

UKJAES

University of Kirkuk Journal
For Administrative
and Economic Science

ISSN:2222-2995 E-ISSN:3079-3521

University of Kirkuk Journal For
Administrative and Economic Science



Ibrahim Omar Salim & Ahmed Rikan A. . Nonparametric Approach in Estimating Longitudinal Data Models. *University of Kirkuk Journal For Administrative and Economic Science* (2025) 15 (4) Part (2):342-350.

Nonparametric Approach in Estimating Longitudinal Data Models

Omar Salim Ibrahim ¹, Rikan A. Ahmed ²

^{1,2} Department of Statistics and Informatics, College of Computer Science and Mathematics-University of Mosul, Mosul, Iraq

omarsalim85@uomosul.edu.iq ¹

rikan.ahmed@uomosul.edu.iq ²

Abstract: We routinely observe longitudinal patterns in medical studies, which track individuals' progress in one or more study variables over time. Most modern modeling techniques assume a parametric dependence of the outcome on time (e.g., linear, quadratic); nonetheless, these dependencies may not be that much dependent on the flow of time. The fixed-effects models are flexible in handling longitudinal data whose pattern is not parametric because the shape of the curve is left to be determined by the data. Although parametric models are one of the most important data analysis tools, they have not been sufficient in some cases.

Therefore, this thesis aims to study methods for estimating nonparametric fixed models for longitudinal data, then compare them to determine which is the best and most efficient, yielding the lowest value for the comparison criteria. These models are then applied to medical data for patients with high blood pressure. To achieve this goal, four methods were employed: Profile Least-Squares, Profile Estimation, Differencing/Transformation Methods, and Marginal Integration. The methods included employing multiple local regression at different degrees and the Gaussian, Cosine, and Epanchnikov kernel functions. Furthermore, a comparison was conducted between the models using five comparison criteria: the mean square error (RMSE), the relative root mean square error (RRSE), the mean relative absolute deviation (MAPE), the mean absolute error (MAE), and the relative absolute error (RAE). The results obtained showed that the Profile estimation method is the best among the methods.

Keywords: longitudinal data, fixed effects model, kernel function, nonparametric estimation methods.

الاسلوب اللامعلمي في تقدير نماذج البيانات الطولية

م.د. عمر سالم ابراهيم ¹، ا.م.د. ريكان عبد العزيز احمد ²

^{1,2} قسم الاحصاء والمعلوماتية، كلية علوم الحاسوب والرياضيات، جامعة الموصل، العراق

المستخلص: نلاحظ الأنماط الطولية بشكل روتيني في الدراسات الطبية، حيث تتابع تطورات الأفراد في متغير واحد أو أكثر من متغيرات الدراسة على مدار فترات زمنية. تفترض العديد من الطرق الحالية للنمذجة وجود علاقة

معلمية بين النتيجة والوقت (مثل: الخطية، والتربيعية)؛ ومع ذلك، قد لا تعكس هذه العلاقات بدقة مع مرور الوقت. توفر نماذج التأثيرات الثابتة دالة مرنة في التعامل مع البيانات الطولية ذات الانماط اللامعلمية، لأنها تسمح للبيانات بتحديد شكل المنحنى. ورغم أن النماذج المعلمية تعد أحد أهم أدوات تحليل البيانات؛ إلا أنها لم تستطع أن تكون كافية في بعض الحالات.

ولهذا تهدف هذه الرسالة الى دراسة طرق تقدير النموذج الثابت اللامعلمي للبيانات الطولية ثم المقارنة بينها لتحديد ايهما الافضل والاكفا الذي يعطي اقل قيمة لمعايير المقارنة ثم تطبيق هذه النماذج على بيانات طبية لاعداد المصابين بضغط الدم ولتحقيق ذلك الهدف تم توظيف أربع طرق هي طريقة

Differencing/Transformation ، طريقة Profile Least-Squares ، طريقة Profile Estimation ، طريقة Marginal Integration وتضمنت الطرق توظيف كل من دوال كيرنل (Gaussian) و (Cosine kernel) ، و (Epanchnikov) وكذلك اجراء مقارنة باستعمال كل من معيار معدل متوسط مربعات الخطأ (RMSE)، الجذر التربيعي النسبي للخطأ (RRSE)، ومعدل متوسط الانحرافات المطلقة النسبية (MAPE)، متوسط الخطأ المطلق (MAE)، الخطأ المطلق النسبي (RAE). وقد تبين من خلال النتائج ان طريقة التقدير Profile هي الافضل من بين الطرق.

