

Using Panel Data Techniques for Social Science Research: an Illustrative Case and Some Guidelines

J. Ramón Gil-García* y Gabriel Puron-Cid*

Recepción: 30 de enero de 2013

Aceptación: 02 de diciembre de 2013

*Centro de Investigación y Docencia Económicas, México.

Correos electrónicos: joseramon.gil@cide.edu y gabriel.puron@cide.edu

Se agradecen los comentarios de los árbitros de la revista.

Usando técnicas para datos tipo panel en la investigación en ciencias sociales: un caso ilustrativo y algunas recomendaciones

Resumen. Se muestran las ventajas y se describen los pasos para aplicar técnicas para datos tipo panel (DP) en ciencias sociales. Se argumenta que los investigadores deberían entender las oportunidades y beneficios que se derivan al usarlas, incluyendo análisis más robustos de fenómenos sociales complejos. Usando un set de datos longitudinal sobre subsidios para capacitación en empresas de Michigan, se presentan dos variedades de análisis estadístico usando métodos DP (efectos fijos y efectos aleatorios). El ejemplo es un caso usado en cursos de análisis cuantitativo en programas de educación superior en los Estados Unidos y alrededor del mundo. Modelos de efectos fijos, efectos aleatorios y regresión tradicional son presentados y discutidos.

Palabras clave: estadística, datos tipo panel, ciencias sociales, métodos cuantitativos, efectos fijos, efectos aleatorios.

Abstract. This paper shows the advantages of applying panel data (PD) techniques and describes its steps in the context of social science research. We argue that researchers should understand the opportunities and benefits of using PD techniques in social sciences, including a more robust and meaningful analysis of complex social phenomena. Using a longitudinal dataset of training subsidies for Michigan firms, this paper shows how to perform two varieties of statistical analysis using PD methods (fixed effects and random effects). The example is a case widely used in quantitative analysis courses in higher education programs in the United States and around the world. Pooled OLS, fixed effects, and random effects models are presented and discussed.

Key words: statistics, panel data, social sciences, quantitative methods, fixed effects, random effects.

Introduction

Social phenomena usually present multiple dimensions in terms of time, context, and the complexity of variables interacting with each other. These effects are captured by social scientists in theory as well as in empirical research. However, recent use of sophisticated statistical techniques along with technological and computational advances have allowed to extend these techniques into a wide range of applications. Controlling for potential variance and covariance among interacting variables across time and units has imposed several

theoretical and methodological challenges for understanding social phenomena that now are possible due to technological and statistical advances.

This article presents these difficulties as an opportunity if the researcher uses panel data techniques (PD). Models using traditional Ordinary Least Squares (OLS) estimations have demonstrated questionable results and the history of regression analysis is riddled with violations of its assumptions (Bickel, 2007; Gil-García, 2008; Gefen, Straub & Boudreau, 2000; Hair *et al.*, 1998). As a result, a vast group of tests and procedures to detect or correct for OLS violations has

been developed. Similar to these adjustments, PD methods attempt to make regression analysis more useful and robust for scientific work.

The purpose of this paper is to illustrate the steps and procedures for conducting PD analysis in the context of social science research. PD methods are best used when the power and versatility of OLS regression and its correctives have been exhausted or when data with certain characteristics are available. Using a longitudinal dataset of training subsidies for Michigan firms, this study applies PD techniques to show its steps and necessary actions in the context of doing research about social phenomena. It also compares the results of PD techniques and OLS with the purpose of highlighting their differences. The dataset involves three-year panel data from a unique survey of manufacturing firms that applied for training grants collected, which were studied by Holzer *et al.* (1993). The case has been previously used in quantitative courses of higher education programs in the US and around the world.

The idea of this article is to explore the benefits and limitations of applying PD methods when studying a particular social phenomenon. In this case, the paper illustrates the applications of PD methods to examine the effects of job training provided by state government grants. The purpose of this government program is to give grants to firms in order to increase job training and consequently the productivity of these firms and their employees' salary. Therefore, the actual hours of training per employee can serve as a proxy of effectiveness. This case could be considered a classical example for analyzing longitudinal and nested effects.

By using this case, this paper shows the steps, benefits, and limitations of two PD techniques (fixed effects and random effects) to demonstrate the main procedures for estimation and significance tests. Fixed effects models capture the effects of the individual characteristics of firms over time. Random effects techniques keep unobserved variables across individuals or time in the error term. The goal of the first technique is to explicitly control for the effects of a particular set of variables, while the goal of the second method is to evaluate the general effect on average across time and units.

Corrections to OLS limitations are possible by using fixed effects techniques that control for time and contextual variables or by applying random effects models that allow for variation across cases and time. In any case, researchers should consider the purpose of their investigation and the nature of the nested problem that is under examination. This article also shows that a proper specification of the model by controlling for fixed or random effects is a more powerful tool than just piling up pooled data for correcting

autocorrelation, deflated standard errors, and underestimated degrees of freedom.

In general, the paper shows an improvement of model specification by using PD techniques and a better understanding of nested social phenomena. In our illustrative case, the effects of government grants on job training, productivity and wages were controlled across different firms and time. The analysis adds insights and complements the findings of the original case by demonstrating the advantages and disadvantages of two PD techniques: fixed effects and random effects.

The paper is organized in four sections, including the foregoing introduction. Section two describes the main characteristics of pooled OLS, fixed effects, and random effects models. It also includes their functional forms, as well as some advantages and disadvantages. Section three presents an illustrative example using data from the Michigan experience and discusses some similarities and differences among the estimates obtained using these three models. Finally, section four highlights some of the most important advantages and disadvantages of these techniques, provides general results, and proposes some practical recommendations as guidance for social science researchers.

1. Panel Data Techniques

Different names exist to refer to panel data: pooled data, pooled time series and cross-sectional data, micropanel data, longitudinal data, and event history analysis, among others (Baltagi, 2008; Greene, 2012; Gujarati, 2003; Wooldridge, 2002). The use of PD models comes from the fact that data used in many social sciences usually combines time series and cross sections of units (Greene, 2012). In other words, data is embedded across multiple contexts and over time in what is called the nested phenomena (Bickel, 2007), where contextual characteristics of units or individuals are embedded over time, or vice versa, and each point in time is entrenched across different characteristics of units (the growth curve). Therefore, PD sets represent the same cross-sectional unit of study surveyed over time (Baltagi, 2008; Gujarati, 2003; Wooldridge, 2002). Originally, PD estimations were used as a way to improve model specification and estimations. Now due to technological and computational progress, researchers and decision makers can apply these methodologies for examination of a particular social problem or series of events across time and contexts.

Today, the techniques have been refined and there are several statistical packages that compute PD models such as NLOGIT, SAS, Stata, SPSS, and LISREL, or using various programming environments like Gauss, Matlab, Mathematica, and

R. Several free access resources are available to researchers to introduce themselves to PD techniques and models as well. For example, Dr. William Greene's website¹ includes class notes, presentations, datasets, and guide examples about PD techniques. The International University of Japan also provides a practical guide to PD models using Stata software² (Park, 2011). Bond (2002)³ from the Institute for Fiscal Studies in UK has also contributed with a guide for micro data methods and practice for dynamic PD models. Princeton University has developed a site for data and statistical services, with one page dedicated to PD models⁴ using Stata. These are resources available on the Web to guide researchers in different social sciences about PD models and techniques.

