

“Natural disaster accounting bias and its equivalence across genetic resource stocks”

AUTHORS	Garth Holloway
ARTICLE INFO	Garth Holloway (2013). Natural disaster accounting bias and its equivalence across genetic resource stocks. <i>Environmental Economics</i> , 4(3)
RELEASED ON	Wednesday, 09 October 2013
JOURNAL	"Environmental Economics"
FOUNDER	LLC "Consulting Publishing Company "Business Perspectives"



NUMBER OF REFERENCES

0



NUMBER OF FIGURES

0



NUMBER OF TABLES

0

© The author(s) 2026. This publication is an open access article.

Garth Holloway (UK)

Natural disaster accounting bias and its equivalence across genetic resource stocks

Abstract

Two sources of bias arise in conventional loss predictions in the wake of natural disasters. One source of bias stems from neglect of accounting for animal genetic resource loss. A second source of bias stems from failure to identify, in addition to the *direct effects* of such loss, the *indirect effects* arising from implications impacting animal-human interactions. The author argues that, in some contexts, the magnitude of bias imputed by neglecting animal genetic resource stocks is substantial. The author shows, in addition, and contrary to popular belief, that the biases attributable to losses in distinct genetic resource stocks are very likely to be the same. The paper derives the formal equivalence across the distinct resource stocks by deriving an envelope result in a model that forms the mainstay of enquiry in subsistence farming and we validate the theory, empirically, in a World-Society-for-the-Protection-of-Animals application.

Keywords: natural disaster, accounting bias, equivalence across genetic resource stocks, direct effects, indirect effects, subsistence farming, animal-human interactions, World-Society-for-the-Protection-of-Animals application, hierarchical Bayesian methodology.

JEL Classifications: Q51, Q54, Q57.

Introduction

Subsistence farming systems comprise a vast swath of humanity (Singh, Squire and Strauss, 1986). One of the most important features of subsistence systems is their dependence on animal genetic resource stocks for their welfare (Campbell and Knowles, 2011). Because animal genetic resource stocks are welfare-enhancing, their existence bestows value within the households that employ them. And because they receive value, households are willing to relinquish units of other productive resources in order to mitigate loss in the welfare-enhancing animal resource stocks. These simple observations can be employed in order to develop robust methodology for valuing livestock loss within livestock dependent households. It is conjectured that such loss may be substantial; may bias relief-effort targets during post-disaster management; and, hitherto neglected in post-disaster social accounting exercises, need to be computed (Livestock Emergency Guidelines and Standards, 2009). These facts motivate a search for the values that subsistence households place on animal inputs. This search is especially important when it is realized that animal inputs promote productivity, enhance the surplus-generating potential of the household and can, as a consequence, promote immersion into markets that are contemporaneously constrained by thinness and instability. We consider the problem of placing formal economic valuations on livestock inputs in the context of a rich data set on milk-market participation by small-holder dairy producers in the Ethiopian highlands (Nicholson, 1997). We take up this search in the context of a familiar framework for household decision-making (Singh, Squire and

Strauss, 1986); formal development of a comprehensive, multi-dimensional hierarchical model (Good, 1980); and algorithmic developments that exploit fully the existence of known full conditional distributions for all of the relevant unknown quantities (Gelfand and Smith, 1990). Hence, the model and its fundamentals have roots set firmly in the computational advances that have been with us since the early 1990's (Gelfand, 2000). However, these advances have yet to be exploited to their full potential for the purpose of valuing animal genetic resource stock losses.

1. Crossbreeding programmes in the Ethiopian highlands and background to the study sites

Crossbreeding of imported animals with indigenous stock began in Ethiopia in 1968 (Kiwuwa, Trail, Kurtu, Worku, Anderson and Durkin, 1983). Since their introduction, performance has been closely monitored at a number of institutions. An extensive set of records on the performance of various crosses has been compiled. Crossbreed animals' enormous potential for increasing per-capita milk production is well-documented. For example, between 1968 and 1977 crosses between the local Arsi and Zebu animals with introduced stock produced significant increases in output, measured in terms of annual, fat-corrected milk yield (kg) per-unit metabolic weight of milking cows (kg). The main findings include sizable enhancement of yields over extant indigenous stock to the order of 75% in some instances. Thus, hybrid-vigor has been important in the Ethiopian highlands since the late 1960's. However, the advantages of increased yields has brought with it some intense demands on management; crossbreed animals are susceptible to local diseases such as anthrax, rinderpest, and blackleg; and adopters have often been required to change management

practices. Consequently, adoption rates have been lowered from what might have been expected given the initial, sizable increases in production yields (Brokken and Seyoum, 1990).

The ‘laboratory’ for investigation is a panel of observations on $N = 68$ sample units (households) at two respective sites (being $N_1 = 35$ at the *Ilu Kura peasant association* and $N_2 = 33$ at the *Mirti peasant association*), each, approximately one hundred miles, in opposite directions, from the capital, Addis Ababa. The panel periodicity is approximately four months, representing three visits in the same production year. At each visit, production output on the day in question is measured, along with a record of the amount of milk sold, the number of local-breed and crossbreed cows milked, and other relevant socio-demographic characteristics of the households. Thus the available data consist of the panel records across the households on selected covariates and on the two response variables ‘output’ and ‘sales.’ The selected covariates include the key livestock variables ‘*Crossbreed*’ and ‘*Local-breed*’ cows; two site-specific dummy variables, ‘*Ilu-Kura*’ and ‘*Mirti*,’ three intellectual-capital accumulators, ‘*Experience*,’ ‘*Education*’ and ‘*Extension*,’ and one sales-related covariate that we deem *a priori* significant in the decision to immerse in the market, namely ‘*Distance*.’ Respectively, *Output* and *Sales* refer to the amounts of fluid milk (in liters) that the household produced, respectively, sold, on the day that the interview was enacted; *Ilu-Kura*, respectively, *Mirti* are binary covariates assuming the value one if the household in question resided within the peasant association, and assuming the value zero, otherwise; *Experience* denotes the number of years of farming experience accumulated by the household head; *Education* refers to the total number of years of formal education accumulated by the household head; *Extension* refers to the number of times in the twelve months prior to the interview that the household was visited by an extension agent discussing either production or marketing activities; and *Distance* represents the amount of time (in minutes) that it takes the household to transport bucketed fluid milk to the milk-cooperative. Additional institutional and geographical background relevant to our investigation is available elsewhere (Staal, Delgado and Nicholson, 1997; Holloway, Barrett and Ehui, 2000; Holloway, Nicholson, Delgado, Ehui and Staal, 2000; Holloway and Ehui, 2001; Holloway, Barrett and Ehui, 2001; Holloway, Dorfman and Ehui, 2001; Holloway, Nicholson, Delgado, Ehui and Staal, 2004; Holloway, Barrett and Ehui, 2005; and Holloway, Teklu and Ehui, 2008).

2. Canonical implementations

Relegating detail to *the Appendix*, we enact loss predictions across the two genetic resource stocks in five steps. The first step derives the formal household production framework upon which the productivity estimates are grounded (Singh, Squire and Strauss, 1986); the second step derives the precise metrics upon which the loss estimates depend through the application of the standard calculus (Chiang and Wainwright, 2005) and the envelope theorem, exploiting Roy’s identity (Roy, 1947); the third step selects covariates for the empirical implementation following extensive models comparisons Chib, 1995); the fourth step generates empirical estimates of the respective genetic-resource stocks using matrix extensions (Bauwens, 1984; Drèze and Richard, 1983) of well-known adaptations of the normal-linear model (Lindley and Smith, 1972); and the fifth and final step generates the loss estimates by exploiting standard results for the normal-linear model and the resulting posterior predictive distributions (Zellner, 1971; Koop, 2003; Koop, Poirier and Tobias, 2008). The covariate specifications that we consider are four, which we refer to as ‘models’ wherein, we define *Model One*, as the model consisting of just the livestock inputs, *Crossbreed* and *Local-breed*; *Model Two*, consisting of the livestock inputs and the two site specific dummy variables, *Ilu-Kura* and *Mirti*; *Model Three*, consisting of the latter specification plus the additional covariates *Experience*, *Education*, *Extension* and *Distance*; and, finally, *Model Four*, consisting of just the livestock covariates, *Crossbreed* and *Local-breed*, and the single, site-specific dummy variable, *Mirti*. Across these ‘models’ we implement, in turn, three respective specifications. *Specification One* consists of the normal linear model (Zellner, 1971; Koop, 2003; Koop, Poirier and Tobias, 2008); *Specification Two* permits the constant terms in the linear model to vary in the usual, hierarchical manner (Koop, 2003; Koop, Poirier and Tobias, 2008; Good, 1980); and *Specification Three* permits all of the covariate coefficients to relate hierarchically (Lindley and Smith, 1972; Koop, 2003; Koop, Poirier and Tobias, 2008). With these three specifications and these four models at hand, the Cartesian product *Models* $\times$ *Specifications* leads to a total of *twelve* credible alternatives with which to assess the losses at interest. The results of the twelve experiments are reported, graphically, in Figure 1 (see Appendix B). The diagonal entries in the figure shaded in red represent the line of best fit among the data; the twelve alternative specifications represent reasonable scatter along the best-fit line and suggest, without discriminating further, that any single representation seems satisfactory as a form of explana-

