

EVALUATING APPROXIMATIONS OF BAYESIAN SOLUTIONS: THE CASE OF FISHER TEST

by

Michel MOUCHART* and Eliana SCHEIHING†

July 1993

Abstract

In this paper, we evaluate from a Bayesian point of view how much information is lost when the sampling process for the 2×2 contingency table is specified conditionally on the two margins as in the exact test of Fisher. We first analyse the general problem of admissible conditioning and next consider the evaluation of the loss of information when a non-admissible conditioning is used for an approximation of the exact posterior distribution. Turning to the Fisher test, three different sampling models are considered and three comparisons are designed between the exact and the approximate posterior distributions. The numerical results obtained through simulation indicate that for a specific range of parameters the loss of information increases with the sample size and decreases with the precision of the *a priori* distribution.

Keywords: Approximate Bayesian Solutions, Admissible conditioning, contingency tables, Fisher Test.

AMS-MOS : Primary : 62F15, Secondary: 62A15, 65C05

Acknowledgement : Financial help from the Belgian Fund for Scientific Research (FNRS) partially supported the work of the first author and is gratefully acknowledged. This version has benefited from many useful discussions with J.M. Rolin.

*Institut de Statistique, Université Catholique de Louvain, Belgium

†Institut de Statistique, Université Catholique de Louvain, Belgium and Instituto de Informática, Universidad Austral de Chile

1 Introduction

In the analysis of 2 x 2 contingency tables, the exact test of Fisher is based on the sampling distribution of the table conditional on the two margins. When learning from the data concerns only a test of independence, Fisher test has been shown to enjoy reasonable sampling properties (such as UMPU test, see Kendall and Stuart(1967), Vol.2 p. 550-552). However when inference is considered more generally, it is also known that the two margins do contain information on the parameter (λ) characterizing the independence (see, for instance, Haber(1989)).

The object of this paper is to evaluate from a Bayesian point of view the admissibility of the conditioning on the two margins. More specifically we want to evaluate how much information is lost when the sampling process is specified only conditionally on the two margins rather than jointly with the two margins. From a Bayesian point of view, this loss of information may be evaluated either through a divergence (or a distance) between the posterior distributions corresponding respectively to a complete and to an incomplete specification of the sampling process, through a comparison of the distributions of the two Bayesian estimators corresponding respectively to a complete and to an incomplete specification of the sampling process, or, finally, through a comparison of the Bayesian risks associated to optimal decision rules based respectively on the complete and on the incomplete specification of the sampling process. Note that these three comparisons correspond to three different levels of the specificity of the underlying decision problem.

It so happens that these comparisons require, as a starting point, a precise specification of the complete sampling process and of the conditions that allow one to separate the inference about the various parameters characterizing the data generating process. Indeed, one of the conclusions of this analysis is that insisting either on the multinomial sampling of the table or on independent binomial samplings implies a loss of information on λ and insisting on no loss of information on λ implies a sampling different from a multinomial one or an independent binomial one, which in turn implies that λ , the parameter characterizing the row/column independence, is likely to be not anymore a parameter sufficient for the process conditional on the margins.

In next section, we analyse the problem of admissible conditioning in general and the evaluation of the loss of information when a not admissible conditioning is specified, while section 3 considers the problem in the particular case of the 2x2 contingency table, with special emphasis on the role of the model specification. Finally section 4 presents numerical results obtained through simulation,

that illustrate the loss of information involved by using a non-admissible conditioning on the two margins.

2 Admissible Conditioning in Bayesian Methods

2.1 Exact Admissibility

In this section we shortly summarize the Bayesian approach to the admissibility of the reductions by conditioning. This is the question whether it is *admissible*, in the sense of losing no information, to specify the sampling process only partially by treating part of the data "as if" it were not random. This summary is short and concentrates on the issues relevant for the problem of conditioning on the two margins. The topic of conditioning is treated most systematically in chapter 3 of Florens, Mouchart and Rolin (1990), see also, Florens and Mouchart (1977,1985) and Dawid (1980).

As usual, we consider that a Bayesian model is simply a (unique) probability measure jointly on the parameters and the observations. If we write x for the observations and θ for the parameters, the Bayesian model may accordingly be represented in terms of densities (with respect to suitable measures) as follows:

$$q(x, \theta) = m(\theta) p(x | \theta) = p(x) m(\theta | x) \quad (1)$$

where as usually

$$m(\theta | x) \propto m(\theta) p(x | \theta) \quad (2)$$

represents the posterior distribution of θ , i.e. the Bayesian learning rule from the data.

Let us now decompose the observations into two components y and z , i.e. $x = (y, z)$ and decompose accordingly the sampling distributions:

$$p(x | \theta) = p(z | \theta) p(y | z, \theta) \quad (3)$$

and therefore

$$m(\theta | x) \propto m(\theta) p(z | \theta) p(y | z, \theta) \quad (4)$$

We are interested in the following question: how far can we "forget" the term $p(z | \theta)$ for the inference on θ ? Clearly if z is *ancillary* (i.e. $z \perp\!\!\!\perp \theta$, so that $p(z | \theta) = p(z)$ and $m(\theta | z) = m(\theta)$) we have:

$$m(\theta | x) \propto m(\theta) p(y | z, \theta) \quad (5)$$

so that the posterior distribution of θ remains unaffected. We say that this reduction of the sampling process by conditioning is *admissible*.

This means that when z is ancillary, inference may proceed "as if" z were not random in the sense that the sampling process needs be specified up to the distribution of the data *conditionally* on z only. Note that Sampling theory, at least under the likelihood principle, and Bayesian theory here coincide on the essential.