The data arranged in panel includes case by case observations for each variable over different observations across time. Time units may be in years, months, weeks, days, or any other time expressions (hours, minutes, or seconds). Time units have different potential expressions depending on the expected variable behavior over time. Researchers may consider different time expressions such as lagged, linear, squared or quadratic, among others representations. A case is an individual observation of a particular indicator or variable from a person, group, firm, organization, city, state or country. An identifier of each case should be included.

In principle, it is possible to estimate time series for each case or cross-sectional regressions for each time unit by using the expressions (1) and (2) correspondingly. These expressions represent simple pooled OLS models. Simple linear regression using a panel data arrangement is called pooled OLS regression. Pooled OLS models only stack observations for each case over time, one on the top of the other, which does not result in distinctions across cases and over time. This estimation disregards the effects over individuals and time, and eventually distorts the true picture of the relationships of variables being studied across cases and over time.

$$Y_t = \alpha + \beta_1 X_{1t} \dots \beta_n X_{nt} + e_t \quad (1)$$

$$t = 1 \dots t$$

$$Y_i = \alpha + \beta_1 X_{1i} \dots \beta_n X_{ni} + e_i \quad (2)$$

$$i = 1 \dots i$$

The results of pooled regression may present statistically significant coefficients; the slope coefficients may indicate expected signs; and, the R^2 value may be reasonably high. As a result, analysts may infer a strong, positive or negative, relationship between Y_{it} and X_{it} . However, the estimation may imply potential auto-correlation in the data, which a

Durbin-Watson statistic can diagnose. When a low Durbin-Watson statistic is found, then there are high chances of auto-correlation or misspecification of the model. For instance, the estimated model may assume that the intercept and slope coefficients values are the same on average across cases (i) or over time (t), but the panel data set does not corroborate this finding. In this case these assumptions are misspecified in the model and there is the potential risk of heteroscedasticity or auto-correlation in the estimations. The results will lead to biased estimates of the variances for each of the estimated coefficients, and therefore statistical tests and confidence intervals will be incorrect (Baltagi, 2008; Gujarati, 2003; Pindyck and Rubinfeld, 1998; Wooldridge, 2002).

As efforts to correct for these errors, researchers started to extend traditional OLS models using dummy variables for each category, characteristic of cases, or time expressions as shown in the expression (3). These types of models are known as Least-Squares Dummy Variable Regression Models (LSDV), which is a variation of fixed effects models.

$$Y_{it} = \alpha + \alpha_1 D_{1it} + \dots + \alpha_n D_{nit} + \beta_1 X_{1it} \dots \beta_m X_{mit} + \mu_{it} \quad (3)$$

$$t = 1 \dots t \quad n = 1 \dots \text{number of dummy variables}$$

$$i = 1 \dots i \quad m = 1 \dots \text{number of independent variables}$$

With the advent of technological sophistication and computational power, modern fixed effects techniques today do not explicitly consider dummy variables in their models. Instead, most statistical packages use probabilistic estimations based on the covariance matrix. PD models are more robust because, in general, they consider full information from all observations across cases and over time in the same dataset. However, these estimations need adequate model specifications to properly capture the effects of covariance among cases and across time. A basic PD model, for fixed effects and random effects, can be written as in (4) and (5) respectively.

$$Y_{it} = \alpha_{1i} + \beta_1 X_{1it} \dots \beta_n X_{nit} + \mu_{it} \quad (4)$$

$$t = 1 \dots t$$

$$i = 1 \dots i$$

$$Y_{it} = \alpha_1 + \beta_1 X_{1it} \dots \beta_n X_{nit} + \omega_{it} \quad (5)$$

$$t = 1 \dots t$$

$$i = 1 \dots i$$

1. <http://pages.stern.nyu.edu/~wgreene/Econometrics/PanelDataEconometrics.htm>

2. http://www.iuj.ac.jp/faculty/kucc625/documents/panel_iuj.pdf

3. <http://www.ifs.org.uk/publications/2661>

4. http://dss.princeton.edu/online_help/stats_packages/stata/panel.htm

As a matter of convention, several authors explain that the subscripted notation i stands for i th cross-sectional unit and t for the t th time period (Baltagi, 2008; Gujarati, 2003; Wooldridge, 2002). It is assumed that there are a maximum of N cross-sectional units and a maximum of T time observations. Cameron and Trivedi (2009), Gujarati (2003) and Greene (2012) also observe that if each cross-sectional unit has the same number of time series observations, then such a panel (set of data) is called a balanced panel. If this is not the case, then the panel is called an unbalanced panel. Gujarati (2003) also indicates that “it is assumed that the X ’s are non-stochastic and the error term follows the classical assumptions: $E(\mu_{it}) \sim N(0, \sigma^2)$ ”.

Gujarati (2003), Judge *et al.* (1985), and Hsiao (1986) discuss two main approaches for PD regression models: (1) fixed effects, and (2) random effects. Each of these approaches presents different modalities of PD estimation. These models are based on functions (4) and (5). These two modalities depend on diverse assumptions about the intercept, the slope coefficients, and the error term. Table 1 compares the simple OLS assumptions with these two modalities of the PD approaches. Gujarati (2003) and Baltagi (2008) discuss several advantages of using PD techniques over traditional OLS regression:

a) PD models take into consideration time and cases simultaneously, whereas other models have the limitation of only expressing these heterogeneities across units or over time.

b) PD models provide more information, more variation, less collinearity among variables, more degrees of freedom, and more efficiency by combining time series and cross-section observations.

c) PD models better express the dynamics of change within social phenomena.

d) PD models better detect and measure the mixed and pure effects across cases and over time.

e) PD models provide better representations of complex social or behavioral models.

f) PD models minimize the bias of aggregating cases into broader categories.

In sum, the techniques of PD models can simultaneously take into account the heterogeneity across cases and over time by controlling the different effects at the level of individual-specific variables. According to many authors (Baltagi, 2008; Gujarati, 2003; Pindyck and Rubinfeld, 1998; Wooldridge, 2002), there are some advantages of using PD models: increasing the sample size; capturing the heterogeneity involved both in cross-section units and time dimensions; testing hypotheses about the presence of heteroscedasticity or autocorrelation, or both; and, finally, they are better suited to study the dynamics of change and complex behavioral models.

These authors have also mentioned that PD models have several estimation and inference problems that need to be addressed. Hsiao (2006: 19) specifically mentions, “The power of panel data to isolate the effects of specific actions,

Table 1. Assumptions for OLS Regression and PD Models.

