

ON FITTING MODE SPECIFIC CONSTANTS IN THE PRESENCE OF NEW OPTIONS IN RP/SP MODELS

Elisabetta CHERCHI

CRiMM - Dipartimento di Ingegneria del Territorio
Facoltà di Ingegneria - Università di Cagliari
Piazza d'Armi, 16 - 09123 Cagliari, Italia
Tel: 39 70 675 5274; Fax: 39 70 675 5261; E-mail: echerchi@unica.it

and

Juan de Dios ORTUZAR

Departamento de Ingeniería de Transporte
Pontificia Universidad Católica de Chile
Casilla 306, Cod. 105, Santiago 22, Chile
Tel: 56-2-686 4822; Fax: 56-2-553 0281; E-mail: jos@ing.puc.cl

ABSTRACT

We treat the problem of fitting alternative specific constants (ASC) in models estimated with a mixture of revealed preference (RP) and stated preference (SP) data to forecast the market shares of new alternatives. This important problem can have non-trivial solutions, particularly when some of the SP alternatives are completely revamped versions of existing ones (i.e. an advanced passenger train replacing a normal railway service). As there is no explicit treatment of this problem in the literature we examined it in depth and illustrated it empirically using data especially collected to analyse mode choice in a corridor to the West of Cagliari. We propose a hopefully useful guide to this art (as no practical recipes seem to serve all purposes). Careful specification of the systematic component of utility functions in RP and SP, including the ASC, serves to illuminate the true nature of the underlying error structure in the different data sets, yielding superior forecasting models.

Keywords: RP/SP models, model specification, error structure, forecasting new options

1. INTRODUCTION

The recommended approach to estimating choice models intended to forecast demand for new alternatives involves pooling (if available) revealed preference (RP) and stated preference (SP) data. The former are based on observations of actual choices and allow to characterise current travel behaviour, while the latter provide information about user preferences for new alternatives or for alternatives that differ radically from existing ones. The combined use of both types of data allow to exploit their respective advantages and to overcome their specific limitations (Ben Akiva and Morikawa, 1990; Bradley and Daly, 1997; Louviere *et al*, 2000).

In joint estimation some attributes may be specified as generic for both the RP and SP alternatives and others specific to each sub-set. Variable specification depends on the analyst's choice; this is based on the type of data containing the variable, where it is measured with greater precision, and on the estimation results. Since alternative specific constants (ASC) allow simple logit (MNL) models to reproduce observed market shares, their specification as generic for the two data sources depends on the specific application context and on the market the analyst is trying to reproduce. Also, since the ASC represent the mean of the error difference and the errors are strictly related to the systematic utility specification, their values and significance are influenced by the degree to which the model is able to reproduce real behaviour. Thus, quite different results may be obtained depending on the way the RP/SP ASC are specified and it is not straightforward to define the best specification. As far as the authors are aware, this issue has not been addressed directly elsewhere.

The rest of the paper is organised as follows. Section 2 analyses some theoretical aspects concerning the inclusion of ASC in joint RP/SP model estimation, with particular reference to their role in the random and systematic utility components and to the MNL property of reproducing actual market shares. Section 3 briefly describes the data used to estimate the models presented in section 4, where many different ASC specifications are tested and analysed, with reference to various hypotheses about the specification of the variables and error terms. Section 5 summarises our main conclusion.

2. ALTERNATIVE-SPECIFIC CONSTANTS IN LOGIT MODELS

Following the classical formulation of random utility models (Domencich and McFadden, 1975), individuals are assumed to choose among several available alternatives ($A_j \in \underline{A}(q)$), where $\underline{A}(q)$ is the set of options available to individual q), associating to each an index of preference (called utility) that depends on its characteristics and those of the individual:

$$U_{qj} = U(\underline{X}_{qj}, \varepsilon_{qj}) \quad (1)$$

where $\underline{X}_{qj}$ is a vector of measurable attributes, and ε_{qj} is a random term typically introduced by the modeller due to his/her inability to obtain perfect information about the individuals. Under the almost universally accepted assumption of additive random terms, utility is specified as the sum of a systematic, representative or observable part (V_{qj}) and a random error (ε_{qj}):

$$U_{qj} = V_{qj}(\underline{X}_{qj}) + \varepsilon_{qj} \quad (2)$$

Once assumptions on the distribution of the random error terms have been made, the core of the model is represented by specifying systematic utility¹ as a function of the parameters θ , to be estimated, and the attributes $\underline{X}_{qj}$, which may be of two types:

- variables representing attributes of the alternative as perceived by each traveller (namely travel times and costs) and characteristics of the individual (mainly socio-economic features)
- alternative specific constants (ASC) that do not depend on the individual nor on the choices made but take the value one for the alternative in which they are included and zero otherwise. ASC are normally interpreted as representing the net influence of all unobserved, or not explicitly included, characteristics of the individual or the alternative in its utility function (Ortúzar and Willumsen, 2001).

As the variables are key to explaining the phenomenon, rightly much effort has been directed towards finding the best way of introducing them in the systematic utility. In particular, the ASC should be included in all alternatives except one, which is left as a reference, because logit models work with the differences of systematic utilities. For the analysis that follows it is interesting to note that this ASC specification is not simply “good practice” but has a theoretical underpinning.

Let us first recall that as utility is a random variable for the modeller s/he can only evaluate the probability that an individual q will choose alternative A_j , the latter being that with the highest utility:

$$p_{qj} = \text{prob}(U_{qj} \geq U_{qi}) \quad \forall A_i \in A(q), \quad i \neq j \quad (3)$$

Now re-arranging equation (2), in order to show the two types of attributes described above, we get:

$$U_{qj} = K_j + V'_{qj}(\underline{X}_{qj}) + \varepsilon_{qj} \quad (4)$$

and substituting into equation (3) we obtain²:

$$p_{qj} = \text{prob}(K_j - K_i + V'_{qj} - V'_{qi} \geq \varepsilon_{qi} - \varepsilon_{qj}) \quad \forall A_i \in A(q), \quad i \neq j \quad (5)$$

where the superscript (') simply indicates those parts of the systematic utility depending upon the explanatory variables alone, i.e. not including the constant K_j .

Note from (2), first, that under the assumption of additive random terms “... *there is no real distinction between shifting the mean of the disturbance of one alternative's utility and shifting the systematic component by the same amount*” (Ben-Akiva and Lerman, 1985). Secondly and as stated in equation (5), since what really matter are the differences in utilities “... *if the mean of alternative j 's disturbance is some amount greater than that of alternative i 's disturbance, we can fully represent the difference by adding that amount to V_{jq} . This implies that as long as one can add constants to the systematic component, the*

¹ Since attributes not explicitly specified in the systematic part of the utility function are included in the error term of the corresponding alternative, there is a strong relation between the systematic utility specification and the error terms.

