

Neighborhood and School Effects on Children's Development and Well-Being*

Jinsook Kim,[†] Alair MacLean,[‡] Anne Pebley,[†] and Narayan Sastry[‡]

[†] UCLA, Los Angeles, CA

[‡] RAND Corporation, Santa Monica, CA

March 2006

Draft, do not quote or cite

Version: March 29, 2006

* The authors gratefully acknowledge support for this research from the National Institute of Child Health and Human Development and the Russell Sage Foundation. Assistance from Chris Peterson, Afshin Rastegar, Claude Setodji, and Stephanie Williamson is greatly appreciated.

1. INTRODUCTION

During the past twenty years, there has been extensive research examining the effects of social environments on many aspects of children's development. There are two, generally separate, literatures on the effects of social context on child development. One focuses on the effects of classrooms and schools (and sometimes the home environment) while the other centers on the effects of neighborhoods and families. Although a few studies have considered both neighborhood and school factors (Webster et al. 1996; Betts & Morell 1999; Card & Rothstein 2006), neighborhood and school effects tend to be examined in separate analyses. The most likely reason that these two literatures have not been integrated is the lack of appropriate data, given the substantial interest in this topic among researchers, policymakers, and the public.

Findings in both the school effects and neighborhood effects literatures are controversial and often contradictory. For example, in a review of the literature on schools, Hanushek (1997) concludes that there is no strong and consistent relationship between school resources (principally financial resources) and school achievement (primarily test scores). However, numerous studies have found significant effects of school-level socioeconomic status and social capital on children's reading and math test scores and on SAT scores (Parcel and Dufur, 2001; Everson and Millsap, 2004). Some studies have also found mixed results of school characteristics on educational attainment and subsequent earnings (Betts, 1996; Betts and Morell, 1999). The neighborhood effects literature suggests that these effects appear in preschool years, but are most consistent for school-age children and that neighborhood effects appear to be stronger for cognitive and achievement outcomes than for behavior and mental health measures (Duncan and Raudenbush, 1999; Pebley and Sastry, 2004). However, a number of studies have found little or no effect of neighborhood characteristics — particularly, poverty or concentrated disadvantage — on children's developmental measures (Pebley and Sastry, 2004; Sampson et al., 2002).

In previous research based on the Los Angeles Family and Neighborhood Survey (L.A.FANS) and the Panel Study of Income Dynamics Child Development Supplement (PSID-CDS), we examined how neighborhoods affected children's reading and problem solving test scores (Sastry et al., 2005; MacLean et al., 2005). These cognitive skills are important because they affect children's chances to do well during both childhood and adulthood. For example, Farkas et al. (1997) show that cognitive skills measured in adolescence and young adulthood significantly affect occupation and income years later, even when work experience, educational attainment, and other factors are held constant. Our analysis of neighborhood effects exploited the availability of multilevel data in both the L.A.FANS and the PSID, including information on multiple children per family and multiple families per neighborhood, as well as an extensive set of measures of parents' socioeconomic status, family background, and skill levels. We used these data and multilevel random effects models to account for unobserved family- and neighborhood level characteristics and to assess the correlation in test scores among siblings and among children from the same neighborhood. Our results for both Los Angeles County and the national PSID sample indicate that, while family membership accounts for a greater proportion of the variance in children's reading and problem-solving test scores, the neighborhood of residence also was significantly associated with these cognitive skills. In particular, we found that neighborhood median family income accounted for a significant portion of the variation in

children's test scores, even when an extensive set of family characteristics — including the mother's reading score — was held constant. Neighborhood income accounted for more of the variation in children's test scores than family income, assets, or mother's years of school. Only the mothers' reading scores were more strongly associated with children's reading and math scores.

Given these results and the likelihood that poor neighborhoods are more likely to have poorer quality schools than wealthier neighborhoods, an obvious question is whether the apparent effects of residential neighborhoods on children's skills acquisition are due to neighborhood conditions themselves or to the fact that children in lower income neighborhoods are more likely to attend poorer schools. More generally, there is a major gap in the neighborhood effects literature regarding the role of schools in shaping children's achievement and a similar gap in the school effects literature regarding the role played by neighborhood factors. The reason for these research gaps is primarily due to the lack of appropriate data. In particular, specific study design features are needed in order to separately identify school and neighborhood effects. The minimum requirement for neighborhood-based studies is for some children in the same neighborhood to attend different schools; similarly, for school-based studies, some children in the same school should live in different neighborhoods. However, the ideal design would have both of these features present — in other words, children would be *cross-classified* by neighborhood and school. An additional, and largely unexplored, question concerns the statistical power available for identifying separate school and neighborhood effects, even with an ideal design. A related issue is the degree to which school catchment areas are conceptually distinct from neighborhoods in the lives of children; if the two largely coincide, it may be difficult to identify separate school and neighborhood effects.

In this paper, we make a start at exploring these various issues. A primary goal is methodological: to assess the ability of researchers to separately identify school and neighborhood effects using data from L.A.FANS and PSID-CDS. These two data sets are in increasingly wide use by demographers and other social scientists to study children's achievement and have many design features that make them promising for separately identifying both school and neighborhood effects. A secondary goal is substantive: to extend our earlier analysis of neighborhood and family effects to examine neighborhood, school, and family effects simultaneously. We use cross-classified multi-level random effects models to assess the relative importance of neighborhood and school effects on children's test scores, holding family characteristics constant. The paper is organized as follows. In the section that follows, we provide an overview of the main conceptual and methodological issues that shape our analysis. In Section 3, we describe the L.A.FANS and PSID-CDS surveys and identify the independent variables and covariates. In the subsequent section we present results from preliminary models estimated using these data.