tion across the observed responses. Some signs of possible bias are evident at larger outlier quantities, but, for the most part, the correlations between the scatter and the line of best fit seem satisfactory. Consequently, we turn to the problem of selecting a 'best' formulation with which to predict loss (Chib, 1995). An extensive model selection exercise suggests that the *a posteriori* preferred explanation for the response data resides in the second specification and in the model with the parsimonious covariate description consisting of the *Local-Breed*, *Crossbreed* and *Mirti* covariates. Indeed, a wide variety of alternative specifications enacted suggests that only these three covariates are consistently significant in explaining milk-output production within the Ethiopian-highlands households. Turning to the important matter of posterior prediction and assessing the actual magnitude of marginal productivities across the two distinct resource stocks, we summarize our essential findings in the contour plots in Figure 2. The cross-breed input leads to marginal productivities about twice the scale of those for the indigenous breed animals; produces estimates that are considerably more precise; but are, nonetheless, significantly correlated across the two, alternative, genetic resource stocks. In the absence of further adaptation, given the marginal-productivity estimates, we are able to construct loss estimates based purely on these inferences about animal productivity. However, additional, relevant sample information is hitherto ignored.

3. Extended implementations

One extensively documented aspect of the institutional setting (Singh, Squire and Strauss, 1986; Staal, Delgado and Nicholson, 1997; de Janvry, Fafchamps and Sadoulet, 1991; Sadoulet and de Janvry, 1995) is the inter-linkage between production and sales decisions. The notion that many households encounter barriers to entry, both pecuniary and non-pecuniary and, therefore, generate observations on milk-market participation in the sales dimension that are censored (Tobin, 1958) at some unobserved threshold motivates additional statistical scrutiny. A set of alternative censored regressions is developed and, relegating all of the essential mathematical details to the *supplementary materials section* we take up the search for an appropriate form in much the same way as we pursued across the single-equation, canonical forms, above. Thus, in like fashion, results for the matrix-Normal extension of the single-equation models evolve and produce the estimates of fit presented in Figure 3. Some minor improvements are apparent from comparison of Figures 1 and 3. The preferred model specification within the matrix system is found to be a standard censored-regression model (Tobin, 1958)

employing standard Bayesian procedures (Chib, 1992) in which the two equations, one pertaining to sales and the other to output, have hierarchical, household-specific constants that are correlated. This formulation produces the marginal productivity distributions reported in Figure 4 and which we use to translate into meaningful aggregates of loss computed on the 'national-annual scale.' Using this calculation, and, recalling that, by virtue of the facts that the marginal products which $\hat{\theta}_l$ and $\hat{\theta}_c$ depict are random variables and, thus, have distributions (recall Figures 2 and 4), so too will the quantity $\hat{\theta}$ which we derive as our estimate of the annual average cost of loss of livestock in US dollars. The joint distribution of the respective genetic stock loss estimates is depicted in Figure 5. Specifically, we determine that catastrophic loss of the order of one-hundred percent in grade indigenous stock and in cross-breed animals result in losses, of approximately \$US 3.84×10^{10} and \$US 3.37×10^{10} , respectively. That these measures represent significant quantities which, if neglected, could seriously bias social accounting exercises is apparent; and that the marginal distributions corresponding to the two dimensions of Figure 5 are quite comparable is also apparent. Thus, when nature occasionally gives cause to summon the social accountant two sources of bias arise in conventional loss predictions in the wake of natural disasters. We have shown these biases to be substantial and we have shown also that there is good reason to believe that the losses computed across distinct genetic resource stocks are approximately the same.

Conclusion

Two sources of bias arise in conventional loss predictions in the wake of natural disasters. One source of bias stems from neglect of accounting for animal genetic resource loss. A second source of bias stems from failure to identify, in addition to the *direct effects* of such loss, the *indirect effects* arising from implications impacting animal-human interactions. We have argued that, in some contexts, the magnitude of bias imputed by neglecting animal genetic resource stocks is substantial. We have shown, in addition, and contrary to popular belief, that the biases attributable to losses in distinct genetic resource stocks are very likely to be the same. We have derived the formal equivalence across the distinct resource stocks by deriving an envelope result in a model that forms the mainstay of enquiry in subsistence farming and we have validated the theory, empirically, in a World-Society-for-the-Protection-of-Animals application in the Ethiopian Highlands. Further work should investigate the results derived herein to a wide and broader set of contexts.

Acknowledgements

In memoriam: Arnold Zellner, January 2, 1927-August 11, 2010. I am eternally grateful to Dale Poirier for sharing my thesis research with, and introducing me to, Arnold Zellner at the inaugural meetings of the International Society for Bayesian Analysis, San Francisco, 1992. I am also indebted to Rodger and Sheila Hayward Smith for their hospitality in Barnes, London, England during phases of this research. The principal contribution that this research reports arises as a contribution to the World Society

for the Protection of Animals (©WSPA | Working towards a world where animal welfare matters and animal cruelty ends). The contribution is dedicated to the author's 'tutors' in statistics, John J. Deely and Wesley O. Johnson. The author is also grateful to Francois Bouchetoux, Donald Lacombe, Mark Laffey, Aman Ramali, Lise Tole and Marinos Tsigas who provided comments on a penultimate draft. And, finally, all data and computer codes used in the analysis are available from request to the author, whom, alone, is responsible for any remaining errors.

References

1. Albert, J., Chib, S. (1993). Bayesian Analysis of Binary and Polychotomous Response Data, *Journal of the American Statistical Association*, 88, pp. 669-79.
2. Bauwens, L. (1984). Bayesian Full Information Analysis of Simultaneous Equations Models Using Integration by Monte Carlo, Springer-Verlag Lecture Notes in Economics and Mathematical Systems, Volume 232, M. Beckman and W. Krelle (eds.). Berlin: Springer-Verlag.
3. Becker, G. (1965). A Theory of the Allocation of Time, *Economic Journal*, 75, pp. 493-517.
4. Brokken, R., Seyoum, S. (1990). Dairy Marketing in Sub-Saharan Africa. Proceedings of a Symposium at ILCA, Addis Ababa, Ethiopia.
5. Campbell, R., Knowles, T. (2011). The economic impacts of losing livestock in a disaster. A report for the World Protection of Animals (WSPA) by Economists at Large, World Society for the Protection of Animals, London.
6. Chiang, A., Wainwright, K. (2005). Fundamental Methods of Mathematical Economics, Fourth Edition. Boston: McGraw Hill International Edition.
7. Chib, S. (1995). Marginal Likelihood from the Gibbs Output, *Journal of the American Statistical Association*, 90, pp. 1313-1321.
8. Chib, S. (1992). Bayes Inference in the Tobit Censored Regression Model, *Journal of Econometrics*, 52, pp. 79-99.
9. Dawid, S. (1998). The infinite regress and its conjugate analysis, in, Bayesian Statistics 3 J.M. Bernardo, M.H. De Groot, D.V.M. Lindley and A.F.M. Smith (eds.), Oxford: Clarendon Press, pp. 95-110.
10. Deaton, A., Muellbauer, J. (1980). Economics and Consumer Behavior. Cambridge: Cambridge University Press.
11. De Janvry, A., Fafchamps, M., Sadoulet, E. (1991). Peasant household behaviour with missing markets: some paradoxes explained, *The Economic Journal*, 101, pp. 1400-1417.
12. Drèze, J., Richard, J.-F. (1983). Bayesian Analysis of Simultaneous Equations Systems in Handbook of Econometrics, in: Z. Griliches & M.D. Intriligator (eds.), Handbook of Econometrics, Edition 1, Volume 1, Chapter 9, Amsterdam: Elsevier, pp. 517-598.
13. Fair, A. (1978). A Theory of Extramarital Affairs, *Journal of Political Economy*, 86, pp. 45-61.
14. Gelfand, A. (2000). Gibbs Sampling, *Journal of the American Statistical Association*, 95, pp. 1300-1304.
15. Gelfand, A., Smith, A.F.M. (1990). Sampling based approaches to calculating marginal densities, *Journal of the American Statistical Association*, 85, pp. 398-409.
16. Good, I. (1980). Some History of the Hierarchical Bayesian Methodology, in J.M. Bernardo, M.H. de Groot, D.V. Lindley and A.F.M. Smith (eds.) Bayesian Statistics, Valencia: University Press, pp. 489-519.
17. Holloway, Garth J., Barrett, Christopher B., and Ehui, Simeon K. (2000). Innovation and Market Creation, *Journal of the American Statistical Association*, Proceedings of the Section on Bayesian Statistical Science, pp. 148-153.
18. Holloway, G.J., C. Nicholson, C. Delgado, S. Ehui and S. Staal (2000). Agroindustrialization through Institutional Innovation: Transactions Costs, Cooperatives and Milk-Market Development in the East African Highlands, *Agricultural Economics*, 23, pp. 279-88.
19. Holloway, G.J. and S.K. Ehui (2001). Demand, Supply and Willingness-to-Pay for Extension Services in an Emerging-Market Setting, *American Journal of Agricultural Economics*, 83, pp. 764-68.
20. Holloway, G.J., C.B. Barrett and S. Ehui (2001). Cross-Bred Cow Adoption and Milk-Market Participation in a Multivariate, Count-Data Framework, Eurostat Special Issue: Bayesian Methods with Applications to Science, Policy and Official Statistics, pp. 233-42.
21. Holloway, G.J., J. Dorfman and S. Ehui (2001). "Tobit Estimation With Unknown Point of Censoring With An Application To Milk-Market Participation in The Ethiopian Highlands", *Canadian Journal of Agricultural Economics*, 49, pp. 293-311.
22. Holloway, G., C. Nicholson, C. Delgado, S. Ehui and S. Staal (2001). How To Make A Milk Market: A Case Study From The Ethiopian Highlands. Socio-Economic and Policy Research Working Paper 28, Addis Ababa: International Livestock Research Institute.
23. Holloway, G. and S. Ehui (2002). Expanding Market Participation Among Smallholder Livestock Producers: A Collection of Studies Employing Gibbs Sampling and Data from the Ethiopian Highlands (1998-2001), Addis Ababa: International Livestock Research Institute Working Paper Series.