Suppose now that we also decompose θ into two components α and β , i.e. $\theta = (\alpha, \beta)$, where α is a nuisance parameter and β is a parameter of interest (i.e. the underlying loss function depends on β only). Therefore we only have interest in the posterior distribution of β which may be evaluated indifferently as:

$$m(\beta | x) = \int m(\alpha, \beta | x) d\alpha \quad (6)$$

$$\propto m(\beta) p(x | \beta) \quad (7)$$

where

$$p(x | \beta) = \int p(x | \alpha, \beta) m(\alpha | \beta) d\alpha \quad (8)$$

may be called a *marginal data density*. Considering again the decomposition of the observations into $x = (y, z)$, we also have:

$$m(\beta | x) \propto m(\beta) p(z | \beta) p(y | z, \beta) \quad (9)$$

and raise a question similar to the previous one: how far can we forget the term $p(z | \beta)$ for the inference on β . Formally, the answer is also similar to the previous answer, namely that if $z \perp\!\!\!\perp \beta$, i.e. $p(z | \beta) = p(z)$, equivalently $m(\beta | z) = m(\beta)$, we also have:

$$m(\beta | x) \propto m(\beta) p(y | z, \beta) \quad (10)$$

and say that z and β are *mutually ancillary*. Clearly, if z is ancillary, z is also mutually ancillary with any function of θ , but the converse is evidently false. Note however that the use of formula (10) implies integrating the conditional sampling distribution of $(y | z)$ with respect to the conditional posterior distribution of α given z and β , more specifically:

$$p(y | z, \beta) = \int p(y | z, \alpha, \beta) m(\alpha | \beta, z) d\alpha \quad (11)$$

unless β is a sufficient parameter of the conditional sampling distribution of $(y | z)$ (i.e. $y \perp\!\!\!\perp \alpha | \beta, z$), in which case $p(y | z, \beta) = p(y | z, \alpha, \beta)$ for any prior distribution and therefore $m(\alpha | \beta, z)$ (which would also be equal to $m(\alpha | \beta, y, z)$), need not be specified for performing the integration underlying (11). When both $z \perp\!\!\!\perp \beta$ and $y \perp\!\!\!\perp \alpha | \beta, z$ are verified, β and z are said to be *mutually exogenous*. When β is a sufficient parameter of the sampling distribution conditional on z , a rather natural way of obtaining the mutual exogeneity of β and z (i.e. thus a stronger condition than the mutual ancillarity condition), is through the structure of a *Bayesian cut* which means that α is furthermore a sufficient parameter of the sampling marginal process of z (i.e. $z \perp\!\!\!\perp \beta | \alpha$) and is also, a priori independent of β ; this is so simply because $z \perp\!\!\!\perp \beta | \alpha$ and $\alpha \perp\!\!\!\perp \beta$ trivially imply $z \perp\!\!\!\perp \beta$. The advantage of that structure is that the first two conditions ($y \perp\!\!\!\perp \alpha | z, \beta$ and $z \perp\!\!\!\perp \beta | \alpha$) only depend on the sampling process and the last condition ($\alpha \perp\!\!\!\perp \beta$) only depends on the prior specification so that the structure of a Bayesian cut is robust with respect to any modification of the prior specification, provided it keeps the prior independence of α and β . Note however that these two requirements (of sufficiency of α and β for the marginal and the conditional sampling processes) do relate the structure of the loss function (depending only on β) and of the sampling process, a not always realistic feature.

2.2 Approximate Admissibility

Suppose now that z and β are not mutually ancillary but that one decides, for some (good or bad) reason to "forget" $p(z | \beta)$ in equation (9), i.e. using (10) as an "approximation" of (9). Thus let us write $m_A(\beta | x)$, as an approximation to $m(\beta | x)$, defined as:

$$m_A(\beta | x) \propto m(\beta) p(y | z, \beta) \quad (12)$$

without having $z \perp\!\!\!\perp \beta$, (in next section, β would also be a sufficient parameter of the sampling conditional process -i.e. $y \perp\!\!\!\perp \alpha | z, \beta$ - but this is not essential for what follows).

A natural question is now "how much information do we lose by overlooking the term $p(z | \beta)$?"

We will consider two ways of evaluating this loss of information. The first one, that is more inference-oriented, is based on the evaluation of the expected divergence (or distance) between the posterior distribution of parameter β , $m(\beta | x)$ and its approximation, $m_A(\beta | x)$, more specifically one may mea-

sure a loss of information L as follows:

$$L = E[D(m_A(\beta | x), m(\beta | x))] = \int D(m_A(\beta | x), m(\beta | x)) p(x) dx \quad (13)$$

where $D(\cdot, \cdot)$ is a divergence (or a distance) between probability measures. Note that $D(m_A(\beta | x), m(\beta | x))$ is a statistic, i.e. a function of the data only and the expectation is taken with respect to the predictive distribution of the exact model. For the sake of interpretation it may be more convenient to use a divergence (or a distance) bounded in the interval $[0,1]$. The second way, that is more decision-oriented, compares the risk functions corresponding to the optimal decision rules under the exact posterior distribution of β and under the approximate posterior distribution of β . Under the loss function $l(\beta, a)$, the optimal decision rule under the exact posterior distribution of β is:

$$a_E^*(x) = \arg \inf_a E_{\beta|x} l(\beta, a) \quad (14)$$

and the corresponding risk function is:

$$\rho(a_E^*) = E_{x,\beta} l(\beta, a_E^*(x)) \quad (15)$$

Similarly the optimal decision rule under the approximate posterior distribution of β is:

$$a_A^*(x) = \arg \inf_a E_{\beta|x}^A l(\beta, a) \quad (16)$$

where

$$E_{\beta|x}^A l(\beta, a) = \int l(\beta, a) m_A(\beta | x) d\beta \quad (17)$$

and the corresponding risk function is:

$$\rho(a_A^*) = E_{x,\beta} l(\beta, a_A^*(x)) \quad (18)$$

Clearly

$$0 \leq \rho(a_E^*) \leq \rho(a_A^*) \quad (19)$$

Therefore we propose as a measure of the loss of information the following comparison between the two risk functions, which may be read as the "Relative Efficiency of the approximation m_A ":

$$RelEff(m_A) = \frac{\rho(a_E^*)}{\rho(a_A^*)} \in [0, 1] \quad (20)$$

If this ratio takes the value $\frac{1}{k}$ this means that the risk of a_A^* is k times greater than the risk of a_E^* , so this expression is a measure of the relative increase of the risk due to the approximation on the posterior distribution.