2. ANALYTICAL ISSUES

There are several important conceptual and methodological issues that shape the analysis of school and neighborhood effects on children's development. In this section, we provide an overview of these issues. When we describe and analyze the L.A.FANS and PSID-CDS data, we assess the strengths and weaknesses of these data sets according to these issues and, ultimately,

to their ability to accurately identify clear school effects that are separate from neighborhood effects.

There are four main obstacles to disentangling the effects of neighborhoods and schools from each other.

- The basic survey design needs to be appropriate for the statistical identification of neighborhood and school effects. First, there must be multiple children per school and per neighborhood. Without this basic requirement, the analyses will attribute what may be variation at the school or neighborhood level to variation at the individual level. Second, at least some children in the same neighborhood should attend different schools or some children in the same school should live in different neighborhoods. Otherwise, the analyses will confound the effects of school with those of neighborhood and vice versa. Variation according to only one of these dimensions may be sufficient. For example, if children in the same neighborhood attend different schools (which is more common in household-based surveys) then it is probably not necessary for children in any school in the sample to come from different neighborhoods. However, if neighborhoods correspond precisely with schools in the sample design (either because there is only one child in each neighborhood and no two children attend the same school, or because all sampled children from the same neighborhood attend the same unique school), then it will be impossible to identify separate school and neighborhood effects. In addition, there may be a threshold beyond which it is impossible to disentangle the effects of schools and neighborhoods, or, more generally, of different levels. For example, if 90 percent of neighborhoods contain only one child and 70 percent of schools contain only one child, this may not provide enough variance to identify the separate effects of the different levels.

Some survey design features help identify effects at one level but not the other. For instance, sampling multiple children per family makes it likely that siblings will attend different schools, which in turn allows for school effects to be identified; however, neighborhood effects cannot be identified unless there are multiple families per neighborhood. With one family per neighborhood, measured family effects will most likely represent the effects of both family and neighborhood. If only a single children per family is sampled, then it is not possible to identify family effects separately from individual effects — although, depending on other design features, both neighborhood and school effects could be independently identified.

In sum, it is crucial there are multiple families per neighborhood in the sample in order to identify neighborhood effects. It is also crucial that children from the same neighborhood attend different schools. Moreover, the variation in schools should be among similar school types (i.e., children of the same age attending different schools) rather than among children of disparate ages attending different schools (e.g., a younger child attending an elementary school and an older child attending a high school) because school effects may then in part reflect the effects of age.

- A major complication in both the school effects and neighborhood effects literatures is the fact that children are not randomly assigned to schools, classrooms, and neighborhoods (Duncan and Raudenbush, 1999; Raudenbush and Willms, 1995). Non-random assignment – also known as selection effects or the endogeneity of neighborhood choice – means that apparent neighborhood and school effects may be due to the fact that families and school administrators have a choice both in where children live and go to school and in whether resources should be invested in child development. Non-random assignment can create or obscure social contextual effects child development. Another related limitation has been failure to control adequately for family characteristics which may affect children’s development (Ginther et al. 2000). Thus, what appear to be neighborhood and school effects may actually be due to compositional differences in family characteristics among neighborhoods and schools.
- Homogeneity matters—if families are sorted into neighborhoods based on each family having similar characteristics, then it becomes difficult to separate neighborhood effects from family effects. At the extreme, if all families (or individuals) in a neighborhood share the same characteristics then, by definition, it is no longer a family (or individual) characteristic but rather a neighborhood characteristic.
- It is important to define what we mean by neighborhood. If schools are an important part of neighborhood social life, then schools may help to define a neighborhood. For example, social networks among children as well as among families may form outside of school based on connections initially made at the school. Although these social effects may now be distinct from what researchers would consider to be school effects, it may not be possible to identify them separately. If neighborhood size is in fact much larger than we expect, then two schools that we think may be drawing from different neighborhoods may in fact be drawing from the same one, which could lead to biased estimates of neighborhood effects and overstated standard errors. On the other hand, if neighborhoods are in reality much smaller than we treat them, then results may again be biased either up or down but in this case standard errors would be understated.

3. DATA

For this study, we used data from the Los Angeles Family and Neighborhood Survey (L.A.FANS) and the Panel Study of Income Dynamics Child Development Supplement (PSID-CDS). Each data set has a number of important advantages for this analysis. L.A.FANS was designed to support multilevel analyses with separate family, neighborhood, and school effects, and incorporates many of the design features discussed above. The survey sampled census tracts, which often include students from multiple public (and sometimes private) schools. By design, the data are geographically more highly clustered than in other surveys and there are thus multiple children per neighborhood and school, as well as multiple children per family. The PSID-CDS has the advantage of being nationally representative and allows us to determine whether our results for Los Angeles are representative more general national patterns. The PSID-CDS and L.A.FANS collected extensive data on families, while L.A.FANS also collected detailed information on neighborhoods. L.A.FANS used many of the items and measures from

PSID-CDS and the core PSID to facilitate comparison between the two data sets. We next describe each data set in more detail.

Los Angeles Family and Neighborhood Survey

The analysis uses individual- and family-level data from the 2000-2001 Wave 1 L.A.FANS, a sample of families in Los Angeles county. In households with children, one child was chosen at random from all household members age less than 18. Within each household with more than one child, a sibling was also selected at random. L.A.FANS-1 interviewed 3,090 households, 3,140 children, and 3,558 adults. Neighborhood-level information was obtained from tract-level data of the 2000 Census and linked to individual cases in the L.A.FANS data. Children age 3 and older completed the subsets of Woodcock Johnson-Revised (WJ-R) standardized assessments (Woodcock & Johnson, 1989). The school data were linked to individual cases in the L.A.FANS data using school information (name and address) collected in the survey.

