

Some Model-based Estimator

Danutė Krapavickaitė and Vilma Nekrašaitė

Vilnius, Lithuania

Survey sampling:

Finite population is a fixed, mammoth parameter

Sampler does a dance around it: picks this, omits that, introducing randomness

Survey design and inference – the study of that dance

Main branches of statistics:

Data is a realization of random variables

Statistician makes inference about probability laws of these random variables

Finite population $\mathcal{U} = \{1, 2, \dots, N\}$

Study variable y : $\{y_1, y_2, \dots, y_N\}$,

Parameter of a population, for example, total $t_y = \sum_{k=1}^N y_k$, fixed, unknown

Sample $i \subset \mathcal{U}$ size n drawn according to SRS and SRW designs:

$$\hat{t}_y = \sum_{k \in i} \frac{n}{N} y_k$$

$$Var_{SRS}(\hat{t}_y) = N^2 \left(1 - \frac{n}{N}\right) \frac{s_y^2}{n}, \quad Var_{SRW}(\hat{t}_y) = N^2 \frac{s_y^2}{n}$$

They differ much if $\frac{n}{N} = \frac{1}{2}$

Conditionality principle used often in statistical inference:

statistical inference is made conditionally

on observed random variables whose probability distribution is known

and not dependent on parameters about which we must make inferences.

Sampling variance is a property of process of selecting the sample and

not of the outcome of that process.

Something is wrong with this principle in survey sampling!

Nothing is fixed in the world.

What statistical model can describe the process

which has generated the finite population?

Other statisticians are interested in estimating the parameters of the model

For the survey statistician when the model is chosen

estimation of model parameters plays some assistant role for the main task:

to make inference about the realized finite population itself.

Model-based approaches:

- Bayesian prediction techniques (the posterior distribution of t_y is calculated when y_k , $k \in 1$ are given) Boltarine and zacks, 1992; Ericson, Ghosh, Pfeferman, Royal
- Fiducial inference (palyginamoji analizė), Kalbfleisch and Sprott, 1992
- Likelihood prediction, Bjornstad, Royall
- Prediction based on general linear models. Valliant, Dortman, Royall, 2000; Chambers

Randomisation is desirable but not necessary nor sufficient for statistical inference.

Valid inference can be done with and without it, and when randomisation is present, it does not necessarily creates the best probabilistic approach for inference.

Question: whether inferences should be based on distributions created by randomization or on probability models arises when interpreting data from

- designed experiments
- clinical trials
- observational studies

Prediction Theory and the General Linear Model

Finite population $\mathcal{U} = \{1, 2, \dots, N\}$

The population **vector of values of the variable** y : $y = (y_1, \dots, y_N)'$,

It is treated as realization of the **random vector** $Y = (Y_1, \dots, Y_N)'$

The goal – to estimate a linear combination of y 's:

$\gamma'Y$, $\gamma = (\gamma_1, \dots, \gamma_N)'$ – vector of constants

Separate cases:

$$\gamma = (1, \dots, 1)' \iff \gamma'y = t_y$$

population total

$$\gamma = \left(\frac{N}{1}, \dots, \frac{N}{1}\right)' \iff \gamma'y = \mu_y$$

population mean

Any sample $i \subset \mathcal{U}$

Let i consists of the first n population elements. Then:

$$\mathbf{Y} = (\mathbf{Y}_i, \mathbf{Y}_{n \setminus i})'$$

$$\mathbf{y} = (\mathbf{y}_i, \mathbf{y}_{n \setminus i})'$$

$$\gamma = (\gamma_i, \gamma_{n \setminus i})'$$

Estimation target

$$\gamma' \mathbf{y} = \gamma'_i \mathbf{y}_i + \gamma'_{n \setminus i} \mathbf{y}_{n \setminus i}$$

is a realization of

$$\gamma' \mathbf{Y} = \gamma'_i \mathbf{Y}_i + \gamma'_{n \setminus i} \mathbf{Y}_{n \setminus i}$$

Estimation of $\gamma' \mathbf{Y} \iff$ prediction of $\gamma'_{n \setminus i} \mathbf{Y}_{n \setminus i}$.

The term "**prediction**" is used in the sense of making statistical guess about the values of $\mathbf{Y}_i$ which we have not seen – not in the sense of forecasting the future values.

Def. 1. A linear estimator of $\theta = \gamma'Y$ is defined as $\hat{\theta} = w'Y$,
 $w = (w_1, \dots, w_n)'$ – vector of coefficients.

Def. 2. The estimation error of an estimator $\hat{\theta} = w'Y$ is
 $\hat{\theta} - \theta = w'Y - \gamma'Y$.

Def. 3. General linear model M for Y :

$$\begin{aligned} E^M(Y) &= X\beta, \\ Var^M(Y) &= V, \end{aligned} \quad \begin{aligned} (1) \\ (2) \end{aligned}$$

$X = (X_{(1)}, \dots, X_{(J)})$ – matrix of auxiliary variables,
 β – vector of unknown parameters,
 V – positive definite variance-covariance matrix of Y .

Def. 4. The estimator $\hat{\theta}$ is unbiased for θ under a model M if $E^M(\hat{\theta} - \theta) = 0$.

Def. 5. The error variance of $\hat{\theta}$ under a model M is $Var^M(\hat{\theta}) = E^M(\hat{\theta} - \theta)^2$.

Def. 6. A linear model unbiased estimator $\hat{\theta}$ of θ which minimizes the error variance $Var^M(\hat{\theta})$ is called best linear unbiased predictor (BLU) of θ .

It's expression and error variance is known.

When X are quantitative variables the estimator $\hat{\beta}$ for β which gives BLU predictor of θ , is the least squares estimator.