² Leaving one of the options i as reference for the ASC specification does not modify the general discussion.

means of the disturbances can be defined as equal to any constant without loss of generality” (Ben-Akiva and Lerman, 1985).

In this sense note that the widely accepted assumption that all errors have zero mean, either in the MNL and its derivatives or in more flexible model structures, is simply the most convenient one. For example, in the MNL model derivation the difference of any pair of IID Gumbel-distributed error components with parameters (η_1, λ) and (η_2, λ) has a mean equal to the difference of the respective distribution modes $(\eta_1 - \eta_2)$. The hypothesis of a null alternative leading to a logistic distribution with zero mean³ has no effect provided a full set of ASC (i.e. N-1 if there are N alternatives) is included in the specification.

Interestingly, the rule⁴ of specifying a full set of ASC only holds within a specific data set. In other words, when estimation is performed using different data sets, each pertaining to a specific subgroup of individuals (as is the case of mixed RP/SP data), there is no relation between the mean of the error terms in the different subsets (i.e. between the mean of the RP and SP error terms) because the data sets are complementary.

An important issue we want to discuss is whether to specify generic or specific RP/SP ASC (obviously for the same alternatives); this leads to the question of whether the mean of the error difference between any pair of alternatives (A_i, A_j) present in the two data sets is equal or not:

$$E(\varepsilon_j^{RP} - \varepsilon_i^{RP}) = E(\varepsilon_j^{SP} - \varepsilon_i^{SP}) \quad (6)$$

If equation (6) does not hold, specific ASC for RP and SP data should be used; conversely a generic ASC could be specified. We will come back to this point later.

Another important aspect which is useful to understand the ASC specification concerns the restrictions usually made on the error terms to ensure that their scale is consistent with that of the V's. The effect of this is that, similarly to any other attribute included in the utility function, the ASC parameter estimates are deflated in some sense by the standard deviation of the error component which is unknown to the modeller. Whatever the density function (f) chosen to describe the error component distribution, the probability of choosing an alternative A_j for an individual q (for simplicity, the index q will be omitted in the following notation), in a binary context, is given by⁵:

$$P_j = f[(V_j - V_i) / \sigma_\varepsilon] \quad \forall A_j \in A(q), \quad i \neq j \quad (7)$$

In joint RP/SP estimation the secondary data set (SP) is multiplied by a factor $\phi = \sigma^{RP} / \sigma^{SP}$ to ensure consistency of scale. Recall that the estimated parameters in an RP model are multiplied by a non-identifiable scale factor $\lambda^{RP} = \pi / \sigma^{RP} \sqrt{6}$ and in a SP model by another scale factor $\lambda^{SP} = \pi / \sigma^{SP} \sqrt{6}$. Note also that the ratio $\sigma^{RP} / \sigma^{SP} = \lambda^{SP} / \lambda^{RP}$.

³ In the MNL model the assumption that random terms have zero mean is equivalent to constraining the distribution mode to zero in the model derivation. The mean of a Gumbel distribution is equal to $(\eta + \gamma/\lambda)$ where γ is Euler's constant and λ the location parameter. Given two independent Gumbel distributed variables with parameters (η_1, λ) and (η_2, λ) the mean of their difference is equal to the difference of their respective distribution modes $(\eta_1 - \eta_2)$.

⁴ We now define this as a “rule” as it is theoretically justified.

⁵ σ_ε stands for the standard deviation of the difference $(\varepsilon_j - \varepsilon_i)$ in equation (7).

If we estimate two separate models using only the RP data in one and only the SP data in the other, we would get the following ASC values:

$$\begin{aligned}\widehat{K}_j^{*RP} &= \widehat{\lambda}^{RP} K_j^{RP} \\ \widehat{K}_j^{*SP} &= \widehat{\lambda}^{SP} K_j^{SP}\end{aligned}\tag{8}$$

where the asterisk refers to the estimated value, while the ASC without asterisk denote the population or observed values. Now, as long as the SP utilities are scaled by the ratio of the two scale parameters, in the joint RP/SP case we would get the following estimates even when we specify generic or specific ASC:

$$\begin{aligned}K_j^{*RP} &= \lambda^{RP} K_j^{RP} \\ K_j^{*SP} &= \lambda^{RP} K_j^{SP}\end{aligned}\tag{9}$$

where the asterisk values again refer to the estimates. Note that we are referring to ASC for simplicity, but the same applies to any variables included in the model.

We call attention to the different notation used in (8) and (9). The $\widehat{K}_j^*$ in (8) indicate ASC estimated with only one dataset (RP or SP only), while the K_j^* in (9) denotes ASC estimated in a joint RP/SP model. Moreover, note that we also use this notation for the scale parameters in (8) and (9); as they depend on the unknown variance of the errors we are not able to explain with the systematic part of utility, they should vary depending on the type of data involved in model estimation. In particular, the model variance depends on the data, on the “true parameters” (i.e. the population parameters) and on the utility form (i.e. linear in the attributes, non-linear, etc.). Therefore the scale parameters in (8) are $\widehat{\lambda}^{RP} = f^{RP}(\underline{\theta}', \underline{X}_{qj}^{RP})$ and $\widehat{\lambda}^{SP} = f^{SP}(\underline{\theta}', \underline{X}_{qj}^{SP})$ ⁶, while in (9), we only have $\lambda^{RP} = f(\underline{\theta}', \underline{X}_{qj}^{SP}, \underline{X}_{qj}^{RP})$ ⁷. Note that this latter result depends only on the fact that in joint RP/SP estimation we impose the condition that some parameters are the same in the RP and SP environments (at least one needs to be generic for both data sets in order to justify the joint estimation), but it is independent from the ratio ϕ (i.e. from the fact that we scale the SP utility). Note also that the same happens for the SP scale parameter; i.e. in a joint RP/SP estimation it would be $\lambda^{SP} = f(\underline{\theta}', \underline{X}_{qj}^{SP}, \underline{X}_{qj}^{RP})$. Finally note that as K_j^{RP} and K_j^{SP} are constant over individuals, they only influence the mean of the error difference but not their variance. Thus, given the specification we impose in a joint RP/SP estimation:

$$\begin{aligned}U_{qj}^{RP} &= K_j^{RP} + V'_{qj}(\underline{X}_{qj}^{RP}) + \varepsilon_{qj}^{RP} \\ \frac{\lambda^{SP}}{\lambda^{RP}} U_{qj}^{SP} &= \frac{\lambda^{SP}}{\lambda^{RP}} \left(K_j^{SP} + V'_{qj}(\underline{X}_{qj}^{SP}) + \varepsilon_{qj}^{SP} \right)\end{aligned}\tag{10}$$

⁶ Where f^{RP} and f^{SP} indicate that the way we can combine attributes and parameters (i.e. the utility form) is generally different in the two dataset. $\underline{\theta}'$ represents the vector of “true parameters”.