Panel Study of Income Dynamics-Child Development Supplement

The Panel Study of Income Dynamics (PSID) is a nationally-representative longitudinal survey that has been conducted since 1968. The baseline survey included a representative sample of 3,000 families and an additional sample of 2,000 low-income families. All of these families, including splitoffs, were followed in subsequent waves. The outcome measures and most of the independent variables are drawn from the second wave of the Child Development Supplement (CDS), which was conducted in 2002. In 1997, PSID respondent families were targeted to provide information for wave 1 of the CDS if they had at least one child who was under 13. There were 2,705 households eligible for inclusion in the CDS, with 2,394 responding (a response rate of 88 percent). If the eligible families had one child, that child was interviewed. If they had two children, both children were interviewed. If they had three or more children, two children were randomly selected for interview. As in the L.A.FANS, the PSID child respondents three years of age and older and their mothers completed subtests of Woodcock-Johnson Revised standardized assessments (Woodcock and Johnson, 1989) to directly assess reading and mathematics skills. As in the L.A.FANS, individual data are linked to census tract and school identifiers.

The current set of analyses is based on a sample with some missing data. In future work, we plan to correct this problem. There are a total of 2,630 children in the 2002 wave of the CDS who had been included in the first wave in 1997. The current set of analyses focus on individuals for whom the PSID provides information about schools. The school data are only provided for children who attended public school. In March 2006, PSID made available data about the characteristics of private schools. However, the following analyses utilize NCES identifiers to account for clustering, and these have not yet been made available for private schools. There are 468 cases (17 percent) that do not include school data. Unfortunately, the PSID data do not identify the reason that school data are not included. Therefore, the school data may be missing because the children went to a private school or because the information about the schools was not ascertained. Because we intend to conduct analyses with covariates (results not in the current version of the paper), we delete cases missing data on a variety of independent

measures. The following analyses are based on a final sample of 1,907 children. This analytic sample represents 72 percent of the full sample, or 88 percent of the cases for which school data are available.

Dependent Variables

The dependent variables in this study are test scores on the Woodcock-Johnson- Revised Tests of Achievement, specifically on reading and problem solving or math-related skills. Both surveys administered the same three tests. The Letter-Word Identification test assesses symbolic learning (matching a picture with a word) and reading identification skills (identifying letters and words). The Applied Problems test measures the subject's skill in analyzing and solving practical mathematics problems, and provides an assessment of mathematics reasoning. Older children, as well as primary caregivers, also completed the Passage Comprehension test. The Passage Comprehension test includes items that require the subject to point to the picture represented by a phrase and items in which the subject reads a short passage and identifies a missing key word. We ultimately plan to link the two waves of CDS data. Therefore, in order to include the broadest age range of children in the analysis, we use the Letter-Word Identification test to assess children's reading skills. We use the Applied Problems test to assess math skills. For both L.A.FANS and PSID-CDS, raw scores were converted to standardized scores and percentile ranks based on the subject's age and a set of national norms (McGrew, Werder, and Woodcock, 1989). Table 1 presents the means and standard deviations for both tests in both surveys.

4. METHODS

Cross-classified models have been fitted using frequentist or Bayesian approaches. The most well-known frequentist approach is a likelihood-based method using iterative generalized least squares (IGLS) algorithm that involves transforming the cross-classified model into a constrained nested model (Browne et al., 2001). According to Rasbash and Goldstein (1994), cross-classified models can be fitted using a pure hierarchical formulation (Goldstein, 2003). For a data with students (level-1 unit) cross-classified by schools and neighborhoods (level-2), the classification with the larger number of units (eg, school) is specified as a standard hierarchical level-2 classification. Next, for the other classification (eg, neighborhood), a dummy explanatory variable (0, 1) for each unit is created to indicate whether the observation belongs to that unit or not. Then, each of these dummy variables is set to have a coefficient random at level 3 with equal variances. The level-2 variance is the sum of the separate classification variances. The covariance for two level-1 units in the same classification is equal to the variance for that classification and the covariance for two level units that do not share either classification is zero. If there is a third classification at level 2, then the third variance is obtained by defining a similar set of dummy variables with coefficients varying at level 4 and variance constrained to be equal (Goldstein, 2003; Rasbash & Goldstein 1994; A User's Guide to MLwiN). However, this approach has limitations in handling data with multiple cross-classifications and large numbers of units due to its great memory demand for intensive computation and numerical instability of the constraining procedures (Browne et al., 2001; Rasbash & Browne 2004; Simonite & Browne, 2003).

A Bayesian approach using the Markov Chain Monte Carlo (MCMC) methods is used as an alternative to frequentist approaches for complicated large cross-classified data (Browne et al., 2001; Simonite & Browne, 2003). The MCMC methods have advantages over the IGLS algorithm used in likelihood-based approaches because of their sampling-based estimation. The MCMC estimation methods generate samples from the joint distribution of all unknown parameters; then use these samples to calculate point and interval estimates for each individual parameter. The Gibbs sampler produces samples from the joint posterior by generating from the conditional posterior distributions of unknown parameters (Rasbash & Browne, 2004). As MCMC methods treat the fixed part of the model and one for each of the random classifications as a random additive term, they do not need to construct the global block-diagonal covariance matrix used in the IGLS algorithm. For instance, to fit a basic two-level cross-classified model (e.g., with neighborhoods and schools), we need to specify 6 sets of unknown parameters: the fixed effects, the neighborhood random effects, the school random effects, the neighborhood variance, the school variance, and the residual variance. Since the MCMC algorithm works on each of these terms separately, the algorithm for a cross-classified model is not especially more complicated than for a hierarchical model (Rasbash & Browne 2004).

