: The Gamma density" Econometrica Vol.42 No.6
- [24] Schwartz (1978) "Estimating the Dimension of a Model". Annals of Statistics. Vol.6 No.2
- [25] Schucany W.R. y Sommers J.P. (1977) Improvement of kernel type density estimators. Journal of the American Statistical Association 66
- [26] Silverman B.W. (1986) Density Estimation for Statistics and Data Analysis Monographs on Statistics and Applied Probability, Chapman and Hall, London, UK

Tabla 4: *MISE* de las estimaciones de la función de densidad sobre 30 muestras de 4000 observaciones c/u por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline de las distribuciones lognormal, gamma y dagum para los valores de sus parámetros de la Tabla 1. La tabla incluye el *MISE* de las estimaciones por los modelos lognormal, gamma y dagum empleando máxima verosimilitud. El *MISE* está calculado en los intervalos [0,200], [0,100] y [100,200]⁹

Año	Método	[0,200]			[0,100]			[100,200]		
		Lognormal	Gamma	Dagum	Lognormal	Gamma	Dagum	Lognormal	Gamma	Dagum
1974	Kernel	0,370	0,052	0,015	0,265	0,042	0,010	0,104	0,010	0,005
	K-ésimo vecino cercano	1,951	0,213	0,032	1,891	0,192	0,028	0,061	0,021	0,004
	Kernel adaptativo	0,468	0,416	0,347	0,128	0,127	0,086	0,339	0,289	0,261
	Logspline	0,006	0,010	0,013	0,003	0,004	0,005	0,003	0,006	0,008
	Lognormal	0,000	3,403	4,188	0,000	1,541	0,980	0,000	1,862	3,207
	Gamma	5,202	0,007	2,722	3,407	0,002	0,607	1,796	0,006	2,116
	Dagum	1,675	1,475	0,009	0,276	0,582	0,004	1,399	0,893	0,005
1986	Kernel	0,751	0,267	0,261	0,332	0,112	0,109	0,419	0,155	0,152
	K-ésimo vecino cercano	1,888	0,418	0,319	1,846	0,410	0,293	0,042	0,008	0,026
	Kernel adaptativo	0,614	0,957	0,361	0,440	0,927	0,232	0,175	0,030	0,129
	Logspline	0,016	0,021	0,076	0,008	0,019	0,016	0,008	0,002	0,060
	Lognormal	0,001	17,415	4,995	0,000	4,055	1,026	0,001	13,360	3,969
	Gamma	19,176	0,001	20,730	8,298	0,001	13,725	10,878	0,001	7,005
	Dagum	3,605	4,406	0,003	0,696	1,295	0,001	2,909	3,111	0,002
1998	Kernel	1,294	0,825	0,564	1,083	0,756	0,221	0,212	0,069	0,343
	K-ésimo vecino cercano	2,469	1,049	0,485	2,296	0,999	0,438	0,172	0,049	0,047
	Kernel adaptativo	1,259	1,960	1,002	1,186	1,942	0,986	0,074	0,017	0,016
	Logspline	0,146	0,514	0,219	0,097	0,472	0,081	0,049	0,042	0,137
	Lognormal	0,005	47,780	20,618	0,001	29,641	15,713	0,004	18,139	4,904
	Gamma	71,659	0,005	52,102	36,587	0,005	18,625	35,072	0,000	33,477
	Dagum	7,849	7,186	0,009	5,210	5,181	0,009	2,639	2,005	0,000
2005	Kernel	1,952	1,344	0,588	1,898	1,269	0,305	0,054	0,076	0,284
	K-ésimo vecino cercano	2,646	0,998	0,661	2,560	0,985	0,557	0,086	0,013	0,104
	Kernel adaptativo	2,094	1,970	1,492	2,045	1,961	1,449	0,049	0,010	0,042
	Logspline	0,257	1,028	0,067	0,134	1,018	0,065	0,123	0,009	0,003
	Lognormal	0,012	53,110	29,742	0,011	40,639	23,592	0,001	12,472	6,150
	Gamma	86,096	0,017	65,233	52,940	0,006	22,196	33,156	0,012	43,037
	Dagum	11,328	9,628	0,006	7,259	7,739	0,004	4,070	1,889	0,003