24. Holloway, G.J., C. Nicholson, C. Delgado, S. Ehui and S. Staal (2004). A Revised Tobit Procedure for Mitigating Bias in the Presence of Non-Zero Censoring, *Agricultural Economics*, 31(1), pp. 97-106.
25. Holloway, G.J., Barrett, C., Ehui, C. (2005). "Bayesian Estimation of the Double Hurdle Model in the Presence of Fixed Costs, *Journal of International Agricultural Trade and Development*, 1(1), pp. 17-28.
26. Holloway, G.J., Ehui, S., Teklu, A. (2008). Bayes Estimates of Distance to Market: Transactions Costs, Cooperatives and Milk-Market Development in the Ethiopian Highlands, *Journal of Applied Econometrics*, 23 (5), pp. 683-696.
27. Holloway, G.J. (2013). Valuing Livestock Loss Within Livestock-dependent Households. Mimeograph, University of Reading.
28. Kiwuwa, G., Trail, J., Kurtu, M., Worku, G., Anderson, F., Durkin, J. (1983). Cross-Breed Dairy Cattle Productivity in Arsi Region, Ethiopia. International Livestock Center for Africa (ILCA) Research Report No. 11, Addis Ababa: ILCA.
29. Koop, G. (2003). *Bayesian Econometrics*, New York: Wiley.
30. Koop, G., Poirier, D., Tobias, J. (2008). *Bayesian Econometric Methods. Econometric Exercises*, Cambridge: Cambridge University Press.
31. Linardakis, M., Dellaportas, P. (2003). Assessment of Athens's metro passenger behavior via a multiranked probit model, *Applied Statistics*, 52, pp. 185-200.
32. Lindley, D., Smith, A.F.M. (1972). Bayes Estimates for the Linear Model, *Journal of the Royal Statistical Society (Series B Methodological)*, 34 (1), p. 41.
33. Livestock Emergency Guidelines and Standards Project (2009). Livestock Emergency Guidelines and Standards, Rugby: Practical Action Publishing.
34. McCulloch, G., Rossi, P. (2000). Reply to Nobile, *Journal of Econometrics*, 99, pp. 347-48.
35. Mood, A., Graybill, F., Boes, D. (1974). *Introduction to the Theory of Statistics*, New York: McGraw-Hill, third edition.
36. F. Nelson (1977). Censored regression models with unobserved stochastic censoring thresholds, *Journal of Econometrics*, 6, pp. 309-327.
37. Nicholson, C. (1997). The Impact of Milk Groups in the Shewa and Arsi Regions of Ethiopia: Project Description, Survey Methodology, and Collection Procedures. Mimeograph, Livestock Policy Analysis Project, International Livestock Research Institute, Addis Ababa.
38. Nobile, A. (2000). Comment: Bayesian multinomial probit models with a normalization constraint, *Journal of Econometrics*, 99, pp. 335-345.
39. Rothenberg, T. (1963). Bayesian Analysis of a Simultaneous Equation System, Economic Institute Report 6315, Erasmus Universiteit, Rotterdam.
40. Roy, R. (1947). La Distribution du Revenu Entre Les Divers Biens, *Econometrica*, 15, pp. 205-225.
41. Sadoulet, E. de Janvry, A. (1995). *Quantitative Development Policy Analysis*. London: The John Hopkins University Press.
42. Singh, I. Squire, L. Strauss, J. (1986). *Agricultural Household Models: Extension, Application and Policy*, Baltimore: John Hopkins University Press.
43. Staal, S., Delgado, C., Nicholson, C. (1997). Smallholder Dairying Under Transactions Costs in East Africa, *World Development*, 25, pp. 779-794.
44. Tobin, J. (1958). Estimation of relationships for limited dependent variables, *Econometrica*, 26, pp. 24-36.
45. Zellner, A. (1971). *An Introduction to Bayesian Inference in Econometrics*, New York: Wiley Classics Library Edition.

Appendix A

The motivating framework underlying conceptual developments is the household production model (Singh, Squire and Strauss, 1986) which is the *sine qua non* of modern development economic investigations and forms a mainstay of models within which costs of time predominate (Becker, 1965; Fair, 1978; Deaton and Muellbauer, 1980). Consisting of a maximand defined over utilities derived from consuming a home-produced good, a market-produced good, and leisure; optimization occurs with choices made subject to the restrictions that expenditure cannot exceed disposable income, that leisure and work combined cannot exceed the total stock of time available to the household, and that physical output cannot exceed its technological possibilities. With the primal model (Chiang and Wainwright, 2005) at hand, we enact the Lagrangean method (Chiang and Wainwright, 2005) and consider changes in endowments and their impact on indirect utility and, consequently, welfare. By considering one of these endowments to be livestock, the Envelope Theorem (Chiang and Wainwright, 2005), Roy's Identity (Roy, 1947), and some further standard calculus (Chiang and Wainwright, 2005) yield the indirect values sought, namely $\Delta\omega = p_a f_\ell(\cdot, \ell) \Delta\ell$. Here, $\Delta\omega$ denotes the 'change in income' consequent upon a 'change in the animal genetic resource stock' or ' $\Delta\ell$ '; ' p_a ' denotes 'the per-unit price of milk output'; ' $f_\ell(\cdot, \ell)$ ' denotes the 'marginal productivity of livestock, ' ℓ ', in the production enterprise ' $f(\ell)$ '; and ' $\Delta\ell$ ' denotes the total change in the livestock resource that is created by the investigator in order to simulate 'catastrophe.' Given each of the elements on the right hand side of the equality ' $\Delta\omega = p_a f_\ell(\cdot, \ell) \Delta\ell$ ', the left-hand side is available immediately. The term ' p_a ' is publicly available at each of the two study sites; the term ' $\Delta\ell$ ' is selected and within the control of the investigator; but the term ' $f_\ell(\cdot, \ell)$ ' must be estimated. Considerable econometric efforts are devoted to the latter objective in the knowledge that these parametric reports may be very sensitive to alternative model specifications. The extensive model-selection search yields a preferred specification which, in turn, yields precise esti-