الكلمات المفتاحية: البيانات الطولية، أنموذج التأثيرات الثابتة، الدالة اللبية، طرق التقدير اللامعلمية.

Corresponding Author: E-mail: omarsalim85@uomosul.edu.iq

Introduction

Longitudinal techniques have been extensively applied in epidemiology, biomedicine, economics and finance among others. Longitudinal study data sets are generally obtained as repeated measurements of the response variable and independent variables across a group of sectors across a time period. A major objective of statistical analyses is to estimate how the independent variables affect the response variable during the study period.

Parametric methods have especially been popular with applied econometricians, and are the most adopted methods even though Longitudinal data models have been particularly popular. When these assumptions are not consistent with the process of data generation (DGP), then the estimators of these assumptions themselves are biased and poorer than the consistent ones. When practitioners apply the diagnostic tests to their parametric models one would easily find Their models get rejected by the data. Thus, they are also able to demand more elastic nonparametric substitutes.

The literature on constructing longitudinal data models has been extensive over the past 50 years or so. An overview of this literature is provided by referring readers to (Arellano, 2003), (Hsiao, 2003), and (Baltagi, 2008). However, the literature devotes little attention to examining underspecified parametric longitudinal data models. The estimation of an underspecified model is often inconsistent, and thus the subsequent statistical inference is invalid. This is why we may also be witnessing a rapid expansion in the literature on nonparametric longitudinal data models over the past 15 years (Parmeter, 2019).

Nonparametric models can give longitudinal data, including the nonparametric fixed effect model and the nonparametric random effect model. Issues of autocorrelation of these nonparametric models and issues of heteroscedasticity of these models are also present. These problems can be dealt with, and the model then estimated and analyzed.

But it is a fact that the preset parametric assumption on the Longitudinal data models may be too strong and the parametric models may be misspecified under certain applied circumstances. This misspecification can result in inconsistent estimates and therefore wrong conclusions can be made. In response to these problems, In recent years, there have been suggestive more adaptable nonlinear modeling models and non-parametric estimation techniques, which allow the Longitudinal data to talk themselves. In this paper, we take special note of two modelling frameworks: nonparametric Longitudinal data models. (Su, 2011).

This research will attempt to research the nonparametric models of longitudinal data and later compare them to identify which is the most suitable and effective of the two models which provides the lowest value to the criteria of comparison and the comparison of some nonparametric estimators by making use of these models on real data within the health area.

1st: Selection Method of Data Panel

In order to take the best model, we can conduct a number of tests, including:

(1) Hausman Test test is a statistical test used to establish which of Fixed Effect or Random Effect is the best model to use. A Result: H0: Choose RE (p (0.05)) H1: Choose FE (p(0.05)) H1: Choose RE (p(0.05))

(2) **Test Lagrange Multiplier** Lagrange multiplier test (LM) A test that depends on a statistic with a single degree of freedom distribution to differentiate between the ensemble model and the fixed and random models. This test is known as the Lagrange multiplier test (LM), and it depends on errors estimated using the ordinary least squares method to test the following hypothesis. Result: H0: Select CE (p> 0.05) H1: Select RE (p <0.05). (Honda, 1985), (Baltagi, 2008)

2nd: Longitudinal Data Models

1- Nonparametric Longitudinal Data Models with Random Effects

Here we are dealing with nonparametric Longitudinal data models of random effects:

$$y_{it} = m(x_{it}) + \alpha_i + u_{it}, i = 1, \dots, n, t = 1, \dots, T, \quad (1)$$

and x_{it} is $p \times 1$ vector variables, $0, \sigma_\alpha^2, u_{jt}$ are (i.i.d.) $(0, \sigma_\alpha^2), u_{jt}$, and $i, j = 1, \dots, n$ and $t = 1, \dots, T$. We observe that some these assumptions are weak and we can also test specification.

