



UNIVERSIDADE FEDERAL DO CEARÁ
CENTRO DE HUMANIDADES
DEPARTAMENTO DE LETRAS VERNÁCULAS
PROGRAMA DE PÓS-GRADUAÇÃO EM LINGUÍSTICA

DANIELE LIMA MIRANDA

SINTAXE, EFEITO COGNATO E NÍVEL DE PROFICIÊNCIA EM L2
NO PROCESSAMENTO DE SINTAGMAS NOMINAIS
POR BILÍNGUES TARDIOS PORTUGUÊS - INGLÊS

FORTALEZA

2025

DANIELE LIMA MIRANDA

SINTAXE, EFEITO COGNATO E NÍVEL DE PROFICIÊNCIA EM L2
NO PROCESSAMENTO DE SINTAGMAS NOMINAIS
POR BILÍNGUES PORTUGUÊS - INGLÊS

Tese apresentada ao Programa de Pós-Graduação em Linguística da Universidade Federal do Ceará, como requisito parcial à obtenção do título de doutor em Linguística. Área de concentração: Aquisição, desenvolvimento e processamento da linguagem.

Orientadora: Profa. Dra. Pamela Freitas Pereira Toassi.

FORTALEZA

2025

Dados Internacionais de Catalogação na Publicação
Universidade Federal do Ceará
Sistema de Bibliotecas
Gerada automaticamente pelo módulo Catalog, mediante os dados fornecidos pelo(a) autor(a)

- M641s Miranda, Daniele Lima.
 Sintaxe, efeito cognato e nível de proficiência em L2 no processamento de sintagmas nominais por bilíngues tardios português-inglês / Daniele Lima Miranda. – 2025.
 197 f. : il. color.
- Tese (doutorado) – Universidade Federal do Ceará, Centro de Humanidades, Programa de Pós-Graduação em Linguística, Fortaleza, 2025.
 Orientação: Profa. Dra. Pamela Freitas Pereira Toassi.
1. Inglês como L2. 2. Psicolinguística do bilinguismo. 3. Processamento sintático bilíngue. 4. Efeito cognato. 5. Nível de proficiência em L2. I. Título.

CDD 410

DANIELE LIMA MIRANDA

SINTAXE, EFEITO COGNATO E NÍVEL DE PROFICIÊNCIA EM L2
NO PROCESSAMENTO DE SINTAGMAS NOMINAIS
POR BILÍNGUES PORTUGUÊS - INGLÊS

Tese apresentada ao Programa de Pós-Graduação em Linguística da Universidade Federal do Ceará, como requisito parcial à obtenção do título de doutor em Linguística. Área de concentração: Aquisição, desenvolvimento e processamento da linguagem.

Aprovada em 23 de abril de 2025.

BANCA EXAMINADORA

Profa. Dra. Pamela Freitas Pereira Toassi (Orientadora)
Universidade de Brasília (UnB)

Profa. Dra. Lídia Amélia de Barros Cardoso
Universidade Federal do Ceará (UFC)

Prof. Dr. José Ferrari Neto
Universidade Federal da Paraíba (UFPB)

Prof. Dr. Ronaldo Magueira Lima Júnior
Universidade de Brasília (UnB)

Prof. Dr. Bernardo Kolling Limberger
Universidade Federal de Ciências da Saúde de Porto Alegre (UFCSPA)

AGRADECIMENTOS

A Deus, em primeiro lugar, por Seu amor, Sua misericórdia e Seus milagres infinitos em minha vida.

À Nossa Senhora Auxiliadora, minha mãezinha no céu, que me cobre com Seu manto e me conduz com Seu cetro desde sempre.

À dona Socorro, minha mãezinha na terra, que me ensinou sobre o amor, sobre os estudos e sobre enfrentar a vida sem medo.

Ao meu companheiro de vida, Mário Filho, com quem partilho muitas alegrias e bênçãos cotidianas.

Aos meus filhos amados, Pedro e Felipe, meus amores incondicionais, que me ensinam tanto em todos os momentos da vida.

Aos meus amigos, Jarinívia, Cícero, Percília, Socorro, Lourdinha, Tobias, Camila e Neudson, pelo companheirismo e pelas gargalhadas.

À minha querida professora e orientadora, Pâmela, pela caminhada desde o Mestrado em Tradução, pelo suporte e paciência, pelas orientações e pelas conversas de incentivo e de força.

Aos meus colegas de pós-graduação e do Plibimult, pelas trocas. Em especial aos meus queridos colegas Leonardo, Daniela, Letícia, John, Marta, Regina e Nirla.

Aos professores Lídia Cardoso, Ferrari Neto, Ronaldo Lima Júnior e Bernardo Limberger, pela honra de participarem da minha banca de defesa.

Sou doutora em Linguística no dia de São Jorge! Que vitória! Que alegria! A vida presta!

RESUMO

Esta tese teve como principal objetivo analisar o processamento dos sintagmas nominais preposicionados com dois substantivos, como em *barra de chocolate*, por bilíngues tardios brasileiros com inglês como L2. Investigamos o custo de processamento das três estruturas sintáticas possíveis em L2: preposição *of*, caso genitivo e inversão de substantivos, incluindo como variáveis o efeito cognato e o nível de proficiência em L2. Os objetivos específicos foram: (1) identificar qual construção sintática em L2 tem o menor custo de processamento, dentre as três construções possíveis; (2) examinar se há efeito cognato no processamento desses sintagmas nominais; e (3) investigar se há relação entre o processamento de sintagmas nominais e o nível de proficiência em L2. Como principais referenciais teóricos deste trabalho tivemos: (1) o processamento sintático bilíngue (Bernolet e Hartsuiker, 2018); (2) os sintagmas nominais (Miranda, 2021); (3) o efeito cognato (Dijkstra *et al.*, 2015; Poort e Rodd, 2017; Soares *et al.*, 2018; Toassi e Pereira, 2019; Antón e Duñabeitia, 2020); e (4) o nível de proficiência em L2 (Casaponsa *et al.*, 2015; Lan *et al.*, 2019). Utilizamos três diferentes métodos experimentais: tarefa de reconhecimento de tradução, tarefa de julgamento de aceitabilidade e tarefa de seleção sintática. Para descrever o perfil dos participantes, utilizamos questionário demográfico e linguístico e teste de conhecimento de vocabulário de L2. Na tarefa de reconhecimento de tradução, os resultados indicaram que a acurácia foi maior e que o tempo de resposta foi menor para os sintagmas com preposição *of* e com inversão de substantivos e que houve efeito de facilitação cognata no tempo de resposta. Na tarefa de julgamento de aceitabilidade, a acurácia também foi maior nas sentenças com preposição *of* e com inversão de substantivos, o nível de proficiência em L2 levou a maiores índices de acurácia, o tempo de resposta foi menor em itens com inversão de substantivos e em itens com cognatos. Na tarefa de seleção sintática, o efeito cognato e o nível de proficiência em L2 influenciaram as opções escolhidas e o tempo de resposta sofreu interferência do nível de proficiência em L2. Os resultados encontrados podem reforçar o modelo de desenvolvimento sintático da L2 para bilíngues tardios de Bernolet e Hartsuiker (2018) e a hipótese de não seletividade das línguas em mentes bilíngues.

Palavras-Chave: inglês como L2; psicolinguística do bilinguismo; processamento sintático bilíngue; efeito cognato; nível de proficiência em L2.

ABSTRACT

The main objective of this doctoral dissertation was to analyze the processing of prepositioned noun phrases with two nouns, as in *barra de chocolate*, by Brazilian late bilinguals with English as L2. We investigated the processing cost of different syntactic structures in L2: preposition of, genitive case and noun inversion, including as variables the cognate effect and the proficiency level in L2. The main thematic axes of this work were: (1) bilingual syntactic processing (Bernolet and Hartsuiker, 2018); (2) noun phrases (Miranda, 2021); (3) the cognate effect (Dijkstra *et al.*, 2015; Poort and Rodd, 2017; Soares *et al.*, 2018; Toassi and Pereira, 2019; Antón and Duñabeitia, 2020); and (4) proficiency level in L2 (Casaponsa *et al.*, 2015; Lan *et al.*, 2019). Three different experimental methods were used: translation recognition task, acceptability judgment task and syntactic selection task. To describe the participants profile, demographical and linguistic questionnaire and L2 vocabulary knowledge test were used. In the translation recognition task, the results indicated that accuracy was greater, and the reaction time was shorter in noun phrases with preposition of and with noun inversions and that there was a cognate facilitation effect on reaction time. In the acceptability judgment task, accuracy was also greater in sentences with preposition of and with noun inversion, the level of proficiency in L2 led to higher accuracy rates, reaction time was shorter in items with noun inversion and in items with cognates. In the syntactic selection task, the cognate effect and the proficiency level in L2 influenced the options chosen and the reaction time was influenced by the proficiency level in L2. The results can reinforce the model of L2 syntactic development for late bilinguals by Bernolet and Hartsuiker (2018) and the hypothesis of non-selectivity of languages in bilingual minds.

Keywords: English as L2; psycholinguist of bilingualism; bilingual syntactic processing; cognate effect; proficiency level in L2.

LISTA DE FIGURAS

Figura 1 - Modelo de desenvolvimento sintático para bilíngues tardios.....	26
Figura 2 - Representação do sintagma nominal preposicionado em L1.....	29
Figura 3 - Representação do sintagma nominal em L2 (<i>of</i>).....	30
Figura 4 - Representação do sintagma nominal em L2 (inversão de substantivos).....	31
Figura 5 - Representação do sintagma nominal em L2 (caso genitivo).....	31
Figura 6 - Relações semânticas - SN1 de SN2.....	33
Figura 7 - Modelo Hierárquico Revisado.....	38
Figura 8 - <i>Bilingual Interactive Activation Plus</i>	42
Figura 9 – Multilink.....	43
Figura 10 - Tela com exemplo de item do experimento 1.....	61
Figura 11 - Tela com exemplo de item do experimento 2.....	71
Figura 12 - Tela com exemplo de item do experimento 3.....	77