Let $\gamma = (1, \dots, 1)'$: $\gamma'y = \gamma'y_i + \gamma'_{n \setminus i} y_{n \setminus i}$.

is equivalent to

$$t_y = \sum_{i \in I} y_i + \sum_{k \in U \setminus I} y_k.$$

Estimating $\gamma'y = t_y \iff$ predicting value $\sum_{k \in U \setminus I} y_k$ of unobserved random variable $\sum_{k \in U \setminus I} Y_k$

Taking various matrices of auxiliary variables X , matrices V and vectors β we obtain BLU predictors coinciding with the estimators of totals known in the design-based theory:

expansion estimator,

linear regression estimator,

stratified expansion estimator,

.....

Valliant, R, Dorfman A. H., and Royall R. M. (2000) *Finite Population sampling and Inference. A prediction Approach*. John Wiley & Sons

Prediction theory under the nonlinear model

R. Vaillant, *JASA*, 1985

Nonlinear model:

$$(3) \quad \mathbf{E}^M(Y_k) = f(\mathbf{X}_k; \psi), \quad k = 1, \dots, N$$

(4)

ψ – vector of J unknown coefficients, model parameter,

$f(\cdot; \cdot)$ – nonlinear function, has at least 3 partial derivatives with respect to ψ .

Remark: if model is such that

$Y_1, \dots, Y_N$ are independent random variables $\implies$

$\mathbf{V}$ is diagonal $\implies$

$Var^M(Y_k)$ is only needed for it in (4).

Population total

$$t_y = \sum_{i=1}^N y_i = \sum_{i \in I} y_i + \sum_{i \in N \setminus I} y_i.$$

When ψ is known, the BLU estimator of t_y is obtained by adding $\sum_{i \in I} y_i$ with BLU predictor of $\sum_{i \in N \setminus I} y_i$.

When ψ is not known – its estimator is inserted into BLU predictor of $\sum_{i \in N \setminus I} y_i$.

$\hat{t}_y$ becomes biased.

Valliant: If asymptotically normal and consistent estimator of the model parameters ψ is available

1) Estimator $\hat{t}_y$ is consistent

2) Expressions for

$$ABias(\hat{t}_y) \approx E_M(\hat{t}_y - t_y), \\ AVar(\hat{t}_y) \approx Var_M(\hat{t}_y - t_y).$$

are obtained

3) Consistent estimator of $Var_M(\hat{t}_y)$ constructed through the consistent estimator of ψ .

Models for study variable with positive values and many zeroes

Let us consider a study variable y , values of which have the properties:

$$y_k \begin{cases} = 0 & \text{or} \\ > 0, \end{cases} k = 1, \dots, N \Rightarrow \text{large}(y) \Rightarrow \text{large}(\widehat{\text{Var}}(t_y)) \text{ large.}$$

Examples:

Area under some kind of crop in the farm, which is not grown up often (rape),
Investment of enterprise for an environmental protection

Design-based estimators are not useful

Model-assisted estimators are not useful – there are no auxiliary variables correlated with the study variable because zero values can not be known in advance.

Model-based estimator

$\mathcal{U}^*$ – superpopulation or a model for y , $\mathcal{U}$ – its *random* realization,

$y_k, k = 1, 2, \dots, N, t_y$ – random variables

$$t_y = \sum_{k \in \mathcal{I}} y_k + \sum_{k \in \mathcal{U} \setminus \mathcal{I}} y_k, \quad \widehat{t_y} = \sum_{k \in \mathcal{I}} y_k + \sum_{k \in \mathcal{U} \setminus \mathcal{I}} \widehat{y_k}$$

Nonlinear prediction $\widehat{y_k} = ?$

1. F. Kalberg, *Journal of Official Statistics*, 2000

$$Y_k = \xi_k Y_k, \quad \tilde{Y}_k = \mathbf{x}_k' \boldsymbol{\beta} + \varepsilon_k, \quad \varepsilon_k \sim \mathcal{N}(0, \sigma^2)$$

$$\xi_k = \begin{cases} 1, & p_k \\ 0, & 1 - p_k \end{cases} \Leftrightarrow \widehat{\boldsymbol{\beta}}, \widehat{\sigma}, \widehat{p_k} \Rightarrow \widehat{y_k} = \widehat{\mathbf{E}_M Y_k}$$

Econometric models

2. Censored regression or Tobit model for a study variable y .

Let there are unobserved random variables

$$Y_k^* = \mathbf{x}_k' \boldsymbol{\beta} + \varepsilon_k, \quad \varepsilon_k \sim \mathcal{N}(0, \sigma^2), \quad k = 1, 2, \dots, N$$

$\varepsilon_k, \varepsilon_l$ independent for $k \neq l$.

Define random variables

$$Y_k = \begin{cases} Y_k^*, & \text{if } Y_k^* > 0, \\ 0, & \text{if } Y_k^* \leq 0, \end{cases} \quad k = 1, 2, \dots, N$$

W.H.Greene. Econometric Analysis. Prentice Hall, Upper Saddle River, 2003.

$$\mathbb{E}_M(Y_k) = f(\mathbf{x}_k; \boldsymbol{\psi}) = \mathbf{x}_k' \boldsymbol{\beta} \Phi\left(\frac{\sigma}{\mathbf{x}_k' \boldsymbol{\beta}}\right) + \sigma \phi\left(\frac{\sigma}{\mathbf{x}_k' \boldsymbol{\beta}}\right), \quad \Phi, \phi \sim \mathcal{N}(0, 1)$$

if $\mathbf{x}_k = (1, x_k)'$:

$$\boldsymbol{\psi} = (\sigma, \alpha_0, \alpha_1), \quad \alpha_0 = \beta_0/\sigma, \quad \alpha_1 = \beta_1/\sigma.$$

$$Var_M(Y_k) = \sigma^2 \Phi(\alpha_0 + \alpha_1 x_k) g_k(\alpha_0, \alpha_1), \quad k = 1, \dots, N$$

T. Amemiya, *Econometrica*, 1973.

Maximum likelihood estimators $\hat{\sigma}$, $\hat{\alpha}_0$, $\hat{\alpha}_1$ of tobit model parameters σ , α_0 , α_1 are consistent and asymptotically normal.