mates of the respective marginal productivities of cross-breed and grade, indigenous stocks. The matrix system producing the estimates in Figure 3 is based on modifications to the normal linear model (Zellner, 1971; Koop, 2003; Koop, Poirier and Tobias, 2008), including extensions to consider multiple decision-making (Drèze and Richard, 1983; Bauwens, 1984), extensions to consider both standard and non-standard censoring mechanisms (Nelson, 1977; Chib, 1992), and extensions of the basic Gibbs-sampling algorithm using the powerful marginal-likelihood-from-the-Gibbs-sampling methodology (Chib, 1995). Finally, in order to translate the household-specific, marginal-product estimates, derived from the econometric estimates at the two sample sites, into a more meaningful aggregate per-annum metric, we compute $\hat{\theta} \equiv \{p_1 \times (N_1/N) + p_2 \times (N_2/N)\} \times \{(\hat{\theta}_1 \times (N_l/N)) + (\hat{\theta}_c \times (N_c/N))\} \times \sigma_{s:pa} \times \sigma_{pa:w} \times \sigma_{w:e} \times \sigma_{eb:us} \times \sigma_{d:y}$, where $\hat{\theta}$ is the loss estimate; p_1 denotes the Ethiopian birr price in the *Ilu Kura* peasant association; N_1 denotes the sample size in the *Ilu Kura* peasant association; p_2 denotes the Ethiopian birr price in the *Mirti* peasant association; N_2 denotes the sample size in the *Mirti* peasant association; N denotes the total sample size; $\hat{\theta}_l$ denotes the marginal productivity of local-breed cows; N_l denotes the total number of local-breed cows employed; $\hat{\theta}_c$ denotes the marginal productivity of cross-breed cows; N_c denotes the number of cross-breed animals employed; N_e denotes the total number of livestock employed; $\sigma_{s:pa}$ denotes the scale factor transforming the sample to the peasant association; $\sigma_{pa:w}$ denotes the scale factor transforming the peasant association to the wereda; $\sigma_{w:e}$ denotes the scale factor transforming the wereda to the Ethiopian geographic aggregate; $\sigma_{eb:us}$ denotes the scale factor transforming Ethiopian birr into US dollars; and, finally, $\sigma_{d:y}$, the ‘temporal aggregate’ denotes the transformation from days into years. Applying these aggregations we arrive at the final economic loss estimates depicted in Figure 5.

Details of the explicit procedures evolve from the following section summary, which subdivides into the respective subsections entitled ‘background,’ ‘notation,’ ‘density development,’ ‘observational equations,’ ‘specifications,’ ‘models,’ ‘priors,’ ‘parameter estimation,’ ‘models comparisons,’ and ‘marginal likelihood computation.’

1. Background

Implementation is facilitated by amending well-known results in three, small, but important, ways. First, we extend to the *matrix-Normal* framework *hierarchical* developments (Drèze and Richard, 1983) defined, specifically, for the vector-Normal-linear model. Second, we modularize the process known as *completing-the-square*, extending the vector-valued Normal form, as it appears, for example, in (Zellner, 1971), to an automated, matrix-Normal presentation. Third, we modify, slightly, the *basic marginal-likelihood identity* as outlined in (Chib, 1995) which forms the mainstay of models comparisons. The first modification facilitates hierarchical developments in the multiple-equation setting; the second modification is convenient for encoding the Gibbs-sampling algorithms, in debugging computing routines and in automating their descriptions; the third modification circumvents problematic integrations defined over latent data during marginal likelihood evaluation. In these contexts, in addition to (Zellner, 1971; Drèze and Richard, 1983; and Chib, 1995), some familiarity with conjugate developments in reduced-form multiple-equations Normal-data systems, as appears, for example, in Bauwens (1984) and in Drèze and Richard (1984) is desirable. Primers for the scalar-Normal and vector-Normal derivations presented within this *Appendix* are Zellner (1971, Koop (2003), and Koop, Poirier and Tobias (2008).

2. Notation

By way of notation we use lower-case Greek and Roman numerals to reference scalar quantities, use emboldened lower-case symbols to reference vectors and use emboldened upper-case symbols to reference matrix quantities. Thus, let $\theta \equiv (\theta_1, \theta_2, \dots, \theta_N)'$ denote a vector of parameters of interest, where “ $'$ ” denotes the ‘transpose’ of the column vector θ , $\pi(\theta)$ denotes the prior probability density function (*pdf*) for θ , and $\pi(\theta|y)$ the posterior pdf for θ , where $y \equiv (y_1, y_2, \dots, y_N)'$ denotes data. Frequently, we reference the data generating model $f(y|\theta)$, which is the likelihood function when viewed as a function of θ and, sometimes, make use of variants of the $f(\cdot|\cdot)$ notation in order to reference particular probability density functions. Occasionally we find it useful to reference just the variable part of the density (integrating constant excluded) in which case we use the symbol ‘ $\propto$ ’ to denote ‘is proportional to.’ In view of the prior-to-posterior conjugacy shared by each model that we consider, we adopt the notational convention employed by (Drèze and Richard, 1983) wherein postscripts indicated “ o ” reflect prior information and postscripts indicated “ * ” reflect posterior information; accordingly $f(\theta|\theta^o) \equiv \pi(\theta)$ and $f(\theta|\theta^*) \equiv \pi(\theta|y)$. Additionally, we will find it useful, to refer separately to the observed responses, which we denote Y ; distinguish between the observed responses and those that are latent, which we denote, Z ; and distinguish, again, between the observed and latent responses and another, we reference, when the observed and latent responses are combined, which we denote V . The exact dimensions of the response quantities, Y , Z and V will become apparent when their model-specific dimensions are defined, subsequently. Finally, we use indices $i = 1, 2, \dots, N$, to reference the households in question, where, we remind the reader, $N = 68$; use $t = 1, 2, \dots, T$ in order to reference periods within the ‘panel,’ in which $T = 3$; and use S to denote the sample collection which is $S = N \times T = 204$.

3. Probability density functions

We use three probability density functions. The first, which we use to model correlation in equation errors, is the *inverted-Wishart distribution*, namely, $f_{M \times M}^W(\Sigma | S, \nu) \equiv 2^{-5M} \times \pi^{-25M(M-1)} \times \prod_{i=1}^M \Gamma(.5(\nu+1-i))^{-1} \times |S|^{.5\nu} \times |\Sigma|^{-.5(\nu+M+1)} \times \exp\{-.5 \text{trace} \Sigma^{-1} S\}$. The second, which we use to model covariate response, is the *matrix-Normal distribution*, which we specify as $f_{N \times M}^{mN}(\Theta | \hat{\Theta}, \Sigma, \Omega) \equiv (2\pi)^{-.5NM} \times |\Sigma|^{-.5N} \times |\Omega|^{-.5M} \times \exp\left\{-.5 \text{trace} \Sigma^{-1} (\Theta - \hat{\Theta})' \times \Omega^{-1} (\Theta - \hat{\Theta})\right\}$. Finally, we make use of the uniform distribution $f_{1 \times 1}^U(\alpha | \beta, \gamma) \equiv (\gamma - \beta)^{-1}$. With reference to $f_{M \times M}^W(\Sigma | S, \nu)$, we emphasize that Σ has dimension $M \times M$; with reference to $f_{N \times M}^{mN}(\Theta | \hat{\Theta}, \Sigma, \Omega)$, we emphasize that Θ has dimension $N \times M$; and with reference to $f_{1 \times 1}^U(\alpha | \beta, \gamma)$, we emphasize that α is a scalar. Occasionally, we reference the multivariate-Normal and the univariate-Normal distributions, which are special cases of $f_{N \times M}^{mN}(\Theta | \hat{\Theta}, \Sigma, \Omega)$, wherein $M = 1$ in the first case and $N = M = 1$, in the second. Frequently, we reference the covariance matrix of dimension $NM \times NM$, corresponding to the *column expansion* of Θ , of dimension $N \times M$, which is $\Sigma \otimes \Omega$, and where ' $\otimes$ ' denotes the Kronecker product. Finally, we make repeated use of the well-known transformation property corresponding to the matrix-Normal (see, for example, 17, p. 72, and the references cited there), namely that, given $f_{N \times M}^{mN}(\Theta | \hat{\Theta}, \Sigma, \Omega)$, and the transformation $\Lambda \equiv A \Theta B$, where A is a $P \times N$ matrix of rank $P \leq N$ and B is an $M \times Q$ matrix of rank $Q \leq M$; then $\Lambda \equiv A \Theta B$ has distribution $f_{P \times Q}^{mN}(\Lambda | A\hat{\Theta}B, B'\Sigma B', A\Omega A')$.

