

MULTINOMIAL PROCESSING TREE MODELS OF RECENT ACCOUNTS OF
THE PROCESSING OF NONCANONICAL WORD ORDER SENTENCES

KASIM BURAK ÇAVUŞOĞLU

BOĞAZİÇİ UNIVERSITY

2022

MULTINOMIAL PROCESSING TREE MODELS OF RECENT ACCOUNTS OF
THE PROCESSING OF NONCANONICAL WORD ORDER SENTENCES

Thesis submitted to the
Institute for Graduate Studies in Social Sciences
in partial fulfillment of the requirements for the degree of

Master of Arts
in
Linguistics

by
Kasım Burak Çavuşoğlu

Bogaziçi University
2022

DECLARATION OF ORIGINALITY

I, Kasım Burak Çavuşoğlu, certify that

- I am the sole author of this thesis and that I have fully acknowledged and documented in my thesis all sources of ideas and words, including digital resources, which have been produced or published by another person or institution;
- this thesis contains no material that has been submitted or accepted for a degree or diploma in any other educational institution;
- this is a true copy of the thesis approved by my advisor and thesis committee at Boğaziçi University, including final revisions required by them.

Signature.....

Date

ABSTRACT

Multinomial Processing Tree Models of Recent Accounts of the Processing of Noncanonical Word Order Sentences

This study investigates the predictions of the retrieval and the parsing accounts of the processing of noncanonical word order sentences by implementing their suggestions in multinomial processing tree (MPT) models (Riefer & Batchelder, 1988). MPT models allow the estimation of probabilities of unobservable and hypothesized cognitive events or states using categorical data, and an analysis of the predictions and the technical properties of these models can provide information as to how accounts of cognitive processes can be improved. Recent literature has identified systematic performance decrease in the agent/patient naming task, among other types of experimental tasks, when comprehenders were faced with noncanonical word order sentences (Ferreira, 2003; Bader & Meng, 2018). The retrieval account (Bader & Meng, 2018; Meng & Bader, 2021), suggests that this decrease in performance is caused by problems with the cue-based retrieval operation triggered by the task probes, whereas the parsing account (Ferreira, 2003; Christianson, Luke & Ferreira, 2010) suggests that the decrease in performance is caused by misinterpretation. We developed multiple MPT models that reflect the assumptions of these two accounts and fitted these to the data from the first experiment of Meng and Bader (2021). We found that the retrieval account is more eligible for adaptation into an MPT structure than the parsing account, and that the parsing account needed deeper revision of its assumptions.

ÖZET

Sıradışı Kelime Sırasına Sahip Cümlelerin İşlemlemesine Dair Yakın Zamandaki Açıklamaların Çokterimli İşleme Ağacı Modelleri

Bu çalışma, sıradışı kelime sırasına sahip cümlelerin işlemlemesinin Geri Erişim ve Sözdizimsel Analiz açıklamalarının önerilerini çokterimli işleme ağacı (MPT) modellerinde (Riefer ve Batchelder, 1988) uygulayarak, bu açıklamaların tahminlerini araştırmaktadır. MPT modelleri, kategorik verileri kullanarak gözlemlenemeyen ve varsayımsal bilişsel olayların veya durumların olasılıklarının tahminine olanak sağlar ve bu modellerin tahminlerinin ve teknik özelliklerinin analizi, bilişsel süreçlere dair açıklamaların nasıl geliştirilebileceği konusunda bilgi sağlayabilir. Güncel literatürde, katılımcıların sıradışı kelime sırasına sahip cümleler ile karşı karşıya kaldıklarında, diğer deneysel görevlerin yanı sıra, yapıcı/etkilenen adlandırma görevinde sistematik performans düşüşü yaşadıkları tespit edilmiştir (Ferreira, 2003; Bader & Meng, 2018). Geri Erişim Açıklaması (Bader & Meng, 2018; Meng & Bader, 2021) performanstaki bu düşüşün, görev yoklayıcıları tarafından tetiklenen işaret-tabanlı alma işlemindeki sorunlardan kaynaklandığını öne sürerken, Sözdizimsel Analiz Açıklaması (Ferreira, 2003; Christianson, Luke & Ferreira, 2010), performanstaki düşüşün yanlış yorumlamadan kaynaklandığını öne sürüyor. Bu iki açıklamanın varsayımlarını yansıtan MPT modelleri geliştirdik ve bunları Meng ve Bader'in (2021) ilk deneyindeki verilere oturttuk. Çalışmamız, Geri Erişim Açıklaması'nın MPT yapısına Sözdizimsel Analiz Açıklaması'ndan daha uygun olduğunu ve Sözdizimsel Analiz Açıklaması'nın varsayımlarında daha derin bir revizyona ihtiyaç duyduğunu gösterdi.

ACKNOWLEDGEMENTS

First and foremost, I am extremely grateful to my supervisor, Pavel Logačev. I have always found his approach to the field of linguistics, his methods of teaching and his views on matters exceptional and inspiring. Learning from and working with him has deeply influenced my approach to science and many other aspects of my life.

Everything I did for this thesis has a lot to do with what I learned from him, and I am sure that the knowledge and skills that he has equipped me with will be of great use to me my entire life in whatever path I choose to follow.

Secondly, I would like to extend my thanks to the members of my thesis committee, Ümit Atlamaz and Umut Özge, whose input following my defense has widened my perspective about the subject and has allowed me to notice the strengths and shortcomings of my work.

Thirdly, I would like to express my gratitude towards the faculty members in Boğaziçi University, Elena Guerzoni, Ömer Demirok, Sumru Özsoy, Stephano Canalis and Nazik Dinçtopal Deniz, from whom I took the courses which developed my competence and identity as a researcher and as a student of linguistics. I feel that the experience of studying for a master's degree that they provided collectively has deeply affected and improved me as a person.

I could not withhold from also expressing my deep gratitude towards the people who have started it all by putting in their valuable time during my undergraduate years to provide extensive encouragement and support to me to pursue a master's degree, my professors from İstanbul Bilgi University, Sevdeğer Çeçen Salamin, Hande Serdar Tülüce, Fatma Nihan Ketrez Sözmen; I will be forever grateful.

I should also mention my head of department, İnci Bilgin Tekin; a faculty member of our department, Aslı Gürer; and my dearest office friend, Özge Öz; who have generously supported me, covered for me, and have very kindly forgiven and corrected my mistakes and shortcomings in our workplace due to the workload induced by this thesis. Your efforts have been pivotal to the completion of this work.

Finally, I thank the boys who have been around me since the first year of primary school; HML, who have been around since my high school years; my band, Minor Illusion, with whom I have been playing for three years now; my stalwart neighbours, the KD; and my friends from the Department of Linguistics at Boğaziçi University for their emotional support throughout the creation of this work.

TABLE OF CONTENTS

CHAPTER 1: INTRODUCTION	1
CHAPTER 2: LITERATURE REVIEW	9
2.1 The parsing account of noncanonical sentence processing	9
2.2 The retrieval account of noncanonical sentence processing.....	17
CHAPTER 3: MULTINOMIAL PROCESSING TREE MODELS	33
CHAPTER 4: MPT MODELS UNDER THE RETRIEVAL ACCOUNT.....	46
4.1 MPT models of plausibility judgment task responses	46
4.2 Joint MPT models of agent/patient naming and plausibility judgement task responses	65
CHAPTER 5: MPT MODELS UNDER THE PARSING ACCOUNT.....	92
5.1 Joint MPT models of agent/patient naming and plausibility judgement task responses	92
CHAPTER 6: GENERAL DISCUSSION	134
CHAPTER 7: CONCLUSION.....	150
REFERENCES.....	152

LIST OF TABLES

Table 1. Levels of the Structure Factor along with their Examples from Ferreira's (2003) Experiments.....	9
Table 2. Levels of the Meaning Factor along with their Examples in all of Ferreira's (2003) Experiments.....	11
Table 3. Agent and Patient Question Examples from Ferreira (2003).....	12
Table 4. Factor Levels of Structure, Meaning and Noun Order along with their Examples from Bader and Meng's (2018) Experiment	21
Table 5. Cue-Match Status when the Task is to Name the Agent of a Sentence for each Cue Posited by Bader and Meng (2018) for some of the Experimental Conditions in their Experiment, along with the Corresponding Percentage of Correct Responses	22
Table 6. Mean Percentages of Correct Plausibility Judgments and Correct Agent/Patient Namings for each Condition in the First Experiment of Meng and Bader (2021).....	26
Table 7. The Conditions and the Corresponding Example Sentence from the Experiment of Cutter, Paterson and Filik (2021).....	31
Table 8. Experimental Conditions and the Corresponding Example Sentences from the First and the Second Experiments of Logacev and Dokudan (2021).....	44
Table 9. Probability Equations Resulting from the Model in Figure 4.....	47
Table 10. Probability Equations Resulting from the Model in	50
Table 11. Estimated Values for each of the Free Parameters when the Model in	51
Table 12. Results from the Plausibility Judgment Task in the Data and the Predictions of the Model in.....	52

Table 13. Estimated Values for each of the Free Parameters when the Model in Figure 7 is Fit to the Data.....	58
Table 14. Results from the Plausibility Judgment Task in the Data and the Predictions of the Model in.....	58
Table 15. New Fixed Parameter Values for the model in Figure 5.....	61
Table 16. Estimated Values for each of the Free Parameters when the Model in Figure 5 with the Fixed Parameter Values in Table 15 is Fit to the Data	61
Table 17. Comparison of Probability Equations that Return the Probability of a 'plausible' Response for the Two Models.....	63
Table 18. Results from the Plausibility Judgment Task in the Data and the Predictions of the Model in.....	64
Table 19. Cue-Match Status of Items from all Conditions in the data. when the Probe Requests the Retrieval of the Agent.....	70
Table 20. Cue-Match Status of Items from all Conditions in the data. when the Probe Requests the Retrieval of the Patient.	71
Table 21. Estimated Values for each of the Free Parameters when the Model in Figure 10 with the Ratios Listed in is Fit to the Data	72
Table 22. Cue-Match Status of Items from all Conditions in the data. as they are Coded for our Second Version of the Model in Figure 10 when the Probe Requests the Retrieval of an Agent.....	78
Table 23. Estimated Values for each of the Free Parameters when the Model in Figure 10 with the Function in (16) is Fit to the Data.....	80
Table 24. Estimated Values for each of the Free Parameters when the Model in Figure 10 with the Function in (17) is Fit to the Data.....	85

Table 25. The Items from the First Experiment of Meng and Bader (2021) from the NO1 Nonreversible or Biased SO Conditions and NO1 Nonreversible or Biased OS and Passive Conditions	98
Table 26. The Items from the First Experiment of Meng and Bader (2021) from the NO1 Symmetrical SO, OS and Passive Conditions.....	102
Table 27. The Items from the First Experiment of Meng and Bader (2021) from the NO2 Nonreversible or Biased SO Conditions	104
Table 28. The Items from the First Experiment of Meng and Bader (2021) from the NO2 Nonreversible or Biased OS and Passive Conditions.....	104
Table 29. The Probability Functions which Return the Probability of a ‘plausible’ Response to the Plausibility Judgment Task and those which Return the Probability of a Correct Response to Agent/Patient Naming for each of the Conditions from the First Experiment of Meng and Bader (2021) According to the First Version of our Model of their Data	105
Table 30. Estimated Values for each of the Free Parameters when the Models in Figures 17-21 with the Functions in Table 29 are Fit to the Data.....	109
Table 31. The Probability Functions which Return the Probability of a ‘plausible’ Response to the Plausibility Judgment Task and those which return the Probability of a Correct Response to Agent/Patient Naming for each of the Conditions from the First Experiment of Meng and Bader (2021) According to the Second Version of our Model of their Data	113
Table 32. Estimated Values for each of the Free Parameters when the Models in Figures 17-21 with the Functions in Table 31 are Fit to the Data.....	114
Table 33. The Probabilities of Correct Agent/Patient Naming and a ‘plausible’ Response to the Plausibility Judgment Task for each Condition from the First	

Experiment of Meng and Bader (2021) as Predicted by the First and the Second Versions of the Parsing Account Model, Referred to as Model 1 and Model 2, Respectively.....	117
Table 34. The Probability Functions which Return the Probability of a ‘Plausible’ Response to the Plausibility Judgment Task and those which Return the Probability of a Correct Response to Agent/Patient Naming for Each of the Conditions from the First Experiment of Meng and Bader (2021) According to the Third Version of our Model of their Data.....	119
Table 35. Estimated Values for each of the Free Parameters when the Models in Figures 17-21 with the Functions in Table 34 are Fit to the Data.....	120
Table 36. The Probability Functions which Return the Probability of a ‘plausible’ Response to the Plausibility Judgment Task and those which Return the Probability of a Correct Response to Agent/Patient Naming for each of the Conditions from the First Experiment of Meng and Bader (2021) According to the Fourth Version of our Model of their Data	126
Table 37. Estimated Values for each of the Free Parameters when the Models in Figures 17-21 with the Functions in Table 36 are Fit to the Data.....	129

LIST OF FIGURES

- Figure 1. Results from the three experiments of Ferreira (2003). Y-axis shows the percentage of correct responses to agent/patient naming, and x-axis shows the structure levels. The meaning levels are represented by shapes and colors. Version 1 is always plausible, and Version 2 is always implausible except for symmetrical sentences where both versions are plausible 13
- Figure 2. Example multinomial processing tree models of the plausibility judgment task used in the second experiment of Bader and Meng (2018). The model at the top models the responses to an implausible item, while the one at the bottom models the responses to a plausible item. 36
- Figure 3. The plots on the left show the probability of a correct response to the plausibility judgment task for sentences as predicted by the model in Figure 2. The plots on the right show the percentage of ‘plausible’ responses in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points. 40
- Figure 4. Our first model of plausibility judgment responses from the first experiment of Meng and Bader (2021) 46
- Figure 5. Our second model of plausibility judgment responses from the first experiment of Meng and Bader (2021). 49
- Figure 6. The plot on the left shows the percentage of ‘plausible’ responses to NO2 sentences as predicted by the model in Figure 5. The plot on the right shows the percentage of ‘plausible’ responses to NO2 sentences in the data. The

percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis.....	53
Figure 7. Our third model of plausibility judgment responses from the first experiment of Meng and Bader (2021)	56
Figure 8. The plot on the left shows the percentage of ‘plausible’ responses as predicted by the model in Figure 7. The plot on the right shows the percentage of ‘plausible’ responses in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points	57
Figure 9. The plot on the left shows the percentage of ‘plausible’ responses as predicted by the model in Figure 5 with the parameter values shown in (18) and (19). The plot on the right shows the percentage of ‘plausible’ responses in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points	63
Figure 10. Our first model of the agent/patient naming task responses in Meng and Bader (2021), which is built over our final model of the plausibility judgment task responses. ‘Correct A./P. N.’ stands for correct agent/patient naming, and the empty set symbol represents either incorrect agent/patient naming or the lack of a response to the agent/patient naming task.....	68
Figure 11. The plot on the left shows the percentage of correct responses to agent naming as predicted by the model in Figure 10 with the parameter values shown in Table 21 and the cue-match ratios shown in Table 19. The plot on	

the right shows the percentage of correct responses to agent naming in the data. The percentage of correct responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 73

Figure 12. The plot on the left shows the percentage of correct responses to patient naming as predicted by the model in Figure 10 with the parameter values shown in Table 21 and the cue-match ratios shown in Table 12. The plot on the right shows the percentage of correct responses to patient naming in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 75

Figure 13. The plot on the left shows the percentage of ‘plausible’ responses as predicted by the model in Figure 10 with the parameter values shown in Table 21. The plot on the right shows the percentage of ‘plausible’ responses in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 76

Figure 14. The plot on the left shows the percentage of correct responses to agent naming as predicted by the model in Figure 10 with the parameter values shown in Table 23. The plot on the right shows the percentage of correct responses to agent naming in the data. The percentage of correct responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on

the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 83

Figure 15. The plot on the left shows the percentage of correct responses to agent naming as predicted by the model in Figure 10 with the parameter values shown in Table 24. The plot on the right shows the percentage of correct responses to agent naming in the data. The percentage of correct responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 86

Figure 16. The plot on the left shows the percentage of ‘plausible’ responses as predicted by the model in Figure 10 with the parameter values shown in Table 24. The plot on the right shows the percentage of ‘plausible’ responses in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 89

Figure 17. Our model of the agent/patient naming and plausibility task responses to the NO1 Nonreversible and Biased SO conditions in Meng and Bader (2021). ‘Correct A./P. N.’ stands for correct agent/patient naming, and the empty set symbol represents either incorrect agent/patient naming or the lack of a response to the agent/patient naming task..... 99

Figure 18. Our model of the agent/patient naming and plausibility task responses to the NO1 Nonreversible and Biased, OS or Passive conditions in Meng and Bader (2021). ‘Correct A./P. N.’ stands for correct agent/patient naming, and

the empty set symbol represents either incorrect agent/patient naming or the lack of a response to the agent/patient naming task.....	100
Figure 19. Our models of the agent/patient naming and plausibility task responses to the symmetrical SO conditions and the symmetrical OS and passive conditions in both noun orders in Meng and Bader (2021). ‘Correct A./P. N.’ stands for correct agent/patient naming, and the empty set symbol represents either incorrect agent/patient naming or the lack of a response to the agent/patient naming task.....	101
Figure 20. Our model of the agent/patient naming and plausibility task responses to the NO2 Nonreversible and Biased SO conditions in Meng and Bader (2021). ‘Correct A./P. N.’ stands for correct agent/patient naming, and the empty set symbol represents either incorrect agent/patient naming or the lack of a response to the agent/patient naming task.....	103
Figure 21. Our model of the agent/patient naming and plausibility task responses to the NO2 Nonreversible and Biased, OS or Passive conditions in Meng and Bader (2021). ‘Correct A./P. N.’ stands for correct agent/patient naming, and the empty set symbol represents either incorrect agent/patient naming or the lack of a response to the agent/patient naming task.....	105
Figure 22. The plot on the left shows the percentage of correct responses to agent/patient naming as predicted by the models in Figures 17-21, with the parameter values shown in Table 30. The plot on the right shows the percentage of correct responses to agent/patient naming in the data. The percentage of correct responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points.....	106

Figure 23. The plot on the left shows the percentage of ‘plausible’ responses to plausibility judgment as predicted by the models in Figures 17-21, with the parameter values shown in Table 30. The plot on the right shows the percentage of ‘plausible’ responses to the plausibility judgment task in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 108

Figure 24. The plot on the left shows the percentage of ‘plausible’ responses to plausibility judgment as predicted by the models in Figures 17-21, with the probability functions shown in Table 31, and the parameter values shown in Table 32. The plot on the right shows the percentage of ‘plausible’ responses to the plausibility judgment task in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 114

Figure 25. The plot on the left shows the percentage of correct responses to agent/patient naming as predicted by the models in Figures 17-21, with the probability functions shown in Table 31, and the parameter values shown in Table 32. The plot on the right shows the percentage of correct responses to agent/patient naming in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 115

Figure 26. The plot on the left shows the percentage of ‘plausible’ responses to plausibility judgment as predicted by the models in Figures 17-21, with the probability functions shown in Table 34, and the parameter values shown in Table 35. The plot on the right shows the percentage of ‘plausible’ responses to the plausibility judgment task in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 121

Figure 27. The plot on the left shows the percentage of correct responses to agent/patient naming as predicted by the models in Figures 17-21, with the probability functions shown in Table 34, and the parameter values shown in Table 35. The plot on the right shows the percentage of correct responses to agent/patient naming in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 123

Figure 28. The plot on the left shows the percentage of ‘plausible’ responses to plausibility judgment as predicted by the models in Figures 17-21, with the probability functions shown in Table 19, and the parameter values shown in Table 35. The plot on the right shows the percentage of ‘plausible’ responses to the plausibility judgment task in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 127

Figure 29. The plot on the left shows the percentage of correct responses to agent/patient naming as predicted by the models in Figures 17-21, with the probability functions shown in Table 36, and the parameter values shown in Table 35. The plot on the right shows the percentage of correct responses to agent/patient naming in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points 130

CHAPTER 1

INTRODUCTION

Recent research on sentence processing has revealed systematic performance effects in certain tasks for comprehenders when they are confronted with unambiguous sentences with noncanonical word order. The noncanonicity of word order referred to here is essentially a reversal in the linear position of the subject and the object of a sentence, with regard to their typical positions in a language. For example, sentences in passive structure and object-cleft sentences in English are considered examples of sentences with noncanonical word order. As opposed to active sentences in English, passive and object-cleft sentences place the semantic object of the verb in a linear position that comes before that of the semantic subject of the verb. Recent research has focused on the thematic roles that the subject and object of a sentence carries and utilized designs that would reveal the understanding of the comprehender of unambiguous noncanonical word order sentences by eliciting their thematic role assignments (Ferreira, 2003; Christianson, Luke & Ferreira, 2010; Bader & Meng, 2018; Meng & Bader, 2021; Cutter, Paterson & Filik, 2021). The tasks that elicited thematic role assignments of the comprehenders in these studies have showed that the comprehenders make errors in identifying thematic roles or perform in a way that is consistent with an assignment of thematic roles that does not match the linguistic input, more often when the task required the comprehension of a noncanonical word order sentence. These findings gave rise to suggestions about how the human parsing mechanism (HPM) forms the interpretation of a sentence, that were categorized into two accounts in the recent sentence processing literature (Bader and Meng, 2018; Meng and Bader, 2021; Cutter, Paterson & Filik, 2021). Using the terminology of

Meng and Bader (2021), we will refer to these two accounts as the parsing account and the retrieval account.

The parsing account of noncanonical sentence processing originated from the study of Ferreira (2003) who developed a task that required the participants to name the ‘do-er’ of an action, or the object which was ‘acted-on’, through comprehension questions that were presented after a participant had read a sentence in the canonical or the noncanonical word order, in both plausible and implausible forms (agent/patient naming task). By comparing the participant performance in this task in four experiments, each contrasting a different set of canonical and noncanonical structures, Ferreira found that the participants made more errors when naming the ‘do-er’ of an action, the agent of the sentence, or the object which was ‘acted-on’, the patient of the sentence, when the sentence they read had noncanonical word order, as opposed to canonical word order, and when the sentence they read was implausible. Ferreira (2003) considered the occasional errors in naming the agent and the patient of the noncanonical word order sentences found in the study as reflecting erroneous interpretation, or misinterpretation, of these sentences and suggested a distinction between the cognitive processes that lead to the interpretation of a sentence that a comprehender ends up with adopting as their final understanding of the sentence’s meaning. According to Ferreira (2003), although most of the time comprehenders interpret sentences correctly through syntactic algorithms which construct detailed sentence representations, it is also sometimes the case that comprehenders resort to algorithms which construct sentence representations faster and more efficiently by making use of frequently found syntactic facts about a language, called heuristics. Ferreira (2003) suggested two heuristics that help the comprehenders in forming interpretations of sentences quickly and efficiently. One of these heuristics is the

NVN (noun-verb-noun) heuristic, which takes advantage of the frequently found word order in terms of phrase categories in English to quickly form a syntactic representation of a sentence by assigning the first noun encountered by the comprehender while reading a sentence to the agent thematic role, and the second noun encountered to the patient thematic role. The second heuristic on the other hand, completely ignores the syntactic structure of a sentence that was read by the comprehender and forms a syntactic representation of the sentence solely based on the lexical properties of the words by assigning the thematic roles of agent and patient based on these properties so that a plausible relationship between the nouns and the verb encountered in the sentence is obtained. Ferreira (2003) suggests that these heuristics work alongside the algorithms that produce accurate and detailed syntactic representations while a sentence is being read, and the interpretation that the comprehender ends up with is the result of some combination of the influences of these processes.

Christianson, Luke and Ferreira (2010) found further evidence of misinterpretation errors in a similar study that made use of the same passive and active items as in Ferreira (2003) but using a slightly different task. The task used in this study was to answer a comprehension question presented in the form of a logical statement by responding either 'yes' or 'no'. For example, for the sentence 'The angler caught the fish.', the comprehension question was 'Catcher = Angler?'. Christianson, Luke and Ferreira, using this task, elicited the thematic role assignment of the comprehender just like in Ferreira (2003), and found that comprehenders performed worse in the task when the sentence was in the passive form, as opposed to the active form, and when the sentence was implausible. Christianson, Luke and Ferreira considered these results to support the claim that heuristics like the NVN

heuristic or the heuristic that made plausible connections between the nouns and the verb in a sentence to form interpretations as suggested by Ferreira (2003) occasionally were the cognitive processes that determined the interpretation that the comprehenders ended up with.

Karimi and Ferreira (2016) considered the claim that the HPM uses both fast and efficient heuristics and algorithms that produce accurate and detailed syntactic representations is in support of the more general theory of sentence processing called Good-Enough Processing (Christianson et al., 2001; F. Ferreira, 2003; F. Ferreira, Ferraro, & Bailey, 2002; F. Ferreira & Patson, 2007; Sanford & Sturt, 2002), and integrated this idea into the theory by further developing these claims. According to Karimi and Ferreira (2016), the HPM is a dual-route mechanism where an algorithmic route and a heuristic route work simultaneously to form interpretations. The algorithmic route uses the algorithms that produce accurate and detailed syntactic representations, while the heuristic route, makes use of the fast and efficient heuristics to form syntactic representations. They further suggested that the result of this dual-route mechanism in the form of an interpretation is highly sensitive to task demands, which determine when during parsing the HPM enters a state of satisfaction with the output, called equilibrium. Under this model of the HPM, the heuristic route starts working on an interpretation before the algorithmic route, and only if equilibrium was not reached through the heuristic route the algorithmic route starts working on an interpretation, while the heuristic route still continues to influence the algorithmic route. Thus, under the account of Karimi and Ferreira (2016) interpretations that resulted from this procedure had influence of both the heuristic and the algorithmic route.

The retrieval account of the performance decrease in tasks that elicited thematic role assignments of the comprehenders when they were confronted with noncanonical word order sentences, was developed by Bader and Meng (2018, 2021). Bader and Meng (2018) using a design similar to that of Ferreira (2003), studied the processing of German sentences in noncanonical word order. Bader and Meng (2018) used both the agent/patient naming task that was used in Ferreira (2003), and a plausibility judgment task to investigate the participant performance in these tasks when the items were in canonical or noncanonical word order. Bader and Meng (2018), just like Ferreira (2003) also had both plausible and implausible versions of these sentences in their study. Their choice of including the plausibility judgment task in a second experiment was based on the assumption that the agent/patient naming task requires a cue-based retrieval operation (Bader & Meng, 2018; see Van Dyke & Johns, 2012, for a review of cue-based retrieval), whereas the plausibility judgment task does not, although both tasks require the correct assignment of the thematic roles of agent and patient in case of either implausible or noncanonical word order sentences in order to be correctly responded to. They found the same effects as in Ferreira (2003) in their experiment which featured the agent/patient naming task, but a different pattern of results as well as different effects of sentence structure and sentence plausibility when the task was to judge the plausibility of a sentence as in their second experiment. Based on these findings, Bader and Meng (2018) suggested that the performance decrease in tasks that elicited thematic role assignments of the comprehenders was not a result of misinterpretation, as it was suggested under the parsing account, but was a result of erroneous retrieval of information about the sentences, which occurred more often when the sentence was either implausible or in the noncanonical word order. Bader and Meng (2018)

proposed four retrieval cues that are defined based on the typical properties of an agent and suggested that the probability of successful retrieval of the agent in the agent/patient naming task is a function of the number of retrieval cues that match the target noun.

However, Bader and Meng's (2018) results could also be explained by the parsing account, in that, under the parsing account, the HPM is highly sensitive to the task demands (Karimi & Ferreira, 2016), and so it could have been that the plausibility judgment task caused systematical changes in the influences that the two processing routes had on the sentence interpretation that the comprehender ends up with, when compared to those caused by the agent/patient naming task. Therefore, Meng and Bader (2021) conducted two experiments in which participants had to complete both the agent/patient naming and the plausibility judgment tasks sequentially after reading a sentence. The two experiments had the same design and tasks as their previous study (Bader & Meng, 2018), and in their first experiment, the participants first judged the plausibility of a sentence and then named the agent or the patient, but in their second experiment, the order of the tasks was reversed. Both experiments yielded similar results to Bader and Meng (2018) for both the agent/patient naming and plausibility judgment tasks, leading Meng and Bader (2021) to argue that the performance decrease in implausible and noncanonical word order sentences could not have been caused by misinterpretations, due to the fact that the participants had to complete both tasks after reading a sentence which yielded a different pattern of results and so it could not be that the participants responded to each task using a different interpretation.

Essentially, the two accounts of noncanonical word order sentence processing differ in where they place the assumed cause of the effects in the time course of the

events that take place in the HPM throughout the process of reading the sentence and responding to the task. The parsing account assumes that there are at least two parallel processes that work with the linguistic input every time a sentence is read, and the product of which process ultimately becomes the interpretation of the sentence is highly sensitive to task demands. Therefore, since what determines the comprehenders task performance has to do with which process ended up providing the interpretation for the sentence in the parsing account, it is possible to say that, for this account, the factors that influence task performance come into play before or at the point when an interpretation of the sentence is obtained. The retrieval account on the other hand, attributes the cause of the effects to processes that come into play after an interpretation is obtained and the task is encountered. The retrieval account, just like the parsing account, also assumes that performance is task dependent, but here, the task has no influence on the interpretation but directly affects the processes related to the retrieval of the information obtained from the linguistic input.

Our study seeks to test the assumptions of these two accounts and to reveal their predictions by making use of multiple multinomial processing tree models (MPTs) (Riefer & Batchelder, 1988; Riefer & Batchelder, 1999; Erdfelder et. al., 2009). Defining accounts of behavioral data as MPT models allows us to estimate probabilities for the unobservable and hypothesized cognitive events or states suggested under the account, and also to generate the predictions of a series of hypothesized cognitive events or states so that these can be compared to the actual data from the experiment to see whether the model is a good account of the data or not. MPT models can also point us into a direction as to how to improve an account so that it can capture more data because the technical reasons why an MPT model cannot account for certain effects in a set of data is usually very clear. In order to

assess the performance of all of our MPT models, we have used the data from the first experiment of Meng and Bader (2021).

Therefore, in order to clarify the predictions that the suggestions under each of the two accounts make, and to explore in what ways the two accounts can be improved or modified, we have created multiple MPT models which follow the suggestions of the two accounts as closely as possible in an MPT model structure and presented and discussed these in this thesis. In the second chapter of this thesis the literature that focused on the processing of noncanonical word order sentences is reviewed, while the third chapter explains the MPT modelling procedure as well as how to assess the predictions of MPT models in detail. The fourth chapter focuses on our MPT models of the retrieval account of noncanonical word order sentence processing, and the fifth chapter focuses on our MPT models of the parsing account. Finally, in chapter six, a summary of our findings from the modelling procedures of the two accounts will be presented, as well as a discussion of the two accounts' ability to capture the data from the first experiment of Meng and Bader (2021).

CHAPTER 2

LITERATURE REVIEW

2.1 The parsing account of noncanonical sentence processing

The parsing account has its roots at the oft-cited study by Ferreira (2003). Ferreira conducted three experiments to compare the processing of English unambiguous sentences with canonical and noncanonical word order. Canonicity of word order referred to in here was determined by two factors, the frequency of the structure and the linear order of the thematic role assignment in the sentence. The frequency of the structure was also considered because, in English, structures where the thematic role of agent belongs to the first noun in terms of linear order are more commonly used than those where the thematic role of agent belongs to the second noun, and any observed effects when the linear order of the thematic role assignment is manipulated, may also have emerged due to the frequency of the structure.

The first experiment compared actives and passives as in Table 1, which differ in terms of both factors. In active sentences in English, the thematic role of agent is assigned before that of patient, and for passives, the order is reversed. In addition, active sentences are more frequently used than passive sentences.

Table 1. Levels of the Structure Factor along with their Examples from Ferreira's (2003) Experiments

Active	The mouse ate the cheese.
Passive	The cheese was eaten by the mouse.
Subject-cleft	It was the mouse who ate the cheese.
Object-cleft	It was the cheese the mouse ate.

The second experiment compared subject-clefts and passives as in Table 1, which differ only in the order of thematic role assignment procedures, as they are

both infrequently used structures. The third experiment compared subject-clefts with object-clefts as in Table 1, which are both infrequent structures that also differed in the order of thematic role assignment.

In addition to the comparison of structures, Ferreira also manipulated the plausibility of the sentences as shown in Table 2. Passive and active sentences each had a nonreversible, biased, and symmetrical version. In the nonreversible versions, one of the two nouns which were assigned a thematic role was animate and the other inanimate. Thus, each sentence of a particular structure in the nonreversible condition could either have an animate agent and an inanimate patient, which resulted in a plausible meaning, or it could have an inanimate agent and an animate patient, which resulted in an implausible meaning. In the biased versions, both nouns were animate, but one was more compatible with the agent role due to general world-knowledge when the verb was considered, and thus resulted in a more plausible meaning than the other. In the symmetrical versions, both nouns were animate, and the verb was selected such that both nouns were compatible with agent and patient roles in terms of general world-knowledge, which resulted in both order of nouns being equally plausible.

In each trial, for all three experiments, participants listened to a sentence and named either the ‘do-er’ of the action (Agent) or the object (Patient) that was acted-on depending on the probe (Agent/patient naming). Filler probes were also used, which asked for either the color of an object, a location mentioned in the sentence, the action, or the time in which the action took place.

The experiments showed that participants made more errors in naming the agent and the patient of the sentences in trials with passive sentences, where the agent followed the patient, compared to those with active sentences or subject-clefts,

where the patient followed the agent. Additionally, participants also made more errors in trials with implausible sentences overall. Moreover, the comparison of subject-clefts and passives, and that of subject-clefts and object-clefts yielded the same results, and thus showed that the errors were not due to the frequency of the structure.

Table 2. Levels of the Meaning Factor along with their Examples in all of Ferreira's (2003) Experiments

Sentence Type	Example
Nonreversible, plausible	The mouse ate the cheese.
Nonreversible, implausible	The cheese ate the mouse.
Biased reversible, plausible	The dog bit the man.
Biased reversible, implausible	The man bit the dog.
Symmetrical, one order	The woman visited the man.
Symmetrical, other order	The man visited the woman.

Based on the findings from the three experiments shown in Figure 1, Ferreira (2003) suggested that HPM forms sentence interpretations by making use of both simple heuristics and syntactic algorithms. What is referred to here by syntactic algorithms is essentially a sentence processing mechanism that works meticulously to produce sentence interpretations that are detailed and faithful to the linguistic input. Heuristics, on the other hand, are essentially shortcuts that favor efficiency over accuracy, which assume that the linguistic input follows what occurs most frequently in a language. Ferreira suggests two heuristics based on the findings that are specific to a language and it is possible for these two heuristics to produce conflicting interpretations of the same sentence. One of these is the NVN (noun-verb-noun) heuristic, which maps the first noun encountered by HPM to the agent thematic role and the second noun to the patient thematic role. In English, this heuristic wrongly interprets the first noun in passives and object-clefts as the agent of the verb. The second heuristic suggested by Ferreira forms a semantic connection

between the noun phrases encountered by the HPM that ignores syntactic positions related to thematic roles to create an interpretation. For an English sentence that features an inanimate noun in the agent thematic role position, such as in the non-reversible conditions in Ferreira's experiment, this heuristic produces the wrong interpretation. According to Ferreira (2003), these two heuristics have a strong impact on what final interpretation the HPM ends up with, but it is unclear how a conflict between the results of the two heuristics would be resolved. Moreover, to Ferreira (2003) the results of the experiment are evidence in favor of a view of HPM where the NVN heuristic has a stronger impact on what the final interpretation will be than the other heuristic, particularly the fact that the participants made more errors while naming the agent and the patient in plausible passive sentences than implausible active sentences.