Fuente: Elaboración propia

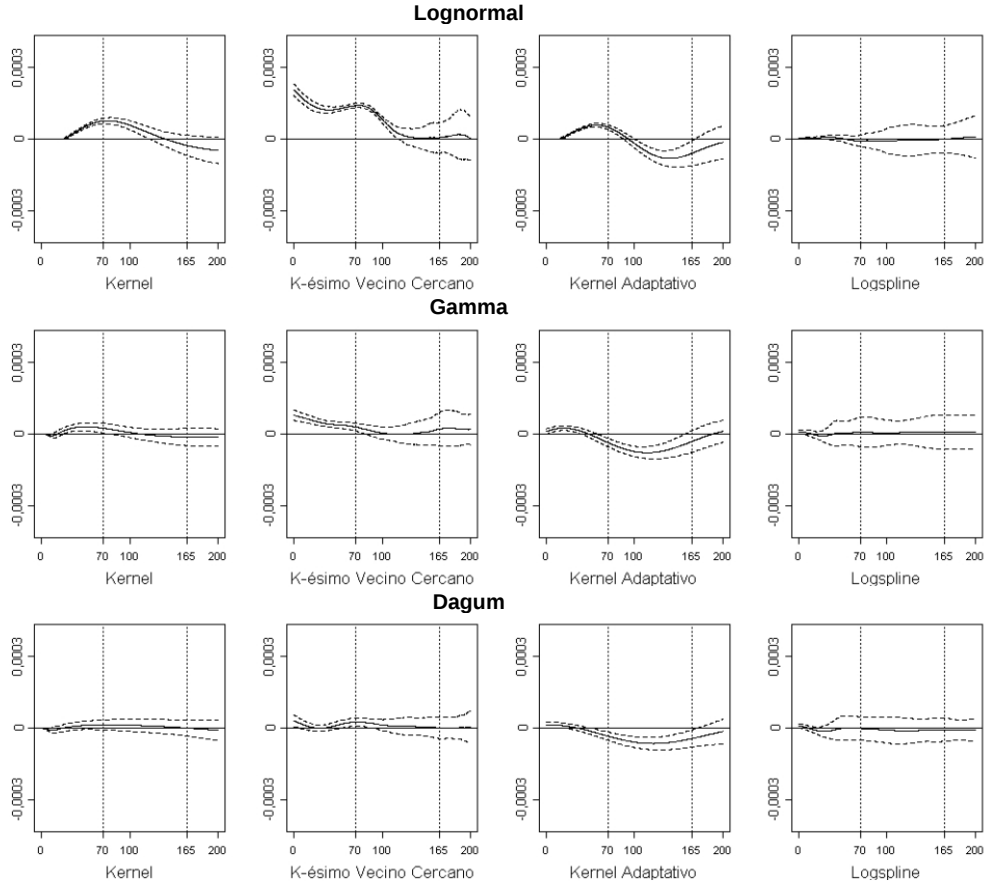
⁹ Los valores de los *MISE* presentados en esta tabla han sido multiplicados por un millón a fines de facilitar la su presentación.

Tabla 5: *MISE* de las estimaciones de la elasticidad-crecimiento de la tasa de pobreza sobre 30 muestras de 4000 observaciones c/u por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline de las distribuciones lognormal, gamma y dagum para los valores de sus parámetros de la Tabla 1. La tabla incluye el *MISE* de las estimaciones de la elasticidad-crecimiento de la tasa de pobreza por los modelos lognormal, gamma y dagum empleando máxima verosimilitud. El *MISE* está calculado en los intervalos [50,125] y [125,200]

Año	Método	[50,200]			[50,125]			[125,200]		
		Lognormal	Gamma	Dagum	Lognormal	Gamma	Dagum	Lognormal	Gamma	Dagum
1974	Kernel	402,647	63,439	36,081	381,047	58,395	32,589	21,600	5,044	3,493
	K-ésimo vecino cercano	187,410	59,627	32,212	169,366	56,169	28,902	18,045	3,457	3,310
	Kernel adaptativo	550,386	109,515	79,414	510,349	100,389	69,170	40,037	9,126	10,244
	Logspline	779,814	71,747	47,572	759,236	68,233	42,939	20,578	3,513	4,634
	Lognormal	0,003	397,180	526,061	0,002	333,346	454,628	0,001	63,834	71,433
	Gamma	563,188	0,011	3,796	456,977	0,006	0,387	106,212	0,005	3,409
	Dagum	304,558	46,207	0,218	280,261	34,154	0,122	24,297	12,053	0,096
1986	Kernel	86,373	8,991	26,492	84,191	8,585	24,467	2,182	0,406	2,026
	K-ésimo vecino cercano	71,562	5,983	23,333	70,996	5,889	22,300	0,565	0,094	1,033
	Kernel adaptativo	145,748	14,079	55,078	144,939	14,001	53,683	0,808	0,078	1,395
	Logspline	117,487	5,703	37,363	116,798	5,628	36,765	0,689	0,075	0,597
	Lognormal	0,006	95,877	50,955	0,004	86,450	49,718	0,001	9,428	1,237
	Gamma	222,891	0,003	106,401	191,142	0,002	60,616	31,750	0,001	45,785
	Dagum	71,003	10,366	0,002	69,906	8,809	0,001	1,097	1,557	0,001
1998	Kernel	4,517	0,316	2,280	4,504	0,303	2,182	0,013	0,013	0,098
	K-ésimo vecino cercano	2,154	0,052	1,307	2,109	0,042	1,292	0,046	0,010	0,015
	Kernel adaptativo	4,123	0,077	2,224	4,100	0,072	2,213	0,023	0,005	0,011
	Logspline	1,744	0,169	1,465	1,723	0,158	1,430	0,020	0,011	0,035
	Lognormal	0,001	13,442	9,456	0,001	13,178	7,280	0,000	0,265	2,176
	Gamma	41,657	0,001	25,141	39,179	0,001	16,668	2,478	0,000	8,473
	Dagum	3,693	1,543	0,008	2,578	1,421	0,006	1,115	0,122	0,002
2005	Kernel	1,771	0,217	0,709	1,771	0,215	0,680	0,001	0,002	0,029
	K-ésimo vecino cercano	0,751	0,035	0,273	0,740	0,030	0,256	0,011	0,005	0,017
	Kernel adaptativo	1,181	0,032	0,396	1,178	0,032	0,388	0,003	0,000	0,007
	Logspline	0,708	0,170	0,326	0,699	0,163	0,326	0,008	0,008	0,000
	Lognormal	0,004	8,114	6,313	0,003	7,591	3,900	0,001	0,523	2,414
	Gamma	25,790	0,002	17,605	25,070	0,002	12,627	0,720	0,000	4,977
	Dagum	2,310	1,334	0,003	1,331	1,045	0,003	0,978	0,289	0,000

