



Executive functions and the generation of “random” sequential responses: A computational account

Richard P. Cooper*

Centre for Cognition, Computation and Modelling, Department of Psychological Sciences, Birkbeck, University of London, United Kingdom

HIGHLIGHTS

- We present an information-processing model of random generation behaviour.
- The model is developed within a more general cognitive architecture.
- The architecture makes explicit the role of executive functions.
- The model reproduces random generation performance in single and dual task conditions.
- Dual task performance is accounted for via interference to executive control.

ARTICLE INFO

Article history:

Received 12 January 2016

Received in revised form

15 April 2016

Keywords:

Executive function

Information processing model

Random generation

Set shifting

Monitoring

Response inhibition

ABSTRACT

When asked to generate sequences of random responses, people exhibit strong and reliable biases in their behaviour. The origins of these biases have been linked to the operation of so-called executive functions through empirical studies varying, e.g., rate of production, modality of response, and (in dual task conditions) secondary task. We present a computational process model of random generation that accounts for a broad range of these empirical effects. The model, which operationalises a previous verbal account of random generation, is grounded in both the cognitive architectures and the executive functions literatures. As such, it instantiates a hypothesis concerning the interaction of multiple distinct executive functions in the generation of complex behaviour. In particular, it is argued on the basis of simulations of empirical findings that three cognitive factors play separable roles in random generation behaviour: cognitive load, which when high exacerbates underlying biases in a generation stage, monitoring, which when impaired results in greater inequality of response usage, and set-shifting, which when impaired results in less frequent switching between response schemas.

© 2016 Elsevier Inc. All rights reserved.

1. Introduction

One goal of much empirical work over the past two decades, in both cognitive psychology and cognitive neuroscience, has been to characterise the operation of the “central executive” in terms of separable functions or processes such as task-setting, working memory maintenance, response selection, monitoring and response inhibition (for reviews, see for example: Aron, Robbins, & Poldrack, 2014; Cooper, 2010; Nee et al., 2013; Vandierendonck, Liefoghe, & Verbruggen, 2010). There is now considerable empirical support for many such functions, and several computational accounts of specific executive functions have been developed

(e.g., task switching: Altmann & Gray, 2008; Brown, Reynolds, & Braver, 2007; Gilbert & Shallice, 2002; Herd, Hazy, Chatham, Brant, & Friedman, 2014; memory maintenance and updating: Ashby, Ell, Valentin, & Casale, 2005; O'Reilly & Frank, 2006; response inhibition: Band, Van Der Molen, & Logan, 2003; Boucher, Palmeri, Logan, & Schall, 2007; Wiecki & Frank, 2013). A limitation of much of this work – though one that is reasonable given the developing nature of the field – is that it generally considers the operation of individual executive functions in isolation on relatively simple executive function tasks. For example, the Boucher et al. (2007) model focuses on a version of the stop-signal task, in which a subject must, on a small proportion of trials normally signalled by a tone, withhold the production of an otherwise routine response such as a button press. Working at a similar level of complexity, Gilbert and Shallice (2002) focus specifically on the process of task set switching involved in switching from word reading to colour naming (and vice versa) in the Stroop task. While this work has advanced our

* Correspondence to: Department of Psychological Sciences, Birkbeck, University of London, Malet Street, London, WC1E 7HX, United Kingdom.

E-mail address: R.Cooper@bbk.ac.uk.

understanding of basic executive functions, behaviour on tasks of greater complexity is likely to require multiple executive functions working in concert in order to appropriately regulate behaviour.

A landmark empirical study specifically aimed at understanding the role of executive functions in more complex tasks is the individual differences study of Miyake et al. (2000). These authors considered the correlational structure of the performance of over 130 individuals on 14 tasks, comprising 9 relatively simple tasks and 5 more complex tasks. The simple tasks included three held primarily to tap the putative executive function of set-shifting, three to tap memory updating and monitoring, and three to tap response inhibition. Subject performance within each subset of three simple tasks was found to correlate more highly than between subsets, as supported by Confirmatory Factor Analysis (CFA). The three factors obtained from this analysis were then used via Structural Equation Modelling (SEM) to determine the role of the three underlying constructs in the more complex tasks. Thus, on the basis of this it was argued that, for example, in the Wisconsin Card Sorting Test the set-shifting function (and this function alone) is critical in limiting the production of perseverative errors. This was in contrast to an alternative hypothesis, namely that such errors arise from a failure in response inhibition.