Allow $\varepsilon_{it} = \alpha_i + u_{it}, \varepsilon_i = (\varepsilon_{i1}, \dots, \varepsilon_{iT})'$ and $\varepsilon_i = (\varepsilon_1, \dots, \varepsilon_n)'$. Then $\Sigma \equiv E(\varepsilon_i \varepsilon_i') = \sigma_\alpha^2 I_T + \sigma_u^2 l_T l_T'$ and $\Omega \equiv E(\varepsilon \varepsilon') = I_n \otimes \Sigma$. We examine the local least-squares (LLSS) approximation of m and first order derivatives, ignoring the variance-covariance matrix. We then, with that matrix, obtain the most effective estimates of m and derivatives of m .

2- Nonparametric Longitudinal data Fixed effect model

Here we assume the following nonparametric fixed effects Longitudinal model of data.

$$y_{it} = m(x_{it}) + \alpha_i + u_{it}, i = 1, \dots, n, t = 1, \dots, T, \quad (2)$$

with covariate (regressor) x_{it} of dimension $p \times 1, m(\cdot)$ unknown smooth, α_i 's fixed effects heterogeneity, u_{it} is i.i.d. zero mean, finite variance σ_u^2 and x_{jt} it are independent with all i, j and t . We suppose $\sum_{i=1}^n \alpha_i = 0$ (so that $\alpha_1 = -\sum_{i=2}^n \alpha_i$) to make an identification. Moreover, x it is purely exogenous, to be simple. We can approximate the model in (3)

$$Y \approx X(x)\delta(x) + D\alpha + U \quad (3)$$

where $\alpha = (\alpha_2, \dots, \alpha_n)'$, $D = (I_n \otimes l_T) d_n, d_n = [-l_{n-1} l_{n-1}]'$, and other notation as above. It is important to note that α has parameters of heterogeneity. (Rodriguez-Poo, 2017)

3rd: Method estimate Nonparametric Longitudinal Data Model with Fixed Effects

1- Profile Least-Squares

profile least squares estimation of the model in Equation (4). Assume that the data is ordered, t being the rapid index. Subsequently, the method profile least squares of begins by assuming. that α_i is known, locally-linearly, $M(z)$ is the estimate of $M(z) = (m(z), h \odot \dot{m}(z)')'$ (i.e. the estimate of $M(z)$ by estimating (Su, 2006).

$$M_\alpha(x) = \arg \min_{M \in \mathbb{R}^{q+1}} (Y - D\alpha - D_x M)' \mathcal{K}_x (Y - D\alpha - D_x M) \quad (4)$$

In which $D = (I_n \otimes i_T) \cdot d_n, d_n = [-i_{n-1}, l_{n-1}]'$, and i_n is a $n \times 1$ vector of ones. D is added to make sure that $N^{-1} \sum_{i=1}^N \alpha_i = 0$. Define the operator of smoothing.

$$S(x) = (D_x' \mathcal{K}_x D_x)^{-1} D_x' \mathcal{K}_x$$

2- Method Differencing Transformation

Consider the first differencing of Equation (4), and results in

$$y_{it} - y_{it-1} = m(x_{it}) - m(x_{it-1}) + \varepsilon_{it} - \varepsilon_{it-1}, i = 1, \dots, N, t = 2, \dots, T$$

(Ullah ,1998) assumed the existence of known derivatives of $m(x)$ and thus, it was possible to determine and estimate the derivatives $m(x)$ using local linear (polynomial) regression only. As an example, given $q = 1$ then write $\Delta z_{it} = z_{it} - z_{it-1}$.
 $\Delta y_{it} = m(x_{it}) - m(x_{it-1}) + \Delta \varepsilon_{it}$.