Fuente: Elaboración propia

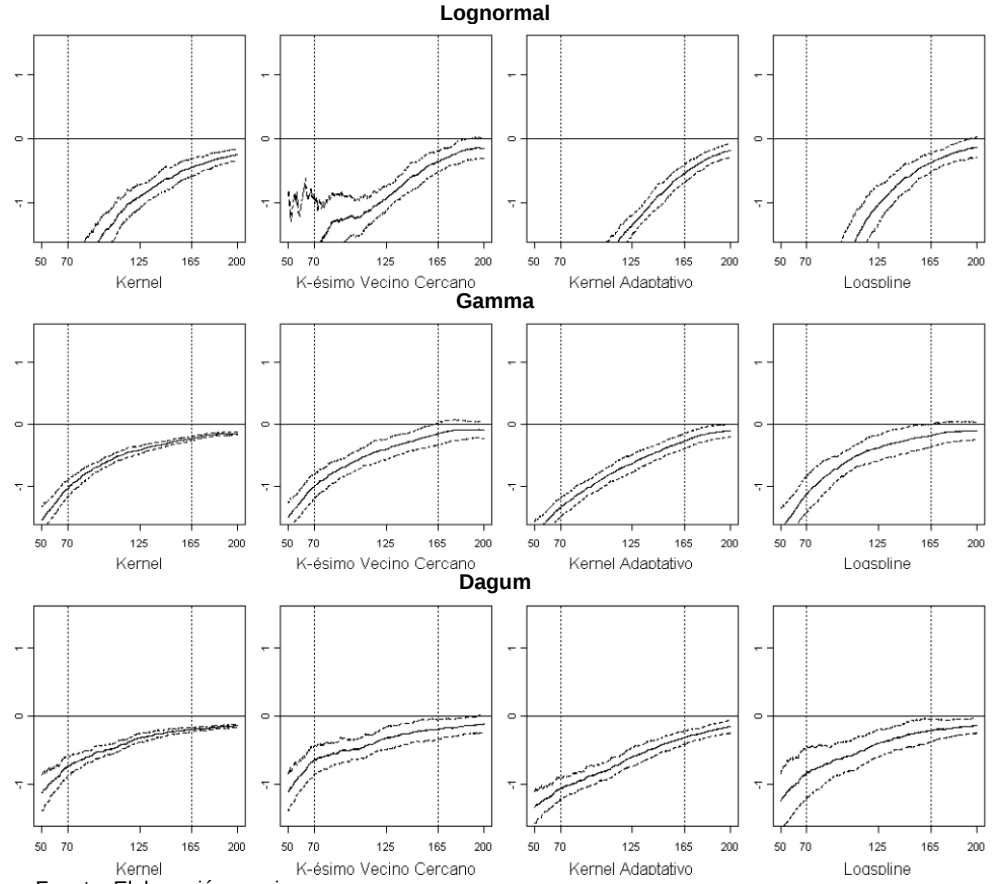
Figura 1.1: Diferencia promedio de las estimaciones de la función de densidad sobre 30 muestras de 4000 observaciones c/u por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline y de las distribuciones modelos lognormal, gamma y dagum $\pm$ un desvío estándar para los valores de sus parámetros en la Tabla 1. correspondientes al ipcf-ae de 1974¹⁰



Fuente: Elaboración propia

¹⁰ La línea sólida representa a $f_i(x) - \frac{1}{30} \sum_{j=1}^{30} \hat{f}_j$ con $i=LN,G,D$ y $j=1,2,3,4$. En tanto que las líneas punteadas indican una desviación estándar de la estimación de f en x .