The data for this study have a cross-classified structure as presented in Figure 1-1. There are two main hierarchies: a hierarchy for residence and a hierarchy for schools. In the first hierarchy, children are nested in a family, which is nested in a neighborhood. The hierarchy for schools is children within schools. Schools are not situated in the residence hierarchy because there is no clear hierarchical relationship between families and schools or between schools and neighborhoods. The crossed structure of the data, the crossing between families and schools and that between neighborhoods and schools, appears when we connect children to schools. Not all children from the same families attend same schools. Some schools have children from more than one neighborhood. For example, in L.A.FANS, the reading score (Letter-Word Identification) analytic sample includes 1933 children from 1377 families nested in 65 neighborhoods. The children are also enrolled in 647 schools. Among the families, 821 have one child, while 556 families have 2 children in the sample. Among the children of two-child families, 646 children attend different schools, while 466 attend the same school with their sibling. Children in a neighborhood attend 6 to 24 different schools (mean 12 schools). Schools have students from 1 to 5 neighborhoods (mean 1.2 neighborhoods). About 16 percent of the 647 schools have students from more than one neighborhood.

Table 2 shows the number of observations at different levels in the PSID and in the L.A.FANS. There are similar numbers of children in the each sample. As the table shows, each survey contains similar numbers of children per family. Nearly two thirds of the families have only one child in them, and thus may be referred to as a “singleton” groups. (With singleton groups, one obviously cannot distinguish between individual and group heterogeneity (Clarke and Wheaton, forthcoming).) Because of the clustering by tract in the L.A.FANS, there are fewer tracts, and therefore fewer schools in that survey than in the PSID. Thus, there is greater clustering of children by schools and tracts in the L.A.FANS than in the national survey. In the L.A.FANS, just under half of the schools have only one child, and none of the tracts have one child. In the PSID, three quarters of the schools are represented by one child, and slightly more than half of the tracts are represented by one child in the survey. Thus, at all levels, the majority of “groups” are singletons containing only 1 observation in the PSID. One potential way of

dealing with cross-classified data to make it more manageable is to create combinations of the data, cross-classified by tract and school. Unfortunately, when these data are combined in this way, slightly more than 90 percent of the combinations have only one observation. Just 32 combinations (1 percent) have three or more observations.

Analysis

Using data from both the L.A.FANS and the PSID, we fitted a series of multilevel models of reading and math scores using the MCMC estimation methods available in MLwiN 2.0 to account for the complex structure of the data with a large number of observations and more than one cross-classification.

The first stage of the analysis is to disaggregate sources of variation in children's test scores by individual, family, neighborhood, and school to provide information on the degree to which each dimension (or level) contributes to variation in the outcomes of interest. The model provides information on the distribution of variation in test scores by individuals, families, neighborhoods, and schools.

In the second stage of the analysis, we examine the association of specific neighborhood and school factors with children's test scores, after controlling for individual and family characteristics.

5. RESULTS

We now turn to presenting the results of our analyses. There are two main sets of findings that we consider: first, the degree of clustering in children's test scores by neighborhood, school, and family and, second, the effects of neighborhood and school covariates on test scores.

Clustering of Children's Test Scores by Neighborhood, School, and Family

The degree of clustering in children's test scores provides an extremely useful measure of the importance of unmeasured factors operating at each particular data level. We determine the degree of clustering by estimating multilevel models and calculating summary measures based on the distribution of the random effects. In particular, the variance of the random effects can be assessed directly or can be used to calculate the intra-class correlation coefficient (ICC). The ICC is a measure of the percentage of the total variance in the outcome (test scores) that is accounted by all unmeasured factors operating at each of the four specific data levels (neighborhood, school, family, and individual).

The most basic multilevel model is one that includes no covariates, other than a constant that returns the overall sample mean. However, by capturing the multilevel cross-classified structure of the data, such a model can provide valuable information about the overall degree of clustering in the outcome at each data level. Through the ICC, this in turn provides an indicator of the importance of unmeasured factors operating at each data level. Since no measured factors are included in the model, the ICC summarizes the effects of *all measurable and unmeasurable*

factors at each level. The ICC thus provides a useful summary of the importance of each data level in accounting for variation in the outcome measure.

We present results from estimating multilevel cross-classified models using MCMC estimation for reading and math test scores using data from L.A.FANS and the PSID-CDS in Table 3. Each column corresponds to a model for a separate outcome and data set. The table shows the estimated constant for the model — which is the sample mean for the outcome — along with the estimated variance components. In the bottom panel of Table 3 we present the ICC estimates. We focus on these measures, because they are the easiest to interpret. Note, however, that the ICCs are calculated from the panel of variances which include the estimated standard errors and indicators of statistical significance.

Looking across the four different models, the most striking finding is the extremely small ICC for schools. For reading and mathematics test scores in both the L.A.FANS and PSID the results indicate that essentially *none* of the variation in children’s test scores is accounted for by school characteristics. On the other hand, neighborhood factors that are independent of schools account for between 9% and 22% of the overall variation in children’s reading and math test scores. The most consistent results is for family factors, which are associated with 25% – 27% of the variation in test scores. Finally, individual level factors represent the remaining source of variation.

The results concerning the complete lack of schools effects in these models are puzzling. Although we may have expected neighborhoods to be somewhat more important based on past research (e.g., Card and Rothstein, 2006) and our own intuition, we did not expect that schools would have no effect at all on variation in children’s test scores even after accounting for clustering by family and neighborhood.

Drawing on our discussion of conceptual issues in identifying school and neighborhood effects in Section 2, we have explored a number of factors that might explain the lack of school effects in two surveys that are similar in content yet distinct in terms of sample design and structure (see Table 2). For example, we examined whether the results differed by age group. Results (not shown) from age-stratified models revealed no consistent effect of school-level factors for either outcome in either data set.