As a particular case, consider the square loss function:

$$l(\beta, a) = (a - \beta)^2 \quad (21)$$

then

$$a_E^*(x) = E(\beta | x) \quad (22)$$

$$a_A^*(x) = E^A(\beta | x) \quad (23)$$

and the corresponding risk functions are:

$$\rho(a_E^*) = E_x V(\beta | x) \quad (24)$$

$$\rho(a_A^*) = E_{x, \beta} [E^A(\beta | x) - \beta]^2 \quad (25)$$

$$= E_x [V(\beta | x) + (E(\beta | x) - E^A(\beta | x))^2] \quad (26)$$

Therefore:

$$RelEff(m_A) = \left(1 + \frac{E_x [(E(\beta | x) - E^A(\beta | x))^2]}{E_x V(\beta | x)} \right)^{-1} \quad (27)$$

which shows that under a quadratic loss function, an approximation will be evaluated according to the ratio between the (predictively) expected exact posterior variance and the (predictively) expected square bias of the approximate posterior expectation. An evaluation finer than $RelEff(m_A)$ would be an analysis of the predictive distribution of the statistics $D(m_A(\beta | x), m(\beta | x))$. A possible interest of such an analysis is to provide a natural background for a bayesian testing of the mutual exogeneity (see e.g. Florens and Mouchart [1989], [1993]).

3 Conditioning on the two margins

3.1 Introduction

Let us now consider that X , the data, are in the form of a 2x2 contingency table, viz. $X = (X_{11} \ X_{12} \ X_{21} \ X_{22})$. We shall make use of the dot notation for the summation, for instance $X_{1.} = X_{11} + X_{12}$, or $X_{.} = X_{11} + X_{12} + X_{21} + X_{22}$. We want to consider how far it is admissible to realize the inference conditionally on the margins, i.e., in the notation of previous section, to assume the exogeneity of $Z = (X_{1.}, X_{.1})$. This analysis requires some care when specifying the sampling

process and the parameter of interest. We shall deal successively with these problems.

It is usual, in the analysis of the contingency tables, to exogeneize the grand total $X_{..}$ on the ground of a non-informative stopping rule for the sample size. Formally, this may be obtained as follows. Let α_0 and θ be sufficient parameters for the sampling distribution of $X_{..}$ and of $(X | X_{..})$ respectively, i.e.:

$$p(X | \theta, \alpha_0) = p(X_{..} | \alpha_0) p(X | X_{..}, \theta) \quad (28)$$

or equivalently

$$X_{..} \perp\!\!\!\perp \theta | \alpha_0 \quad (29)$$

and

$$X \perp\!\!\!\perp \alpha_0 | X_{..}, \theta \quad (30)$$

The mutual exogeneity of θ and $X_{..}$, i.e. $X_{..} \perp\!\!\!\perp \theta$ along with (30) is usually obtained by an hypothesis of a Bayesian cut, i.e. the conditions (29),(30) and the prior independence between θ and α_0 , i.e.

$$\theta \perp\!\!\!\perp \alpha_0 \quad (31)$$

This implies that once the parameter of interest is a function of θ , the inference may proceed conditionally on $X_{..}$, a situation we systematically assume from now on. In other words, in what follows, we never specify the way the sample size is generated, we only assume that any unknown parameter embodied in this process is a priori independent of the parameter characterizing the process generating the table conditionally on the sample size. We shall see later that this assumption might be less innocent than the first appearance might suggest.

In order to evaluate the exogeneity of the margins $(X_{1.}, X_{.1})$ in the process conditional on $X_{..}$ one has to be more specific on the sampling process and on the parameter of interest. In what follows, the parameter of interest, denoted $\lambda = l(\theta)$, is the parameter characterizing the statistical association between the row and column criteria, a particular value of which would imply independence. We now consider successively several sampling alternatives conditional on $X_{..}$.

3.2 Multinomial sampling

A multinomial sampling distribution may be obtained as the sampling distribution of the sufficient statistic under a generalized Bernoulli sampling with

non-informative stopping. More specifically let

$$W_l = (W_{11,l}, W_{12,l}, W_{21,l}, W_{22,l}) \in S_4 = \{w \in \{0, 1\}^4 \mid \sum_j w_j = 1\} \quad (32)$$

be distributed as independent Bernoulli of order 4, $W_l \sim I.B_4(\theta)$, i.e.

$$p(W_l = w \mid \theta) = \theta_{11}^{w_{11}} \theta_{12}^{w_{12}} \theta_{21}^{w_{21}} \theta_{22}^{w_{22}} \mathbb{1}_{\{w \in S_4\}} \quad (33)$$

with $\theta \in \mathcal{S}_4 = \{\theta \in [0, 1]^4 \mid \sum_j \theta_j = 1\}$. Clearly under independent sampling with non-informative stopping, the statistic:

$$X = (X_{ij})_{1 \leq i, j \leq 2} \quad X_{ij} = W_{ij, \cdot} = \sum_{1 \leq l \leq n} W_{ij,l} \quad (34)$$

where $n = X_{..}$ is the sample size, is a sufficient statistic (conditionally on $X_{..}$), the distribution of which is:

$$p(X = x \mid X_{..} = n, \theta) = \frac{n!}{\prod_{1 \leq i, j \leq 2} x_{ij}!} \prod_{1 \leq i, j \leq 2} \theta_{ij}^{x_{ij}} \quad (35)$$

The parameter of interest is defined as the usual cross product ratio:

$$\lambda = \frac{\theta_{11}\theta_{22}}{\theta_{12}\theta_{21}} \quad , \quad (36)$$

the value 1 of which characterizes the row-column independence.

In order to evaluate the admissibility of the conditioning on $Z = (X_{1.}, X_{.1})$ for the inference on λ , we notice that $(X_{11}, Z, X_{..})$ characterizes the table. The sampling distribution conditional on Z and $X_{..}$ boils therefore down to the sampling distribution of $(X_{11} \mid Z, X_{..})$, a non central hypergeometric distribution of parameter λ (see e.g. Agresti [1991], Cornfield [1956], Fisher [1935a]):

$$\begin{aligned} p(X_{11} = x_{11} \mid x_{1.}, x_{.1}, x_{..}, \theta) &= \frac{\lambda^{x_{11}}}{\prod_{i,j} x_{ij}!} \left(\sum_{x_{11} \in V} \frac{\lambda^{x_{11}}}{\prod_{i,j} x_{ij}!} \right)^{-1} \\ &= \lambda^{x_{11}} \binom{x_{1.}}{x_{11}} \binom{x_{.1} - x_{1.}}{x_{.1} - x_{11}} \left[\sum_{u \in V} \lambda^u \binom{x_{1.}}{u} \binom{x_{.1} - x_{1.}}{x_{.1} - u} \right]^{-1} \end{aligned} \quad (37)$$

where

$$V = \{x \in \mathbb{N} \mid \max\{0, x_{1.} + x_{.1} - x_{..}\} \leq x \leq \min\{x_{1.}, x_{.1}\}\} \quad (38)$$