⁷ In the case of joint RP/SP estimation the scale parameter is that of RP because we scale the SP utilities. However, in joint RP/SP estimation the RP scale parameter is not necessarily equal to the RP scale parameter inherent to using only RP data, as by adding more information we should reduce the error variance.

as in a joint model the SP utilities are scaled by the ratio $\sigma^{RP}/\sigma^{SP} = \lambda^{SP}/\lambda^{RP} = \phi$, the SP constants we would obtain from estimating (10) should be:

$$\phi K_j^{*SP} = \frac{\lambda^{SP}}{\lambda^{RP}} (\lambda^{RP} K_j^{SP}) = \lambda^{SP} K_j^{SP} \quad (11)$$

To clarify the above result note that when we scale one data set (the SP set in our case) and impose equality between the two variances, we usually refer to the variances at each environment (i.e. before combining the data). However, as illustrated above, both variances (in RP and SP) should vary when we combine the two data sets. Therefore the ϕ parameter that guarantees consistency of scale in equation (10) is the ratio between the RP and SP scale parameters in the RP/SP environment, i.e. $\phi = \frac{\lambda^{SP}}{\lambda^{RP}}$.

Given the above result the estimated mean of the distribution for each data set in a joint RP/SP estimation should be:

$$\begin{aligned} E(\varepsilon_{qj}^{RP} - \varepsilon_{qi}^{RP}) &= \lambda^{RP} K_j^{RP} = K_j^{*RP} \\ E(\varepsilon_{qj}^{SP} - \varepsilon_{qi}^{SP}) &= \frac{\lambda^{SP}}{\lambda^{RP}} K_j^{*SP} = \frac{\lambda^{SP}}{\lambda^{RP}} \underbrace{(\lambda^{RP} K_j^{SP})}_{K_j^{*SP}} = \lambda^{SP} K_j^{SP} \end{aligned} \quad (12)$$

Note that although in the SP utilities the RP scale parameter cancels out (see equations (11) and (12)) if we multiply the SP-ASC (estimated using the mixed RP/SP data set) by the ratio of the scale parameters, we should obtain SP-ASC estimates which are different from those we would have obtained had we used only the SP data, i.e. $\hat{K}_j^{*SP} = \hat{\lambda}^{SP} K_j^{SP}$. How different they are depends on the influence of each data set in the estimation of the RP/SP generic parameters. The same happens if generic RP/SP ASC are specified and, with very little difference, for any type of variables.

Further knowledge on how to introduce ASC into RP/SP alternatives can be obtained by examining the use of the estimated models for prediction purposes. As the ASC play an important role in reproducing market shares the decision on how to specify them in a mixed RP/SP model depends on what are the different alternatives expected to predict.

Unfortunately, this is neither a simple matter nor a clear one. In fact, no theory exists to provide specific guidance on this issue; in particular (i) no *a priori* knowledge is available for deciding what are the different alternatives expected to predict and/or whether in prediction mode the alternatives will be the same; (ii) there seems not to be agreement on how to specify models to cope with these problems and (iii) “*although perhaps obvious, the whole procedure of passing from estimation to prediction in mixed RP/SP modelling does not appear to have been reported before*” (Ortúzar, 2000).

Although the process of data enrichment (i.e. the joint RP/SP estimation) dates back almost 15 years (Morikawa, 1989) we have found that very few papers tackle the problem of ASC specification. Following the suggestion of one referee we discussed this point through private communications with several experts who had worked with RP/SP models in practice. There seems to be a general agreement (Hensher, 2002; Louviere, 2004; Train, 2004) that the state of practice in using ASC in a joint estimation is:

1. Always use RP-ASC for existing alternatives. Do not rescale when using for prediction.
2. Always use SP-ASC for new alternatives but rescale when moving into the RP domain.
3. Calibrate a specific RP-ASC (or constrained generic RP/SP ASC, if there are common alternatives in the two sub-sets) to reproduce RP shares in the absence of new alternatives.

Some points are worth highlighting regarding the above general “wisdom”. Firstly, it should be noted that the market shares estimated with only one set of data should not be equal, in general, to the market shares estimated using the mixed RP/SP data. For simplicity we will first discuss this point in the case of a MNL. As it is well-known, a MNL with just a full set of ASC (i.e. even without explanatory variables) will replicate *exactly* the market shares of the sample used to calibrate the model⁸. Moreover, in a MNL specified with only ASC if all individuals have the same choice set the observed market share for a given alternative will equal the probability of choosing that alternative (Ortúzar and Willumsen, 2001). Thus, in the case of a MNL model estimated with only SP data, say, the SP observed market share for alternative A_j should be equal to:

$$\widehat{MS}_j^{SP} = \widehat{P}_j^{SP} = \frac{e^{\hat{\lambda}^{SP} K_j^{SP}}}{\sum_i e^{\hat{\lambda}^{SP} K_i^{SP}}} \quad (13)$$

Instead, if a joint RP/SP estimation was used the SP market share should equal:

$$MS_j^{SP} = \frac{e^{\phi(\lambda^{RP} K_j^{SP})}}{\sum_i e^{\phi(\lambda^{RP} K_i^{SP})}} = \frac{e^{\frac{\lambda^{SP}}{\lambda^{RP}}(\lambda^{RP} K_j^{SP})}}{\sum_i e^{\frac{\lambda^{SP}}{\lambda^{RP}}(\lambda^{RP} K_i^{SP})}} = \frac{e^{\lambda^{SP} K_j^{SP}}}{\sum_i e^{\lambda^{SP} K_i^{SP}}} \quad (14)$$

Since in a joint RP/SP estimation all the estimated parameters are scaled by the RP scale factor, the SP-ASC should be multiplied by the ratio $\lambda^{SP}/\lambda^{RP}$ in order to reproduce the SP market shares when introduced into the RP model for prediction. Therefore, the market shares estimated from a joint RP/SP estimation will differ from the market shares estimated using only one set of data. This is true for both the RP and SP market shares.

Now, as only RP data represent observations of behaviour that have actually taken place, Daly and Rohr (1998) suggest estimating a joint RP/SP model using specific (preferably) or generic RP/SP ASC to determine values for all the model coefficients; then, they suggest adjusting the RP-ASC for the existing alternatives using the RP data, and constrain the other coefficients to the values estimated in the model calibration process to match the base-year observations. In the presence of new alternatives, they suggest estimating the SP-ASC using SP data alone, but to guarantee consistency with the existing alternatives it is

⁸ Note that for any model structure (including the ML) we can compute the market shares estimated by the model as $\mathcal{Z}p_{ij}$; but this will be, in general, different from the market shares observed in the sample used to calibrate the model. The only exception is an MNL with a full set of ASC.

necessary to constrain the values of the ASC for observed modes to be equal to the RP-ASC estimated with RP data alone.