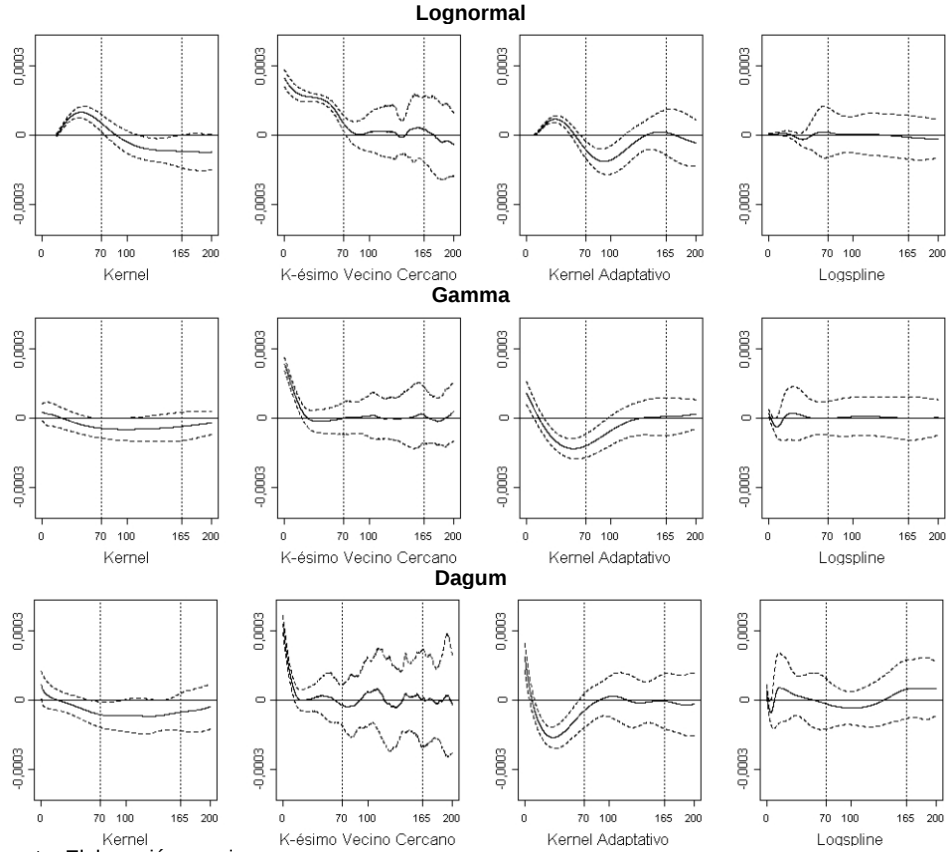
Figura 1.2: Diferencia promedio de las estimaciones de la elasticidad-crecimiento de la pobreza sobre 30 muestras de 4000 observaciones c/u por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline y de las distribuciones modelos lognormal, gamma y dagum $\pm$ un desvío estándar para los valores de sus parámetros en la Tabla 1. correspondientes al ipcf-ae de 1974¹¹



Fuente: Elaboración propia

¹¹ La línea sólida representa a $\eta_i(x) - \frac{1}{30} \sum_{j=1}^{30} \hat{\eta}_j(x)$ con $i=LN,G,D$ y $j=1,2,3,4$. Por su parte, las líneas punteadas indican una desviación estándar de la estimación de η_H en x .

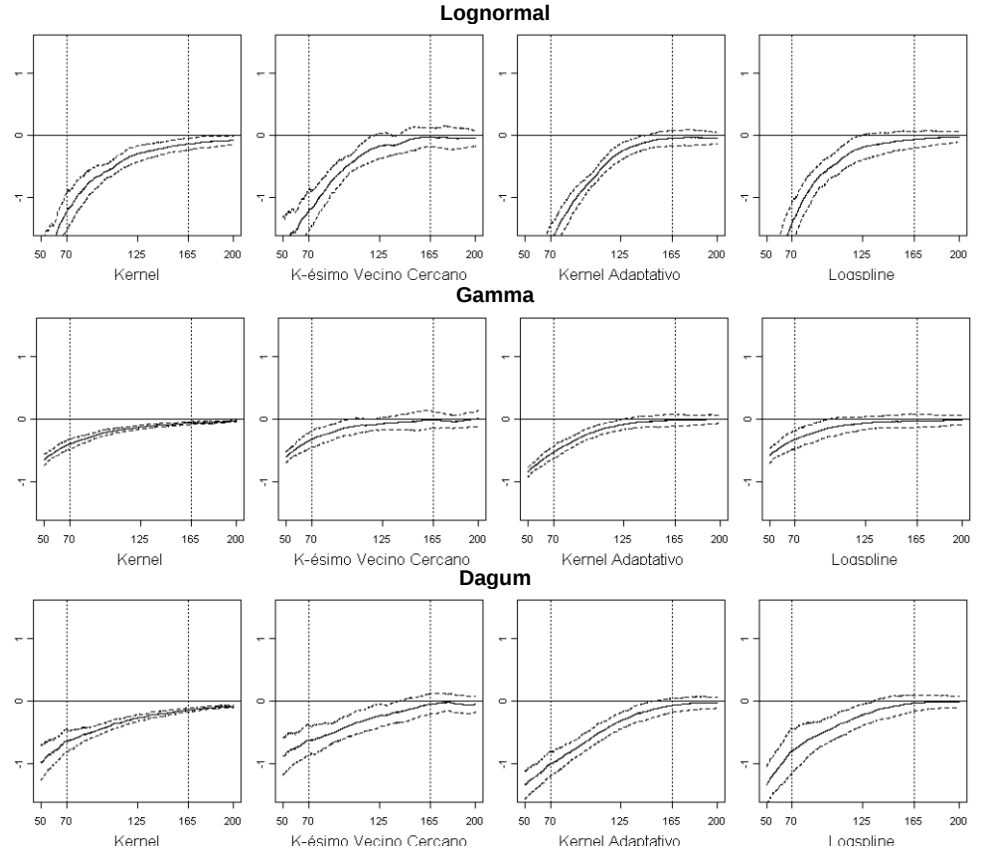
Figura 2.1: Diferencia promedio de las estimaciones de la función de densidad sobre 30 muestras de 4000 observaciones c/u por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline y de las distribuciones modelos lognormal, gamma y dagum $\pm$ un desvío estándar para los valores de sus parámetros en la Tabla 1. correspondientes al ipcf-ae de 1986¹²