We estimate

$$\widehat{y_k^{(tobit)}} = \widehat{\mathbf{E}_M Y_k} = \mathbf{x}_k' \widehat{\boldsymbol{\beta}} \Phi \left(\frac{\widehat{\sigma}}{\widehat{\mathbf{x}_k' \boldsymbol{\beta}}} \right) + \sigma \phi \left(\frac{\widehat{\sigma}}{\widehat{\mathbf{x}_k' \boldsymbol{\beta}}} \right),$$

$$\widehat{t_y^{(tobit)}} = \sum_{k \in I} y_k + \sum_{k \in N \setminus I} \widehat{y_k^{(tobit)}}$$

Using results of Valliant we obtain expressions for

$$ABias(\widehat{t_y^{(tobit)}})$$

$$AVar(\widehat{t_y^{(tobit)}})$$

$$\widehat{Var}(\widehat{t_y^{(tobit)}})$$

SAS procedure LIFEREG: $\widehat{\boldsymbol{\beta}}$, $\widehat{\sigma}$.

The decisions can be separated:

to have positive values of Y_k or not to have?

how big will be Y_k if "yes"?

3. Censored regression model with the unobserved stochastic threshold (Heckman model)

Maddala G. S. Limited-dependent and qualitative variables in econometrics. Cambridge University Press, 1983.

Two unobserved random variables:

$$Y_k^* = \mathbf{x}_{(1)}^k \boldsymbol{\theta}_{(1)} + \varepsilon_k, \quad (5)$$

$$Z_k^* = \mathbf{x}_{(2)}^k \boldsymbol{\theta}_{(2)} + u_k, \quad (6)$$

$$\begin{pmatrix} \varepsilon_k \\ u_k \end{pmatrix} \sim \mathcal{N} \left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} \sigma_\varepsilon^2 & \sigma_{\varepsilon u} \\ \sigma_{\varepsilon u} & \sigma_u^2 \end{pmatrix} \right)$$

$$I_k = \begin{cases} 1, & \text{if } Z_k^* > 0, \\ 0 & \text{otherwise.} \end{cases} \quad Y_k = \begin{cases} Y_k^*, & \text{if } I_k = 1, \\ 0 & \text{otherwise.} \end{cases}$$

If both equations in (5) coincide $\Rightarrow$ Heckman $\equiv$ Tobit.

$$\mathbf{E}_M(Y_k) = f(\mathbf{x}_{(1)}^k, \mathbf{x}_{(2)}^k; \boldsymbol{\theta}_{(1)}^k, \boldsymbol{\theta}_{(2)}^k) = \Phi(\mathbf{x}_{(2)}^k \boldsymbol{\theta}_{(2)}^k) \left(\mathbf{x}_{(1)}^k \boldsymbol{\theta}_{(1)}^k + \sigma_{\varepsilon} \frac{\Phi(\mathbf{x}_{(2)}^k \boldsymbol{\theta}_{(2)}^k)}{\phi(\mathbf{x}_{(2)}^k \boldsymbol{\theta}_{(2)}^k)} \right)$$

Estimation of the model parameters

$$\widehat{y}_k^{(Heckman)} = \widehat{\mathbf{E}}_M(Y_K), \quad \widehat{t}_j^{(Heckman)} = \sum_{i \in I} y_i + \sum_{k \in \mathcal{U} \setminus I} \widehat{y}_k^{(Heckman)}$$

Other generalizations of nonlinear model (3), (4) when residuals are not normally distributed – skew distributions in economics.

Simulation

$N = 1\,000$ size finite opulation generated

$$\mathbf{x}_k = (1, x_k)', \quad x_k \sim \mathcal{N}(\mu, \sigma^2), \quad \varepsilon_k \sim \mathcal{N}(0, \sigma^2)$$

$$Y_k^* = \beta_0 + \beta_1 x_k + \varepsilon_k \Rightarrow Y_k$$

$M = 1\,000$ replicates:

Simple random sample of size n

Estimates: *design-based* and *tobit model-based*

$$\bar{t}_y^{(RS)} = \frac{n}{N} \sum_{k \in I} y_k = \bar{t}_y^{(tobit)} + \sum_{k \in N \setminus I} \tilde{y}_k^{(tobit)}$$

$$\tilde{y}_k^{(tobit)} = f(x_k; \hat{\sigma}, \hat{\beta}_0, \hat{\beta}_1), \quad k \in N \setminus I$$

Notations: $\widehat{\eta}_m$ and $\widehat{Var}(\widehat{\eta}_m)$ – replicates of the estimates η :

$$\widehat{\eta} = \widehat{t}_{(SRS)}^{\eta} \quad \text{or} \quad \widehat{\eta} = \widehat{t}_{(tobit)}^{\eta}$$

$$\bar{\eta} = \frac{1}{M} \sum_{m=1}^M \widehat{\eta}_m,$$

$$\begin{aligned} \widehat{Bias}_{emp}(\widehat{\eta}) &= \bar{\eta} - t, \\ \widehat{Var}_{emp}(\widehat{\eta}) &= \frac{1}{M} \sum_{m=1}^M (\widehat{\eta}_m - \bar{\eta})^2, \\ \widehat{MSE}_{emp}(\widehat{\eta}) &= \widehat{Bias}_{emp}(\widehat{\eta})^2 + \widehat{Var}_{emp}(\widehat{\eta}), \\ RMSE_{emp}(\widehat{\eta}) &= \sqrt{\widehat{MSE}_{emp}(\widehat{\eta})}, \end{aligned}$$