However, it is likely that in the presence of new alternatives the RP-ASC of the existing ones will not be representative anymore of the observed market shares, as the introduction of a new alternative should modify the market. Therefore, all the ASC (i.e. both of the existing and of the new alternatives) should be re-estimated in order to reproduce the new market shares, and this will be normally different from those at the base year. If one is confident that the SP data can reproduce correctly the market shares of the population in the design year, then the ASC (both for the existing and for the new alternatives) should be adjusted to match the SP market shares. Unfortunately, there is general agreement that SP data is not necessarily consistent with the market shares of new alternatives; therefore, in most cases the new market shares will be unknown and the only guide left to fitting ASC is to rely on the estimation results (i.e. best fit), as long as the microeconomic conditions of the model are satisfied.

A more involved situation may arise in the case of SP designs based on changes with respect to a real situation. Consider the case of an alternative which has the same label (e.g. train) in RP and SP, but in some scenarios of the latter case it represents a very different alternative to that of the base year (i.e. scenarios involving “structural changes” such as variations in the number of stations or parking places, frequency, type of vehicle, etc.). In this case, a major problem is deciding whether the SP alternative is different from the RP one and what is the degree of difference. Common practise states that if one could argue that in forecasting all alternatives will be exactly the same as in the base year, then one could specify the same constants for them or at least test empirically for a generic specification. Otherwise, different constants should be estimated, each representing the specific market share.

It should also be noted first that even if generic RP/SP ASC are specified the estimated SP market share for a given alternative is never equal to the RP market share because the ratio of the scale parameters multiplies the SP set. Thus, following the example in equations (13)-(14), for a given specification with a generic RP/SP ASC for alternative A_j we estimate $K_j^* = \lambda^{RP} K_j$ and the RP and SP market shares for the same alternative A_j are given respectively by:

$$MS_j^{RP} = P_j^{RP} = \frac{e^{\lambda^{RP} K_j}}{\sum_i e^{\lambda^{RP} K_i}} \quad (17)$$

and

$$MS_j^{SP} = \frac{e^{\phi(\lambda^{RP} K_j)}}{\sum_i e^{\phi(\lambda^{RP} K_i)}} = \frac{e^{\frac{\lambda^{SP}}{\lambda^{RP}}(\lambda^{RP} K_j)}}{\sum_i e^{\frac{\lambda^{SP}}{\lambda^{RP}}(\lambda^{RP} K_i)}} = \frac{e^{\lambda^{SP} K_j}}{\sum_i e^{\lambda^{SP} K_i}} \quad (18)$$

which are obviously different. Note that if we constrain the ASC of an alternative to be the same in RP and SP, we assume that the two “real” market shares are equal. But due to the different sources of data used in estimation it seems more plausible to get two different

estimated values for the market shares of seemingly the same alternative in both cases. Note also that when generic RP/SP ASC are used, the difference in market shares is implicitly determined by the value of the scale factor $\phi = \sigma^{RP}/\sigma^{SP}$ (or $\lambda^{SP}/\lambda^{RP}$).

Notwithstanding, it is important to remark that in a full specification (i.e. including the ASC and the rest of the explanatory variables), imposing equality between RP and SP-ASC does not imply assuming the two “real” market shares are equal, but only that the mean of the difference of the unobserved part of the utilities are equal. Obviously the validity of the statement depends on the specification of the RP and SP systematic utilities. But this point is important as it relaxes the use of generic RP/SP ASC also for alternatives whose market shares in RP and SP are not exactly the same. The point is even more important, because estimation of ASC with mixed RP/SP data is required precisely when and because some alternatives are not the same in prediction as in the base year. In fact, if all the alternatives do not change from the base year, then best practice is to adjust the ASC with RP data.

When it is not clear if some alternatives remain the same after the changes proposed in the SP design, Ben-Akiva (2004) considers that using a RP scaled SP-ASC is questionable. If there are solid *a-priori* reasons to think that the new version of an alternative has an advantage over the existing one, he suggests calculating lower and upper prediction bounds for the new alternative: the lower bound using the RP-ASC and the upper bound using the RP scaled new SP-ASC. The test here would be the difference between the RP scaled new SP-ASC and the existing RP-ASC. If, as expected, this difference is positive (in case of negative difference, just stop) then divide it by the absolute value of the RP or RP scaled (in-vehicle) travel time coefficient. But “...only use this value if this ‘new alternative bonus’ is reasonable (e.g., less than 20% of the average travel time on the RP alternative). If an unreasonably large bonus is obtained then it could be better to consider external evidence to set a ‘reasonable’ upper bound” (Ben Akiva, 2004).

However, Train (2004) argues that it may not be a good idea to try and reach consensus on a general “right” way to specify ASC in RP/SP joint estimation. Rather, each case should be considered individually. Therefore, setting or not equality constraints on the constants for equal alternatives in the RP and SP sets, constraining the ASC to the values estimated using only one data set, scaling the SP-ASC, or not scaling them but calculate a range of variation for the market shares, all these approaches could be potentially correct and valid. Although we basically agree with this position, the previous discussion seems to offer some basic guidelines for the “subjective personal considerations”. These are as follows:

1. If the RP and SP alternatives are exactly the same, then the ASC should be adjusted to match the market shares of the base year.
2. If the SP data include new alternatives (i.e. not present in the base year) and if one truly believes that these data reproduce correctly the market shares of the population in forecasting, then the ASC (both for the existing and for the new alternatives) should be adjusted to match the SP market shares.
3. Conversely, if the market shares to match are unknown then we should rely on estimation results, i.e. as long as the above theoretical analyses are satisfied it might be useful to draw further considerations on the ASC specification from the model that provides the best statistical fit.
4. Finally, if the SP design implies substantial changes, such that alternatives sharing the same label could represent new options, and there is uncertainty on to what extent they are actually different, then best fit and analyst’s judgment, seem to be the only guide.

Now, regardless the way the ASC are specified (i.e. generic or specific), depending on the results for each specific context the application of a mixed RP/SP model in forecasting implies some limitations on the scenarios to be tested (Cherchi *et al.*, 2002). In particular:

1. If specific RP/SP ASC are estimated forecasts can only be made for scenarios involving structural characteristics not inferior to those described in the SP design, and in that case the rescaled SP-ASC should be used.
2. If specific RP/SP ASC are estimated and a scenario not involving structural changes is considered, the RP-ASC should be used.
3. Finally, if constrained generic RP/SP ASC are estimated, scenarios involving structural changes (for those alternatives with constrained RP/SP ASC) should not be tested, unless we get ASC with a fairly close value from estimation.

3. DATA BANK

The data for our empirical application was collected in 1998 and incorporated three kinds of surveys (Cherchi and Ortúzar, 2002). A qualitative survey using focus groups to gain a good understanding of the phenomenon, a RP survey describing current trips and a SP survey to evaluate the introduction of radical improvements to an existing alternative.