One of the five complex executive tasks used by Miyake et al. (2000) was random number generation. In random generation tasks, which are widely used in executive function research (e.g., Baddeley, Emslie, Kolodny, & Duncan, 1998; Cooper, Wutke, & Davelaar, 2012; Jahanshahi, Dirnberger, Fuller, & Frith, 2000; Jahanshahi et al., 1998; Jahanshahi, Saleem, Ho, Dirnberger, & Fuller, 2006; Peters, Gieshrecht, Jellicic, & Merckelbach, 2007; Proios, Asaridou, & Brugger, 2008; Towse, 1998; Towse, Towse, Saito, Maehara, & Miyake, 2016), subjects are asked to produce a sequence of “random” responses. More precisely, subjects are required to produce responses where each successive response is independent of his or her previous responses, or equivalently, where each successive response cannot be predicted with greater than chance accuracy from the subject's previous responses. Random generation tasks typically yield multiple dissociable dependent measures (as discussed in detail below; see also Towse & Neil, 1998), and Miyake et al. argue on the basis of their SEM analysis that different dependent measures reflect the efficacy of different executive functions. This paper aims to evaluate the account of executive involvement in random generation performance given by Miyake et al. and, more generally, to explore the interaction of executive functions on a task with multiple dependent measures. We present a computational account of random generation derived from the verbal model of Baddeley et al. (1998) and grounded in a cognitive architecture in which executive functions have explicit roles. The model, which is a development of that of Sexton and Cooper (2014), is evaluated both against the results of Miyake et al. and against results from a dual task study of Cooper et al. (2012) where the primary task was random generation.

The remainder of this paper is structured as follows. We begin by providing an overview of previous findings from research on random generation tasks together with a re-analysis of the dual-task interference data from Cooper et al. (2012). We then present a computational model of performance on a random generation task. The model builds upon the verbal model of Baddeley et al. (1998), embedding it within a cognitive architecture based on the Contention Scheduling/Supervisory System approach of Norman and Shallice (1986). Critically, the Supervisory System part of the model draws also on the executive functions investigated by Miyake et al. (2000) and the decomposition of supervisory processes outlined by Shallice and Burgess (1996). Subsequent sections report simulation studies that demonstrate that the model can capture both frequently reported biases in random generation and the interference patterns arising from concurrent performance

of different secondary tasks. The general discussion focuses on two broad sets of issues: the role of key parameters of the model and their relation to the efficacy of executive functions; and the model's architecture as an elaboration of the verbal theories on which it is based, with a specific focus on the relationship between the model's architecture and other established cognitive architectures (including ACT-R: Anderson, 1993, 2007, and Soar: Laird, 2012; Newell, 1990).

2. Random generation

2.1. Dependent measures and standard effects

Random generation tasks have a long history within information processing psychology research (see Tunes, 1964, and Wagenaar, 1972, for early reviews). Behaviour is typically assessed through measures of randomness calculated from the sequence of responses produced by the subject. The degree of randomness of a sequence cannot be characterised with a single measure and many different measures have been considered. Thus, if each successive response in a sequence is equally likely but independent of the previous responses then one would expect, over the long run, that each response would occur equally often. The degree of response equality is frequently quantified in information-theoretic terms by the redundancy, or R , score:

$$R = 100 \times \left(1 - \frac{\log_2 n - \frac{1}{n} \sum_i n_i \cdot \log_2 n_i}{\log_2 a} \right) \quad (1)$$

where n is the number of responses in the sequence, a is the size of the response set, and n_i is the number of times response i is produced.

The redundancy score, which is a linear transformation of the Shannon entropy (Shannon, 1948) of the multiset of elements in a sequence, ranges from 0 to 100 and is 0 when each response is equally frequent and 100 when all responses are identical (i.e., when a single response is repeated).

A highly unpredictable sequence will have a low redundancy score, but cycling through all possible responses will similarly produce a low redundancy score. It is therefore necessary to also consider the relative frequency of response pairs (i.e., of bigrams). Again, in an ideal random sequence of infinite length each bigram should be equally frequent. Bigram equality may be quantified by calculating a redundancy score for bigrams (rather than individual responses) using Eq. (1). An alternative approach is the RNG score of Evans (1978):

$$\text{RNG} = \frac{\sum_{ij} n_{ij} \cdot \log_2 n_{ij}}{\sum_i n_i \cdot \log_2 n_i} \quad (2)$$

where n_i is the frequency of response i , and n_{ij} is the frequency of response i followed by response j .

The RNG score scales the entropy of bigrams by the entropy of elements. It attains a maximum of 1 when an element is perfectly predicted by its predecessor (i.e., bigram probabilities are zero or one). Lower values reflect greater equality of bigram usage. Note, however, that even if a sequence has perfect equality of bigram usage, it may still be predictable at the level of trigram (or higher) statistics. Further measures of randomness are therefore required. Indeed, Towse and Neil (1998) report 11 measures that have been used to quantify randomness, while Towse and Valentine (1997) consider 16.

Download English Version:

<https://daneshyari.com/en/article/6799288>

Download Persian Version:

<https://daneshyari.com/article/6799288>

[Daneshyari.com](https://daneshyari.com)