4. Observational equations

Each of the estimating models we consider can be specified as variants of the observational equation

$$F = G \times \mathcal{G} + U, \quad (\text{A.1})$$

where $F \equiv (f_1, f_2, \dots, f_M)$ denotes an $S \times N$ collection of responses, $G \equiv (g_1, g_2, \dots, g_R)$ is an $N \times R$ collection of observations on known covariates, $\mathcal{G} \equiv (\phi_1, \phi_2, \dots, \phi_M)$ is an $R \times M$ collection of unknown covariate responses and $U \equiv (u_1, u_2, \dots, u_M)$ is an $S \times M$ collection of random disturbances. In addition, and retained throughout as a *maintained hypothesis*, we assume that the disturbance matrix U has the matrix-Normal distribution, $f_{N \times M}^{mN}(U | O_{S,M}, \Sigma, I_S)$, where $O_{S,M}$ denotes the $S \times M$ -dimensional null matrix and where I_S denotes the S -dimensional identity matrix.

5. Specifications

The total number of variants that we consider is *twelve*, being the Cartesian product *Specifications* $\times$ *Models* in which 'Specifications' refers to assumptions about the hierarchical structures and in which 'Models' refers to assumptions about the censoring mechanisms. On the first count, *Specification One*, assumes that hierarchical structures are non-existent; in this case $G \equiv X \equiv (x_1, x_2, \dots, x_K)$ of dimension $S \times K$ defines a (dense) covariate matrix of observations on relevant covariates; the corresponding response matrix is $\mathcal{G} = \Psi \equiv (\psi_1, \psi_2, \dots, \psi_M)$, of dimension $K \times M$; and we place investigator-specific priors on the distributions for Σ and Ψ . We detail the prior information used subsequent to introducing the remaining specifications and the various censoring mechanisms. *Specification Two* introduces *hierarchical constants*, in which case G is partitioned into $G \equiv [W X]$ where $W \equiv (w_1, w_2, \dots, w_M) \equiv I_N \otimes t_T$, of dimension $S \times N$, is a binary matrix of unit vectors; $X \equiv (x_1, x_2, \dots, x_K)$ retains dimension $S \times K$, as before; we partition $\mathcal{G}$ as $\mathcal{G} \equiv [\Xi \Psi]$, of dimension $S \times M$, where $\Xi \equiv (\xi_1, \xi_2, \dots, \xi_M)$ has dimension $N \times M$; we assume, in the hierarchical spirit, that Ξ , in turn, evolves according to $f_{N \times M}^{mN}(\Xi | H\Gamma, \Sigma, \text{Cov}\Xi_0)$, where $H \equiv I_M \otimes t_T$, $\text{Cov}\Xi_0$ denotes an investigator-induced component of the covariance matrix $\Sigma \otimes \text{Cov}\Xi_0$ corresponding to the *column expansion* of Ξ ; and we place investigator-induced priors on the matrices Σ , Γ and Ψ . Finally, *Specification Three*, assumes that all of the regression coefficients evolve hierarchically, in which case we restructure G into $P \equiv \text{block-diagonal}\{X\}$ where $\text{block-diagonal}\{X\} \equiv \text{diagonal}\{X_1', X_2', \dots, X_N'\}'$ is the $NT \times NK$ block-diagonal arrangement of the household-common components of X with typical element, X_i , a matrix of dimension $T \times K$; we redefine $\mathcal{G}$ as $\mathcal{G} \equiv \Delta$ where $\Delta \equiv (\delta_1, \delta_2, \dots, \delta_M)$, of dimension $NK \times M$; we assume that Δ , in turn, evolves according to $f_{NK \times M}^{mN}(\Delta | Q\Psi, \Sigma, \text{Cov}\Delta_0)$, $Q \equiv t_N \otimes I_K$, $\text{Cov}\Delta_0$ denotes an investigator-induced component of the covariance matrix $\Sigma \otimes \text{Cov}\Delta_0$ corresponding to the *column expansion* of Δ ; and we place investigator-induced priors on the matrices Σ and Ψ .

6. Models

Turning to specializations of the separate specifications, we enact four distinct formulations, henceforth referred to as Models in order to accommodate the censoring which is prevalent within the data. For this purpose, introduce $T \equiv \{\tau_{ij}, i = 1, 2, \dots, N, j = 1, 2, \dots, T\}$ and consider specializations across the four, respective forms. In the first specialization,

which we refer to as *Model One*, we assume that the thresholds are non-existent and, thus, ignore the presence of censoring. In the second formulation, *Model Two*, we enact conventional censoring as in (Tobin, 1957) and define $\tau_{ij} = 0$, for all $i = 1, 2, \dots, N$ and $j = 1, 2, T$. In the third formulation, which is *Model Three*, we enact a random censoring threshold, which is $\tau_{ij} = \tau$, common for all $i = 1, 2, \dots, N$ and $j = 1, 2, \dots, T$. And in the final formulation, *Model Four*, we enact, conditional censoring wherein $\tau_{ij} \equiv z_{ij} = x_{ij}'\beta_{m+1} + u_{ijm+1}$, for all $i = 1, 2, \dots, N$ and $j = 1, 2, \dots, T$, where x_{ij} denotes covariate information and β_{m+1} denotes the corresponding vector of response coefficients. In this latter setting three features of the sampling environment warrant comment. First, unlike *Models One, Two and Three*, the *Model Four* censoring threshold is agent-and-period specific; it allows for the possibility of cross-equation correlation; requires implementation of an additional equation containing the latent censoring thresholds; and extends important, preceding work, most notably the random-censoring-threshold (Nelson, 1977) and, more recently, the single-equation Bayesian implementation of the Tobit regression of (Chib, 1992). Finally, we emphasize that, whereas, *Models One, Two and Three* contain estimating systems of dimension $S \times M$, under *Model Four* the system assumes dimension $S \times (M + 1)$.

7. Prior probability density functions

In conceptualizing priors propriety is necessary in order to enable model comparison. We adopt the approach of employing ‘weak but proper’ forms and additional comment is relevant. Whereas conjugacy imposes restrictions on the cross-equation covariances of Γ and Ψ , conjugacy of itself does not impose stringent *a priori restrictions* on the distribution of Σ . Hence, we allow Σ to evolve *a priori* according to $f_{M \times M}^{iW}(\Sigma | S_0, v_0)$, $S_0 = I_M \times 10^2$, $v_0 = M + 2$, which is indeed ‘weak’ but ‘proper.’ Second, although the restriction that the conjugate prior covariance matrices conditioned by Σ revert to ‘restrictive’ ‘block-diagonal structures is well documented (Drèze and Richard, 1983) – see especially, their comments on page 541; an indication they attribute to (Rothenberg, 1963) – those criticisms are less relevant here, for three reasons. First, by view of the fact that we place a priori weight equally across *twelve models* the restrictive feature of any one form is mitigated. Second, the inclusion of hierarchical components permit sufficient variability across coefficient columns that they further mitigate these concerns. Third, conjugacy endows the posterior quantities with an attractive feature that considerably facilitates models comparison. Thus, the prior assumptions that we invoke for the relevant response matrices are $f_{K \times M}^{mN}(\Psi | \Psi_0, \Sigma, \text{Cov}\Psi_0)$, where $\Psi_0 = 0_{KM}$, $\text{Cov}\Psi_0 = I_K \times 10^2$ and $f_{K \times M}^{mN}(\Gamma | \Gamma_0, \Sigma, \text{Cov}\Gamma_0)$, $\Gamma_0 = 0_{1M}$, $\text{Cov}\Gamma_0 = I_1 \times 10^2$, which are also ‘weak’ but ‘proper.’