LISTA DE GRÁFICOS

Gráfico 1 – Acurácia do experimento 1.....	91
Gráfico 2 – Histograma do tempo de resposta (ms) do experimento 1.....	92
Gráfico 3 - Acurácia por construção sintática do experimento 1.....	93
Gráfico 4 - Tempo de resposta (ms) por construção sintática do experimento 1.....	94
Gráfico 5 – Acurácia por condição cognata do experimento 1.....	96
Gráfico 6 – Tempo de resposta (ms) por condição cognata do experimento 1.....	97
Gráfico 7 – Acurácia por quantidade de cognatos do experimento 1.....	98
Gráfico 8 – Tempo de resposta (ms) por quantidade de cognatos do experimento 1.....	99
Gráfico 9 – Acurácia por condição sintática e cognata do experimento 1.....	102
Gráfico 10 – Tempo de resposta (ms) por condição sintática e cognata do experimento 1.....	104
Gráfico 11 - Probabilidade de acerto no experimento 1.....	106
Gráfico 12 - Previsão do tempo de resposta do experimento 1.....	109
Gráfico 13 – Acurácia do experimento 2.....	111
Gráfico 14 – Tempo de resposta (ms) do experimento 2.....	112
Gráfico 15 – Acurácia por construção sintática do experimento 2.....	114
Gráfico 16 – Tempo de resposta (ms) por construção sintática do experimento 2.....	115
Gráfico 17 – Acurácia por condição cognata do experimento 2.....	116
Gráfico 18 – Tempo de resposta (ms) por condição cognata do experimento 2.....	117
Gráfico 19 – Acurácia por quantidade de cognatos do experimento 2.....	118
Gráfico 20 – Tempo de resposta por quantidade de cognatos do experimento 2.....	119
Gráfico 21 – Acurácia por condição sintática e cognata do experimento 2.....	121
Gráfico 22 – Tempo de resposta (ms) por condição sintática e cognata do experimento 2.....	123
Gráfico 23 - Probabilidade de acerto no experimento 2.....	125
Gráfico 24 - Previsão do tempo de resposta (ms) do experimento 2.....	128
Gráfico 25 - Tempo de resposta (ms) do experimento 3.....	130
Gráfico 26 - Tempo de resposta (ms) por condição cognata do experimento 3.....	131
Gráfico 27 - Seleção sintática do experimento 3.....	132
Gráfico 28 - Tempo de resposta (ms) por seleção sintática do experimento 3.....	133
Gráfico 29 - Seleção sintática por condição cognata do experimento 3.....	135
Gráfico 30 - Tempo de resposta (ms) por seleção sintática e cognata do experimento 3...	136
Gráfico 31 - Probabilidade de seleção da preposição of no experimento 3.....	138
Gráfico 32 - Probabilidade de seleção do caso genitivo no experimento 3.....	139

Gráfico 33 - Probabilidade de seleção da inversão de substantivos no experimento 3.....	140
Gráfico 34 - Previsão do tempo de resposta (ms) do experimento 3.....	142

LISTA DE TABELAS

Tabela 1 – Lista de estímulos do experimento 1.....	62
Tabela 2 – Grau de similaridade dos cognatos.....	65
Tabela 3 – Lista de estímulos do experimento 2.....	72
Tabela 4 – Lista de estímulos do experimento 3.....	77
Tabela 5 – Grau de similaridade dos cognatos do experimento 3.....	80
Tabela 6 - Gênero dos participantes.....	85
Tabela 7 – Idade dos participantes.....	85
Tabela 8 – Semestre dos participantes.....	85
Tabela 9 – Mão utilizada para escrever.....	86
Tabela 10 – Nível inicial de estudo de inglês na escola.....	86
Tabela 11 – Motivos para estudar inglês.....	87
Tabela 12 – Frequência de leitura em inglês.....	88
Tabela 13 – Proficiência dos participantes.....	88
Tabela 14 – Status geral do experimento 1.....	91
Tabela 15 – Tempo de resposta do experimento 1.....	91
Tabela 16 – Status por construção sintática do experimento 1.....	93
Tabela 17 – Tempo de resposta (ms) por construção sintática do experimento 1.....	94
Tabela 18 – Status por condição cognata do experimento 1.....	95
Tabela 19 – Tempo de resposta (ms) por condição cognata do experimento 1.....	96
Tabela 20 – Status por quantidade de cognatos do experimento 1.....	97
Tabela 21 – Tempo de resposta (ms) por quantidade de cognatos do experimento 1.....	98
Tabela 22 – Status por condição sintática e cognata do experimento 1.....	100
Tabela 23 – Tempo de resposta (ms) por condição sintática e cognata do experimento 1.....	103
Tabela 24 – Modelo Misto Linear Generalizado (glmer) do experimento 1.....	105
Tabela 25 – Modelo Misto Linear (lmer) do experimento 1.....	107
Tabela 26 – Status geral do experimento 2.....	111
Tabela 27 – Tempo de resposta do experimento 2.....	112
Tabela 28 – Status por construção sintática do experimento 2.....	113
Tabela 29 – Tempo de resposta (ms) por construção sintática do experimento 2.....	114
Tabela 30 – Status por condição cognata do experimento 2.....	115
Tabela 31 – Tempo de resposta (ms) por condição cognata do experimento 2.....	116
Tabela 32 – Status por quantidade de cognatos do experimento 2.....	118

Tabela 33 – Tempo de resposta (ms) por quantidade de cognatos do experimento 2.....	119
Tabela 34 – Status por condição sintática e cognata do experimento 2.....	120
Tabela 35 – Tempo de resposta (ms) por condição sintática e cognata do experimento 2.	122
Tabela 36 - Modelo Misto Linear Generalizado (glmer) do experimento 2.....	124
Tabela 37 - Modelo Misto Linear (lmer) do experimento 2	126
Tabela 38 - Tempo de resposta (ms) do experimento 3.....	129
Tabela 39 - Tempo de resposta (ms) por condição cognata do experimento 3.....	131
Tabela 40 - Seleção sintática do experimento 3.....	132
Tabela 41 - Tempo de resposta (ms) por seleção sintática do experimento 3.....	133
Tabela 42 - Seleção sintática por condição cognata do experimento 3.....	134
Tabela 43 - Tempo de resposta (ms) por seleção sintática e cognata do experimento 3.....	135
Tabela 44 - Modelo Logístico Multinomial (polytomous) do experimento 3.....	137
Tabela 45 - Modelo Misto Linear (lmer) do experimento 3.....	141

SUMÁRIO

1	INTRODUÇÃO	16
2	FUNDAMENTAÇÃO TEÓRICA.....	21
2.1	A interação das línguas na mente bilíngue.....	23
2.2	Processamento sintático bilíngue	24
2.3	Sintagmas nominais preposicionados com dois substantivos	27
2.4	Efeito cognato.....	38
2.5	Nível de proficiência em L2	47
3	METODOLOGIA.....	51
3.1	Objetivos.....	51
3.1.1	<i>Objetivo geral</i>	51
3.1.2	<i>Objetivos específicos</i>	52
3.2	Questões e hipóteses da pesquisa.....	52
3.3	Participantes	53
3.4	Desenho do estudo e coleta de dados.....	54
3.5	Análise e processamento de dados	55
3.6	Documentos, instrumentos e materiais	55
3.6.1	<i>TCLE e TALE</i>	55
3.6.2	<i>Questionário demográfico e linguístico</i>	56
3.6.3	<i>Teste de conhecimento de vocabulário</i>	56
4	OS EXPERIMENTOS	59
4.1	Experimento 1	59
4.1.1	<i>Tarefa de reconhecimento de tradução</i>	59
4.1.2	<i>Desenho e corpus do experimento 1</i>	60
4.2	Experimento 2	69
4.2.1	<i>Tarefa de julgamento de aceitabilidade</i>	70
4.2.2	<i>Desenho e corpus do experimento 2</i>	71
4.3	Experimento 3	75
4.3.1	<i>Seleção sintática</i>	76
4.3.2	<i>Desenho e corpus do experimento 3</i>	76
5	RESULTADOS	84
5.1	Perfil dos participantes	84
5.1.1	<i>Questionário demográfico e linguístico</i>	84

5.1.2	<i>Teste de conhecimento de vocabulário</i>	88
5.2	Experimento 1	89
5.2.1	<i>Análise geral</i>	90
5.2.2	<i>Análise por construção sintática</i>	92
5.2.3	<i>Análise por condição cognata</i>	95
5.2.4	<i>Análise por quantidade de cognatos</i>	97
5.2.5	<i>Análise por sintaxe e quantidade de cognatos</i>	99
5.2.6	<i>Análise inferencial do experimento 1</i>	104
5.3	Experimento 2	109
5.3.1	<i>Análise geral</i>	111
5.3.2	<i>Análise por construção sintática</i>	113
5.3.3	<i>Análise por condição cognata</i>	115
5.3.4	<i>Análise por quantidade de cognatos</i>	117
5.3.5	<i>Análise por sintaxe e quantidade de cognatos</i>	120
5.3.6	<i>Análise inferencial do experimento 2</i>	123
5.4	Experimento 3	128
5.4.1	<i>Análise geral</i>	129
5.4.2	<i>Análise por seleção sintática</i>	132
5.4.3	<i>Análise por seleção sintática e quantidade de cognatos</i>	134
5.4.4	<i>Análise inferencial do experimento 3</i>	137
6	DISCUSSÃO DOS RESULTADOS	144
6.1	Tarefa de reconhecimento de tradução	145
6.2	Tarefa de julgamento de aceitabilidade	148
6.3	Tarefa de seleção sintática	152
6.4	Comparação dos três experimentos	154
7	CONSIDERAÇÕES FINAIS	157
	REFERÊNCIAS	162
	APÊNDICE A - TERMO DE CONSENTIMENTO LIVRE E ESCLARECIDO (TCLE)	172
	APÊNDICE B - TERMO DE ASSENTIMENTO LIVRE E ESCLARECIDO	176
	APÊNDICE C - QUESTIONÁRIO DEMOGRÁFICO E LINGUÍSTICO E LINGUÍSTICO	179
	APÊNDICE D - ITENS DISTRATORES DOS EXPERIMENTOS 1 E 2 ...	180

APÊNDICE E - FREQUÊNCIA DOS SUBSTANTIVOS DOS EXPERIMENTOS 1 E 2	182
APÊNDICE F - NÚMEROS DE CARACTERES DOS ITENS EXPERIMENTAIS DOS EXPERIMENTOS 1 E 2	186
APÊNDICE G - ITENS DISTRATORES DO EXPERIMENTO 3	190
APÊNDICE H - FREQUÊNCIA DOS SUBSTANTIVOS DO EXPERIMENTO 3	192
APÊNDICE I - NÚMEROS DE CARACTERES DOS ITENS EXPERIMENTAIS DO EXPERIMENTO 3	195

1 INTRODUÇÃO

O processamento de L2 costuma apresentar alguns desafios para bilíngues tardios, especialmente em relação a algumas propriedades sintáticas específicas que podem se apresentar distintamente entre a L1 e a L2 do usuário. Este é o caso do processamento de sintagmas nominais preposicionados com dois substantivos por bilíngues tardios, que são brasileiros e que têm inglês como L2.

Com relação aos objetos de pesquisa, apresentamos inicialmente três definições importantes para este estudo: bilíngues, bilíngues tardios e sintagmas nominais preposicionados com dois substantivos.

Para a definição de bilíngues, consideramos o que é postulado por Grosjean e Li (2013), que argumentam que ser bilíngue é conhecer e utilizar duas línguas ou dialetos no cotidiano. E para a definição de bilíngues tardios, levamos em consideração o que propõe Wei (2000), que defende que bilíngues tardios são aqueles indivíduos que adquirem L2 depois dos 13 anos.

