

PANEL DATA MODELS
(or Longitudinal or Time-Series/Cross-Section)
and an INTRODUCTION TO SIMULATION-BASED INFERENCE

1 Preliminary Issues:

1.1 At least double-indexed data

Panel or Longitudinal or Time-series/Cross-section Data are such where a unit of observation s subsumes at least two indices/dimensions of sampling. E.g.,

$$\begin{array}{llll} s = 1, \dots, S & & & \\ y_s & x'_s & \epsilon_s & s = it \text{ where } \begin{array}{ll} i = 1, \dots, N & \text{cross-section side} \\ t = 1, \dots, T_i & \text{time-series side} \end{array} \end{array}$$

NB: Throughout our discussion, we will focus on “Large N, small T” asymptotics with $N \rightarrow \infty$ while $\max_i T_i \approx$ small and finite.

1.2 Organization of the data — three alternatives with two dimensions:

1.2.1 t “fastest”:

$$\{y_s\} = \begin{pmatrix} y_1 \\ \vdots \\ y_s \\ \vdots \\ y_S \end{pmatrix} = \begin{pmatrix} y_{11} \\ \vdots \\ y_{it} \\ \vdots \\ y_{NT_N} \end{pmatrix} = \begin{pmatrix} y_{11} \\ y_{12} \\ \vdots \\ y_{1t} \\ \vdots \\ y_{1T_1} \\ \hline y_{21} \\ y_{22} \\ \vdots \\ y_{2t} \\ \vdots \\ y_{2T_2} \\ \hline \vdots \\ y_{N1} \\ \vdots \\ y_{NT_N} \end{pmatrix} \cdots \begin{pmatrix} IID \end{pmatrix} \cdots \begin{pmatrix} TID \end{pmatrix}$$

1.2.2 *i* “fastest”:

$$\{y_s\} = \begin{pmatrix} y_1 \\ \vdots \\ y_s \\ \vdots \\ y_S \end{pmatrix} = \begin{pmatrix} y_{11} \\ \vdots \\ y_{it} \\ \vdots \\ y_{NT_N} \end{pmatrix} = \begin{pmatrix} y_{11} \\ y_{21} \\ \vdots \\ y_{i1} \\ \vdots \\ y_{N1} \\ \hline y_{12} \\ y_{22} \\ \vdots \\ y_{i2} \\ \vdots \\ y_{N2} \\ \hline \vdots \\ y_{N1} \\ \vdots \\ y_{NT_N} \end{pmatrix} \dots \begin{pmatrix} IID \end{pmatrix} \dots \begin{pmatrix} TID \end{pmatrix}$$

1.2.3 data organized as they come but double-indexed ID variables:

$$\begin{aligned}
 \{y_s\} &= \begin{pmatrix} y_1 \\ \vdots \\ y_s \\ \vdots \\ y_S \end{pmatrix} \dots \begin{pmatrix} iid(1) \\ \vdots \\ iid(s) \\ \vdots \\ iid(S) \end{pmatrix} \dots \begin{pmatrix} tid(1) \\ \vdots \\ tid(s) \\ \vdots \\ tid(S) \end{pmatrix} \\
 &= \text{Sx1 vector } y \dots \text{Sx1 vector } IID \dots \text{Sx1 vector } TID
 \end{aligned}$$

1.3 (3) Balanced ($T_i = T$) vs. Unbalanced Data Sets (T_i varies with i)

Balanced: $S = N \times T$:

$$\begin{pmatrix} y_1 \\ \vdots \\ y_s \\ \vdots \\ y_S \end{pmatrix} = \begin{pmatrix} y_{11} \\ y_{12} \\ \vdots \\ y_{1t} \\ \vdots \\ y_{1T} \\ \hline y_{21} \\ y_{22} \\ \vdots \\ y_{2t} \\ \vdots \\ y_{2T} \\ \hline \vdots \\ y_{N1} \\ \vdots \\ y_{NT} \end{pmatrix}$$

Unbalanced: $S = \sum_{i=1}^N T_i :$

$$\begin{pmatrix} y_1 \\ \vdots \\ y_s \\ \vdots \\ y_S \end{pmatrix} = \begin{pmatrix} y_{11} \\ y_{12} \\ \vdots \\ y_{1t} \\ \vdots \\ y_{1T_1} \\ \hline y_{21} \\ y_{22} \\ \vdots \\ y_{2t} \\ \vdots \\ y_{2T_2} \\ \hline \vdots \\ y_{N1} \\ \vdots \\ y_{NT_N} \end{pmatrix}$$

1.3.1 (3b) (related issue) Use PADDING with Missing Data Code (MDC) — Then every Unbalanced PDS becomes Balanced

New single constant $T = \max_i T_i$.