Note that in the equation (37), each x_{ij} is a function of $(x_{11}, x_{1.}, x_{.1}, x_{..})$; therefore this formula *can not* be simplified by $\prod_{i,j} x_{ij}!$. Thus λ is a sufficient parameter for the distribution of the table conditional on the grand total $X_{..}$ and the margins $Z = (X_{1.}, X_{.1})$:

$$X \perp\!\!\!\perp \theta \mid \lambda, Z, X_{..} \quad (39)$$

However, conditionally on $X_{..}$, a Bayesian cut between the parameter of interest, λ , and the margins $Z = (X_{1.}, X_{.1})$ is impossible because the sampling distribution of the margins identifies *all* the parameters θ , indeed:

$$p(X_{1.} = x_{1.}, X_{.1} = x_{.1} \mid X_{..}, \theta) = X_{..}! \theta_{12}^{x_{1.}} \theta_{21}^{x_{.1}} \theta_{22}^{X_{..} - x_{1.} - x_{.1}} \sum_{x_{11} \in V} \frac{\lambda^{x_{11}}}{\prod_{i,j} x_{ij}!} \quad (40)$$

Therefore θ is a minimal sufficient parameter of the process generating $(Z \mid X_{..})$ and no prior distribution can satisfy the condition of a (conditional) Bayesian cut, namely $\lambda \perp\!\!\!\perp \theta \mid X_{..}$, because λ is a function of θ .

As the Bayesian cut is not possible, one may ask whether some prior distribution would allow Z and λ to be mutually exogenous (conditionally on $X_{..}$), this is condition (39) along with the condition of (conditional) mutual ancillarity:

$$Z \perp\!\!\!\perp \lambda \mid X_{..} \quad (41)$$

Note that (41) implies

$$Z \perp\!\!\!\perp \lambda \quad (42)$$

because of the first cut (equations (29), (30) et (31)).

Condition (41) means that

$$p(Z \mid \lambda, X_{..}) = \int p(Z \mid \theta, X_{..}) m(\theta_1 \mid \lambda) d\theta_1 \quad (43)$$

where $p(Z \mid \theta, X_{..})$ is given in (40) and θ_1 is defined for instance as $\theta_1 = \left(\theta_{22}, \frac{\theta_{21}}{\theta_{21} + \theta_{22}}\right)$, does not depend on λ .

As an example, let us assume that θ is distributed a priori as a Dirichlet distribution with parameter $a = (a_{ij})$ i.e.

$$m(\theta) = \frac{\Gamma(a_{..})}{\prod_{i,j} \Gamma(a_{ij})} \prod_{i,j} \theta^{a_{ij}-1} \mathbf{1}_{\{\theta \in \mathcal{S}_4\}} \quad (44)$$

It is shown, in Proposition 1 of the appendix, that in such a case:

$$p(Z = (x_{1.}, x_{.1}) \mid \lambda, X_{..}) = c X_{..}! \frac{E[(1 - Y + \lambda Y)^{-(a_{1.} + x_{1.})}]}{E[(1 - W + \lambda W)^{-a_{1.}}]} \sum_{x_{11} \in V} \frac{\lambda^{x_{11}}}{\prod_{i,j} x_{ij}!} \quad (45)$$

where

$$\begin{aligned} c &= \frac{\Gamma^2(a_{..})}{\Gamma^2(a_{..} + x_{..})} \frac{\Gamma(a_{1.} + x_{1.})\Gamma(a_{2.} + x_{2.})\Gamma(a_{.1} + x_{.1})\Gamma(a_{.2} + x_{.2})}{\Gamma(a_{1.})\Gamma(a_{2.})\Gamma(a_{.1})\Gamma(a_{.2})} \\ Y &\sim \beta(a_{.1} + x_{.1}, a_{.2} + x_{.2}) \\ W &\sim \beta(a_{.1}, a_{.2}) \end{aligned}$$

Therefore under a Dirichlet prior specification for θ , it may be seen that no specification of a is likely to make $p(Z \mid \lambda, X_{..})$ independent of λ . The computations, in the appendix, suggest that $p(Z \mid \lambda, X_{..})$ could be independent of λ only under very particular prior (non Dirichlet) specification, a theme not to be pursued in this paper leaving open the question whether there exists a prior specification $m(\theta_1 \mid \lambda)$ implying (41). The conclusion is therefore that under multinomial sampling, conditioning on the margins is generally not an admissible reduction.

3.3 Independent Binomial Sampling

In this case, the 2x2 contingency table is generated by two independent binomial samplings corresponding, say, to each column. More specifically, we again assume the first cut (29), (30), (31) and consider accordingly the grand total $X_{..}$ as exogenous. Let us write again θ for the parameter sufficient for the process generating the 2x2 contingency table, conditionally on $X_{..}$ and decompose now θ as follows:

$$\theta = (\alpha, \beta) = (\alpha, \beta_1, \beta_2) \quad (46)$$

where α is the parameter sufficient for the process generating $(X_{.1} \mid X_{..})$, more specifically α is defined by the condition:

$$X_{.1} \perp\!\!\!\perp \theta \mid \alpha, X_{..} \quad (47)$$

and β is the parameter sufficient for the process generating the table conditionally on $(X_{.1}, X_{..})$, i.e.:

$$(X_{11}, X_{12}) \perp\!\!\!\perp \theta \mid \beta, X_{.1}, X_{..} \quad (48)$$