The Taylor expansion of $m(x_{it}), m(x_{it-1})$ at the point x , therefore, is first-order and the result is.

$$\begin{aligned}\Delta y_{it} &= m(x) + m'(x)(x_{it} - x) - (m(x) + m'(x)(x_{it-1} - x)) + \Delta \varepsilon_{it} \\ &= m(x) - m(x) + m'(x)((x_{it} - x) - (x_{it-1} - x)) + \Delta \varepsilon_{it} \\ &= \Delta x_{it} m'(x) + \Delta \varepsilon_{it}.\end{aligned}$$

It is the local linear estimator of $m'(x)$, which is suggested by (Lee , 2014). which is as follows:

$$\hat{m}'(x) = \frac{\sum_{i=1}^N \sum_{t=2}^T K_{itxh} \Delta x_{it} \Delta y_{it}}{\sum_{i=1}^N \sum_{t=2}^T K_{itxh} \Delta x_{it}^2}. \quad (5)$$

The operation of this estimator can be illustrated by observing that the ordinary in transformation calculates the individual specific mean of a variable z as $\bar{z}_i = T^{-1} \sum_{t=1}^T z_{it}$. (Lee , 2014) substitute the fixed weights of $1/T$ by weights on the kernel which result into individual-specific weighted means that are locally weighted.

$$\tilde{z}_i = \frac{\sum_{t=1}^T K_{itzh} z_{it}}{\sum_{s=1}^T K_{iszh}} = \sum_{t=1}^T w_{itzh} z_{it} \quad (6)$$

The locally defined by the notation $z_{it}^* = z_{it} - \tilde{z}_i$ it bar.

$$\hat{m}'_{LWTLc}(x) = \frac{\sum_{i=1}^N \sum_{t=2}^T K_{itxh} x_{it}^* y_{it}^*}{\sum_{i=1}^N \sum_{t=2}^T K_{itxh} x_{it}^{*2}} \quad (7)$$

These local differentiating and transformation techniques are easy to use.

3- method Profile Estimation

n methods proposed by (Ullah ,1998) (Lee ,2014) assume estimation of the Equation (8) as a fixed-effects model, using the first difference , according to period 1:

$$\tilde{y}_{it} \equiv y_{it} - y_{i1} = m(x_{it}) - m(x_{i1}) + \varepsilon_{it} - \varepsilon_{i1}. \quad (8)$$

Although this estimator might be applied with period $t - 1$, which is more usual, we do as they do. The advantage of such a conversion is that, in fact, the fixed effects are eliminated, and, when the transferences are exogenous $E[m(x_{it})] = E(y_{it})$. The issue where it has been expressed in Equation (8) is the occurrence of $m(x_{it})$ and $m(x_{i1})$. The key point of (Henderson , 2005) is using the variance structure of $\varepsilon_{it} - \varepsilon_{i1}$ when building the estimator, much like (Wang ,2003) in the random effects context. (Fan, 2005).

4- method Marginal Integration

This is proposed by (Qian ,2012). To provide a description of their estimator we initially redefine the first differencedl so as follows:

$$\Delta y_{it} \equiv y_{it} - y_{it-1} = m(x_{it}) - m(x_{it-1}) + \Delta \varepsilon_{it}. \quad (9)$$

the estimation in Equation (9) by .

$$\Delta y_{it} = m(x_{it}, x_{it-1}) + \Delta \varepsilon_{it}$$

then eliminating x_{it-1} to get an estimator of $m(x_{it})$. This advantage is that any standard nonparametric estimator can be employed, e.g. local-polynomial leastsquares.