Downing (2015) apresenta duas definições de sintagmas nominais voltadas à semântica e à sintaxe, respectivamente. Para a autora, semanticamente, sintagmas nominais são palavras ou grupos de palavras que se referem a algo ou a alguém, a abstrações, a atividades, a emoções e a qualidades. Sintaticamente, trata-se de uma unidade linguística que tem um núcleo que pode ser acompanhado por determinantes ou modificadores. Estes modificadores nominais podem vir antes ou após o núcleo. Downing (2015) explica que o pré-modificador descreve ou classifica o núcleo do sintagma nominal e o pós-modificador acrescenta alguma informação sobre esse núcleo, apresentando sua identificação ou sua definição.

Descobrimos, durante as leituras realizadas para esta tese, que a estrutura sintática que investigamos pode receber diferentes designações. Em Resende *et al.* (2020), os autores se referem a este tipo de construção sintática como sintagma nominal com relação de posse entre os substantivos. Mantivemos esta nomenclatura em minha dissertação (Miranda, 2021). Em Santos e Alonso (2022), as autoras esclarecem que este tipo de construção sintática abarca um campo maior de relações semânticas, que não se limitam a posse, definindo a referida estrutura como construção relacional binominal SN1 de SN2. Em outros estudos, como em De Cat *et al.* (2015), por exemplo, esta mesma estrutura sintática recebe nomenclatura diferente: palavras compostas. Nas pesquisas de Fiorentino e Poeppel (2007) e Li *et al.* (2017), a definição para palavras compostas pode se referir também a expressões formadas por uma única palavra,

que deriva de outras duas, como em guarda-costas (*bodyguard*) ou palavra-chave (*keyword*), para exemplificar.

Nesta tese, a decisão foi de utilizar o termo sintagma nominal preposicionado com dois substantivos. Seleccionamos o termo sintagma nominal porque esta definição parece ser mais comum do que palavras compostas, mas optamos por especificar ainda mais de que tipo de sintagma nominal se trata: aquele que é preposicionado em L1 e é constituído por dois substantivos que são ligados por uma preposição (*de* e suas variações *da* ou *do*), como em *porta da casa*, *barra de chocolate* e *cadeira do escritório*.

Dessa forma, o processamento de sintagmas nominais preposicionados com dois substantivos, como em *barra de chocolate*, constituem o principal tema desta tese, que tem viés psicolinguístico. As línguas dos participantes desta pesquisa, o português, que é a língua materna ou língua nativa (a partir de agora, L1) e o inglês, que é a língua não nativa¹ (daqui em diante, L2), possuem uma construção sintática em comum para a representação dessas estruturas, pois ambos os idiomas podem expressar a relação polissêmica entre os substantivos, através do uso da preposição *de* (em L1) e *of* (em L2). Por exemplo, *barra de chocolate* em L1 pode ser expressa em L2 como *bar of chocolate*. Neste caso, as estruturas sintáticas são semelhantes e a ordem das palavras se mantém. Entretanto, enquanto em português o sintagma nominal preposicionado com dois substantivos é exclusivamente representado pelo uso da preposição *de* e suas variações (*da* e *do*), em inglês há mais duas construções possíveis: o caso genitivo (como em *chocolate's bar*²) e a inversão de substantivos (como em *chocolate bar*).

É possível observar nas salas de aula de inglês como L2 que a possibilidade da representação variada na língua não-nativa da relação polissêmica estabelecida entre os dois substantivos que formam o sintagma nominal vai passando por transformações no processamento e no uso, de forma natural e gradual, à medida que o contato com a L2 vai se intensificando. Como duas das três construções sintáticas possíveis não existem em L1, algumas dificuldades costumam aparecer. Uma das consequências é o uso excessivo da construção que mais se assemelha à L1, especialmente em níveis iniciais de aprendizagem da L2. Isso pode indicar influência interlinguística da L1 na L2, assim como uma possível comparação entre os idiomas, como proposto no modelo de desenvolvimento sintático da L2 para bilíngues tardios, de Bernolet e Hartsuiker (2018).

A investigação dos sintagmas nominais preposicionados teve início em minha

¹ Nesta tese, não fizemos distinção entre segunda língua, língua estrangeira ou língua adicional, conforme Lima Júnior *et al.* (2021).

² Esta construção não é aceita na maioria dos contextos possíveis, mas está prevista na interlíngua dos bilíngues.

dissertação (Miranda, 2021), tendo como principal motivação a investigação científica do que era possível observar em minha prática de sala de aula de L2. Assim, esta tese dá continuidade à investigação iniciada em minha dissertação, com a expectativa de abordar algumas limitações detectadas e de ampliar alguns aspectos do estudo inicial. Então, resolvemos incluir nesta investigação dos sintagmas nominais preposicionados por dois substantivos o efeito cognato e o nível de proficiência em L2. Minha tese é de que as diferentes estruturas sintáticas em L2 para representar os sintagmas nominais preposicionados com dois substantivos têm custos de processamento diferentes, podendo haver interferência do efeito cognato e do nível de proficiência em L2.

Para Dijkstra e Heuven (2002), os cognatos são palavras de diferentes idiomas que se sobrepõem consideravelmente em forma e significado. A palavra *piano*, para dar um exemplo, é representada da mesma forma em inglês e em português, sendo considerada cognata perfeita, pois, além de possuir a mesma grafia nos dois idiomas, possui também o mesmo significado. Portanto, podemos deduzir que a palavra *piano* seria internalizada e compreendida sem muito esforço cognitivo por um brasileiro falante de PB que tenha contato com esse item lexical pela primeira vez. Assim, os cognatos podem funcionar como uma ponte entre dois idiomas, facilitando a conscientização de semelhanças linguísticas, o que pode gerar benefícios como a identificação de ligações entre o vocabulário e os aspectos sintáticos.

Há três classificações diferentes para os cognatos, segundo Dijkstra *et al.* (2010): cognatos idênticos, muito semelhantes e não idênticos. Os cognatos idênticos são aqueles que se sobrepõem ortograficamente e semanticamente, como *chocolate*, em português, que em inglês é *chocolate*. Os cognatos muito semelhantes são diferentes entre si apenas em uma ou duas letras, como em *museu* em português, que em inglês é *museum*. Já os cognatos não idênticos possuem mais do que duas letras diversas entre si, como em *identificação* em L1 e *identification* em L2. Nesta pesquisa, utilizaremos os três tipos de cognatos.

Para explicar os efeitos de facilitação cognata dentro da estrutura e da arquitetura do léxico mental bilíngue, há alguns modelos que defendem a ativação mental de uma palavra em uma língua e a ativação simultânea de sua tradução correspondente no outro idioma, tais quais: Kroll e Stewart (1994) e seu Modelo Hierárquico Revisado; o modelo *Bilingual Interactive Activation* (BIA, daqui por diante) de Dijkstra e Heuven (1998) e seus aprimoramentos, BIA+, de Dijkstra e Heuven (2002), BIA –d, de Grainger *et al.* (2010), e BIA+S, de Casaponsa *et al.* (2020); e o modelo *Multilink*, de Dijkstra *et al.* (2018).

O efeito cognato é comumente chamado de efeito de facilitação cognata nas pesquisas linguísticas, devido à sua constante comprovação de viabilização de menor custo

interlinguístico. Dijkstra *et al.* (2015), Poort e Rodd (2017), Soares *et al.* (2018), Toassi e Pereira (2019) e Antón e Duñabeitia (2020) são algumas das pesquisas que detectaram que os cognatos podem facilitar o processamento de L2.

A mensuração do nível de proficiência em L2 é um dispositivo metodológico bastante comum nos estudos sobre aquisição e processamento de L2, como em Lan *et al.* (2019), que identificaram que o processamento de sintagmas nominais é dependente do nível de proficiência dos bilíngues. Dentro da Psicolinguística do Bilinguismo, a proficiência em L2 pode ser compreendida como uma dimensão cognitiva capaz de moldar diferencialmente a arquitetura do processamento da linguagem. Em outras palavras, a proficiência em L2 pode estar relacionada à disponibilidade de representações gramaticais da L2 e estar ligada à fluidez de acesso a essas representações nos eventos de processamento linguístico em tempo real.

Souza (2019) explica que a noção de proficiência em L2 é remissiva ao campo da testagem de habilidades linguísticas. Dessa maneira, trata-se de um conceito que indica clara conotação psicométrica, uma vez que se trata de um construto ou um traço latente. Neste estudo, o nível de proficiência em L2 foi medido por um teste de conhecimento de vocabulário de inglês com 40 itens, que foi baseado no teste de vocabulário do *Shipley Institute of Living Scale* (SILS), e adaptado pela professora Pamela Toassi (UnB), em parceria com os professores Justin Lauro (*Barrys University*) e Bernhard Angele (*Universidad de Madrid*).

Com esta tese, objetivamos analisar o processamento de sintagmas nominais preposicionados com dois substantivos por bilíngues tardios brasileiros, com inglês como L2. De forma mais detalhada, objetivamos: (1) identificar qual construção sintática em L2 tem menor custo de processamento; (2) examinar como o efeito cognato pode interferir no processamento desses sintagmas nominais; e (3) investigar qual a relação entre o custo de processamento de sintagmas nominais e o nível de proficiência em L2.

As perguntas de pesquisa são: (1) dentre as três construções sintáticas possíveis em L2 (preposição *of*, caso genitivo e inversão de substantivos) para a representação de sintagmas nominais preposicionados com dois substantivos, qual tem menor custo de processamento?; (2) como os cognatos interferem no processamento de sintagmas nominais em L2?; e (3) como o nível de proficiência em L2 afeta o processamento de sintagmas nominais?

As hipóteses que consideramos para responder as perguntas de pesquisa são: (1) o custo cognitivo da estrutura com a preposição *of* deve ser menor, porque se assemelha mais com a L1, com base em Bernolet e Hartsuiker (2018), que apontam que a compreensão da L2 pode ser impulsionada pela sintaxe da L1, especialmente em estágios iniciais de aquisição da L2; (2) os cognatos interferem no processamento de L2, aumentando a acurácia e diminuindo

o tempo de resposta de diferentes itens lexicais e estruturas sintáticas, tendo como referências Dijkstra *et al.* (2015), Poort e Rodd (2017), Soares *et al.* (2018), Toassi e Pereira (2019) e Antón e Duñabeitia (2020), que detectaram que os cognatos podem facilitar o processamento de L2; e (3) os bilíngues com menor nível de proficiência em inglês como L2 apresentam menos acurácia e precisam de mais tempo para processar os sintagmas nominais, com base no modelo de aquisição de L2 para bilíngues tardios de modelo de aquisição de L2 para bilíngues tardios, que indica diferentes processamentos de L2 para diferentes estágios de proficiência, de Bernolet e Hartsuiker (2018) e na pesquisa feita por Lan *et al.* (2019), que identificou que o processamento de sintagmas nominais é dependente do nível de proficiência em L2 dos bilíngues.