Christianson, Luke and Ferreira (2010) in a similar study used the same plausible and implausible active and passive sentences as Ferreira (2003). However, the experimental procedure, including the tasks, were different in this study. Participants listened to the sentences to later answer a comprehension question which was comprised of a logical statement about the agent, or the patient of the action mentioned in the sentence, as shown in Table 3. Participants responded either 'yes' or 'no' to the question. Immediately following the response, the participants were shown a drawing which either depicted a plausible or realistic scene, or an implausible or clearly fictitious scene, and the participants were prompted to orally describe the picture.

Table 3. Agent and Patient Question Examples from Ferreira (2003)

Sentence	Agent Question	Patient Question
The angler caught the fish.	Catcher = Angler?	Catchee = Fish?

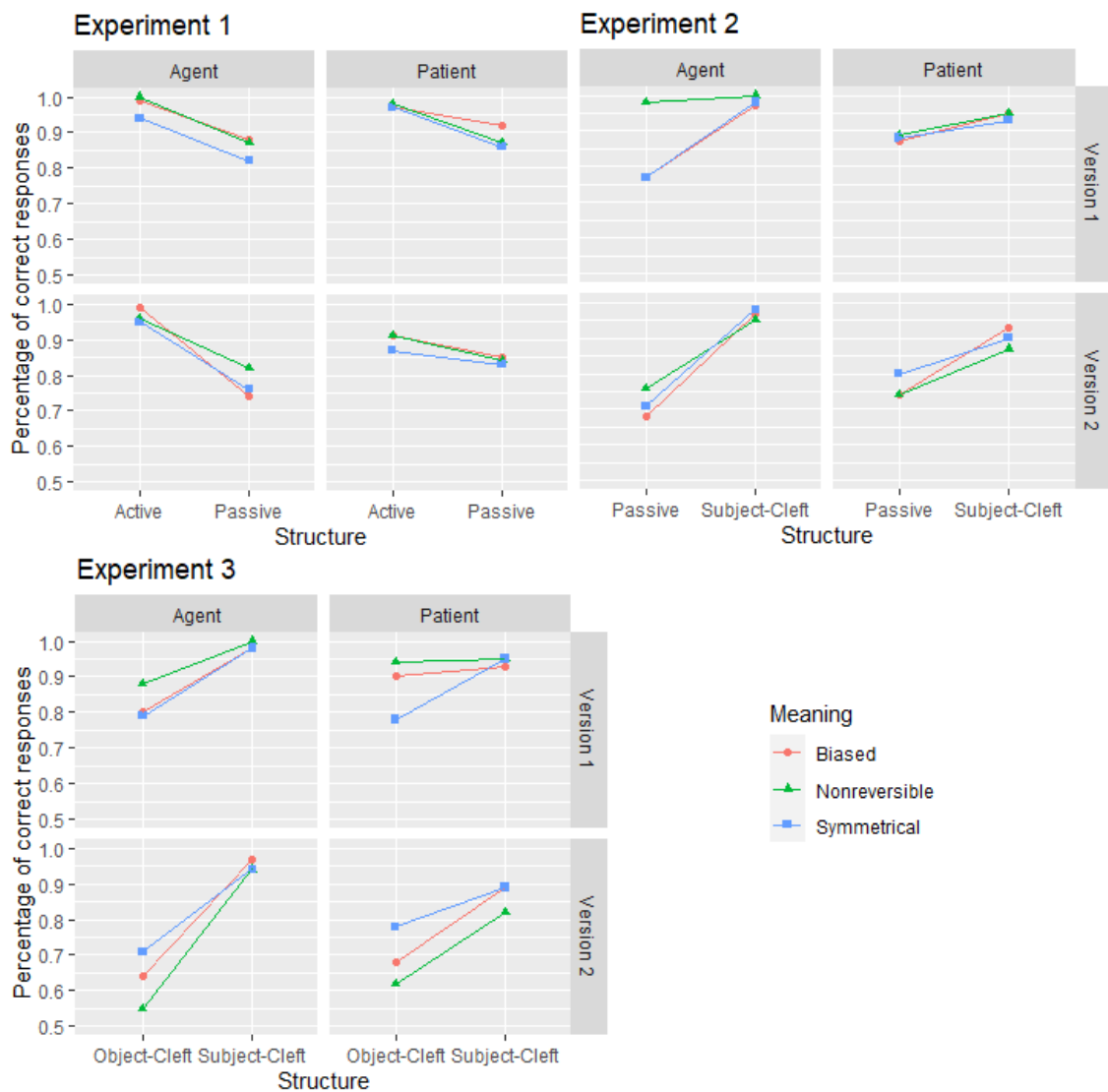


Figure 1. Results from the three experiments of Ferreira (2003). Y-axis shows the percentage of correct responses to agent/patient naming, and x-axis shows the structure levels. The meaning levels are represented by shapes and colors. Version 1 is always plausible, and Version 2 is always implausible except for symmetrical sentences where both versions are plausible

The results of the study showed an (i) effect of voice, whereby passive sentences were responded to significantly less accurately than active sentences, (ii) an effect of plausibility, whereby plausible sentences were responded to more accurately than implausible sentences, and (iii) an effect of question type, whereby agent questions were answered more accurately than patient questions. In addition, an interaction between question type and voice was found, in that accuracy was the

highest for agent questions of active sentences, and it was the lowest for patient questions of passive sentences. Moreover, the participants' oral descriptions of the pictures shown after the comprehension questions were such that pictures of implausible scenes following active sentences lead to more descriptions in the passive voice, while pictures of plausible scenes following passive sentences lead to more descriptions in the passive voice.

Based on the findings of the comprehension question task, Christianson, Luke and Ferreira (2010) conclude that the results are further evidence in favor of a language processing model where two different processing routes work in parallel, one of which is more accurate but less efficient than the other and produces an interpretation faithful to the linguistic input, while the other uses simple heuristics as described in Ferreira (2003) to quickly arrive at an interpretation that will be inaccurate when the sentence structure is non-canonical or the event described is implausible. In addition, they suggest the finding that the pictures describing implausible events following active sentences leading to more descriptions in the passive voice shows that the participants tried to make the depicted event more believable by changing the structure of the sentence while preserving the order of the nouns, and thus switching their thematic roles. They further suggest that this supports the claim that there is a mechanism at play during sentence processing, which produces interpretations that are inconsistent with the linguistic input but are consistent with the world-knowledge of the comprehender.

The parsing account of task performance effects found in agent/patient naming follows from the more general language processing theory of Good-Enough Processing (Christianson et al., 2001; F. Ferreira, 2003; F. Ferreira, Ferraro, & Bailey, 2002; F. Ferreira & Patson, 2007; Sanford & Sturt, 2002). This theory

suggests that the HPM is strongly task-dependent and versatile, in that it produces representations of linguistic input that are detailed to the level that is required by the task at hand. The detail mentioned here refers to both the completeness and the accuracy of the representation.

For example, Swets, Desmet, Clifton, and Ferreira (2008), investigated the comprehension of and the reading times in sentences that feature relative clause attachment ambiguity. They used three types of sentences where the attachment of the relative clause was either fully ambiguous, disambiguated to attach to the first noun encountered, or disambiguated to attach to the second noun encountered. Moreover, they used three types of comprehension question conditions in which the participants were either always asked a question about the relative clause, or always asked a superficial question, or only occasionally asked a superficial question. The results of the experiment showed that reading times were modulated by the type of questions asked, in that superficial questions lead to faster reading times for ambiguous sentences while a different pattern of reading times was found for these sentences when the questions were about the relative clause attachment. In addition, the participants took longer to respond to the attachment questions of ambiguous sentences than disambiguated sentences. Based on these findings, Swets et al. concluded that it might be sometimes the case that the relative clause was not attached to any nouns when both nouns are suitable candidates for attachment when a sentence is read but is only attached after it is found out by the participant that the task requires it to be attached. An interpretation such as the one mentioned here, where the relative clause is not attached to any noun, is an example of an incomplete syntactic representation under the theory of Good-Enough Processing.

On the other hand, it is also possible under the theory of Good-Enough Processing that an interpretation of a sentence is downright inaccurate. For example, the lower accuracy in naming the agent and the patient in passive and object-cleft sentences along with implausible sentences in Ferreira (2003), and the lower accuracy in answering comprehension questions in passive and implausible sentences in Christianson, Luke and Ferreira (2010) is attributed to the inaccuracy of the interpretations, or misinterpretation.

The task-dependency of language processing as explained by Good-Enough Processing is the result of a dual-route mechanism that is engaged any time a sentence is read (F. Ferreira, 2003; Swets, Desmet, Clifton, & Ferreira, 2008; Christianson, Luke & Ferreira, 2010; Karimi & Ferreira, 2016). The two routes that are at work are called the heuristic route and the algorithmic route. The heuristic route works with semantic information to rapidly create an interpretation using simple heuristics such as the NVN heuristic, which maps the first noun encountered to the agent role and the second to the patient role, or a heuristic which forms interpretations solely based on the plausibility of the semantic relations of the nouns encountered to the verb¹, completely ignoring the syntactic structure of the sentence. The algorithmic route on the other hand, works with syntactic information and always forms the correct interpretation provided that the sentence was perceived correctly. Both of these routes are engaged simultaneously when a sentence is read, but the heuristic route is quicker to produce an interpretation, while the algorithmic route, which works with more complex syntactic information, takes longer to produce an interpretation. Problems like misinterpretation or incomplete sentence representations are a result of failure in the integration of the information produced

¹ Semantic-association heuristic henceforth.

through the two routes. Karimi & Ferreira (2016) explain this failure of integration of information from the two routes with the Online Cognitive Equilibrium Hypothesis. Under this hypothesis, a criterion for equilibrium is set depending on the task-at-hand by the comprehender. Because the heuristic route is faster to produce an interpretation, if the criterion is low enough, the HPM ends up with a misinterpretation or an incomplete parse when confronted with sentences where the structure conflicts with the assumptions of the NVN or the semantic-association heuristic. However, when the task demands the criterion for equilibrium to be higher, more time is allocated for processing and thus it is also possible for an interpretation to be produced through the algorithmic route, which in turn, lowers the chance that the misinterpretation or the incomplete parse surfaces.

As the parsing account of noncanonical sentence processing follows from the theory of Good-Enough Processing (F. Ferreira, 2003; Swets, Desmet, Clifton, & Ferreira, 2008; Christianson, Luke & Ferreira, 2010; Karimi & Ferreira, 2016), and the online cognitive equilibrium hypothesis further details the suggestions of this theory (Karimi & Ferreira, 2016), the implementation of the parsing account in this thesis will include the Online Cognitive Equilibrium Hypothesis as a part of the parsing account.

2.2 The retrieval account of noncanonical sentence processing

The retrieval account was suggested by Bader and Meng (2018, 2021) as an alternative account of the same performance effects observed in the studies of noncanonical sentence processing that led to the development of the parsing account. Bader and Meng (2018) in their first experiment used an agent/patient naming task similar to that which was used by Ferreira (2003) to compare task performance for

simple active sentences (SO sentences), object-initial active sentences (OS sentences), and passive sentences in German as shown in Table 4. For both OS and passive sentences in German, the noun with the thematic role of patient comes first, whereas for SO sentences, the agent is encountered first. They also manipulated the plausibility of the sentences similarly by having nonreversible, biased, and symmetrical versions of sentences of each structure. Just as with Ferreira's study, one of the nouns in the nonreversible versions of the sentences was animate, and the other was inanimate. Thus, when the order of the nouns was reversed, the sentence became implausible. Moreover, in the biased versions, both nouns were animate, but one was more compatible with the agent role; and in the symmetrical versions, both nouns were animate and equally compatible with the agent role with regard to general world-knowledge.

In addition, in German, nominative and accusative case are overtly marked in NPs with masculine nouns. As can be seen in the examples in Table 3, when such an NP is marked with nominative, the determiner is 'der', whereas when it is marked with accusative, the determiner is 'den'. This can potentially affect agent/patient naming task performance by allowing the comprehender to distinguish between the two nouns more easily, making the correct thematic role assignment more probable (Bader & Meng, 2018).

Just as with Ferreira's study (2003), the results of the experiment showed that participants made more errors in naming the agent and the patient in sentences where either the linear order of the agent and patient thematic roles was reversed (Passive and OS sentences), or if the sentence was implausible (Noun Order 2). This pattern of errors is incompatible with an account that suggests a facilitative effect of German overt case marking on agent/patient naming task performance because such an

account would predict more accurate naming for OS sentences than passive sentences, due to the fact that in OS sentences, the object NP, and so the patient noun, is marked with accusative, whereas in passive sentences, the patient noun is marked with nominative (Bader & Meng, 2018).

Both the parsing and the retrieval account are in-line with the results of this experiment. The parsing account predicts that for passive and OS sentences, comprehenders will sometimes settle with the interpretation that results from the NVN heuristic or the semantic-association heuristic, both of which make wrong predictions for these types of sentences with regard to thematic roles as described before. The retrieval account, on the other hand, predicts that the errors are caused by difficulty of the cue-based retrieval process triggered by the agent/patient naming task instead of a misinterpretation (Bader & Meng, 2018; see Van Dyke & Johns, 2012, for a review of cue-based retrieval). When the agent or the patient probe is encountered, retrieval cues that encode the typical properties of an agent or a patient are matched against the items in memory, and the item with the highest match score, also depending on the weights of the cues, has the highest probability of being successfully retrieved². The typical properties of an agent or a patient mentioned here include semantic properties such as animacy, as well as the syntactic properties of a noun that can be assigned the thematic role of an agent or a patient. Bader and Meng (2018), develop their analysis of the results of the first experiment by assuming the four cues: plausibility, position, function, and category. The first cue, ‘Plausibility’, describes whether a noun is a ‘plausible’ ‘do-er’ of an action or not. For example, for the verb ‘cook’, the plausibility cue will match with an animate noun. The other three

² Under the cue-based retrieval theory of sentence comprehension, decay also affects the probability of retrieval for an item (Lewis, Vasishth, & Van Dyke, 2006); however, this property of the theory is not included in the analysis of Bader and Meng (2018, 2021) for the processing of non-canonical word order sentences which we call the retrieval account here.

cues describe syntactic properties of a noun. The ‘Position’ cue will match with a noun that is in the typical agent position, the ‘Category’ cue will match with a noun phrase only, and the ‘Function’ cue will match with the noun that syntactically functions as the agent of the sentence. For example, for a passive sentence, the ‘Function’ cue matches with the semantic object of the sentence, because under this account, this noun is assumed to be syntactically functioning as the agent of the sentence.

Bader and Meng (2018) provide the information about the cue-match status of some of the conditions in their experiment and relate these to the percentage of correct responses obtained for each condition as a result of the agent naming task. Table 5 shows the cue-match status for these conditions. They suggest that the accuracy in a condition is closely related to the number of retrieval cues matching the target noun. As can be seen in Table 4, for the SO and OS conditions, this relation is fairly apparent. However, for the passive conditions, accuracy is higher than expected under this analysis despite the low number of cues that match the target noun. To address this issue, Bader and Meng suggest a special status for the Category cue, pointing out that the visibility of the preposition ‘by’ could have allowed this cue to be more distinctive. Their suggestion regarding the category cue can also be understood in terms of cue-weights under the cue-based retrieval theory of sentence comprehension, that is, it can be said that the Category cue is weighted more heavily than the other three cues, which increases the contribution of a match with the target noun to the probability of its retrieval (Van Dyke, J. A., & Lewis, 2003; Lewis, Vasishth, & Van Dyke, 2006; Van Dyke & Johns, 2012).

Table 4. Factor Levels of Structure, Meaning and Noun Order along with their Examples from Bader and Meng's (2018) Experiment

Structure	Meaning	Noun Order	Example Sentence	English Translation
SO	Nonreversible	1	Der Koch hat den Topf gereinigt.	'The chef cleaned the pan'
		2	Der Topf hat den Koch gereinigt.	'The pan cleaned the chef'
	Biased	1	Der Koch hat den Braten ruiniert.	'The chef ruined the roast'
		2	Der Braten hat den Koch ruiniert.	'The roast ruined the chef'
	Symmetrical	1	Der Vater hat den Onkel umarmt.	'The father hugged the uncle'
		2	Der Onkel hat den Vater umarmt.	'The uncle hugged the father'
OS	Nonreversible	1	Den Topf hat der Koch gereinigt.	'The pan, the chef cleaned.'
		2	Den Koch hat der Topf gereinigt.	'The chef, the pan cleaned.'
	Biased	1	Den Braten hat der Koch ruiniert.	'The roast, the chef ruined'
		2	Den Koch hat der Braten ruiniert.	'The chef, the roast ruined'
	Symmetrical	1	Den Onkel hat der Vater umarmt.	'The uncle, the father hugged.'
		2	Den Vater hat der Onkel umarmt.	'The father, the uncle hugged.'
Passive	Nonreversible	1	Der Topf wurde vom Koch gereinigt.	'The pan, by the chef, was cleaned.'
		2	Der Koch wurde vom Topf gereinigt.	'The chef, by the pan, was cleaned.'
	Biased	1	Der Braten wurde vom Koch ruiniert.	'The roast, by the chef, was ruined.'
		2	Der Koch wurde vom Braten ruiniert.	'The chef, by the roast, was ruined.'
	Symmetrical	1	Der Onkel wurde vom Vater umarmt.	'The uncle, by the father, was hugged.'
		2	Der Vater wurde vom Onkel umarmt.	'The father, by the uncle, was hugged.'

In the second experiment, in which the same materials and design was used, Bader and Meng (2018) used a plausibility judgment task instead of the agent/patient naming task. Their reasoning for this decision was that the processes induced by the

agent/patient naming task were different from those induced by the plausibility judgment task. Both the agent/patient naming task and the plausibility judgment task required correct identification of the agent and patient, because in order for the participants to correctly determine the plausibility of a sentence in the nonreversible and biased versions, they need to check if there is an animate noun that is compatible with the verb with regard to general world-knowledge at the agent position. However, the two tasks differed in terms of the processes induced by them under the retrieval account in that the plausibility judgment task did not require the successful retrieval of the agent and the patient, and also the verb, after the sentence was read, but instead only required the retrieval of the information that the sentence was plausible or not (e.g., Isberner & Richter, 2013).

Table 5. Cue-Match Status when the Task is to Name the Agent of a Sentence for each Cue Posited by Bader and Meng (2018) for some of the Experimental Conditions in their Experiment, along with the Corresponding Percentage of Correct Responses

Experimental Condition (Agent Probe)	Syntactic Cues				Accuracy
	Plausibility	Position	Function	Category	
Nonreversible-SO-NO1	target	target	target	-	97%
Nonreversible-PS-NO1	target	competitor	competitor	target	94%
Nonreversible-OS-NO1	target	competitor	target	-	89%
Symmetrical-SO-NO1	-	target	target	-	94%
Symmetrical-PS-NO1	-	competitor	competitor	target	76%
Symmetrical-OS-NO1	-	competitor	target	-	62%
Nonreversible-SO-NO2	competitor	target	target	-	85%
Nonreversible-PS-NO2	competitor	competitor	competitor	target	77%
Nonreversible-OS-NO2	competitor	competitor	target	-	54%

The results of the experiment showed that participants did not make more errors in judging the plausibility of OS and passive sentences where the order of the agent and patient is reversed compared to SO sentences in the nonreversible and biased conditions, but only did so in the symmetrical conditions, where both noun orders resulted in plausible meaning. Bader and Meng reported that the results were in conflict with the assumptions of the parsing account in that if an incorrect interpretation of the sentence was obtained due to the NVN strategy heuristic, then the plausibility judgment accuracy should have suffered for OS and passive sentences in the plausible conditions, unless the task did not induce the use of such a heuristic. Therefore, in order for the parsing account to explain the results of the second experiment, it must be the case that the plausibility judgment task does not trigger the use of the NVN heuristic suggested by the parsing account.

In addition, Bader and Meng (2018), prior to the experiments mentioned above, collected offline plausibility ratings from different participants for the same sentences used in the experiments. The absolute plausibility ratings obtained from the offline measure closely matched the percentage of ‘plausible’ responses to each condition in their second experiment in which an online plausibility judgment task was used, and the relative plausibility ratings obtained from the offline measure strongly correlated with accuracy in their first experiment where the agent/patient naming task was used. The relative plausibility rating refers to the plausibility rating of a plausible sentence, or Noun Order 1 (NO1) sentence, relative to its implausible, or Noun Order 2 (NO2), counterpart. A correlation in this regard shows that more errors were made in the agent/patient naming task when the experimental item was less plausible than its counterpart with reversed noun order. Bader and Meng

consider this finding additional evidence of the two tasks inducing different cognitive processes as such a correlation was only found for the agent/patient naming task.

Since both the retrieval account and the parsing account predict task-specific differences, the fact that there was a different pattern of errors in the first experiment of Bader and Meng (2018), where an agent/patient naming task was used, and their second experiment, where a plausibility judgment task was used, cannot be considered evidence that can distinguish between the two accounts. To address this issue, Meng and Bader (2021) conducted two experiments where the same materials and design was used as Bader and Meng (2018), but in these experiments, after each sentence was read, participants had to both judge plausibility and name either the agent or the patient of the sentence. This paradigm is assumed to be eligible to distinguish between the two accounts because the parsing account assumes that the communicative goal, or in this case, the task, determines the extent of the influence that the heuristic and the algorithmic route have on the final interpretation, hence the effect of the task must come into play while the comprehender is reading the sentence under this account; whereas, under the retrieval account, because it is the task that determines whether cue-based retrieval is engaged or not, the effect of the task comes into play when the task probe is encountered. For example, under the parsing account, since the interpretation that the comprehender ends up with must be the same for both tasks, if the interpretation is plausible, then the response to an agent probe must be the animate noun in the nonreversible condition, and a ‘plausible’ response must be obtained for the plausibility judgment task. On the other hand, under the retrieval account, a failure in retrieval of the animate noun in the nonreversible condition can cause the agent probe to elicit the wrong noun while a ‘plausible’ response can still be obtained for the plausibility judgment task. This

distinction in the assumptions of the two accounts allows the paradigm to accomplish distinguishing between them.

In the first experiment of Meng and Bader (2021), the participants first judged the plausibility of the sentence and then orally named either the agent or the patient of the same sentence. Table 6 shows the results of the experiment. For the plausibility judgment task, the correct response to all NO2 conditions was ‘implausible’, hence the overall low accuracy in plausibility judgments for biased sentences in NO2 conditions reflected the moderately plausible status of these sentences. More importantly, OS sentences were judged with less accuracy than SO and passive sentences across all meaning and noun order levels. The agent/patient naming task, on the other hand, showed effects of noun order and structure, replicating the findings of Ferreira (2003) and the first experiment of Bader and Meng (2018). There were more errors in naming the agent and the patient in all NO2 conditions except the symmetrical condition where both orders were plausible, which shows that plausibility was a significant factor in determining agent/patient naming accuracy. Moreover, agent/patient naming was the most accurate in SO conditions, less accurate in passive conditions except for the nonreversible versions, and even less accurate in OS conditions overall.

As was discussed before, the parsing account predicts a close relationship between the results for both tasks because the same interpretation must be used for both of the tasks, and it is assumed that any errors in both tasks are caused by misinterpretation. Further analysis of the results by Meng and Bader (2021) revealed that it was not the case, at least for the nonreversible sentences, that a correct judgment of plausibility lead to correct naming of agent and the patient. However, the analysis also revealed that incorrect plausibility judgments lead to more

agent/patient naming errors. While the former finding conflicts with the assumptions of the parsing account, the latter finding can be explained by both accounts. Incorrect plausibility judgments reflect misinterpretations under the parsing account, and it is in-line with the parsing account that the misinterpretations also lead to incorrect naming of the agent and the patient. The retrieval account, on the other hand, explains this in terms of comprehension failure: when the HPM fails to construct an interpretation that is in-line with the linguistic input, the comprehender fails to accurately complete both tasks.

Table 6. Mean Percentages of Correct Plausibility Judgments and Correct Agent/Patient Namings for each Condition in the First Experiment of Meng and Bader (2021)

Experimental Condition	Plausibility Judgment Accuracy	Agent Naming Accuracy	Patient Naming Accuracy
Nonreversible-SO-NO1	0.97	0.97	0.90
Nonreversible-PS-NO1	0.95	0.94	0.91
Nonreversible-OS-NO1	0.87	0.89	0.84
Biased-SO-NO1	0.93	0.95	0.96
Biased-PS-NO1	0.93	0.87	0.93
Biased-OS-NO1	0.77	0.74	0.77
Symmetrical-SO-NO1	0.92	0.94	0.94
Symmetrical-PS-NO1	0.91	0.76	0.87
Symmetrical-OS-NO1	0.82	0.62	0.60
Nonreversible-SO-NO2	0.93	0.85	0.77
Nonreversible-PS-NO2	0.95	0.77	0.88
Nonreversible-OS-NO2	0.77	0.54	0.60
Biased-SO-NO2	0.61	0.91	0.91
Biased-PS-NO2	0.62	0.76	0.77
Biased-OS-NO2	0.42	0.61	0.56
Symmetrical-SO-NO2	0.94	0.90	0.89
Symmetrical-PS-NO2	0.89	0.81	0.86
Symmetrical-OS-NO2	0.83	0.61	0.63

Moreover, the fact that OS sentences were judged with less accuracy than SO and passive sentences across all meaning and noun order levels, conflicts with the suggestion of the parsing account that the errors are due to misinterpretation caused

by the NVN heuristic or the semantic-association heuristic. This is because for the symmetrical sentences, a change in the order of the two nouns does not yield a difference in plausibility but the participants still made more errors for these sentences. This should not be so under the parsing account because even when the NVN heuristic categorizes the first noun encountered by the HPM as the agent of sentence or the semantic-association heuristic categorizes the noun that is more plausible as an agent as the agent of the sentence, for symmetrical sentences, since both nouns are plausible agents, the response under the parsing account should be ‘plausible’, but it was not.

On the other hand, to account for the low accuracy in plausibility judgments for OS sentences, Meng and Bader (2021) suggest that the decrease in performance for these sentences is due to their information-structural markedness in the lack of an appropriate discourse context. In order to test their suggestion, Bader and Meng conducted additional analyses including only the SO and OS sentences in the nonreversible and symmetrical conditions and found that there was only an effect of structure, and no interaction between structure and meaning or a three-way interaction between structure, meaning and noun order in the plausibility judgment data; whereas from the agent/patient naming data, an interaction between structure and meaning as well as a three-way interaction between structure, meaning and noun order was obtained. Therefore, Bader and Meng concluded that the performance decrease for the OS sentences in the agent/patient naming task and the plausibility judgment task had different sources.

Meng and Bader (2021) conducted a second experiment which used the same design and materials as their first experiment with only the SO and OS structure levels and the nonreversible meaning level. They chose to focus on these levels of

meaning and structure because the biased sentences yielded results similar to the nonreversible sentences, and the passive sentences yielded results similar to the SO sentences. Moreover, in their second experiment, participants named the agent or the patient first, and judged the plausibility of the sentence second, instead of the other way around as with their first experiment.

The second experiment, just like the first, showed significant main effects of structure and noun order for the plausibility judgment task, with the SO sentences being judged more accurately than OS sentences, and the NO1 sentences, which were plausible, being judged more accurately than the NO2 sentences, which were implausible. For the agent/patient naming task, effects of structure and noun order were found as well as an interaction between them. Agent/patient naming was more accurate for SO sentences than for OS sentences, and more accurate for NO1 sentences than for NO2 sentences. Moreover, OS sentences in NO2 yielded significantly lower accuracy in agent/patient naming than OS sentences in NO1, but this pattern was not apparent in SO sentences.

Taken together, the findings from the two experiments of Meng and Bader (2021) support the view that HPM creates sentence representations that are faithful to the linguistic input regardless of the task at hand, and the task-specific differences are related to the processes that the task induces, which operate on an accurate representation of the sentence. The participants, after each sentence was read, completed both the plausibility judgment task, and the agent/patient naming task, and different patterns of accuracy results were obtained for the two tasks, regardless of their order of completion. Structure, plausibility and meaning effects were found for both the plausibility judgment task and the agent/patient naming task, but interactions between the three factors were only found for the agent/patient naming

task. Moreover, it was often not the case that correct plausibility judgments lead to correct agent/patient naming for nonreversible sentences, supporting the view that errors cannot be attributed to misinterpretation alone as this should cause the comprehender to fail in both tasks.

Cutter, Paterson and Filik (2021), in another study focusing on the processing of noncanonical word order sentences, used self-paced reading to uncover the effects of word order Canonicity on the reading of a follow-up sentence. Implausible English sentences in the canonical word order (Subject-cleft) or in the noncanonical word order (Object-cleft) as in Table 7 were presented with either a follow-up sentence that was consistent in terms of meaning with the actual meaning of the first sentence (Algorithmically Consistent) or was consistent with a plausible version of the sentence where the subject and object nouns switched places (Good-Enough Consistent). The terminology used to represent the conditions regarding the properties of the follow-up sentence is from the parsing account, with the term ‘Algorithmically Consistent’ referring to the processing route that works with syntactic information and always forms the correct interpretation provided that the sentence was perceived correctly, and the term ‘Good-Enough Consistent’ referring to the processing route that uses the NVN heuristic or the semantic-association heuristic to derive an interpretation of the sentence as suggested by the parsing account. Cutter, Paterson and Filik argue that this paradigm allows the assessment of how implausible noncanonical word order sentences are processed through a more naturalistic task than the agent/patient naming and the plausibility judgment tasks used in the previous research of the same phenomenon (Ferreira, 2003; Christianson, Luke & Ferreira, 2010; Bader & Meng 2018; Meng & Bader, 2021).

Under the parsing account, it is expected that in some portion of the trials in the noncanonical first sentence conditions, reading times will be longer in the penultimate and final sentence regions if the follow-up is Algorithmically Consistent, due to the use of the NVN or the semantic-association heuristic (Cutter, Paterson & Filik, 2021). This is because the NVN heuristic will categorize the first noun encountered as the agent of the sentence, ‘king’ in the example in Table 7, or the semantic-association heuristic will do the same because it is the more plausible agent considering the verb ‘execute’, and so the meaning of the interpretation formed for the first sentence will conflict with the meaning of the Algorithmically Consistent follow-up, which states something impossible given the heuristically formed interpretation of the first sentence. However, if slow-downs only occur at the penultimate and final sentence regions of the Good-Enough Consistent follow-up conditions, regardless of the canonicity of the first sentence, then misinterpretation of the noncanonical word order sentences due to use of heuristics could not have been the case. Nevertheless, such a result can be explained by the retrieval account if it is assumed that the reading of a follow-up sentence does not induce the use of the same retrieval processes as the agent/patient naming task used in the previous studies because under this account, it is assumed that the HPM always computes a representation of the sentence that is faithful to the linguistic input.

The analysis of Cutter, Paterson and Filik (2021) of the reading times from the penultimate region of the sentences showed evidence against both accounts in that when the first sentence was in the canonical order, Good-Enough Consistent follow-ups were read faster, and when the first sentence was in the noncanonical order, Algorithmically Consistent follow-ups were read faster. However, reading times from the final region of the sentences showed evidence against the parsing

account in that an effect of follow-up type, with the Good-Enough Consistent follow-ups showing longer reading times, was found but no interaction between the canonicity of the sentence and the follow-up type was found.

Table 7. The Conditions and the Corresponding Example Sentence from the Experiment of Cutter, Paterson and Filik (2021)

Condition	Example Sentence
Canonical first sentence, Algorithmically Consistent follow-up sentence	It was the peasant that executed the king. Afterwards, the peasant rode back to the countryside.
Noncanonical first sentence, Algorithmically Consistent follow-up sentence	It was the king that was executed by the peasant. Afterwards, the peasant rode back to the countryside.
Canonical first sentence, Good-Enough Consistent follow-up sentence	It was the peasant that executed the king. Afterwards, the king rode back to his castle.
Noncanonical first sentence, Good-Enough Consistent follow-up sentence	It was the king that was executed by the peasant. Afterwards, the king rode back to his castle.

Cutter, Paterson and Filik (2021) suggest that the results are against the parsing account and that they can be explained by the retrieval account. Firstly, they argue that the use of the semantic-association heuristic, which maps the lexical items into the agent and the patient thematic roles according to the verb used in the sentence with no regard of the syntactic structure, should have led to a speed-up in the reading of the Good-Enough Consistent sentences but such an effect was not found. Secondly, they argue that the task-dependency of the HPM under the parsing account could not be argued to explain the speed-up for the Algorithmically Consistent sentences regardless of sentence canonicity because reading a follow-up sentence could not have motivated the participants to adopt a criterion for equilibrium (Karimi & Ferreira, 2016) which would lead to more interpretations resulting from the algorithmic processing route than naming the agent or the patient of a sentence. In other words, they argue that the task of naming the agent or the

patient of a sentence is more likely to elicit interpretations from the algorithmic route than reading a follow-up sentence. On the other hand, they argue that the retrieval account can explain the results, suggesting that the reading task that they used, because it is more naturalistic than the agent/patient naming task, may have provided the participants with more retrieval cues so that the thematic relationships were easier to retrieve for the participants.

As the retrieval account of noncanonical sentence processing follows from the cue-based retrieval theory of sentence comprehension (Van Dyke, J. A., & Lewis, 2003; Lewis, Vasishth, & Van Dyke, 2006; Van Dyke & Johns, 2012), the implementation of the retrieval account in this thesis will include the suggestions under this theory as part of the retrieval account.

CHAPTER 3

MULTINOMIAL PROCESSING TREE MODELS

The aim of this thesis is to create multiple multinomial processing tree (MPT) models of the first experiment of Meng and Bader (2021) that are as faithful as possible to the suggestions of the parsing and the retrieval account regarding the processing of noncanonical word order sentences, and to explore the predictions of the two accounts if they were expressed as multinomial processing tree models. This chapter aims to explain what MPT models are, the steps of creating MPT models of data, and why this practice can be helpful in understanding the phenomena of concern in this thesis.

MPT modelling is a method of experimental data analysis that allows estimation of probabilities of unobservable and hypothesized cognitive events or states using categorical data (Riefer & Batchelder, 1988; Riefer & Batchelder, 1999; Erdfelder et. al., 2009). MPT modelling has been widely used in the field of cognitive psychology to model behavioral data (see Erdfelder et. al., 2009, for a review), including studies of sentence processing (Logacev & Dokudan, 2021; Paape, Avetisyan, Lago, & Vasishth, 2021). Each MPT model is specific to the experimental paradigm and data that it models. This is because, in an MPT model, the frequency of the specific outcomes possible in an experimental paradigm are thought to occur as a result of a specific set of cognitive events or states, where each outcome must have a unique set of cognitive events or states that lead to it. In other words, an MPT model models the experimental paradigm that it is tied to and so should not be considered a generalized language processing model. The hypothesized cognitive events or states featured in an MPT model, although their suggestion must

be theoretically justified, are specific to the experimental task that they model, hence they provide probabilities of the outcomes from the experiment and the cognitive events or states that lead to those outcomes in that experiment only. Moreover, as the MPT models are fit to the data from the experiment that they model, the probabilities of the outcomes and the hypothesized cognitive events or states are contingent on it. However, these probabilities can still help in improving generalized language processing models by showing how much of the data in question can be explained by certain combinations and configurations of cognitive events or states assumed in a generalized language processing model.

In order to model experimental data using this technique, it must be possible to organize every possible outcome of each trial from the experiment into a finite number of discrete and observable categories (Riefer & Batchelder, 1988). For example, the possible outcomes of the plausibility judgment task used in Bader and Meng's (2018) second experiment, can be conceptualized to be limited to a 'plausible' response and an 'implausible' response³. Once an exhaustive list of the possible outcome categories is obtained, any number of latent cognitive events or states can be postulated to occur on the way to the outcomes, provided that for each state or event, there is only one state or event that is hypothesized to complement it such that there can occur no other state or event at that stage of processing. This is because in an MPT model, if the probability of an event or state occurring is ' p ', then the probability of its complement occurring must be ' $1 - p$ '.