Finally, we shall also assume the prior independence of α and β :

$$\alpha \perp\!\!\!\perp \beta \quad (49)$$

or, equivalently under the first cut:

$$\alpha \perp\!\!\!\perp \beta \mid X_{..} \quad (50)$$

This hypothesis ensures a Bayesian cut between β and the column total $X_{.1}$ conditionally on $X_{..}$. Thus for an inference on any function of β , $(X_{.1}, X_{..})$ may be considered as exogenous. The first cut, along with (48) and (49), also implies:

$$\beta \perp\!\!\!\perp (X_{.1}, X_{..}) \quad (51)$$

The independent sampling of columns may now be written as:

$$X_{11} \perp\!\!\!\perp X_{12} \mid X_{.1}, X_{..}, \theta \quad (52)$$

and β_j , $j = 1, 2$ are now the parameters sufficient for the process generating the j th column conditionally on its total (remember that $X_{.2} = X_{..} - X_{.1}$: conditionally on $X_{..}$, knowing $X_{.1}$ is therefore equivalent to knowing $X_{.2}$), more specifically the β_j 's are defined by the conditions:

$$X_{1j} \perp\!\!\!\perp \theta \mid X_{.1}, X_{..}, \beta_j \quad j = 1, 2 \quad (53)$$

Finally, we assume that each column is generated as an independent binomial sampling:

$$(X_{1j} \mid X_{.j}, X_{..}, \theta) \sim Bi(X_{.j}, \beta_j) \quad j = 1, 2 \quad (54)$$

i.e.

$$P(X_{1j} = x_{1j} \mid X_{.j}, X_{..}, \theta) = \binom{X_{.j}}{x_{1j}} \beta_j^{x_{1j}} (1 - \beta_j)^{X_{.j} - x_{1j}} \quad (55)$$

where $X_{.2} = X_{..} - X_{.1}$.

Remark: If furthermore $(X_{.1} \mid X_{..}, \alpha)$ is also assumed to follow a binomial distribution, then $(X \mid X_{..}, \theta)$ has also a multinomial distribution as in the previous section, but this is not important for what follows.

The parameter of interest is now defined as

$$\lambda = \frac{\beta_1(1 - \beta_2)}{\beta_2(1 - \beta_1)} \quad (56)$$

The value $\lambda = 1$, meaning $\beta_1 = \beta_2$, is often read as a condition of "homogeneity" and implies, again, a row-column independence of the 2x2 table for whatever process generating the column totals (i.e. for any distribution of $(X_{\cdot 1} | X_{\cdot}, \alpha)$).

The problem is again to check whether the inference on λ , a function of β only, may be done, without loss of information, conditionally on the two totals $Z = (X_{\cdot 1}, X_{1 \cdot})$. We follow the same route as in the previous section and first check whether the sampling distribution generating $(X_{11} | Z, X_{\cdot}, \theta)$ depends on λ only. This is easily seen by noticing that the sampling distribution of the 2x2 table conditionally on $(X_{\cdot 1}, X_{1 \cdot})$ is an independent product of two binomial distributions with parameters $(\beta_1, X_{\cdot 1})$ and $(\beta_2, X_{1 \cdot} - X_{11})$ respectively, from which one concludes that the sampling distribution of the 2x2 table conditionally on $(Z, X_{\cdot})$ is again a non-central hypergeometric distribution of parameter λ , exactly as in (37); thus condition (39) is again satisfied even if the column total $(X_{\cdot 1} | X_{\cdot})$ is not binomially generated.

Conditionally on $X_{\cdot}$, a Bayesian cut between the parameter of interest, λ , and the margins $Z = (X_{\cdot 1}, X_{1 \cdot})$ is again impossible. Indeed, the sampling distribution generating $(X_{11} | X_{\cdot 1}, X_{1 \cdot})$ depends only on β and is given by

$$p(X_{11} = x_{11} | X_{\cdot 1}, X_{1 \cdot}, \beta) = \frac{X_{\cdot 1}! (X_{1 \cdot} - X_{11})! (1 - \beta_1)^{X_{\cdot 1} - x_{11}} (1 - \beta_2)^{X_{1 \cdot} - x_{11} - x_{11}}}{\sum_{x_{11} \in V} \prod_{i,j} \frac{\lambda^{x_{ij}}}{x_{ij}!}} \quad (57)$$

Thus, the marginal distribution of the margins, $Z = (X_{\cdot 1}, X_{1 \cdot})$ (given the grand total $X_{\cdot}$) identifies both β_1 and β_2 so that $\theta = (\alpha, \beta)$ is a minimal sufficient parameter of the process generating $(Z | X_{\cdot})$ and no prior distribution can satisfy the condition of a Bayesian cut, namely $\lambda \perp\!\!\!\perp \theta | X_{\cdot}$, because λ is a function of β and therefore of θ .

Given the impossibility of a Bayesian cut, one may still ask whether some prior specification could allow Z and λ to be mutually exogenous, i.e. (39) along with (41), or equivalently (42) under the first cut. Now (43) may be decomposed as follows

$$p(Z | \lambda, X_{\cdot}) = p(X_{\cdot 1} | \lambda, X_{\cdot}) p(X_{1 \cdot} | \lambda, X_{\cdot}, X_{\cdot 1}) \quad (58)$$

The question whether $p(Z | \lambda, X_{\cdot})$ can possibly be independent of λ may be considered in two steps. The first term, $p(X_{\cdot 1} | \lambda, X_{\cdot})$ does not depend on λ under (47) and (49) because λ is a function of β only, i.e. (47) and (49) jointly imply $X_{\cdot 1} \perp\!\!\!\perp \lambda | X_{\cdot}$. Let us consider the second term

$$p(X_{1 \cdot} | \lambda, X_{\cdot}, X_{\cdot 1}) = \int p(X_{1 \cdot} | \beta_0, X_{\cdot}, X_{\cdot 1}) m(\beta_0 | \lambda) d\beta_0 \quad (59)$$

where β_0 is a function of λ : $\beta_0 = b(\lambda)$ such that (λ, β_0) is a (smoothly) bijective transformation of β .

As an example, when β'_j 's are a priori distributed as independent Beta variates with parameters (a_j, b_j) (as is the case when θ has a priori a Dirichlet distribution), it is shown in Proposition 2 of the appendix that

$$p(X_{1.} = x_{1.} \mid \lambda, X_{..}, X_{.1}) = c \frac{X_{.1}! X_{.2}!}{\sum_{x_{11} \in V} \prod_{i,j} \frac{\lambda^{x_{11}}}{x_{ij}!}} \cdot \frac{E[(1 - Y + \lambda Y)^{-(a_2 + b_2 + X_{.2})}]}{E[(1 - W + \lambda W)^{-(a_2 + b_2)}]} \quad (60)$$

where

$$\begin{aligned} c &= \frac{\Gamma(a_{.} + b_{.})}{\Gamma(a_{.})\Gamma(b_{.})} \frac{\Gamma(a_{.} + x_{1.})\Gamma(b_{.} + x_{2.})}{\Gamma(a_{.} + b_{.} + X_{..})} \\ Y &\sim \beta(a_{.} + x_{1.}, b_{.} + x_{2.}) \\ W &\sim \beta(a_{.}, b_{.}) \end{aligned}$$

For the same reason as for the multinomial sampling, one may conclude that under independent binomial sampling, conditioning on the margins is generally not an admissible reduction.