Type of Approach	PD Model	Intercept α	Slope Coefficients β	Error Term μ	Assumptions
OLS Regression	Unspecified Regression	Constant _{it}	Constant _{it}	$E(\mu_{it}) \sim N(0, \sigma^2)$	All coefficients constant disregarding time and individuals.
Fixed Effects (FEM)	Pooled Regression	Constant _{it}	Constant _{it}	$E(\mu_{it}) \sim N(0, \sigma^2)$	Assume that the <i>intercept</i> and <i>slope coefficients</i> are constant across time and space, and the error term captures differences over time and individuals.
	Least-Squares Dummy Variable Regression Model (LSDV) or Covariance Model	Variable _t	Constant _{it}	$E(\mu_{it}) \sim N(0, \sigma^2)$	The slope coefficients are constant but the <i>intercept</i> varies across <i>individuals</i> .
		Variable _i	Constant _{it}		The slope coefficients are constant but the <i>intercept</i> varies over <i>time</i> .
		Variable	Constant _{it}		The slope coefficients are constant but the <i>intercept</i> varies across <i>both</i> individuals and over time.
		Variable _t	Variable _t		All coefficients (the <i>intercept</i> as well as <i>slope coefficients</i>) vary over <i>individuals</i> .
		Variable	Variable		The <i>intercept</i> and slope coefficients vary over <i>both</i> individuals and time.
Random Effects (REM)	Generalized Least Squares (GLS)	Variable	Variable	$\varepsilon_i \sim N(0, \sigma_\varepsilon^2)$ $\mu_{it} \sim N(0, \sigma_\mu^2)$	The <i>intercept</i> is the <i>grand mean</i> across cases and over time. No correlations and autocorrelations between and within error terms.

Source: assumptions taken from Gujarati (2003: 640).

treatments or more general policies depends critically on the compatibility of the assumptions of statistical tools with the data generating process". In particular, Hsiao (2006) indicates that choosing a proper method for exploiting the richness and unique properties of the panel involves three disadvantages: the availability of data, generating specification problems, omitted variable bias, bias induced by dynamic structures, and the efficiency of estimates across and within cross-sectional units and time. Hsiao (1986) has also pointed out that the main limitations of PD models are mainly in the data generating process, such as inadequate sampling of population of interest, response rates, question design errors, distortion by deliberating answering behavior, frequency or temporal lapses of answering, and time references or periods for answering, etc.

1. 1. Fixed Effects Models (FEM)

Baltagi (2008), Gujarati (2003), and Wooldridge (2002) point out that this modality known as "fixed effects models" (FEM) or "covariance models" allows the intercept to differ across cases, but not over time (time invariant). FEM assume that the slope coefficients are constant while the intercept varies across cross-sectional units. This type of approach takes into account individuality by letting the intercept vary across cases while slope coefficients are assumed to be constant across firms.

At the beginning, FEM applications were called "Least-Squares Dummy Variable" (LSDV) models. This initial type of model took their name from the common practice of including dummy variables (the dummy-variable trap). By specifying dummy variables for each of the "characteristics" of the cases, the estimation provides differential intercept effects (fixed effects or differential intercept dummies' effects). In order to apply this technique, the equations (1) and (2) were re-written as in (3) where D_{ni} are dichotomic ("1" if characteristic is present and "0" otherwise).

The LSDV models can also be applied to "time effects" in the sense that dummy variables are functions of time units, not cases or characteristics. This modality shows shifts over time, such as changes in performance, technology, and other contextual factors applied in the same case. Similar considerations and tests should be considered in order to deal with the differential dummy specification. A more complex application of the LSDV model may consider dummies for both "case" effects and "time" effects. Baltagi (2008), Gujarati (2003), and Park (2011) indicate that this type of model assumes that the intercept and the slope coefficients are different across cases. By reconsidering the equation (3), the specification also introduces differential

slope dummies (by multiplying each dummy by each of the X variables). The γ 's represent the number of these interactions. The resulting equation is as follows:

$$Y_{it} = \alpha + \alpha_1 D_{1i} + \dots + \alpha_n D_{ni} + \beta_1 X_{1it} \dots \beta_m X_{mit} + \gamma_1 (D_{1i} X_{1it}) + \dots + \gamma_w (D_{ni} X_{mit}) + \mu_{it} \quad (6)$$

$t = 1 \dots t$ $n = 1 \dots$ number of dummy variables
 $i = 1 \dots i$ $m = 1 \dots$ number of independent variables
 $w = 1 \dots$ number of interactions

If any of these interaction results are statistically significant, it will mean that there is a significant difference between groups. For example, if the interaction between D_{1i} and X_{1it} is statistically significant, then the net effect of a particular case i over time is $(\beta_1 + \gamma_1)$, suggesting that X_1 is different than the comparison case. If the differential intercept and all the differential slope coefficients are statistically significant, we conclude that all cases are different from the comparison case. Therefore, there is no support for pooled regression estimation due to significant differences among cases. In other words, the data is not "poolable" (Gujarati, 2003; Park, 2011; Wooldridge, 2002).

The limitation of the LSDV technique is that depending on the number of characteristics being studied, the number of degrees of freedom increases. A second limitation is that estimations are linearly estimated. Previously, analysts had to pay attention to what is specified in the differential dummy model. One option was to set a case or characteristic that will be specified as the intercept (α_1), which the rest of the dummy variables will be implicitly compared to. The other option was to set dummy variables for all possible cases or characteristics, but it is necessary to drop the common intercept.

Further computational advances have allowed for more robust estimations of FEM specifications by adding the subscripted notation i in the intercept that allows for variation across cases. Therefore, new FEM representations modified the equations (1) and (2) into (4). By adding the subscript i on the intercept term, the equation (4) suggests variation on intercepts across cases. This variation in the intercept is assumed a priori in the specified model because there are differences across cases attributable to special characteristics of individuals, due to their contexts or another grouping category. Several authors have also recognized other modalities of the traditional FEM approach (Baltagi, 2008; Greene, 2012; Gujarati, 2003; Park, 2011). For example, re-writing the equation by adding to the intercept the subscripts i and t , we suggest that the intercept of each case varies, but also that each case is also time variant.

The equation (7) below shows this addition:

$$\begin{aligned} Y_{it} &= \alpha_{1it} + \beta_1 X_{1it} \dots \beta_n X_{nit} + \mu_{it} \\ t &= 1 \dots t \\ i &= 1 \dots i \end{aligned} \quad (7)$$

These advances imply several advantages for social scientists. The estimated coefficients and the level of significance of these estimates may not be so different than the OLS estimates. Similarly, p values and t coefficients may either be controversial or different from previous OLS calculations. However, some advantages of PD models over the OLS estimation are of particular interest to researchers. First, by allowing variation in the intercept, more information is provided about the characteristics across cases, and it is possible to detect potential differences in the intercepts that may be due to unique features of the cases and their contexts. This analytical tool is critical for social investigators that attempt to isolate common implications of contextual effects on social behavior or phenomena.