Figure 2 shows the visual representation of a possible MPT model of the plausibility judgment task used in the first experiment of Bader and Meng (2018)

³ The lack of a response is also a possible outcome but since no such event is recorded in the data, the event or outcome of no response need not be featured in the list of all possible outcomes from the experimental trials.

where the observable outcomes are marked by sharp corners and the hypothesized cognitive states that lead to these outcomes are marked by rounded corners. In Figure 1, there are two models, one for each type of item, discriminated by their plausibility. It can be seen from the visual representation that the cognitive state of ‘comprehension’ and ‘guessing’ complement each other, such that the occurrence of one implies that the other did not. The same is true for the cognitive states ‘guess plausible’ and ‘guess implausible’. This relationship between the complementary cognitive states that belong to the same stage of an MPT model is reflected by their probabilities of occurrence. In the visual representation in Figure 1, the probability of a cognitive state occurring is denoted by the symbol or equation on the line that connects it to the previous cognitive state to its left. For example, the probability of ‘comprehension’ occurring is represented by ‘ c ’ and the probability of ‘guessing’ occurring is represented by ‘ $1-c$ ’. In a similar fashion, the probability of being in a state of guessing ‘plausible’ is ‘ g_p ’, and the probability of being in a state of guessing ‘implausible’ is ‘ $1-g_p$ ’.

In an MPT model, the probability of each unique outcome is the product of the probabilities of the cognitive events or states that lead to it, and if it is possible in the MPT model for both of the two cognitive events or states that complement each other to result in the same outcome, then the probability of that outcome is the sum of the two product equations, each featuring one or the other of the complementary cognitive events or states. For example, for the model of an implausible item in Figure 2, an ‘implausible’ response can both be obtained following a state of ‘comprehension’ or ‘guessing’, whereas a ‘plausible’ response, can only be obtained following a state of ‘guessing’. Since an ‘implausible’ response can be obtained as a result of both of the two complementary cognitive states, the probability of an

‘implausible’ response according to this model is the sum of two product equations⁴ as shown in (1-3), while the equation that returns the probability of a ‘plausible’ response features no summation since only one of the two complementary cognitive states lead to it, as shown in (2-4).

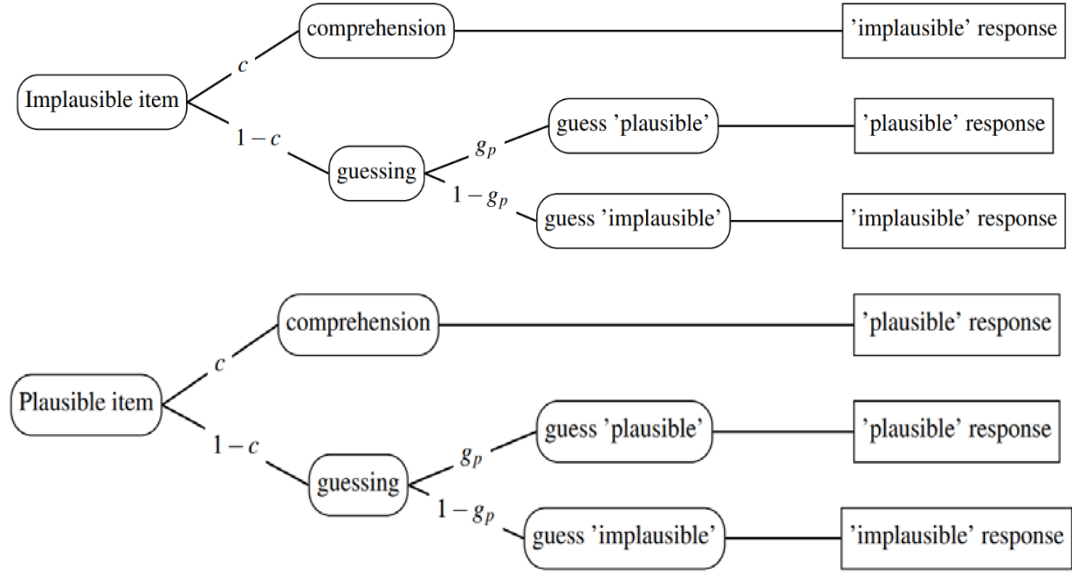


Figure 2. Example multinomial processing tree models of the plausibility judgment task used in the second experiment of Bader and Meng (2018). The model at the top models the responses to an implausible item, while the one at the bottom models the responses to a plausible item.

- (1) $\Pr(\text{'implausible' response} \mid \text{Implausible item}) = c + (1-c) \times (1 - g_p)$
- (2) $\Pr(\text{'plausible' response} \mid \text{Implausible item}) = (1-c) \times g_p$
- (3) $\Pr(\text{'plausible' response} \mid \text{Plausible item}) = c + (1-c) \times g_p$
- (4) $\Pr(\text{'implausible' response} \mid \text{Plausible item}) = (1-c) \times (1 - g_p)$

As such, the two models shown in Figure 2, express assumptions about the processes that lead to each of the two outcomes possible in the plausibility judgment task. To illustrate the implications that even such a simple model may have, we can

⁴ Because no other cognitive state was postulated in the MPT model in Figure 1 following ‘comprehension’, the product equation of this path features only the parameter c .

consider its assumptions separately. Firstly, both models assume that guessing does not happen in the case of comprehension. While this may seem trivial to express at first, this assumption, for example, could be ignoring possible events such as forgetting the information about the sentence which would allow the comprehender to respond to the plausibility judgment accurately as the plausibility judgment probe is encountered. Secondly, the model assumes that the comprehension of the item always leads to a correct response, as for an implausible item, the comprehension path leads to an ‘implausible’ response and for a plausible item, it leads to a ‘plausible’ response. This assumption of the model is also non-trivial in that it calls for a definition of comprehension that does not allow, for example, shallow processing or underspecified interpretations (Ferreira, 2003; Christianson, 2016; Karimi & Ferreira, 2016). Thus, it is an essential part of creating an MPT model to consider its assumptions and the implications these may have about the processes that it models.

The equations that give the probabilities of the outcomes in an MPT model can be used to estimate the probabilities of the individual cognitive states or events that lead to them. This is called parameter estimation as the probability of each cognitive state or event occurring is a parameter of the equation that returns the probability of an outcome. Parameters can be estimated by fitting the model to data using a goodness-of-fit statistic and an optimization technique. Although there are several other ways to estimate parameters (Riefer & Batchelder, 1988; Erdfelder et al., 2009; Logacev & Dokudan, 2021), for all MPT models discussed in this thesis, Maximum-Likelihood Estimation (Cox and Hinkley, 1974) was used for goodness-of-fit, and the Nelder-Mead algorithm was used for optimization (Nelder & Mead,

1965). The models were implemented, the likelihood functions were defined and the optimization was carried out in R (R Core Team, 2021).

The quality of the MPT modelling that probabilities for individual cognitive states or events can be obtained takes it beyond the general-purpose statistical modelling techniques such as analysis of variance (ANOVA) or log-linear models which usually do not allow the measurement of underlying cognitive processes, but instead only provide the information that whether these underlying cognitive processes act in conjunction to create a difference between experimental conditions (Riefer & Batchelder, 1988).

After an estimation of the parameters of an MPT model are obtained, these parameters values are plugged in to the model to generate predictions. For example, after getting estimations for the parameters c and g_p in the model in Figure 2, these estimated values for the two parameters are plugged in to the equations in (1) and (3) to generate the predictions of the model. This procedure provides us with the probability of an implausible response given an implausible item, and the probability of a plausible response given a plausible item, which is collectively the probability of a correct response to a plausibility judgment trial. This predicted probability of a correct response can then be compared to the actual percentage of correct responses obtained in the experiment to assess the model's goodness-of-fit to the data. In this regard, a numeric value for the goodness-of-fit can be obtained by making use of fit statistics (Erdfelder et. al., 2009). However, because the MPT models suggested in this thesis can be easily compared in terms of goodness-of-fit to the data visually after generating model predictions, and because such a comparison of these models already reveal sufficient information about the data, a numeric measure for the goodness-of-fit is not used in this thesis. For example, the plot in Figure 3 shows the

predictions for the model in Figure 2, when the model was fit to the plausibility judgment data from the first experiment of Meng and Bader (2021). It is clear just from the plot that this model is unable to account for the differences between the levels of structure as well as meaning in the data, even when considering only the differences that were found significant in Meng and Bader's analysis. The technical reason why the model is unable to account for these differences has to do with the number of parameters that the model has, or its flexibility. Because the number of parameters in the model, that is two, is too little for the equations in (1-4) to result in a variety of values, each of which is close to one of the percentage of correct answers obtained from a condition in the data, when a single value obtained through estimation is plugged into the place of each parameter in the equation, it is impossible for this model to account for all of the significant differences found between the conditions.

It is important to know the technical reason why a model is unable to predict the data at hand because this can provide hints for the modeler as to how a model can be improved. In the case of the model in Figure 2, the technical problem tells us that we need to make further assumptions about the cognitive states or events that lead to the outcomes from the task. One way to address this issue is to suggest additional cognitive states or events that would introduce new parameters into the model as well as the equations that give the probability of each outcome. When the new parameters provide the model with the flexibility to adjust for the variety of results obtained from an experiment, the technical problem with the model will be solved. However, since the goal of modelling is not solely getting good fits, but to discover the feasibility of assumptions about the cognitive processes that lead to the outcomes

from an experiment, introduction of new cognitive states or events, and thus also parameters, into the model should be done with theoretical considerations.

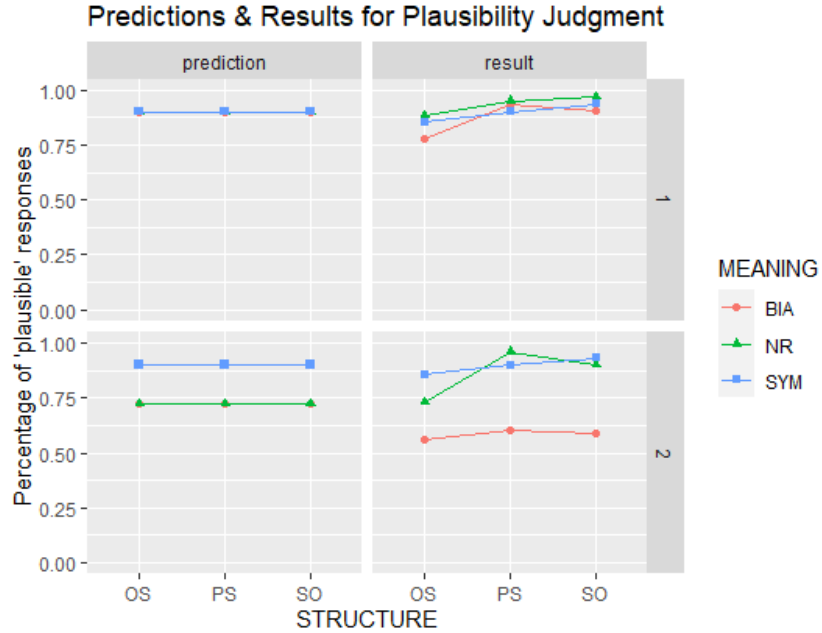


Figure 3. The plots on the left show the probability of a correct response to the plausibility judgment task for sentences as predicted by the model in Figure 2. The plots on the right show the percentage of ‘plausible’ responses in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points.

The entire procedure of creating a model, parameter estimation, generating predictions from the model and comparing these to the data, and revising a model is an iterative process that should be carried out in the light of theoretical considerations. For example, a theoretical suggestion about the processing of biased sentences which tells that the comprehension of these types of sentences, regardless of its structure, is more difficult than the comprehension of nonreversible and symmetrical sentences would lead to the introduction of a separate c parameter for the biased sentences in the model in Figure 2. The introduction of this new parameter would result in the equations in (5) and (6) which give the probability of a correct plausibility judgment response to biased sentences.

$$(5) \quad \text{Pr}(\text{'implausible' response} \mid \text{Biased implausible item}) = c_{bias} + (1 - c_{bias}) \times (1 - g_p)$$

$$(6) \quad \text{Pr}(\text{'plausible' response} \mid \text{Biased plausible item}) = c_{bias} + (1 - c_{bias}) \times g_p$$

When combined with the equations in (1) and (3) which give the probability of a correct plausibility judgment response to the sentences in the remaining meaning levels, nonreversible and symmetrical, predictions of greater variety than which would be obtained before the introduction of the new parameter to the model in Figure 2 can be obtained. Thus, the flexibility of the model in Figure 2 would be improved, and the predictions would fit the data more closely, allowing the model to account for more of the significant differences found between the conditions. However, little can be said about the theoretical motivation behind suggesting that the processing of biased sentences is more difficult regardless of their structure than the processing of nonreversible and symmetrical sentences, hence also about the theoretical motivation behind introducing a separate c for these sentences. Therefore, although such a revision to the model improves its goodness-of-fit, as the revision itself is not theoretically motivated, it is hard to say that the model with the new parameter offers insights of the value that are needed to learn more about the processes being studied.

On the contrary, when the entire modelling procedure is carried out with theoretical considerations in mind, the procedure can provide valuable insights into what assumptions can be made about the cognitive processes that are being studied, which are not entirely clear at a first glance from verbal explanations of data usually found in research articles that suggest the existence of cognitive events or states based on observed response frequencies.

For example, Logacev and Dokudan (2021) in a study of relative clause attachment preferences in ambiguous sentences, used data from two experiments which had the same design but were in different languages, the first being English and the second Turkish. In both experiments, participants answered comprehension questions like ‘Did the maid/princess/son scratch in public?’ after reading sentences like which are shown in Table 8, which elicited their relative clause attachment preferences. In both experiments, questions about unambiguous sentences, which are the sentences from the N1 and N2 attachment conditions, were answered notably inaccurately, with only %79 accuracy in the first experiment, and %66.5 in the second experiment. For the questions about ambiguous sentences, %58 of the responses in the first experiment were compatible with the second noun, while %58 of the responses in the second experiment were compatible with the first noun. The inaccuracy found for the unambiguous sentences led Pavel and Dokudan to question the apparent preference for N2 attachment found in the first experiment and the preference for N1 attachment found in the second experiment. To investigate whether the suggestion that there is N2 attachment preference in English and N1 attachment preference in Turkish is questionable, Pavel and Dokudan suggested two MPT models, the second of which assumed an extra cognitive event, namely recollection certainty/uncertainty, and thus an extra parameter, the probability equations for the outcomes of which are shown in (6-8) and (9-11), respectively.

$$(7) \quad \Pr(\text{'yes'} \mid N1) = a + (1 - a) \cdot g$$

$$(8) \quad \Pr(\text{'yes'} \mid N2) = (1 - a) \times g$$

$$(9) \quad \Pr(\text{'yes'} \mid \text{Ambiguous}) = a \times h + (1 - a) \times g$$

$$(10) \quad \Pr(\text{'yes'} \mid N1) = r_1 + (1 - r_1) \times g$$

$$(11) \quad \Pr(\text{'yes'} \mid N2) = (1 - r_2) \times g$$

$$(12) \quad \Pr(\text{'yes'} \mid \text{Ambiguous}) = h \times \Pr(\text{'yes'} \mid N1) + (1 - h) \cdot \Pr(\text{'yes'} \mid N2)$$

In the first MPT model (6-8), the parameter a stands for the probability that the participant is in an attentive state, the parameter g stands for the probability that the participant guesses ‘yes’ as a response to the attachment question, the parameter h stands for the probability that the participant will disambiguate an ambiguous relative clause towards attachment to the first noun. In the second MPT model (9-11), parameters h and g have the same meaning, while the parameters r_1 and r_2 stand for the probability that a correct representation of the sentence was generated and that this was recalled correctly when the probe is encountered for attachment to the first noun and attachment to the second noun respectively.

After fitting both models to the data from the first and the second experiments, Logacev and Dokudan (2021) found that the first model predicted more ‘yes’ responses to N1 questions and more ‘no’ responses to N2 questions than was found in the first experiment which was studying English, and that it predicted more ‘yes’ responses to N2 questions and more ‘no’ responses to N1 questions than was found in the second experiment which was studying Turkish. However, the second model, did not overestimate the proportions of these responses and deviated less from the data, and thus was a better fit to the data.

Because the second model (9-11) had an extra parameter, as discussed before, it was more flexible. Therefore, the fact that it fit the data better than the first model is not surprising. Although the addition of the new parameter, which stood for the cognitive states of recollection certainty/uncertainty is highly motivated by a theoretical background, in that it has been suggested in a number of studies which used detection or recognition paradigms (Erdfelder et. al., 2009), Logacev and Dokudan (2021) also performed the PSIS-LOO-CV (Vehtari et al., 2016) test to

assess the two model's out-of-sample performance. This test penalizes models that have additional flexibility, and so provides a balanced measure of their performance. The results of the test showed that the second model was preferable.

Table 8. Experimental Conditions and the Corresponding Example Sentences from the First and the Second Experiments of Logacev and Dokudan (2021)

Experiment 1	Ambiguous Attachment	The maid of the princess who scratched herself in public was terribly humiliated.
	N1 Attachment	The son of the princess who scratched himself in public was terribly humiliated.
	N2 Attachment	The son of the princess who scratched herself in public was terribly humiliated.
Experiment 2	Ambiguous Attachment	<i>Dün akşam, [birbirini döven]_{RC}</i> Yesterday evening, each other hit <i>[futbolcu-lar-ın]_{N1} [hayran-lar-ı]_{N2}</i> footballer-PL-GEN fan-PL-POSS <i>stadyumu hemen terk etti.</i> stadium immediately leave did.
	N1 Attachment	<i>Dün akşam, [birbirini döven]_{RC}</i> Yesterday evening, each other hit <i>[futbolcu-lar-ın]_{N1} [hayran-ı]_{N2}</i> footballer-PL-GEN fan.SG-POSS <i>stadyumu hemen terk etti.</i> stadium immediately leave did.
	N2 Attachment	<i>Dün akşam, [birbirini döven]_{RC}</i> Yesterday evening, each other hit <i>[futbolcu-nun]_{N1} [hayran-lar-ı]_{N2}</i> Footballer.SG-GEN fan-PL-POSS <i>stadyumu hemen terk etti.</i> stadium immediately leave did.

In conclusion, MPT modelling is a useful method of experimental data analysis which can be applied to categorical data and can also shed light on the underlying cognitive processes that result in the data. MPT models allow the estimation of probabilities of unobservable and hypothesized cognitive events or states from which predictions of hypotheses about the cognitive processes that lead to the outcomes from an experiment can be obtained. MPT modelling is also an iterative procedure whereby models can be created from which predictions are

obtained to be compared to data, and the comparison results can be used to revise the model, thus allowing the comparison of different hypotheses about the same data. Finally, the iterative procedure of model creation and revision should be done with theoretical considerations as the goal of modelling is to investigate the performance of theoretically motivated hypotheses about the underlying mechanisms that lead to data, rather than to create the model that best fits the data.

CHAPTER 4

MPT MODELS UNDER THE RETRIEVAL ACCOUNT

4.1 MPT models of plausibility judgment task responses

4.1.1 Retrieval account model 1

Although we aim to eventually model the plausibility judgement and agent/patient naming task responses jointly, to simplify the modelling process, we decided to start by modelling only the plausibility judgment task responses. We tried to create a model that would match the assumptions of the retrieval account as suggested by Meng and Bader (2021) as closely as possible while also taking into account that in some of the trials in their experiments the responses obtained from the participants may have simply been guesses.

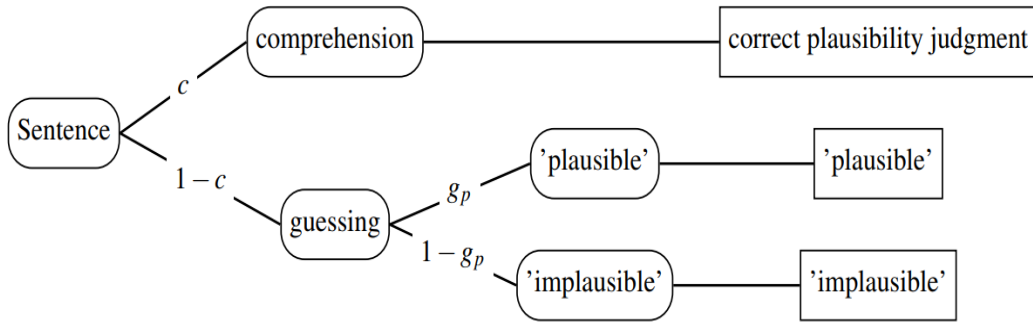


Figure 4. Our first model of plausibility judgment responses from the first experiment of Meng and Bader (2021)

In order to capture the responses that may have been guesses and to separate the assumptions of the retrieval account about the HPM from these trials, we first introduced the model in Figure 4. There are two parameters in this model. Parameter c represents the probability that an interpretation for the sentence that is faithful to the linguistic input was created (henceforth, ‘comprehension’), and the parameter g_p

represents the probability of a guess resulting in a ‘plausible’ response in the plausibility judgment task. This is a very simple model which assumes that guesses only happen when the sentence was not comprehended, and that comprehension always results in a correct plausibility judgment response. By looking at the visual representation of the model, we can derive the equations that give the probability of a correct plausibility judgment for each of the experimental conditions under the assumptions of the model. While deriving the equations, it must be considered which plausibility judgment response to a given experimental condition is correct. In the experiment of Meng and Bader (2021), the correct response to all of the NO1 conditions is ‘plausible’, and the correct response to all of the NO2 conditions, except for the symmetrical conditions, is ‘implausible’. This procedure reveals that the model does not predict any difference in the probability of a correct plausibility judgment between structure or meaning levels as in Table 9, where the equations for some of the experimental conditions are presented.

Table 9. Probability Equations Resulting from the Model in Figure 4

Experimental Condition	Pr(Correct Plausibility Judgment)
Nonreversible, NO1, Active SO	$c + (1 - c) \times g_p$
Nonreversible, NO2, Active SO	$c + (1 - c) \times (1 - g_p)$
Nonreversible, NO1, Passive	$c + (1 - c) \times g_p$
Nonreversible, NO1, Active OS	$c + (1 - c) \times g_p$
Biased, NO1, Active SO	$c + (1 - c) \times g_p$
Biased, NO2, Active SO	$c + (1 - c) \times (1 - g_p)$
Symmetrical, NO1, Active SO	$c + (1 - c) \times g_p$
Symmetrical, NO2, Active SO	$c + (1 - c) \times g_p$

However, it is evident from the data from the experiments of Meng and Bader (2021) that both noun order and the structure of a sentence affects the accuracy of plausibility judgments. Moreover, the accuracy differences between the NO1 and NO2 conditions are not attributed to comprehension failure under the retrieval

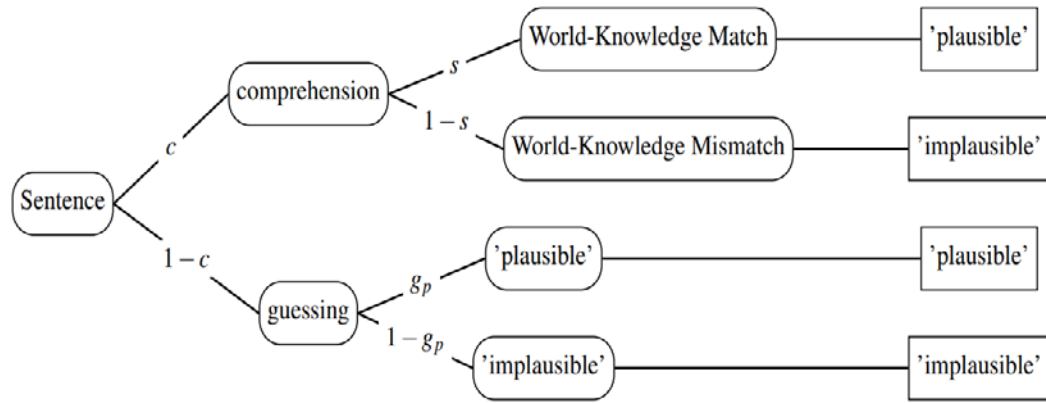
account as the model in Figure 4 suggests. Therefore, it is clear, as expressed in Chapter 2, before estimating parameters and comparing the model's predictions, that the model will be insufficient in both explaining the data and reflecting the assumptions of the retrieval account.

4.1.2 Retrieval account model 2

In order to start addressing this problem with the model, we created the model in Figure 5 with regard to the significant effect of noun order⁵ found in the experiments of Meng and Bader (2021). They attribute the cause of this significant effect of noun order to the plausibility of the sentences rather than comprehension failure as was predicted by our first model. Moreover, the finding that the offline measure of the plausibility of the biased and nonreversible sentences yielded average plausibility scores similar to the percentage of 'plausible' responses obtained in their experiments for the same sentences supports this claim (Bader & Meng, 2018). Therefore, it is possible to say that NO2 biased and nonreversible sentences are implausible because of what people know about the world. Thus, we found it fitting to suggest that there might be a process whereby the plausibility of a sentence is judged with regard to world-knowledge after successful comprehension, where the probability of a successful match with world-knowledge is similar to the offline plausibility score of the sentence. To implement this idea within the model, we added the parameter s , which represents the probability of an interpretation matching with the comprehender's world-knowledge. In addition, we have changed the possible outputs of the model so that they show which path leads to a 'plausible' or

⁵ This effect only reaches significance when symmetrical sentences are ignored because in the symmetrical condition, both noun orders are result in plausible sentences; hence, they are called symmetrical.

‘implausible’ response rather than a correct or incorrect plausibility judgment. This is because, in this way, it can easily be determined which of the new states added to the model, World-Knowledge Match or Mismatch, leads to a ‘plausible’ or ‘implausible’ response.



Conditions	Value of s
Nonreversible, NO1	$s = 1$
Nonreversible, NO2	$s = 0$
Biased, NO1	$s = 1$
Biased, NO2	$s = s_{biased}$
Symmetrical, NO1	$s = 1$
Symmetrical, NO2	$s = 1$

Figure 5. Our second model of plausibility judgment responses from the first experiment of Meng and Bader (2021).

Added to the assumptions of the first model, the model in Figure 5 assumes that once comprehension is successful, and so an interpretation faithful to the linguistic input is obtained, a ‘plausible’ response that is not a guess can only be provided when the interpretation matches the world-knowledge of the comprehender. The fact that in this model a process that determines whether the sentence matches the world knowledge of the comprehender can only occur after an interpretation faithful to the linguistic input is obtained is also in-line with the assumption of the retrieval account that the accuracy of plausibility judgments result from postinterpretive processes. Because there are only two possible outcomes in the

model in Figure 5, ‘plausible’ and ‘implausible’, the equation (13) that returns the probability of a ‘plausible’ answer for all conditions is the same.

$$(13) \quad \text{Pr}(\text{‘plausible’ Response}) = (c \times s) + (1 - c) \times g_p$$

However, the model as such predicts that the probability of a ‘plausible’ answer for all conditions is also the same, which does not reflect what the data shows at all. In order for the predictions to match the data more closely, we must take into account that nonreversible and biased NO1 sentences are highly plausible, and NO2 biased sentences are somewhat more plausible than NO2 nonreversible sentences, whereas NO2 nonreversible sentences are highly implausible. One way to have our model reflect that is to fix the value of s for some of the conditions. Since s is the probability of a match between an interpretation and the comprehender’s world-knowledge, and we can assume that this probability is very low for nonreversible NO2 sentences, and very high for nonreversible NO1 and all symmetrical sentences, we thought it acceptable as an initial step to fix the value of s to 0 for nonreversible NO2 conditions, and to 1 for nonreversible NO1 and all symmetrical sentences. For the biased sentences however, s remains a free parameter, the value of which will be estimated when the model is fit to the data. Plugging in the fixed values for the parameter s results in the equations shown in Table 10.

Table 10. Probability Equations Resulting from the Model in Figure 5

Experimental Condition	Pr(‘plausible’ Response)
Nonreversible, NO1	$(c \times 1) + (1 - c) \times g_p = c + (1 - c) \times g_p$
Nonreversible, NO2	$(c \times 0) + (1 - c) \times (1 - g_p) = (1 - c) \times (1 - g_p)$
Biased, NO1	$(c \times 1) + (1 - c) \times g_p = c + (1 - c) \times g_p$
Biased, NO2	$(c \times s_{biased}) + (1 - c) \times (1 - g_p)$
Symmetrical, NO1	$(c \times 1) + (1 - c) \times g_p = c + (1 - c) \times g_p$
Symmetrical, NO2	$(c \times 1) + (1 - c) \times g_p = c + (1 - c) \times g_p$

It is evident from the equations that the model will predict different probabilities of a ‘plausible’ response for NO2 biased and nonreversible conditions, as well as predicting a difference between all NO1 and 2 conditions except for the symmetrical conditions, which matches the data more closely as we needed. Moreover, an additional prediction made by the model due to the fixed values for the s parameter is that a ‘plausible’ response to a nonreversible NO2 item can only be provided in the form of a guess, which we believe does not go against any assumption of the retrieval account as obtaining an interpretation that is faithful to the linguistic input for these sentences should almost always lead to the correct plausibility judgment, which is an ‘implausible’ response, under the retrieval account.

Table 11. Estimated Values for each of the Free Parameters when the Model in Figure 5 is Fit to the Data

Parameter	Meaning of the Parameter	Estimated Value
c	Probability of comprehension	~ 0.76
s_{bias}	Probability of a biased NO2 item matching the world knowledge of the comprehender	~ 0.37
g_p	Probability of guessing ‘plausible’	~ 0.57

When the model is fit to the data, the values for each of the free parameters are estimated to be around the values shown in Table 11. When the functions in Table 10 are evaluated with the parameter values in Table 11 the predicted percentage of ‘plausible’ responses for each condition are as shown in Table 12 and Figure 6. Here, we can see that the model predicts percentages of ‘plausible’ responses that better match the average percentages obtained in the experiment for each meaning level.

Table 12. Results from the Plausibility Judgment Task in the Data and the Predictions of the Model in Figure 5 with the Values from Table 11 Plugged in

Order	Meaning	Structure	Results	Predictions
O1	nonreversible	SO	0.965035	0.901786
O1	nonreversible	Passive	0.951389	0.901786
O1	nonreversible	OS	0.888889	0.901786
O1	biased	SO	0.919014	0.901786
O1	biased	Passive	0.929078	0.901786
O1	biased	OS	0.785455	0.901786
O1	symmetrical	SO	0.943262	0.901786
O1	symmetrical	Passive	0.902878	0.901786
O1	symmetrical	OS	0.847584	0.901786
O2	nonreversible	SO	0.092527	0.129619
O2	nonreversible	Passive	0.042254	0.129619
O2	nonreversible	OS	0.257246	0.129619
O2	biased	SO	0.406475	0.420978
O2	biased	Passive	0.405797	0.420978
O2	biased	OS	0.450909	0.420978
O2	symmetrical	SO	0.932624	0.901786
O2	symmetrical	Passive	0.903915	0.901786
O2	symmetrical	OS	0.843066	0.901786

However, it is also evident that this model falls short of accounting for the differences between structure levels in the data. In addition, a comparison of the estimated parameter values in Table 11 with the results from the experiment in Table 12, suggest that some portion of the resulting percentages of ‘plausible’ responses to each condition are not due to comprehension or a biased NO2 item matching the world knowledge of the comprehender. The estimated parameter values suggest this because the estimated average probability of comprehension, that is 0.76, is lower than the percentage of ‘plausible’ responses obtained for most of the conditions, as can be seen in Table 12, but the predictions of the model match the results more closely than the average probability of comprehension due to the estimated value of the average probability of guessing ‘plausible’, that is 0.57, and that is just over 0.5, which shows that the model predicts a small bias for guessing ‘plausible’. A similar analysis leads also to the conclusion that the estimated probability of a biased NO2

item matching the world knowledge of the comprehender, that is 0.37, is also compensated by the same bias for guessing ‘plausible’ for the model predictions to match the results from the experiment more closely.

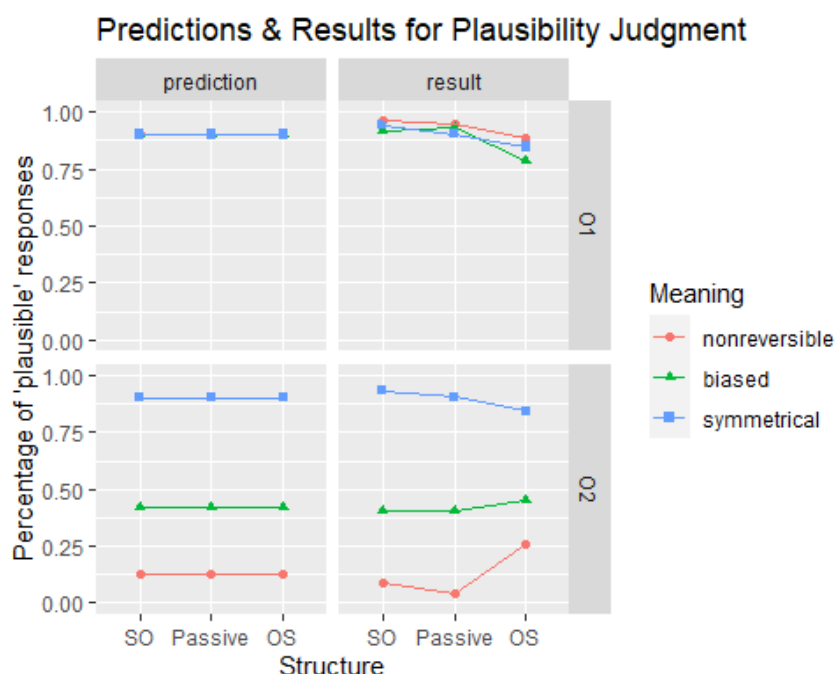


Figure 6. The plot on the left shows the percentage of ‘plausible’ responses to NO2 sentences as predicted by the model in Figure 5. The plot on the right shows the percentage of ‘plausible’ responses to NO2 sentences in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

4.1.3 Retrieval account model 3

The results from the experiments of Meng and Bader (2021) also show that OS sentences were judged less accurately than passive and SO sentences on average across all meaning levels (nonreversible, biased, and symmetrical), for both NO1 and NO2 sentences. Meng and Bader attributed the decrease in accuracy for OS sentences in plausibility judgments to the information-structural markedness of these sentences, arguing that the lack of an appropriate discourse context causes the acceptability of these sentences to drop. To accommodate this assumption into our

model, we related the need for an appropriate discourse context to license OS sentences to the concept of pragmatic well-formedness, or, in other words, to the felicity status of an OS sentence in isolation. Therefore, we added a process to the model whereby the pragmatic well-formedness of a sentence is judged following successful comprehension, and a match between the resulting interpretation and the comprehender's world-knowledge. The parameter f in the model in Figure 7 represents the probability of a sentence being judged as felicitous, or pragmatically well-formed.

Added to the assumptions of the first two models, the model in Figure 7 assumes that a 'plausible' response that is not a guess can only be obtained if the sentence was judged as being felicitous, and that the felicity of a sentence is only judged if the interpretation of the sentence is compatible with the world-knowledge of the comprehender.