3.4 Sampling Conditional on the Margins

A natural idea is to consider a sampling model where the table is actually generated conditionally on the two margins $Z = (X_{.1}, X_{1.})$. This situation is exemplified in the celebrated example of tea-tasting where a lady is tested for her ability to judge whether tea or milk was poured in the cup first (see Fisher (1935b)).

We again suppose the first cut (29),(30) and (31), making $X_{..}$ exogenous for the inference on any function of θ , the parameter (minimal) sufficient for the process generating $(X \mid X_{..})$. Conditioning on the two margins would be admissible under a second cut, namely under the condition

$$\delta \perp\!\!\!\perp \lambda \mid X_{..} \quad (61)$$

or equivalently, under the first cut:

$$\delta \perp\!\!\!\perp \lambda \quad (62)$$

where δ is a parameter sufficient for the $(Z \mid X_{..})$ -DGP:

$$Z \perp\!\!\!\perp \theta \mid \delta, X_{..} \quad (63)$$

and λ is a parameter sufficient, as in (37), for the $(X | Z, X_{..})$ -DGP. Under these definitions, $\theta = (\lambda, \delta)$ but the analysis of the multinomial sampling and of the independent binomial samplings show that under (61), the sampling distribution of $(X | X_{..})$ may not be multinomial, and the sampling distributions of $(X_{1j} | X_{1.}, X_{..})$ may not be independent binomial. It is therefore fitting to ensure a proper understanding of the sampling procedure.

Let us consider, along with the tea-tasting experiment, another example. Let $X_{..}$ be the number of students in a classroom and let these students be classified into {failure, success} in a mathematics exam (row criterion) and into {type A, type B} for the schooling system followed the year before (column criterion). Here, the $(X_{1.} | X_{..})$ -DGP would reflect a recruitment structure of the school and the $(X_{.1} | X_{..})$ -DGP would reflect the teacher's academic excellence policy. Note that in this case one could accept that $X_{1.}$ and $X_{.1}$ be independent (conditionally on $X_{..}$) whereas in the tea-testing experiment the joint $(X_{1.}, X_{.1} | X_{..})$ -DGP would be degenerate, in the sense that $P(X_{1.} = X_{.1} | X_{..}) = 1$, once the tea-taster is told, before guessing, the number of cups of each type. In both examples, the row-column independence should be characterized as a property of the $(X_{11} | X_{1.}, X_{.1}, X_{..})$ -DGP without any restriction on the $(X_{1.}, X_{.1} | X_{..})$ -DGP. In this conditional process, the condition of row-column independence is a condition that each individual is assigned to the row and to the column criterion as if he had been selected randomly, in the sense of a hypergeometric sampling. This is equation (37) with $\lambda = 1$, i.e. (see Plackett (1974)):

$$p(X_{11} = x_{11} | x_{1.}, x_{.1}, x_{..}, \theta) = \frac{\prod_{i,j} \frac{1}{x_{ij}!}}{\sum_{x_{11} \in V} \prod_{i,j} \frac{1}{x_{ij}!}} \quad (64)$$

$$= \frac{\begin{pmatrix} x_{1.} \\ x_{11} \end{pmatrix} \begin{pmatrix} x_{.2} \\ x_{21} \end{pmatrix}}{\begin{pmatrix} x_{..} \\ x_{.1} \end{pmatrix}} = \frac{\begin{pmatrix} x_{.1} \\ x_{11} \end{pmatrix} \begin{pmatrix} x_{.2} \\ x_{12} \end{pmatrix}}{\begin{pmatrix} x_{..} \\ x_{1.} \end{pmatrix}}$$

The nature of this biconditional sampling is such that the joint process of the table can not anymore be a counting process of independent individual units, as was the case of the multinomial or of the independent binomial samplings. For instance, in the tea-testing case, a possible sampling scheme for a tea-taster aware of his (her) absolute inability of recognizing which liquid has been poured first, would be to assign randomly each cup selected in the order of a random

permutation until he faces the constraint $X_{\cdot 1} = X_{1\cdot}$. Note that in such a case, which particular probability is used in the randomization is irrelevant as far as it is independent of the true type of the cup but the joint probability of all cups cannot be independent, if only because of the constraint $X_{\cdot 1} = X_{1\cdot}$ (i.e. the assignation of the $n - X_{1\cdot}$ last ones is determined by the joint assignation of the $X_{1\cdot}$ first ones).

Thus, in the case of this conditional sampling, where the $(X_{11} \mid X_{\cdot 1}, X_{1\cdot}, X_{\cdot\cdot})$ -distribution represents an actual sampling process, the hypergeometric distribution (64) happens to represent the only distribution displaying the row-column independence; in other words, this hypothesis of independence is a simple (a one-point) hypothesis in the space of $(X_{11} \mid X_{\cdot 1}, X_{1\cdot}, X_{\cdot\cdot})$ -distributions but leaves completely free the $(X_{\cdot 1}, X_{1\cdot}, X_{\cdot\cdot})$ -distribution. Furthermore, the alternative hypothesis should be specific to any actual situation. For instance, the alternative hypothesis could model the (optimal control) behavior of a tea-taster trying sequentially each cup with a view to exert her guessing ability or the teacher's policy when he would be willing to differentiate his evaluation according to the origin of his student. In such circumstances there is no reason to believe that the non-central hypergeometric distribution (37) would be the only reasonable alternative. Also, one could expect that the sampling model under the alternative hypothesis would be parametrized with more than one parameter.