8. Parameter estimation strategy

Despite complications, the estimation retains the same basic simplicity inherent in all Normal-data models. Moreover, because all of the essential vector-valued results in (Lindley and Smith, 1972) extend intuitively and readily to the matrix form, the entire posterior analysis would be available in closed-form, if not for censoring. In the presence of censoring we adopt a Gibbs-sampling estimation strategy and the execution is routine. Here, we outline the basic estimation algorithms and detail amendments required to incorporate censoring. In terms of *Specification One*, *Model One* represents the basic matrix-Normal formulation wherein $f_{S \times M}^{mN}(Y | X\Psi, \Sigma, I_S)$, $f_{M \times M}^{iW}(\Sigma | S_0, v)$, and $f_{K \times M}^{mN}(\Psi | \Psi_0, \Sigma, \text{Cov}\Psi_0)$, comprise the joint distribution for the data and the parameters. It follows that the joint posterior for the unknown quantities is defined by the conjugate distributions $f_{M \times M}^{iW}(\Sigma | S_*, v_*)$ and $f_{K \times M}^{mN}(\Psi | \Psi_*, \Sigma, \text{Cov}\Psi_*)$ and that the fully conditional distributions underlying the Gibbs-sampling algorithm are $f_{M \times M}^{iW}(\Sigma | S_*, v_*)$ and $f_{K \times M}^{mN}(\Psi | \Psi_*, \Sigma, \text{Cov}\Psi_*)$, where, $S_* \equiv (\Psi - \Psi_0)' \text{Cov}\Psi_0^{-1} (\Psi - \Psi_0) + (Y - X\Psi)' I_S (Y - X\Psi)$, $v_* \equiv K + S$, $\text{Cov}\Psi_* \equiv (X' I_S X + \text{Cov}\Psi_0^{-1})^{-1}$ and $\Psi_* \equiv \text{Cov}\Psi_* (X' I_S X + \text{Cov}\Psi_0^{-1} \Psi_0)^{-1}$. On the first count, the definitions of S_* and v_* follow, straight-forwardly, from the definition of the inverted-Wishart distribution. On the second, count, the definitions $\text{Cov}\Psi_*$ and Ψ_* are available from well-known results (Zellner, 1971; Bauwens, 1983; Drèze and Richard, 1984) for the matrix-Normal model. Nevertheless, we explicate developments in order to introduce a procedure automating repeated derivations for the response coefficients. Each of the matrix-Normal coefficient matrices, for example, Θ , has a fully-conditional, data-dependent form derivable in terms of generic matrices A, B, C, D and E , namely, $f(\Theta) \propto \exp\{-.5 \text{trace}(A - B\Theta)' C (A - B\Theta) \Sigma^{-1}\} \times \exp\{-.5 \text{trace}(\Theta - D)' E (\Theta - D) \Sigma^{-1}\}$. Thus, completing the square in similar fashion to the vector-Normal (Zellner, 1971, p. 381; 19, pp. 5 and 6) we have

$$f^{mN}(\Theta) \propto \exp\{-.5 \text{trace}(\Theta - \Theta_*)' \text{Cov}\Theta_*^{-1} (\Theta - \Theta_*) \Sigma^{-1}\} \propto f^{mN}(\Theta | \Theta_*, \Sigma, \text{Cov}\Theta_*), \quad (\text{A.2})$$

where $\text{Cov}\Theta_* \equiv (B'CB + E)^{-1}$ and $\Theta_* \equiv (B'CB + E)^{-1} (B'CA + ED)$. Thus, one can confirm that $\text{Cov}\Psi_* \equiv (B'CB + E)^{-1}$ and $\Psi_* \equiv (B'CB + E)^{-1} (B'CA + ED)$ where $A \equiv Y, B \equiv X, C \equiv I_S, D \equiv \Psi_0$ and $E \equiv \text{Cov}\Psi_0^{-1}$, as conjectured, and that the formal, Gibbs-sampling algorithm required for implementing *Specification-One, Model-One* consists of consecutive draws from $f_{M \times M}^{iW}(\Sigma | S_*, v_*)$ and $f_{K \times M}^{mN}(\Psi | \Psi_*, \Sigma, \text{Cov}\Psi_*)$. In the *conventional censoring* extension of *Model One*, *Model Two* employs V in place of Y and contains the additional draw from $f_{1 \times 1}^{tN}(z_{ijk} | \mu_{ijk}, \sigma_{ijk}, (-\infty, 0])$ for $\{ijk\} \in c \equiv \{ijk | y_{ijk} = 0\}$, where c denotes the censor set for households $i = 1, 2, \dots, N$ in time periods $j = 1, 2, \dots, T$ and with respect

to output quantity k ; and from $f_{1 \times 1}^{tN}(z_{ijk} | \mu_{ijk}, \sigma_{ijk}, (-\infty, 0])$ denotes the (fully conditional) truncated-Normal, univariate distribution, with mean μ_{ijk} and standard deviation σ_{ijk} , defined on the interval $(-\infty, 0]$ for the scalar latent quantity z_{ijk} . Here the univariate-Normal draws are retrievable by direct application of the decompositions in (Zellner, 1971, pp. 381-382) and direct, one-to-one draws are available from the *probability-integral transform* as outlined, for example, in (Albert and Chib, 1993, p. 202). Under *Specification One, Model Three*, we require an additional draw for the random censoring threshold, which is a (scalar) uniform distribution $f_{1 \times 1}^U(\tau | \tau_{\min\#}, \tau_{\max\#})$, where (exploiting similarities between the posterior distributions for the ordered-probit bin boundaries in (Albert and Chib, 1993) and the fully conditional posterior distribution for the threshold parameter, τ) we deduce that $\tau_{\min\#} \equiv \max\{\max\{z_{ijk}, \{ijk\} \in c\}, \tau_{\min 0}\}$ and $\tau_{\max\#} \equiv \min\{\min\{y_{ijk}, \{ijk\} \notin c\}, \tau_{\max 0}\}$. Here, parameters $\tau_{\min 0}$ and $\tau_{\max 0}$ comprise the support for τ within the prior pdf $f_{1 \times 1}^U(\tau | \tau_{\min 0}, \tau_{\max 0})$ and $f_{1 \times 1}^{tN}(z_{ijk} | \mu_{ijk}, \sigma_{ijk}, (-\infty, \tau])$ denotes the (fully conditional) truncated-Normal distribution for the latent data. Finally, *Model Four* introduces the *conditional censoring threshold* such that from $f_{1 \times 1}^{tN}(z_{ijk} | \mu_{ijk}, \sigma_{ijk}, (-\infty, z_{ijM+1}])$ now denotes the relevant truncated-Normal distribution and the following details are emphasized. First, introduction of the threshold introduces one additional entire column of latent data into V . Accordingly, in the absence of restrictions on the parameter space, the regression model is unidentified. A convenient remedy (borrowed from probit analysis) is to restrict one of the diagonal terms in Σ to one and derive the fully conditional distributions for the remaining, non-restricted components of Σ . The idea for this modification stems from an investigation of ‘infinite regression’ (Dawid, 1998), with further modifications and refinements arising in the context of multinomial-probit estimation (Nobile, 2000; McCulloch and Rossi, 2000), with, finally, explicit derivations for the draw for Σ outlined in an appendix to an application on transport choice (Linardakis and Dellaportas, 2003). Finally, in addition to the covariance restriction, one must draw the latent data in the $M + 1^{\text{st}}$ column of V , which are made according to $f_{1 \times 1}^{tN}(z_{ijM+1} | \mu_{ijk}, \sigma_{ijk}, [z_{ijk}, +\infty))$, for $\{ijk\} \in c$, and $f_{1 \times 1}^{tN}(z_{ijM+1} | \mu_{ijk}, \sigma_{ijk}, (-\infty, z_{ijk}])$, for $\{ijk\} \notin c$. Turning to *Specification Two*, the basic model has component conditional distributions consisting of five forms, namely $f_{S \times M}^{mN}(Y | W\Xi, X\Psi, \Sigma, I_S)$, $f_{S \times M}^{mN}(\Psi | \Psi_0, \Sigma, \text{Cov}\Psi_0)$, $f_{S \times M}^{mN}(\Xi | H\Gamma, \Sigma, \text{Cov}\Xi_0)$, $f_{1 \times M}^{mN}(\Gamma | \Gamma_0, \Sigma, \text{Cov}\Gamma_0)$, and $f_{M \times M}^{iW}(\Sigma | S_0, v)$. The joint posterior is defined by the conjugate distributions $f_{M \times M}^{iW}(\Sigma | S_*, v_*)$, $f_{1 \times M}^{mN}(\Gamma | \Gamma_*, \Sigma, \text{Cov}\Gamma_*)$, $f_{N \times M}^{mN}(\Xi | \Xi_*, \Sigma, \text{Cov}\Xi_*)$ and $f_{K \times M}^{mN}(\Psi | \Psi_*, \Sigma, \text{Cov}\Psi_*)$ and the fully conditional distributions required for implementation are $f_{M \times M}^{iW}(\Sigma | S_{\#}, v_{\#})$, $f_{1 \times M}^{mN}(\Gamma | \Gamma_{\#}, \Sigma, \text{Cov}\Gamma_{\#})$, $f_{N \times M}^{mN}(\Xi | \Xi_{\#}, \Sigma, \text{Cov}\Xi_{\#})$, and $f_{K \times M}^{mN}(\Psi | \Psi_{\#}, \Sigma, \text{Cov}\Psi_{\#})$, where $S_{\#} \equiv (\Gamma - \Gamma_0)' \text{Cov}\Gamma_0^{-1} (\Gamma - \Gamma_0) + (\Xi - H\Gamma)' \text{Cov}\Xi_0^{-1} (\Xi - H\Gamma) + (\Psi - \Psi_0)' \text{Cov}\Psi_0^{-1} (\Psi - \Psi_0) + (Y - W\Xi - X\Psi)' I_S (Y - W\Xi - X\Psi)$, $v_{\#} \equiv 1 + N + K + S$, $\text{Cov}\Gamma_{\#} \equiv (B'CB + E)^{-1}$ and $\Gamma_{\#} \equiv (B'CB + E)^{-1} (B'CA + ED)$ where $A \equiv \Xi$, $B \equiv H$, $C \equiv \text{Cov}\Xi_0$, $D \equiv \Gamma_0$ and $E \equiv \text{Cov}\Gamma_0^{-1}$, $\text{Cov}\Xi_{\#} \equiv (B'CB + E)^{-1}$ and $\Xi_{\#} \equiv (B'CB + E)^{-1} (B'CA + ED)$ where $A \equiv Y - W\Xi$, $B \equiv X$, $C \equiv I_S$, $D \equiv \Psi_0$ and $E \equiv \text{Cov}\Psi_0^{-1}$. The conventionally-censored, randomly-censored and conditionally-censored regressions (respectively, *Model Two*, *Model Three* and *Model Four*) are executed in similar fashion to the manner described under *Specification One*. *Specification Three* consists of the distributional components $f_{S \times M}^{mN}(Y | P\Delta, \Sigma, I_S)$, $f_{KN \times M}^{mN}(\Delta | Q\Psi, \Sigma, \text{Cov}\Delta_0)$, $f_{K \times M}^{mN}(\Psi | \Psi_0, \Sigma, \text{Cov}\Psi_0)$, and $f_{M \times M}^{iW}(\Sigma | S_0, v_0)$; the joint posterior contains components $f_{M \times M}^{iW}(\Sigma | S_{\#}, v_{\#})$, $f_{K \times M}^{mN}(\Psi | \Psi_{\#}, \Sigma, \text{Cov}\Psi_{\#})$ and $f_{NK \times M}^{mN}(\Delta | \Delta_{\#}, \Sigma, \text{Cov}\Delta_{\#})$; and the corresponding fully conditional distributions are $f_{M \times M}^{iW}(\Sigma | S_{\#}, v_{\#})$, $f_{K \times M}^{mN}(\Psi | \Psi_{\#}, \Sigma, \text{Cov}\Psi_{\#})$ and $f_{NK \times M}^{mN}(\Delta | \Delta_{\#}, \Sigma, \text{Cov}\Delta_{\#})$, where $S_{\#} \equiv (\Psi - \Psi_0)' \text{Cov}\Psi_0^{-1} (\Psi - \Psi_0) + (\Delta - Q\Psi)' \text{Cov}\Delta_0^{-1} (\Delta - Q\Psi) + (Y - P\Delta)' I_S (Y - P\Delta)$, $v_{\#} \equiv K + NK + S$, $\text{Cov}\Delta_{\#} \equiv (B'CB + E)^{-1}$ and $\Delta_{\#} \equiv (B'CB + E)^{-1} (B'CA + ED)$ where $A \equiv Y$, $B \equiv P$, $C \equiv \text{Cov}\Delta_0$, $D \equiv Q\Psi$ and $E \equiv \text{Cov}\Delta_0^{-1}$, $\text{Cov}\Psi_{\#} \equiv (B'CB + E)^{-1}$ and $\Psi_{\#} \equiv (B'CB + E)^{-1} (B'CA + ED)$ where $A \equiv \Delta$, $B \equiv Q$, $C \equiv \text{Cov}\Delta_0^{-1}$, $D \equiv \Psi_0$ and $E \equiv \text{Cov}\Psi_0^{-1}$. Finally, as above, the *conventionally-censored*, *randomly-censored* and *conditionally censored* regressions (respectively, *Model Two*, *Model Three* and *Model Four*) are executed in similar fashion to the manner described in *Specification One*.