Para o alcance dos objetivos propostos, realizamos três experimentos com bilíngues tardios que possuem diferentes níveis de proficiência em L2. Os métodos experimentais foram: uma tarefa de reconhecimento de tradução, uma tarefa de julgamento de aceitabilidade e uma tarefa de seleção sintática. Utilizamos como documentos protocolares os termos de consentimento e assentimento, teste de conhecimento de vocabulário em L2 e questionário demográfico e linguístico.

Esta pesquisa está dividida em sete capítulos. Este é o capítulo 1, a introdução, que faz as considerações iniciais a respeito deste estudo, apresentando uma síntese do trabalho que foi desenvolvido e como ele está organizado. No capítulo 2, apresentamos os principais referenciais teóricos sobre o processamento sintático bilíngue, os sintagmas nominais, o efeito cognato e o nível de proficiência em L2. No capítulo 3, a metodologia é apresentada, trazendo informações acerca dos participantes, dos métodos e dos instrumentos que foram utilizados. No capítulo 4, os experimentos que foram aplicados e os métodos experimentais utilizados estão descritos. Os resultados das análises dos dados coletados são apresentados no capítulo 5 e no capítulo 6, apresentamos as discussões dos resultados. Finalmente, no capítulo 7, apresentamos as considerações finais. As referências bibliográficas e os apêndices estão especificados ao final deste trabalho.

2 FUNDAMENTAÇÃO TEÓRICA

A tarefa de aprender um novo idioma é normalmente associada à memorização de um novo vocabulário. Para Fernández e Cairns (2011), aprender as estruturas sintáticas da nova língua é igualmente importante. Desta maneira, para se comunicar em L2, tanto o léxico quanto a sintaxe são importantes, pois os itens lexicais devem ser acessados e organizados sintaticamente para que o bilíngue se comunique de forma satisfatória.

Mas como diferentes línguas interagem na mente de bilíngues? A Psicolinguística pode nos ajudar a responder esta pergunta porque é a ciência que estuda a interseção entre a Linguística e a Psicologia e tem como principal interesse o processamento e as representações do conhecimento que estão implícitas na habilidade de usar a linguagem, conforme Cowles (2011). Além disso, essa ciência também se interessa em como a linguagem interage com outros aspectos da cognição humana, caracterizando-se como uma área de investigação interdisciplinar, como afirmam Grosjean e Li (2013). A Psicolinguística busca compreender como as pessoas adquirem a língua, como ela é utilizada para falar e entender as outras pessoas e como ela é representada e processada no cérebro, segundo Fernández e Cairns (2011). Para Maia (2018), a Psicolinguística é a ciência que investiga a aquisição a aprendizagem e o processamento da linguagem verbal e se utiliza da pesquisa experimental e da análise de dados quantitativos para melhor compreender os processos cognitivos da linguagem. Dessa forma, somente através do estudo da Psicolinguística do Bilinguismo é possível a compreensão do processamento e do uso de mais de uma língua na mente humana.

A Psicolinguística do Bilinguismo estuda como as pessoas aprendem e usam as línguas que dominam, bem como os processos linguísticos e cognitivos envolvidos neste fenômeno. Os psicolinguistas estudam os mecanismos e os processos cognitivos que estão envolvidos na compreensão, na produção, na aquisição e no uso da linguagem, como pontuam Godfroid e Hopp (2023). Estes pesquisadores compartilham a crença de que desvendar a arquitetura e os processos mentais dos falantes é fundamental para compreender como ocorre o processamento da L2.

Para Tokowicz (2015), pesquisar a operação do sistema linguístico bilíngue inclui investigar cada conceito ligado aos vários itens lexicais das línguas que o bilíngue utiliza. Assim, como os bilíngues conseguem transitar entre as línguas que domina, mesmo que com algum esforço, fenômenos como o *code-switching*³ e o uso de empréstimos linguísticos

³ O *code-switching* é um fenômeno que consiste na alternância entre as línguas de um usuário e é comum em bilíngues altamente proficientes.

podem ser frequentes e se tornam características da comunicação bilíngue.

A Hipótese do Duplo Monolíngue, proposta por Saer (1923), é uma das primeiras definições reportadas na literatura para caracterizar o bilinguismo⁴. Esta hipótese caracterizou o indivíduo bilíngue como dois monolíngues em uma mesma pessoa. Uma década depois, Bloomfield (1933) afirmava que os bilíngues eram aqueles que controlavam duas línguas como um nativo.

Essas concepções tradicionais que propagavam o bilíngue com uma visão monolíngue perduraram por muitas décadas. Entretanto, o modelo de perfeição da competência linguística e o desempenho equivalente nas línguas que fala, em comparação aos falantes monolíngues, começou a ser questionado, como explicam Finger e Ortiz-Preuss (2018). Desta maneira, esse entendimento mítico de bilinguismo perfeito não encontra mais respaldo nos estudos mais recentes.

Grosjean (2008) considera que um sujeito bilíngue não pode ser tido como um “monolíngue em duas línguas”. Isto quer dizer que uma pessoa bilíngue é um todo que não pode ser dividido em duas partes separadas. Para Grosjean e Li (2013), o uso e a fluência da L2 determinam que um falante é bilíngue. Mas os autores afirmam que, como os bilíngues utilizam suas línguas com diferentes propósitos, em diferentes domínios da vida e para realizar coisas diferentes, é natural que seu nível de fluência na L2 dependa da necessidade de uso dessa língua. Assim, muitos bilíngues são mais fluentes em uma de suas línguas do que na outra. Grosjean e Li (2013), então, definem que ser bilíngue é utilizar duas ou mais línguas (ou dialetos) no dia a dia.

Para classificar os tipos de bilinguismo, Wei (2000) lista como critérios a idade, o contexto de aquisição da L2, o grau de uso da L2 e o nível de conhecimento da L2. Considerando a idade, as autoras classificam os bilíngues em quatro tipos: (1) precoce, que é quando a aquisição de L2 ocorre entre 0 e 12 anos; (2) tardio, quando a aquisição acontece depois dos 13 anos; (3) simultâneo, quando há a aquisição de duas línguas desde o nascimento; e (4) sucessivo, quando há a aquisição de uma língua e depois de outra.

Para o mesmo autor, o contexto de aquisição da L2 também pode ser chamado de composto ou coordenado. O contexto de aquisição é composto quando as duas línguas são adquiridas simultaneamente e, muitas vezes, no mesmo contexto. Já o contexto de aquisição coordenado de duas línguas é aquele que ocorre em contextos separados, de formas diferentes.

Com relação ao grau de uso das línguas, o autor classifica o bilinguismo como

⁴ Embora os termos bilinguismo e multilinguismo não sejam sinônimos, não trataremos desta distinção neste estudo, como em Grosjean e Li (2013).

dominante ou recessivo. A primeira classificação se dá quando há alto nível de uso e proficiência em uma das línguas, enquanto a segunda se refere ao pouco uso de uma das línguas, o que pode gerar dificuldade de compreensão nessa língua pouco utilizada.

Para Wei (2000), há três tipos de bilinguismo, se levarmos em consideração o nível de conhecimento das línguas: (1) produtivo, que é quando há domínio de produção das habilidades oral ou escrita na L2; (2) receptivo, quando o domínio da L2 está voltado à compreensão ou percepção oral ou escrita; e (3) incipiente, que é quando o bilíngue está em fase inicial de aquisição da L2, ou seja, a língua não nativa ainda está em desenvolvimento.

Neste trabalho, levamos em consideração a definição de Grosjean e Li (2013), que argumentam que ser bilíngue é conhecer e utilizar duas línguas ou dialetos no cotidiano. Além disso, retomamos no capítulo sobre a metodologia as diversas características de um bilíngue que foram apresentadas aqui, a fim de melhor descrever os participantes desta pesquisa.

2.1 A interação das línguas na mente bilíngue

Na busca pelo entendimento de como diferentes línguas interagem em um só cérebro, muitas questões podem ser levantadas. Um dos principais debates ocorre em torno da ativação seletiva ou não seletiva das línguas em bilíngues. Isto ocorre porque quando alguém que usa mais de um idioma escuta ou lê em suas línguas, palavras dos idiomas que ele domina competem por ativação e seleção na mente bilíngue.

Por um lado, Fernández e Cairns (2011) esclarecem que os modelos de ativação seletiva da língua sugerem que a seleção de uma entrada lexical afeta exclusivamente a entrada lexical correspondente na língua pretendida. Nessa visão, o acesso lexical bilíngue ocorreria através de processos semelhantes ao acesso lexical monolíngue durante a seleção lexical.

Por outro lado, os modelos de ativação não seletiva da língua defendem que os sistemas dos dois léxicos dos bilíngues são ativados, em que o processamento da L2 não é completamente autônomo. Ao contrário, o processamento da L2 é afetado em diferentes níveis pela L1, que foi adquirida anteriormente. Dessa maneira, o acesso lexical de bilíngues seria um processo complexo, porque a seleção de um dado conceito poderia ativar duas entradas lexicais, no mínimo. Schwartz e Kroll (2006) debatem a alternância de códigos em bilíngues, questionando como usuários de duas línguas diferem homógrafos interlinguísticos ou falsos cognatos, por exemplo. Para essa distinção, as autoras chamam atenção da função do contexto no acesso lexical, bem como do grau de compreensão da linguagem.

É comum que pesquisas que investigam a interação entre as duas línguas de um

bilíngue utilizem palavras que compartilhem características lexicais nas duas línguas como estímulos. Schwartz e Kroll (2006) explicam que o raciocínio por trás dessas pesquisas é de que se os bilíngues ativam uma única língua, as palavras que compartilham propriedades em duas línguas não apresentam efeitos. Nesse caso de ativação paralela das duas línguas, palavras controle⁵ seriam processadas de modo diferente das palavras que compartilham características lexicais, o que poderia ser verificado pela acurácia no desempenho ou pelo tempo de processamento.

Há evidências consistentes de uma ativação não seletiva das línguas de um bilíngue e de um processamento paralelo, por meio de estudos de reconhecimento de palavras, como pontuam Schwartz e Kroll (2006). Entretanto, a maioria dessas pesquisas foi realizada com palavras isoladas como estímulo. Por outro lado, quando as palavras estão no contexto de uma sentença, há evidência de que o sistema linguístico pode operar de forma mais seletiva em relação à língua. No entanto, esse efeito parece surgir mais tardiamente no processamento da sentença, uma vez que os processos iniciais de acesso lexical não descartariam totalmente uma ativação inicial não seletiva dos candidatos dos dois idiomas, seguida pela inibição de um deles.

2.2 Processamento sintático bilíngue

Sob a perspectiva gerativa, a sintaxe é o centro da cognição da linguagem humana, uma vez que retira informações do léxico e da morfologia para alimentar os sistemas fonológicos e semântico com representações linguísticas, conforme explica Kenedy (2013). A aquisição da sintaxe de uma língua é um processo complexo e multifacetado, pois é necessário compreender o significado dos itens lexicais, formar categorias sintáticas, aprender e aplicar o conjunto de regras sintáticas, estabelecendo a correta relação que deve ser estabelecida entre os elementos que formam uma sentença, como explicam Bernolet e Hartsuiker (2018). Para crianças, a aprendizagem costuma ocorrer mais rapidamente e sem instruções explícitas, enquanto para bilíngues tardios de L2, a aquisição de sintaxe requer muito esforço e atenção, e muitas vezes algum tipo de instrução explícita.