1.3.2 (3c) (related issue) DROP OBSERVATIONS to make Balanced

Example: new single constant $T = \min_i T_i$.

(4) Lagged variables in Panel Data

$$\begin{aligned}
 LAG1 \begin{pmatrix} y_0 \\ \vdots \\ y_{s-1} \\ \vdots \\ y_{S-1} \end{pmatrix} &= \begin{pmatrix} y_{10} \\ y_{11} \\ \vdots \\ y_{1,t-1} \\ \vdots \\ y_{1,T_1-1} \\ \text{---} \\ y_{1T_1} \\ y_{21} \\ \vdots \\ y_{2,t-1} \\ \vdots \\ y_{2,T_2-1} \\ \text{---} \\ \vdots \\ y_{N-1,T_N} \\ \vdots \\ y_{N,T_N-1} \end{pmatrix} = \begin{pmatrix} MDC \\ y_{11} \\ \vdots \\ y_{1,t-1} \\ \vdots \\ y_{1,T_1-1} \\ \text{---} \\ y_{1T_1} \\ y_{22} \\ \vdots \\ y_{2,t-1} \\ \vdots \\ y_{2,T_2-1} \\ \text{---} \\ \vdots \\ y_{N-1,T_N} \\ \vdots \\ y_{N,T_N-1} \end{pmatrix} \quad \text{vs.} \quad XTLAG1 \begin{pmatrix} y_0 \\ \vdots \\ y_{s-1} \\ \vdots \\ y_{S-1} \end{pmatrix} = \begin{pmatrix} y_{10} \\ y_{11} \\ \vdots \\ y_{1t} \\ \vdots \\ y_{1,T_1-1} \\ \text{---} \\ y_{20} \\ y_{22} \\ \vdots \\ y_{2,t-1} \\ \vdots \\ y_{2,T_2-1} \\ \text{---} \\ \vdots \\ \text{---} \\ y_{N0} \\ \vdots \\ y_{N,T_N-1} \end{pmatrix} = \begin{pmatrix} MDC \\ y_{11} \\ \vdots \\ y_{1t} \\ \vdots \\ y_{1,T_1-1} \\ \text{---} \\ MDC \\ y_{22} \\ \vdots \\ y_{2,t-1} \\ \vdots \\ y_{2,T_2-1} \\ \text{---} \\ \vdots \\ \text{---} \\ MDC \\ \vdots \\ y_{N,T_N-1} \end{pmatrix}
 \end{aligned}$$

In sum, the LAG1 variable will contain a single Missing Value, whereas the XTLAG1 variable will contain N Missing Values.

1.4 (5) Linear vs. Nonlinear models (additive vs nonadditive, index vs general)

	$s = 1, \dots, S$	$i = 1, \dots, N$ and $t = 1, \dots, T_i$
<i>Linear</i>	$y_s = x'_s \beta + \epsilon_s$	$y_{it} = x'_{it} \beta + \epsilon_{it}$
<i>Additively Nonlinear Index</i>	$y_s = f(x'_s \beta) + \epsilon_s$	$y_{it} = f(x'_{it} \beta) + \epsilon_{it}$
<i>Additively Nonlinear</i>	$y_s = g(x'_s, \beta) + \epsilon_s$	$y_{it} = g(x'_{it}, \beta) + \epsilon_{it}$
<i>Non-additively Nonlinear</i>	$y_s = h(x'_s, \beta, \epsilon_s)$	$y_{it} = h(x'_{it}, \beta, \epsilon_{it})$

1.5 (6) Combination of (1) and (3): Endogenous Data Availability

NB: even an apparently Linear model is in fact Nonlinear if Endogenous Data Availability — Distinction between Latent and (observed) Limited Dependent Variables.