5- Kernel smoothing methods

In the case of longitudinal data, the functional coefficients are functions of the time variable and supposed to be smooth. The following modeling structures are generally used in nonparametric estimates of curves: In order to gain a modest idea of the smoothing methods most used, I will make a brief analysis of these various techniques using simple cross-sectional, i.i.d. data without $X(t)$. The modeling structures used to get an estimate of curves are usually: (Henderson, 2015)

$$\mathcal{Y} = \{(Y_i, t_i)^\top : i = 1, \dots, n\}$$

In a regression model underneath the i.i.d of $(Y(t), t)^\top$ the goal statistical is to approximate the smooth mean $\mu(t) = E(Y(t) | t)$ under the same statistic.

$$Y(t) = \mu(t) + \epsilon(t),$$

where $\epsilon(t) \sim N(0, \sigma^2)$. There are local and global smoothing techniques of estimating $\mu(t)$. Two of the most popular methods of applying the techniques of local smoothing using each known kernel in the field are kernel and local methodology. Nadaraya Watson estimator of $\mu(t)$ can be obtained by minimizing the local least squares in regard to $\mu(t)$ (Zhou, 2010).

$$\hat{\mu}_K(t) = \frac{\sum_{i=1}^n \left\{ Y_i K \left[\frac{(t - t_i)}{h} \right] \right\}}{\sum_{i=1}^n \left\{ K \left[\frac{(t - t_i)}{h} \right] \right\}}, \quad (10)$$

with $K(\cdot)$ a kernel function, usually a non-negative probability density function and h a positive bandwidth. The estimator will operate with units whose $K(\cdot)$ is in a band of h of t . A few common instances of a kernel function (Kernel(statistics)), which has restricted or unbounded support are as follows:

Table (1): Kernel functions used

Epanechnikov kernel	$K_E(u) = \frac{3}{4}(1 - u^2)1_{ u \leq 1}$
Gaussian kernel	$K_G(u) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}u^2}$
Cosine kernel	$K_{\text{Cos}}(u) = \frac{\pi}{4} \cos\left(\frac{\pi}{2}u\right) 1_{ u \leq 1}$

The above functions are used since the data is not single (numerical image) and so the functions can handle more than one variable. (Köhler, 2014).

The optimality property of the Epanechnikov kernel as the minimizer of the asymptotic MSE of the kernel estimators, has been proved under some rather weak general asymptotic assumptions (Bosq, 2012). Whereas various kernel functions lead to various local weights of the kernel estimators, the theoretical and practical behaviour of kernel estimators have been shown by asymptotic derivation, simulation, and practice to be largely dependent on the bandwidth decisions, as opposed to the shape of the kernel functions (Hardle, 2004). (Chad Schafer, 2014)

Bandwidth selection The bandwidth selection is key to the kernel smoothing to achieve a sufficient kernel estimator of $\mu(t)$. A small bias and a large variance make an estimator whose bandwidth value generates an overspecialized curve (Kohler et al., 2014). (Wu, 2000) suggested that the bandwidth would be chosen by applying cross-validation methods. In case of longitudinal data, it is more appropriate to drop the entire i th unit rather than merely the i th observation since repeated measurements within a unit are normally correlated, and the correlation pattern is often quite uninformed in practice.

I suppose that I have a longitudinal sample, which is represented by $\{Y_{ij}, t_{ij} : i = 1, \dots, n; j = 1, \dots, n_i\}$. Suppose that $\hat{\mu}(t)$ is a kernel estimator of $\mu(t)$. n units are selected at random out of a population, where a unit level can be either an individual or a cluster/family, and n_i is the

amount of repeat measurements of the i th unit. A unit The 0 -optimal bandwidth will be determined by minimizing 0 MSE of $\hat{\mu}(t)$. LUCV bandwidth of $\hat{\mu}(t)$ is selected by the following three procedures.

- 1- Find estimator $\hat{\mu}^{(-i)}(t)$
- 2- foretell the outcome of the i th unit at time t_{ij} .
- 3- Define the LUCV score of $\hat{\mu}(t)$ by

$$LUCV(h) = \sum_{i=1}^n \sum_{j=1}^{n_i} [Y_{ij} - \hat{\mu}^{(-i)}(t_{ij})]^2 \quad (11)$$

Then h_{LUCV} is the best bandwidth which is determined by decreasing $CV(h)$ to all the positive values of $h > 0$. (Rodriguez-Poo, 2017).