Quando as duas línguas do bilíngue são processadas sintaticamente de formas diferentes, algumas questões podem surgir. Um dos objetos de investigação que são estudados comumente é de que os bilíngues possam processar sintaticamente a L2 por meio de sua L1, como explicam Schwartz e Kroll (2006). Outra possibilidade seria de que sejam desenvolvidas

⁵ Christensen *et al.* (2013) esclarecem que palavras controle são itens que não estão dentro do mesmo escopo das palavras investigadas em um estudo experimental, para que seja possível a identificação dos efeitos das variáveis exploradas, o que viabiliza resultados mais confiáveis.

estratégias de processamento sintático exclusivas a quem fala mais de uma língua. Essa estratégia seria utilizar uma mistura de construções das duas línguas, de acordo com as preferências dos bilíngues, que podem ir mudando conforme a proficiência em L2 vai se tornando mais alta.

Ademais, o conhecimento de conceitos e sistemas sintáticos pode ajudar no aprendizado de um segundo idioma, mas também pode atrapalhar os bilíngues tardios. Isto significa que o conhecimento da sintaxe da L1 pode levar a suposições fundamentadas para a compreensão e a produção da sintaxe em L2, o que pode exigir pensamento e esforço mais conscientes do que o processamento de tentativa e erro que caracteriza a aquisição da sintaxe na L1.

Bernolet e Hartsuiker (2018) questionam como as representações na segunda língua se desenvolvem nesses bilíngues tardios e qual é a relação entre a sintaxe da L1 e da L2, o que se conecta a uma das questões que mais despertam interesse nos estudos da Psicolinguística do Bilinguismo: como os bilíngues representam e processam suas línguas mentalmente. Há um único sistema na mente bilíngue para todas as línguas ou há um sistema separado para cada uma delas?

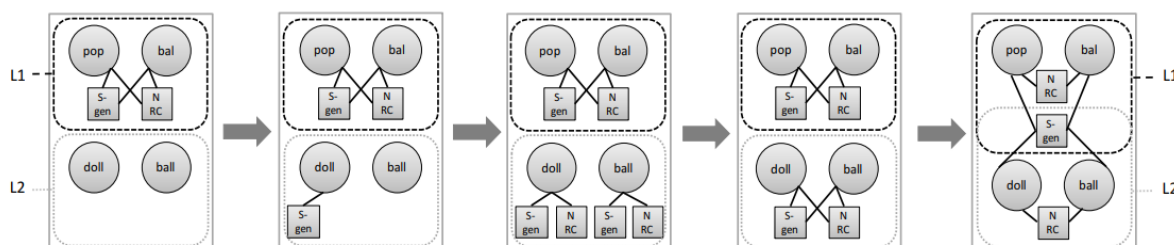
Pickering e Branigan (1998) apresentam um modelo que defende que bilíngues compartilham informações sintáticas entre as línguas, o que foi reforçado por Hartsuiker *et al.* (2004). As estruturas sintáticas seriam representadas por nós combinatórios, que são conectados a outros nós, que representam cada palavra. Essas conexões estariam armazenadas em uma única representação das línguas que o bilíngue utiliza, desde que as línguas sejam suficientemente semelhantes.

A suposição de que as línguas pudessem ter representações sintáticas compartilhadas evoluiu posteriormente para um modelo de sintaxe compartilhada em bilíngues tardios, que foi desenvolvido por Hartsuiker e Bernolet (2017). De acordo com estes autores, os bilíngues compartilham realmente informações sintáticas, mas o fazem em diferentes níveis. Ou seja, as representações linguísticas de bilíngues tardios são mais integradas ou menos integradas, de acordo com o nível de desenvolvimento da L2 desses bilíngues.

Baseados em pesquisas anteriores, Bernolet e Hartsuiker (2018) desenvolveram um modelo de desenvolvimento sintático da L2 para bilíngues tardios, que contém cinco estágios. Utilizando como L1 o holandês e como L2 o inglês, o modelo mostra as trajetórias de aprendizagem para o caso genitivo da L2 (s-gen: *a boneca do menino*) e para o sintagma nominal pós-modificado em L2 (NRC: *a bola que é vermelha*), estruturas que podem ser combinadas com os substantivos holandeses *pop* (boneca) e *bal* (bola). Vejamos a

representação deste modelo na figura 1:

Figura 1 - Modelo de desenvolvimento sintático para bilíngues tardios



Fonte: Bernolet e Hartsuiker (2018), pág. 208.

O primeiro painel da figura 1 representa a primeira etapa de aquisição de estruturas sintáticas e mostra que nesta fase as palavras são representadas mentalmente sem conexão devida com as estruturas sintáticas. A sintaxe da L1 já está totalmente representada no primeiro painel, que representa a fase inicial de aquisição de ambas as estruturas. Nesta etapa inicial, a aquisição da L2 começa com o aprendizado de representações flexíveis sem conexões firmes com informações sintáticas. Um bilíngue com holandês como L1 e inglês como L2 pode ter aprendido substantivos simples como *doll* (boneca) e *boy* (menino), mas pode não saber quais construções em inglês podem ser usadas para expressar que a boneca é propriedade do menino, por exemplo. Se esse aprendiz em fase inicial quiser expressar uma relação possessiva entre os substantivos, ele normalmente usa seu conhecimento sobre estruturas genitivas na L1. Ou seja, nesta fase, como as representações ainda não estão formadas, os bilíngues tendem a transferir construções sintáticas da L1 para a L2, o que faz com que a compreensão da L2 seja impulsionada pelas preferências sintáticas da L1.

Os painéis 2, 3 e 4 mostram o desenvolvimento das estruturas sintáticas na L2. O segundo painel mostra que, à medida em que tem mais contato com a L2, o bilíngue começa a aproximar as conexões entre palavras e representações sintáticas específicas. Esta seria a segunda fase de desenvolvimento da L2 para bilíngues tardios. O terceiro estágio do desenvolvimento sintático da L2 para bilíngues tardios é representado no painel 3. Para Bernolet e Hartsuiker (2018), é a partir desta fase que a informação sintática começa a se tornar mais abstrata e, consequentemente, as representações sintáticas da L2 se desenvolvem de forma separada das representações sintáticas da L1. Primeiramente, as representações sintáticas são conectadas ao léxico da L2, para que as estruturas sintáticas sejam adquiridas em L2. O quarto painel representa o estágio seguinte do desenvolvimento sintático de L2. Quando os bilíngues

atingem esta fase, as representações sintáticas estão mais integradas, especialmente em construções que são mais frequentes em L2.

A última etapa do modelo proposto por Bernolet e Hartsuiker (2018) é atingida por bilíngues altamente proficientes em L2. O último painel mostra que os bilíngues tardios alcançaram nível satisfatório de desenvolvimento sintático. Em outras palavras, os autores defendem que a abstração sintática na L2 está diretamente ligada ao nível de proficiência do bilíngue.

Este modelo proposto por Bernolet e Hartsuiker (2018) assume que algumas construções sintáticas são compartilhadas entre as línguas do bilíngue, mas isto não ocorre com todas as estruturas, diferentemente do que foi apresentado em Hartsuiker e Bernolet (2017). Em casos de estruturas sintáticas que não existem na L1 do bilíngue ou em que não há correspondência na L1, as representações sintáticas são realizadas de forma específica na L2. Por outro lado, as estruturas sintáticas que existem nos dois idiomas do bilíngue seriam integradas. Isto quer dizer que, no processamento da L1 ou da L2, estas construções seriam ativadas nas duas línguas simultaneamente.

Bernolet e Hartsuiker (2018) esclarecem que para a produção de uma nova estrutura sintática em L2, os bilíngues podem utilizar duas estratégias diferentes: transferindo a sintaxe da L1 ou imitando a sintaxe de bilíngues mais competentes. Consequentemente, essas sentenças em L2 que são repetições que são exatamente iguais ou minimamente editadas de frases anteriores, só são produzidas imediatamente após um bilíngue mais proficiente utilizar uma expressão que servirá como exemplo.

Faz-se necessário dar destaque ao fato de que há um número considerável de estudos que investigam a possível influência de aspectos lexicais e fonológicos da L1 para a L2. Hansen Edwards e Zampini (2008), Pavlenko (2009) e Elgort *et al.* (2015) são alguns exemplos de pesquisas que exploram a influência da L1 na L2, considerando interações lexicais e fonológicas.

2.3 Sintagmas nominais preposicionados com dois substantivos

Kenedy (2013) explica que a sintaxe é um sistema computacional da linguagem humana, na qual a palavra é uma das menores engrenagens, que formam sintagmas. Os sintagmas, por sua vez, formam frases. Ou seja, o sintagma é um conjunto de unidades que é manipulado pelo sistema computacional como se fosse uma peça única. Esse conjunto pode ser

constituído exclusivamente por unidades lexicais e, também, pode ser constituído por outras unidades sintagmáticas.

Desta forma, quando precisamos comunicar algo e um único item lexical não é suficiente para suprir esta referência, é necessário acrescentar mais termos para tornar mais específica esta palavra mais geral, como, por exemplo, barra de chocolate. Não se trata de qualquer barra, mas uma específica, aquela que é feita de chocolate. Assim, sintagmas nominais, que são unidades linguísticas formadas por um núcleo e seus modificadores, são utilizados. Downing (2015) afirma que esses núcleos normalmente são substantivos e os modificadores podem aparecer antes ou depois deles. Quando há um pré-modificador, ele se caracteriza como determinante e quando há um pós-modificador, ele acrescenta alguma informação sobre o núcleo, apresentando sua identificação ou sua definição.

Em definição similar à de Downing (2015), Huddleston *et al.* (2022) conceituam os sintagmas nominais como estruturas sintáticas formadas por uma parte principal, que é o núcleo, e seus dependentes. Estes últimos elementos podem ser externos ou internos à parte principal. Para os autores, os elementos pré-modificadores são considerados internos, enquanto os pós-modificadores são externos. Downing (2015) explica que os elementos pós-modificadores do núcleo podem ser representados de várias formas, como por meio de sintagmas preposicionados, sintagmas adjetivais, apostos e pronomes reflexivos.

Neste estudo, investigamos os sintagmas nominais cujos núcleos são formados por um substantivo e por um elemento dependente pós-modificador, que também é um substantivo. Em outras palavras, pesquisamos os sintagmas nominais preposicionados com dois substantivos, como no exemplo dado no início deste capítulo: barra de chocolate.