The RP survey considered data on trips actually made and the socio-economic features of the travellers. It was conducted for a sample of 300 families living in a corridor to the West of Cagliari, randomly extracted from the telephone directory. The 524 responses obtained yielded a total of 1,840 reported trips. However, only 40.7% of these had taken place in the corridor of interest so the useful sample was reduced to 748 observations. Eventually, and after many tests (to ensure that only people with a choice were included), a final sample of 338 observations was left for model estimation, reflecting the modal choice context among Car, Bus and Train users. The latter were clearly the minority of trips in the corridor.

A project by the local rail authority envisaged upgrading the local railway line serving the corridor to a commuter train service, increasing not only speed and frequency, but also the number of stations inside the corridor. Thus, a SP survey was designed primarily to test the introduction of this new passenger service, considerably different from the existing one. A stated choice experiment between the proposed new train service and the mode currently used by respondents was conducted for 90% of the 338 individuals. The bus users design included four variables at three levels (Time, Cost, Frequency and Comfort) and the car users design considered three variables at three levels (Time, Cost and Frequency) and one variable (Comfort) at two levels. The SP games were customised to ensure greater realism and the attribute levels were determined as a percentage variation of the values declared by the respondents of the RP survey (Cherchi and Ortúzar, 2002).

Four variables at three levels allow to estimate interactions between Cost, Frequency and Travel time, accounting for their non additive effects in the analysis. To estimate three two-term interactions with four variables, 27 hypothetical situations should be used (Louviere *et al.*, 2000), so a block design was employed reducing to nine the number of situations presented to each respondent (most experts agree that this number can easily be handled by a non-trained person without undue stress, Bates, 1998; Caussade *et al.*, 2005). The order of presentation of situations inside each block was randomised and a second check was made to ensure that no dominant cases ensued. A final sample (i.e. mixed RP/SP data set) of 45% train users, 32% car users and 24% bus users for a total of 1,396 observations was finally used in model estimation.

4. RP/SP MODELS WITH DIFFERENT ASC SPECIFICATIONS

Before estimating a mixed RP/SP model, good practice advises that separate models for each data set should be estimated first to detect which attributes are candidates to be generic to both sets (Louviere *et al*, 2000). When generic RP/SP variables are postulated for joint estimation we implicitly assume that there are no true divergences among them in both data sets apart from scale. Assuming the population parameters to be indeed equal and that the differences observed are mainly caused by the scale factors, we would get:

$$\theta^{RP} = (\lambda^{SP} / \lambda^{RP}) \theta^{SP} \quad (19)$$

To test this hypothesis we proceeded to estimate RP and SP models separately and then compared their coefficient values taking into account the scale problem involved. Table 1 shows the parameters estimated using only RP data (in the first column) and only SP data (second column), while the third column reports the ratio between the values in the previous two columns. These models, as well as all subsequent models in the paper (where a mixed RP/SP specification was adopted), were estimated using an Expenditure Rate (ER) specification (Jara-Díaz and Farah, 1987). This produced better results than the classic Train and McFadden (1978) Wage Rate specification for our data bank (containing only a small percentage of self-employed).

Table 1: Initial model estimation results

	θ (RP)	θ (SP)	Ratio RP/SP
Travel time PT	-0.03082 (-1.1)	-0.03932 (-1.6)	0.769
Travel time Car	-0.06949 (-3.0)	-0.1157 (-3.9)	0.595
Walking time	-0.07831 (-2.8)	-0.01869 (-1.9)	4.1 ***
Cost/g	-0.02514 (-3.9)	-0.0104 (-1.3)	2.5 **
Frequency	0.1747 (1.2)	0.2845 (3.2)	0.607
Comfort1	-1.6 (-2.1)	-2.153 (-9.5)	0.762
Comfort2	-0.645 (-1.2)	-1.037 (-6.6)	0.621
Transfer	-0.8322 (-1.2)	-0.6733 (-2.3)	1.24 *

Five out of eight ratios (θ^{RP}/θ^{SP}) in Table 1 are approximately equal to 0.7; this tell us that those attributes are good candidates for generic estimation in the RP and SP data sets⁹. The variables which present true divergences apart from scale are marked with asterisks. For these, their estimated parameters could even have different sign and most probably should not be put as common, but separate coefficients estimated for them in both data sets as recommended by Louviere *et al*. (2000). However, as the parameters for Cost/g and Transfer have a low t-test in one data set, a generic RP/SP specification could still be beneficial and improve estimation.

⁹ Note that even if several parameters showed the same ratio 0.7, this does not mean that we whould expect 0.7 to be the value of ϕ that best fits the data in a joint RP/SP estimation (see the discussion in page 5).

In fact, when Cost/g was allowed to be specific for both sets its SP parameter got a very low significance (t-test rejected at less than 81%) and also became positive when interaction terms were not included in the utility function. However, in models with both generic or specific RP/SP Cost/g parameters, the correlation between Bus and Train appears to depend on the presence or absence of interaction terms. The latter effect also appeared when Transfer was allowed to have specific RP and SP parameters; but those models were never superior to models estimated with Transfer being specified with a generic parameter in both data sets.

As the case of Walking time was clearly more involved, we decided to study the problem in more depth (Table 2). For example, model NL2 where Walking time is specific for the RP and SP alternatives, dominates model NL1 where the same variable is treated as generic (all parameters in NL2 have higher significance than those in NL1, including the interaction terms and the structural parameter ϕ_1 ; a likelihood ratio test is rejected at the 99.99% level). Notwithstanding, the RP Car-ASC in model NL2 is negative; this is unexpected as it does not reflect the condition that the present Car alternative is superior to the present Bus service. Interestingly, this same effect is observed every time that Walking time is specific in the RP and SP alternatives (see also model NL5) but it does not occur if the variable is treated as generic (see models NL1 and NL4). Indeed, Walking time is important in explaining SP utilities; for example if it is included only in the RP alternatives (model NL3), the overall fit clearly worsens (NL2 is better than NL3 at the 95% level).

In our SP design the Walking time variable was omitted under the assumption that car users did not perceive it (Cherchi and Ortúzar, 2002). We just saw that if Walking time is not included in the SP Car alternative results improved. Model NL4 is close to model NL1; all t-tests are the same with the exception of that for Walking time which is much higher than in NL1, but the overall test of fit is also much better. However, model NL5 is still far better than model NL4 (a likelihood ratio test of equivalence is rejected at the 99.86% level). This confirms that Walking time is definitely perceived differently in reality than in the scenarios tested in the SP design. Note that although the RP-ASC for Car is negative (i.e. counterintuitive) in model NL5, its significance is null; so, we can conclude that this model reflects the best way to specify Walking time.

Two further things are worth mentioning regarding our analysis of Walking time. Firstly, the results in Table 2 refer to a specification where all ASC are RP/SP specific; however, similar results were found for different ASC specifications (Cherchi and Ortúzar, 2001). Secondly, note that our main finding (i.e. that it is better not to include Walking time in the SP Car alternative) was only revealed by the joint RP/SP specification. In fact, a model estimated using only SP data without Walking time in the Car alternative achieved better log-likelihood than the original SP model used in Table 1. However, when we calculated the ratio between the RP and SP parameters using the new SP model, we found that Walking time still showed the greater discrepancy for both data sets (the ratio θ^{RP}/θ^{SP} was 3.06; while the other ratios were almost identical to those in Table 1).