Fuente: Elaboración propia

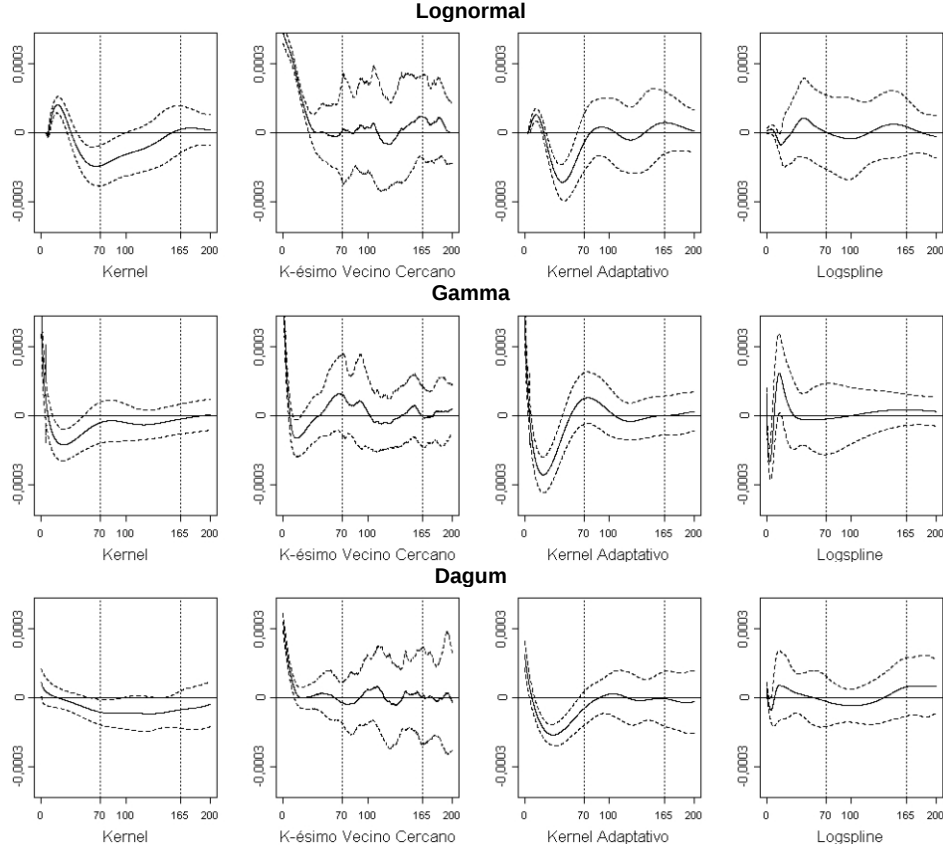
¹² La línea sólida representa a $f_i(x) - \frac{1}{30} \sum_{j=1}^{30} \hat{f}_j$ con $i=LN,G,D$ y $j=1,2,3,4$. En tanto que las líneas punteadas indican una desviación estándar de la estimación de f en x .

Figura 2.2: Diferencia promedio de las estimaciones de la elasticidad-crecimiento de la pobreza sobre 30 muestras de 4000 observaciones c/u por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline y de las distribuciones modelos lognormal, gamma y dagum $\pm$ un desvío estándar para los valores de sus parámetros en la Tabla 1. correspondientes al ipcf-ae de 1986¹³



¹³ La línea sólida representa a $\eta_i(x) - \frac{1}{30} \sum_{j=1}^{30} \hat{\eta}_j(x)$ con $i=LN,G,D$ y $j=1,2,3,4$. Por su parte, las líneas punteadas indican una desviación estándar de la estimación de η_H en x .