Investigated:

Dependence of $RMSE$ on n , σ , percentage of zeroes in the population.

$$\begin{aligned} \widehat{Bias}(\widehat{\eta}) &= \frac{1}{M} \sum_{m=1}^M \widehat{Bias}(\widehat{\eta}_m) \quad \text{for } \widehat{\eta} = \widehat{t}_{(tobit)}, \\ \widehat{Var}(\widehat{\eta}) &= \frac{1}{M} \sum_{m=1}^M \widehat{Var}(\widehat{\eta}_m), \\ \widehat{MSE}(\widehat{\eta}) &= \widehat{Bias}(\widehat{\eta})^2 + \widehat{Var}(\widehat{\eta}); \\ RMSE(\widehat{\eta}) &= \sqrt{\widehat{MSE}(\widehat{\eta})}. \end{aligned}$$

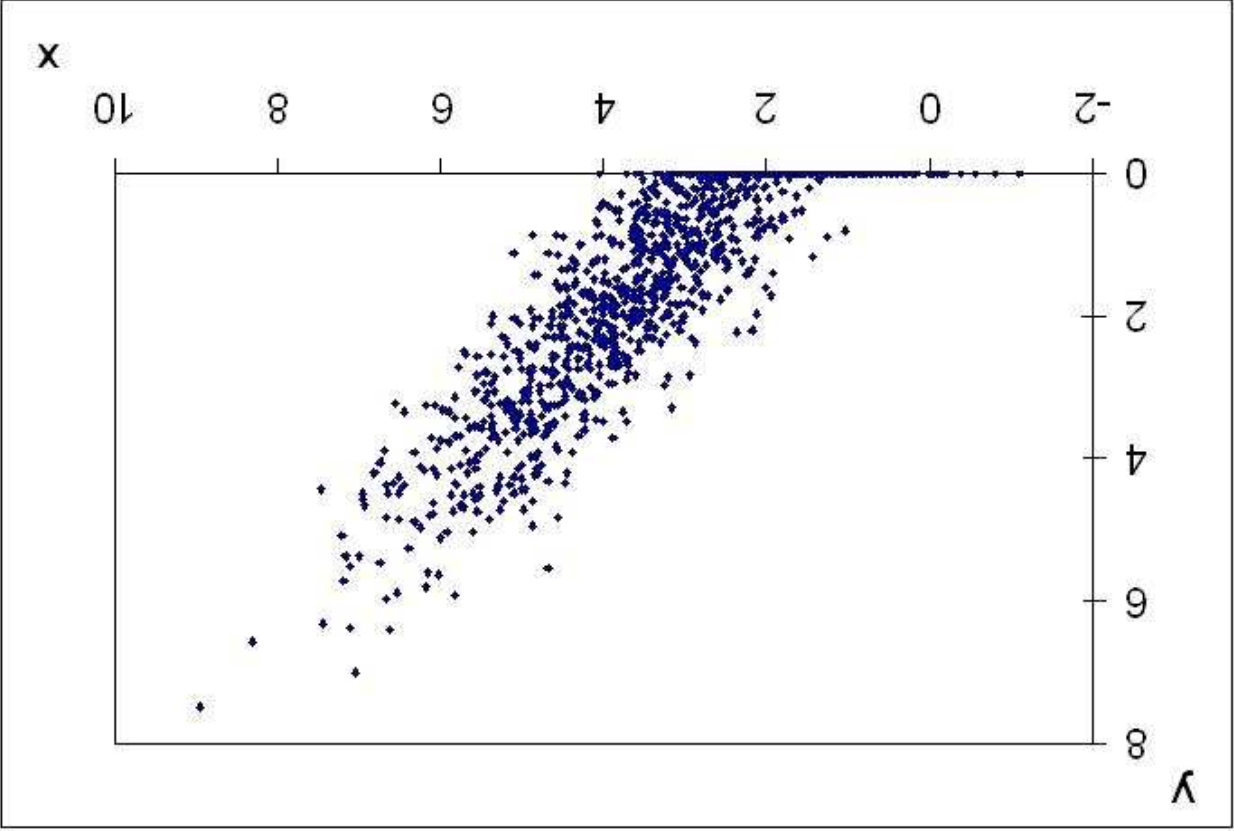
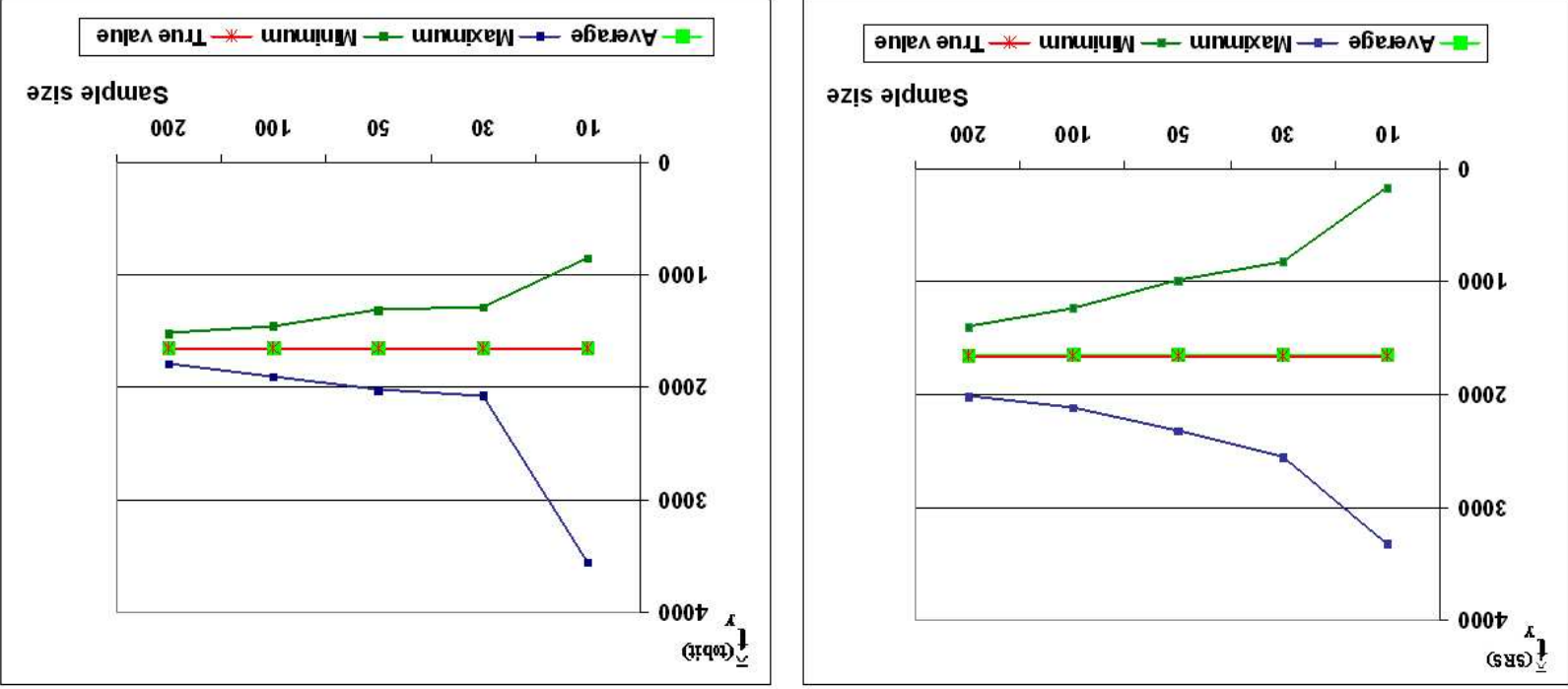