In conclusion, the only situation where the sampling model conditional on the two margins does not lead to lose information is a case where the modelling of the non-independent alternative should be context specific, because the framework of counting independent units is not anymore relevant. In particular, the non-central hypergeometric distribution does not appear anymore as a natural modelling of the $(X_{11} \mid X_{\cdot 1}, X_{1\cdot}, X_{\cdot\cdot})$ -DGP. Therefore the prior information on that parameter and the form of the critical region should also be context specific.

4 Some numerical Results

4.1 Simulation Design and Methods

In this section, we want to evaluate numerically how much information is lost when the inference on λ is based on the sampling process conditional on the margins, considered as an approximation, rather than on the actual sampling process, considered as an exact one. More specifically we want to evaluate the

respective roles of the sample size and the relative precision of the prior information.

As mentioned in section 2, one may compare either the two posterior distributions, the exact one and the approximate one, by means of a (predictively) expected divergence (or distance) or the two Bayesian risk functions, by means of the relative efficiency. As those quantities are not available analytically one has to resort to numerical comparisons through simulation methods. As a matter of fact, the numerical and the statistical problems involved in the comparison of the posterior distributions are vastly different from those involved in the comparison of the risks functions. In order to keep this paper within a reasonable size we have found more useful to concentrate on the presentation of numerical results relative to the comparison of posterior distributions only, leaving for a "further work in progress" a deeper comparison of the two risks functions.

Thus, we want to evaluate an expected discrepancy between both distributions, as written in (13):

$$L = E[D(m_A(\lambda | x), m(\lambda | x))] \quad (65)$$

where the expectation is taken with respect to the exact predictive distribution and $D(\cdot, \cdot)$ is a distance or a divergence on the space of the distributions on λ . We consider two specifications for $D(\cdot, \cdot)$. The first one is the entropy divergence defined by:

$$\begin{aligned} D_e(m_A(\lambda | x), m(\lambda | x)) &= \int \ln\left(\frac{m(\lambda | x)}{m_A(\lambda | x)}\right) m(\lambda | x) d\lambda \\ &\in (0, \infty) \end{aligned} \quad (66)$$

The second one is the Hellinger distance:

$$D_H(m_A(\lambda | x), m(\lambda | x)) = \int (\sqrt{m(\lambda | x)} - \sqrt{m_A(\lambda | x)})^2 d\lambda \quad (67)$$

$$\begin{aligned} &= 2(1 - \int \sqrt{m(\lambda | x) m_A(\lambda | x)} d\lambda) \\ &\in [0, 2] \end{aligned} \quad (68)$$

In these experiments we simulate both cases of a multinomial sampling with Dirichlet prior and independent binomial samplings with independent Beta priors. For the multinomial sampling, the simulation consider 1000 drawings of $\theta = (\theta_{ij})$, $j = 1, 2$ distributed as a Dirichlet distribution of parameter $n_0 p_0$, where p_0 is fixed to $(0.25, 0.25, 0.25, 0.25)$ and for each drawing of θ , we draw a multinomial distribution of parameter (n, θ) , which constitutes the table X .

For the independent binomial sampling, the simulation consider 1000 drawings of two independent Beta distributions with the same parameter $n_0 p_1$, where p_1 is fixed to (0.5, 0.5) and for each one of them, let θ_1 and θ_2 be their values, a draw of two independent Binomial distributions of parameters (n, θ_1) and (n, θ_2) respectively, from which we constitute the table X .

For the two measures of comparison, we consider their behavior in relation to the values of two parameters: n the sample size and n_0 the relative precision of the prior information. We consider the following values of these parameters:

Multinomial Sampling	n	10, 15, 20, 25, 30, 35, 40, 45, 50
	n_0	10, 15, 20, 25, 30, 35, 40, 45, 50
Ind. Binomial Sampling	n	5, 10, 15, 20, 25, 30, 40, 50
	n_0	5, 10, 15, 20, 25, 30, 40, 50

In the multinomial sampling it is shown in Proposition 3 of the appendix that:

$$m(\lambda|x) = \frac{\lambda^{b_4-1} I_1(\lambda, x)}{\int_0^\infty \lambda^{b_4-1} I_1(\lambda, x) d\lambda} \quad (69)$$

where

$$\begin{aligned}
I_1(\lambda, x) &= \int_0^1 \delta^{b_1-1} (1-\delta)^{b_2-1} (1-\delta+\lambda\delta)^{-b_3} d\delta \\
b_1 &= x_{\cdot 1} + a_{\cdot 1} \\
b_2 &= x_{\cdot 2} + a_{\cdot 2} \\
b_3 &= x_{1\cdot} + a_{1\cdot} \\
b_4 &= x_{11} + a_{11} \\
(a_{11}, a_{12}, a_{21}, a_{22}) &= n_0 p_0
\end{aligned}$$

and

$$m_A(\lambda|x) = \frac{\lambda^{c_4-1} I_2(\lambda) p(y|z, \lambda)}{\int_0^\infty \lambda^{c_4-1} I_2(\lambda) p(y|z, \lambda) d\lambda} \quad (70)$$

where

$$\begin{aligned}
p(y|z, \lambda) &= p(x_{11}|x_{1\cdot}, x_{\cdot 1}, \lambda) \\
&= \frac{\lambda^{x_{11}}}{\prod_{i,j} x_{ij}!} \left(\sum_{x_{11} \in V} \frac{\lambda^{x_{11}}}{\prod_{i,j} x_{ij}!} \right)^{-1} \\
I_2(\lambda) &= \int_0^1 \delta^{c_1-1} (1-\delta)^{c_2-1} (1-\delta+\lambda\delta)^{-c_3} d\delta \\
c_1 &= a_{\cdot 1} \\
c_2 &= a_{\cdot 2} \\
c_3 &= a_{1\cdot} \\
c_4 &= a_{11} \\
(a_{11}, a_{12}, a_{21}, a_{22}) &= n_0 p_0
\end{aligned}$$