Second, there is a potential improvement in the level of significance due to the fact that the explained variation and consequently the standard deviations are more accurate and properly estimated than the pooled OLS version. This analytical instrument is important for social scientists that attempt to properly estimate the influence of different indicators presumed to represent multi-faceted structures or components of social reality. In other words, FEM models allow for controlling certain aspects of social reality that interact with the study subjects in a more efficient way.

Third, the R^2 and the Durbin-Watson statistic values, as a consequence of a better estimation of the explained variation, may increase substantially, thereby suggesting evidence of misspecification of models. This tool proves critical for evaluating a model's specifications and hypothesized variables as influential and therefore necessary to be included in the model. Consideration of the R^2 is not enough on its own due to the addition of dummy variables in the model (3) or the variation in the intercept across cases and time in the model (6). Using a formal F -test can better compare the efficiency among models. In other words, the F -test is used to compare unrestricted versus restricted models. A traditional OLS specification may serve as the "restricted" model (R) because it imposes the common intercept across cases. The "unrestricted" models (UR) may be LSDV or another FEM. The F -test uses the following formula (8) for this evaluation (Baltagi, 2008; Gujarati, 2003; Park, 2011; Wooldridge, 2002):

$$\begin{aligned} F &= [(R^2_{UR} - R^2_R)/(k - 1)] / [(1 - R^2_{UR})/(n - l)] \quad (8) \\ k &= \text{total number of additional variables added in } UR \text{ model} \\ n &= \text{total number of observations} \\ l &= \text{total number of variables in the } UR \text{ model} \end{aligned}$$

In any case, analysts should evaluate and test the level of efficiency of each model (using both the R^2 and F -test), but first it is critical to consider the assumptions of the type of effects, in terms of the cases, time, or both, that prevail in the data set or that need to be evaluated in the study. According to these assumptions, the logical next step is to incorporate feedback into the specification of the model for better fit. This evaluation represents a powerful tool for social researchers that need guidance for their theoretical representations, operationalization of variables, and methodological instruments applied in research.

1. 2. Random Effects Models (REM)

REM or error components models (ECM) are based on the idea that the intercept (α_{1i}) is a random variable with a mean value of α_1 (without subscript i) and the error term (ε_i) with mean value of zero and variance of σ_ε^2 expressing the following:

$$\alpha_{1i} = \alpha_1 + \varepsilon_i \quad i = 1 \dots i \quad (9)$$

The assumption behind this specification is that all cases are drawn from a larger universe and that they have in common a *grand mean* value for the intercept (α_1). Individual differences are reflected in the error term (ε_i). By substituting (9) on (4).

$$\begin{aligned} Y_{it} &= \alpha_1 + \beta_1 X_{1it} \dots \beta_n X_{nit} + \varepsilon_i + \mu_{it} \\ &= \alpha_1 + \beta_1 X_{1it} \dots \beta_n X_{nit} + \omega_{it} \\ t &= 1 \dots t \\ i &= 1 \dots i \\ \text{where:} \\ \omega_{it} &= \varepsilon_i + \mu_{it} \end{aligned} \quad (10)$$

The two components of the composite error term are ε_i that refers to the individual-specific error (cross-section), and the μ_{it} that refers to the combined time series and cross-section error. The assumptions of the composite error term ω_{it} are as follow:

$$\begin{aligned} \varepsilon_i &\sim N(0, \sigma_\varepsilon^2) \\ \mu_{it} &\sim N(0, \sigma_\mu^2) \\ E(\varepsilon_i \mu_{it}) &= 0 \quad E(\varepsilon_i \varepsilon_j) = 0 \quad (i \neq j) \\ E(\mu_{it} \mu_{is}) &= E(\mu_{it} \mu_{jt}) = E(\mu_{it} \mu_{js}) = 0 \quad (i \neq j; t \neq s) \end{aligned} \quad (11)$$

Therefore

If $E(\varepsilon_i) = 0$, then the $\text{var}(\omega_{it}) = \sigma_\varepsilon^2 + \sigma_\mu^2$

In general, these assumptions suggest no correlations between error terms and no autocorrelation across both cross-section and time series units. The ε_i is latent, but unobservable. That advance over FEM specifications allows social researchers to estimate general baselines and variation of the variables under examination and to isolate social phenomena from complex and spurious relationships due to other effects. This tool is critical for social scientist that usually battle to control their study subjects from anything other than what it is specified in the model.

2. An Illustrative Example: Revisiting the Michigan Experience

For illustrative purposes, this article uses a dataset originally used by Holzer *et al.* (1993) to show the main characteristics of two PD models. The estimations of the PD statistical methods are compared with the results of Pooled OLS for demonstration purposes. Using the longitudinal dataset of training subsidies for Michigan firms, called the Michigan Experiment, this study examines the effects of a state government grant program for job training in the actual hours of training per employee (as a proxy of effectiveness), the productivity of the firm, and the salary of employees.

The dataset is considered a PD arrangement and was collected using a survey across various manufacturing firms over three years. This section includes two parts. The first part includes a brief description of the Michigan Experiment and some relevant descriptive statistics of this dataset.

The second subsection involves the comparative analysis between pooled OLS and Panel Data estimations using the original hypotheses tested by Holzer *et al.* (1993) and implemented using SAS. Each section discusses advantages and limitations of applying PD methods.

2.1. The Michigan Experience: a Brief Description

The original work of Holzer *et al.* (1993), called the Michigan Experiment, is used in this article as an example to apply PD models and to highlight the differences between Pooled OLS and PD techniques in social science research. The data comes from the original questionnaire of the Michigan Experiment used by Holzer *et al.* (1993). The questionnaire was given to representatives of the firms that

applied for the grants in the Michigan Job Opportunity Back-Up-Grade Program (MOJB) during the years 1987-1989. The grants were designed to cover the direct costs of training. In this period of operation, the MOJB administered over 400 grants to firms in the state, with an average grant amount of \$16 000 each. The questionnaire obtained a response rate of 39%, which is considered an acceptable rate for this type of method (Bryman, 2004). The data selected for this study includes a total of 471 observations from 157 firms over 3 years.⁵

Holzer *et al.* (1993) were concerned about two potential sources of bias selection in the questionnaire: the potential auto-selection bias from the firms that applied to the grants and whether the person who responded to the survey was representative of the firm. According to Holzer *et al.* (1993), both potential biases were not significant in overstating or understating the responses. The questions asked include the number of formal job training hours, the reasons for training, the quantity of employees in job training, and some productivity measures such as scrap rate and rework rate. The levels of sales, employment, wages, union membership, and labor relation practices at the firm were also covered.

Table 2 shows the descriptive statistics of the variables examined in this paper for all firms and for the two potential categories of firms (those that received grants and those that did not). In order to analyze the level of potential auto-selection bias, Table 3 compares over time the estimates for all three groupings. The results show significant differences in all variables between firms that received grants and firms that did not receive grants. These differences are evidence of

Table 2. Descriptive statistics of dependent and independent variables (Std. Dev.).