O substantivo dependente modificador dos sintagmas nominais preposicionados com dois substantivos pode se apresentar de diferentes formas na língua portuguesa e na língua inglesa. Na L1 deste estudo, esse elemento aparece sempre como pós-modificador. Os sintagmas nominais preposicionados com dois substantivos em PB são codificados exclusivamente pela preposição *de* e suas variações *da* (*de + a*) ou *do* (*de + o*). Para exemplificar, vejamos as sentenças apresentadas por Miranda (2021):

$$(1) \quad \text{Subst1} + \text{de} + \text{Subst2}$$

Exemplo: barra de chocolate

(2) $S_{\text{ubst}1} + \text{da} + S_{\text{ubst}2}$

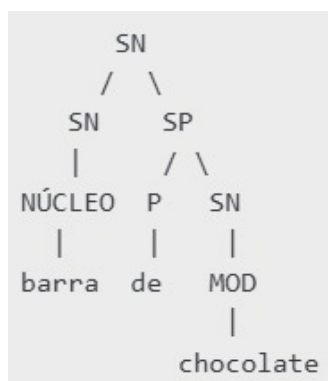
Exemplo: janela da casa

(3) $S_{\text{ubst}1} + \text{do} + S_{\text{ubst}2}$

Exemplo: cadeira do escritório

Não há outra possibilidade de representação sintática da construção desse sintagma nominal, sendo esta a única alternativa possível em língua portuguesa, conforme explicam Resende *et al.* (2020). Dessa forma, os sintagmas nominais preposicionados com dois substantivos podem ser representados arboreamente da seguinte forma, exclusivamente:

Figura 2 - Representação do sintagma nominal preposicionado em L1



Fonte: autoria própria

Na representação arbórea da figura 2, SN é o sintagma nominal, SP é o sintagma preposicionado e P é a preposição (*de*, *da* ou *do*). O sintagma nominal *barra* é o núcleo e o sintagma nominal *chocolate* é o modificador. A representação sintática dos sintagmas nominais preposicionados com dois substantivos em L1 não varia em organização estrutural.

Por outro lado, na L2 deste estudo, o substantivo dependente do núcleo do sintagma nominal pode aparecer como elemento pós-modificador ou como pré-modificador, dependendo da construção sintática utilizada. As possíveis diferenças sintáticas entre as línguas costumam ser desafiadoras para bilíngues, especialmente em estágios iniciais de aprendizagem. Os sintagmas nominais preposicionados com dois substantivos podem ser ativados em inglês por meio de três construções sintáticas distintas, distribuídas de forma parcial complementar. Vejamos as três possibilidades para representar *barra de chocolate* em PB, para exemplificar:

$$(4) \quad S_{\text{ubst}1} + \text{of} + S_{\text{ubst}2}$$

Exemplo: *bar of chocolate*

$$(5) \quad S_{\text{ubst}2} + S_{\text{ubst}1}$$

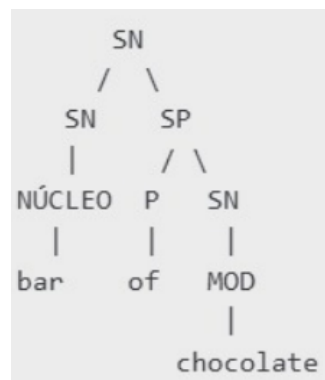
Exemplo: *chocolate bar*

$$(6) \quad S_{\text{ubst}2} 's + S_{\text{ubst}1}$$

Exemplo: *chocolate's bar*⁶

A representação arbórea das três construções sintáticas possíveis para a representação dos sintagmas nominais preposicionados com dois substantivos em língua inglesa se apresentam de três formas distintas. Vejamos a representação arbórea da primeira possibilidade sintática:

Figura 3 - Representação do sintagma nominal em L2 (*of*)



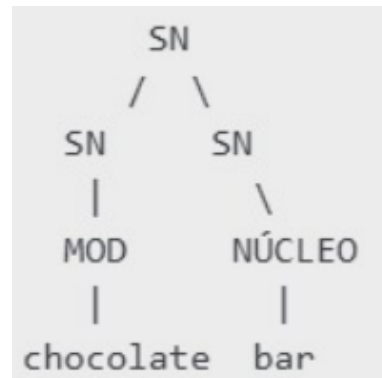
Fonte: autoria própria

Na representação arbórea da figura 3, SN é o sintagma nominal, SP é o sintagma preposicionado e P é a preposição (*of*). O sintagma nominal *bar* é o núcleo e o sintagma nominal *chocolate* é o modificador. Esta é a opção sintática que mais se assemelha à L1, pois a estrutura da sintaxe se mantém, incluindo a ordem dos itens lexicais, com o elemento modificador (*chocolate*) após o núcleo (*bar*). Murphy (2015) afirma que esta construção pode ser usada em L2 para representar os sintagmas nominais preposicionados com dois substantivos em todas as circunstâncias, apesar de não ser a estrutura mais utilizada nas situações reais.

⁶ Esta construção sintática não é considerada gramatical, mas está prevista na interlíngua de aprendizes bilíngues.

Vejamos a representação arbórea da segunda possibilidade sintática em L2:

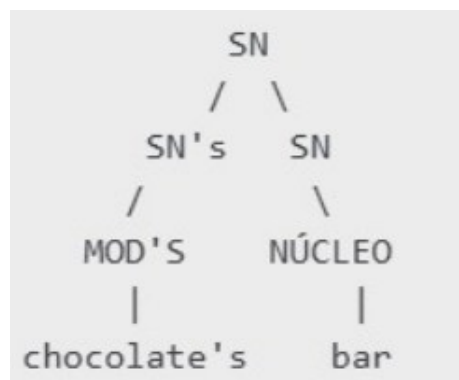
Figura 4 - Representação do sintagma nominal em L2 (inversão de substantivos)



Fonte: autoria própria

Na representação arbórea da figura 4, SN é o sintagma nominal. O sintagma nominal *bar* é o núcleo e o sintagma nominal *chocolate* é o modificador. A construção sintática (5), *chocolate bar*, pode ser representada arboreamente desta forma, pois os substantivos que formam o sintagma nominal foram invertidos, tão somente. Esta estrutura se diferencia da L1, porque o elemento modificador (*chocolate*) surge antes do núcleo do sintagma nominal (*bar*). Esta construção pode ser usada em L2 para representar os sintagmas nominais preposicionados com dois substantivos em todas as circunstâncias em que o substantivo modificador for inanimado. Além disso, parece ser a estrutura mais utilizada nas situações reais, sendo associada à fluência e à naturalidade, conforme mostram as experiências em salas de aula de L2. Vejamos a representação arbórea da terceira possibilidade sintática em L2:

Figura 5 - Representação do sintagma nominal em L2 (caso genitivo)



Fonte: autoria própria

Na representação arbórea da figura 5, SN é o sintagma nominal. O sintagma nominal *bar* é o núcleo e o sintagma nominal *chocolate* é o modificador. A partícula 's é

utilizada entre os substantivos para expressar o caso genitivo. A construção sintática (6), *chocolate's bar*, pode ser representada arboreamente desta forma, em que foi utilizado 's no segundo substantivo, que fica seguido do primeiro substantivo, nesta ordem. Ou seja, o elemento modificador do sintagma nominal se apresenta antes do núcleo, o que também causa diferença sintática, comparando com a L1. Murphy (2015) esclarece que esta construção sintática deve ser utilizada para indicar relações de posse, familiares, de propósito ou de origem, quando o substantivo modificador é animado, como em *John's pen* (A caneta de João) e em *women's rights* (direitos das mulheres).

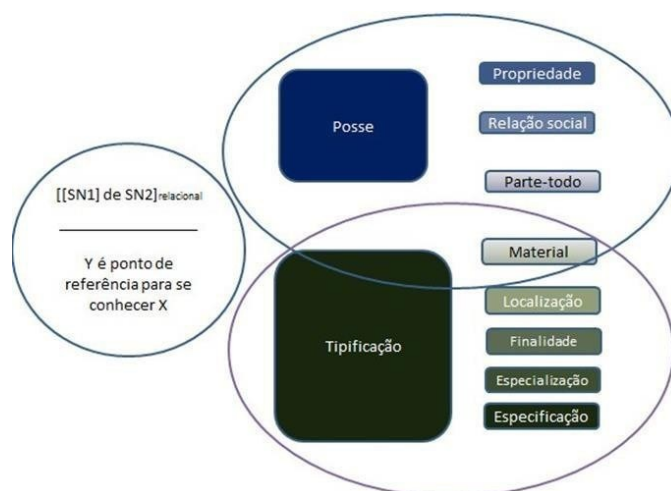
Consideramos conveniente abordar esta terceira possibilidade, mesmo sabendo que esta construção não resulta em gramaticalidade ou aceitabilidade em todas as situações, porque essa terceira construção costuma ser bastante utilizada por bilíngues tardios, em virtude da supergeneralização da regra do genitivo, fenômeno observado em salas de aula de inglês como L2. Romano (2016) explorou o caso genitivo em falantes nativos e não-nativos, constatando maior grau de dificuldade em falantes não-nativos.

As três representações arbóreas em L2 apresentadas anteriormente indicam que a organização sintática pode variar, inclusive no número de itens lexicais necessários para a representação dos sintagmas nominais preposicionados com dois substantivos. Enquanto as representações arbóreas das figuras 3 (uso da preposição *of*) e 5 (caso genitivo) utilizam três itens lexicais, que precisam ser organizados para a representação sintática dos sintagmas nominais, a representação arbórea da figura 4 (inversão dos substantivos) utiliza apenas dois itens lexicais para esta representação. Esta observação pode nos fazer refletir sobre a interferência do número de itens lexicais no custo de processamento linguístico.

Apesar desta pesquisa estar voltada à investigação das estruturas sintáticas que podem ser utilizadas para representar os sintagmas nominais preposicionados com dois substantivos, consideramos pertinente abordar os variados sentidos que essas construções podem estabelecer, pois estas considerações também podem contribuir para a compreensão do processamento mental bilíngue.

Para Santos e Alonso (2022), os elementos integrantes dos sintagmas nominais ativam relações semânticas e são tratados como X e Y, em que Y representa normalmente o ponto de referência para se conhecer X, entende-se que o falante elabora uma conexão mental entre X e Y. Os vários sentidos que se associam a esta construção podem ser abarcados por dois domínios diferentes e apresentam valores mais ou menos prototípicos em relação a cada um dos domínios, conforme mostra o esquema proposto pelas autoras:

Figura 6 - Relações semânticas - SN1 de SN2



Fonte: Santos e Alonso (2022)

Santos e Alonso (2022) identificam os sintagmas nominais preposicionados com dois substantivos como SN1 de SN2. A relação entre os dois referentes expressos por SN1 e SN2 na construção SN1 de SN2 encontra fundamento na própria natureza associativa do pensamento humano e pode ser atestada em diferentes formações do PB, com diferentes sentidos. Por exemplo, em *livro de João*, a relação entre os substantivos se caracteriza pela posse, pois o livro é de João. Em *cadeira de praia*, a relação é de lugar-fim, porque a finalidade da cadeira é sua utilização na praia. No exemplo *chapéu de palha*, o sentido construído entre os substantivos é material, pois palha se refere ao material de que o chapéu é feito. No sintagma nominal *banqueta de bar*, a relação é de especificação, porque não se trata de qualquer banqueta, é uma banqueta específica, aquela que é de bar.