9. Models comparison strategy

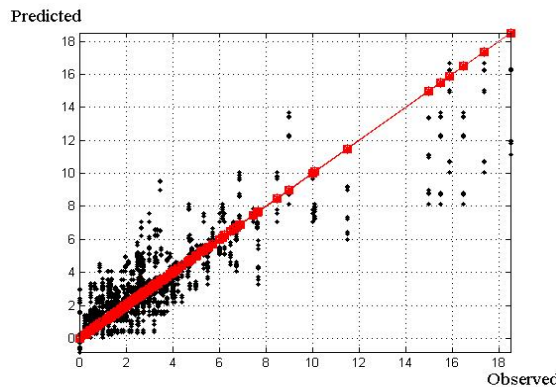
The essential input in models comparison is the *marginal likelihood*, estimates of which demand numerical methods due to the presence of censoring. The technique preferred by us is a generalization of the Gibbs-sampling technique proposed by (Rothenberg, 1963) who shows that a robust estimate of the marginal likelihood is available via simple extensions of the basic Gibbs algorithm used in parameter estimation. However, we find it convenient and considerably more efficient to execute models comparisons using the marginal distributions for the data as opposed to the fully conditional data distributions one typically employs. Because, in the generalization to the matrix-Normal, the vector-Normal contributions in Drèze and Richard (1984, equation (4), page 5) go through in the same way, the marginal distributions are easily obtained. Under *Specification One* the marginal distribution for Y remains

$f_{S \times M}^{mN}(Y | X\psi, \Sigma, I_S)$; under *Specification Two* the marginal distribution is $f_{S \times M}^{mN}(Y | WH\Gamma + X\psi, \Sigma, I_S)$; and under *Specification Three* the marginal distribution for the data is $f_{S \times M}^{mN}(Y | PQ, \Sigma, I_S)$. Thus, *Specification One* is enacted according to the Gibbs-sampling algorithm already discussed. Under *Specification Two* we draw, respectively, from $f_{M \times M}^{iW}(\Sigma | S_{\#}, v_{\#})$, $f_{1 \times M}^{mN}(\Gamma | \Gamma_*, \Sigma, Cov\Gamma_*)$, and $f_{K \times M}^{mN}(\psi | \psi_*, \Sigma, Cov\psi_*)$, where $S_{\#} \equiv (\Gamma - \Gamma_0)' Cov\Gamma_0^{-1} (\Gamma - \Gamma_0) + (\Psi - \Psi_0)' Cov\Psi_0^{-1} (\Psi - \Psi_0) + (Y - WVT - X\psi)' (I_S + WCov\Xi_0 W')^{-1} (Y - WVT - X\psi)$, $v_{\#} \equiv 1 + K + S$, $Cov\Gamma_{\#} \equiv (B'CB + E)^{-1}$ and $\Gamma_{\#} \equiv (B'CB + E)^{-1} (B'CA + ED)$ where $A \equiv Y$, $B \equiv WH$, $C \equiv (I_S + WCov\Xi_0 W')^{-1}$, $D \equiv \Gamma_0$, $E \equiv Cov\Gamma_0^{-1}$; and $Cov\Psi_{\#} \equiv (B'CB + E)^{-1}$ and $\Psi_{\#} \equiv (B'CB + E)^{-1} (B'CA + ED)$ where $A \equiv Y - WH\Gamma$, $B \equiv X$, $C \equiv (I_S + WCov\Xi_0 W')^{-1}$, $D \equiv \Psi_0$ and $E \equiv Cov\Psi_0^{-1}$. Finally, under *Specification Three* the Gibbs-sampling algorithm consists of draws from $f_{M \times M}^{iW}(\Sigma | S_{\#}, v_{\#})$, and $f_{K \times M}^{mN}(\psi | \psi_*, \Sigma, Cov\psi_*)$, where $S_{\#} \equiv (\Psi - \Psi_0)' Cov\Psi_0^{-1} (\Psi - \Psi_0) + (Y - PQ\Psi)' (I_S + PCov\Delta_0 P')^{-1} (Y - PQ\Psi)$, $v_{\#} \equiv K + S$, $Cov\Psi_{\#} \equiv (B'CB + E)^{-1}$ and $\Psi_{\#} \equiv (B'CB + E)^{-1} (B'CA + ED)$ where $A \equiv Y$, $B \equiv PQ$, $C \equiv (I_S + PCov\Delta_0 P')^{-1}$, $D \equiv \Psi_0$ and $E \equiv Cov\Psi_0^{-1}$.