Our leading hypothesis is, however, that the (lack of) results for schools effects on test scores reflects a shortcoming in terms of statistical power together with a school effect that is relatively modest in reality. Clues regarding the precision with which the school effects are estimated are shown in Figure 2, which presents various diagnostic tests for the Markov-chain Monte Carlo estimation of the neighborhood random effects (in the top panel) and of the school random effects (in the bottom panel) in a multilevel cross-classified model of children’s reading scores based on the L.A.FANS data. The results reveal a poorly behaved school random effect that displays very high levels of autocorrelation (see right and left middle panels) and a parameter estimate near the boundary of the parameter space (see top right panels and note that variances cannot be negative). In contrast, the diagnostic parameters for the neighborhood effects in the top panel show what well-behaved results look like — there are modest levels of

autocorrelation in the chains, the distribution of the parameter estimates is symmetric (see top row) and precision of the parameter estimate is high (reflected in the remaining diagnostics).

To examine whether small samples might be driving these results, we estimated a model based on pooled data from both the L.A.FANS and the PSID-CDS-2. This essentially doubled the sample size but did not lead to more precise or well-behaved parameter estimates.

In circumstances when there appears to be insufficient power to estimate models of interest, the best solution is to obtain more data. Because both the L.A.FANS and the PSID-CDS are longitudinal studies, it is in fact possible to obtain additional data and strengthen the statistical power of the parameter estimates. Nevertheless, it was a striking finding to us that the PSID-CDS and especially the L.A.FANS were not sufficiently large to obtain precise and trouble-free estimates of school effects in the presence of large and easily identifiable family and neighborhood effects.

Effects of Neighborhood and School Covariates on Test Scores

We are in the process of verifying the findings described above regarding the very small and underpowered school effects by estimating models of test scores incorporating specific school-level covariates. Among the covariates we are examining are those describing socioeconomic status of schools and pupils, levels of financing and investment, and average test scores. Table 4 presents the means, standard deviations and proportions of the covariates. Table 5 presents estimates from the multilevel models with covariates using data from both the PSID and L.A.FANS.

6. CONCLUSIONS

We began this paper by noting that there has been relatively little research to date on disentangling the effects of schools from the effects of neighborhoods in shaping children's academic outcomes. This situation has occurred despite the considerable interest in each of the two separate research literatures that examines separately the effects of schools and the effects of neighborhoods. Our results suggest that either school effects are very small or else that the size of our data sets is too small to obtain parameter estimates with sufficient statistical power. One obvious solution is to collecting more data — either through new waves of existing surveys or through new, larger surveys. However, a more definitive way to assess whether even these new larger data sets will allow us to estimate school effects with sufficient precision is to conduct a simulation study.

Another alternative is to pursue estimates based on other data sets that have a fundamentally different design than L.A.FANS and PSID-CDS. For example, school-based studies may provide an opportunity to examine both school and neighborhood effects and, possibly, family effects as well.

We are in the process of actively pursuing these different research alternatives in order to obtain a better and more precise understanding of how social context affects children's academic achievement.

References

A User's Guide to MLwiN. 2004.

Betts, JR and D Morell (1999) "The Determinants of Undergraduate Grade Point Average: The Relative Importance of Family Background, High School Resources, and Peer Group Effects" *The Journal of Human Resources*, 34(2): 268-293.

Bronfenbrenner, U. (1986). Ecology of the family as context for human development. *Developmental Psychology*, 22(6): 723-742.

Browne, W. J., Draper, D., Goldstein, H. and Rasbash, J. (2002). Bayesian and likelihood methods for fitting multilevel models with complex level-1 variation. *Computational Statistics and Data Analysis*, **39**: 203-225.

Browne, W. J., Goldstein, H. and Rasbash, J. (2001). Multiple membership multiple classification (MMMC) models. *Statistical Modelling* **1**: 103-124

Coleman, J.S. (1988). Social capital in the creation of human capital. *American Sociological Review Supplement*, 94: S95-S120.

Everson, HT and RE Millsap (2004) "Beyond Individual Differences: Exploring School Effects on SAT Scores" *Educational Psychologist*, 39(3), 157-172.

Ginther, D., Haveman, R., & Wolfe, B. (2000). Neighborhood attributes as determinants of children's outcomes: how robust are the relationships. *Journal of Human Resources*, 35(4):603-642.

Goldstein H. 2003. Multilevel Statistical Models 3rd edition. Chapter 8. Hodder Arnold. London, UK.
Available at http://www.ats.ucla.edu/stat/examples/msm_goldstein/.

Hanushek, EA (1997) "Assessing the Effects of School Resources on Student Performance: An Update" *Educational Evaluation and Policy Analysis*, 19(2): 141-164.

Hox J. 2002. Multilevel Analysis Techniques and Applications. pp.123-138. Lawrence Erlbaum Associates, Inc. Publishers. Mahwah, New Jersey.

Leventhal, T & Brooks-Gunn, J. (2000). The neighborhoods they live in: the effect of neighborhood residence on child and adolescent outcomes. *Psychological Bulletin*, 126(2):309-337.

Parcel, TL and MJ Dufur (2001) "Capital at Home and School: Effects on Student Achievement" *Social Forces*, 79(3): 881-911.

Rasbash, J. and Browne, W. J. (2001). Modelling Non-Hierarchical Structures. In Leyland, A. H. and Goldstein, H. (Eds.), *Multilevel Modelling of Health Statistics*, 93-105. Chichester: John Wiley & Sons.

Rasbash J.& Browne WJ. 2004. Non-Hierarchical Multilevel Models. To appear in De Leeuw, J and Kreft, I.G.G. (Eds.), *Handbook of Quantitative Data Analysis*, Kluwer Press.

Rasbash J. & Goldstein H. 1994. Efficient analysis of mixed hierarchical and cross-classified random structures using a multilevel model. *Journal of Educational and Behavioral Statistics*. Vol. 19, No. 4, pp.337-350

Raudenbush, SW and JD Willms (1995) "The Estimation of School Effects" *Journal of Educational and Behavioral Statistics*, 20(4): 307-3

Sewell, W.H., & Armer, J.M. (1966). Neighborhood context and college plans. *American Sociological Review*, 31(2):159-168.