All results in Table 2 and following, were obtained using non-linear specifications of utility (Cherchi and Ortúzar, 2002); thus, almost all specifications include interactions between Time, Cost and Frequency¹⁰. The interaction terms are SP specific and the

¹⁰ We only report here models with interaction terms and a ER formulation, but many specifications were tested, such as linear utilities, utilities with a cost squared variable and also with a Wage Rate specification (Cherchi and Ortúzar, 2002).

Early/late and Car/licences variables were included only in the RP alternatives. The former was omitted from the SP alternatives because the new SP Train had varying frequency and thus the variable was no longer relevant (Cherchi and Ortúzar, 2002). The Car/licences attribute is a socio-economic (SE) variable that was not considered in the SP experiment. In fact, it only included Travel time, Frequency, Cost and Comfort; while Walking time and Transfer were treated as implicit variables (i.e. not defined in the experimental design).

Table 2. Testing generic/specific RP/SP variable specification

Attributes	NL1	NL2	NL3	NL4	NL5
Travel time PT	-0.0625 (-2.0)	-0.0731 (-2.3)	-0.0748 (-2.3)	-0.0688 (-2.2)	-0.0716 (-2.3)
Travel time Car	-0.1554 (-2.7)	-0.1946 (-3.1)	-0.1973 (-3.1)	-0.1624 (-2.7)	-0.1942 (-3.1)
Walking time(RP)	--	-0.1942 (-2.9)	-0.1951 (-2.9)	--	-0.1940 (-2.9)
Walking time (SP)	--	-0.0183 (-1.4)	--	--	-0.0366 ⁽²⁾ (-2.0)
Walking time	-0.0266 (-1.8)	--	--	-0.0524 ⁽²⁾ (-2.2)	--
Cost/g	-0.0237 (-2.3)	-0.0338 (-2.6)	-0.0343 (-2.6)	-0.0271 (-2.4)	-0.0335 (-2.6)
Frequency	0.3124 (2.4)	0.3999 (2.6)	0.4000 (2.6)	0.3588 (2.6)	0.4004 (2.6)
Comfort 1	-2.565 (-3.0)	-3.022 (-3.5)	-3.034 (-3.5)	-2.900 (-3.2)	-3.001 (-3.5)
Comfort 2	-1.252 (-2.7)	-1.500 (-3.3)	-1.507 (-3.2)	-1.418 (-2.9)	-1.490 (-3.2)
Transfer	-0.7985 (-1.9)	-1.023 (-2.2)	-0.9215 (-2.0)	-0.9744 (-2.1)	-1.086 (-2.3)
Early/Late	-0.1160 (-1.7)	-0.1610 (-1.8)	-0.1608 (-1.8)	-0.1204 (-1.6)	-0.1613 (-1.8)
Car/Licences	6.742 (2.0)	12.63 (2.4)	12.76 (2.4)	7.990 (2.2)	12.60 (2.4)
TravelTime*fare	0.00084 (2.3)	0.00118 (2.7)	0.00121 (2.7)	0.00096 (2.4)	0.00118 (2.7)
Travel Time*freq	-0.00308 (-1.0)	-0.00466 (-1.2)	-0.00459 (-1.2)	-0.00363 (-1.1)	-0.00472 (-1.2)
K_car (RP)	0.4070 (0.3)	-0.0657 (-0.0)	-0.0053 (-0.0)	0.3315 (0.2)	-0.0662 (-0.0)
K_train (RP)	-2.453 (-3.6)	-1.572 (-2.0)	-1.589 (-2.0)	-2.410 (-3.5)	-1.551 (-2.0)
K_car (SP)	0.7962 (1.1)	1.569 (1.5)	1.706 (1.6)	0.5164 (0.7)	1.252 (1.3)
K_train (SP)	-0.8762 (-2.7)	-1.120 (-2.9)	-1.251 (-3.1)	-0.8299 (-2.6)	-0.9916 (-2.8)
ϕ_1 (EMU) ⁽¹⁾	0.6141 (1.40)	0.3255 (5.52)	0.3214 (5.60)	0.5237 (2.15)	0.3266 (5.47)
ϕ_2 (SP scale factor) [...] ⁽¹⁾	0.8154 (2.8) [0.64]	0.6891 (3.4) [1.51]	0.6847 (3.3) [1.54]	0.7163 (3.1) [1.21]	0.6941 (3.4) [1.48]
Log likelihood	-743.88	-736.95	-738.10	-740.29	-735.17
$\rho^2(C)$	0.1215	0.1297	0.1284	0.1258	0.1318
Sample size	1,396	1,396	1,396	1,396	1,396

(1) t-test with respect to one; (2) Car Walking time introduced only in the RP alternatives;

Returning to the problem of fitting ASC in a RP/SP specification, many equally valid hypotheses can be advanced in principle as pointed out in section 2. Since the only different alternative in both data sets is the Train, the only *a priori*¹¹ constraint on ASC

¹¹ This is an *a priori* condition that must be checked with the statistical results from estimation.

specification is that the Train constant should not be the same for the RP and SP datasets. As usual, at least one alternative¹² has to be left as reference.

Table 3 shows the results for several nested logit (NL) models with different ASC specifications. All were estimated using the mixed RP/SP data set with generic Walking time not included in the SP Car alternative. The first three refer to different ways of specifying the ASC, but all attributes were estimated jointly:

1. Model NL6 included a generic RP/SP ASC for Car and left Bus as reference alternative (as in all models of Table 2). This specification assumes the same “real” market shares for both modes but it is possible to obtain different estimated values for seemingly the same option in both cases, due to the different sources of data as noted in section 2.
2. Model NL7 specified a generic RP/SP ASC for Car and two specific ASC, one for RP Bus and another for the SP Train service, so we assume that only Car has the same “real” market share. Thus, the only difference with model NL6 is the alternatives left as reference in the RP environment. However, as we will see later, this is not costless.
3. Model NL8 was specified under the assumption that both Car and Train were the same for the RP and SP data sets (so two constrained RP/SP ASC were estimated) leaving the Bus as reference.

First note that all the models in Table 3 can be compared (because they are restricted versions) with model NL5 in Table 2, where all the ASC are specific. Note that neither the significance of the variables nor the maximum likelihood value change with the ASC specification. This occurs because the selection of the reference alternative has usually no effect on the model other than to shift the values of the estimated constants, preserving their difference (Ben-Akiva and Lerman, 1985). But, as we will discuss later, in a joint RP/SP estimation some odd results can occur.