The predicted probability of a 'plausible' response for every condition for the model in Figure 7 is given by the equation in (14) where the value of the parameter s is determined by the condition.

$$(14) \quad \text{Pr}(\text{'plausible' Response}) = (c \times s \times f) + (1 - c) \times g_p$$

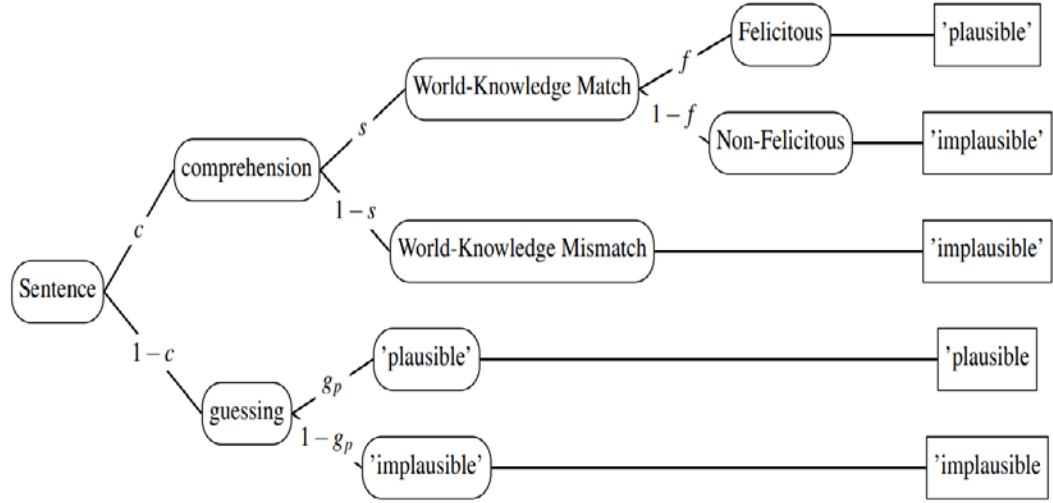
With the model as such, we are unable to predict the differences in the percentages between the structure levels. With regard to the markedness effect unique to the OS structure suggested by Bader and Meng (2018, 2021), we found it fitting to resort to fixing the value of the parameter f for all conditions except OS conditions to 1 and leave f as a free parameter for the OS conditions as shown in Figure 7 as a solution to this problem with the model. Fixing the f parameter's value to 1 for the SO and Passive conditions means that the sentences belonging to these conditions will always be judged as felicitous. Clearly this does not reflect the

reality, but this fixation is sufficient for the purpose of getting the model to reflect the assumption that felicity suffers for OS sentences that are isolated from an appropriate discourse context.

When the model is fit to the data, the values for each of the parameters are estimated to be around the values shown in Table 13. When the function in (14) is evaluated with the parameter values in Figure 7 and Table 13 the predicted percentage of ‘plausible’ responses for each condition are as shown in Figure 8 and Table 14. We can see that the model successfully predicts the significant differences between the levels of structure for all NO1 conditions and NO2 symmetrical conditions. However, for the NO2 nonreversible conditions the predicted percentage of ‘plausible’ responses is the same for all structure levels; and for the NO2 biased OS condition, the predicted percentage of ‘plausible’ responses is lower than the SO and Passive structures; both of which do not match the differences between the structures in the data.

The technical reason why the model in Figure 7 predicts equal percentages of ‘plausible’ responses for the nonreversible NO2 conditions has to do with the fixed value that is assigned to the s parameter for these conditions. Because in the equation in (14) that returns the probability of a ‘plausible’ response for all conditions, the parameter f is multiplied by the parameters s and c , and because the value of s for the nonreversible NO2 conditions is 0 as shown in Figure 7, evaluating the function in (14) for these conditions results in the nullification of the parameters c and f , in other words, these parameters have no effect on the predicted outcome probability of a ‘plausible’ response. Therefore, we can say that the model predicts that the only way a ‘plausible’ response can be obtained for the NO2 nonreversible conditions, regardless of whether the sentence is in the OS structure or not, is through guesses,

thus the predicted probability of a ‘plausible’ response for all of these conditions are equal.



Conditions	Value of s	Value of f
Nonreversible, NO1, SO	$s = 1$	f
Nonreversible, NO2, SO	$s = 0$	f
Nonreversible, NO1, Passive	$s = 1$	f
Nonreversible, NO2, Passive	$s = 0$	f
Nonreversible, NO1, OS	$s = 1$	$f = f_{os}$
Nonreversible, NO2, OS	$s = 0$	$f = f_{os}$
Biased, NO1, SO	$s = 1$	f
Biased, NO2, SO	$s = s_{biased}$	f
Biased, NO1, Passive	$s = 1$	f
Biased, NO2, Passive	$s = s_{biased}$	f
Biased, NO1, OS	$s = 1$	$f = f_{os}$
Biased, NO2, OS	$s = s_{biased}$	$f = f_{os}$
Symmetrifal, NO1, SO	$s = 1$	f
Symmetrical, NO2, SO	$s = 1$	f
Symmetrical, NO1, Passive	$s = 1$	f
Symmetrical, NO2, Passive	$s = 1$	f
Symmetrical, NO1, OS	$s = 1$	$f = f_{os}$
Symmetrical, NO2, OS	$s = 1$	$f = f_{os}$

Figure 7. Our third model of plausibility judgment responses from the first experiment of Meng and Bader (2021)

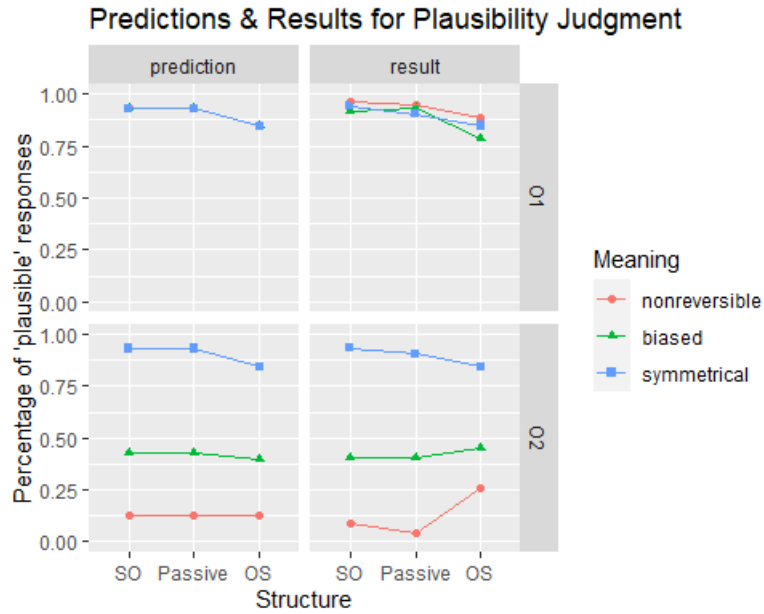


Figure 8. The plot on the left shows the percentage of ‘plausible’ responses as predicted by the model in Figure 7. The plot on the right shows the percentage of ‘plausible’ responses in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

More importantly, the model in Figure 7 predicts a lower percentage of ‘plausible’ responses for the NO2 biased OS condition than NO2 biased SO and Passive conditions, whereas the results show that NO2 biased OS condition sentences were judged more plausible than NO2 SO and Passive conditions. The technical reason why this is so, has to do with effect that all of the NO1 OS conditions as well as the NO2 symmetrical OS condition have on the optimization function which fits the model to the data, thus estimating the parameters. Firstly, because the f parameter is set to 1 for all non-OS conditions, which is the highest probability value possible, there is no way for the optimization function to adjust the value of the parameter that is special to the OS conditions, which is f_{os} , to become higher than the value of f in the rest of the conditions and thus predict a higher probability of a ‘plausible’ response to the OS conditions than for the SO and passive conditions. Secondly, this impossibility cannot be compensated further than a certain

degree by the optimization function through the adjustment of the rest of the parameters because in the results from all of the conditions, apart from the NO2 nonreversible OS condition, OS sentences are judged less ‘plausible’ than the rest of the structures.

Table 13. Estimated Values for each of the Free Parameters when the Model in Figure 7 is Fit to the Data

Parameter	Meaning of the Parameter	Estimated Value
c	Probability of comprehension	~ 0.80
S_{bias}	Probability of a biased NO2 item matching the world knowledge of the comprehender	~ 0.37
f_{os}	Probability of an OS sentence being judged felicitous.	~ 0.89
g_p	Probability of guessing ‘plausible’	~ 0.65

Table 14. Results from the Plausibility Judgment Task in the Data and the Predictions of the Model in Figure 5 with the Values from Table 13 Plugged in

Order	Meaning	Structure	Results	Predictions
O1	nonreversible	SO	0.965035	0.930482
O1	nonreversible	Passive	0.951389	0.930482
O1	nonreversible	OS	0.888889	0.844116
O1	biased	SO	0.919014	0.930482
O1	biased	Passive	0.929078	0.930482
O1	biased	OS	0.785455	0.844116
O1	symmetrical	SO	0.943262	0.930482
O1	symmetrical	Passive	0.902878	0.930482
O1	symmetrical	OS	0.847584	0.844116
O2	nonreversible	SO	0.092527	0.130166
O2	nonreversible	Passive	0.042254	0.130166
O2	nonreversible	OS	0.257246	0.130166
O2	biased	SO	0.406475	0.430168
O2	biased	Passive	0.405797	0.430168
O2	biased	OS	0.450909	0.397794
O2	symmetrical	SO	0.932624	0.930482
O2	symmetrical	Passive	0.903915	0.930482
O2	symmetrical	OS	0.843066	0.844116

However, when we compare the estimated value of the parameter g_p for the model in Figure 5, which was 0.57, to the one for the model in Figure 7, which is

0.65, it is evident that the optimization function compensated for the impossibility mentioned earlier by adjusting the probability of a ‘plausible’ guess to become higher, because this would increase the predicted probability of a ‘plausible’ response to all structures in the NO2 nonreversible and biased conditions, thus also bringing the probability of a ‘plausible’ response to the NO2 nonreversible and biased OS conditions closer to the percentages in the data for the same conditions. This can be confirmed by looking at the predicted probabilities of a ‘plausible’ response to the NO2 nonreversible and biased conditions in Table 14, which show that the predictions are higher for the SO and passive sentences in these conditions than the results for the same conditions. To not leave it unmentioned, we can also say that the model in Figure 7 predicts a stronger bias for guessing ‘plausible’ than the model in Figure 5.

This brings into question the explanation provided by Meng and Bader (2021) for the observed lower accuracy in plausibility judgment for OS sentences. Meng and Bader, as was formerly mentioned, suggest that the lower accuracy scores for OS sentences are a result of the information structural markedness of these sentences in the lack of an appropriate discourse context. While creating the model in Figure 7, we took this suggestion to mean that the pragmatic well-formedness of OS sentences should be lower than SO and passive sentences, due to the lack of a discourse context for all sentences in the experiment, and thus made the adjustments to the f parameter which we just discussed. However, as shown in the former paragraph, it is impossible for such a model to account for the higher percentage of ‘plausible’ responses to the NO2 nonreversible and biased sentences. In other words, the fact that the correct response to the trials with NO2 nonreversible and biased sentences was ‘implausible’ in the first experiment of Meng and Bader, and so the observed

lower accuracy scores for OS sentences actually reflect more ‘plausible’ responses, makes it impossible for a model that takes the suggestion of Meng and Bader about the information-structural markedness of OS sentences to mean that these sentences are less plausible when there is the lack of an appropriate discourse context.

Therefore, we must accept that the information structural markedness of the OS sentences in the lack of an appropriate discourse context does not make the sentences less plausible as we have implemented in the model in Figure 7, and that some other property of the OS sentences, that is perhaps still related to their information structural markedness, results in inaccurate plausibility judgments for these sentences.

4.1.4 Retrieval account model 4

Another way to understand the explanation of Bader and Meng for the observed lower accuracy in plausibility judgment for OS sentences is to assume that OS sentences are harder to comprehend than the SO and Passive sentences. It is possible to implement this suggestion in a model if we remove the f parameter and go back to the model in Figure 5, while also assuming a different value for c in all OS conditions as shown in Table 15. When the model in Figure 5 with the parameter values shown in Table 15 is fit to the data, the values for each of the free parameters are estimated to be around the values shown in Table 16.

When the function in Table 15 is evaluated with the parameter values in Table 15 and Table 16, the predicted percentage of ‘plausible’ responses for each condition are as shown in Figure 9 and Table 18. We can see that the model can now successfully predict the differences between the levels of structure for all conditions, as well as predicting lower accuracy for OS sentences across conditions, which is

reflected by the higher percentage of ‘plausible’ responses predicted for OS sentences in NO2 nonreversible and biased conditions, and the lower percentage of ‘plausible’ responses predicted for OS sentences in all symmetrical and NO1 conditions.

Table 15. New Fixed Parameter Values for the model in Figure 5

$\text{Pr}(\text{'plausible' Response}) = (c \times s) + (1 - c) \times g_p$		
Conditions	Value of c	Value of s
Nonreversible, NO1, SO	c	$s = 1$
Nonreversible, NO2, SO	c	$s = 0$
Nonreversible, NO1, Passive	c	$s = 1$
Nonreversible, NO2, Passive	c	$s = 0$
Nonreversible, NO1, OS	$c = c_{OS}$	$s = 1$
Nonreversible, NO2, OS	$c = c_{OS}$	$s = 0$
Biased, NO1, SO	c	$s = 1$
Biased, NO2, SO	c	$s = s_{biased}$
Biased, NO1, Passive	c	$s = 1$
Biased, NO2, Passive	c	$s = s_{biased}$
Biased, NO1, OS	$c = c_{OS}$	$s = 1$
Biased, NO2, OS	$c = c_{OS}$	$s = s_{biased}$
Symmetrical, NO1, SO	c	$s = 1$
Symmetrical, NO2, SO	c	$s = 1$
Symmetrical, NO1, Passive	c	$s = 1$
Symmetrical, NO2, Passive	c	$s = 1$
Symmetrical, NO1, OS	$c = c_{OS}$	$s = 1$
Symmetrical, NO2, OS	$c = c_{OS}$	$s = 1$

Table 16. Estimated Values for each of the Free Parameters when the Model in Figure 5 with the Fixed Parameter Values in Table 15 is Fit to the Data

Parameter	Meaning of the Parameter	Estimated Value
c	Probability of comprehension for SO and Passive sentences	~0.84
c_{OS}	Probability of comprehension for OS sentences	~0.61
s_{bias}	Probability of a biased NO2 item matching the world knowledge of the comprehender	~0.37
g_p	Probability of guessing ‘plausible’	~0.57

Moreover, because the value of c for the OS conditions was implemented as a separate free parameter, c_{OS} , in this model, the model was still able to predict the

higher percentage of ‘plausible’ responses for the OS sentences in the NO2 nonreversible condition despite the fact that the value of the s parameter being fixed to 0. This is not because the same nullification of the c parameter that the previous model suffered from did not happen, but because the predicted probability of comprehension failure, denoted by the term ‘ $(1-c)$ ’, was different for OS sentences in these conditions, as is apparent from the comparison of the equations that give the probability of a ‘plausible’ response for the NO2 nonreversible OS condition between the two models as shown in Table 17. However, the assumption that the only way a ‘plausible’ response can be obtained from NO2 nonreversible conditions is through guesses still holds for the model.

The estimated values for the model parameters in Table 17 show a notable difference between the probabilities of comprehension for the SO and passive sentences on one hand, and for the OS sentences on the other, with it being as high as 0.84 for SO and passive sentences, and as low as 0.61 for OS sentences. It is up to question whether the effect that the lack of an appropriate discourse context has on the comprehension of OS sentences as suggested by Meng and Bader (2021) should be as large as predicted by this model. Future studies could address this question by making use of a self-paced reading paradigm similar to that used by Cutter, Paterson and Filik (2021), but by putting German OS, SO and passive sentences in the follow-up position and manipulating the first sentences so that there are both conditions where there is an appropriate discourse context and where there is not for the OS sentences. The large effect of a lack of discourse context on the probability of comprehension for OS sentences found in this thesis should show up as an effect on the reading times for follow-up sentences in such a study, if the assumption that the comprehension of OS sentences suffers from a lack of an appropriate discourse

context holds truth. In addition, it should also be considered that the effect of discourse context on the comprehension of sentences could be language-specific, which would suggest that the magnitude of the effect in question here may have to do with the status of OS sentences in the German language.

Table 17. Comparison of Probability Equations that Return the Probability of a 'plausible' Response for the Two Models

Model Identity	Experimental Condition	Pr('plausible' Response)
Model 3, Figure 7	Nonreversible, NO2, OS	$(c \times 0) + (1 - c) \times g_p = (1 - c) \times g_p$
Model 2, Figure 5, With the values from Table 16 plugged in	Nonreversible, NO2, OS	$(cos \times 0) + (1 - cos) \times g_p = (1 - cos) \times g_p$

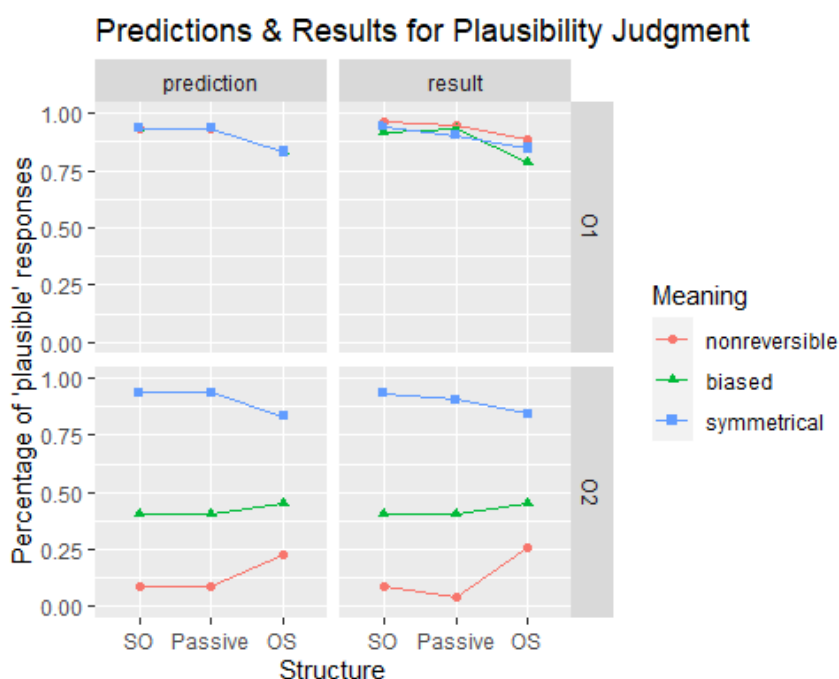


Figure 9. The plot on the left shows the percentage of 'plausible' responses as predicted by the model in Figure 5 with the parameter values shown in (18) and (19). The plot on the right shows the percentage of 'plausible' responses in the data. The percentage of 'plausible' responses are shown on the y-axis and the levels of the factor 'structure' are shown on the x-axis. The levels of the factor 'meaning' are represented by the color of the lines and the shape of the points

Table 18. Results from the Plausibility Judgment Task in the Data and the Predictions of the Model in Figure 5 with the Values from Table 16 Plugged in

Order	Meaning	Structure	Results	Predictions
O1	nonreversible	SO	0.965035	0.935823
O1	nonreversible	Passive	0.951389	0.935823
O1	nonreversible	OS	0.888889	0.832582
O1	biased	SO	0.919014	0.935823
O1	biased	Passive	0.929078	0.935823
O1	biased	OS	0.785455	0.832582
O1	symmetrical	SO	0.943262	0.935823
O1	symmetrical	Passive	0.902878	0.935823
O1	symmetrical	OS	0.847584	0.832582
O2	nonreversible	SO	0.092527	0.085941
O2	nonreversible	Passive	0.042254	0.085941
O2	nonreversible	OS	0.257246	0.224193
O2	biased	SO	0.406475	0.405463
O2	biased	Passive	0.405797	0.405463
O2	biased	OS	0.450909	0.452923
O2	symmetrical	SO	0.932624	0.935823
O2	symmetrical	Passive	0.903915	0.935823
O2	symmetrical	OS	0.843066	0.832582

In conclusion, this final MPT model of the plausibility judgment data from the Meng and Bader (2021) study is sufficient for the purpose of this thesis. The modelling procedure has firstly revealed that we can assume a postinterpretive process whereby the plausibility of a sentence is judged with regard to world-knowledge, which only takes place following successful comprehension to account for the differences between the percentages of ‘plausible’ responses for the different levels of meaning obtained from the first experiment of Bader and Meng. Secondly, the procedure has also revealed that we cannot assume a postinterpretive process whereby the pragmatic well-formedness of a sentence is judged, and so the plausibility of the sentence is determined following successful comprehension, to account for the percentages of ‘plausible’ responses obtained for the OS condition in the data. Finally, the procedure has revealed that the information structural markedness of the OS sentences, as suggested by Meng and Bader (2021), could not

have reduced the plausibility of these sentences, but that assuming that OS sentences are harder to comprehend than SO and Passive sentences, possibly due to the same property of the OS sentences, can explain the data, provided that there is a large effect on the probability of comprehension of OS sentences when there is a lack of an appropriate discourse context.

4.2 Joint MPT models of agent/patient naming and plausibility judgement task responses

As was mentioned before, we aim to model the plausibility judgment and agent/patient naming task responses jointly. Similar to what we did so far with the MPT models of plausibility judgment, we will discuss the suggestions of Bader and Meng (2018, 2021) that explain the agent/patient naming data under the retrieval account, and try to integrate these into our models.

4.2.1 Retrieval account model 5

Under the retrieval account, the task performance effects found in the experiment reflect retrieval errors rather than a wrong interpretation of the sentence. Therefore, for a start, we can assume that comprehension must occur before retrieval takes place in an MPT model of agent/patient naming under the retrieval account. To this end, we have added the parameter r , which is the probability of successful retrieval of either the patient or the agent, depending on the probe, to our final model of plausibility judgment task responses, as can be seen in Figure 10.

The model in Figure 10 assumes that a correct agent/patient naming response can only be provided if comprehension is achieved. Hence, it also assumes that a correct agent/patient naming response is impossible through a guess. Moreover,

because correct retrieval of the agent and the patient is possible in the case of both world-knowledge match and mismatch, the s parameter, which is the probability of an interpretation matching with the comprehender's world-knowledge, has no effect on the agent/patient naming outcome, as is evident from the equation in (15). The fact that this is so reflects the suggestion under the retrieval account that incorrect plausibility judgments can also lead to correct agent patient naming.

$$(15) \quad \text{Pr}(\text{Correct Naming}) = (c \times s \times r) + (c \times (1-s) \times r) = (c \times r)$$

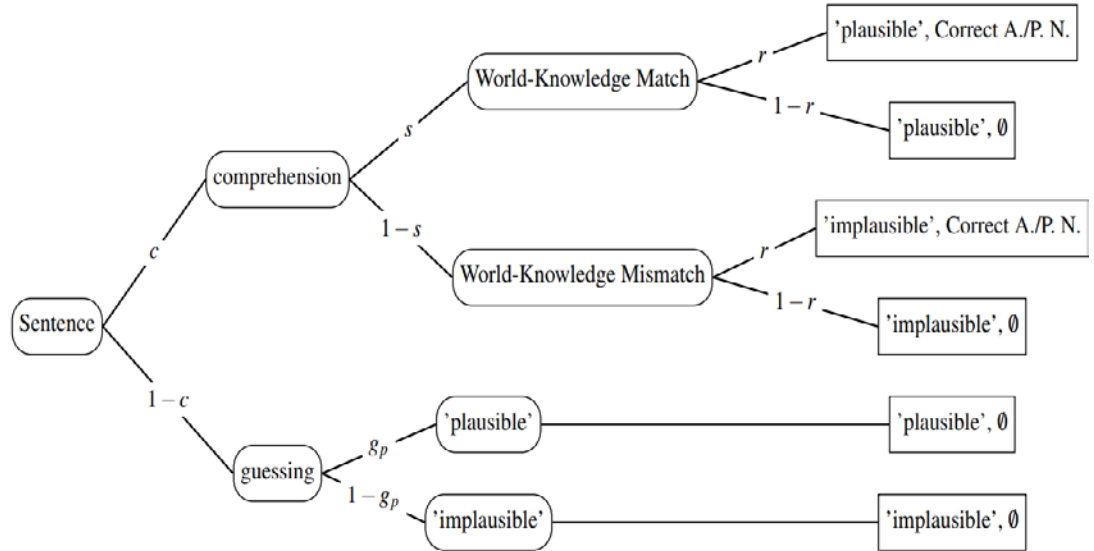
However, it is also evident from the equation in (15) that the model as such would fail to predict the significant differences in accuracy between the structures across all conditions reported in Meng and Bader (2021)'s experiments. In order to overcome this shortcoming of the model, we will start with the suggestion of Bader and Meng (2018) that success in the retrieval of the agent and the patient is probabilistically determined by the degree of match between the retrieval cues and the attributes of the two nouns competing for retrieval. Bader and Meng (2018) suggest four cues that reflect the typical properties of an agent. The first cue, 'Plausibility', describes whether a noun is a 'plausible' 'do-er' of an action or not. For example, for the verb 'cook', the plausibility cue will match with an animate noun. The other three cues describe syntactic properties of a noun. The 'Position' cue will match with a noun that is in the typical agent position, the 'Category' cue will match with a noun phrase only, and the 'Function' cue will match with the noun that syntactically functions as the agent of the sentence. For example, for a passive sentence, the 'Function' cue matches with the semantic object of the sentence, because under this account, this noun is assumed to be syntactically functioning as the agent of the sentence. Table 19 shows the cue-match status of items from all

conditions in the data. for the four cues suggested when the probe requests the retrieval of the agent.

A simple way to understand in terms of an MPT model the notion of a degree of match between the retrieval cues and the attributes of the two competing nouns, is to think of the concept in terms of the ratio of cues matching the target noun to the total number of available cues, with the cues that match both nouns counted as half of a match with the target. The numbers are calculated as such: each cue that matches the target noun contributes plus 1 towards the total dividend, each cue that matches both of nouns contributes plus 0.5 towards the total dividend, and each cue that matches only the competitor noun has no contribution towards the total dividend, or contributes plus 0 towards the total dividend, and after the contribution of each cue is calculated, the total dividend is divided by 4, that is the total number of available cues as suggested by Meng and Bader (2021). The ratio that results from such an approach is listed for each condition of the first experiment of Meng and Bader (2021) under the ‘Ratio’ column of Table 19. For example, for a NO1 nonreversible SO condition item, three cues, namely Plausibility, Position and Function, match the target noun while the Category cue matches both nouns, hence the total dividend for this condition is 3.5, and when this number is divided by 4, that is the total number of available cues, the ratio that results from this calculation is 0.88 when rounded up.

Our first trial with the model in Figure 10, assumes that the value of r , that is the probability of successful retrieval of the probed noun, the agent or the patient, is exactly the ratio that was explained in the former paragraph and is also shown for every condition in Table 19, along with the accuracy scores obtained for each of these conditions. This is a very basic way to understand the suggestion of Meng and Bader (2021), and it does not require the estimation of the parameter r but starting

from this basic assumption can help us understand both the model and the data better, and so, can show us how the model can be improved.



Conditions	Value of c	Value of s
Nonreversible, NO1, SO	c	$s = 1$
Nonreversible, NO2, SO	c	$s = 0$
Nonreversible, NO1, Passive	c	$s = 1$
Nonreversible, NO2, Passive	c	$s = 0$
Nonreversible, NO1, OS	$c = c_{OS}$	$s = 1$
Nonreversible, NO2, OS	$c = c_{OS}$	$s = 0$
Biased, NO1, SO	c	$s = 1$
Biased, NO2, SO	c	$s = s_{biased}$
Biased, NO1, Passive	c	$s = 1$
Biased, NO2, Passive	c	$s = s_{biased}$
Biased, NO1, OS	$c = c_{OS}$	$s = 1$
Biased, NO2, OS	$c = c_{OS}$	$s = s_{biased}$
Symmetrical, NO1, SO	c	$s = 1$
Symmetrical, NO2, SO	c	$s = 1$
Symmetrical, NO1, Passive	c	$s = 1$
Symmetrical, NO2, Passive	c	$s = 1$
Symmetrical, NO1, OS	$c = c_{OS}$	$s = 1$
Symmetrical, NO2, OS	$c = c_{OS}$	$s = 1$

Figure 10. Our first model of the agent/patient naming task responses in Meng and Bader (2021), which is built over our final model of the plausibility judgment task responses. ‘Correct A./P. N.’ stands for correct agent/patient naming, and the empty set symbol represents either incorrect agent/patient naming or the lack of a response to the agent/patient naming task

Before fitting the model in Figure 10 to both the plausibility judgment and the agent/patient naming data. We must also consider the cue-match status of every condition when the probe requests the retrieval of the patient. To calculate the cue-match ratio for these conditions, we used the same method as with the agent probe, with the extra assumption that if a certain cue is not available for a particular structure, this cue will not be added to the total number of cues, which determines the divisor of the equation used for calculating the cue-match ratio. This extra assumption is made because it is assumed under the account of Meng and Bader (2021) that the function cue matches with the noun that is in the syntactic position of the agent in a passive sentence, which is the noun that is the semantic object of the sentence, hence it must be assumed that the syntactic position for the object in a passive sentence is left empty, or there remains only a trace of the moved constituent, that is the semantic object of the sentence, in this position. Therefore, in our model, for every condition that features a passive sentence, we assume that the Function cue is not among the list of the available cues, and so for these conditions, the divisor used in the calculation of the cue-match ratio, is three, instead of four. The ratios that result from this calculation are listed for every condition from the first experiment of Meng and Bader, along with the accuracy score obtained for each of these conditions, in Table 19 and Table 20.

When the model in Figure 10 is fit to both the plausibility judgment and the agent/patient naming data from the first experiment of Meng and Bader (2021), using the ratios listed in Table 19 and Table 20 as the exact values of the parameter r for the corresponding conditions, the estimated parameter values are as shown in Table 21.

Table 19. Cue-Match Status of Items from all Conditions in the data. when the Probe Requests the Retrieval of the Agent

Experimental Condition (Agent Probe)	Syntactic Cues				Accuracy	Ratio
	Plausibility	Position	Function	Category		
Nonreversible-SO-NO1	target	target	target	-	0.97	0.88
Nonreversible-PS-NO1	target	competitor	competitor	target	0.94	0.50
Nonreversible-OS-NO1	target	competitor	target	-	0.89	0.63
Biased-SO-NO1	target	target	target	-	0.95	0.88
Biased-PS-NO1	target	competitor	competitor	target	0.87	0.50
Biased-OS-NO1	target	competitor	target	-	0.74	0.63
Symmetrical-SO-NO1	-	target	target	-	0.94	0.75
Symmetrical-PS-NO1	-	competitor	competitor	target	0.76	0.38
Symmetrical-OS-NO1	-	competitor	target	-	0.62	0.50
Nonreversible-SO-NO2	competitor	target	target	-	0.85	0.63
Nonreversible-PS-NO2	competitor	competitor	competitor	target	0.77	0.25
Nonreversible-OS-NO2	competitor	competitor	target	-	0.54	0.38
Biased-SO-NO2	competitor	target	target	-	0.91	0.63
Biased-PS-NO2	competitor	competitor	competitor	target	0.76	0.25
Biased-OS-NO2	competitor	competitor	target	-	0.61	0.38
Symmetrical-SO-NO2	-	target	target	-	0.90	0.75
Symmetrical-PS-NO2	-	competitor	competitor	target	0.81	0.38
Symmetrical-OS-NO2	-	competitor	target	-	0.61	0.50

Table 20. Cue-Match Status of Items from all Conditions in the data. when the Probe Requests the Retrieval of the Patient.

Experimental Condition (Agent Probe)	Syntactic Cues				Accuracy	Ratio
	Plausibility	Position	Function	Category		
Nonreversible-SO-NO1	target	target	target	-	0.91	0.88
Nonreversible-PS-NO1	target	competitor	competitor	target	0.91	0.67
Nonreversible-OS-NO1	target	competitor	target	-	0.84	0.63
Biased-SO-NO1	target	target	target	-	0.96	0.88
Biased-PS-NO1	target	competitor	competitor	target	0.93	0.67
Biased-OS-NO1	target	competitor	target	-	0.77	0.63
Symmetrical-SO-NO1	-	target	target	-	0.94	0.75
Symmetrical-PS-NO1	-	competitor	competitor	target	0.87	0.50
Symmetrical-OS-NO1	-	competitor	target	-	0.60	0.50
Nonreversible-SO-NO2	competitor	target	target	-	0.77	0.63
Nonreversible-PS-NO2	competitor	competitor	competitor	target	0.88	0.33
Nonreversible-OS-NO2	competitor	competitor	target	-	0.60	0.38
Biased-SO-NO2	competitor	target	target	-	0.91	0.63
Biased-PS-NO2	competitor	competitor	competitor	target	0.77	0.33
Biased-OS-NO2	competitor	competitor	target	-	0.56	0.38
Symmetrical-SO-NO2	-	target	target	-	0.89	0.75
Symmetrical-PS-NO2	-	competitor	competitor	target	0.86	0.50
Symmetrical-OS-NO2	-	competitor	target	-	0.63	0.50

Table 21. Estimated Values for each of the Free Parameters when the Model in Figure 10 with the Ratios Listed in is Fit to the Data

Parameter	Meaning of the Parameter	Estimated Value
c	Probability of comprehension for SO and Passive sentences	~ 0.95
c_{OS}	Probability of comprehension for OS sentences	~ 0.85
S_{bias}	Probability of a biased NO2 item matching the world knowledge of the comprehender	~ 0.43
g_p	Probability of guessing 'plausible'	~ 0.36

When the function in (15) is evaluated with the parameter values in Table 21 and the ratios listed in Table 19 and Table 20, the predicted percentage of correct agent naming and patient naming responses for each condition are as shown in Figure 11 and Figure 12, respectively, and the predicted percentage of 'plausible' responses to all conditions are as shown in Figure 13. As can be seen in Figure 11 and Figure 12, the model in Figure 10 performs poorly overall in predicting the agent/patient naming results. Although the relationship between the percentage of correct agent/patient naming responses to SO and OS conditions, that is, more correct responses to SO conditions than OS conditions across all meaning levels, is predicted by the model, the model for agent naming predicts a percentage of correct responses to passive conditions across all meaning levels so low that the predicted percentage of correct responses to OS conditions is higher than the passive conditions. Moreover, the model also does poorly in predicting the absolute percentages of correct responses to all conditions. However, it can also be seen in Figure 12 that, for patient naming, the relationship between the percentage of correct responses to each of the structure levels in the results; where SO conditions have the highest percentages, passive conditions the next highest, and OS conditions the

lowest; is predicted by the model, at least for the NO1 conditions and the symmetrical conditions in both noun orders.

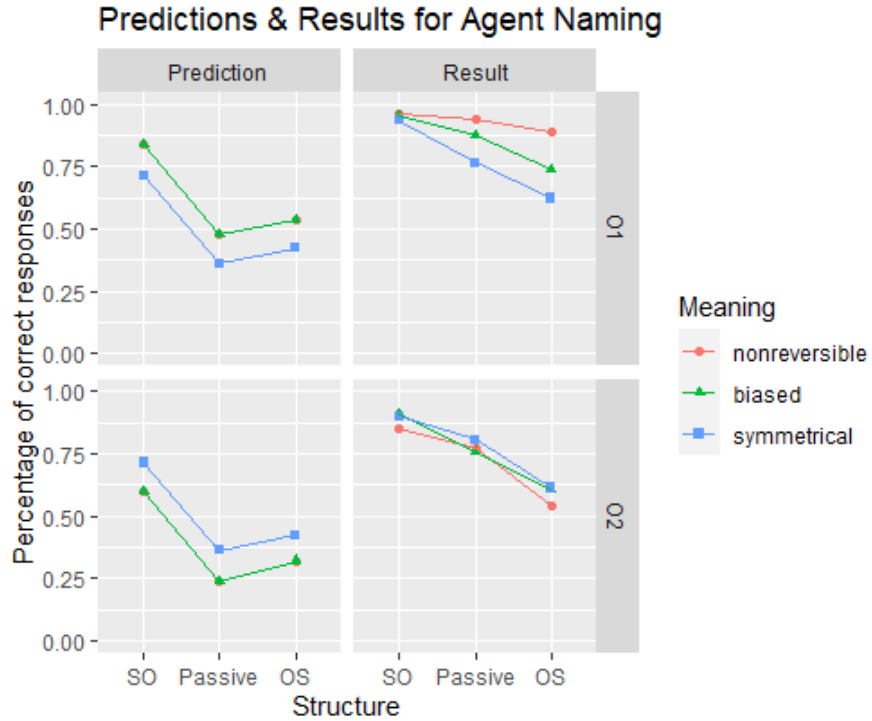


Figure 11. The plot on the left shows the percentage of correct responses to agent naming as predicted by the model in Figure 10 with the parameter values shown in Table 21 and the cue-match ratios shown in Table 19. The plot on the right shows the percentage of correct responses to agent naming in the data. The percentage of correct responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

The technical reason why the model predictions for agent/patient naming are as shown in Figure 11 and Figure 12 becomes clear when the probability function in (15) and our calculation formula for the cue-match ratios for all conditions is considered. The probability function in (15), as explained before, can be reduced to the term ‘ $(c \times r)$ ’, because both of the cognitive states of World-Knowledge Match and World-Knowledge Mismatch can lead to correct retrieval of the agent and the patient. Moreover, as also mentioned before, in this first trial of the model, the value for the parameter r for each condition is equal to the cue-match ratio for that

condition as shown in Table 19 and Table 20. Therefore, the model's predictions for the percentage of correct responses to agent/patient naming are the probability of comprehension multiplied by that condition's cue-match ration, which leaves little room for the optimization function to adjust the parameters so that the predictions match the results more closely, as the only parameters it can estimate are the parameters c and c_{os} , which must also be adjusted to match the results from the plausibility judgment data. With these considerations in mind, it also becomes clear that the reason why the model is able to predict the relationship between the percentage of correct responses across structure levels for the NO1 conditions in the patient naming data is the fact that the cue-match ratios for NO1 passive conditions in patient naming are slightly higher than those for the NO1 OS conditions. This is because, as mentioned before, the Function cue for passive conditions in patient naming is not considered among the available cues, due to the assumption that the syntactic position for the patient thematic role in passive sentences is unoccupied or occupied only by the trace of the noun which moved to the syntactic position of the agent.