Variables	All Firms across Years	Firms That	
		Received Grants	Did Not Receive Grant
Number of Observations	471	66	405
A. H. T. E.	14.97 (25.71)	42.94 (37.03)	9.98 (19.32)
Scrap Rate (per 100)	3.84 (6.01)	3.13 (4.52)	3.99 (6.28)
Rework Rate (per 100)	3.47 (5.46)	3.36 (3.93)	3.5 (5.76)
Annual Sales	6 116 327.0 (7 912 602.6)	6 651 157.4 (9 577 663.0)	6 035 442.1 (7 643 649.7)
Employment	59.32 (74.12)	70.82 (96.67)	57.39 (69.63)
Annual Average Salary	18 872.82 (6 703.16)	19 278.97 (6 845.45)	18 806.28 (6 687.28)
Union	0.1975 (0.40)	0.2424 (0.43)	0.1901 (0.39)

Note: A. H. T. E.: Annual Hours of Training per Employee.

Source: adapted from Holzer *et al.* (1993).

5. There is a difference in the sample size between the dataset we are using in this paper and the one used by Holzer *et al.* (1993). It is not clear why this difference exists, but this does not affect the illustrative purpose of the example.

the nested effects and the potential positive effects from applying PD models. Table 3 also indicates the annual hours of training per employee over the three-year period between firms that received grants and firms that did not receive grants. The results also point out important differences across the firms who received and did not receive grants over time. From this table it seems that grants have great influence on the annual hours of training per employee. On average, firms that receive grants did increase their hours of training compared with the firms that did not receive grants in a particular year. However, this increase is not a long-run effect. For example, firms that received grants in 1988 increased the amount of hours of training during that year but reduced them after in 1989. Finally, by using Holzer *et al.* (1993)'s study as an example, this article attempts to illustrate the main steps and methods for PD. We warn the readers that there are no critical statistics nor reliable datasets available from the original work of Holzer *et al.* (1993) that can be used for diagnosing potential heteroskedasticity and autocorrelation such as Durbin-Watson and Chow tests, among others (Drukker, 2003; Wooldridge, 2010).⁶ However, Holzer *et al.* (1993) tested autocorrelation using *F-test* in order to pool the data across years.⁷ Therefore, we do not include the results of these tests, but we strongly recommend these tests before conducting a PD analysis.

2. 2. Panel Data Estimations

As in the original study from Holzer *et al.* (1993), this article uses the data from the Michigan Experience to explain the impact of governmental grants in the hours of job training per employee, the productivity of the firm, and the salary of employees. Social researchers may hypothesize direct effects between government grants on job training, productivity of the firm, and salary of employees across time and compare these effects between two categories of firms that received

or did not receive grants. We decided to use the same three hypotheses as in the original study:

a) There is a direct positive relationship between government grants for job training and the actual number of hours of training per employee that firms offer (H_1).

b) There are positive indirect effects of this training in the productivity of the firm and the wages growth of their employees. So, two hypotheses are related to this increase of training:

- Training is positively related with some measures of productivity of the firm (H_2).

- Training has a positive effect on average salary of employees (H_3).

Each of these hypotheses may represent different lines of research, research questions, and model specifications that can be tackling using PD models. In order to statistically test our hypotheses, we estimate our explanatory models using four different specifications: Pooled Ordinary Least Squares (OLS), fixed effects (FEM), random effects (REM), and something that we called expanded fixed effects (XFEM), which are fixed effects models that take into consideration other variables that vary over time. Other characteristics of firms are also included in the analysis.

2. 2. 1. Job training grants and training per employee

Following the original paper, some regression equations were re-estimated to formally test for the effects of receiving grants on job training, productivity (represented by scrap rate), and wages growth (represented by the annual average salary of the employees). The first hypothesis is represented by the equation (12) that portrays the relationship between hours of job training per employee, grants, and other firm characteristics:

$$\Delta TR_{jt} = a + a_1 \text{GRANT}_{jt} + a_2 X_{jt} + u_{1jt} \quad (12)$$

where j and t represents the firm and time, respectively; TR denotes the hours of training per employee (ΔTR represents the log version of TR); GRANT is a dummy variable representing whether or not the firm received a grant; and the X variables are controls such as sales, employment, union membership, and individual characteristics of the firms that are going to be captured in the fixed and random effects models. Table 4 presents the results

Table 3. Hours of training per employee, by year and receipt of grant (Std. Dev.).

Description	1987	1988	1989
Firms receiving grants in 1988	7.67 (19.58)	37.98 (36.95)	10.06 (19.47)
Firms receiving grants in 1989	6.12 (11.31)	9.32 (17.36)	50.65 (36.22)
Firms that did not receive grants	8.89 (17.40)	9.67 (18.18)	11.59 (22.38)

Source: adapted from Holzer *et al.* (1993).

6. For more information, see Stata procedures to test for heteroskedasticity and autocorrelation in panel data models:

<http://www.stata.com/support/faqs/statistics/panel-level-heteroskedasticity-and-autocorrelation/>.

7. See discussion in the first column and footnote 24 in p. 632 about pooling the data by years in Holzer *et al.* (1993).

of the pooled OLS regression and the FEM and REM estimations for the relationship between hours of job training per employee, grants, and other firm characteristics.

From these results it is possible to identify across our four different specifications that a government grant has a statistically significant coefficient. In the pooled OLS and REM specifications, government grants seem to duplicate the amount of training. In both FEM and XFEM models government grants have a greater effect on the amount of training, almost triple. In OLS estimates, some specific characteristics of the firms underestimated the actual effect of the grant on training. This underestimate is an indication of structural differences across cases and therefore FEM computations seem to be more appropriate. However, most of the variables capturing some characteristics of the firms are not significant, with the exception of "sales" that is statistically significant at the five percent level.

The FEM, REM, and XFEM models present higher R^2 than the Pooled OLS estimation, which indicates improvements in the goodness of fit of the model by using Panel Data models. As mentioned before, the FEM and XFEM models present high R^2 due to the large number of dummy variables used for its computations (one dummy variable for each year). This is a potential disadvantage for social researchers that attempt to claim higher levels of robustness and soundness of their findings.

By applying the Hausman test, it is possible to infer that the FEM and REM models are equally good. However, the difference in the estimated value of government grants between them is not small (12 and 4, respectively). There is also an important advantage for social investigators that look to evaluate the direct effects of a set of variables of interest at a particular time by controlling for other contextual or time effects.

By using dummy variables instead, social researchers need to be cautious in interpreting the intercept across models. In the Pooled OLS, FEM, and XFEM models, dummies are used over years. In particular, the original study from Holzer *et al.* (1993) considers the year 1987 as the baseline for the analysis. However, researchers need to select a year with average behavior of the variables under study in order to

prevent any possible bias in the results caused from contingencies or events happening at that time. This sensitivity is an important limitation of traditional methods using OLS estimates. In our example, there is no indication that the year 1987 potentially created this type of bias, but researchers do need to pay attention to these issues.