Fig. 1. Population

$\mathbf{x}_k = (1, x_k)'$, $x_k \sim \mathcal{N}(3.49, 1.56)$, $\varepsilon_k \sim \mathcal{N}(0, 0.8)$, $\theta = (-2.42, 1.1)$, 20% zeroes

$$Y_k^* = -2.42 + 1.1x_k + \varepsilon_k \Rightarrow Y_k$$

Fig. 2. Dependence of $\hat{t}_{ij}^{(SRS)}$ (left) and $\hat{t}_{ij}^{(toit)}$ (right) on the sample size n



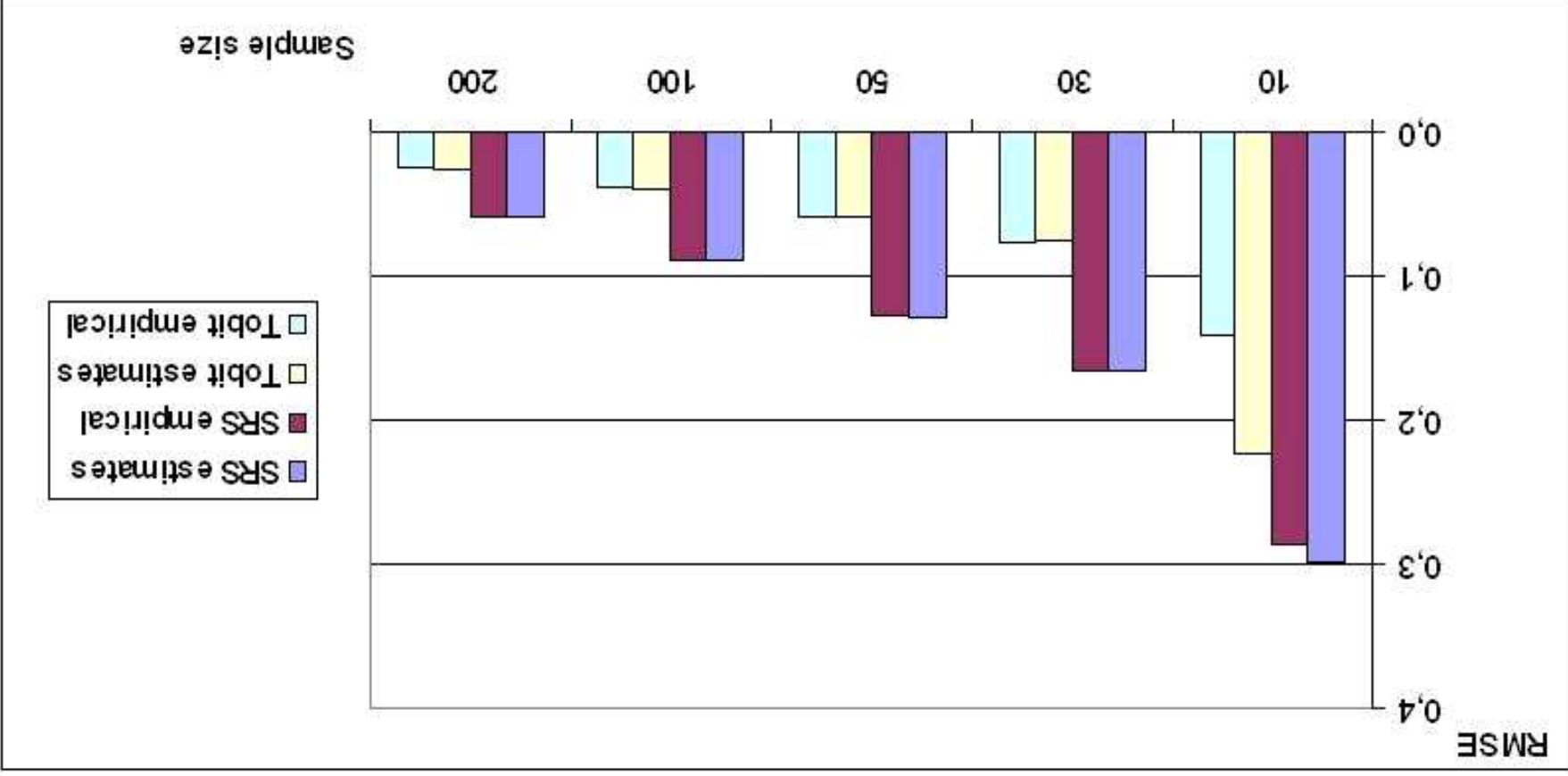


Fig. 3. Dependence of $RMSE(\hat{t}_y)$ on the sample size n

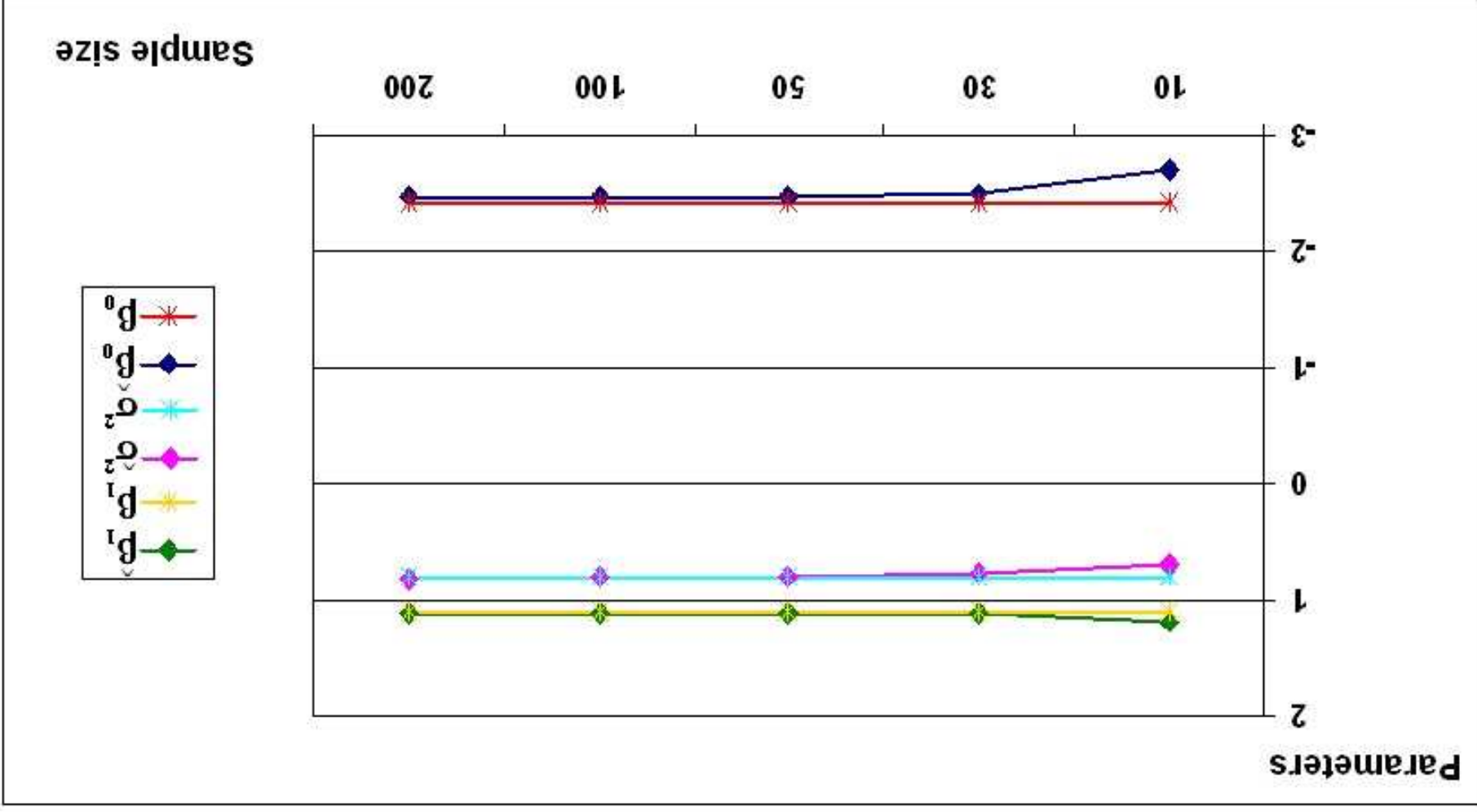


Fig. 4. Dependence of $\hat{\sigma}_2^2$, $\hat{\beta}_0$, $\hat{\beta}_1$ on the sample size n