As can be seen from Figure 13, the model in Figure 10, with the parameter values in Table 21, also performs somewhat worse than our final model of the plausibility judgment data which is discussed in the previous section of this chapter. The model in Figure 10, fails to predict the increase in the percentage of 'plausible' responses to the NO2 biased OS condition when compared to the SO and passive conditions of the same noun order and meaning levels, and it underestimates the magnitude of the increase in the same regard for the NO2 nonreversible OS condition when compared to the SO and passive conditions of the same noun order and meaning levels.

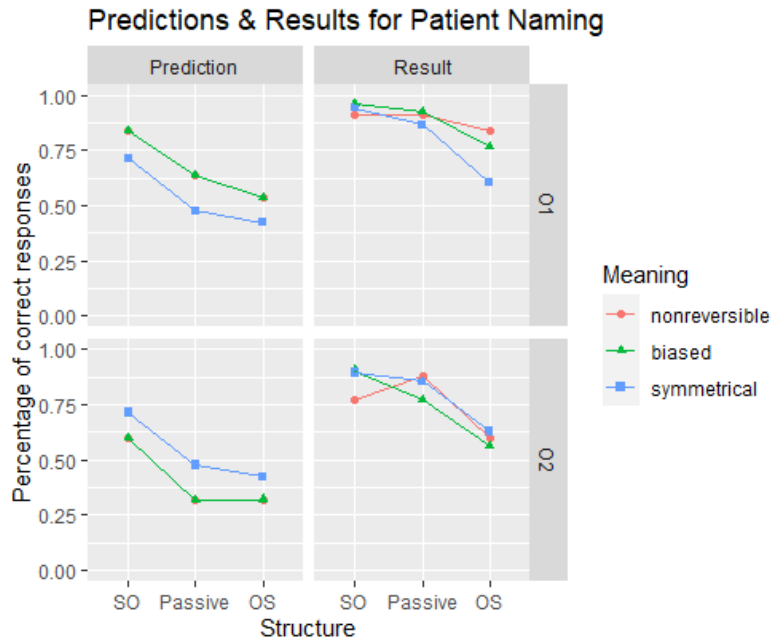


Figure 12. The plot on the left shows the percentage of correct responses to patient naming as predicted by the model in Figure 10 with the parameter values shown in Table 21 and the cue-match ratios shown in Table 12. The plot on the right shows the percentage of correct responses to patient naming in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

The technical reason why this is so has also to do with the model being fit to both the plausibility judgment and the agent/patient naming data. As was mentioned before, the optimization function can only adjust the parameters c and c_{OS} while fitting the model to the agent/patient naming data because equation that returns the predictions for agent/patient naming only features these two parameters and the fixed values of cue-match ratios, and since the parameter c_{OS} is special to the OS conditions, the optimization function must have adjusted this parameter so that the predicted percentage of correct responses to agent/patient naming are closer to what the data shows.

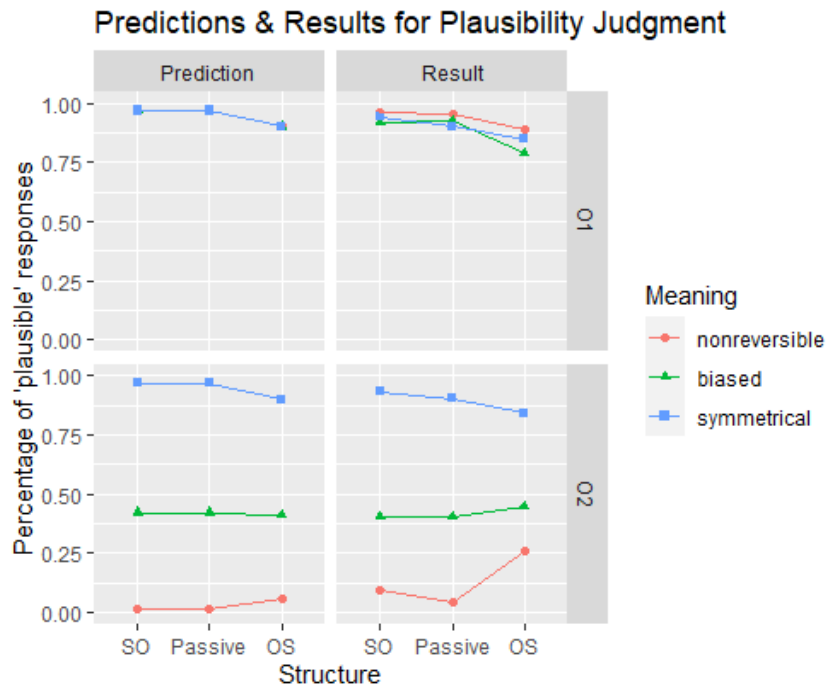


Figure 13. The plot on the left shows the percentage of ‘plausible’ responses as predicted by the model in Figure 10 with the parameter values shown in Table 21. The plot on the right shows the percentage of ‘plausible’ responses in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

The same did not happen for the passive conditions in agent naming because the equation that returns the predictions for the passive conditions use the parameter c , which, despite being estimated by the optimization function to be as high as 0.95 as shown in Table 21, cannot, in a way, bridge the gap between the cue-match ratios for the passive and OS conditions in agent naming, since those for the passive conditions are considerably lower than the ones for the OS conditions. The adjustment of the parameter c_{os} to the agent/patient naming data, on the other hand, has caused it to be estimated to be as high as 0.85, which is much higher than the estimated value of the same parameter for our final model of plausibility judgment discussed in the previous section of this chapter, that is 0.61. As a result, the model underestimated the magnitude of increase in the percentage of ‘plausible’ responses

to NO2 nonreversible OS condition and failed to predict the increase in the same regard for NO2 biased OS condition.

4.2.2 Retrieval account model 6

To address the issue of the predicted percentage of correct responses for passive conditions in agent naming being lower than that for the OS conditions in the first version of the model in Figure 10, we took into account the suggestion of Bader and Meng (2018) that the visibility of the preposition ‘by’ in the prepositional phrase which features the semantic agent of a passive sentence could have allowed the Category cue to be more distinctive. This suggestion can also be understood in terms of cue-weights under the cue-based retrieval theory of sentence comprehension, that is, it can be said that the category cue is weighted more heavily than the other three cues, which increases the contribution of a match with the target noun to the probability of its retrieval (Van Dyke, J. A., & Lewis, 2003; Lewis, Vasishth, & Van Dyke, 2006; Van Dyke & Johns, 2012). Therefore, for the next version of the model in Figure 10, instead of using the cue-match ratio calculation method that we used in the first version, we coded the cue-match status of each condition into the model with regard to the cue-match statuses shown in Table 19, so that we can integrate a linear model into the model in Figure 10 that provides an estimation of the r parameter by estimating the weights of each of the four cues.

The reason why we use the cue-match statuses shown in Table 19 and not those in Table 20 is that for the rest of the versions of the model in Figure 10, we decided to fit the model to the plausibility judgment and the agent naming data only because, as mentioned before, the total number of available cues are different for passive sentences when the probe requests the retrieval of the patient, and this adds

too much complexity to the modelling procedure to be in the scope of this thesis, and because we think that a model that provides a good account of the agent naming data should provide sufficient insight about the processing of noncanonical word order sentences for this thesis.

Table 22. Cue-Match Status of Items from all Conditions in the data. as they are Coded for our Second Version of the Model in Figure 10 when the Probe Requests the Retrieval of an Agent

Experimental Condition (Agent Probe)	Syntactic Cues			
	Plausibility	Position	Function	Category
Nonreversible-SO-NO1	1	1	1	0.5
Nonreversible-PS-NO1	1	0	0	1
Nonreversible-OS-NO1	1	0	1	0.5
Biased-SO-NO1	1	1	1	0.5
Biased-PS-NO1	1	0	0	1
Biased-OS-NO1	1	0	1	0.5
Symmetrical-SO-NO1	0.5	1	1	0.5
Symmetrical-PS-NO1	0.5	0	0	1
Symmetrical-OS-NO1	0.5	0	1	0.5
Nonreversible-SO-NO2	0	1	1	0.5
Nonreversible-PS-NO2	0	0	0	1
Nonreversible-OS-NO2	0	0	1	0.5
Biased-SO-NO2	0	1	1	0.5
Biased-PS-NO2	0	0	0	1
Biased-OS-NO2	0	0	1	0.5
Symmetrical-SO-NO2	0.5	1	1	0.5
Symmetrical-PS-NO2	0.5	0	0	1
Symmetrical-OS-NO2	0.5	0	1	0.5

We chose to code the cue-match statuses for the conditions from the first experiment of Meng and Bader (2021) as shown in Table 22. The value ‘0.5’ was used to indicate that a cue matches both nouns because, in this way, the assumption of the cue-based retrieval theory of sentence comprehension that when a retrieval cue matches multiple items in memory, interference occurs and so the contribution of the said cue to the total activation of an item reduces (e.g., Lewis, Vasishth, & Van Dyke, 2006) will be implemented in a basic way in our model. Therefore, this new version of the model assumes that there are four cues that contribute to the

probability of retrieval of the agent in the agent naming task that each have different weights and so contribute in different degrees to the probability of retrieval, and it also assumes that when a cue matches both nouns, the contribution of the cue to the probability of retrieval is halved.

(16) $\Pr(\text{Correct Agent Naming}) = (c \times r)$ where;

$$\begin{aligned} r_{condition} = & w_{\text{plausibility}} \times (\text{plausibility cue-match status}) + \\ & w_{\text{position}} \times (\text{position cue-match status}) + \\ & w_{\text{function}} \times (\text{function cue-match status}) + \\ & w_{\text{category}} \times (\text{category cue-match status}) \end{aligned}$$

To estimate weights for each of the four cues, we integrated the function in (16) to the model in Figure 10. In this way, the parameter r , which is the probability of successful retrieval of the agent, will be different for each of the conditions, and will be determined by the sum of the match status of the cue multiplied by the weight of that cue as shown in (16). However, because r is a probability, it must be constrained between the values 0 and 1. Therefore, while estimating the weights, we must consider them relative to each other, and so the sum of all weight terms must be 1. In order to achieve this, we used the method of unit simplex inverse transform (Betancourt, 2012), whereby the values of a vector with ‘ n ’ unconstrained values are mapped by a function according to their proportions into values between 0 and 1, and the sum of these now-constrained values are subtracted from 1 to obtain an additional value, the entire procedure of which results in a vector of ‘ $n+1$ ’ values, the sum of which is 1. We coded this function in R (R Core Team, 2021) and integrated it into the likelihood function which was used to estimate parameters. Inside the likelihood function, the parameters which represented the weights for the plausibility, $w_{\text{plausibility}}$, position, w_{position} , and function cues, w_{function} , were inserted into the unit simplex

inverse transformation function to obtain estimates of all of the weight parameters.

The results of the estimation are as listed in Table 23.

Table 23. Estimated Values for each of the Free Parameters when the Model in Figure 10 with the Function in (16) is Fit to the Data

Parameter	Meaning of the Parameter	Estimated Value
c	Probability of comprehension for SO and Passive sentences	~0.96
c_{OS}	Probability of comprehension for OS sentences	~0.81
S_{bias}	Probability of a biased NO2 item matching the world knowledge of the comprehender	~0.40
g_p	Probability of guessing 'plausible'	~0.34
$w_{plausibility}$	Relative weight of the Plausibility Cue	~0.027
$w_{position}$	Relative weight of the Position Cue	~0.012
$w_{function}$	Relative weight of the Function Cue	~0.316
$w_{category}$	Relative weight of the Category Cue	~0.643

By looking at Table 23 we can see that the estimated weights for the plausibility, $w_{plausibility}$, and the position cues, $w_{position}$, are very close to 0, while the estimated weights for the function, $w_{function}$, and category cues, $w_{category}$, are quite high, with the category cue estimated to have the highest weight. This does not look right from the perspective of Bader and Meng's (2018) suggestion regarding the cues that contribute to the probability of retrieval of the agent, and it is apparent from Figure 15 that the model does poorly at predicting the absolute percentages of correct responses to all conditions from the first experiment of Meng and Bader (2021). However, this model is better than the previous model, which used cue-match ratios to determine the value of r for each condition in that it does not underestimate the percentage of correct responses to the passive conditions in agent naming, hence it can predict the relationship between the percentages of correct responses to the structure levels where the percentage of correct responses to the OS conditions are lower than that for the SO and passive conditions.

The reason why this model does better in predicting the relationship between the percentage of correct responses to the structure levels is that the contrasts shown in Table 22, together with the estimated weights in Table 23 have resulted in values of the parameter r that have allowed the probability equation in (16) to produce results that better match the data. It can also be said that the additional weight parameters have contributed to the flexibility of the model. However, the model still does poorly in both reflecting the suggestion of Bader and Meng (2018) about the cues that contribute to the probability of retrieval of the agent and predicting the absolute percentages of correct responses to agent naming for all conditions.

To understand why this is so, we can examine the cue-match status contrasts between the conditions and consider why the estimated weights for the Function and Category cues are so high while the estimated weights for the Plausibility and Position cues are near zero as shown in Table 23. Firstly, the cue-match status or contrast value for the Plausibility cue in all NO1 conditions is 1, and in all NO2 conditions is 0, except for the symmetrical conditions in these two noun orders. Therefore, we can say that this cue distinguishes between the noun order levels.

However, if we look at the agent naming results from the first experiment of Meng and Bader (2021) shown in Figure 14, we can see that the factor noun order does not cause any significant change in the percentage of correct responses obtained in agent naming, and a main effect of noun order was not reported for their first experiment by Meng and Bader. Similarly, the cue-match status or contrast value for the Position cue in all SO conditions is 1, while it is 0 for all conditions in the other two structure levels, but neither such an effect of structure that singles out the SO sentences was reported, nor such an effect is visible in the experiment results shown in Figure 14. On the contrary, the Function and Category cues provide a contrast

between the SO and OS conditions on one hand, and the passive conditions on the other. All Function cue-match status values for SO and OS sentences are 1, while those for passive sentences are 0, and all Category cue-match status values for SO and OS sentences are 0.5, while those for passive sentences are 1. Furthermore, we know from our first trial with the model in Figure 10, where we used cue-match ratios, that the number of cues that match the retrieval cues for passive conditions is the lowest and so a model that assumes that the probability of retrieval of the agent in agent naming is determined by the number of cues that match the retrieval cues in that condition underestimates the probability of retrieval of the agent in passive conditions. All of these facts together, reveal why the estimated values for the weights of the Plausibility and the Position cues were near zero, while those for the Function and Category cues were so high. The latter two cues provide a contrast between SO and OS conditions on one hand, and passive conditions on the other that, with the right weight estimations for these two cues, can allow the model to compensate for the low number of matching cues in the passive conditions, so that the model's predictions better align with the data that it is fit to.

There is also a technical reason why the model does poorly at predicting the absolute percentages of correct responses to all conditions. To understand this, we need to consider the unit simplex inverse transformation method used to constrain the relative weights of the cues between 0 and 1, and the cue-match status values. Because the equation in (16) which returns the probability of retrieval for a condition is the sum of the match status of the cues multiplied by the weight of the corresponding cue, and because the estimated values for the cue-weights all add up to 1 and can at most be multiplied by 1 in the equation, the result of which is the exact weight of that cue, it is impossible for the equation to result in values for r that

are closer to 1 than what we got from the parameter estimation procedure of this model. The discovery of the technical reason why there is such an issue provides us with a way to improve the model. If there was another parameter in the equation that returns the probability of retrieval for a condition which had the same value for all conditions and which was not multiplied by any weight values, it would be possible for the equation to return values that are closer to 1.

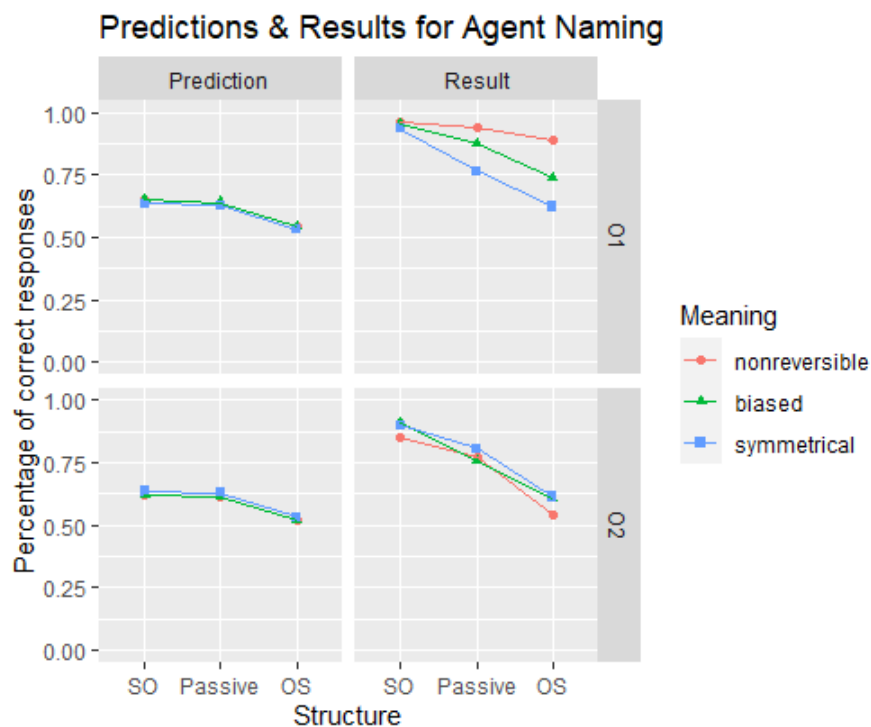


Figure 14. The plot on the left shows the percentage of correct responses to agent naming as predicted by the model in Figure 10 with the parameter values shown in Table 23. The plot on the right shows the percentage of correct responses to agent naming in the data. The percentage of correct responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

4.2.3 Retrieval account model 7

To address the issue with the second version of the model in Figure 10 that it is impossible for the equation that returns the probability of retrieval for a condition to return values closer to 1, and possibly to also address the issue with the same model

that the estimated cue weights for the Plausibility and Position cues are near zero, we thought of adding a ‘base activation’ parameter to the equation that returns the probability of retrieval for a condition in a third version of the model in Figure 10, such that the probability function will now be as shown in (17).

(17) $\text{Pr}(\text{Correct Agent Naming}) = (c \times r)$ where;

$$r_{\text{condition}} = r_{\text{base}} + w_{\text{plausibility}} \times (\text{plausibility cue-match status}) + w_{\text{position}} \times (\text{position cue-match status}) + w_{\text{function}} \times (\text{function cue-match status}) + w_{\text{category}} \times (\text{category cue-match status})$$

The assumption that there is a base amount of activation that contributes to the probability of retrieval of the agent that is the same for every condition is also in-line with the cue-based retrieval theory of sentence comprehension (e.g., Lewis, Vasishth, & Van Dyke, 2006) in that it is assumed under this theory that the total activation of a memory item is the sum of that item’s base activation and spreading activation which is determined by a function of the cue-weights (Dotlačil, 2021). Therefore, the third version of the model, added to the assumptions of the previous model, assumes that, added to the activation that each cue contributes to the probability of retrieval of the agent, there is a base activation value for the agent noun in every condition that is the same for all conditions. Apart from this new assumption, every other property of this version of the model is the same as the previous version of it.

When the model in Figure 10 is fit to the agent naming data from the first experiment of Meng and Bader (2021) with the cue-match status values in Table 22, and the function that returns the value for the parameter r in (17), the estimated

values for the parameters are as shown in Table 24. When these values are plugged into the probability functions for the same model, the predicted percentage of correct responses in agent naming are as shown in Figure 15, and the predicted percentage of ‘plausible’ responses in plausibility judgment are as shown in Figure 16.

Table 24. Estimated Values for each of the Free Parameters when the Model in Figure 10 with the Function in (17) is Fit to the Data

Parameter	Meaning of the Parameter	Estimated Value
c	Probability of comprehension for SO and Passive sentences	~0.90
c_{OS}	Probability of comprehension for OS sentences	~0.69
S_{bias}	Probability of a biased NO2 item matching the world knowledge of the comprehender	~0.38
g_p	Probability of guessing ‘plausible’	~0.48
r_{base}	Base activation value	~0.874
$w_{plausibility}$	Relative weight of the Plausibility Cue	~0.036
$w_{position}$	Relative weight of the Position Cue	~0.033
$w_{function}$	Relative weight of the Function Cue	~0.032
$w_{category}$	Relative weight of the Category Cue	~0.023

It is apparent by examining Figure 15 that this version of the model performs better at predicting the absolute percentages of correct responses to agent naming for all conditions than the previous version of the model. Moreover, the relationship between the structure levels in terms of the percentage of correct responses to agent naming is accurately predicted by the model, with the predicted percentage of correct responses to the SO conditions being higher than that for the passive conditions, and the percentage of correct responses to the passive conditions being higher than that for the OS conditions. However, this version of the model is still unable to predict important facts about the data from the first experiment of Meng and Bader (2021). Meng and Bader found an interaction between the structure and meaning factors, as well as a three-way interaction between structure, meaning and noun order, which is also apparent when the agent naming results shown in Figure 15 are examined. The

two-way interaction between structure and meaning shows itself in the NO1 condition, as the percentage of correct responses decrease more drastically across structure levels from SO to OS for symmetrical conditions than for biased conditions, and they decrease more drastically from SO to OS for biased conditions than for nonreversible conditions in this noun order. The three-way interaction, on the other hand, is revealed when we consider that this effect is not present in NO2 conditions.

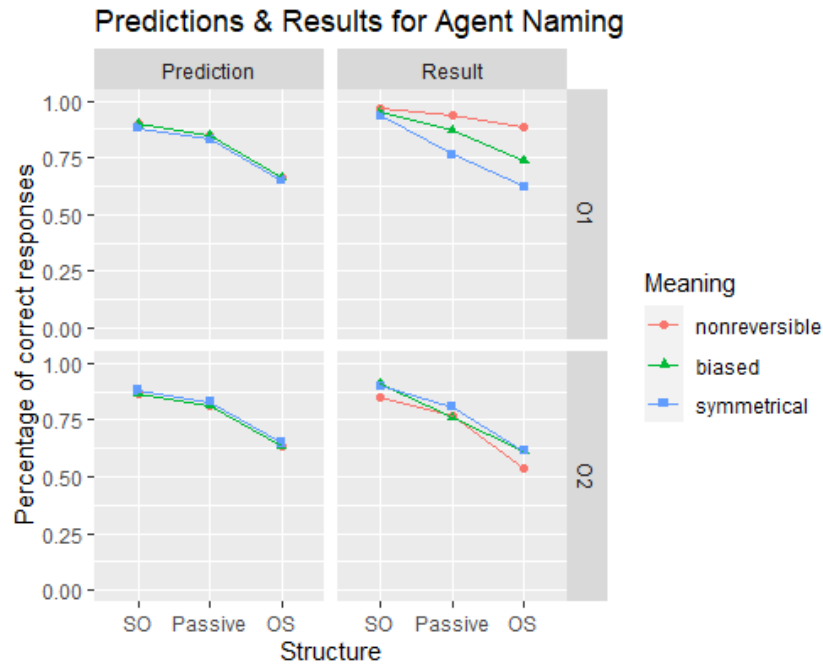


Figure 15. The plot on the left shows the percentage of correct responses to agent naming as predicted by the model in Figure 10 with the parameter values shown in Table 24. The plot on the right shows the percentage of correct responses to agent naming in the data. The percentage of correct responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

Regarding the plausibility judgment predictions, it is clear from Figure 16 that this version of the model does not fail, like the previous models, in predicting the increases in the percentage of ‘plausible’ responses to OS sentences in NO2 nonreversible and biased conditions, although the fit is visibly worse for the NO2

nonreversible OS condition than that of our final model of plausibility judgment data discussed in the previous section of this chapter.

We can begin assessing the performance of this version of the model by looking at the parameter estimates in Table 24. Although this version of the model does notably better than the previous versions, the parameter estimates paint a picture that is different from what is suggested by Bader and Meng (2018). Bader and Meng (2018) suggested that the Category cue could have a special status among the set of cues that they suggested due to the visibility of the preposition ‘by’ which could have allowed this cue to be more distinctive. However, the parameter estimates from this version of the model show that the relationship between the structure levels in terms of the percentage of correct responses to agent naming can be predicted even when it is assumed that the Category cue has the lowest relative weight among the same set of suggested cues when the cue-match status values in Table 22 are assumed for all conditions as well. To explain the technical reason why this is so, we must also refer to the parameter estimates of the second version of the model. We discussed that the reason why the estimates for the weight of the Function and Category cues in the second version of the model were much higher than the other two cues, which were estimated to be near zero, is that the number of cues that match the retrieval cues for passive conditions is the lowest, and that this situation must be compensated by adjusting the weight of the cues that distinguish the passive conditions from the other structure levels. It seems, from the parameter estimates for the current version of the model shown in Table 24 that an adjustment to the base activation parameter is sufficient to carry out the same compensation. The technical reason why this is so also requires that we go back to the discussion about the parameter c_{os} that we had under the first version of the model. Because for all versions of the model, the

equation that returns the probability of a correct response to agent naming in OS conditions is $(c_{os} \times r)$, and it is $(c \times r)$ for the rest of the structure levels, when c_{os} is sufficiently lower than c , which is required for the model predictions to closely match the plausibility judgment data as discussed in the previous section of this chapter, the probability of a correct agent naming response can be lower for OS conditions than for passive conditions, the occurrence of which compensates for the little number of cues that match the retrieval cues for passive conditions by pulling the predicted percentage of correct responses to OS conditions even lower despite the greater number of cues that match the retrieval cues for the OS conditions. As this is the technical reason why the predictions of this version of the model are as such, we must also state that all versions of the model assume that the difficulty of comprehension in the lack of an appropriate discourse context for OS sentences, which was our understanding of the suggestion made by Bader and Meng (2021) regarding the information-structural markedness of OS sentences, also reduces the probability of a correct response to agent naming for OS sentences, which was also suggested by Bader and Meng (2021), although they claimed that this contributed less to the error rate in OS conditions than misinterpretation effects in the form of agent-patient reversal.

As mentioned before, another issue with this version of the model is that it cannot predict the two-way interaction between the structure and meaning factors, and the three-way interaction between structure, meaning and noun order found by Meng and Bader (2021) in their first experiment. The technical reason why this is so can be explained through a discussion of what kinds of revisions made to the model could allow it to predict these interaction effects.

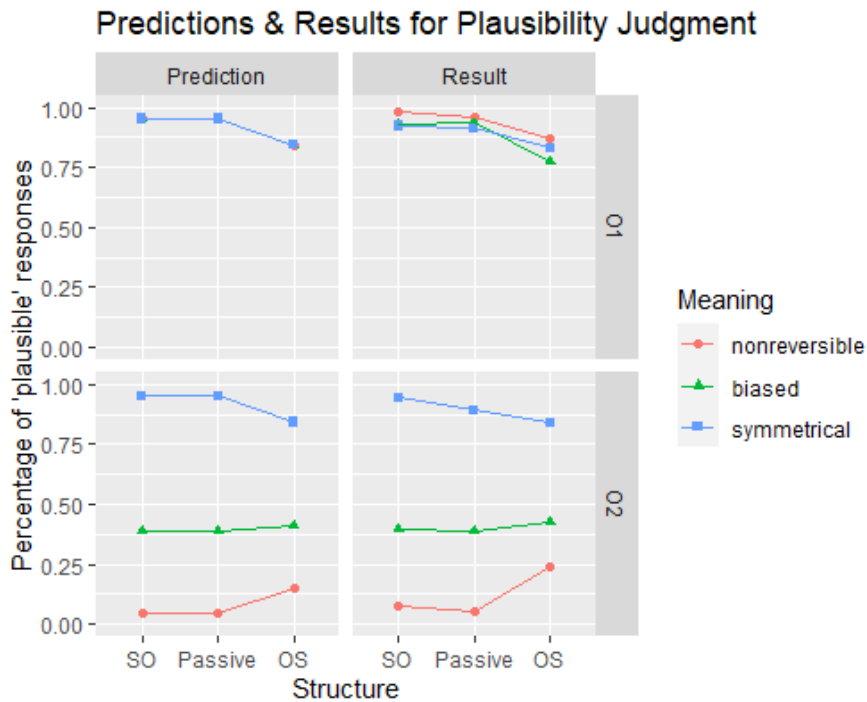


Figure 16. The plot on the left shows the percentage of ‘plausible’ responses as predicted by the model in Figure 10 with the parameter values shown in Table 24. The plot on the right shows the percentage of ‘plausible’ responses in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

The technically simple way to improve the model so that it can account for the interaction effects would be to postulate at least two new cognitive states and arrange the outcomes such that these new cognitive states would distinguish all structure and meaning levels, which would allow the model to account for the two-way interaction between the structure and meaning factors, and to postulate an additional cognitive state that further distinguishes noun order levels, which would allow the model to also account for the three-way interaction between structure, meaning and noun order. Our model actually has a parameter that distinguishes the meaning levels, that is the s parameter, which is fixed to 0 in all NO2 conditions except the symmetrical and biased conditions in this noun order and fixed to 1 in all NO1 conditions, but this parameter has no effect on the probability of a correct

response to agent naming because both the cognitive states of World-Knowledge Match and Mismatch can lead to the correct retrieval of the agent in our model as shown in Figure 10. Revising the model such that only World-Knowledge Match leads to a correct response to agent naming, added to the two new cognitive states that distinguish the structure and meaning levels that would have to be postulated in order for the model to account for the two-way interaction mentioned, can also potentially allow the model to account for the three-way interaction between structure, meaning and noun order. However, none of these solutions are theoretically motivated, and since each new cognitive state would introduce a new parameter into the model, the model can, as a result, become too flexible, and potentially able to account for almost any kind of data, the occurrence of which would eliminate the benefits of the MPT modelling procedure in terms of the insights that it can provide about the cognitive processes being studied.

Another way to improve the model so that it can predict the interaction effects would be to introduce interaction terms into the equation that returns for each of the conditions the probability of successful retrieval of the agent noun, the parameter r . The equation that returns the probability of successful retrieval of the agent noun is a linear equation, hence it is impossible for it to allow the model to predict the interaction effects on its own. In other words, only a nonlinear function can allow the model to predict the interaction effects without any change to the number of cognitive states assumed in the model. Parker (2019) found evidence for nonlinear combination of retrieval cues and suggested a nonlinear model of cue-based retrieval based on evidence found in a study of antecedent-reflexive dependencies in English. Therefore, implementing such an idea in our model may also be considered to be in-line with the cue-based retrieval structure that we

assumed in it. However, such an implementation is beyond the scope of this thesis and will be left for future work to explore.

In conclusion, this final MPT model of both the plausibility judgment and the agent naming data from the Meng and Bader (2021) study is sufficient for the purpose of this thesis. The modelling procedure has revealed that what we can assume under the retrieval account about the cognitive processes that are at work in the plausibility judgment task that we discussed in the previous section of this chapter can also be assumed in a model of both the plausibility judgment and the agent naming data, provided that the model of both tasks assumes a cue-based retrieval procedure which is unaffected by whether sentence meaning matches the world-knowledge of the comprehender or not, and in which the probability of successful retrieval of the agent is determined by both the base activation level of all items in memory and the degree of match between the retrieval cues that vary in their weights and the item's features. The modelling procedure has also revealed that an MPT model of agent naming for noncanonical word order sentences that assumes that the probability of successful retrieval of the agent noun for OS sentences is negatively affected by the difficulty of comprehension for these sentences in the lack of an appropriate discourse context, can predict the agent naming and the plausibility judgment results obtained in the first experiment of Meng and Bader (2021). Finally, the modelling procedure has revealed that, in order for an MPT model that is faithful to the assumptions of the retrieval account to predict the interaction effects found between the structure, meaning and noun order levels in the first experiment of Meng and Bader, it must be assumed that the retrieval cues combine in a nonlinear fashion.

CHAPTER 5

MPT MODELS UNDER THE PARSING ACCOUNT

5.1 Joint MPT models of agent/patient naming and plausibility judgement task responses

Our MPT models of the parsing account will be fundamentally different from those we created for the retrieval account, due to the assumption that an error in both the agent/patient naming task and the plausibility judgment task reflects misinterpretation errors under the parsing account, instead of a problem with the retrieval of a memory item as the retrieval account suggests. For example, ‘comprehension’, or the state of a comprehender having constructed an interpretation that is faithful to the linguistic input, was a cognitive state suggested in all of our models which followed the suggestions of the retrieval account, a cognitive state which was followed by other cognitive states in the model, including the retrieval of the agent and the patient. However, we cannot suggest such a cognitive state for a model that follows the suggestions of the parsing account because under the parsing account, ‘comprehension’, in the sense that any interpretation having been created which could either be faithful or unfaithful to the linguistic input unlike the retrieval account suggests, is not a cognitive state that is followed by other such states or event, but rather the final outcome of a series of cognitive events after which a comprehender responds to the task. Hence, the retrieval account, because it suggests the outcomes from a task are a result of cognitive processes following comprehension, is also referred to as the postinterpretive account, while the parsing account suggests that all errors are misinterpretation errors.

However, our suggestion, that is neither suggested under the retrieval nor the parsing account, that guessing, that is due to the absence of an interpretation or inattentiveness, as suggested by Logacev and Dokudan (2021), could lead to correct plausibility judgments, but never correct agent/patient naming, can still be integrated into a model that follows the suggestions of the parsing account. This is because, we believe that the suggestions of the parsing account regarding the cognitive processes, the occurrence of which are modulated by the task demands under the parsing account, that lead to either a correct interpretation or misinterpretation, do not conflict with the suggestion that a comprehender may sometimes do not pay attention the task at hand or has a perception problem with the linguistic input. For example, under the parsing account, there are two routes which work with the linguistic input simultaneously in order to create an interpretation, the results of which could either lead to a correct interpretation or misinterpretation, but for the engagement of both of these routes or them resulting in an interpretation, we believe, presupposes that the comprehender is attentive to the task at hand. Moreover, our assumption that guessing may lead to a correct plausibility judgment, but never to correct agent/patient naming, could even be argued to be more in-line with the suggestions of the parsing account. This is because in the studies that gave rise to the parsing account, a plausibility judgment task was never analyzed in a way that would influence the suggestions of the account, but instead, the assumptions of the parsing account about noncanonical word order sentence processing were formulated and adapted by Meng and Bader (2021) to the plausibility judgment task, in virtue of the fact that the parsing account follows from the more general theory of processing that is Good-Enough Processing (Christianson et al., 2001; F. Ferreira, 2003; F. Ferreira, Ferraro, & Bailey, 2002; F. Ferreira & Patson, 2007; Sanford & Sturt, 2002) which

suggests that all language comprehension is subject to its suggestions. Furthermore, in the first experiment of Meng and Bader (2021), the data from which we use to fit all MPT models discussed in this thesis, the agent/patient naming task required the participants to orally name the agent or the patient, whereas the plausibility judgment task only required the participants to choose either ‘yes’ or ‘no’, which can be argued to make the task more eligible for guessing. Therefore, in our MPT models of the parsing account, we will retain the parameter g_p , which is the probability of guessing ‘plausible’ in the lack of comprehension, or attentiveness in this case.