2. 2. 2. Grants, training and productivity in the firm

For the second hypothesis, the original study from Holzer *et al.* (1993) specified the model suggested in equation (13). This second equation represents the relationship between the effect of job training on productivity of the firm (measured by using the scrap rate - SR) and some other variables of the firms:

$$\Delta SR_{jt} = b + b_1 TR_{jt} + b_2 X_{jt} + u_{1jt} \quad (13)$$

where j and t represents the firm and time, respectively; TR denotes the hours of training per employee; SR is the scrap rate of each firm (ΔSR represents the log version of SR); and the X variables are controls such as sales, employment, union membership, and other individual characteristics of the firms. Table 5 presents the results of the pooled OLS regression and the FEM and REM estimations for the relationship between productivity of the firm, hours of job training per employee, and other firm characteristics.

Table 4. $\Delta \log$ (Hours of Job Training per Employee).

Variable	Pooled OLS	Fixed Effects Model	Random Effects Model	Expanded Fixed Effects Model
Constant	-2.918 (-1.005)	0.143 (0.153)	-0.0738 (-0.82)	-4.716 (-0.2)
Job Training Grant	2.297*** (10.156)	3.098*** (12.886)	2.418*** (13.10)	3.121*** (10.199)
Dummy for 1988	—	0.228 (1.372)	—	0.129 (0.492)
Dummy for 1989	-0.178 (-0.948)	—	—	—
Log (Sales)	0.347** (2.027)	—	—	0.682 (1.278)
Log (Employment)	-0.251 (-1.356)	—	—	-0.616 (-0.743)
Log (Average Salary)	-0.143 (-0.514)	—	—	-0.33 (-0.14)
Union	-0.0275 (-0.107)	—	—	0.365 (0.248)
R^2	0.368	0.677	0.411	0.655
F-statistic or Hausman Test	18.312***	2.00***	19.86***	1.662***

Note: T-statistics are in parenthesis under coefficient values. The level of significance is noted by * for 10%, ** for 5%, *** for 1%.

Source: adapted from Holzer *et al.* (1993).

Despite the fact the results for the FEM model show some effect of training on the scrap rate, the increase in hours of training per employee does not represent a strong determinant of productivity of the firm. These results confirm the idea of important differentials due to contextual or inherent business practices across firms or over time. In this way, PD models prove to be useful for social researchers that precisely look for potential controls for the effects of time and categories of units embedded in the social phenomena under scrutiny. In particular, it seems that some effects of past training are stronger than the present training. In the FEM model, we can see that the amount of training in the last year presents a better level of statistical significance on the scrap rate (ΔSR) than the present year. This finding might tell us that training has an inter-temporal behavior on the scrap rate that a simple OLS specification may fail to capture. The training provided during the year has a different effect on the scrap rate than the grant received last year. In other words, a maturity effect of past training is present that impacts current levels of productivity.

No control variables are significant across models, which is another advantage for social researchers that also investigate other potential contextual effects or spurious relationships in the model. The R^2 across models improved by using PD techniques. For FEM and XFEM compared to

the pooled OLS and REM models, there is an improvement of R^2 , an indication of the presence of critical differentials across firms or over time. Again, this improvement benefits social researchers who evaluate model fit and the adequacy of model specification or set of variables included in the models. In particular social researchers may apply Hausman tests to evaluate different model specifications. In our example, the Hausman test shows us that FEM specification presents a better fit due to these differentials that are not random.

2. 2. 3. Grants, training, and the annual average salary of employees

The last hypothesis shows the effect of job training on wages growth and some other variables of the firms. The equation (14) applied PD techniques in the original study of Holzer *et al.* (1993), which is as follows:

$$\Delta WW_{jt} = b + b_1 TR_{jt} + b_2 X_{jt} + u_{1jt} \quad (14)$$

where j and t represents the firm and time, respectively; TR denotes the hours of training per employee; WW is the annual average salary; and the X variables are controls such as sales, employment, union membership, and other

individual characteristics of the firms. Table 6 shows the resulting estimates from the different specifications. As we can see, it appears that the effect of training on wages is mediated by differences in the effects across firms and over time. In both the FEM and XFEM models the coefficients of the lagged hours of training per employee are significant and have a negative sign. This situation suggests that training at first has a positive effect on wage growth in periods of government grants, but over time this effect “normalizes”. According to Holzer *et al.* (1993), this finding corroborates the durability of the effects of training over time found in previous studies (Brown 1990; Osterman, 1988; Lillard and Tan, 1986; Lynch, 1989; Mincer, 1989). In the REM model, only the coefficient of hour of training per employee is statistically significant, suggesting the presence of effects that are random across firms and over time.

Table 5. $\Delta \log$ (Scrap Rate in Michigan Manufacturing Firms).

Variable	Pooled OLS	Fixed Effects Model	Random Effects Model	Expanded Fixed Effects Model
Constant	0.0824 (0.032)	-0.000123 (0.00)	-0.156 (-1.38)	1.062 (0.079)
Log (Hour of training per employee)	-0.0531 (-1.159)	0.0749 (1.104)	-0.04 (-0.95)	0.130 (0.155)
Log (Lag hour of training per employee)	0.0432 (0.878)	0.220* (2.012)		0.136 (0.992)
Dummy for 1988	—	0.3** (2.084)		0.179 (0.978)
Dummy for 1989	-0.102 (-0.797)	—		—
Log (Sales)	0.0396 (0.234)			-0.238 (-0.366)
Log (Employment)	-0.0608 (-0.325)			-0.067 (-0.077)
Log (Average Salary)	0.0683 (-0.325)			0.338 (0.293)
Union	0.0629 (0.397)			
R^2	0.036	0.684	0.01	0.63
F-statistic or Hausman Test	0.415	1.895**	0.26	1.276

Note: T-statistics are in parenthesis under coefficient values. The level of significance is noted by * for 10%, ** for 5%, *** for 1%.

Source: adapted from Holzer *et al.* (1993).

The dummy variable for the following year (1988) is significant; perhaps because the data for salaries in Holzer *et al.*'s (1993) original work is not deflated. However, the sign is negative, suggesting some erosion effect over the durability of training over time. In any case, the evidence shows the presence of contextual or individual differences across cases and over time.

In the REM and Pooled OLS models the R^2 are small. This situation suggests that there are no random effects that on average influence the estimations. On the contrary, the presence of individual specific conditions or maturity effects over time plays an important role across firms. In the FEM and XFEM models the R^2 are high, indicating that these models are more robust than the Pooled OLS and REM models. However, it is important to consider that the FEM and XFEM models use dummy variables for each of the firms. As a result the R^2 may be potentially higher than the other models due to the number of degrees of freedom. In contrast to this potential effect of dummy variables, the Hausman test indicates that the FEM and XFEM specifications are more robust than the REM and Pooled OLS models.

The estimate for sales is statistically significant in the pooled OLS model, which was also found significant in the first relationship between government training and the number of training hours per employee. This finding simply confirms the logical path from the aid provided by government through training grants to the increase of job training per employee and its derivative expected impact of wage growth. Union membership is also significant in the XFEM model. This finding is negative, representing the potential specific individual differences between unionized workers and non-unionized employees. These differences may potentially be translated at the level of firms, where some firms have more unionized workers than other firms. Holzer *et al.*'s (1993) original paper does not involve further analysis, but these questions are open for future research. In any case, the use of PD models raises these differences as significant. In this case, union membership explains certain parts of the difference in salary across units and over time that are not significant for the number of hours of training per employee or for the productivity of the firm.