Simonite, V., and Browne, W. J. (2003). Estimation of a large cross-classified multilevel model to study academic achievement in a modular degree course. *Journal of the Royal Statistical Society, A*. 166: 119-134.

St. John, NH and Lewis R. 1971. The influence of school racial context on academic achievement. *Social Problems*. Vol. 19, No.1, 68-79.

Webster WJ, Mendro RL, Orsak TH, Dash Weerasinghe Dallas Public Schools. 1996. The applicability of selected regression and hierarchical linear models to the estimation of school and teacher effects. Paper presented at the Annual Meeting of the National council on Measurement in Education. New York, NY, April 9-11.

Wilson, W.J. (1987). *The Truly Disadvantaged*. Chicago: University of Chicago Press.

Wilson, W.J. (1996). *When Work Disappears: The World of the New Urban Poor*. New York: Alfred A. Knopf.

Betts, JR (1996) "Is There a Link between School Inputs and Earnings? Fresh Scrutiny of an Old Literature" in G. Burtless (ed.) *Does Money Matter? The Effects of School Resources on Student Achievement and Adult Success*. Washington, DC: The Brookings Institution.

Table 1. Summary Statistics for Children’s Reading and Mathematics Achievement in L.A.FANS and PSID-CDS

Measure	L.A.FANS		PSID-CDS	
	Mean	(Std. Dev.)	Mean	(Std. Dev.)
Reading Score	103.2	(18.8)	103.0	(18.9)
Observations	1,933		1,904	
Math Score	102.9	(17.4)	101.4	(16.9)
Observations	1,928		1,907	

Table 2. Multilevel Data Structure for L.A.FANS and PSID-CDS

Measure	L.A.FANS				PSID-CDS		
	Mean (Std. dev.)	Pct. cases with just 1	Maximum		Mean (Std. dev.)	Pct. cases with just 1	Maximum
Children							
Total	1,933				1,907		
Per family	1.40 (0.49)	60%	2		1.38 (0.49)	62%	2
Per school	2.99 (3.32)	48	26		1.33 (0.73)	75	9
Per tract	29.74 (7.73)	0	63		1.62 (1.07)	55	19
Families							
Total	1,377				1,384		
Per school	2.38 (2.41)	56%	16		1.16 (0.57)	89%	7
Per tract	21.18 (5.25)	0	44		1.18 (0.64)	89	10
Schools							
Total	647				1,436		
Per family	1.23 (0.42)	77%	2		1.21 (0.41)	79%	2
Per tract	8.91 (3.29)	5	20		1.34 (0.62)	72	6
Tracts							
Total	65				1,176		
Per school	1.20 (0.50)	84%	5		1.10 (0.36)	92%	4

Table 3. Multilevel Cross-Classified Random Effects Models for Children’s Reading and Math Achievement: L.A.FANS and PSID-CDS

Covariate	L.A.FANS		PSID-CDS	
	Reading score	Math score	Reading score	Math score
Constant	103.2*** (0.82)	102.8*** (1.03)	103.3*** (0.50)	103.2*** (0.82)
Variance				
Neighborhood	30.8*** (8.25)	56.6*** (12.58)	77.9*** (24.43)	30.8*** (8.25)
School	0.7 (1.80)	3.2 (4.32)	3.2 (6.42)	0.7 (1.80)
Family	95.7*** (13.40)	78.5*** (11.24)	92.0*** (25.64)	95.7*** (13.40)
Child	230.1*** (13.03)	171.7*** (10.35)	185.5*** (11.73)	230.1*** (13.03)
Total	357.3	310.0	357.3	357.3
ICC				
Neighborhood	9%	18%	22%	9%
School	0	1	1	0
Family	27	25	26	27
Child	64	56	52	64
Observations	1,933	1,928	1,910	1,907

Note: Models estimated using Markov-chain Monte Carlo estimation procedure. Standard errors in parentheses. *** $p < .01$

Table 4. Summary statistics for the base analytic samples

Variable	<u>PSID/CDS-II</u>		<u>LAFANS</u>
	Mean	SD	Mean(std. dev.) or percent
Child age (years)	12.1	(3.6)	10.6 (3.6)
Child sex			
Male	50%		50%
Female	50%		50%
Birthweight (kilograms)	3.1	(0.7)	3.39 (0.63)
Race/ethnicity			
Latino	7%		60%
African-american	45%		10%
White	44%		21%
Asian	1%		7%
Other	3%		2%
Observations (children)	1,907		1,939
Mother's reading score	92.0	(14.6)	85.0 (17.9)
Reading score missing	24%		
Mother's schooling (years)	13.0	(2.0)	11.5 (4.4)
Schooling missing	5%		
Family income (\$)	60,059.0	(77,553.8)	
Non-housing assets (\$)	104,566.6	(1,161,739.0)	
Wealth (\$)	139,600.6	(1,191,025.0)	
Log family income	10.6	(0.9)	9.09 (3.24)
Log non-housing assets	9.6	(1.8)	
Log wealth	9.9	(2.1)	
Observations (families)	1,384		1,382
Tract median family income (\$10,000)	4.7	(2.1)	4.49
Tract immigrant concentration index	0.0	(0.9)	0.02 (0.99)
Tract residential stability score	0.0	(0.9)	0.04 (0.97)
Tract racial/ethnic diversity score	0.5	(0.2)	
Observations (tracts)	1,176		65
Percent black	31%		
Percent hispanic	12%		
Percent free/reduced lunch	43%		
Pupil/teacher ratio	17.4	(5.7)	
Teacher/student ratio	0.06		0.05 (0.27)
Number of students	886.0	(672.0)	1,148 (971)
District revenue per student	8,922.6	(2,115.9)	
Observations (schools)	1,436		650