Before presenting our models, a review of the suggestions of the parsing account about the processing of noncanonical word order sentences is due. Firstly, under the parsing account, sentence processing occurs in a dual-route mechanism, one of which is the algorithmic route and the other, the heuristic route (F. Ferreira, 2003; Swets, Desmet, Clifton, & Ferreira, 2008; Christianson, Luke & Ferreira, 2010; Karimi & Ferreira, 2016). The heuristic route works with semantic information to rapidly create an interpretation using two simple heuristics; the NVN heuristic, which maps the first noun encountered to the agent role and the second to the patient role, and the semantic-association heuristic, which forms interpretations solely based on the plausibility of the semantic relations of the nouns to the verb of the sentence, completely ignoring syntactic structure. The algorithmic route on the other hand, works with syntactic information and always forms the correct interpretation provided that the sentence was perceived correctly. Both of these routes are engaged simultaneously when a sentence is read, but the heuristic route is quicker to produce an interpretation, while the algorithmic route, which works with more complex syntactic information, takes longer to produce an interpretation. Problems like misinterpretation or incomplete sentence representations are a result of failure in the

integration of the information produced through the two routes. Under the parsing account, the HPM is strongly task-dependent and versatile, in that the influence each of these two routes have on the final interpretation is determined by what the communicative task requires, and so, the sentence interpretations that are produced are detailed to the level that is required by the task at hand.

Implementing all of these ideas in a single MPT model of both the plausibility judgment and agent/patient naming task results from the first experiment of Meng and Bader (2021), like we did with the retrieval account, is a difficult task from an MPT modelling perspective, due to how it approaches and structures latent cognitive processes. However, as with the retrieval account, we will try to stay as faithful as possible from an MPT modelling perspective to the suggestions of the parsing account.

As per our understanding of the suggestions of the parsing account, all of the cognitive events or states suggested under this account are conceived in a way to bring about correct or incorrect responses to a task. For example, the algorithmic route, under the parsing account, is suggested as always producing the correct interpretation, and hence always resulting in the correct response to a task (e.g., Ferreira, 2003; Karimi and Ferreira, 2016). Because of the way the first experiment of Meng and Bader (2021) was designed, and because we believed we had to think in terms of correct and incorrect responses, we created 6 MPT models, shown in Figures 17-21, each modelling a set of unique conditions from the experiment that lead to the same outcomes according to how the processes were conceptualized within the model. For ease of reference, we will refer to the combination of these 6 MPT models as a single MPT model of the entire experiment, as only the combination of them would cover all of the conditions in the experiment, and as the

likelihood function that we used to fit the models to the data from the experiment uses all 6 of these MPT models for the probability functions because we wanted a model of the entire data.

Finally, we will not have a separate model, as we did with the retrieval account, for the plausibility judgment data for the parsing account, due to the fact that the parsing account of noncanonical sentence processing originated from a study of agent/patient naming data. In other words, all of the models presented in this chapter will be joint models of both the plausibility judgment and the agent/patient naming data from the first experiment of Meng and Bader (2021).

5.1.1 Parsing account model 1

All of our models of a set of conditions, shown in Figures 17-21, assume the same cognitive states since the combination of them makes our model of the entire data from the first experiment of Meng and Bader (2021). As can be seen in Figures 17-21, we assume that the first step of the MPT model is a contrast between the cognitive states of ‘Heuristic Route’ and ‘guessing’. The reason why the model begins with the heuristic route, and only after this route is engaged that the algorithmic route can be engaged, is a reflection of our understanding of the suggestion of Karimi and Ferreira (2016) that the parsing of a sentence begins with the heuristic route, which is faster than the algorithmic route, and that only if equilibrium is not achieved through the heuristic route that the algorithmic route completes creating an interpretation in their account. This suggestion is also reflected in the step after the cognitive state of ‘Heuristic Route’ is achieved in our model in that the cognitive state of ‘Heuristic Equilibrium’ contrasts with that of ‘Algorithmic Route’. Moreover, since, under the parsing account, we can talk about whether an

interpretation is the result of the semantic-association heuristic or the NVN heuristic only if equilibrium is achieved through the heuristic route, the cognitive events of ‘Semantic Association Heuristic’ and ‘NVN Heuristic’ follow the cognitive state of ‘Heuristic Route Equilibrium’ in our model, as shown in Figures 17-21. Furthermore, the cognitive state of ‘guessing’ contrasts with that of ‘Heuristic Route’ in our model because this model’s conceptualization of guessing is that it occurs as a result of inattention to the probe or the task, as discussed before.

Before the model is fit to the data, it is important that we consider the assumptions about and outcomes from the semantic-association and NVN heuristics for each model of a set of conditions, since the algorithmic route and guessing lead to the same outcomes for every model of a set of conditions, with the algorithmic route always leading to the correct plausibility judgment and agent/patient naming, and the guessing always leading to an incorrect agent/patient naming response.

We can start addressing the individual models of a set of conditions, with the model shown in Figure 17, which lists the outcomes to the NO1 Nonreversible or Biased SO conditions. Table 25 shows the items from these conditions. Firstly, As can be seen in Figure 17, the semantic-association heuristic leads to a ‘plausible’ response and the correct agent/patient naming in these conditions. This is because, the semantic-association heuristic will assign the thematic role of agent to the noun, ‘Koch’ or ‘chef’ in Table 25, in the interpretation that it creates, since it is the more plausible agent of the verb from among the two nouns; and this assignment will result in the correct agent/patient naming because this noun is the actual agent of the sentence, and in a ‘plausible’ response, because the semantic-association heuristic always builds a plausible interpretation due to its nature as suggested under the parsing account. Secondly, as can be seen in Figure 17, the NVN heuristic also leads

to a ‘plausible’ response and the correct agent/patient naming in these conditions.

This is because the NVN heuristic will assign the thematic role of agent to the first noun it encounters, which is also ‘Koch’ or ‘chef’ in Table 25, the actual agent of the sentence, in the interpretation that it creates; thus, this heuristic will also lead to a ‘plausible’ response to the plausibility judgment task and the correct agent/patient naming.

Table 25. The Items from the First Experiment of Meng and Bader (2021) from the NO1 Nonreversible or Biased SO Conditions and NO1 Nonreversible or Biased OS and Passive Conditions

Structure	Meaning	Noun Order	Example Sentence	English Translation
SO	Nonreversible	1	Der Koch hat den Topf gereinigt.	‘The chef cleaned the pan’
	Biased	1	Der Koch hat den Braten ruiniert.	‘The chef ruined the roast’
OS	Nonreversible	1	Den Topf hat der Koch gereinigt.	‘The pan, the chef cleaned.’
	Biased	1	Den Braten hat der Koch ruiniert.	‘The roast, the chef ruined’
Passive	Nonreversible	1	Der Topf wurde vom Koch gereinigt.	‘The pan, by the chef, was cleaned.’
	Biased	1	Der Braten wurde vom Koch ruiniert.	‘The roast, by the chef, was ruined.’

Figure 18 lists the outcomes to the NO1 Nonreversible or Biased OS or Passive conditions and the items from these conditions are shown in Table 25. As can be seen in Figure 18, the semantic-association heuristic leads to a ‘plausible’ response and the correct agent/patient naming in these conditions. This is because the semantic-association heuristic will assign the thematic role of agent to the noun, ‘Koch’ or ‘chef’ in Table 25, which is the actual agent of the sentence; and the interpretation that results from this thematic role assignment will be plausible as this noun is the plausible agent to the verb. The NVN heuristic, on the other hand, leads to an ‘implausible’ response as well as the incorrect agent/patient naming in these

conditions. This is because the NVN heuristic will assign the thematic role of agent to the first noun it encounters, which is ‘Topf’ or ‘pan’ in Table 25, which is not the actual agent of the sentence, and the interpretation that results from this assignment is implausible, as in the interpretation ‘the pan’, which is an inanimate noun, would be cleaning ‘the chef’, or ‘the roast’, which despite being an animate noun, is the less plausible agent, would be ruining ‘the chef’ as in Table 25.

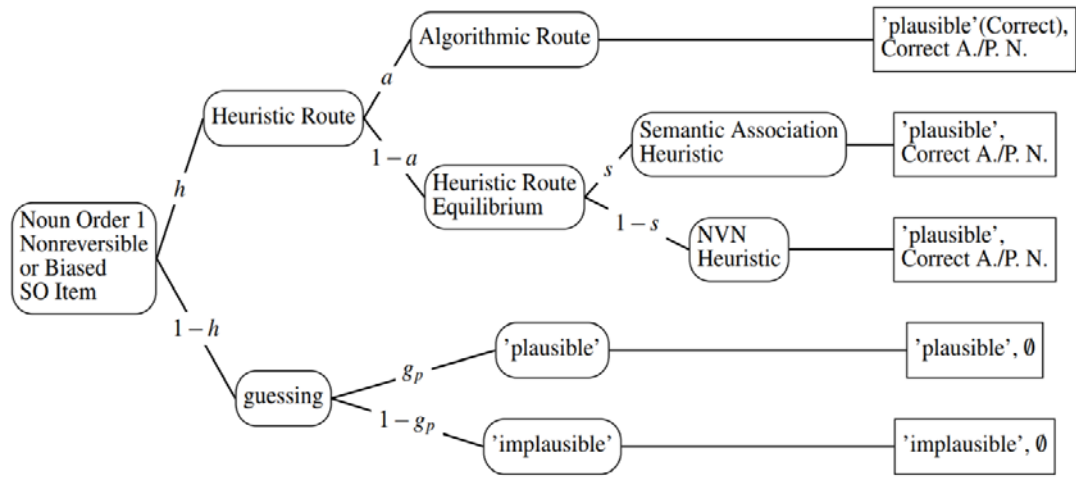


Figure 17. Our model of the agent/patient naming and plausibility task responses to the NO1 Nonreversible and Biased SO conditions in Meng and Bader (2021). ‘Correct A./P. N.’ stands for correct agent/patient naming, and the empty set symbol represents either incorrect agent/patient naming or the lack of a response to the agent/patient naming task

Figure 19 lists the outcomes to the symmetrical SO and symmetrical OS and passive conditions, and the items from these conditions are shown in Table 26. You will notice that for all symmetrical conditions, the semantic-association heuristic, along with the parameter that represents its probability of occurrence, is not on the visual representation of the model. We chose to model the conditions this way because in the symmetrical conditions, as can be seen in Table 26, both nouns are plausible agents to the verb. Hence, we cannot determine exactly, despite having to do so due to how MPT models are structured, what noun will be assigned the thematic role of agent by a heuristic that ignores the syntactic structure entirely and

creates an interpretation of the sentence based on which of the two nouns is a more plausible agent to the verb, and so, we cannot determine whether the response to the agent/patient naming task will be correct or incorrect as a result of the final interpretation being formed by the semantic-association heuristic. This is equivalent to fixing the value of the parameter s , which is the probability that it was the semantic-association heuristic which created the final interpretation of the heuristic route, to zero for the symmetrical conditions, similar to what we did with the models discussed in the retrieval account chapter.

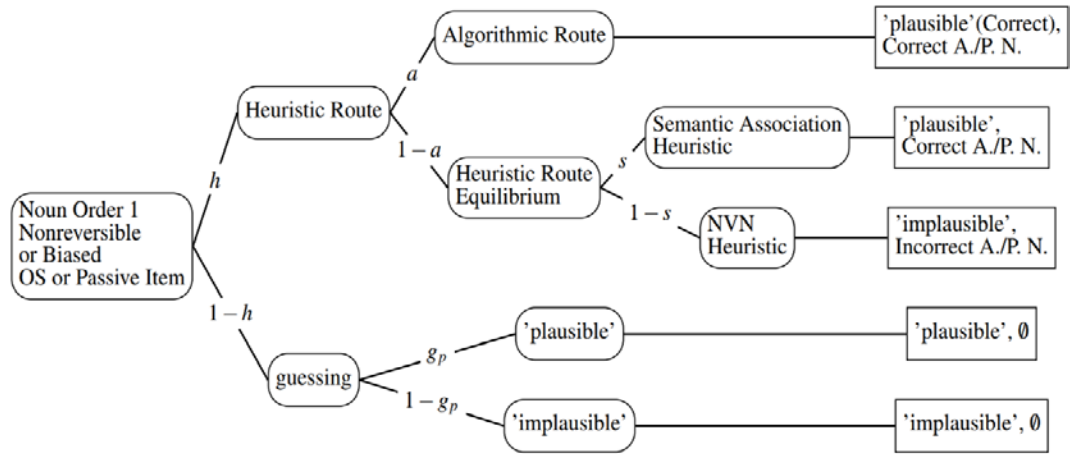


Figure 18. Our model of the agent/patient naming and plausibility task responses to the NO1 Nonreversible and Biased, OS or Passive conditions in Meng and Bader (2021). ‘Correct A./P. N.’ stands for correct agent/patient naming, and the empty set symbol represents either incorrect agent/patient naming or the lack of a response to the agent/patient naming task

When we look at the models in Figure 19, we can see that the NVN heuristic leads to a ‘plausible’ response and the correct agent/patient naming in the symmetrical SO conditions in both noun orders, while it leads to the incorrect agent/patient naming but still a ‘plausible’ response in the symmetrical OS and passive conditions in both noun orders. The former is because the NVN heuristic in symmetrical SO conditions in both noun orders assigns the thematic role of agent to the first nouns it encounters, ‘Vater’ or ‘father’ in NO1, and ‘Onkel’ or ‘uncle’ in

NO2 in Table 26, which are the actual agents of the corresponding sentences in these conditions. On the other hand, the reason why the NVN heuristic leads to the incorrect agent/patient naming but still a ‘plausible’ response in the symmetrical OS and passive conditions is that the NVN heuristic will wrongly assign the thematic role of agent to the first nouns it encounters, ‘Onkel’ or ‘uncle’ in NO1, and ‘Vater’ or ‘father’ in NO2 in Table 26, which are not the actual agents of the corresponding sentences, but this assignment still results in a plausible interpretation due to both nouns being plausible agents to the verb in the symmetrical conditions.

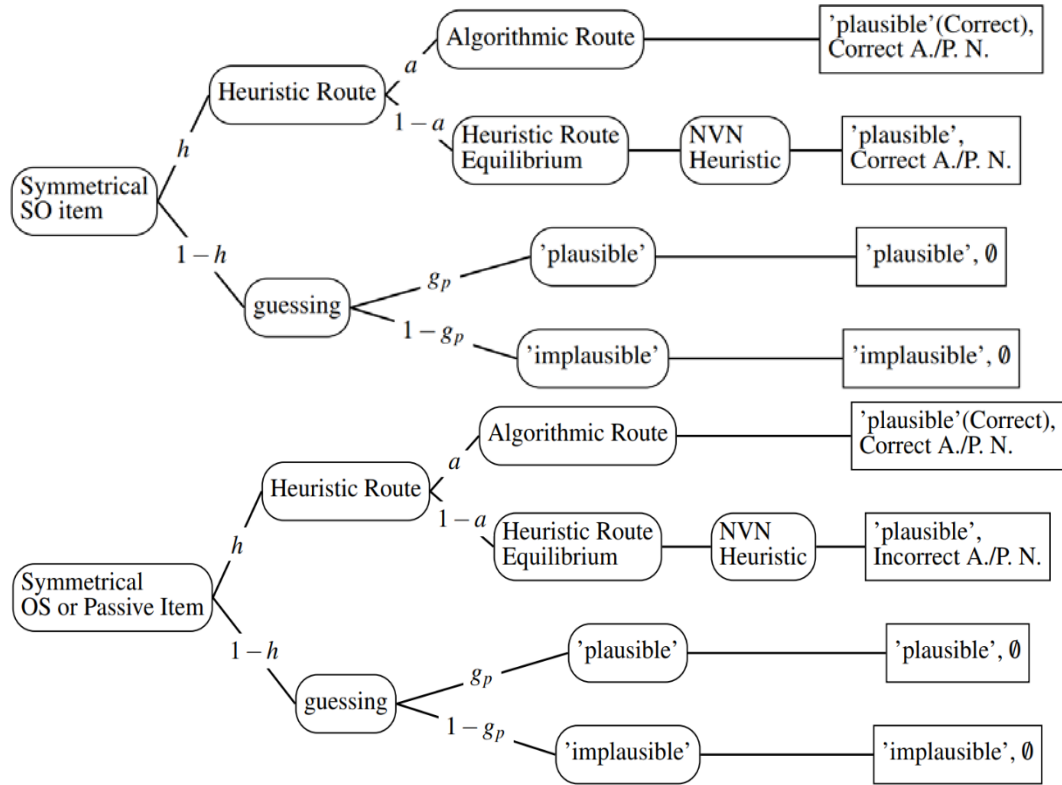


Figure 19. Our models of the agent/patient naming and plausibility task responses to the symmetrical SO conditions and the symmetrical OS and passive conditions in both noun orders in Meng and Bader (2021). ‘Correct A./P. N.’ stands for correct agent/patient naming, and the empty set symbol represents either incorrect agent/patient naming or the lack of a response to the agent/patient naming task

Table 26. The Items from the First Experiment of Meng and Bader (2021) from the NO1 Symmetrical SO, OS and Passive Conditions

Structure	Meaning	Noun Order	Example Sentence	English Translation
SO	Symmetrical	1	Der Vater hat den Onkel umarmt.	‘The father hugged the uncle’
OS	Symmetrical	1	Den Onkel hat der Vater umarmt.	‘The uncle, the father hugged.’
Passive	Symmetrical	1	Der Onkel wurde vom Vater umarmt.	‘The uncle, by the father, was hugged.’
SO	Symmetrical	2	Der Onkel hat den Vater umarmt.	‘The uncle hugged the father’
OS	Symmetrical	2	Der Vater wurde vom Onkel umarmt.	‘The father, by the uncle, was hugged.’
Passive	Symmetrical	2	Den Vater hat der Onkel umarmt.	‘The father, the uncle hugged.’

Figure 20 lists the outcomes to the NO2 Nonreversible and Biased SO conditions and the items from these conditions are shown in Table 27. As can be seen from Figure 20, the semantic-association heuristic leads to a ‘plausible’ response and the incorrect agent/patient naming, while the NVN heuristic leads to an ‘implausible’ response and the correct agent/patient naming in these conditions. The reason why the semantic-association heuristic leads to the incorrect agent/patient naming is that this heuristic will assign the agent thematic role to the plausible agent of the verb, ‘Koch’ or ‘chef’ in Table 27, which is not the actual agent of the sentence; and the reason why the semantic-association heuristic leads to a ‘plausible’ response again is that this heuristic will always construct a plausible interpretation of a sentence. The NVN heuristic, on the other hand, will assign the agent thematic role to the first nouns it encounters, ‘Topf’ or ‘pan’ in nonreversible conditions, and ‘Braten’ or ‘roast’ in biased conditions in Table 27, which are the actual agents of the corresponding sentences, thus creating an interpretation which results in a correct

agent/patient naming response and an ‘implausible’ response to the plausibility judgment task.

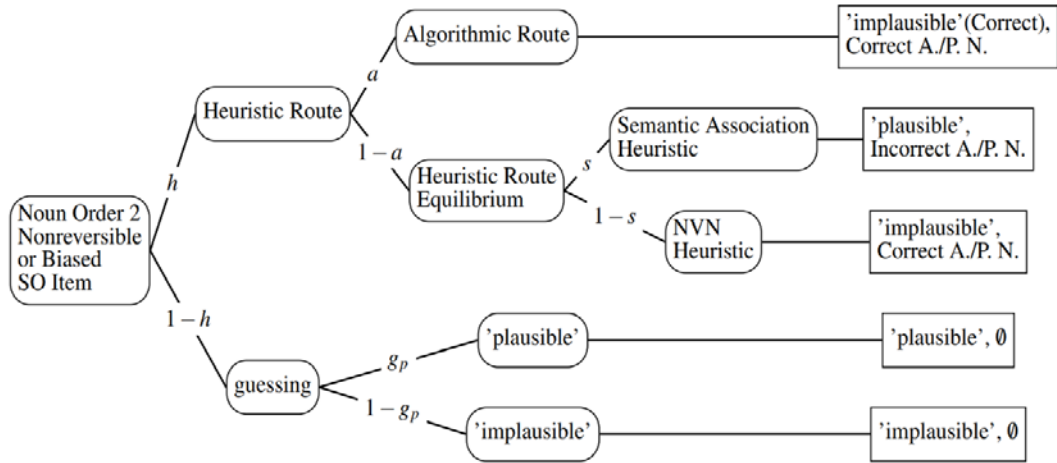


Figure 20. Our model of the agent/patient naming and plausibility task responses to the NO2 Nonreversible and Biased SO conditions in Meng and Bader (2021). ‘Correct A./P. N.’ stands for correct agent/patient naming, and the empty set symbol represents either incorrect agent/patient naming or the lack of a response to the agent/patient naming task

Finally, Figure 21 lists the outcomes to the NO2 Nonreversible and Biased OS and passive conditions and the items from these conditions are shown in Table 28. As can be seen from Figure 21, both the semantic-association heuristic and the NVN heuristic lead to a ‘plausible’ response and the incorrect agent/patient naming in these conditions. The reason why both heuristics lead to the incorrect agent/patient naming and a ‘plausible’ response is because the first noun encountered in the sentences from these conditions is also the plausible agent to the verb, ‘Koch’ or ‘chef’ in Table 28, which is not the actual agent of the sentences.

Before fitting the model to the data, we should also consider the descriptions of its parameters and the probability functions that result in the outcomes from each of the conditions. Our parsing account model, in its first version, has 4 parameters. The parameter h represents the probability of the heuristic route being engaged; hence the term $(1-h)$ represents the probability of guessing. The parameter g_p

represents the probability of guessing ‘plausible’, as with the models previously presented in this thesis. The parameter a represents the probability of the algorithmic route being engaged; thus, the term $(1-a)$ represents the probability of equilibrium being reached through the heuristic route. Finally, the parameter s represents the probability of the heuristic route’s interpretation being a result of the semantic-association heuristic; therefore, the term $(1-s)$ represents the probability of the heuristic route’s interpretation being a result of the NVN heuristic.

Table 27. The Items from the First Experiment of Meng and Bader (2021) from the NO2 Nonreversible or Biased SO Conditions

Structure	Meaning	Noun Order	Example Sentence	English Translation
SO	Nonreversible	2	Der Topf hat den Koch gereinigt.	‘The pan cleaned the chef’
	Biased	2	Der Braten hat den Koch ruiniert.	‘The roast ruined the chef’

Table 28. The Items from the First Experiment of Meng and Bader (2021) from the NO2 Nonreversible or Biased OS and Passive Conditions

OS	Nonreversible	2	Den Koch hat der Topf gereinigt.	‘The chef, the pan cleaned.’
	Biased	2	Den Koch hat der Braten ruiniert.	‘The chef, the roast ruined’
Passive	Nonreversible	2	Der Koch wurde vom Topf gereinigt.	‘The chef, by the pan, was cleaned.’
	Biased	2	Der Koch wurde vom Braten ruiniert.	‘The chef, by the roast, was ruined.’

The probability functions which return the probability of a ‘plausible’ response to the plausibility judgment task and those which return the probability of a correct response to agent/patient naming for each of the conditions from the first experiment of Meng and Bader (2021) are listed in Table 29. It can be recognized by comparing Table 29 with the visual representations of the models in Figures 17-21 that some terms are omitted from the probability equations. This is because the

equations were simplified by removing the unnecessary terms. For example, the sum of ‘(h x a)’ and ‘(h x (1-a))’ is equal to ‘h’, therefore, if a probability function featured the sum of these terms, only the simplified version was listed in Table 29.

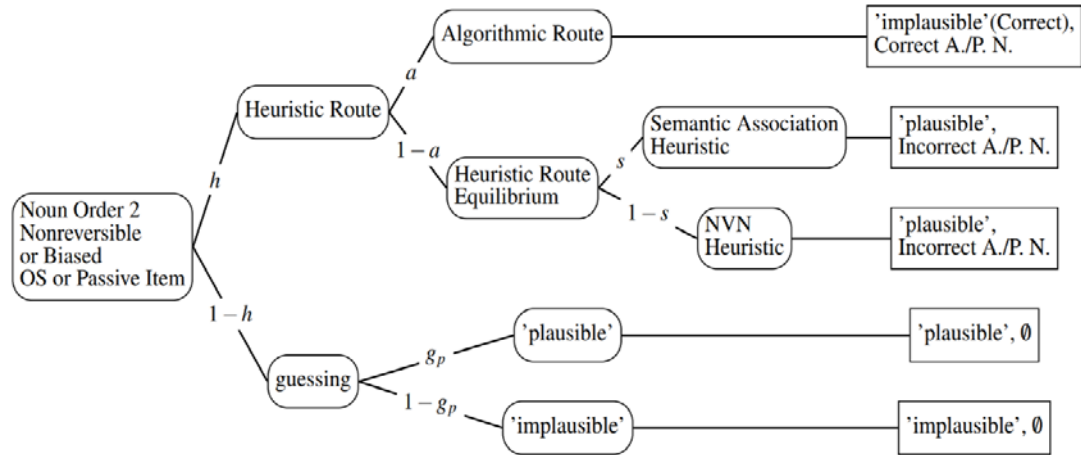


Figure 21. Our model of the agent/patient naming and plausibility task responses to the NO2 Nonreversible and Biased, OS or Passive conditions in Meng and Bader (2021). ‘Correct A./P. N.’ stands for correct agent/patient naming, and the empty set symbol represents either incorrect agent/patient naming or the lack of a response to the agent/patient naming task

Table 29. The Probability Functions which Return the Probability of a ‘plausible’ Response to the Plausibility Judgment Task and those which Return the Probability of a Correct Response to Agent/Patient Naming for each of the Conditions from the First Experiment of Meng and Bader (2021) According to the First Version of our Model of their Data

Noun Order	Structure	Meaning	Pr(‘plausible’ response)	Pr(Correct Agent/Patient Naming)
1	SO	Nonreversible or Biased	$h + (1-h) \times g_p$	h
1	OS & PS	Nonreversible or Biased	$h \times a + h \times (1-a) \times s + (1-h) \times g_p$	$h \times a + h \times (1-a) \times s$
1	SO	Symmetrical	$h + (1-h) \times g_p$	h
1	OS & PS	Symmetrical	$h + (1-h) \times g_p$	$h \times a$
2	SO	Nonreversible or Biased	$h \times (1-a) \times s + (1-h) \times g_p$	$h \times a + h \times (1-a) \times (1-s)$
2	OS & PS	Nonreversible or Biased	$h \times (1-a) + (1-h) \times g_p$	$h \times a$
2	SO	Symmetrical	$h + (1-h) \times g_p$	h
2	OS & PS	Symmetrical	$h + (1-h) \times g_p$	$h \times a$

When the model is fit to the plausibility judgment and the agent/patient naming data from the first experiment of Meng and Bader (2021), the resulting parameter estimates are as shown in Table 30, and when these estimates are plugged into the probability functions listed in Table 29, the resulting model predictions for the agent/patient naming data are as shown in Figure 22, and the model predictions for the plausibility judgment data are as shown in Figure 23.

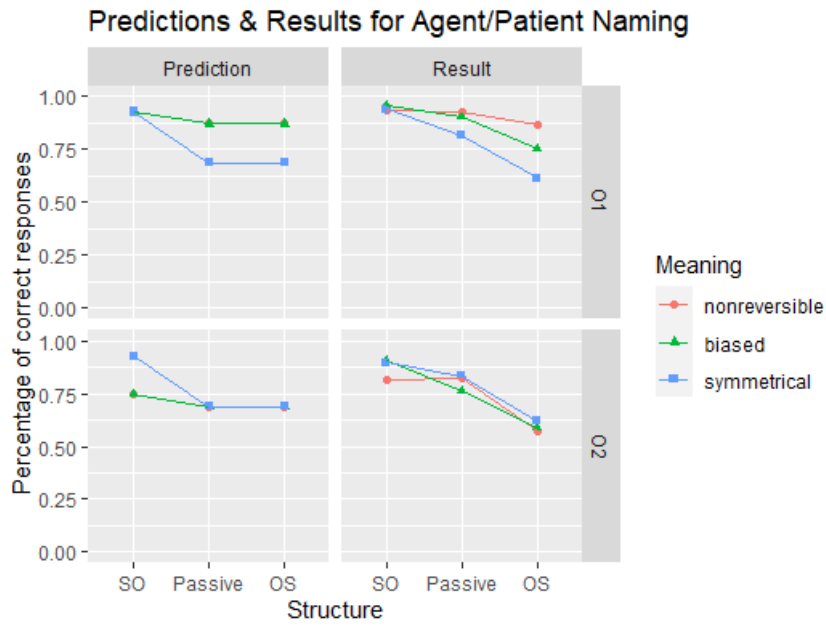


Figure 22. The plot on the left shows the percentage of correct responses to agent/patient naming as predicted by the models in Figures 17-21, with the parameter values shown in Table 30. The plot on the right shows the percentage of correct responses to agent/patient naming in the data. The percentage of correct responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

When we compare the model’s predictions for the agent/patient naming task and the results in Figure 22, we can see that the model does a fair job at predicting the absolute values of the percentage of correct responses to agent/patient naming, and successfully predicts that the SO conditions yielded more correct responses to agent/patient naming across all meaning levels. Moreover, the model is also able to

predict the two-way interaction between structure and meaning found in Meng and Bader's (2021) analysis of the agent/patient naming results. This can be spotted if the predicted percentage of correct responses to agent patient naming for each of the meaning levels in NO1 SO conditions is compared those in Noun Order OS and passive conditions: the predicted percentage of correct responses in the former are equal while those in the latter are different with the symmetrical conditions yielding the lowest percentage of correct responses in the NO1 OS and passive conditions.

However, the model fails in predicting some crucial facts about the data. Firstly, and most importantly, as can be seen in Figure 22, the model is unable to predict a difference between the structure levels of passive and OS. The technical reason why this is so has to do with the structure of the model, which resulted from the suggestions of the parsing account as discussed before, which are the theoretical reasons why the model predicts no difference between passive and OS conditions. Because both of these sentence structures are in the noncanonical word order, neither an outcome of the semantic-association heuristic nor that of the NVN heuristic can predict a difference between these two sentence structures. The NVN heuristic will always assign the agent thematic role to the first noun it encounters, which is not the actual agent of the sentence for both of these sentence structures, and the semantic-association heuristic will always assign the agent thematic role to the more 'plausible' noun with regard to the verb of the sentence, which is always the actual agent in all NO1 conditions for both of these structures and always the patient in all NO2 conditions for both of these structures. Therefore, in order for an MPT model to predict a difference between the percentage of correct responses to agent/patient naming in passive and OS conditions, additional assumptions about the heuristics used by the HPM or additional assumptions about the entire parsing procedure under

the parsing account must be made. These will be explored in the fourth version of this model.

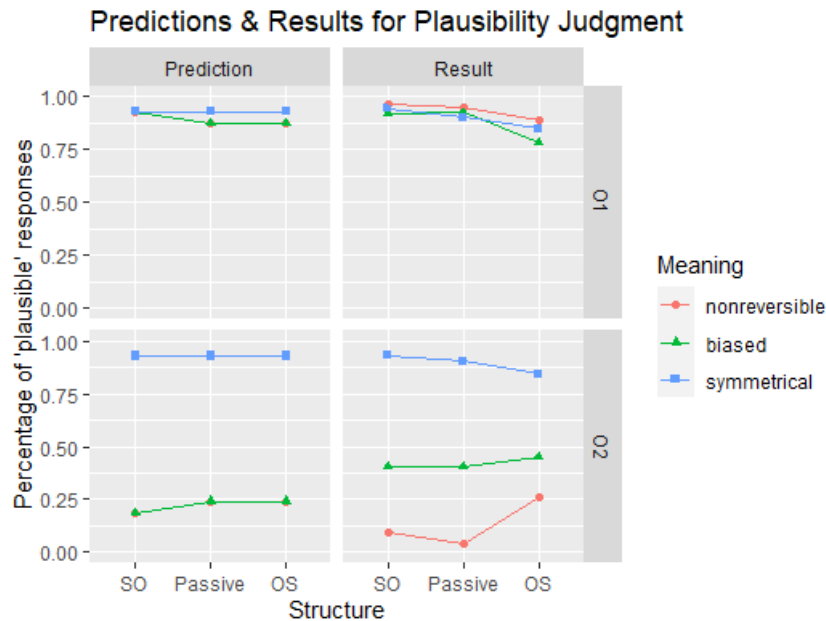


Figure 23. The plot on the left shows the percentage of ‘plausible’ responses to plausibility judgment as predicted by the models in Figures 17-21, with the parameter values shown in Table 30. The plot on the right shows the percentage of ‘plausible’ responses to the plausibility judgment task in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

When we compare the model’s predictions for the plausibility judgment task and the results in Figure 23, we can see that the model does a fair job at predicting the absolute values of the percentage of ‘plausible’ responses to all NO1 conditions and the NO2 symmetrical condition. However, the model fails to predict the difference in the percentage of ‘plausible’ responses across all structure levels between the nonreversible and biased conditions in NO2. This is again because of the structure of the model and the cognitive states we assumed in it, which is a formalization of the suggestions of the parsing account, which has no suggestions that directly refer to a difference in the percentage of ‘plausible’ responses obtained from a plausibility judgment task between these two meaning levels. This is no

wonder if we consider that in the studies that lead to the development of the parsing account, a plausibility judgment task was never analyzed in the manner it was in the study of Meng and Bader (2021), and that a main effect of meaning was only found in the plausibility judgment task data of Meng and Bader, but not in their agent/patient naming data. Therefore, again, in order for an MPT model to predict a difference between the percentage of ‘plausible’ responses to plausibility judgment in nonreversible and biased conditions, additional assumptions about the heuristics used by the HPM or additional assumptions about the entire parsing procedure under the parsing account must be made. What additional assumptions can be made in this regard will not be explored in this thesis, as the complications that these kinds of assumptions would introduce to the modelling procedure are beyond the scope of this thesis.

Table 30. Estimated Values for each of the Free Parameters when the Models in Figures 17-21 with the Functions in Table 29 are Fit to the Data

Parameter	Meaning of the Parameter	Estimated Value
h	Probability of the heuristic route being engaged	~0.93
a	Probability of the algorithmic route being engaged	~0.74
s	Probability of the heuristic route’s interpretation being a result of the semantic-association heuristic	~0.76
g_p	Probability of guessing ‘plausible’	~0.013

When we look at the parameter estimates in Table 30 we can see that the probability of the heuristic route being engaged, that is h , is very high, which means that the probability of a guess is estimated to be around 0.07, which is not much different from our final model of the retrieval account as the probability of comprehension, which contrasted with the probability of a guess in the retrieval account models, was as high as 0.90. On the other hand, we can see in Table 30 that

the estimated value for the probability of guessing ‘plausible’, that is g_p , is almost zero and that is suspiciously low. This means that the model predicts that a guess practically always results in an ‘implausible’ response. The technical reason why this is so has to do with the semantic-association heuristic always resulting in a ‘plausible’ response and the estimated value for the probability of the heuristic route’s interpretation being a result of the semantic-association heuristic, that is s as shown in Table 30. Because of the percentage of ‘plausible’ responses to the plausibility judgment in the data are high for those conditions that match the outcome from the semantic-association heuristic in terms of both its agent/patient naming and plausibility judgment predictions, the optimization function adjusted the g_p parameter so that it compensates for the lack of ‘implausible’ responses in a model where the semantic association heuristic mostly leads to what the data shows. This explanation will be demonstrated in the third version of the model, which does not feature the s parameter, or assume the cognitive state of ‘Semantic Association Heuristic’.