4th: Data description and results analysis

In this section, we will discuss the calculation of nonparametric models for longitudinal data using the estimation methods mentioned in the theoretical aspect of this study. Then, we will compare these methods using comparison criteria to determine the best model and also determine the most efficient method.

The study was conducted on health data in Iraq containing longitudinal data for following 150 Iraqi patients with hypertension in 3 hospitals in Baghdad over a period of 6 months. This data represents the effect of medications on blood pressure in patients with hypertension. The sample consisted of 450 observations. The dependent variable Y was the blood pressure of each patient at each measurement time, while the independent variables were six variables: x1 Time (three measurements for each patient before treatment, three months after, and six months after treatment); x2 Patient ID; x3 Type of medication used (constant for each individual); x4 Age (constant for each individual); x5 Body Mass Index (BMI) (changes over time); and x6 Number of times of exercise per week. These data were obtained from the Central Statistical Organization and health reports.

In the first step, the type of longitudinal data model was identified. The Honda method of Lagrange's multiple test was used to determine the suitability of the data to the clustering model, in order to test the null hypothesis that the clustering model fits the studied data. In the event that the data does not fit the clustering model, then a comparison must be made between the random effects and fixed effects models through the Hausman test, in order to decide on the null hypothesis Random effects model is more suitable compared to the fixed effects model. The results are indicated and shown in Tables (2) below:

Table (2): Lagrange's double test and Hausman's test for data on the number of people with high blood pressure

(p-value)	Test statistic value	The test
1.2 e-11	11.379	Lagrange multiplier test for health data
0.0416	0.204	Hausman test for health data

Table (2), which concerns the health data tests, shows that the (p-value) of the Lagrange multiplier test reached (1.2 e-11), This is lower than the level of significance (5%). This implies that the null hypothesis should be rejected, i.e. the group model is inappropriate in the context of health data. Meanwhile, the Hausman statistical test probability value (p) of the data was (0.0416), which is less than the significance level. This prompts us to reject the null hypothesis, i.e. the fixed model is more suitable than the random model in representing the data. Therefore, the nonparametric estimators of the nonparametric fixed effects model will be addressed.

Due to the inconsistency of the units of measurement for the studied data, they were converted to the standard format. (standardization) $Z = \frac{X-\mu}{\sigma}$

After identifying the nature of the studied data and describing the appropriate longitudinal data model for that data, which is the fixed model, the estimation methods that were discussed in the theoretical aspect will be applied to the data, as the parameters of the fixed non-parametric model were estimated for both sides and the results were placed in Table (3)

Table (3): Estimated values of the parameters of the fixed model for the data on the number of people with high blood pressure

Parameter	β_1	β_2	β_3	β_4	β_5	β_6
	0.063	0.045	0.022	0.34	0.021	7.24

In order to avoid confusion between the different estimation methods, the names and details of these methods have been placed in Tables (4) below.

Table (4): Models used in estimation

Model Name	Method
Non Parametric 1	NW, Profile Least-Squares, Gaussian
Non Parametric 2	NW, Profile Least-Squares, Epanechnikov
Non Parametric 3	NW, Profile Least-Squares, Cosine
Non Parametric 4	NW, Differencing/Transformation Methods, Gaussian
Non Parametric 5	NW, Differencing/Transformation Methods, Epanechnikov
Non Parametric 6	NW, Differencing/Transformation Methods, Cosine
Non Parametric 7	NW, Profile, Gaussian
Non Parametric 8	NW, Profile, Epanechnikov
Non Parametric 9	NW, Profile, Cosine
Non Parametric 10	NW, Marginal Integration, Gaussian
Non Parametric 11	NW, Marginal Integration, Epanechnikov
Non Parametric 12	NW, Marginal Integration, Cosine

To compare these different estimation methods and to indicate the best model, comparison criteria were used Root Mean Squared Error (RMSE), Root Relative Squared Error (RRSE), Average Relative Deviation Error (MAPE), Mean Relative Deviation Error (MAE), and Relative Deviation Error (RAE). (Mehdiyev, N., 2016)(Bai, 2009).