Furthermore, in the independent Binomial sampling, we have from the Proposition 4 of the appendix that:

$$m(\lambda|x) = \frac{\lambda^{d_4-1} I_3(\lambda, x)}{\int_0^\infty \lambda^{d_4-1} I_3(\lambda, x) d\lambda} \quad (71)$$

where

$$\begin{aligned}
I_3(\lambda, x) &= \int_0^1 \beta_0^{d_1-1} (1-\beta_0)^{d_2-1} (1-\beta_0+\lambda\beta_0)^{-d_3} d\beta_0 \\
d_1 &= x_{1\cdot} + a_{\cdot} \\
d_2 &= x_{2\cdot} + b_{\cdot} \\
d_3 &= X_{\cdot 1} + a_1 + b_1 \\
d_4 &= x_{11} + a_1 \\
(a_1, b_1) &= (a_2, b_2) = n_0 p_1
\end{aligned}$$

and

$$m_A(\lambda|x) = \frac{\lambda^{\epsilon_4-1} I_4(\lambda) p(y|z, \lambda)}{\int_0^\infty \lambda^{\epsilon_4-1} I_4(\lambda) p(y|z, \lambda) d\lambda} \quad (72)$$

where

$$p(y|z, \lambda) = p(x_{11}|x_{1\cdot}, x_{\cdot 1}, \lambda)$$

$$\begin{aligned}
&= \frac{\lambda^{x_{11}}}{\prod_{i,j} x_{ij}!} \left(\sum_{x_{11} \in V} \frac{\lambda^{x_{11}}}{\prod_{i,j} x_{ij}!} \right)^{-1} \\
I_4(\lambda) &= \int_0^1 \delta^{\epsilon_1-1} (1-\delta)^{\epsilon_2-1} (1-\delta+\lambda\delta)^{-\epsilon_3} d\delta \\
e_1 &= a, \\
e_2 &= b, \\
e_3 &= a_1 + b_1 \\
e_4 &= a_1 \\
(a_1, b_1) &= (a_2, b_2) = n_0 p_1
\end{aligned}$$

Consequently, the evaluation of the divergences (or distances) between the two distributions ((66) and (67)) require the numerical integration on the rectangle $(0, 1) \times (0, \infty)$. The numerical integrations are carried out by the gaussian quadrature method applied in a first step to the interval $(0, 1)$ and in a second step to the interval $(0, \infty)$. Because of the high cost of these double numerical integrations, we have considered a (n, n_0) -grid just fine enough for providing insights about the general behavior of the information loss in the (n, n_0) -space.

All the simulation have been performed with programs written in C language which call subroutines in Fortran language for the pseudo-random number generation (Ranlib package, Brown and Lovato (1991)) and for the numerical integration (Quadpack routines, Piessens et al (1983)). The figures were made in Splus 3.1. Some illustrative results along with the main conclusions are presented in next section, more results are presented in appendix B.

4.2 Results

4.2.1 Evaluating the loss of information

For both the multinomial sampling and the independent binomial sampling, the evaluation of the Hellinger distance and the Entropy divergence gave roughly similar results for the range of sample sizes and precision of prior information we have simulated. More specifically let us consider the loss of information $L(n, n_0)$, defined in (65), as a function of n , the sample size, and n_0 , the precision of the posterior distribution.

Figures 1, 2, 3 and 4 (b) suggest that $L(n, n_0)$ behaves approximately as:

$$L(n, n_0) \sim h\left(\frac{n}{n_0}\right) \quad (73)$$

where $h : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ is such that $h' > 0$ or equivalently: $L(n, \alpha(c)n_0) = c$ with $\alpha : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ is such that $\alpha' > 0$. This suggests that $L'_n > 0$ and $L'_{n_0} < 0$, implying in particular that the loss of information does *not* vanish asymptotically (in n). The graphs also suggest that $L''_{n,n} < 0$, i.e. for any n_0 , $L(n, n_0)$ is concave in n (Figures 1,2,3 and 4 (c)) and $L''_{n_0,n_0} > 0$ i.e. for any n , $L(n, n_0)$ is convex in n_0 (Figures 1,2,3 and 4 (d)) and eventually tends to 0 when $n_0 \rightarrow \infty$ (in which case the prior distribution becomes dogmatic and remains unrevised by the sample, whether the revision is "exact" or "approximate"). Thus, in this case, the approximation (12) can be qualified as a "small sample approximation" in the sense that the error of approximation tends to vanish when the prior precision tends to dominate the sample precision (more precisely $(n/n_0) \rightarrow 0$) or when $n_0 \rightarrow \infty$ for fixed n whereas this error does not tend to zero, but rather tends to increase when the sample size $n \rightarrow \infty$ with fixed prior precision n_0 .

4.2.2 Remarks on the simulation accuracy

Let us also mention that the graphs relative to the Hellinger distance appears to display a more regular pattern than those relative to the entropy divergence, particularly for the case of independent binomial sampling. These irregularities are due to simulation problems. These graphs are based on simulation of size $N = 1000$; previous trials with $N = 100$ revealed dramatic irregularities. In general terms, one may suspect that the multinomial sampling generate more variations than the independent binomial sampling because this one may be obtained as a conditioning (on one margin) of the first one. One may also suspect more variation for the Entropy divergence because it is not compact-valued as is the Hellinger distance.

Tables 1 to 4 in appendix B, along with figure 5 may also help getting more precise an insight on the simulation problem. Let us write D for the simulated divergence or distance. Tables 1 to 4 give the main characteristics of the simulation for each configuration of n and n_0 . The mean values ("E(D)") have been plotted and commented in figures 1 to 4. The column "Gamma(D)", defined as:

$$\gamma = \frac{E(D - \bar{D})^3}{(E(D - \bar{D})^2)^{\frac{3}{2}}} \quad (74)$$

indicates a systematic positive skewness. The penultimate column gives a simulation error estimated by the ratio $\sqrt{(\frac{V(D)}{N})}$. This suggest confidence intervals of width roughly between 0.01 and 0.02, wider for the Entropy than for the

Hellinger distance. The noise to signal ratio, or coefficient of variation, is given in the last column and illustrated in figure 5 in the form of scatter diagrams between the standard deviations and averages with the main diagonal representing a ratio equal to 1. Note that there is a clear positive association between the standard deviation and the average (i.e. higher average is associated to higher dispersion). Thus one should expect less precise (and more irregular simulations) contour curves for regions of higher values of D . Furthermore, the graphs present a concave rather than linear relationship between the average and the dispersion, with a narrower and closer to the main diagonal relationship for low than for large values and for the Hellinger distance than for the Entropy divergence. Note also that the coefficient of variation is, in general, higher for the Entropy than for the Hellinger distance, particularly for the independent binomial sampling.