5.2.2 Parsing account model 2

As discussed at various points within this thesis, the parsing account is an account of the processing of noncanonical word order sentences that follows the more general theory of sentence processing of Good-Enough Processing (Christianson et al., 2001; F. Ferreira, 2003; F. Ferreira, Ferraro, & Bailey, 2002; F. Ferreira & Patson, 2007; Sanford & Sturt, 2002). One of the most fundamental suggestions of the theory of Good-Enough Processing is that the HPM is task-dependent and versatile in that the cognitive processes that are engaged in creating mental representations of sentences are highly dependent on what the task-at-hand is. The first version of our model of the parsing account, as it was thought of as a model of both the plausibility judgment

and the agent/patient naming tasks from the first experiment of Meng and Bader (2021), did not implement this fundamental idea under the theory of Good-Enough Processing, in the sense that the probabilities of occurrence for each of the cognitive states assumed by our model are the same for both tasks. Therefore, we wanted to create a second version of our parsing account model which implements this idea, in the hope that some of the problems with the predictions of the first version of the model are solved.

Considering our understanding of the parsing account's suggestions and what our understanding has led to in terms of decisions about MPT modelling, discussed in more detail in the previous section of this chapter, we found it most fitting to the suggestions of the parsing account to assume a separate a parameter, which is the probability that the algorithmic route is engaged and its outcome will determine the responses to the tasks, for the plausibility judgment task. This is because the suggestion under the theory of Good-Enough Processing about the task-dependency and versatility of the HPM is that the more demanding a task is, the more likely it is that the interpretation created through the algorithmic route will determine the final interpretation of the comprehender for the sentence in-question (Karimi & Ferreira, 2016), and we believe that it is safe, under the parsing account, as with the retrieval account, to assume that the agent/patient naming task is more demanding than the plausibility judgment task. Moreover, assuming a separate parameter h , which is the probability that the heuristic route is engaged, which, under the parsing account, is always the first of the two processing routes to start (Karimi & Ferreira, 2016), for the two tasks would also translate, in our model, to that the probability of a guess, the cognitive event which contrasts with that of 'Heuristic Route' in our model, is also different for the two tasks. Such an assumption, we believe, is not in the spirit of the

parsing account, due to the fact that under this account, the task demands are not suggested to influence the attentiveness of the comprehender, but only to influence the cognitive processes that come into play during parsing. Furthermore, assuming a separate a parameter for each task in our model would also mean that the probability of equilibrium through the heuristic route, the cognitive event which contrasts with that of ‘Algorithmic Route’ in our model, is different for the two tasks, which is also in-line with the suggestion under the theory of Good-Enough Processing that the more demanding a task is, the harder it is for the HPM to reach equilibrium.

However, a significant problem with suggesting a separate parameter a for each task with regard to the suggestions of the parsing account discussed in the former paragraph while modelling the data from the first experiment of Meng and Bader (2021) is that both the plausibility judgment and the agent/patient naming tasks were completed by the participants after the sentence was read, and hence an interpretation of the sentence was already created. Therefore, assuming that there are different probabilities of the algorithmic route being engaged for the two tasks, cannot follow directly in a model of the data from Meng and Bader’s (2021) first experiment from the suggestion under the theory of Good-Enough Processing that the task demands determine what cognitive processes will influence the final interpretation of the comprehender. Nevertheless, Meng and Bader’s (2021) data showed effects that were similar to those found in Bader and Meng (2018), a study in which the two tasks of plausibility judgment and agent/patient naming were used in separate experiments, and only one of them was used for each experiment. Hence, we believe it is safe to explore in a model of the data from the first experiment of Meng and Bader (2021) whether the introduction of a separate a parameter for each task would improve the model’s predictions.

Table 31. The Probability Functions which Return the Probability of a ‘plausible’ Response to the Plausibility Judgment Task and those which return the Probability of a Correct Response to Agent/Patient Naming for each of the Conditions from the First Experiment of Meng and Bader (2021) According to the Second Version of our Model of their Data

Noun Order	Structure	Meaning	Pr(‘plausible’ response)	Pr(Correct Agent/Patient Naming)
1	SO	Nonreversible or Biased	$h + (1-h) \times g_p$	h
1	OS & PS	Nonreversible or Biased	$h \times a_1 + h \times (1-a_1) \times s + (1-h) \times g_p$	$h \times a_2 + h \times (1-a_2) \times s$
1	SO	Symmetrical	$h + (1-h) \times g_p$	h
1	OS & PS	Symmetrical	$h + (1-h) \times g_p$	$h \times a_2$
2	SO	Nonreversible or Biased	$h \times (1-a_1) \times s + (1-h) \times g_p$	$h \times a_2 + h \times (1-a_2) \times (1-s)$
2	OS & PS	Nonreversible or Biased	$h \times (1-a_1) + (1-h) \times g_p$	$h \times a_2$
2	SO	Symmetrical	$h + (1-h) \times g_p$	h
2	OS & PS	Symmetrical	$h + (1-h) \times g_p$	$h \times a_2$

Therefore, we implemented a separate a parameter for each task, a_1 for the plausibility judgment task, and a_2 for the agent/patient naming task, into the models shown in Figures 17-21. The probability functions that return the probability of a ‘plausible’ response to the plausibility judgment task, and those that return the probability of a correct response to the agent/patient naming task for each of the conditions resulting from this revision of the model are as listed in Table 31.

When the model is fit to the plausibility judgment and the agent/patient naming data from the first experiment of Meng and Bader (2021), the resulting parameter estimates are as shown in Table 32, and when these estimates are plugged into the probability functions listed in Table 31, the resulting model predictions for the agent/patient naming data are as shown in Figure 24, and the model predictions for the plausibility judgment data are as shown in Figure 25.

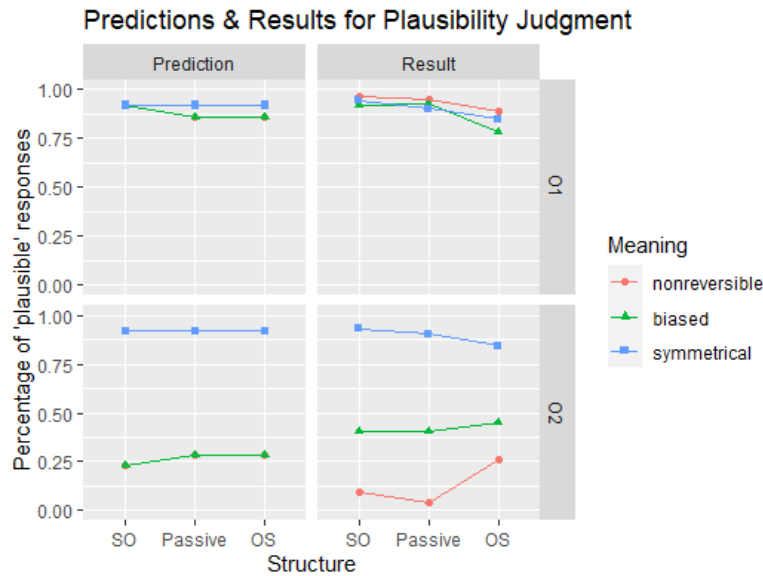


Figure 24. The plot on the left shows the percentage of ‘plausible’ responses to plausibility judgment as predicted by the models in Figures 17-21, with the probability functions shown in Table 31, and the parameter values shown in Table 32. The plot on the right shows the percentage of ‘plausible’ responses to the plausibility judgment task in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

Table 32. Estimated Values for each of the Free Parameters when the Models in Figures 17-21 with the Functions in Table 31 are Fit to the Data

Parameter	Meaning of the Parameter	Estimated Value
h	Probability of the heuristic route being engaged	~ 0.92
a_1	Probability of the algorithmic route being engaged in the plausibility judgment task	~ 0.68
a_2	Probability of the algorithmic route being engaged in the agent/patient naming task	~ 0.78
s	Probability of the heuristic route’s interpretation being a result of the semantic-association heuristic	~ 0.79
g_p	Probability of guessing ‘plausible’	~ 0.00002

Looking at the parameter estimates in Table 32, we can see that the probability of the heuristic route being engaged, parameter h , is around the same value which was estimated for the first version of our parsing account model, which

was 0.93. This shows that the probability of a guess, which is $(1 - h)$, is also similar in this version of the model to what was in the first version of the model.

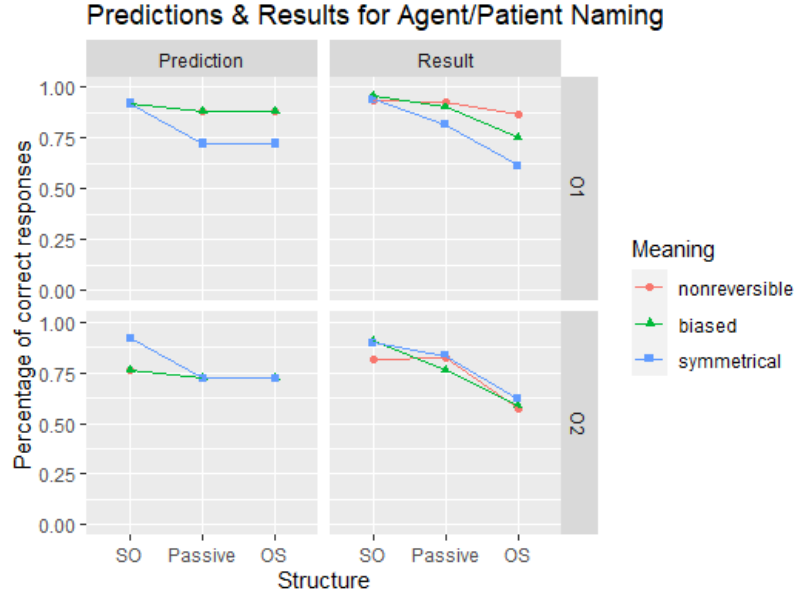


Figure 25. The plot on the left shows the percentage of correct responses to agent/patient naming as predicted by the models in Figures 17-21, with the probability functions shown in Table 31, and the parameter values shown in Table 32. The plot on the right shows the percentage of correct responses to agent/patient naming in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

Moreover, we can see from Table 32 that the probability of the algorithmic route being engaged in the agent/patient naming task, the parameter a_1 , as estimated by this model is higher than that for the plausibility judgment task, the parameter a_2 . This estimation is in-line with the suggestion under the theory of Good-Enough Processing that the more demanding a task is, the more likely it is that the interpretation created through the algorithmic route will determine the final interpretation of the comprehender for the sentence in-question. Furthermore, the probability of guessing plausible, g_p , for this version of the model is estimated to be even lower than that for the first version of the model, which could already be

considered practically zero, which tells us that the model has not improved in that regard with the introduction of the new parameters.

Even though the parameter estimates show that the model predicts that the probability of the algorithmic route being engaged in the agent/patient naming task, a_1 , is higher than that for the plausibility judgment task, the parameter a_2 , which is in-line with the suggestions of the parsing account as discussed before, looking at Figures 24 and 25, we see that the predictions have not improved at all. The model can still predict that the SO conditions yielded more correct responses to agent/patient naming across all meaning levels, and the two-way interaction between structure and meaning in the NO1 conditions in agent/patient naming, and it still fails to predict a difference between the structure levels of passive and OS in both tasks, and a difference in the percentage of ‘plausible’ responses across all structure levels between the nonreversible and biased conditions in NO2.

Table 33 shows the predicted probabilities of correct agent/patient naming and a ‘plausible’ response to the plausibility judgment task for this and the first versions of the model, which are referred to as Model 1 and Model 2, respectively. Through a comparison of the predicted probabilities in Table 33, it is clear that the changes are negligible. Therefore, we will abandon the idea of separate probabilities of the algorithmic route being engaged in the agent/patient naming task and the plausibility judgment task in the next versions of our parsing account model.

Table 33. The Probabilities of Correct Agent/Patient Naming and a ‘plausible’ Response to the Plausibility Judgment Task for each Condition from the First Experiment of Meng and Bader (2021) as Predicted by the First and the Second Versions of the Parsing Account Model, Referred to as Model 1 and Model 2, Respectively

Order	Meaning	Structure	Pr(Correct agent/patient naming)		Pr('plausible' response)	
			Model 1	Model 2	Model 1	Model 2
1	nonreversible	SO	0.93	0.92	0.93	0.92
1	nonreversible	Passive	0.87	0.88	0.87	0.86
1	nonreversible	OS	0.87	0.88	0.87	0.86
1	biased	SO	0.93	0.92	0.93	0.92
1	biased	Passive	0.87	0.88	0.87	0.86
1	biased	OS	0.87	0.88	0.87	0.86
1	symmetrical	SO	0.93	0.92	0.93	0.92
1	symmetrical	Passive	0.69	0.72	0.93	0.92
1	symmetrical	OS	0.69	0.72	0.93	0.92
2	nonreversible	SO	0.74	0.76	0.18	0.23
2	nonreversible	Passive	0.69	0.72	0.24	0.29
2	nonreversible	OS	0.69	0.72	0.24	0.29
2	biased	SO	0.74	0.76	0.18	0.23
2	biased	Passive	0.69	0.72	0.24	0.29
2	biased	OS	0.69	0.72	0.24	0.29
2	symmetrical	SO	0.93	0.92	0.93	0.92
2	symmetrical	Passive	0.69	0.72	0.93	0.92
2	symmetrical	OS	0.69	0.72	0.93	0.92

5.2.3 Parsing account model 3

As discussed under the subsection in this chapter for the first version of our parsing account model, the reason why the estimated value for the probability of guessing ‘plausible’, that is g_p , is almost zero in the first two versions of our parsing account model may have to do with the theoretical assumptions about the semantic-association heuristic. The semantic association heuristic, as suggested under the parsing account, assigns the thematic role of agent to the more plausible of the two nouns with regard to the verb of the sentence. Because this is so, the semantic-association heuristic always leads to the correct agent/patient naming as well as the

correct plausibility judgment in all NO1 conditions except the symmetrical OS and passive conditions, whereas the NVN heuristic leads to errors in both tasks in all OS and passive conditions, regardless of Noun Order. As a result of these outcomes from each of the two heuristics, because the semantic-association heuristic is actually able to account for a larger portion of the data, the optimization function estimated the probability of the heuristic route's interpretation being a result of the semantic-association heuristic, that is s , higher than its counterpart, the probability of the heuristic route's interpretation being a result of the NVN heuristic, that is $(1 - s)$. However, because this also results in a higher predicted percentage of 'plausible' responses, due to the suggested nature of the semantic-association heuristic, the optimization function estimated the probability of a 'plausible' guess to be near zero, to compensate for the lack of predicted 'implausible' responses, so that the model better fits the data. While the probability of a 'plausible' guess is estimated to be near zero, it should also be noted that this does not have a very significant impact on the model's predicted percentage of 'plausible' responses, due to the fact that the probability of a guess, that is $(1 - h)$ was estimated to be as low as ~ 0.07 for the first two versions of the model.

In order to possibly address this problem with the model, we thought of appealing to the suggestion of Ferreira (2003) that the NVN heuristic has a stronger impact on the final interpretation of a sentence than the semantic-association heuristic. Because our parsing account model, with the parameters from the first two versions, predicts that the probability of the heuristic route's interpretation being a result of the semantic-association heuristic, that is s , is higher than its counterpart, the probability of the heuristic route's interpretation being a result of the NVN heuristic, that is $(1 - s)$, our parsing account model can be considered to be against the

suggestion of Ferreira (2003). Therefore, exploring what our parsing account model would predict if the only heuristic available for the heuristic route was the NVN heuristic, may provide valuable insights about the suggestions of the parsing account.

The probability functions resulting from removing the cognitive event of ‘Semantic Association Heuristic’ and thus also the s parameter, from the models in Figures 17-21, results in the probability functions shown in Table 34 for each of the conditions from the first experiment of Meng and Bader (2021).

Table 34. The Probability Functions which Return the Probability of a ‘Plausible’ Response to the Plausibility Judgment Task and those which Return the Probability of a Correct Response to Agent/Patient Naming for Each of the Conditions from the First Experiment of Meng and Bader (2021) According to the Third Version of our Model of their Data

Noun Order	Structure	Meaning	Pr(‘plausible’ response)	Pr(Correct Agent/Patient Naming)
1	SO	Nonreversible or Biased	$h + (1 - h) \times g_p$	h
1	OS & PS	Nonreversible or Biased	$h \times a + (1 - h) \times g_p$	$h \times a$
1	SO	Symmetrical	$h + (1 - h) \times g_p$	h
1	OS & PS	Symmetrical	$h + (1 - h) \times g_p$	$h \times a$
2	SO	Nonreversible or Biased	$(1 - h) \times g_p$	h
2	OS & PS	Nonreversible or Biased	$h \times (1 - a) + (1 - h) \times g_p$	$h \times a$
2	SO	Symmetrical	$h + (1 - h) \times g_p$	h
2	OS & PS	Symmetrical	$h + (1 - h) \times g_p$	$h \times a$

When the model is fit to the plausibility judgment and the agent/patient naming data from the first experiment of Meng and Bader (2021), the resulting parameter estimates are as shown in Table 35, and when these estimates are plugged into the probability functions listed in Table 34, the resulting model predictions for the plausibility judgment data are as shown in Figure 26, and the model predictions for the agent/patient naming data are as shown in Figure 27.

Table 35. Estimated Values for each of the Free Parameters when the Models in Figures 17-21 with the Functions in Table 34 are Fit to the Data

Parameter	Meaning of the Parameter	Estimated Value
h	Probability of the heuristic route being engaged	~ 0.73
a	Probability of the algorithmic route being engaged	~ 0.83
g_p	Probability of guessing 'plausible'	~ 0.88

Looking at the estimated parameter values in Table 35, we can see that the probability of guessing 'plausible', that is g_p , as estimated by this version of our parsing account model is much higher than the estimations from the previous versions of the same model in the absence of the s parameter, which represented the probability of the heuristic route's interpretation being a result of the semantic-association heuristic in the previous versions of the model. This confirms our predictions about why the probability of a guess in the previous versions of our model was near zero, which were discussed in the beginning of this subsection. Moreover, the probability of the heuristic route being engaged, that is h , which contrasts with the probability of a guess, that is $(1 - h)$, in all versions of our model, is estimated to be lower than the estimates for the same parameter in the previous versions of our model. This is because of the model's assumption that, for a 'plausible' guess to occur, the comprehender must first enter into a state of guessing. Thus, for the high probability of a 'plausible' guess, as estimated by the model, as shown in Table 35, to come into effect, the probability of a guess, that is $(1 - h)$, must be higher and so the probability of the heuristic route being engaged, that is h , must be lower than what it was for the previous versions of the model. Furthermore, the fact that the estimated probability of the algorithmic route being engaged in this version of the model, that is a , is higher for this version of the model than the

previous versions, in which the estimated values for a were around 0.74, along with the fact that the estimated probability of a guess, that is $(1 - h)$, is much higher for this version of the model than the previous versions, supports our suggestion that the semantic-association heuristic is more effective than the NVN heuristic in predicting the data from Meng and Bader's (2021) first experiment, when a model of the parsing account, such as the one we created in this chapter, is assumed.

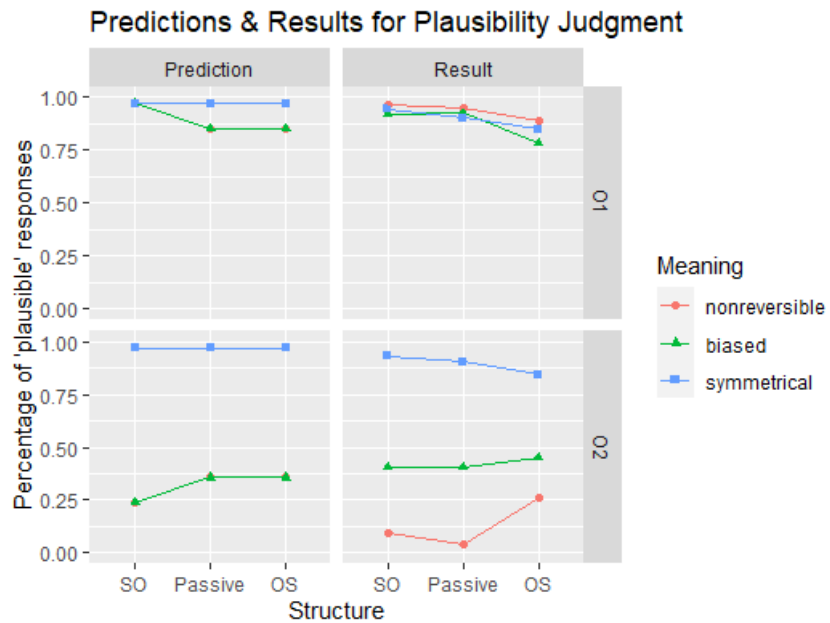


Figure 26. The plot on the left shows the percentage of ‘plausible’ responses to plausibility judgment as predicted by the models in Figures 17-21, with the probability functions shown in Table 34, and the parameter values shown in Table 35. The plot on the right shows the percentage of ‘plausible’ responses to the plausibility judgment task in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

Comparing the predictions of this version of the model with the results from the first experiment of Meng and Bader (2021), both of which are shown in Figures 25 and 26, reveals that this version of the model does not do much worse than the previous versions of the model. This version of the model, like the previous versions, can predict that the SO conditions yielded more correct responses to agent/patient

naming across all meaning levels, as can be seen in Figure 27, and it can predict that the percentage of ‘plausible’ responses to nonreversible and biased SO conditions in NO2 is lower than that to the OS conditions in the same meaning and noun order levels, as can be seen in Figure 26. Moreover, like the previous versions of the model, this model fails to predict the difference in the percentage of ‘plausible’ responses across all structure levels between the nonreversible and biased conditions in NO2, and it fails to predict any differences between the structure levels of passive and OS in both tasks. However, it is also evident from Figure 26 that this version of the model, unlike the previous versions of the model, predicts a higher percentage of ‘plausible’ responses for nonreversible and biased conditions than the previous versions across all structure levels. Moreover, as can be seen in Figure 27, this version of the model, unlike the previous versions, fails to predict the two-way interaction between structure and meaning found in Meng and Bader’s (2021) analysis of the agent/patient naming results.

Although the suggestion that the only heuristic responsible for creating interpretations in the heuristic route is the NVN heuristic, as implemented in this version of the model, is in conflict with the suggestions of the parsing account, the exploration of the predictions of such a model has revealed that it is unlikely that the NVN heuristic is more dominant than the semantic-association heuristic within the heuristic route as suggested by Ferreira (2003) if an MPT model of noncanonical word order sentence processing which follows the suggestions of the parsing account such as the one we suggested in this thesis is actually reflecting those suggestions.

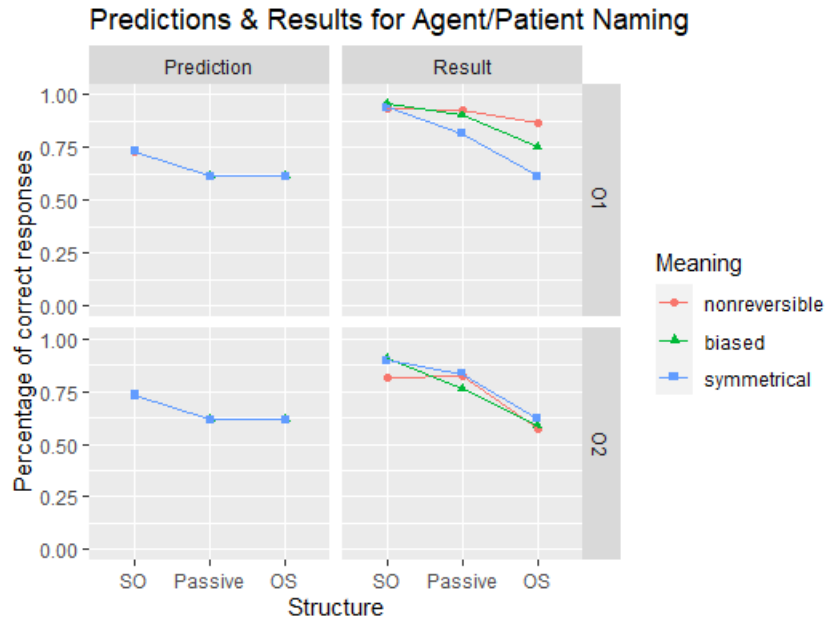


Figure 27. The plot on the left shows the percentage of correct responses to agent/patient naming as predicted by the models in Figures 17-21, with the probability functions shown in Table 34, and the parameter values shown in Table 35. The plot on the right shows the percentage of correct responses to agent/patient naming in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

5.2.4 Parsing account model 4

Our trials with the previous versions of the parsing account model shown in Figures 17-21, showed that for an MPT model of the parsing account, such as the one we created in this thesis, to account for certain significant effects found in the first experiment of Meng and Bader (2021), additional assumptions must be made about the heuristics used by the HPM or about the entire parsing procedure under the parsing account. This is because, the suggestions of the parsing account that we had not implemented in the first version of the model, like the task-dependency of the HPM, or a more heavily weighted NVN heuristic, did not improve the predictions of our parsing account model at all.

Therefore, the only way we can proceed to investigate the performance of an MPT model of the parsing account is to go beyond the suggestions of the parsing

account. At least two ways of doing this were revealed through the examination of the performance of the previous versions of our parsing account model. The first option to improve the model was revealed by the incapacity of our model to account for the differences in the percentage of ‘plausible’ responses between nonreversible and biased conditions in NO2 across all structure levels. We explained under the first subsection of this chapter that the model can only account for such an effect without postulating new cognitive events by assuming a different parameter value for a parameter for the nonreversible and biased conditions, but the complication that this brings into the modelling procedure makes such a revision to the model beyond the scope of this thesis.

Another option to improve the model was revealed by the incapacity of our model to account for any differences between the structure levels of OS and passive. Meng and Bader (2021) found a main effect of structure in their plausibility judgment data, and a two-way interaction between structure and meaning in their agent/patient naming data. Although our parsing account model could not predict any differences between the meaning levels of nonreversible and biased, as explained in the former paragraph, the first version of it was able to predict the two-way interaction between structure and meaning found in the agent/patient naming data from the first experiment of Meng and Bader (2021). Therefore, if the model is also able to predict the differences between the structure levels of OS and passive, it will be a much better account of both the plausibility judgment and the agent/patient naming data from the experiment.

We thought the best way to revise the model so that it can account for the differences between the structure levels of OS and passive, while still remaining in the spirit of the parsing account, was to assume that the probability of the algorithmic

route being engaged and hence the heuristic route's interpretation being abandoned, the parameter a in our model, has different values for the OS and passive conditions. This can be considered to be theoretically motivated if we synthesize the suggestion of Meng and Bader (2021) about the information-structural markedness of the OS sentences in the lack of an appropriate discourse context, with the suggestion under the theory of Good-Enough Processing that the more demanding a task is, the harder it is for the HPM to reach equilibrium (Karimi & Ferreira, 2016). We believe it would not conflict with the suggestions of the Online Cognitive Equilibrium Hypothesis to suggest that when the comprehender encounters an object in the first noun position of the sentence and realizes that there was no appropriate discourse context (Karimi & Ferreira, 2016), the criterion for equilibrium is pushed further, or simply, equilibrium becomes harder to reach. Since equilibrium becoming harder to reach translates to more probability of the final interpretation of the comprehender resulting from the algorithmic route, such an assumption can directly be implemented in our MPT model of the parsing account by assuming a different parameter a , which is the probability of the algorithmic route being engaged and hence the heuristic route's interpretation being abandoned, for the OS and passive conditions.

Therefore, this new version of the model does not assume any new cognitive events or states apart from the ones shown in Figures 17-21 but assumes a separate parameter a for the OS conditions, that is a_{os} . The probability functions that return the probability of a 'plausible' response to plausibility judgment and that of a correct response to agent/patient naming for each of the conditions from the first experiment of Meng and Bader (2021) that result from this configuration are shown in Table 36.

When the model is fit to the plausibility judgment and the agent/patient naming data from the first experiment of Meng and Bader (2021), the resulting parameter estimates are as shown in Table 37, and when these estimates are plugged into the probability functions listed in Table 36, the resulting model predictions for the plausibility judgment data are as shown in Figure 27, and the model predictions for the agent/patient naming data are as shown in Figure 28.

Table 36. The Probability Functions which Return the Probability of a ‘plausible’ Response to the Plausibility Judgment Task and those which Return the Probability of a Correct Response to Agent/Patient Naming for each of the Conditions from the First Experiment of Meng and Bader (2021) According to the Fourth Version of our Model of their Data

Noun Order	Structure	Meaning	Pr(‘plausible’ response)	Pr(Correct Agent/Patient Naming)
1	SO	Nonreversible or Biased	$h + (1 - h) \times g_p$	h
1	Passive	Nonreversible or Biased	$h \times a + (1 - h) \times g_p$	$h \times a$
1	OS	Nonreversible or Biased	$h \times a_{os} + (1 - h) \times g_p$	$h \times a_{os}$
1	SO	Symmetrical	$h + (1 - h) \times g_p$	h
1	Passive	Symmetrical	$h + (1 - h) \times g_p$	$h \times a$
1	OS	Symmetrical	$h + (1 - h) \times g_p$	$h \times a_{os}$
2	SO	Nonreversible or Biased	$(1 - h) \times g_p$	h
2	Passive	Nonreversible or Biased	$h \times (1 - a) + (1 - h) \times g_p$	$h \times a$
2	OS	Nonreversible or Biased	$h \times (1 - a_{os}) + (1 - h) \times g_p$	$h \times a_{os}$
2	SO	Symmetrical	$h + (1 - h) \times g_p$	h
2	Passive	Symmetrical	$h + (1 - h) \times g_p$	$h \times a$
2	OS	Symmetrical	$h + (1 - h) \times g_p$	$h \times a_{os}$

As can be seen in Figure 28, the predictions for the plausibility judgment data from this version of the model show that, with the added a_{os} parameter, the model can now predict the differences between the structure levels of passive and OS, although the predicted percentage of ‘plausible’ responses to OS conditions in NO2

relative to those of the SO and passive conditions in the same noun order seems too high, if we remember that the model is unable to distinguish between the meaning levels of nonreversible and biased, it becomes clear why this was so. Because the collective number of ‘plausible’ responses to OS conditions in the nonreversible and biased meaning level in NO2 were much higher than those for the SO and passive conditions in the same noun order, the optimization function adjusted the parameters so that the predicted percentage of ‘plausible’ responses for all structure levels in nonreversible and biased conditions is closer to the average percentage of ‘plausible’ responses to nonreversible and biased conditions across all structure levels in NO2 in the results.

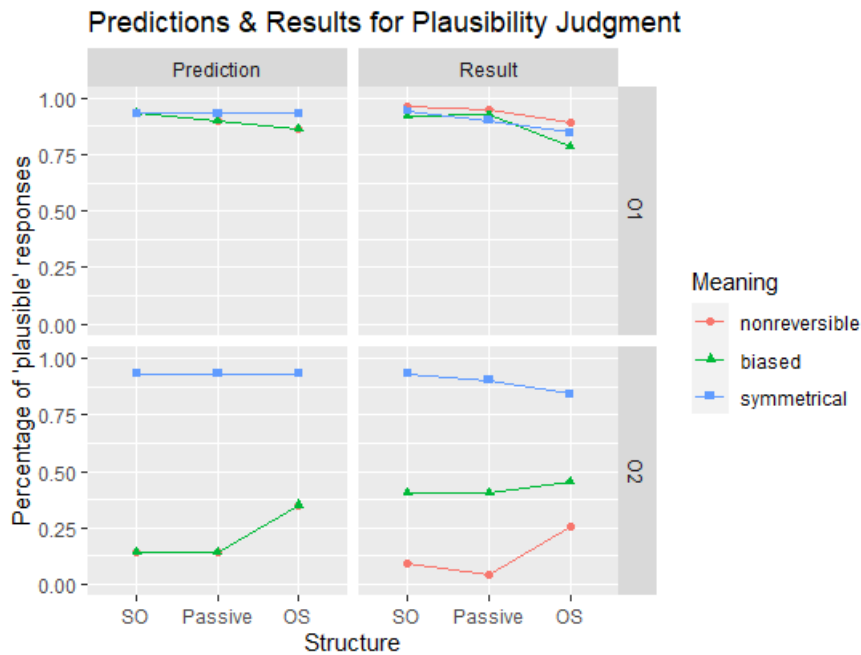


Figure 28. The plot on the left shows the percentage of ‘plausible’ responses to plausibility judgment as predicted by the models in Figures 17-21, with the probability functions shown in Table 19, and the parameter values shown in Table 35. The plot on the right shows the percentage of ‘plausible’ responses to the plausibility judgment task in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

However, as can also be seen in Figure 28, this version, like the previous versions of the model, is still unable to predict a difference between the passive and OS conditions in the symmetrical meaning level. The technical reason why this is so becomes clear upon an inspection of the probability functions for the outcomes shown in Table 36: the probability functions that return the percentage of ‘plausible’ responses to plausibility judgment are the same for all symmetrical conditions. This is because we excluded the cognitive state of ‘Semantic Association Heuristic’ from the parts of the model that modelled the symmetrical conditions. The reasoning behind the exclusion of the cognitive state that represented the semantic-association heuristic was that, in the symmetrical conditions, both nouns are plausible agents to the verb, and so, it is impossible to determine exactly what noun will be assigned the thematic role of agent by a heuristic that ignores the syntactic structure entirely and creates an interpretation of the sentence based on which of the two nouns is a more plausible agent to the verb. In other words, it is impossible to determine the outcome of the semantic-association heuristic in terms of the response to agent/patient naming in symmetrical sentences. Nevertheless, it is possible to create a separate model for the plausibility judgment task which still features the cognitive state that represents the semantic-association heuristic, and one would be theoretically motivated to do so, due to the suggested existence of this heuristic under the parsing account. However, although it is not a complicated procedure to implement this, such a revision to our model of the parsing account is not included in this thesis.

As can be seen in Figure 29, the predictions for the agent/patient naming data from this version of the model show that, with the added a_{os} parameter, the model can now predict the differences between the structure levels of passive and OS in agent/patient naming as well. Moreover, it can also predict the two-way interaction

between structure and meaning and the three-way interaction between structure, meaning and noun order found in the first experiment of Meng and Bader (2021), which reveals itself when we consider that the two-way interaction between structure and meaning does not exist in NO2 conditions in both the results and the predictions.

Although the predictions of the model in terms of percentage of ‘plausible’ responses to plausibility judgment and percentage of correct responses to agent/patient naming align with the actual experiment results fairly well, the estimated values for the model’s parameters have significant implications that run counter to the suggestions of the Online Cognitive Equilibrium Hypothesis (Karimi & Ferreira, 2016).

Table 37. Estimated Values for each of the Free Parameters when the Models in Figures 17-21 with the Functions in Table 36 are Fit to the Data

Parameter	Meaning of the Parameter	Estimated Value
h	Probability of the heuristic route being engaged	~0.92
a	Probability of the algorithmic route being engaged in SO and passive conditions	~0.81
a_{os}	Probability of the algorithmic route being engaged in OS conditions	~0.63
s	Probability of the heuristic route’s interpretation being a result of the semantic-association heuristic	~0.79
g_p	Probability of guessing ‘plausible’	~0.10

When we look at the parameter values in Table 37, we can see that the estimated value for the probability of the heuristic route being engaged, that is h , is almost equal to the first version of our parsing account model, leaving the probability of a guess, that is $(1 - h)$, at 0.07. On the other hand, the estimated value for the probability of the algorithmic route being engaged in SO and passive conditions, that is a , is a bit higher than the that for the parameter represented by the same symbol in

the first version of the model, which is due to the fact that the definition is limited to SO and passive conditions in this version. Moreover, the estimated value for the probability of a ‘plausible’ guess, that is g_p , is a lot higher than what it was for the previous versions of the model, although it is perhaps still too low to eliminate the suspicion about the model’s predictions. The reason why the estimated probability of a ‘plausible’ guess is so low, again, has to do with the high probability of the heuristic route’s interpretation being a result of the semantic-association heuristic, that is s , as explained in the previous subsections of this chapter.