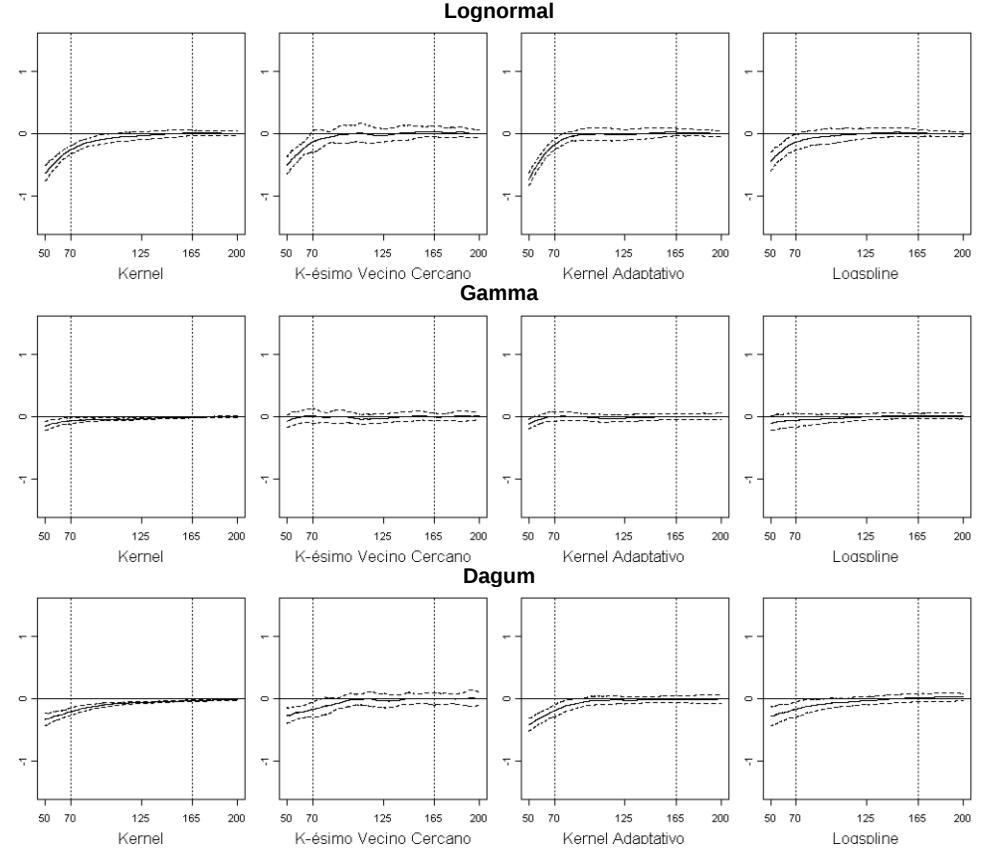
Figura 3.1: Diferencia promedio de las estimaciones de la función de densidad sobre 30 muestras de 4000 observaciones c/u por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline y de las distribuciones modelos lognormal, gamma y dagum $\pm$ un desvío estándar para los valores de sus parámetros en la Tabla 1. correspondientes al ipcf-ae de 1998¹⁴



Fuente: Elaboración propia

¹⁴ La línea sólida representa a $f_i(x) - \frac{1}{30} \sum_{j=1}^{30} \hat{f}_j$ con $i=LN,G,D$ y $j=1,2,3,4$. En tanto que las líneas punteadas indican una desviación estándar de la estimación de f en x .

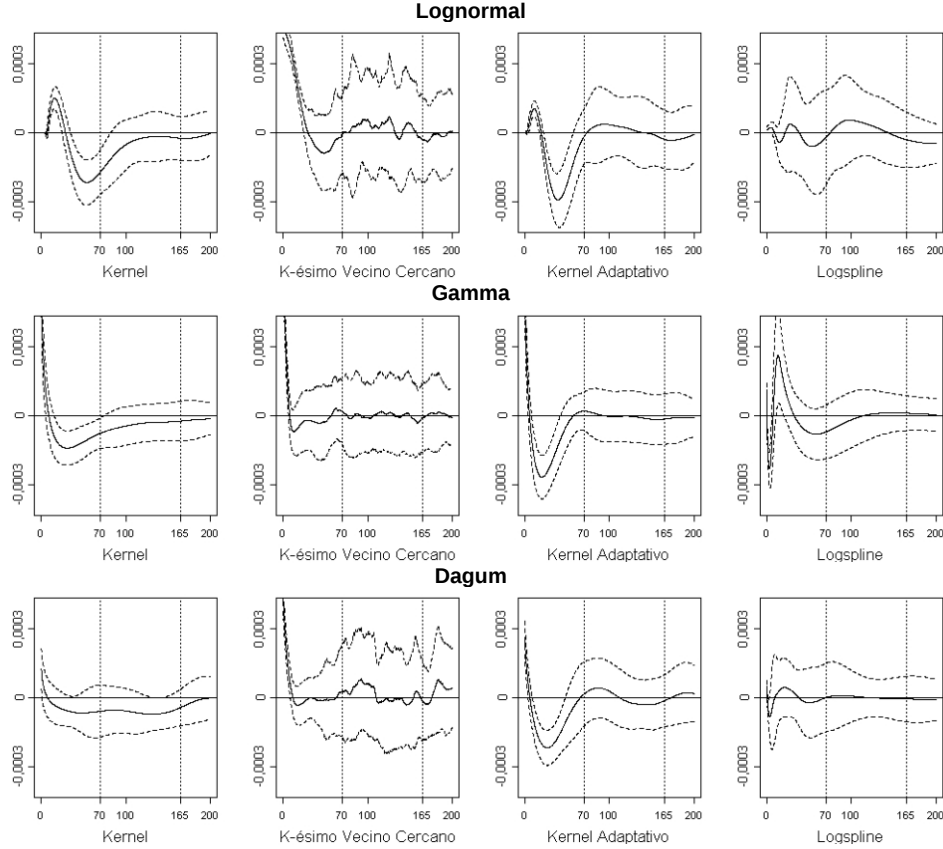
Figura 3.2: Diferencia promedio de las estimaciones de la elasticidad-crecimiento de la pobreza sobre 30 muestras de 4000 observaciones c/u por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline y de las distribuciones modelos lognormal, gamma y dagum $\pm$ un desvío estándar para los valores de sus parámetros en la Tabla 1. correspondientes al ipcf-ae de 1998¹⁵



Fuente: Elaboración propia

¹⁵ La línea sólida representa a $\eta_i(x) - \frac{1}{30} \sum_{j=1}^{30} \hat{\eta}_j(x)$ con $i=LN,G,D$ y $j=1,2,3,4$. Por su parte, las líneas punteadas indican una desviación estándar de la estimación de η_H en x .