As relações semânticas foram investigadas também em Miranda *et al.* (2023). Neste trabalho, as três construções sintáticas em L2 para representar o sintagma nominal preposicionado com dois substantivos foram pesquisadas e os resultados apontam que não houve interferência do tipo de domínio conceitual no *continuum* posse-tipificação. Isso significa que o aspecto semântico não se sobrepôs ao aspecto sintático.

Além da possibilidade das diferentes relações semânticas entre os elementos das construções SN1 de SN2, ao mesmo tempo, há assimetria entre seus referentes. SN2 (modificador) é tomado como referência para a interpretação de SN1 (núcleo ou alvo), conforme Langacker (2003, 2009) e Calabrese (2011).

A escolha por esta construção se justifica justamente pela possível diferença entre a sintaxe da L1 e da L2 dos bilíngues que participarão deste estudo. Também, pela minha experiência no magistério, é possível observar que a aprendizagem de sintagmas nominais

preposicionados com dois substantivos se dá de forma implícita, uma vez que os materiais didáticos não costumam apresentar um tópico voltado exclusivamente à explicação de como esses sintagmas podem ser representados em inglês. Isso ocorre porque os materiais didáticos são elaborados para atender aprendizes de diversas partes do mundo, com diversos idiomas como L1 e essa diferença interlinguística nem sempre ocorre ou não ocorre da mesma maneira do que em português. Assim, trata-se de uma aquisição normalmente baseada no uso, como resultado da experiência do falante com a L2 ao longo de sua vida, como afirma Tomasello (2000). À medida que acontecem as experiências linguísticas, o processo de transformação das representações linguísticas vai ocorrendo de forma implícita. Em outras palavras, os sintagmas nominais são aprendidos de maneira incidental, sem consciência explícita do que está sendo aprendido.

Guimarães (2020) afirma que a cada vez que o aprendiz produz ou compreende a língua, seu sistema linguístico armazena informações sobre os itens linguísticos em questão, fazendo ajustes das representações já existentes às novas informações que são encontradas no uso, o que perdura por toda a vida do falante. O que o cérebro faz é refletir sobre o que é necessário para se ter sucesso na compreensão e na produção da fala. Para isso, ele compara o que foi compreendido ou produzido com o que foi planejado mentalmente, mede a discrepância entre a intenção e o que foi realizado e ajusta o sistema, por fim.

Chang *et al.* (2006) argumentam que a base da aquisição de classes de palavras e estruturas sintáticas é o ajuste feito pelos falantes toda vez que seu planejamento linguístico não é bem-sucedido. Dessa maneira, o sistema linguístico atualiza a distribuição do item em questão, adicionando informações sobre as circunstâncias em que ele ocorreu, para que haja precisão na próxima tentativa. É dessa forma que a proficiência em L2 vai se desenvolvendo, por meio das representações da L2 na mente bilíngue, que evoluem da especificidade de itens e da língua, para categorias, a partir das quais é possível fazer generalizações a respeito da gramática da L2, como Bernolet *et al.* (2013) apresentam. É por isso que, logo nos primeiros contatos com inglês como L2, o bilíngue em fase inicial de aprendizagem percebe que se usa mais comumente *phone number*, e não *number of phone*, por exemplo, para representar número de telefone em inglês. Isso ocorre porque à medida que o aprendiz processa mais construções em L2 e aumenta seu repertório lexical, ele é capaz de identificar mais de uma possibilidade para representar o sintagma nominal preposicionado com dois substantivos ou de fazer restrições de uso de determinada estrutura sintática. Quanto mais processamentos em L2, mais dados o aprendiz tem para realizar ajustes em seu sistema linguístico, o que explica que a proficiência depende diretamente do contato com a L2.

Durante a revisão literária, detectamos que a estrutura sintática que investigamos pode surgir em pesquisas científicas com diferentes nomenclaturas. Em Resende et al. (2020) e em minha dissertação (Miranda, 2021), este tipo de construção sintática é citada como sintagma nominal com relação de posse entre os substantivos. Já Santos e Alonso (2022) definem a referida estrutura como construção relacional binominal SN1 de SN2, porque as autoras esclarecem que este tipo de construção sintática abarca um campo maior de relações semânticas, que não se limitam a posse. Em alguns estudos, percebemos que esta mesma estrutura sintática recebe outra nomenclatura: palavras compostas, como na pesquisa de De Cat *et al.* (2015). Em outros estudos, a definição para palavras compostas pode se referir também a expressões formadas por uma única palavra, que deriva de outras duas, como em guarda-costas (*bodyguard*) ou palavra-chave (*keyword*), como em Fiorentino e Poeppel (2007) e em Li *et al.* (2017).

Nesta tese, decidimos utilizar a nomenclatura sintagma nominal preposicionado com dois substantivos. Decidimos pelo termo sintagma nominal porque esta definição parece ser mais comum do que palavras compostas, mas optamos por especificar ainda mais de que tipo de sintagma nominal se trata: aquele que é preposicionado em L1 e é constituído por dois substantivos, como em *porta da casa*, *barra de chocolate* e *cadeira do escritório*.

Vejamos com mais detalhes os dois estudos que são o ponto de partida desta pesquisa: Resende *et al.* (2020) e Miranda (2021). Na pesquisa de Resende *et al.* (2020) foram testados 20 participantes brasileiros, com inglês como L2, com nível de proficiência intermediário ou avançado. A idade média dos participantes foi de 33,7 anos, com desvio padrão de 5,6. O objetivo desta investigação foi verificar a influência do sistema de tradução automática do Google Tradutor.

Primeiramente, os participantes realizaram traduções de sintagmas nominais, como em os cabos dos talheres são coloridos. Foram utilizadas 104 imagens e palavras abaixo das imagens, para dar suporte de vocabulário e evitar que o não conhecimento das palavras interferisse nos resultados, já que o foco desse estudo era a estrutura sintática. Esta primeira tarefa serviu como pré-teste e as respostas dos participantes foram gravadas.

Em seguida, os participantes foram solicitados a realizar outra tarefa de tradução. Nessa segunda tarefa, foi utilizado o paradigma de priming, a fim de ser verificado se a construção sintática da primeira atividade poderia sofrer alterações em sua seleção. Foram 26 traduções realizadas, que incluíam 2 pares de traduções prime (realizadas a partir do Google Tradutor) e 1 tradução alvo.

Os resultados obtidos indicaram que os participantes realizaram mais traduções

com acurácia seguindo o paradigma de priming exposto no Google Tradutor na segunda tarefa do que na primeira. As escolhas pela construção preposicionada com uso de *of* decaíram de 38,42% na primeira atividade para 10,5% na segunda atividade. Já o uso dos substantivos invertidos que foi utilizado nas traduções prime sofreu acréscimo, saindo de 33,42% na primeira atividade para 55% na segunda atividade.

Miranda (2021) pesquisou 20 bilíngues brasileiros que fizeram a tradução de sintagmas nominais preposicionados com dois substantivos. Primeiramente, foi analisado que estrutura sintática os bilíngues utilizariam dentre as três possibilidades em L2 em um pré-teste com 10 itens. Os resultados mostraram que a construção preposicionada foi utilizada em 35 respostas (17,5%), o caso genitivo foi usado em 78 respostas (39%) e a inversão de substantivos foi utilizada em 87 respostas (43,5%). Ou seja, houve certa distribuição das respostas entre as construções com genitivo e com substantivos invertidos. Para responder o pré-teste, os participantes levaram em média 3 minutos e 18 segundos, com desvio padrão de 1 minuto e 19 segundos. A mediana foi de 2 minutos e 45 segundos. O participante que realizou todo o pré-teste em menor tempo precisou de 1 minuto e 36 segundos e aquele que precisou levou mais tempo gastou 6 minutos e 22 segundos.

Em seguida, os participantes foram submetidos a uma tarefa de tradução da L1 para a L2 com 15 itens, que utilizou o paradigma de *priming*. Nesse caso, as sentenças *prime* utilizaram os substantivos invertidos, como em *house door* para representar porta da casa. Os participantes seguiram a mesma estrutura exibida na tradução *prime* em 13,47 (89,8%) das 15 traduções realizadas, em média, com desvio padrão de 2,09 (13,93%). O mínimo de itens com *prime* foi 9 (60%). Ainda, 6 (30%) participantes seguiram o priming sintático em todas as suas respostas. A média de tempo utilizado para traduzir os 15 itens com acurácia para a L2 foi de 5 minutos e 54 segundos, com desvio padrão de 2 minutos e 34 segundos. A mediana ficou em 5 minutos e 30 segundos. O participante que realizou as 15 traduções com acurácia em menor tempo utilizou 3 minutos e 10 segundos e o que levou o maior tempo precisou de 9 minutos e 21 segundos.

No pós-teste, os participantes precisaram selecionar mais uma vez qual estrutura sintática eles julgavam ser a mais adequada em L2 para representar os sintagmas nominais e foram utilizadas as mesmas sentenças do pré-teste, com a ordem dos itens aleatorizada. Os resultados mostraram que a construção preposicionada foi utilizada em 11 (5,5%) respostas, o caso genitivo foi usado em 32 (16%) respostas e a inversão de substantivos foi utilizada em 157 (78,5%) respostas. Nesta fase da pesquisa, a escolha pelo sintagma nominal não preposicionado que foi utilizado nas sentenças *prime* da tarefa de tradução alcançou quase 80% das respostas,

enquanto as outras duas estruturas possíveis foram selecionadas em um pouco mais de 20% das seleções de respostas dos participantes. Para selecionar as construções que julgavam adequadas, os participantes levaram em média 1 minuto e 23 segundos, com desvio padrão de 50 segundos. A mediana foi de 1 minuto e 17 segundos. O participante que realizou todo o pós-teste em menor tempo precisou de 54 segundos e aquele que precisou levou mais tempo gastou 2 minutos e 35 segundos.

Os resultados deste estudo indicaram que as sentenças *prime* mostraram efeito, porque os participantes foram capazes de reutilizar a mesma construção em suas traduções. Esse efeito permaneceu no pós-teste, fase em que as sentenças *prime* não eram mais exibidas simultaneamente. Além disso, o custo do processamento linguístico foi menor no pós-teste, porque os participantes foram capazes de selecionar mais vezes a construção *prime* em menor tempo.

Em outras palavras, os resultados obtidos em Miranda (2021) confirmaram que à medida que o aprendiz foi aumentando seu repertório lexical e processando mais construções em L2, ele se tornou capaz de identificar mais de uma possibilidade para representar o sintagma nominal preposicionado com dois substantivos e foi fazendo restrições de uso de determinada estrutura sintática. Ao final, os bilíngues foram capazes de realizar os ajustes em seu sistema linguístico, a fim de utilizar a construção não preposicionada que apareceu como sentença *prime*. Os resultados dessa pesquisa foram publicados posteriormente em Miranda e Toassi (2023).