Modelling Framework: Sample Selection or Selectivity or Endogenous Data Availability or Endogenous Attrition
Two-equation Latent variables model:

$$\begin{aligned}y_{it}^* &= x_{it}'\beta + \epsilon_{it} \\ d_{it}^* &= z_{it}'\gamma + u_{it}\end{aligned}$$

Observation LDV Rule:

$$\begin{aligned}D_{it} &= \begin{cases} 1 & \text{iff } d_{it}^* = z_{it}'\gamma + u_{it} > 0 \\ 0 & \text{iff } d_{it}^* = z_{it}'\gamma + u_{it} \leq 0 \end{cases} \quad \text{and} \\ y_{it} &= \begin{cases} y_{it}^* & \text{iff } d_{it}^* = z_{it}'\gamma + u_{it} > 0 \\ MDC & \text{iff } d_{it}^* = z_{it}'\gamma + u_{it} \leq 0 \end{cases}\end{aligned}$$

NB: Distinction between Censored Selectivity and Truncated Selectivity:

Selectivity with Censoring

$$y_{it} = \begin{cases} y_{it}^* & \text{iff } D_{it} = 1 \\ MDC & \text{iff } D_{it} = 0 \end{cases} \quad \text{and} \\ D_{it}, x_{it}, \text{ and } z_{it} \text{ always observed}$$

Selectivity with Truncation

$$y_{it} = \begin{cases} y_{it}^* & \text{iff } D_{it} = 1 \quad \text{and} \\ D_{it}, x_{it}, \text{ and } z_{it} \text{ observed *only* when } D_{it} = 1 \end{cases}$$

NB: Fundamental Point: If u_{it} & ϵ_{it} are *not* *independent*, then

$$\begin{aligned} E(y_{it}|X) &\neq x'_{it}\beta \quad \text{and} \quad E(y_{it}|X, Z) \neq x'_{it}\beta \quad \text{*BUT*} \\ E(y_{it}|X, Z) &= g(x'_{it}, z'_{it}, \delta) \end{aligned}$$

where the parameter vector δ is related to $\beta, \gamma, \sigma_\epsilon^2, \sigma_u^2$, and $\rho_{\epsilon u}$.

1.6 (7) Types of variables w.r.t. i and t indices:

$$\begin{array}{ccccc} x_s^j = x_{it}^j & vs. & z_s^j = z_i^j & vs. & w_s^j = w_t^j \\ \text{default} & & \text{time-invariant} & & \text{individual-invariant (e.g., economy-wide/macro)} \end{array}$$

1.7 (8) Error-Components/Factor-Analytic structures:

1.7.1 Error-components with single time-invariant factor:

$$\epsilon_s = \epsilon_{it} = \alpha_i + \nu_{it} = \alpha_s + \nu_s$$

NOTE: α_i is termed the “unobserved persistent heterogeneity”.

Basic assumptions:

$$\alpha_i \sim ?(0, \sigma_\alpha^2) \\ iid \text{ over } i$$

$$\nu_{it} \sim ?(0, \sigma_\nu^2) \\ iid \text{ over } i \text{ and } t$$

and $\alpha_i, \nu_{\ell t}$ independent/uncorrelated for all i, ℓ, t

NB: Key conclusion: $VCov(\epsilon|regressors)$ is a Block-Diagonal matrix with Diagonal blocks equal to:

$$\begin{pmatrix} \sigma_\alpha^2 + \sigma_\nu^2 & \sigma_\alpha^2 & \cdots & \sigma_\alpha^2 \\ & \sigma_\alpha^2 + \sigma_\nu^2 & \ddots & \vdots \\ & & \ddots & \sigma_\alpha^2 \\ & & & \sigma_\alpha^2 + \sigma_\nu^2 \end{pmatrix}$$

and Off-Diagonal blocks between individuals i and n equal to $0_{T_i \times T_n}$. This is called the “equi-correlated” error components model.