3. Discussion and Implications

This section summarizes the results of this paper and presents some concluding remarks and suggestions for future research within this topic. First, it highlights the general results and then elaborates on the comparison between the two PD techniques used in this paper. Finally, this section provides some implications for social science research.

3.1. General results

In general, PD methods seem to be more effective in correcting autocorrelation, deflated standard errors, and underestimated degrees of freedom. In our example, FEM and REM estimations were more robust than simple pooled OLS results. Pooled OLS estimation disregard the effects over time and across firms. Therefore, the pooled OLS results did not properly depict the relationships of the variables being studied across cases and over time.

Our example shows how OLS estimates underestimated some of the effects that the characteristics of the firms had on how the grant impacted training, and consequently the effects of this training on the productivity of the firm and employee's salary. In other words, the pooled OLS model did not capture the presence of differences across cases or over time, which is caused by the potential autocorrelation or heteroscedasticity risks in piling up the data. This produces biased estimations of variances, standard

Table 6. Annual Average Salary of Employees.

Variable	Pooled OLS	Fixed Effects Model	Random Effects Model	Expanded Fixed Effects Model
Constant	18 388.87*** (19.565)	18 133.57*** (20.132)	18 326.73*** (29.61)	18 226.88*** (14.89)
Hour of training per employee	15.68 (0.819)	4.093 (0.7)	20.77*** (3.74)	6.765 (1.016)
Lag hour of training per employee	-6.063 (-0.250)	-23.82** (-2.328)		-24.58** (-2.180)
Dummy for 1988	—	-1 414.51*** (-7.409)		-1 467.13*** (-6.622)
Dummy for 1989	1 298.42 (1.298)	—		—
Sales	0.000407*** (4.077)			0.0000812 (1.254)
Employment	-43.35 (-3.388)			-9.079 (-0.784)
Union	-179.027 (-0.133)			-2 864.02** (-2.141)
R^2	0.095	0.985	0.038	0.984
F-statistic or Hausman Test	3.315***	58.9***	0.01	54.213***

Note: T-statistics are in parenthesis under coefficient values. The level of significance is noted by * for 10%, ** for 5%, *** for 1%.

Source: adapted from Holzer *et al.* (1993).

deviations and coefficients that eventually affect the inferences of intercept and coefficients and their corresponding levels of significance.

In the case of the first hypothesis, the results from the pooled OLS model present statistically significant coefficients with expected signs (in particular sales) as in FEM and REM estimates. However, by using PD models social science researchers have some advantages. First, FEM and REM results are more accurate and robust than OLS estimates without losing the level of significance. Second, R^2 values for FEM and REM are higher than the pooled OLS value. In this case, the FEM model seems to be more appropriate for controlling these effects between two different groups of firms (those that received the government grant for training and those that did not). The FEM, REM, and XFEM models in this example presented higher R^2 than the Pooled OLS estimation, an indicator of the improvements in the model's goodness of fit by using PD models. Another advantage of applying PD models in this case, either FEM or REM, is that variables capturing some of the firm's characteristics over time were considered while maintaining the levels of significance and the model's goodness of fit. Pooled data and models using dummies are usually also a disadvantage in terms of degrees of freedom. In these cases, PD models also benefit social researchers by using the full range of information contained in the covariance matrix that is necessary for computing estimates and maintaining the soundness of their findings. Hausman tests may be used to evaluate models and effects between variables, which is important for social science researchers whose goal is to evaluate the effects of a set of variables across time or units, since OLS estimates need additional steps to control for possible bias of time or characteristics of units.

In the case of the second hypothesis, the results of the pooled OLS model are not statistically significant. In this case it is clear that the FEM model is more robust and its estimated coefficients are statistically significant. Again, the advantage of using FEM in this case was important for exploring potential effects of firms or potential differentials across firms over time. In the case of the third hypothesis, both FEM and REM presented statistically significant estimates for coefficients and accurate standard deviations that allow for better inference work. Only the intercept and the coefficient for sales were significant in the pooled OLS model and these computations could be biased because of the pile up of data about firms over time. The R^2 values for the FEM and REM models improved considerably from previous versions. In both the FEM and REM models the coefficients were significant and had negative signs.

PD estimations help researchers to properly estimate the durability of the effects of training over time, as well as to estimate the characteristics of the firms in terms of the levels of training, sales, employment, and union membership.

3. 2. Comparison between FEM and REM

Gujarati (2003) and Wooldridge (2002) list some differences between the FEM and REM approaches. These differences were corroborated in our example. In FEM, each category of cases has its own "fixed" intercept value across cases. In our example, the effect of time on grants was captured through year dummies for 1988 and 1989. In addition the effects of some characteristics of the firms in terms of volume of sales, employment, and union membership were controlled using the FEM technique. In REM, the intercept represents the *grand mean* of all cases and the ε_i captures the random deviation of the mean value as the individual differences across cases. If $\sigma_\varepsilon^2 = 0$, then there is an indication of no differences across cases (cross-sectional). In our example this random deviation is significant, which indicates that there are important differences across cases. Therefore, a pooled regression specification is not suitable for the analysis by simply combining all cross-sectional and time series data into one dataset.

In order to evaluate the differences and similarities between FEM and REM, a researcher can apply the Hausman test. Another method is to evaluate the correlation matrix of error terms ω_{it} . If error terms ω_{it} are correlated with each other, then the most appropriate method is REM because it allows for a random variation over time. In our example the first method was enough to consider the FEM model as the most suitable PD technique for the dataset.

Researchers may apply other methods to decide between the FEM and REM approaches. The type of approach depends on the assumptions about the correlations between cross-section specific error terms, ε_i , and the regressors. If ε_i and X 's are not correlated, REM may be appropriate; otherwise, FEM may be more suitable. The main assumption of REM is that the ε_i is drawn randomly from a large population, but this may not happen in the data. So, Judge *et al.* (1985) suggest additional criteria. If the number of time series data is large and the number of cross-sectional units is small, there will likely be similar estimations between REM and FEM. In the contrary case (large cross-sectional units and small time series data), the estimates between REM and FEM will differ. This is the case of our example where $n = 471$ firms and only two points in time are available (years 1988 and 1989).

If the assumption of differentials is stronger, then FEM is more appropriate. If the sample is still regarded as random, then REM is more appropriate. In our example, there is evidence of these differentials between firms that received and did not receive grants (see Table 2). In this case, the Hausman test also serves to test the null hypothesis that the FEM and REM estimates are significantly similar (Hausman, 1978). If the null hypothesis is rejected, then FEM is more appropriate (see the *F-test* in formula (8) for this evaluation).