Table (5): Values of comparison criteria between the estimated models for data on the number of people with high blood pressure.

Model	RMSE	RRSE	MAE	MAPE	RAE	Best
1	0.30847	0.27661	0.35823	0.16584	0.22502	4
2	0.11345	0.03724	0.19569	0.07190	0.06432	1
3	0.90468	0.89156	0.95920	0.92129	0.85360	9
4	0.39291	0.30503	0.39439	0.23025	0.30725	5
5	0.28643	0.16511	0.29562	0.09577	0.08953	2
6	0.95428	0.91868	0.96428	0.97111	0.91370	10
7	0.29621	0.19247	0.30583	0.02572	0.12728	3
8	0.45782	0.40628	0.45738	0.38257	0.38838	6
9	0.98265	0.92626	0.98427	0.97938	0.98455	11
10	0.51378	0.45217	0.49859	0.42740	0.45301	7
11	0.75101	0.60982	0.65731	0.62930	0.55825	8
12	0.99418	0.95653	0.99286	0.97985	0.99368	12

It is clear from Table 5 above, which concerns the comparison criteria estimated for the different models of data on the number of people with blood pressure, that the non-parametric model (non-parametric 7), estimated according to the method of the Nadaraya-Watson method using the Gaussian function, is the most efficient model among all models, as it gave the lowest values for the three comparison criteria. The non-parametric model estimated according to the method of the Nadaraya-Watson method using the Epanechnikov function, i.e., the non-parametric model (Model 2), came in second place, while the non-parametric model (Model 5), estimated according to the Nadaraya-Watson method weighted with the Epanechnikov function, came in third place. The non-parametric model (Model 1), estimated according to the Nadaraya-Watson method weighted with the Gaussian function, came in fourth place for all comparison criteria, while the non-parametric model (Non Parametric 4), estimated according to the method of the The Nadaraya-Watson Gaussian model ranked fifth, while the nonparametric model estimated by the Nadaraya-Watson Epanechnikov method ranked sixth, Although the non-parametric model (Model 10) which was estimated through the Nadaria-Watson method with a Gaussian function was the 7th-ranked, the non-parametric model (Model 11) which was estimated through the Nadaria-Watson method with a Gaussian function was the 8th-ranked among all comparison criteria. The non-parametric model (Model 3) estimated using the Nadaria-Watson method with a cosine function ranked ninth and the non-parametric model (Model 6) estimated with the Nadaria-Watson method with a Gaussian function ranked seventh.

Conclusions:

The random model is not suitable for analyzing and studying data on the number of people with high blood pressure according to the Honda Lagrange multiplier test. The appropriate model for this study is the fixed-effects model according to the Hausman test. Through the comparison criteria values, it became clear that the most efficient model to describe the data of the number of people with blood pressure is the non-parametric model (7) estimated according to the Nadara-Watson method weighted by the Gaussian function.

The Gaussian function is more efficient than the Epanchincov function in estimating the fixed longitudinal data model for all non-parametric estimation methods. While it was proven that the Cosine function was the worst function in estimation. From what was presented in points above, it is clear that the most efficient model that explains the nature of the relationship is the model estimated using the Nadaraja-Watson method weighted by the Gaussian function.