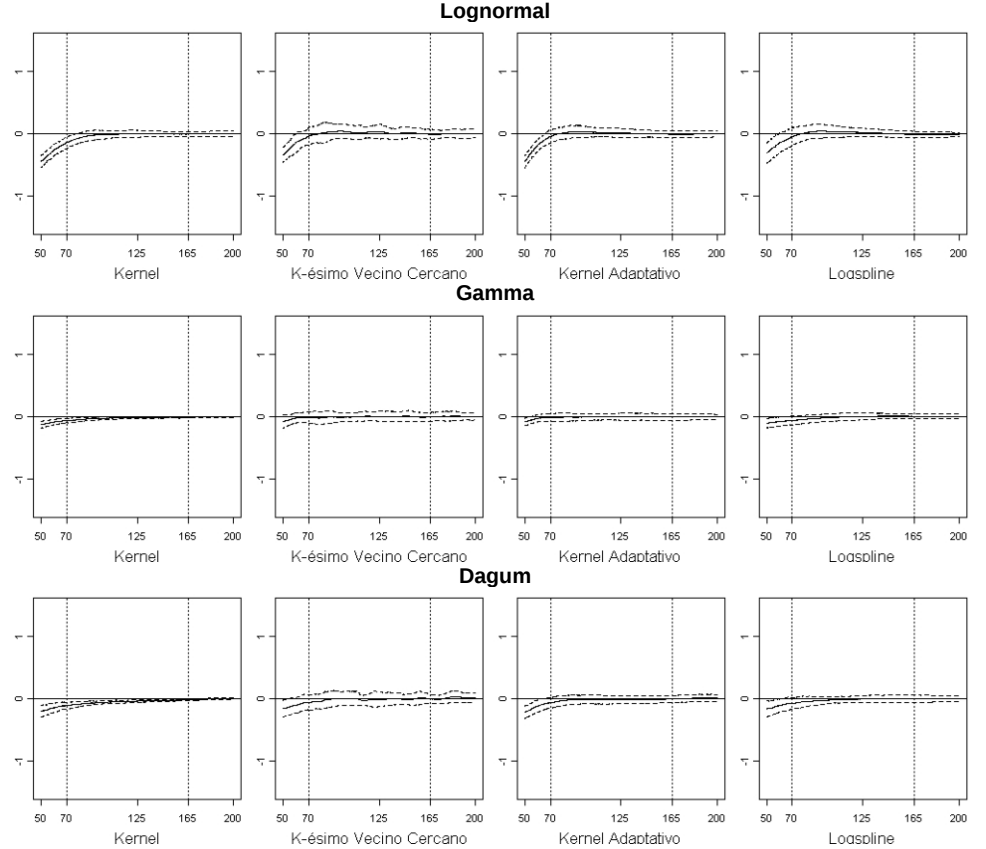
Figura 4.1: Diferencia promedio de las estimaciones de la función de densidad sobre 30 muestras de 4000 observaciones c/u por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline y de las distribuciones modelos lognormal, gamma y dagum $\pm$ un desvío estándar para los valores de sus parámetros en la Tabla 1. correspondientes al ipcf-ae de 2005¹⁶



Fuente: Elaboración propia

¹⁶ La línea sólida representa a $f_i(x) - \frac{1}{30} \sum_{j=1}^{30} \hat{f}_j$ con $i=LN,G,D$ y $j=1,2,3,4$. En tanto que las líneas punteadas indican una desviación estándar de la estimación de f en x .

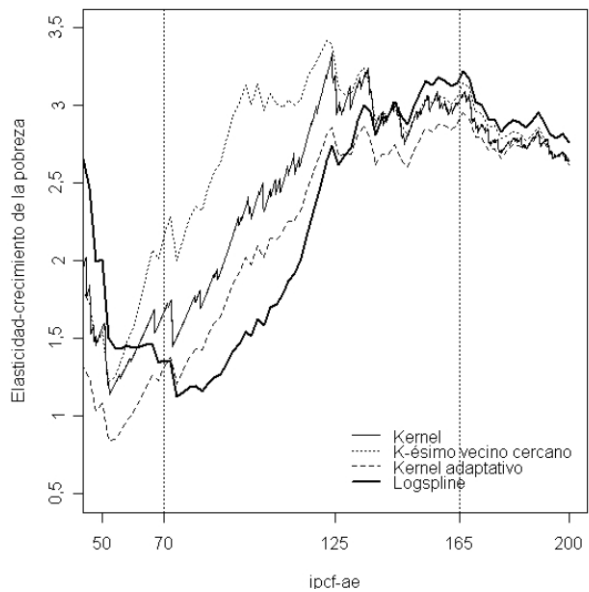
Figura 4.2: Diferencia promedio de las estimaciones de la elasticidad-crecimiento de la pobreza sobre 30 muestras de 4000 observaciones c/u por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline y de las distribuciones modelos lognormal, gamma y dagum $\pm$ un desvío estándar para los valores de sus parámetros en la Tabla 1. correspondientes al ipcf-ae de 2005¹⁷