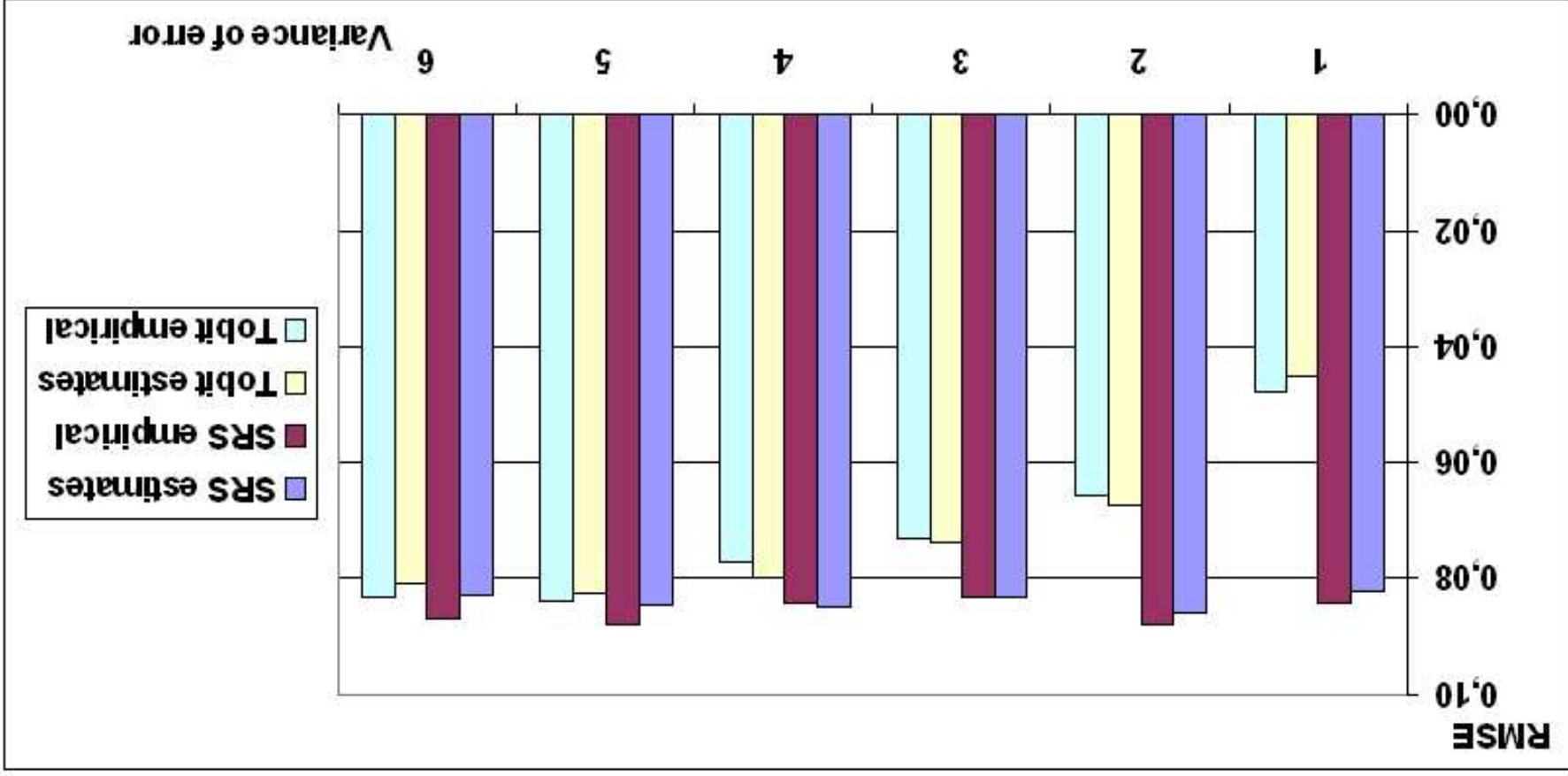


Fig. 5. Dependence of $RMS E(\hat{t}_y)$ on the variance of error σ ($\sigma \downarrow \Rightarrow \beta_0 \downarrow$)