Sometimes the results of comparing estimators vary depending on the comparison criteria. The comparison criteria (AMSE) were more stable and accurate than the comparison criteria (AMAE), Root Mean Squared Error (RMSE), Root Relative Squared Error (RRSE), Mean Absolute Error (MAE), Relative Absolute (RAE)

References

- 1- Arellano, M. (2003). Panel Data Econometrics. Oxford: Oxford University Press.
- 2- Bai, J. (2009). Panel data models with interactive fixed effects. *Econometrica* 77, 1229-1279.
- Chad Schafer and Larry Wasserman, (2014) . " Tutorial on Nonparametric Inference With R " Carnegie Mellon University working paper .
- 3- Baltagi, B. H. (2008). Econometric Analysis of Panel Data. 4th ed. West Sussex: John Wiley & Sons.
- 4- Bosq, D. (2012). *Nonparametric statistics for stochastic processes: estimation and prediction* (Vol. 110). Springer Science & Business Media.
- 5- Cho, H., & Kim, S. (2017). Model specification test in a semiparametric regression model for longitudinal data. *Journal of Multivariate Analysis*, 160, 105-116
- 6- Fan, J. and T. Huang (2005). "Profile likelihood inferences on semiparametric varying-coefficient partially linear models". *Bernoulli* 11, 1031-1057.
- 7- Härdle, W. (2004). *Nonparametric and semiparametric models*. Springer Science & Business Media.
- 8- Henderson, D. J., & Ullah, A. (2005). A nonparametric random effects estimator. *Economics Letters*, 88(3), 403-407
- 9- Henderson, D. J., & Ullah, A. (2015). Nonparametric estimation in a one-way error component model: a Monte Carlo analysis. In *Statistical Paradigms: Recent Advances and Reconciliations* (pp. 213-237).
- 10-Honda, Y.,(1985)." Testing the error components model with nonnormal disturbances" , review of economic studies. 52(4),681-690.
- 11-Hsiao, C. (2003). Analysis of Panel Data. 2nd ed. Cambridge: Cambridge University Press.
- Hu, Z., N. Wang and R. J. Carroll. 2004. Profile-kernel versus backfitting in the partially linear models for longitudinal/clustering data. *Biometrika* 91: 251-262.
- 12-Köhler, M., Schindler, A., & Sperlich, S. (2014). A review and comparison of bandwidth selection methods for kernel regression. *International Statistical Review*, 82(2), 243-274.
- 13-Lee, Y., & Mukherjee, D. (2008). New nonparametric estimation of the marginal effects in fixed-effects panel models: An application on the environmental Kuznets curve. *Available at SSRN 1119315*.
- 14-Lee, Y., & Mukherjee, D. (2014). Nonparametric estimation of the marginal effect in fixed-effect panel data models: an application on the environmental Kuznets curve. *Unpublished manuscript, Syracuse University*.
- 15-Mehdiyev, N., Enke, D., Fettke, P., & Loos, P. (2016). Evaluating forecasting methods by considering different accuracy measures. *Procedia Computer Science*, 95, 264-271
- 16-Parmeter, C. F., & Racine, J. S. (2019). Nonparametric estimation and inference for panel data models. *Panel data econometrics*, 97-129.
- 17-Rodriguez-Poo, J. M., & Soberon, A. (2017). Nonparametric and semiparametric panel data models: Recent developments. *Journal of Economic Surveys*, 31(4), 923-960.
- 18-Su, L., & Ullah, A. (2006). Profile likelihood estimation of partially linear panel data models with fixed effects. *Economics Letters*, 92(1), 75-81.
- 19-Su, L., & Ullah, A. (2011). Nonparametric and semiparametric panel econometric models: estimation and testing. *Handbook of empirical economics and finance*, 455-497.
- 20-Ullah, A. and N. Roy. 1998. Nonparametric and semiparametric econometrics of panel data. In *Handbook of Applied Economic Statistics*, ed. A. Ullah and D. E. A. Giles, 579-604. New York: Marcel Dekker.
- 21-Wang, N. (2003). "Marginal nonparametric kernel regression accounting for within-subject correlation". *Biometrika* 90, 43-52.
- 22-Wu, C. O., & Chiang, C. T. (2000). Kernel smoothing on varying coefficient models with longitudinal dependent variable. *Statistica Sinica*, 433-456.
- 23-Zhou, Z., & Wu, W. B. (2010). Simultaneous inference of linear models with time varying coefficients. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 72(4), 513-531.