Error-components with two factors (one time-, one individual-invariant):

$$\epsilon_s = \epsilon_{it} = \alpha_i + \zeta_t + \nu_{it} = \alpha_s + \zeta_s + \nu_s$$

.where

$$\alpha_i \sim ?(0, \sigma_\alpha^2) \\ iid \text{ over } i$$

$$\nu_{it} \sim ?(0, \sigma_\nu^2) \\ iid \text{ over } i \text{ and } t$$

$$\zeta_t \sim ?(0, \sigma_\zeta^2) \\ iid \text{ over } t$$

and $\alpha_i, \nu_{it}, \zeta_t$ mutually independent/uncorrelated for all i, ℓ, t, q

The $VCov(\epsilon|regressors)$ matrix has a similar block structure with $\sigma_\alpha^2 + \sigma_\nu^2 + \sigma_\zeta^2$ on the main diagonal, and either σ_α^2 , σ_ζ^2 , or $\sigma_\alpha^2 + \sigma_\zeta^2$ in the elements of the off-diagonal blocks depending on the values of i, ℓ, t .

2 Random Effect “vs.” Fixed Effects

Common misconception: the approaches are frequently thought of as *alternative* DGPs. A much more appropriate framework is to think of them as the *same* DGP, but alternative Estimation Approaches

Common DGP with one-factor error-components model as in (1.8) above:

$$y_{it} = x'_{it}\beta + z'_i\gamma + \epsilon_{it} = x'_{it}\beta + z'_i\gamma + \alpha_i + \nu_{it}$$

RE Approaches: in *RED*:

$$y_{it} = x'_{it}\beta + z'_i\gamma + \epsilon_{it} = [x'_{it}\beta + z'_i\gamma] + [\alpha_i + \nu_{it}]$$

FE Approaches in *BLACK*:

$$y_{it} = x'_{it}\beta + z'_i\gamma + \epsilon_{it} = (x'_{it}\beta + z'_i\gamma + \alpha_i) + (\nu_{it})$$

FE-(BLACK): The four classic regression assumptions A1, A2, A3, A4 take the form:

A1	no perfect multicollinearity among the regressors X and Z	$rank(X, Z) = k_x + k_z$
A2	linear additive model	$y = X\beta + Z\gamma + \epsilon$
A3	regressor exogeneity	X and Z exogenous w.r.t. ϵ
A4	$VCov(error regressors)$	$VCov(\epsilon X, Z)$

RE-[RED]: Now the four classic regression assumptions A1, A2, A3, A4 take the form: (D is the full set of N variable intercepts dummies, one for each individual)

A1	no perfect multicollinearity among the regressors X and D	$rank(X, D) = k_x + k_z + N$ NB : Z is dropped since perfectly collinear with D
A2	linear additive model	$y = X\beta + Z\gamma + \epsilon = X\beta + D\alpha + \nu$
A3	regressor exogeneity	X and D exogenous w.r.t. ν (no Z regressors)
A4	$VCov(error regressors)$	$VCov(\nu X, D)$

2.1 *FE-TYPE estimators: the α_i 's are eliminated through suitable transformation or conditioned upon or estimated through sufficient statistics

Key fact: Parameters estimated (either explicitly or implicitly): β (k_x) and $a_1, \dots, a_N$ (N), σ_ν^2 (1)

2.1.1 FE1: FD

***Apply OLS on FD model:

$$\Delta y_{it} = \Delta x'_{it}\beta + 0 + 0 + \Delta \nu_{it}$$

NB1: No estimates of γ are possible by the approach since Z has dropped out.

NB2: $\Delta \nu_{it}$ is a non-invertible MA(1) process, with known parameter -1 . Hence OLS will not be BLUE and we will need to calculate Robust SEs/VCovs

2.1.2 FE2: Quasi-differencing/Within

***Apply OLS on Quasi-Differenced model:

$$Qy = QX\beta + QZ\gamma + Q\alpha + Q\nu = QX\beta + Q\nu$$

where Qy has typical element

$$\{Qy\}_{it} = y_{it} - \bar{y}_{i\cdot} \equiv y_{it} - \sum_{t=1}^{T_i} y_{it}$$

Consequently, the Q transformation eliminates all time-invariant terms — in particular α and Z .