10. Marginal likelihood computation

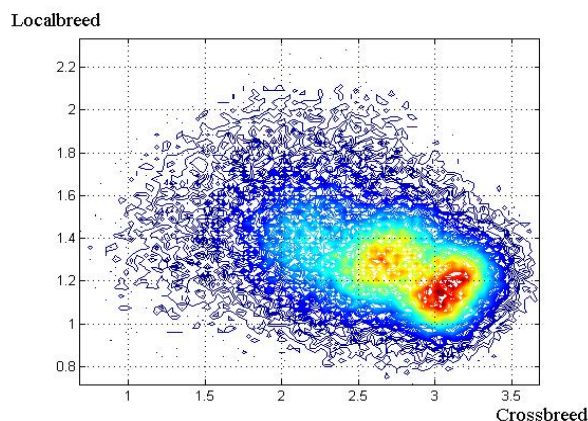
Defining the non-observed data, collectively, as Z ; defining the parameters, collectively, as Θ ; and defining the observed data, collectively, as Y ; the joint distribution for all quantities, $f(\Theta, Z, Y)$, can be written, alternatively, as $f(\Theta, Z, Y) = f(\Theta, Z|Y) \times f(Y)$, $f(Y, \Theta, Z) = f(Y, \Theta|Z) \times f(Z)$, or $f(Z, Y, \Theta) = f(Z, Y|\Theta) \times f(\Theta)$. Writing the joint density in this fashion is important because it emphasizes the important feature of the estimation yielding alternative strategies for computing the *marginal likelihood for the data*, $f(Y)$, which is the essential input into models comparisons. More specifically, using the facts that $f(\Theta, Z) = f(Z|\Theta) \times f(\Theta)$ and $f(\Theta, Z|Y) = f(Z|\Theta, Y) \times f(\Theta|Y)$, we can write $f(Y) = f(Y|\Theta, Z) \times f(Z|\Theta) \times f(\Theta) \div f(Z|\Theta, Y) \div f(\Theta|Y)$ and, more usefully, $f(Y) = f(Y, Z|\Theta) \times f(\Theta) \div f(\Theta|Z, Y) \div f(Z|Y)$. The strategy that we adopt here is the familiar one of integrating over the latent data so that, on the computationally convenient logarithmic scale, our estimating equation, for the marginal likelihood, is $\ln f(Y) = \ln f(Y|\Theta^*) + \ln f(\Theta^*) - \ln f(\Theta^*|Y)$, which is the multiple-equation analog of text equation (5). In our setting, each of the components on the right-hand side, except the last component, is available in closed form. The quantity, $f(Y|\Theta^*)$, is the matrix-Normal density, evaluated at the point $\Theta = \Theta^*$, integrated over the latent data, Z ; and the quantity $f(\Theta^*|Y)$ is simply the conjugate posterior distribution for the parameters, evaluated at the point $\Theta = \Theta^*$, and once, again, integrated over the latent data, Z . By noting that $f(Y, Z|\Theta) = f(Y|Z, \Theta) \times f(Z|\Theta)$, and that $f(Z, \Theta|Y) = f(\Theta|Z, Y) \times f(Z|Y)$, we note that appropriate *Monte Carlo* estimates of $f(Y|\Theta^*)$ and $f(\Theta^*|Y)$ are, respectively, $\hat{f}(Y|\Theta^*) \equiv \frac{1}{G} \sum_{g=1}^G f(Y|Z^{(g)}, \Theta^*)$ where $Z^{(1)}, Z^{(2)}, \dots, Z^{(G)}$ denote draws marginally from the density $f(Z|\Theta)$ and $\hat{f}(\Theta^*|Y) \equiv \frac{1}{G} \sum_{g=1}^G f(\Theta^*|Z^{(g)}, Y)$ where $Z^{(1)}, Z^{(2)}, \dots, Z^{(G)}$ denote draws from the density $f(Z|Y)$. Thus, a robust estimate of the marginal likelihood in the presence of latent data is available by extending procedures outlined in Chib (1998). Finally, all of the algorithmic developments are implemented in MATLAB© version 7.10 with the Statistics Toolbox© installed on a modest hardware platform and the entire computer code and the data are available from the authors upon request.

Appendix B



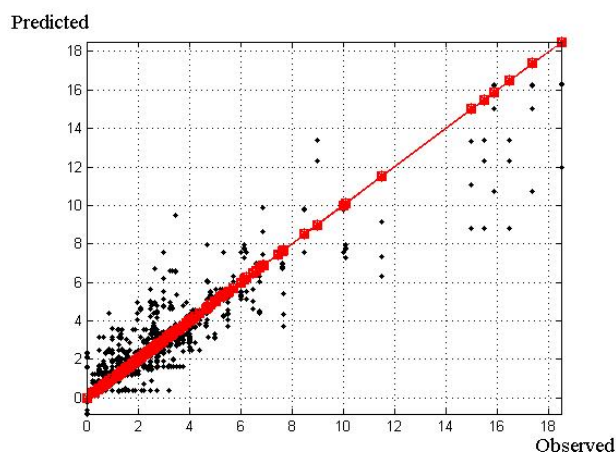
Notes: The grey dotted entry depicts the line of perfect fit. The grey squares along the grey dotted line depict the observations on milk output from each of the households at each period in the panel. The black dots depict predictions on milk output obtained from the three canonical statistical forms.

Fig. 1. Single-equation posterior predictions: milk output (liters per household per day)



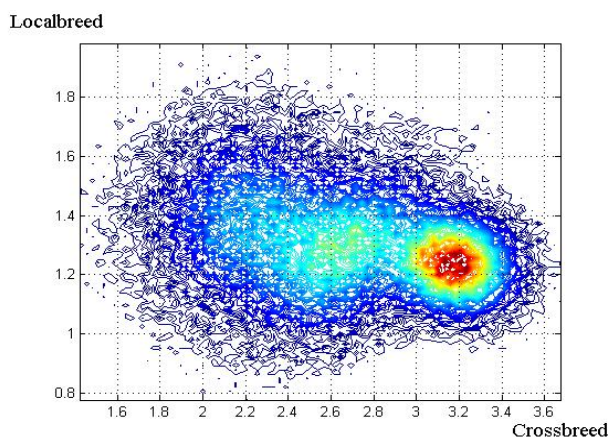
Notes: Marginal products attributable to crossbreed animals are depicted horizontally and marginal products attributable to indigenous-breed animals are depicted vertically. Gradual intensity of grey spectrum depicts gradual intensity of mass. Posterior means coordinates estimate (2.62, 1.30); posterior medians coordinates estimates (2.69, 1.28); ninety-five percent highest posterior density coordinates estimates ([1.54, 3.33], [0.98, 1.77]).

Fig. 2. Single-equation, marginal-product contours: milk output increments (liters of milk per household per day) per unit increment in livestock (crossbreed and indigenous-breed cows)



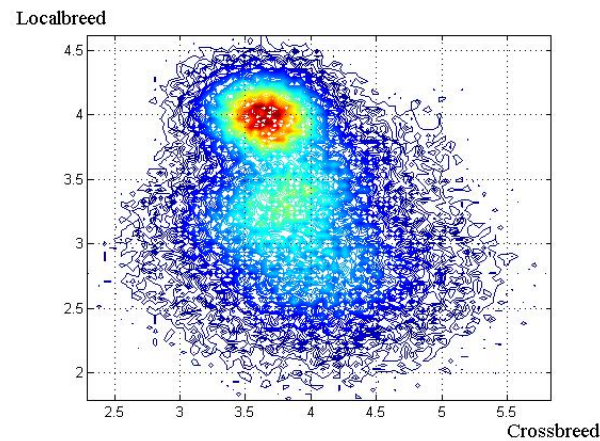
Notes: The grey dotted entry depicts the line of perfect fit. The grey squares along the grey dotted line depict the observations on milk output from each of the households at each period in the panel. The black dots depict predictions on milk output obtained from the three extended statistical forms.

Fig. 3. Multiple-equation posterior predictions: milk output (liters per household per day)



Notes: Marginal products attributable to crossbreed animals are depicted horizontally and marginal products attributable to indigenous-breed animals are depicted vertically. Gradual intensity of grey spectrum depicts gradual intensity of mass. Posterior means coordinates estimate (2.68, 1.30); posterior medians coordinates estimates (2.68, 1.29); ninety-five percent highest posterior density coordinates estimates ([1.86, 3.38], [1.04, 1.64]).

Fig. 4. Multiple-equation, marginal-product contours: milk output increments (liters of milk per household per day) per unit increment in livestock (crossbreed and indigenous-breed cows)



Notes: Values of economic loss attributable to crossbreed animals are depicted on the horizontal axis and values of economic loss attributable to indigenous-breed animals are depicted vertically. Gradual intensity of grey spectrum depicts gradual intensity of mass. Posterior means coordinates estimate (3.84, 3.37); posterior medians coordinates estimates (3.79, 3.37); ninety-five percent highest posterior density coordinates estimates ([3.06, 4.84], [4.84, 4.25]).

Fig. 5. Multiple-equation, catastrophic-livestock loss contours: present economic values ($\text{\$US} \times 10^{10}$) per unit increment in 1997 Ethiopian national herds (all crossbreed cows and all indigenous-breed cows)