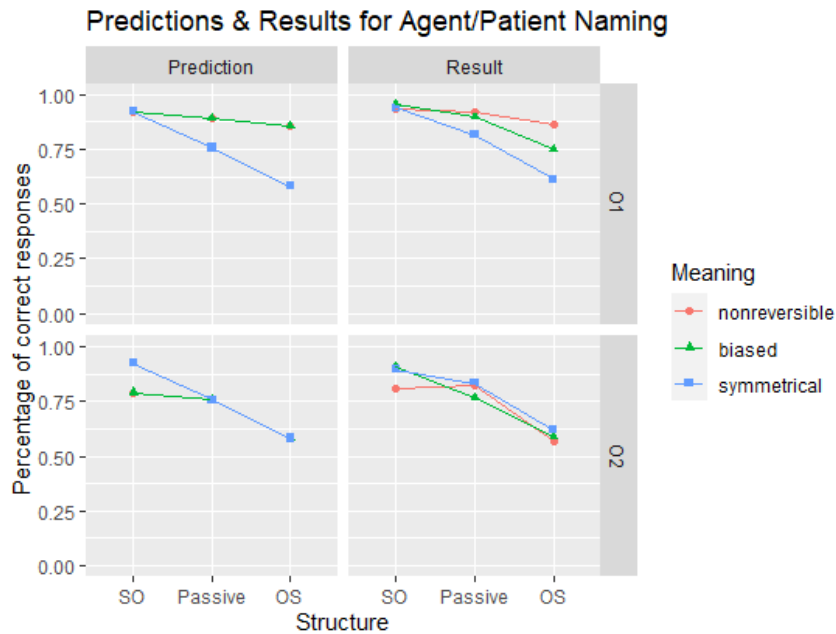


Figure 29. The plot on the left shows the percentage of correct responses to agent/patient naming as predicted by the models in Figures 17-21, with the probability functions shown in Table 36, and the parameter values shown in Table 35. The plot on the right shows the percentage of correct responses to agent/patient naming in the data. The percentage of ‘plausible’ responses are shown on the y-axis and the levels of the factor ‘structure’ are shown on the x-axis. The levels of the factor ‘meaning’ are represented by the color of the lines and the shape of the points

However, although the predictions of the model match the data closely, the value for the probability of the algorithmic route being engaged in OS conditions, that is a_{os} , is estimated to be considerably lower than a , which runs counter to our

intentions when we introduced the parameter a_{os} into the model. As explained in the beginning of this subsection, we introduced this new parameter into the model by appealing to the suggestion of Meng and Bader (2021) about the information-structural markedness of the OS sentences in the lack of an appropriate discourse context, and connected this suggestion with the suggestion under the theory of Good-Enough Processing that the more demanding a task is, the harder it is for the HPM to reach equilibrium (Karimi & Ferreira, 2016). Under the Online Cognitive Equilibrium Hypothesis (Karimi & Ferreira, 2016), when it is harder for the HPM to reach equilibrium due to task demands, it also means that it is more likely that the final interpretation of a comprehender is the result of the algorithmic route, and so, we introduced a separate parameter a for OS conditions, with the intention of implementing this idea in this version of our model. However, this version of our parsing account model predicts exactly the reverse: if we assume that Meng and Bader’s suggestion about the OS sentences is true, which is motivated by our MPT models of the retrieval account, then the probability of the algorithmic route being engaged in OS conditions, that is a_{os} , must be higher than the probability of the algorithmic route being engaged in SO and passive conditions, that is a , but as can be seen in Table 37, the estimated value of a_{os} is considerably lower than that of a . The technical reason why the estimates for these parameters are so, is that the algorithmic route always creates an interpretation of the sentence that leads to a correct response to both of the tasks, and the task accuracy for both tasks is always the lowest for the OS conditions in the data. However, because the suggestion of Meng and Bader clearly expresses higher difficulty for the processing of OS sentences, and because our MPT model of the retrieval account also predicted higher difficulty of comprehension for the OS conditions, we believe that the current version of our

parsing account model constitutes evidence against the Online Cognitive Equilibrium Hypothesis (Karimi & Ferreira, 2016).

In conclusion, this final MPT model of both the plausibility judgment and the agent naming data from the first experiment of Meng and Bader (2021) is sufficient for the purpose of this thesis. Our MPT model of the parsing account assumed the following: the heuristic processing route is only engaged when the comprehender is attentive, otherwise the comprehender resorts to guessing, which can result in the correct plausibility judgment but never the correct agent/patient naming; only if equilibrium through the heuristic route is not achieved, then the algorithmic route is engaged; the resulting interpretation from the heuristic route can either be the result of the NVN heuristic or the semantic-association heuristic; and that the probability of the algorithmic route being engaged in SO and passive conditions is different from that for the OS sentences. The modelling procedure has revealed that an MPT model of the parsing account which makes these assumptions can predict the significant effects of noun order and meaning found in the plausibility judgment data in the first experiment of Meng and Bader (2021), as well as the effect of structure found in their additional analysis of the plausibility judgment data which included only the SO and OS structure levels. Moreover, the procedure has also revealed that such an MPT model can also account for the two-way interaction between structure and meaning, as well as the three-way interaction between structure, meaning and noun order found in the agent/patient naming data from the first experiment of Meng and Bader (2021). However, the modelling procedure has also revealed that, even if we assume that the two tasks of plausibility judgment and agent/patient naming were used in separate experiments, with only one being used for each experiment in the data from the first experiment of Meng and Bader (2021), assuming that the two tasks have

different demands and hence have different probabilities of engagement for the algorithmic route, does not improve the fit of a model such as the one we presented here to the data. Moreover, the procedure has revealed that the suggestion of Ferreira (2003) that the NVN heuristic is more dominant than the semantic-association heuristic in determining the final interpretation created by the heuristic route is unlikely to be the case. Finally, the modelling procedure has revealed that, if we assume that the information-structural markedness of the OS sentences in the lack of an appropriate discourse context increases the difficulty of processing for these sentences, and implement this idea as a separate parameter for the probability of the algorithmic route being engaged for OS conditions, an MPT model of the parsing account, such as the we discussed in this chapter, predicts a higher probability of the algorithmic route being engaged for OS conditions than that for the SO and passive conditions, contrary to the suggestion under the Online Cognitive Equilibrium Hypothesis that the more demanding a task is, the more likely it is that the final interpretation of a comprehender is the result of the algorithmic route.

CHAPTER 6

GENERAL DISCUSSION

We have created multiple MPT models by following the suggestions under the parsing and the retrieval accounts of noncanonical word order sentence processing as closely as possible in an MPT model. The procedure of devising MPT models for the two accounts, getting the predictions of these models, and revising the models with regard to their performance in accounting for the data from the first experiment of Meng and Bader (2021) has provided us with important insights about both accounts including a clearer understanding of their suggestions. This chapter will summarize our findings throughout the entire modelling procedure for each of the accounts and discuss their strengths and weaknesses as revealed by the procedure.

The retrieval account of noncanonical word order sentence processing (Bader & Meng, 2018; Meng & Bader, 2021) suggests that the decrease in performance observed for the task of agent/patient naming in the studies of noncanonical word order sentences (Ferreira, 2003; Christianson, Luke & Ferreira, 2010; Bader and Meng, 2018; Meng and Bader, 2021; Cutter, Paterson & Filik, 2021) are a result of retrieval errors. Retrieval errors happened more often when the task required naming the agent and the patient of a sentence with noncanonical word order or an implausible sentence because the degree of match between the agent or patient noun's features and the retrieval cues in these sentences were lower than those for the canonical word order or plausible sentences. Meng and Bader (2018, 2021) also observed effects of sentence plausibility and sentence structure, particularly for the sentences in the Object-Subject (OS) word order, when the task was to judge the plausibility of a sentence, and suggested that the errors found in the plausibility

judgment task for sentences where the reversal of the linear position of the agent and the patient in a sentence resulted in an implausible (nonreversible sentences) or somewhat plausible (biased) meaning reflected the actual degree of the plausibility of these sentences, and that the errors found in the plausibility judgment task for OS sentences were due to the information structural markedness of these sentences in the absence of an appropriate discourse context.

We started developing models of the retrieval account by integrating the suggestions of the retrieval account about the errors found in the plausibility judgment task into an MPT structure and fit these models to the plausibility judgment data from the first experiment of Meng and Bader (2021). We assumed two cognitive states in the first version of our plausibility judgment model under the retrieval account, comprehension and guessing. These two cognitive states contrasted in our model, such that the model assumed that only one of the two could happen in each trial. In this first version, comprehension always led to the correct plausibility judgment response, and guessing could result in either a ‘plausible’ or ‘implausible’ response. This model was unable to predict the effect of plausibility found in the data, which led to our suggestion of a cognitive process whereby the degree of a match between the information conveyed by the sentence and the comprehenders world-knowledge is judged which followed the event of comprehension in our second version of the retrieval account model. This version of the model was able to predict the plausibility effect but was still unable to predict the errors found in the plausibility judgment task for OS sentences. We developed a third version of the model of the plausibility judgment data under the retrieval account by implementing the suggestion of Meng and Bader (2021) about the information structural markedness of OS sentences in the absence of an appropriate discourse context,

through the introduction of a cognitive process into the model whereby the pragmatic well-formedness of a sentence is judged. This version of the model predicted less ‘plausible’ responses to Object-Subject word order sentences with implausible noun order compared to passives and Subject-Object (SO) word order sentences, but the reverse was found in the data. Therefore, we considered the same suggestion of Meng and Bader (2021) about the OS sentences as conveying that there is a general difficulty of comprehension when these sentences are not in an appropriate discourse context, and so removed the cognitive process whereby the pragmatic well-formedness of a sentence is judged from the model, and instead assumed separate probabilities of comprehension for SO and passive sentences on one hand, and OS sentences on the other in a fourth version of the plausibility judgment model. This final version of the plausibility judgment model was able to account for all of the effects found in the plausibility judgment data from the first experiment of Meng and Bader (2021).

Because we wanted to model the data from the agent/patient naming and the plausibility judgment tasks from the first experiment of Meng and Bader (2021) jointly, we used our final model of the plausibility judgment data under the retrieval account as a starting point for our joint model of the two tasks and added a new parameter to the model which represented the probability of a successful retrieval of the agent and the patient. In our first version of the joint model, we took the suggestion of Meng and Bader (2021) that the probability of successful retrieval of the agent and the patient noun is a function of the degree of match between the agent or patient noun’s features and the retrieval cues in a simple way and calculated for each condition in their experiment the ratio of the cues that match the target noun to the total number of available cues, which determined exactly the probability of a

successful retrieval of the agent and the patient for that condition. This first version of the joint model underestimated the percentage of correct responses to all conditions in agent/patient naming and predicted a lower percentage of correct responses to passive items than those to OS sentences. In order to address this problem with the model, we implemented the suggestion of Meng and Bader (2021) that the difference in the phrase category of the phrase that included the agent noun for passive may have contributed to the distinctiveness of the memory item that represented this phrase, and so rendered its retrieval easier, by introducing cue-weights into the joint model in a third version of the model. The version of the joint model with cue-weights still underestimated the percentage of correct responses to all conditions in agent naming and had problematic estimates for the weights of some of the cues. To address this problem with the model, we implemented the suggestion under the cue-based retrieval theory of sentence comprehension (Van Dyke, J. A., & Lewis, 2003; Lewis, Vasishth, & Van Dyke, 2006; Van Dyke & Johns, 2012) that the total activation of a memory item is the sum of that item's base activation and spreading activation which is determined by a function of the cue-weights (Dotlacil, 2021), and introduced a base activation parameter into the function that determined the probability of successful retrieval of the agent. This final version of our MPT model of the retrieval account could predict all effects found in the first experiment of Meng and Bader (2021), including those found in the plausibility judgment data, except for the interaction effects found in the agent/patient naming data. We explained that it was impossible for our model of the retrieval account to predict the interaction effects because we assumed a single parameter that represented the probability of a successful retrieval of the agent instead of assuming additional cognitive events or states which affected retrieval, and because in our model, a linear

function determined this parameter. We suggested on this basis that the idea of a nonlinear combination of retrieval cues, such as the one suggested by Parker (2019) must be implemented in our MPT model for it to account for the interaction effects if we are not to assume any additional cognitive events or states which affect retrieval in our MPT model. Finally, we did not propose a model in which Parker's (2019) idea was implemented because such a model would be beyond the scope of this thesis.

The parsing account of noncanonical word order sentence processing suggests that the decrease in performance observed for the task of agent/patient naming in the studies of noncanonical word order sentences (Ferreira, 2003; Christianson, Luke & Ferreira, 2010; Karimi & Ferreira, 2016) are a result of misinterpretation. Under the parsing account, most of the time comprehenders interpret sentences correctly through syntactic algorithms which construct detailed sentence representations, but sometimes misinterpretations arise because comprehenders resort to algorithms which construct sentence representations faster and more efficiently by making use of frequently found syntactic facts about a language, called heuristics (Ferreira 2003). The algorithms that produce detailed syntactic representations work alongside heuristics in a dual-route parsing mechanism, and their influence on the resulting interpretation is determined by the task-demands (Karimi & Ferreira, 2016). Two heuristics operate within the heuristic route of the parsing mechanism: the NVN heuristic and the semantic-association heuristic. The NVN heuristic always assigns the agent thematic role to the first noun encountered by the HPM in the sentence, while the semantic-association heuristic ignores syntactic structure entirely and assigns the agent thematic role to the noun that forms the most plausible agent-action relationship with the verb. It is also

suggested under the parsing account that task demands determine when during parsing the HPM enters a state of satisfaction with the output that is called equilibrium (Karimi & Ferreira, 2016). Parsing always begins through the heuristic route and only if the HPM does not reach equilibrium the algorithmic route begins working on an interpretation while the heuristic route continues to work, and so the interpretations that result from this mechanism are influenced by the two routes to a degree that is determined by when equilibrium is reached.

To implement the suggestions of the parsing account in an MPT model, we assumed four cognitive states each of which had an additional cognitive state as their counterpart. The first of the cognitive states represented the beginning of the heuristic route, which contrasted with a state of guessing that represented the times when the participant was inattentive to the task, so that only one of them could happen in any given trial. In our MPT model of the parsing account, it was assumed that the algorithmic route was engaged only after heuristic route began and only if equilibrium through the heuristic route was not achieved. The algorithmic route always resulted in the correct response to both the agent/patient naming and plausibility judgment tasks, whereas the outcome from the heuristic route was either the interpretation resulting from the semantic-association heuristic or the NVN heuristic. After fitting this model to the data from the first experiment of Meng and Bader (2021) we found that although the model was able to predict that the SO conditions yielded more correct responses to agent/patient naming across all meaning levels, and the two-way interaction between structure and meaning found in agent/patient Meng and Bader's (2021) study, it was unable to predict that the OS conditions yielded the lowest percentage of correct responses to agent/patient naming when compared to SO and passive conditions of the same meaning and structure

levels. This inability of the model to differentiate between the OS and passive conditions extended to the plausibility judgment predictions as well. Moreover, the model was unable to predict any significant differences between the nonreversible and biased conditions found in the plausibility judgment data.

Instead of addressing these problems with the predictions of our parsing account model right-away with our own suggestions, we resorted to implementing some of the further suggestions under the parsing account. The first suggestion we tried to implement was the task-dependency of the HPM under the parsing account. The suggestion under the parsing account was that the more demanding a task is, the harder it is for the HPM to reach equilibrium (Karimi & Ferreira, 2016), and we found it safe to assume that the agent/patient naming task is more demanding than the plausibility judgment task. We implemented this idea in our model by assuming different probabilities of the algorithmic route being engaged, or different probabilities of equilibrium through the heuristic route, as these two cognitive states contrasted in our model, without assuming any additional cognitive states or events. This assumption conflicted with the fact that both tasks were completed sequentially after reading a single sentence in the first experiment of Meng and Bader (2021), because it should be impossible for the HPM to adopt different thresholds for equilibrium with regard to the task at hand when the comprehender has to complete both tasks after reading a sentence. Nevertheless, we found it could still provide valuable insights about the performance of this suggestion with regard to its predictions if we assumed that the tasks were completed separately because the results from both tasks in the first experiment of Meng and Bader (2021) were similar to those in Bader and Meng's (2018) study where each of the two tasks were used in a separate experiment. This version of our parsing account model estimated

the probability that the algorithmic route is engaged to be higher for agent/patient naming task than that for the plausibility judgment task as we and the suggestion under the parsing account (Karimi & Ferreira, 2016) predicted, but the predictions for both tasks were only minimally improved as this version of the model could not predict any more of the effects found in Meng and Bader (2021) than the first version of the model. Moreover, as with the previous versions of our parsing account model, this version's estimate for the probability of a plausible guess was nearly zero. Therefore, we abandoned the idea that there could be different probabilities of the algorithmic route being engaged for the two tasks in our next versions of the parsing account model.

Another further suggestion of the parsing account that we tried implementing in our model was the suggestion of Ferreira (2003) that the NVN heuristic is more heavily weighted than the semantic-association heuristic in determining the outcome of the heuristic route. We implemented this idea by assuming that the semantic-association heuristic did not exist as a heuristic used within the heuristic route, and so, all of the outcomes from the heuristic route were the outcomes of the NVN heuristic in this third version of our model. This new version of the predicted that the probability of a guess is as high as 0.28, and the probability of a 'plausible' guess is as high as 0.88, which showed that the model compensated for the lack of a heuristic that always resulted in a 'plausible' response to the plausibility judgement task such as the semantic-association heuristic by assuming a high probability of a guess and a plausible guess. Moreover, with this implementation, the model's predictions became worse in terms of both the absolute values of the predicted percentages and the number of significant effects found in the data it can account for. Therefore, we suggested that it was unlikely that the NVN heuristic was more dominant than the

semantic-association heuristic in determining the outcome of the heuristic route and abandoned this implementation for the next version of our parsing account model.

In our fourth and final version of the parsing account model, in order to have the model account for more data, we went beyond the suggestions of the parsing account, and implemented the suggestion of Meng and Bader (2021) that the information structural markedness of OS sentences causes additional difficulty of comprehension in the absence of an appropriate discourse context, in a way that we found was in the spirit of the parsing account. We thought it could be suggested under the parsing account that the additional difficulty of comprehension for the OS sentences is reflected in the probability that the algorithmic route is engaged, or the probability that equilibrium is not reached through the heuristic route during parsing for sentences in this structure, because the parsing account suggests that the probability of an interpretation being the result of the algorithmic route is higher when the task is more demanding. Therefore, we assumed that there are different probabilities of the algorithmic route being engaged in SO and passive conditions on one hand, and OS conditions on the other. This assumption improved the model predictions significantly in that the model was now able to account for the differences between the OS and passive conditions. For the plausibility judgment data, this final version of the model was able to predict that the OS conditions had the highest percentage of ‘plausible’ responses in the implausible biased and nonreversible conditions, while also predicting that the percentage of ‘plausible’ responses were the lowest for the OS sentences in the plausible conditions. For the agent/patient naming data, this final version of the model, as with the first version of the model, could account for the two-way interaction between structure and meaning found in the first experiment of Meng and Bader (2021), but this version, unlike the

first version, extended this interaction effect to the OS structure as well, thus improving the fit. Furthermore, this final version of the parsing account model was able to account for the three-way interaction between structure, meaning and plausibility found in the data by predicting that the two-way interaction found in the plausible conditions was not available in the implausible conditions. However, the parameter estimates for this final version of the model were such that the probability of the algorithmic route being engaged in OS conditions was notably lower than that in SO and passive conditions, which is contrary to what the parsing account suggests. The parsing account suggests that the probability of an interpretation being the result of the algorithmic route should be higher when the task is more demanding, and we suggested on the basis of this idea that if comprehension of OS sentences becomes more difficult in the absence of an appropriate discourse context, then the probability of an interpretation being the result of the algorithmic route should be higher in the OS conditions. However, as mentioned earlier, the model predicted the exact reverse.

The modelling procedure for the retrieval account and the parsing account both have revealed valuable insights about the accounts and showed us what additional assumptions or modifications to their suggestions each account should make in order to account for the data from the first experiment of Meng and Bader (2021) but some of these insights perhaps deserve further discussion.

Assuming that our final MPT models for both accounts reflect the assumptions of the two accounts fairly accurately, we can say that the findings of this study suggest that the retrieval account is more eligible for adaptation into an MPT structure, and that its suggestions about the processing of noncanonical word order sentences are more eligible for further clarification or revision in order for it to

become a better account of the data. We suggest this on the basis of what we found throughout the modelling procedure for both accounts.

Firstly, our suggestion that there is a cognitive process whereby the plausibility of a sentence is judged with regard to world-knowledge after successful comprehension has significantly improved the predictions of the MPT model of the retrieval account. This suggestion is in-line with the suggestions of the retrieval account in that, under this account, the source of the performance effects found for both tasks are postinterpretive processes, or cognitive processes that come into play after an interpretation is formed (Bader & Meng, 2018; Meng & Bader, 2021), and the cognitive process that we suggested in our MPT model can also occur only after successful comprehension. Moreover, the retrieval account's suggestion that problems in retrieval of the agent or patient are the cause of the errors found in the agent/patient naming task for the noncanonical word order sentences is also in-line with this cognitive process that we suggested because we integrated this process into our MPT model in a way such that the retrieval operation is unaffected by world-knowledge match or mismatch. Furthermore, the fact that the cognitive process we suggested happens before the retrieval operation in our MPT model, is also in-line with the suggestion under the retrieval account that judging the plausibility of a sentence does not require a retrieval operation (Bader & Meng, 2018; Meng & Bader, 2021).

Secondly, although our assumption of an additional cognitive process whereby the pragmatic well-formedness of a sentence is judged, an assumption which we made with regard to the suggestion of Meng and Bader (2021) about the information structural markedness of OS sentences, did not result in better predictions for our MPT model of the retrieval account, taking this suggestion of

Meng and Bader as conveying a general difficulty of comprehension for OS sentences in the lack of an appropriate discourse context, which does not go against any of the suggestions under the retrieval account and can arguably be what was actually meant by the authors when they made the suggestion, and implementing it as such in our model has significantly improved the model's predictions.

Thirdly, assuming in our MPT model of the retrieval account that the probability of successful retrieval is a function of the number of cues that match the features of a memory item, the relative weights of these cues, and the base activation level of that item has significantly improved the model's predictions. Such an assumption about the probability of successful retrieval is in-line with the retrieval account because the suggestions of this account of noncanonical word order sentence processing are derived from the suggestions of the cue-based retrieval theory of sentence comprehension (Van Dyke, J. A., & Lewis, 2003; Lewis, Vasishth, & Van Dyke, 2006; Van Dyke & Johns, 2012), the suggestions of which are in-line with our assumptions in our MPT model of the retrieval account about the probability of successful retrieval of a memory item.

Finally, our final model of the retrieval account was unable to predict any of the interaction effects found in agent/patient naming data from the first experiment of Meng and Bader (2021), although our final model of the parsing account was able to do so. However, the reason why our final model of the retrieval account could not predict the interaction effects was that the function that returned the probability of successful retrieval in our model was a linear function, and implementing the idea of retrieval cues combining in a non-linear fashion, such as what Parker (2019) suggests, was beyond the scope of this thesis. Nevertheless, the idea of non-linear cue combination is both implementable in an MPT structure and in the spirit of the

cue-based retrieval theory of sentence comprehension (Van Dyke, J. A., & Lewis, 2003; Lewis, Vasishth, & Van Dyke, 2006; Van Dyke & Johns, 2012; Parker, 2019).

Therefore, we believe that the retrieval account is eligible for adaptation into an MPT structure, and that its suggestions about the processing of noncanonical word order sentences are eligible for further clarification or revision in order for it to become a better account of the effects found in studies of these types of sentences (Ferreira, 2003; Christianson, Luke & Ferreira, 2010; Bader and Meng, 2018; Meng and Bader, 2021; Cutter, Paterson & Filik, 2021).

On the contrary, throughout the modelling procedure for the MPT models of the parsing account we have encountered more obstacles to developing an MPT model that thoroughly accounts for the data from the first experiment of Meng and Bader (2021), although our final model of the parsing account was able to account for the interaction effects found for the agent/patient naming task in Meng and Bader (2021), while our final model of the retrieval account was not.

Firstly, under the parsing account, it is suggested that the algorithms that produce detailed syntactic representations work alongside heuristics in a dual-route parsing mechanism, and their influence on the resulting interpretation is determined by the task-demands (Karimi & Ferreira, 2016). Moreover, it is suggested under the account that both routes of processing continually influence each other until equilibrium is reached and an interpretation emerges from the HPM (Karimi & Ferreira, 2016). The only way we could implement these suggestions in an MPT model was to assume that the outcome of each trial is either the result of the algorithmic route or the heuristic route, and the outcome of the heuristic route is either the interpretation of the NVN heuristic or the semantic-association heuristic. This is because, in an MPT model, each outcome must have a unique set of cognitive

events or states that lead to it, and it must be possible to organize every possible outcome of each trial from the experiment into a finite number of discrete and observable categories (Riefer & Batchelder, 1988). However, from an MPT modelling perspective, it is unclear, in the suggestion of the parsing account, what is meant by the statements about the two processing routes continually influencing each other's output, or by the statements about the final interpretation of the HPM being influenced by both routes to varying degrees, depending on when equilibrium is reached during the parsing process. Therefore, we cannot be as sure as we were with the retrieval account about whether our MPT model of the parsing account reflects the assumptions of the parsing account about noncanonical word order sentence processing.

Secondly, our trials of implementing the suggestions of the parsing account about the task-dependency of the HPM, which stated that the more demanding a task is, the harder it is for the HPM to reach equilibrium, and about the weights of the NVN heuristic and the semantic-association heuristic within the heuristic route (Karimi & Ferreira, 2016) have revealed that these suggestions have provided little to no improvement to the predictions of our MPT model of the parsing account. We first tried assuming different probabilities of the algorithmic route being engaged for the agent/patient naming and the plausibility judgment task, and the estimated probabilities were such that this probability was higher for the agent/patient naming task than for the plausibility judgment task. However, this assumption did not lead to any improvement in the model's predictions. We then tried assuming that the semantic-association heuristic did not exist in the heuristic route, because the probability of a 'plausible' guess was estimated to be absurdly low in the first version of our model and found that this caused the estimate for the probability of a

guess to be absurdly high and estimate for the probability of a ‘plausible’ guess to also be very high. This showed that it is unlikely that the NVN heuristic is more dominant than the semantic-association heuristic within the heuristic route as suggested by Ferreira (2003), because, in addition to what this version of our MPT model of the parsing account showed, the semantic-association heuristic was clearly a better predictor of the actual results than the NVN heuristic in all other versions of the model.

Finally, we tried to go beyond the suggestions of the parsing account while still remaining in its spirit and synthesized Meng and Bader’s (2021) suggestion about the difficulty of comprehension in the lack of an appropriate discourse context for OS sentences, with the suggestion under the parsing account that the more demanding a task is, the harder it is for the HPM to reach equilibrium (Karimi & Ferreira, 2016). It follows from this suggestion of the parsing account that the difficulty in comprehension for the OS sentences should translate into a higher probability of the final interpretation of the HPM being a result of the algorithmic route. However, when we implemented this idea by assuming that the probability of the algorithmic route being engaged for the OS conditions is different than that for the SO and passive conditions, we found that the estimated probability of the algorithmic route being engaged for the OS conditions is notably lower than that for the SO and passive conditions, although the model’s predictions had improved significantly. In other words, the parameter estimates from this final version of our MPT model of the parsing account showed exactly the reverse of what the parsing account would suggest about the task-dependency of the HPM, although it was a better account of the data than the other versions of the model.

The process of creating MPT models of the retrieval and parsing accounts of noncanonical word order sentence processing have allowed us to organize and clarify the suggestions under both accounts in a model structure, such that we were able to extract information about the underlying cognitive processes suggested under the two accounts in the form of probabilities of occurrence when these suggestions are put together in a single cognitive model of a trial from the first experiment of Meng and Bader (2021). The process of creating and revising MPT models of both accounts has also allowed us to determine how each account can revise their suggestions about noncanonical word order sentence processing and what additional assumptions can be made for each account in order for them to become better predictors of the data. The findings throughout the entire modelling procedure have led us to suggest that the retrieval account of noncanonical word order sentence processing is more suitable for adaptation into an MPT structure than the parsing account, and that the MPT models of the retrieval account are more eligible for improvement than those of the parsing account, due to the theoretical background of the retrieval account.

CHAPTER 7

CONCLUSION

In conclusion, MPT modelling (Riefer & Batchelder, 1988; Riefer & Batchelder, 1999; Erdfelder et. al., 2009) of the data from the first experiment of (Meng & Bader, 2021) in which an agent/patient naming and a plausibility judgment was used to assess the comprehension of canonical and noncanonical word order sentences has allowed us to evaluate the explanatory power of the suggestions under the retrieval (Bader & Meng, 2018; Meng & Bader, 2021) and the parsing accounts (Ferreira, 2003; Christianson, Luke & Ferreira, 2010; Karimi & Ferreira, 2016) of noncanonical sentence processing.

Moreover, because adapting a verbal account of data into an MPT structure obligates breaking into pieces the suggested cognitive events and states under the said account, and the assumption of an order of occurrence for these, we were able to detect the clarity issues within these accounts. Our analysis has revealed that the retrieval account of noncanonical sentence processing has less clarity issues in this regard and therefore is more suitable for an MPT structure than the parsing account.

In addition, the analysis has revealed that the generalized language processing theory which the retrieval account is based on, the cue-based retrieval theory of sentence comprehension (Van Dyke, J. A., & Lewis, 2003; Lewis, Vasishth, & Van Dyke, 2006; Van Dyke & Johns, 2012) provides the theoretical background to devise MPT models that can fully explain the data used in this thesis, whereas the generalized language processing theory which the parsing account is based on, the Good-Enough approach to language comprehension, does not. Therefore, we conclude that the retrieval account in its current state and potential is a better account

of the data in-question, and future research that use the MPT modelling method to model noncanonical sentence processing data should focus on developing MPT models based on this account.

REFERENCES

- Bader, M., & Meng, M. (2018). The misinterpretation of noncanonical sentences revisited. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 44(8), 1286–1311. <https://doi.org/10.1037/xlm0000519>
- Batchelder, W. H., & Riefer, D. M. (1999). Theoretical and empirical review of multinomial process tree modeling. *Psychonomic Bulletin & Review*, 6(1), 57–86. <https://doi.org/10.3758/bf03210812>
- Betancourt, M. (2012). Cruising the simplex: Hamiltonian Monte Carlo and the Dirichlet distribution. *AIP Conference Proceedings*. <https://doi.org/10.1063/1.3703631>
- Christianson, K., Hollingworth, A., Halliwell, J. F., & Ferreira, F. (2001). Thematic roles assigned along the Garden Path Linger. *Cognitive Psychology*, 42(4), 368–407. <https://doi.org/10.1006/cogp.2001.0752>
- Christianson, K., Luke, S. G., & Ferreira, F. (2010). Effects of plausibility on structural priming. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 36(2), 538–544. <https://doi.org/10.1037/a0018027>
- Cutter, M. G., Paterson, K. B., & Filik, R. (2021). Online representations of non-canonical sentences are more than good-enough. *Quarterly Journal of Experimental Psychology*, 75(1), 30–42. <https://doi.org/10.1177/17470218211032043>
- Dotlačil, J. (2021). Parsing as a cue-based retrieval model. *Cognitive Science*, 45(8). <https://doi.org/10.1111/cogs.13020>
- Erdfelder, E., Auer, T.-S., Hilbig, B. E., Aßfalg, A., Moshagen, M., & Nadarevic, L. (2009). Multinomial Processing Tree Models. *Zeitschrift Für Psychologie / Journal of Psychology*, 217(3), 108–124. <https://doi.org/10.1027/0044-3409.217.3.108>
- Ferreira, F. (2003). The misinterpretation of noncanonical sentences. *Cognitive Psychology*, 47(2), 164–203. [https://doi.org/10.1016/s0010-0285\(03\)00005-7](https://doi.org/10.1016/s0010-0285(03)00005-7)

- Ferreira, F., & Patson, N. D. (2007). The ‘good enough’ approach to language comprehension. *Language and Linguistics Compass*, 1(1-2), 71–83. <https://doi.org/10.1111/j.1749-818x.2007.00007.x>
- Ferreira, F., Bailey, K. G. D., & Ferraro, V. (2002). Good-enough representations in language comprehension. *Current Directions in Psychological Science*, 11(1), 11–15. <https://doi.org/10.1111/1467-8721.00158>
- Isberner, M.-B., & Richter, T. (2013). Can readers ignore implausibility? evidence for nonstrategic monitoring of event-based plausibility in language comprehension. *Acta Psychologica*, 142(1), 15–22. <https://doi.org/10.1016/j.actpsy.2012.10.003>
- Karimi, H., & Ferreira, F. (2016). Good-enough linguistic representations and online cognitive equilibrium in Language Processing. *Quarterly Journal of Experimental Psychology*, 69(5), 1013–1040. <https://doi.org/10.1080/17470218.2015.1053951>
- Lewis, R. L., Vasishth, S., & Van Dyke, J. A. (2006). Computational principles of working memory in sentence comprehension. *Trends in Cognitive Sciences*, 10(10), 447–454. <https://doi.org/10.1016/j.tics.2006.08.007>
- Logacev, P., & Dokudan, N. (2021). A multinomial processing tree model of RC Attachment. *Proceedings of the Workshop on Cognitive Modeling and Computational Linguistics*. <https://doi.org/10.18653/v1/2021.cmcl-1.4>
- Meng, M., & Bader, M. (2021). Does comprehension (sometimes) go wrong for noncanonical sentences? *Quarterly Journal of Experimental Psychology*, 74(1), 1–28. <https://doi.org/10.1177/1747021820947940>
- Nelder, J. A., & Mead, R. (1965). A simplex method for function minimization. *The Computer Journal*, 7(4), 308–313. <https://doi.org/10.1093/comjnl/7.4.308>
- Paape, D., Avetisyan, S., Lago, S., & Vasishth, S. (2021). Modeling misretrieval and feature substitution in agreement attraction: A computational evaluation. *Cognitive Science*, 45(8). <https://doi.org/10.1111/cogs.13019>
- Parker, D. (2019). Cue combinatorics in memory retrieval for Anaphora. *Cognitive Science*, 43(3). <https://doi.org/10.1111/cogs.12715>

- R Core Team (2021). R: A language and environment for statistical computing. R Foundation for Statistical Computing, Vienna, Austria. URL <http://www.R-project.org/>
- Riefer, D. M., & Batchelder, W. H. (1988). Multinomial modeling and the measurement of cognitive processes. *Psychological Review*, 95(3), 318–339. <https://doi.org/10.1037/0033-295x.95.3.318>
- Sanford, A. J., & Sturt, P. (2002). Depth of processing in language comprehension: Not noticing the evidence. *Trends in Cognitive Sciences*, 6(9), 382–386. [https://doi.org/10.1016/s1364-6613\(02\)01958-7](https://doi.org/10.1016/s1364-6613(02)01958-7)
- Swets, B., Desmet, T., Clifton, C., & Ferreira, F. (2008). Underspecification of syntactic ambiguities: Evidence from self-paced reading. *Memory & Cognition*, 36(1), 201–216. <https://doi.org/10.3758/mc.36.1.201>
- Van Dyke, J. (2003). Distinguishing effects of structure and decay on attachment and repair: A cue-based parsing account of recovery from misanalyzed ambiguities. *Journal of Memory and Language*, 49(3), 285–316. [https://doi.org/10.1016/s0749-596x\(03\)00081-0](https://doi.org/10.1016/s0749-596x(03)00081-0)
- Van Dyke, J. A., & Johns, C. L. (2012). Memory interference as a determinant of language comprehension. *Language and Linguistics Compass*, 6(4), 193–211. <https://doi.org/10.1002/lnc3.330>
- Vehtari, A., Gelman, A., & Gabry, J. (2016). Practical bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*, 27(5), 1413–1432. <https://doi.org/10.1007/s11222-016-9696-4>