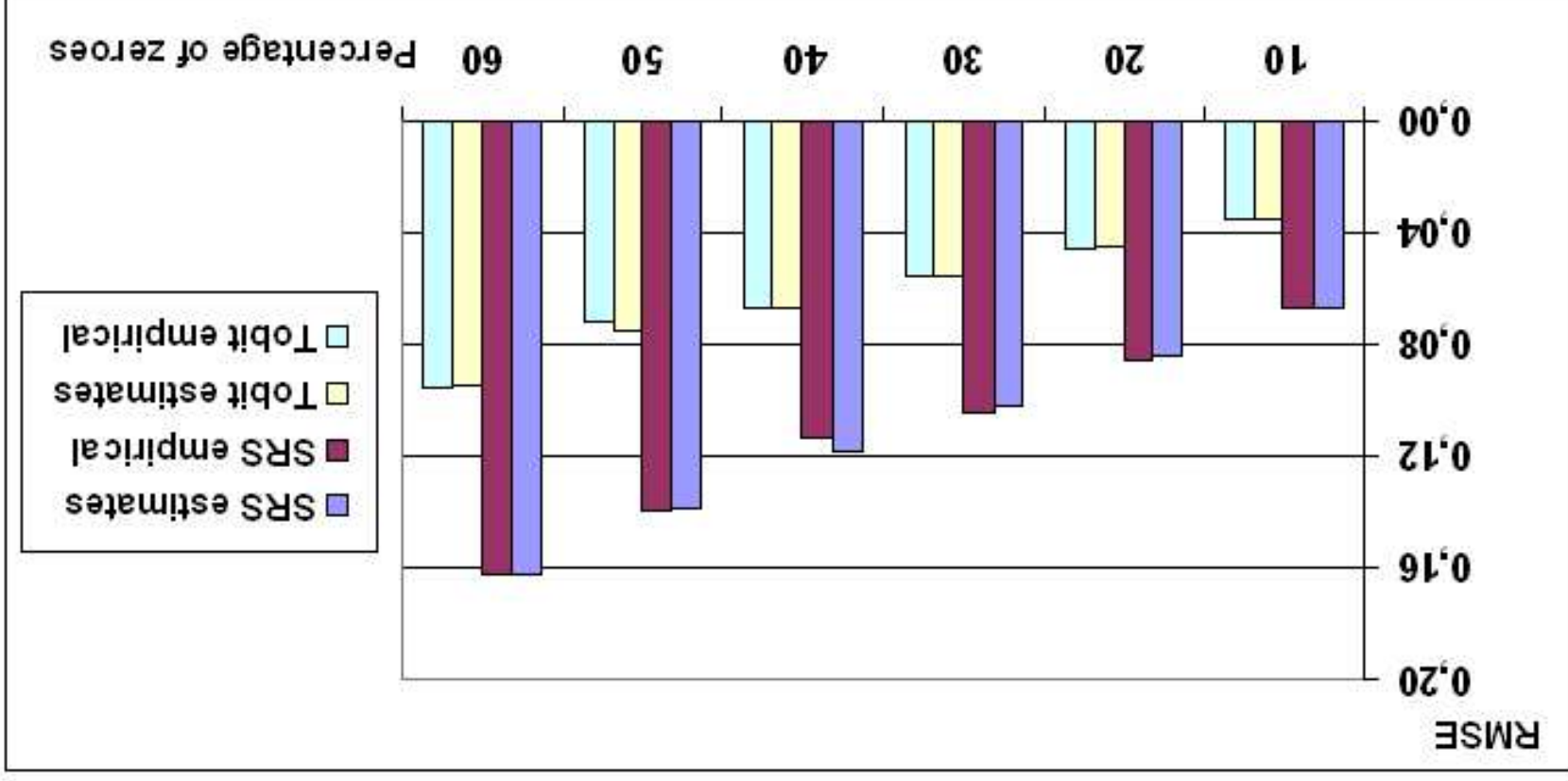


Fig. 6. Dependence of $RMSE(\hat{t}_y)$ on the percentage of zeros in the population (zeros $\downarrow \Rightarrow \mu \uparrow$)

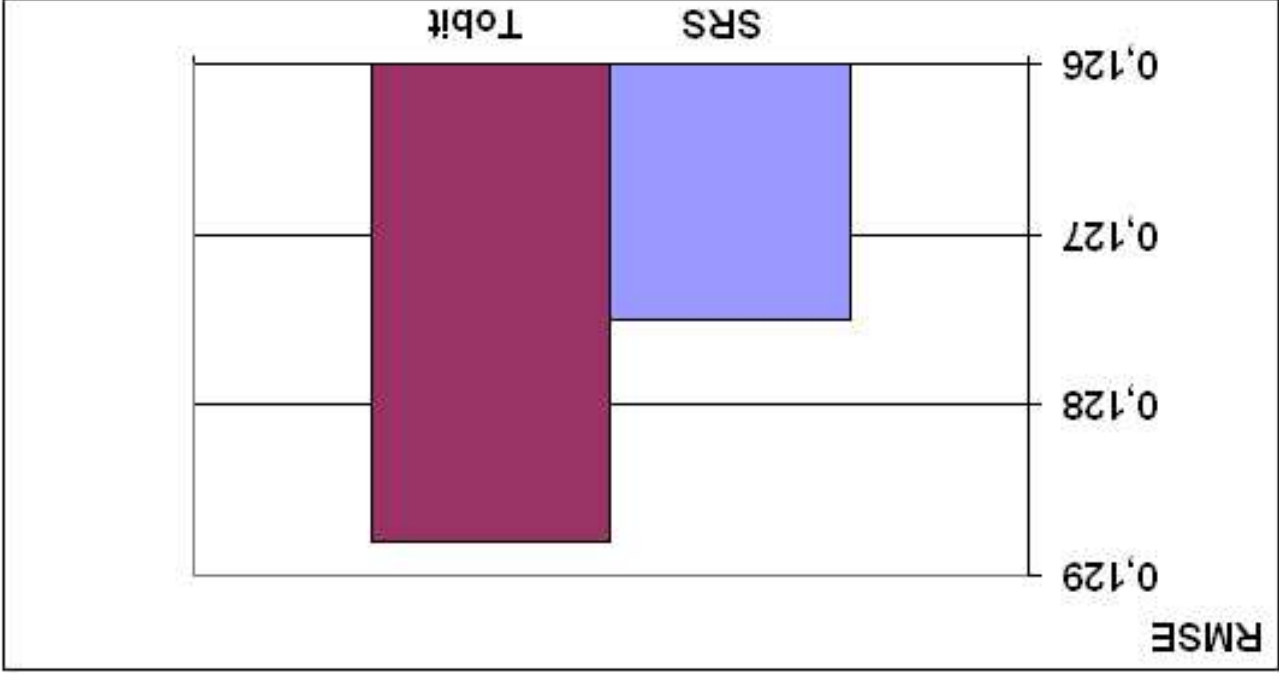


Fig. 7. The case when $R MSE(t_{y}^{tobit}) > R MSE(t_{y}^{SRS})$

$$\mathbf{x}_k = (1, x_k)', \quad x_k \sim \mathcal{N}(1, 2), \quad \varepsilon_k \sim \mathcal{N}(0, 6), \quad \theta = (0.5, 0.3), \quad 45\% \text{ zeroes}$$

$$y_k^* = 0.5 + 0.3x_k + \varepsilon_k \Rightarrow y_k$$

Conclusions

For small sample size ($n = 10$) average $RMSE$ and empirical $RMSE$ of model-based estimator can differ much.

The efficiency of the model-based estimator does not show here any dependence on the sample size or percentage of zeroes in the population.

The efficiency of the model-based estimator depends on the variance of the error σ .

Acknowledgements

Support: Nordic Council of Ministers

Library facilities: Umeå university

Thank you for your attention!