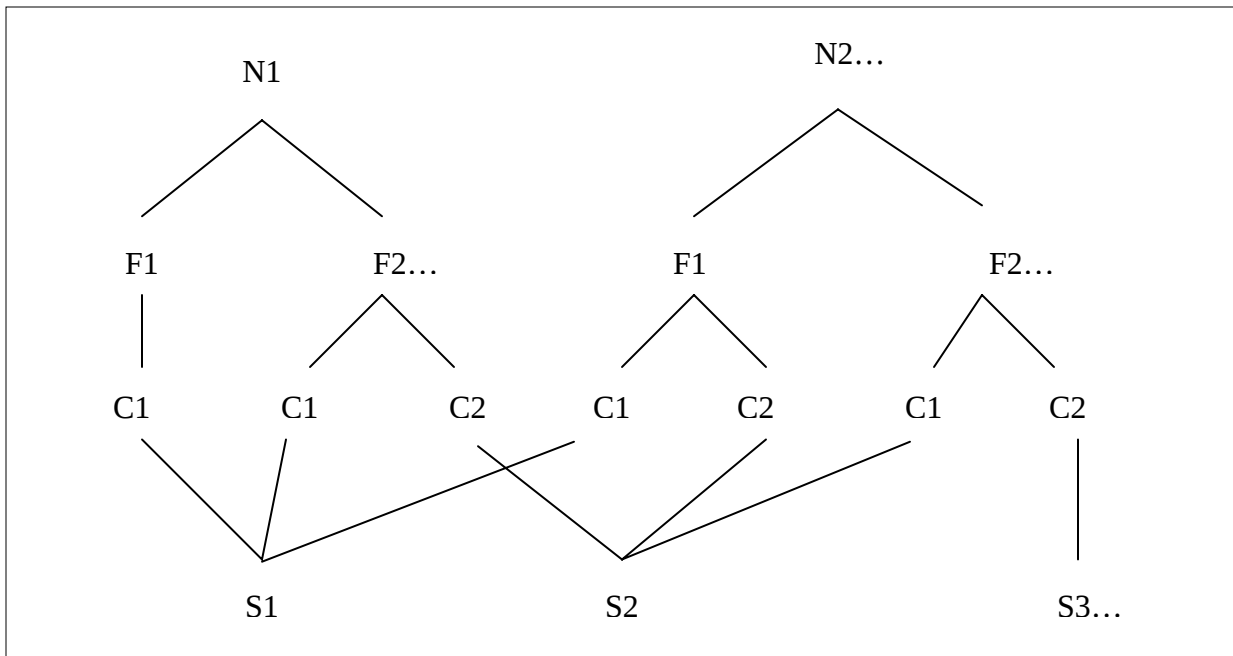
Table 5. Multilevel models with covariates

	<u>L.A.FANS</u>		<u>PSID-CDS-II</u>			
	Reading	Mathematics	Reading		Math	
Sample size	1,933	1,928	1,907		1,905	
Intercept (SE)	73.6 (8.6)	92.0 (7.7)	53.21	(7.23)	68.10	(6.18)
Predictor (SE)						
Individual level						
Age (years)	-0.28 (0.15)	-0.54 (0.13)	-0.71	(0.13)	-0.43	(0.12)
Gender (ref: male)	2.74 (0.78)	-0.80 (0.67)	3.97	(0.76)	-0.06	(0.68)
Latino	-3.45 (1.52)	-5.04 (1.40)	-0.58	(2.76)	-2.67	(2.37)
African American	-2.30 (1.84)	-4.44 (1.67)	2.83	(1.24)	-6.34	(1.07)
Asian	1.86 (2.01)	2.48 (1.81)	-0.88	(4.07)	8.84	(3.48)
Other	0.83 (3.13)	0.09 (2.80)	-0.04	(2.53)	-2.54	(2.18)
Spanish version for test	11.2 (1.32)	-6.29 (1.19)				
Birth weight in kg	0.21 (0.65)	1.12 (0.57)	1.35	(0.62)	1.07	(0.55)
Family level						
Both parents home	1.92 (0.95)	1.93 (0.84)				
PCG's PC score	0.22 (0.03)	0.16 (0.03)	0.19	(0.03)	0.09	(0.03)
Number of children	0.21(0.35)	-0.12 (0.31)				
Log family income	0.29 (0.14)	0.33 (0.13)	1.42	(0.62)	1.42	(0.53)
Immigrated before 90	4.19 (1.26)	1.67 (1.15)				
Immigrated after 90	5.79 (1.53)	2.62 (1.35)				
PCG's years of schooling	0.28 (0.13)	0.31 (0.11)	1.37	(0.25)	0.86	(0.21)