Comparing the three models in Table 3 with model NL5 in Table 2 from a statistical point of view, it is evident that specifying generic ASC RP/SP (model NL8) or specific for the RP and SP data (model NL5) does not have any effect on the results. If models NL5, NL6, NL7 and NL8 are compared through the likelihood ratio test, the null hypothesis of equivalence is accepted at the 60% level. Therefore we would be inclined to select model NL8 as the preferred one. Moreover, if we look at each parameter individually, all (except that for Travel time PT) are more significant in model NL8 than in the other models, and especially more significant than in model NL5 (higher t-test). This is certainly an important issue in forecasting.

¹² Note that since the two data sets are independent, the rule that we must leave one alternative as reference should apply to each environment. However, this is only true if a full set of specific ASC for RP and SP are used. As long as at least one ASC is specified as RP/SP generic, only one alternative for both data sets (i.e. one among all the RP and SP alternatives) can be left as reference as this is enough for the joint RP/SP model to be estimable. However, it is not clear how the two datasets (RP or SP) interact to estimate the ASC.

Table 3. Model estimation results: testing generic/specific RP/SP ASC

Attributes	NL6	NL7	NL8
Travel time PT	-0.0694 (-2.2)	-0.0715 (-2.3)	-0.0619 (-2.2)
Travel time Car	-0.1956 (-3.1)	-0.1946 (-3.1)	-0.2043 (-3.2)
Walking time (RP)	-0.1923 (-2.9)	-0.1957 (-3.1)	-0.2123 (-3.5)
Walking time (SP) ⁽²⁾	-0.03741 (-2.0)	-0.0365 (-2.0)	-0.0369 (-2.0)
Cost/g	-0.0341 (-2.6)	-0.0335 (-2.6)	-0.03593 (-2.7)
Frequency	0.4166 (2.7)	0.4002 (2.6)	0.4719 (3.3)
Comfort 1	-3.033 (-3.5)	-3.009 (-3.5)	-3.106 (-3.6)
Comfort 2	-1.501 (-3.2)	-1.490 (-3.2)	-1.541 (-3.3)
Transfer	-1.092 (-2.3)	-1.090 (-2.3)	-1.148 (-2.4)
Early/Late	-0.1609 (-1.8)	-0.1629 (-1.8)	-0.1920 (-2.3)
Car/Licences	11.05 (2.8)	12.95 (3.3)	12.19 (3.2)
TravelTime*fare	0.00120 (2.7)	0.00118 (2.7)	0.00127 (2.9)
Travel Time*freq	-0.00526 (-1.4)	-0.00471 (-1.2)	-0.00692 (-2.1)
K_car (RP+SP)	1.316 (1.3)	1.263 (1.3)	1.677 (1.8)
K_train (RP+SP)	--	--	-1.044 (-2.8)
K_train (RP)	-1.499 (-1.9)	--	--
K_bus (RP)	--	1.535 (2.0)	--
K_train (SP)	-0.9996 (-2.8)	-0.8219 (-2.6)	--
ϕ_1 (EMU) ⁽¹⁾	0.3234 (5.55)	0.3244 (5.67)	0.2988 (6.74)
ϕ_2 (SP scale factor) [...] ⁽¹⁾	0.6890 (3.4) [1.52]	0.6941 (3.4) [1.47]	0.6704 (3.5) [1.71]
Log-likelihood	-735.3258	-735.1743	-735.5511
$\rho^2(C)$	0.1316	0.1318	0.1467
Sample size	1,396	1,396	1,396

(1) t-test with respect to one; (2) Car Walking time introduced only in the RP alternatives

The analysis can also be helped by checking the estimated market shares depending on whether generic or specific RP/SP ASC are estimated. We first examine the results of different ASC specifications in more depth, through the estimated means of the error components differences. Table 4 reports the results for model NL5 in Table 2 and for all the specifications in Table 3.

Table 4 - Estimated means of the error component difference

	NL5	NL6	NL7	NL8
K_car - K_bus (RP)	-0.0662	1.3160	-0.2720	1.6770
K_train - K_bus (RP)	-1.5510	-1.4990	-1.5350	-1.0440
K_car - K_bus (SP)	0.8690	0.90671	-0.1888	1.1243
K_train - K_bus (SP)	-0.6883	-0.6887	-0.6867	-0.6999

First note that the models perform quite differently depending on their ASC specification; this is not surprising as the means of the error differences estimated depend on the specification adopted. However, the analysis in Table 4 enables to check if the distribution of the modal shares in the two data sets appears correct. As can be seen, only models NL6 and NL8 reproduce correctly a positive difference between Car and Bus in both the RP and SP cases, as well as a negative difference between Train and Bus, bigger in the RP case than in SP (in the SP scenarios Train was much improved). The worst case is model NL7 where an unexpected opposite result appears. Thus, it seems that if we change the base reference for the ASC specification, although the model does not change some odd results may occur. In the case of model NL5, the RP-ASC for Car has a t-test almost equal to zero, which means that the ASC value of -0.0662 is not statistically different from zero. This is still not a good result, as it is unlikely that there are no differences between Car and Bus in RP other than those explained by the attributes.

Table 5 compares the market shares predicted with the above models (excluding NL7), by showing the percentage variation between the market share simulated with the models and that predicted after implementing a simple change in a key variable. The first three cases use the ASC estimated with mixed RP/SP models whilst the last two refer to predictions made re-estimating the ASC using only RP data (Case 4) and only SP data (Case 5), and fixing the other parameters to take the values estimated in the mixed RP/SP model. In particular, in Case 1 the RP-ASC estimated with the mixed RP/SP model are used in forecasting. In Case 2 we assume that the “true” market is represented by the SP alternatives, then the SP-ASC estimated with the mixed RP/SP model are used for predictions after scaling them because they were moved into the RP domain. Case 3 is a mixed situation as we used the RP-ASC for Car and the SP scaled ASC for Train.

Table 5 - % variation in demand predictions

	Case 1			Case 2			Case 3	Case 4	Case 5	
	NL5	NL6	NL8	NL5	NL6	NL8	NL8	NL8	NL8	
bus	-0,6%	-0,6%	-0,7%	-0,7%	-0,6%	-0,6%	-0,6%	-0,4%	-0,7%	-20% Cost train
train	2,2%	2,2%	1,7%	1,7%	1,6%	1,8%	1,6%	2,8%	1,6%	
car	0,0%	0,0%	0,0%	0,0%	0,0%	0,0%	0,0%	0,0%	-0,1%	
bus	-1,5%	-1,6%	-1,9%	-1,9%	-1,7%	-1,6%	-1,8%	-1,1%	-1,8%	-50% Cost train
train	6,0%	5,9%	4,5%	4,6%	4,5%	5,0%	4,5%	7,6%	4,2%	
car	0,0%	0,0%	0,0%	-0,1%	0,0%	0,0%	0,0%	0,0%	-0,3%	
bus	-2,3%	-2,4%	-3,0%	-2,9%	-2,4%	-2,2%	-2,4%	-1,5%	-2,2%	-20% Tv train
train	9,3%	9,0%	7,1%	6,9%	6,1%	6,9%	6,1%	10,2%	5,2%	
car	0,0%	0,0%	-0,1%	-0,1%	-0,1%	0,0%	0,0%	0,0%	-0,4%	
bus	-6,1%	-6,2%	-7,7%	-7,4%	-6,1%	-5,7%	-6,1%	-3,9%	-5,6%	-50% Tv train
train	24,7%	23,7%	18,0%	17,6%	15,8%	17,9%	15,7%	27,7%	13,4%	
car	-0,1%	-0,1%	-0,2%	-0,2%	-0,2%	-0,1%	-0,1%	-0,1%	-1,1%	