NB1: No estimates of γ are possible by the approach since Z has dropped out.

NB2: The transformation Q is idempotent (and symmetric, hence a projection matrix). Therefore, the

$$VCov(\nu|X) = Q\sigma_\nu^2 I_{NT}Q' = \sigma_\nu^2 Q$$

which is *singular* (it has deficient rank). Therefore its generalized inverse will be *itself* and so the GLS estimator to take into account the non-spherical distribution of ν will be *identical* to plain OLS! To see this formally:

$$\begin{aligned}
\text{plain OLS} & : \hat{\beta}_{FE2} = \hat{\beta}_W = ((QX)'(QX))^{-1} (QX)'(Qy) \\
\text{GLS} & : \left((QX)' (VCov(\nu|X))^{geninv} (QX) \right)^{-1} (QX)' (VCov(\nu|X))^{geninv} (Qy) \\
& = ((QX)'Q(QX))^{-1} (QX)'Q(Qy) = \hat{\beta}_{FE2} = \hat{\beta}_W
\end{aligned}$$

NB3: The FE2 model is *numerically* *identical* to the Variable Intercepts OLS model:

$$y = X\beta + D\alpha + \nu$$

because by the Frisch-Waugh-Lovell theorem, linear regression partitioning gives that:

$$\begin{aligned}
\hat{\beta}_{VIols} & = ((M_D X)'(M_D X))^{-1} (M_D X)'(M_D y) : M_D \equiv I_{NT} - D(D'D)^{-1}D' = Q \\
& = ((QX)'(QX))^{-1} (QX)'(Qy) = \hat{\beta}_{FE2} = \hat{\beta}_W \\
\{\hat{\alpha}_{VIols}\}_i & = \bar{y}_{i\cdot} - \bar{x}'_{i\cdot} \hat{\beta}_{FE2}
\end{aligned}$$

2.2 *RE-TYPE estimators:

Key fact: Parameters estimated: $\beta(k_x)$, $\gamma(k_z)$, $\sigma_\alpha^2(1)$, and $\sigma_\nu^2(1)$

Consider model

$$y = [X\beta + Z\gamma] + [\alpha + \nu] = [X\beta + Z\gamma] + [\epsilon] \equiv W\theta + \epsilon$$

RE1: pooled OLS

$$\hat{\theta}_{RE1} = \begin{pmatrix} \hat{\beta}_{RE1} \\ \hat{\gamma}_{RE1} \end{pmatrix} = (W'W)^{-1}W'y$$

NB: This will *not* be BLUE and its *Robust* SEs/VCov must be calculated to allow for the Clustering exhibited by the *block-diagonal* $VCov(\epsilon|X, Z) \equiv \sigma_\epsilon^2\Omega$.

RE2: "the RE"-GLS estimator

$$\hat{\theta}_{RE2} = \hat{\theta}_{REglS} = \begin{pmatrix} \hat{\beta}_{REglS} \\ \hat{\gamma}_{REglS} \end{pmatrix} = (W'\Omega^{-1}W)^{-1}W'\Omega^{-1}y$$

NB1: This estimator will be BLUE and will have the correct SEs/VCov.

NB2: In 1972, Fuller and Battese showed that calculating Ω^{-1} , which is computationally burdensome, can be avoided. Instead, the rotation $\Omega^{-1/2'}$ yields the equivalent very straightforward expressions:

$$\begin{aligned} \Omega^{-1/2'}y &= \{y_{it} - \lambda_i\bar{y}_{i.}\} \\ \Omega^{-1/2'}X &= \{x_{it} - \lambda_i\bar{x}_{i.}\} \\ \Omega^{-1/2'}Z &= \{(1 - \lambda_i)z_i\} \\ \text{where } \lambda_i &= 1 - \sqrt{\frac{\sigma_\nu^2}{\sigma_\nu^2 + T_i\sigma_\alpha^2}} \end{aligned}$$

Hence the RE2-GLS estimator can be obtained by applying plain OLS on the $\Omega^{-1/2'}$ -transformed variables.

© Vassilis Hajivassiliou, LSE 2012–2023