Table 5. Multilevel models with covariates

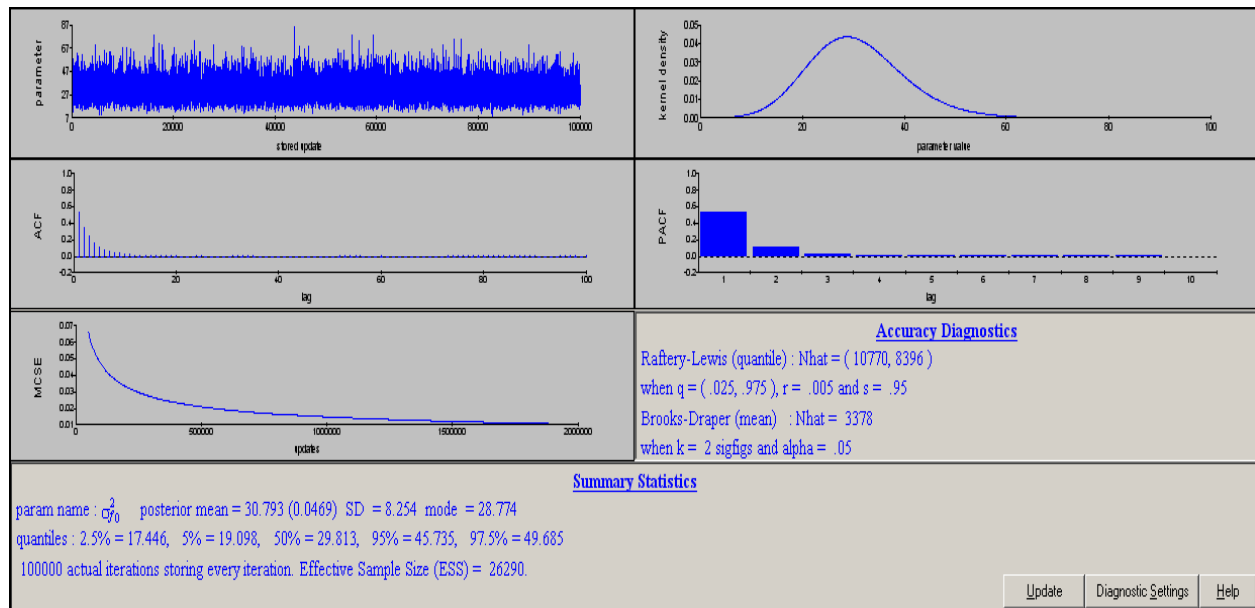
	<u>L.A.FANS</u>		<u>PSID-CDS-II</u>			
	Reading	Mathematics	Reading		Math	
School level						
Sector (ref: private)	-1.13 (7.04)	-9.84 (6.18)				
Teacher/student ratio	-27.9 (25.4)	-32.2 (23.1)	-0.07	(0.08)	-0.05	(0.07)
Size (number of students)	-1.03 (0.56)	-1.07(0.50)	0.00	(0.00)	.000	(0.00)
School mean API score	0.54 (0.51)	0.45 (0.46)				
No API score reported	2.62 (6.90)	-9.37 (6.08)				
Percent receiving free-lunch			-4.45	(2.12)	-2.32	(1.85)
Neighborhood level						
Median Income	0.79 (0.36)	1.19 (0.34)	0.08	(0.03)	0.14	(0.03)
Immigrant concentration	0.98 (0.97)	1.17 (0.96)	0.37	(0.74)	-0.11	(0.64)
Residential stability	-0.50 (0.69)	-1.19 (0.67)	0.73	(0.54)	0.84	(0.46)
Diversity			3.87	(2.60)	0.01	(2.23)
Random part (SE)						
Neighborhood	7.3 (3.8)	9.0 (3.3)	15.15	(17.69)	3.33	(5.21)
School	0.4 (0.9)	0.02 (0.02)	5.07	(7.41)	12.30	(10.83)
Family	63.9 (11.4)	61.2 (9.1)	82.91	(22.25)	43.22	(11.43)
Individual	226.2 (12.1)	165.7 (9.1)	186.46	(12.37)	163.50	(11.97)
Bayesian DIC*	16379	15814	15402		15110	

Figure 1. Multilevel Structure of the L.A.FANS Data, Distinguishing Neighborhood, School, Family, and Individuals

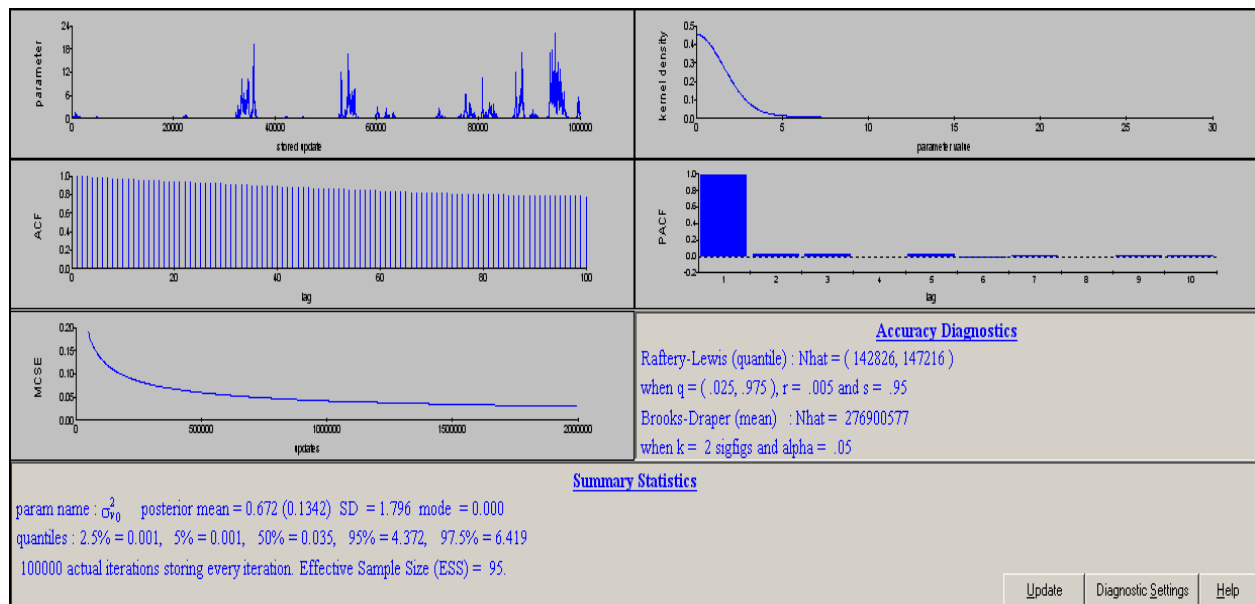


Note: N=neighborhood; F=family; C=child; S=school.

Figure 2. Markov-chain Monte Carlo Diagnostic Plots for Estimated Variances of Neighborhood and School Random Effects From Multilevel Cross-Classified Random Effects of Children's Reading Score in L.A.FANS



A. Variance of Neighborhood Random Effect



B. Variance of School Random Effect