Fuente: Elaboración propia

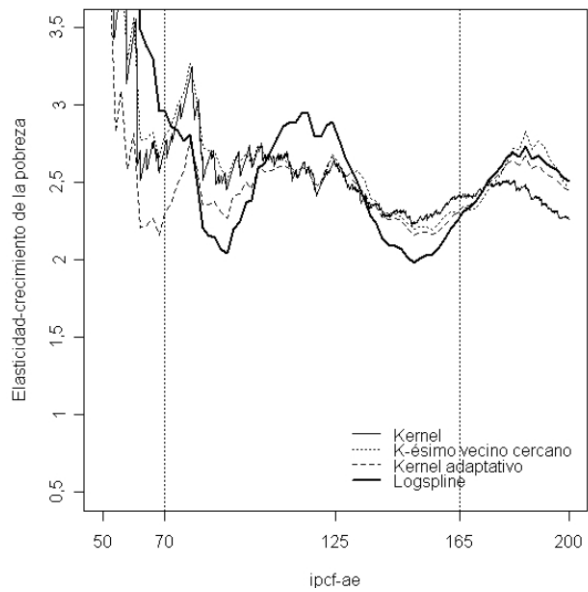
¹⁷ La línea sólida representa a $\eta_i(x) - \frac{1}{30} \sum_{j=1}^{30} \hat{\eta}_j(x)$ con $i=LN,G,D$ y $j=1,2,3,4$. Por su parte, las líneas punteadas indican una desviación estándar de la estimación de η_H en x .

**Figura 5.1: Elasticidad-crecimiento de la pobreza estimada por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline
Ipcf-ae en pesos (a precios de 1999)
Gran Buenos Aires, 1974**



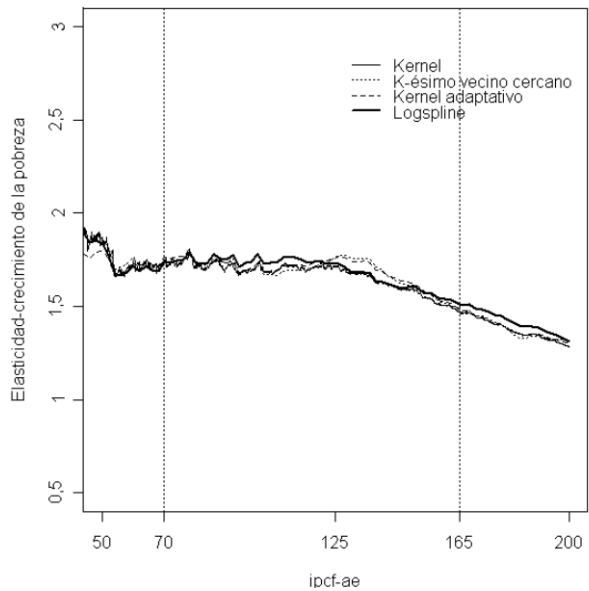
Fuente: Elaboración propia en base a EPH-INDEC

**Figura 5.2: Elasticidad-crecimiento de la pobreza estimada por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline
Ipcf-ae en pesos (a precios de 1999)
Gran Buenos Aires, 1986**



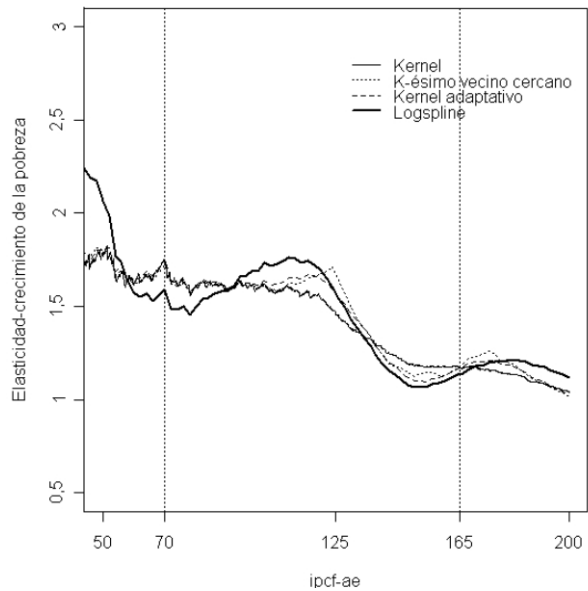
Fuente: Elaboración propia en base a EPH-INDEC

**Figura 5.3: Elasticidad-crecimiento de la pobreza estimada por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline
Ipcf-ae en pesos (a precios de 1999)
Gran Buenos Aires, 1998**



Fuente: Elaboración propia en base a EPH-INDEC

**Figura 5.4: Elasticidad-crecimiento de la pobreza estimada por kernel, k-ésimo vecino cercano, kernel adaptativo y logspline
Ipcf-ae en pesos (a precios de 1999)
Gran Buenos Aires, 2005**



Fuente: Elaboración propia en base a EPH-INDEC