In any case, researchers should pay attention to three critical components of their study: *a)* the objective of the study established by the research question; *b)* the theoretical or conceptual framework that hypothesizes causal relationships; and *c)* the nature of the data or indicators that operationalize the variables in this framework. By considering these components, researchers will evaluate the model that better fits their investigation needs.

Conclusion

Social phenomena are complex and multidimensional. Traditional statistical and regression techniques have proved to be limited and vulnerable to certain violations to its assumptions. Panel data (PD) methods attempt to make a more useful and robust analysis for scientific work. PD methods have made important technical progress within the last few years and they are increasingly being used in many social sciences. In the future, the advantages that they provide will surely make them the method of choice for the analysis of several complex social problems. The purpose of this paper was to illustrate the steps and procedures for conducting PD analysis in the context of social science research.

The paper articulated an example of how PD methods can be more effective in correcting for autocorrelation, deflated standard errors, and underestimated degrees of freedom. The FEM and REM estimations applied in the example proved to be more robust than OLS results. In sum, using PD models in the example showed some advantages: *a)* for applying inferential statistics, FEM and REM estimates were more accurate and robust than OLS estimates without losing statistical significance; *b)* R^2 values for FEM and REM were higher than the pooled OLS values, demonstrating some improvement of the model's goodness of fit by using PD models; and *c)* variables capturing some of the firm's characteristics over time were considered, maintaining the levels of significance and the model's goodness of fit.

Once the benefits of applying PD techniques were shown, researchers were informed about different scenarios for applying PD techniques. Section 3.2 described the differences between FEM and REM for social science research. In particular, researchers should consider three aspects for deciding which PD technique would be the most appropriate: *a)* the objective of the study established by the research question; *b)* the theoretical or conceptual framework that hypothesizes causal relationships; and *c)* the nature of the data or indicators that operationalize the variables in this framework.

Finally, contemporary social science research usually demands the analysis of panel data consisting of multiple units observed at multiple time periods. Today the availability of large data sets, including cross-sectional and time-series, and recent developments in computing have pushed towards the combination of sophisticated statistical techniques like panel data analysis. These techniques have become increasingly popular and important analytical tools in social and behavioral sciences and we think they will become even more important in the near future. However, the application of PD models in social research needs to be discussed and researchers should be informed about its benefits and challenges. For instance, these techniques have the potential to allow more accurate and robust prospective analysis when compare with multiple regression or other less sophisticated statistical techniques.

This article provides a discussion on PD techniques, focusing on two models of panel data analysis (FEM and REM). Variations of panel data models, such as pooled cross sections or pooled time-series, were also briefly discussed. This article includes an empirical example in order to illustrate the use of PD models and techniques. However, there is the need to develop more examples across different social science disciplines to more deeply explore the advantages and limitations of PD techniques and future research should do this. It is also necessary to continue this discussion along with its corresponding implications and guidance for statistical software and procedures capable of performing panel data analysis. In the near future, more technical advances and guidance will surely exist and be available for interested researchers. Similar attempts have been conducted with other sophisticated techniques in particular fields of study. For example, in the field of electronic government several techniques have been explained in recent publications with the purpose of informing and guiding social science researchers. Some examples of these efforts are Gil-Garcia (2008) for Structural Equation Modeling and Partial Least Squares techniques or Gil-Garcia *et al.* (2009) for Web-based surveys.



- Baltagi, B. H. (2008). *Econometric analysis of panel data*. New York: John Wiley and Sons.
- Bickel, R. (2007). *Multilevel analysis for applied research. It's Just Regression*. New York, U.S.: The Guilford Press.
- Bond, S. (2002). *Dynamic panel data models: a guide to micro data methods and practice*. The Institute for Fiscal Studies, Department of Economics, UCL-cemmap working paper CWP09/02.
- Brown, C. (1990). Empirical evidence on private training, in Roland Ehrenberg (ed.) *Research in labor economics*. Greenwich, Conn.: JAI Press.
- Bryman, A. (2004). *Social research methods*. New York, U.S.: Oxford University Press.
- Cameron, A. C. & Trivedi, P. K. (2009). *Microeconometrics using stata*. TX: Stata Press.
- Drukker, D. M. (2003). Testing for serial correlation in linear panel-data models. *Stata Journal*, 3, 168-177.
- Gefen, D., Straub, D. W. & Boudreau, M. C. (2000). Structural equation modeling and regression: guidelines for research practice. *Communications of the Association for Information Systems*, 4.
- Gil-Garcia, J. R. (2008). Using partial least squares in digital government research, in G. D. Garson & Mehdi Khosrow-Pour (eds). *Handbook of research on public information technology*. Hershey, P. A.: IGI Global.
- Gil-Garcia, J. R., Berg, S. A., Pardo, T. A., Burke, G. B. & Guler, A. (2009). conducting web-based surveys of government practitioners in social sciences: practical lessons for e-government researchers. *42nd Hawaii International Conference on System Sciences*.
- Greene, W. H. (2012). *Econometric analysis* (4th ed.). Englewood Cliffs, N. J.: Prentice-Hall.
- Gujarati, D. N. (2003). *Basic Econometrics*, (4th ed.). New York, N. Y.: McGraw-Hill.
- Hair, J. F. Jr., Anderson, R. E., Tatham, R. L. & Black, W. C. (1998). *Multivariate data analysis*. Upper Saddle River, N. J.: Prentice Hall.
- Hausman, J. A. (1978). Specification tests in econometrics. *Econometrica*, 46, 1251-1271.
- Holzer, H. J., Block, R.N., Cheatham, M. & Knott, J. H. (1993). Are training subsidies for firms effective? The Michigan experience. *Industrial and labor relations review*, 46, 625-636.
- Hsiao, C. (1986). *Analysis of panel data*. Cambridge: Cambridge University Press.
- Hsiao, C. (2006). Panel data analysis-advantages and challenges, IEPR Working Paper 06.49. Institute of Economic Policy Research, University of Southern California.
- Judge, G. G., Hill, R. C., Griffiths, W. E., Lutkepohl, H. & Lee, L. C. (1985). *Introduction to the theory and practice of econometrics*, (2nd ed.). New York: John Wiley & Sons.
- Lillard, L. & Hong, T. (1986). *Private sector training: who gets it and what are its effects?* Los Angeles: Rand Corporation.
- Lynch, L. (1989). Private Sector Training and Its Impact on the Earnings of Young Workers. *NBER Working Paper*. National Bureau of Economic Research.
- Mincer, J. (1989). Job training: costs, returns, and wage profiles. *NBER Working Paper*. National Bureau of Economic Research.
- Osterman, P. (1988). *Employment futures*. New York: Oxford Press.
- Park, H. M. (2011). *Practical guides to panel data modeling: a step-by-step analysis using stata*. Tutorial Working Paper. Graduate School of International Relations, International University of Japan.
- Pindyck, R. S. & Rubinfeld, D. L. (1998). *Econometric models and economic forecasts*, Boston: Irwin McGraw-Hill.
- Wooldridge, J. M. (2010). *Econometric analysis of cross section and panel data*. Cambridge: MIT Press.