First, note that the variations in demand predicted in Case 4 are always smaller than those in Cases 1-3 for Car and Bus (i.e. the unchanged alternatives) and bigger for Train; the opposite effect appears in Case 5. Note also that models NL5 and NL6 yield, in general, larger variations than model NL8. Some of these variations (for example a 25% increase in the demand for Train after a reduction of 50% in travel time) appear suspicious.

The above analyses ratifies the selection of model NL8 as the preferred specification. Note though that apart from the theoretical issues discussed previously, there are no general rules to justify the choice of one specification over another. In particular, there is no particular advantage in having generic rather than specific ASC in forecasting. However, and especially if an alternative is the same in both environments (i.e. as Car in our case), having a model where a generic RP/SP ASC specification gives no worse results than a specification with specific RP/SP ASC represents a clear practical advantage, since it allows to avoid the problem of what Car ASC to use in forecasting the Car market share.

5. CONCLUSIONS

In this paper we have tackled the problem of ASC specification in a mixed RP/SP context, and in particular the forecasting implications of using generic or specific constants for both data sets. Techniques for pooling different data sources such as these have experienced many advances since their first application 15 years ago. Now it is safe to say that they constitute the recommended, and even most used, approach to estimate modal and route choice models in the transport field. However, although many implications of mixed RP/SP estimation are well known, important issues in specification and in the process of passing from estimation to prediction do not appear to have been reported before.

Since the ASC represents the mean of the error difference and as the error components are strictly related to the systematic utility specification, the ASC values and their significance are often deflated by the degree to which each model is able to reproduce the real structure of the errors. However, as the latter is unknown fitting ASC in a joint RP/SP forecasting model is not a straightforward task.

Our theoretical and empirical analyses allowed us to draw a set of guidelines for introducing ASC into RP/SP alternatives in a forecasting model. We tested their usefulness and added more insight using real data. We also noted that not only the ASC specification should be carefully evaluated, but as long as theoretical underpinnings are satisfied, we can rely on estimation results (i.e. use the models providing the best statistical fit).

ACKNOWLEDGEMENTS

We are indebted to David Hensher and Joffre Swait for their comments which proved of valuable assistance and a source of inspiration in gaining a deeper insight into the more challenging issues discussed here. We also thank Moshe Ben-Akiva, Andrew Daly, Jordan Louviere and Kenneth Train for their comments regarding their experience in using ASC in RP/SP forecasting. Finally we wish to thank two anonymous referees for helping us to improve the paper. However, and as usual, all errors are our sole responsibility. We are grateful to FONDECYT for having supported this line of research in Chile through several projects (1020981 and 1050672).

REFERENCES

- Bates, J.J. (1998) Reflections on stated preference: theory and practice. In J. de D. Ortúzar, D.A. Hensher and S.R. Jara-Díaz (eds.), *Travel Behaviour Research: Updating the State of Play*. Pergamon, Oxford.
- Ben-Akiva, M.E. (2004). Private communication.

- Ben-Akiva, M.E. and Lerman, S.R. (1985) *Discrete Choice Analysis: Theory and Application to Travel Demand*. The MIT Press, Cambridge, Mass.
- Ben-Akiva, M.E. and Morikawa, T. (1990) Estimation of travel demand models from multiple data sources. *Proceedings 11th International Symposium on Transportation and Traffic Theory*. Yokohama, Japan.
- Bradley, M.A. and Daly, A.J. (1997) Estimation of logit choice models using mixed stated-preference and revealed-preference information. In P.R. Stopher and M.E. Lee-Gosselin (eds.), *Understanding Travel Behaviour in an Era of Change*. Pergamon, Oxford.
- Caussade, S., Ortúzar, J. de D., Rizzi, L.I. and Hensher, D.A. (2005) Assessing the influence of design dimensions on stated choice experiment estimates. *Transportation Research* 39B, 621-640.
- Cherchi, E., Meloni, I. and Ortúzar, J. de D. (2002) Policy forecasts involving new train services: application of mixed RP/SP models with interaction effects. *XII Panamerican Conference on Transport and Traffic Engineering*. November 2002, Quito, Ecuador.
- Cherchi, E. and Ortúzar, J. de D. (2001) On fitting mode specific constants in the presence of new options in RP/SP models. *Working Paper 82*, Department of Transport Engineering, Pontificia Universidad Católica de Chile.
- Cherchi, E. and Ortúzar, J. de D. (2002) Mixed RP/SP models incorporating interaction effects: modelling new suburban train services in Cagliari. *Transportation* 29, 371-395.
- Daly, A. and Rohr, C. (1998) Forecasting demand for new travel alternatives. In T. Gärling, T. Laitila and K. Westin (eds.), *Theoretical Foundations for Travel Choice Modelling*. Pergamon, Amsterdam.
- Domencich, T. and McFadden, D. (1975) *Urban Travel Demand - A Behavioural Analysis*. North Holland, Amsterdam.
- Hensher, D.A. (2002). Private communication.
- Jara-Díaz, S.R. and Farah, M. (1987) Transport demand and user's benefits with fixed income: the goods/leisure trade-off revisited. *Transportation Research* 21B, 165-170.
- Louviere, J.J. (2004). Private communication.
- Louviere, J.J., Hensher, D.A. and Swait, J.D. (2000) *Stated Choice Methods: Analysis and Application*. Cambridge University Press, Cambridge.
- Morikawa, T. (1989) *Incorporating Stated Preference Data in Travel Demand Analysis*. PhD Dissertation, Department of Civil Engineering, MIT.
- Ortúzar, J. de D. (2000) Modelling route and multimodal choices with revealed and stated preference data. In J. de D. Ortúzar (ed.), *Stated Preference Modelling Techniques*. Perspectives 4, PTRC, London.
- Ortúzar, J. de D. and Willumsen, L.G. (2001) *Modelling Transport*. Third Edition, John Wiley & Sons, Chichester.
- Train, K.E. (2004). Private communication.
- Train, K.E. and McFadden, D. (1978) The goods/leisure trade-off and disaggregate work trip mode choice models. *Transportation Research* 12, 349-353.