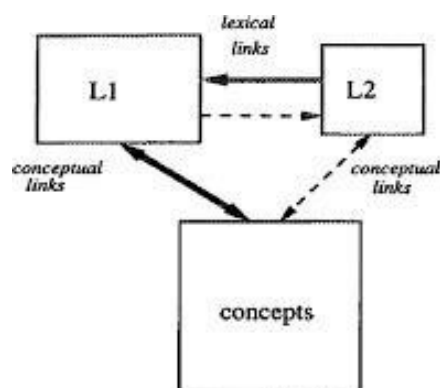
Neste estudo atual, ampliamos a investigação dos sintagmas nominais preposicionados com dois substantivos, incluindo o efeito cognato e o nível de proficiência em L2. A inclusão da investigação acerca de possível efeito cognato se justifica porque o processamento linguístico normalmente é facilitado por esses itens lexicais semelhantes em ortografia e significado, o que pode causar menor custo linguístico. Além disso, o efeito cognato, comumente chamado de efeito de facilitação cognata, pressupõe coativação linguística e dá suporte à hipótese de não seletividade das línguas em mentes bilíngues, o que é defendido por muitos estudos, tais quais Hartsuiker *et al.* (2004), Hartsuiker e Pickering (2008), Hartsuiker *et al.* (2016) e Hwang *et al.* (2018), para citar estudos referenciais. Já a inclusão da análise do nível de proficiência em L2 se justifica porque, como já citado anteriormente e de acordo com minha prática de sala de aula de inglês como L2, é possível observar a variação de processamento e de uso de diferentes estruturas sintáticas para representar os sintagmas nominais preposicionados com dois substantivos conforme progride o contato com a L2.

2.4 Efeito cognato

O processo de aprendizagem de novos itens lexicais é um importante aspecto a ser considerado na aquisição de L2. Por isso, muitos pesquisadores colocam o léxico como centro da compreensão e da produção da linguagem humana, como Hunt e Beglar (2005) afirmam. Para Toassi e Pereira (2019), aprender novas palavras em L2 é um processo sucessivo e cumulativo que vincula novas formas lexicais a representações de conceitos, que já possuem conexão com formas de palavras da L1. Para se comunicar em uma nova língua que está em processo de aquisição, é necessário conhecer um número razoável de palavras. Por isso, a aprendizagem de vocabulário costuma ser uma parte muito importante dentro do processo de aprendizagem de L2.

Para Souza (2020), o modelo mais influente e debatido na Psicolinguística do Bilinguismo para representar a organização do léxico bilíngue é provavelmente o Modelo Hierárquico Revisado (*The Revised Hierarchical Model*), proposto por Kroll e Stewart (1994). Segundo este modelo, nos estágios iniciais de aquisição da L2, a representação conceitual ocorre por meio do léxico da L1. Isto significa que o léxico da L2 dependeria do léxico da L1. Além disso, o modelo hierárquico revisado prevê ligações mais fortes entre o léxico da L1 e o léxico da L2 do que a conexão entre a L2 e a L1, conforme mostra a figura:

Figura 7 - Modelo Hierárquico Revisado



Fonte: Kroll e Stewart (1994)

O modelo proposto por Kroll e Stewart (1994) indica que a tradução da L2 para a L1 seria mais fácil de ser realizada do que da L1 para a L2. No entanto, o modelo prevê que as conexões lexicais da L1 para a L2 se fortalecem à medida que o aprendiz avança em sua proficiência em L2, o que causa menor dependência das associações entre forma e significado de itens lexicais da L2, como explica Souza (2020).

Dijkstra *et al.* (2018) enfatizam que o Modelo Hierárquico Revisado se concentra na atividade de tradução, que ocorre mesmo que inconscientemente repetidamente no cotidiano de bilíngues. No entanto, esse modelo não contempla totalmente palavras semelhantes e homógrafos interlinguísticos.

A variação tipológica entre duas línguas pode modular o processo de aprendizagem de L2, já que o grau de similaridade interlinguística parece ser um dos principais aspectos considerados de forma intuitiva pelos aprendizes na percepção de facilidade ou dificuldade, sobretudo com relação à dimensão lexical, conforme debatem Antón e Duñabeitia (2020). Contudo, considerando que a variação tipológica é multifatorial e multidimensional, fenômenos que podem ocorrer em todos os níveis linguísticos oferecem a possibilidade de observar de que maneira os aprendizes lidam com as diferenças e com as semelhanças tipológicas na construção de sua aprendizagem.

Há algumas plataformas que calculam o grau de similaridade entre idiomas. Uma delas é o sítio <http://www.elinguistics.net/>. A proximidade genética entre idiomas calculada por esta plataforma traz a seguinte classificação: entre 1 e 30, as línguas são altamente relacionadas; entre 30 e 50, os idiomas são relacionados; entre 50 e 70, as línguas são relacionadas remotamente; entre 70 e 78, os idiomas são muito remotamente relacionados; e entre 78 e 100 pontos, não há relação reconhecida entre as línguas.

Levando em consideração os dois idiomas desta pesquisa, de acordo com o sítio citado no parágrafo anterior, a relação de proximidade entre o português e o inglês é de 56,5 pontos, dentro de uma escala de 100 pontos. Considerando o grau de proximidade proposto por este sítio, a L1 e a L2 deste estudo são consideradas línguas relacionadas remotamente.

Focando apenas na L2 desta pesquisa, Green (2020) afirma que cerca de 60% do vocabulário da língua inglesa tem origem latina ou grega. Também, muitos desses itens lexicais apresentam equivalência ortográfica com outros vocábulos que fazem parte das línguas românicas, como o espanhol e o português.

Estes dados podem indicar que algumas palavras em inglês com mesma raiz ou mesma origem que algumas palavras em português podem ser mais facilmente reconhecidas interlinguisticamente. Entretanto, é necessário ressaltar que muitas dessas palavras evoluíram ao longo do tempo e apresentam algumas mudanças.

A busca por semelhanças entre idiomas é uma estratégia comum para facilitar a aprendizagem de novos vocabulários. Hall *et al.* (2009) ponderam que para a aquisição de um novo item lexical deve haver uma ativação de uma entrada dessa nova palavra na memória lexical, fazendo uma conexão da representação ortográfica com seu significado. Para crianças,

este processo normalmente se dá de forma automática e inconsciente. Entretanto, à medida em que crescemos, esta ativação vai deixando de ser automática, gerando maiores desafios para aprendizes tardios de L2.

Com relação a aspectos cognitivos, vale ressaltar que Costa *et al.* (2006) explicam que a similaridade linguística pode potencializar a competitividade entre as línguas, provocando uma sobrecarga no sistema atencional dos bilíngues, o que pode afetar a resolução de conflitos. Em outras palavras, a similaridade interlinguística pode potencializar a ocorrência de interferência da L1 na L2 em mentes bilíngues, uma vez que o acesso lexical pode apresentar conflitos no processo de decisão das representações lexicais, que são ativadas em seus idiomas. Ainda segundo Costa *et al.* (2006), quanto maior a quantidade de léxico compartilhado entre os dois idiomas de um bilíngue, maior será a quantidade de atenção a ser dispensada para conseguir selecionar adequadamente a palavra no idioma-alvo e evitar interferências da língua não-alvo. A hipótese da sobrecarga atencional ligada à similaridade linguística, que foi proposta por Costa *et al.* (2006), foi evidenciada no estudo de Ortiz-Preuss (2011), para exemplificar.

Na pesquisa realizada por Ortiz-Preuss (2011), foram pesquisados 23 bilíngues com português como L1 e espanhol como L2 e bilíngues com espanhol como L1 e português como L2, com o objetivo de analisar o acesso lexical e a produção de fala, analisando efeitos de interferência interlinguística e o efeito cognato. Os participantes realizaram duas tarefas de nomeação de figuras. Nos itens com cognatos, os resultados mostraram que o tempo de resposta foi mais rápido, em comparação com itens sem cognatos, o que evidencia menor custo de processamento. No entanto, a acurácia para estes mesmos itens foi menor, o que evidencia que nem sempre a similaridade linguística está ligada à facilidade de aquisição da L2. Ortiz-Preuss (2011) acredita que a diminuição da acurácia pode ser em decorrência da sobrecarga atencional, provocada pela similaridade lexical.

Os resultados apresentaram o indício de que nem sempre a similaridade interlinguística pode estar associada à facilidade de aquisição de L2, mas pondera que a acurácia mais baixa em itens com cognatos pode ter sido ocasionada pela sobrecarga atencional.

Um dos principais fenômenos que envolvem a semelhança entre as línguas é o efeito cognato, que é comumente chamado de efeito de facilitação cognata nas pesquisas linguísticas, devido à sua constante comprovação de viabilização de menor custo interlinguístico. Ademais, Roembke *et al.* (2024) confirmam que os cognatos são uma ferramenta essencial para investigar como mais de uma língua pode interagir durante o processamento da linguagem.

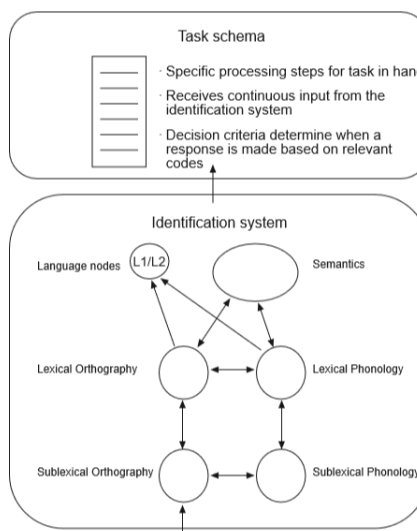
Durante a aprendizagem de vocabulário de uma língua, a semelhança mais marcante

que se pode observar entre as palavras de dois ou mais idiomas são os itens lexicais chamados de cognatos. Para Dijkstra e Heuven (2002), os cognatos são palavras de diferentes idiomas que se sobrepõem consideravelmente em forma e significado. Por exemplo, a palavra piano é representada da mesma forma em inglês e em português, sendo considerada cognata perfeita, pois, além de possuir a mesma grafia nos dois idiomas, possui também o mesmo significado. Portanto, podemos deduzir que um brasileiro falante de PB, que está aprendendo inglês como L2, compreenderia e internalizaria a palavra piano sem muito esforço cognitivo ao ter contato com esse item lexical pela primeira vez. Dessa maneira, os cognatos podem funcionar como uma ponte entre dois idiomas, facilitando a conscientização de semelhanças linguísticas, o que pode gerar benefícios como a identificação de ligações entre o vocabulário e os aspectos sintáticos. Para Crystal (2008), os cognatos são palavras que derivam da mesma fonte historicamente. Este contexto etimológico compartilhado faz com que os cognatos sejam semelhantes na forma como são pronunciados e escritos em diferentes línguas.

Com relação aos tipos de cognatos, Dijkstra *et al.* (2010) considera três classificações diferentes: cognatos idênticos, muito semelhantes e não idênticos. Os cognatos idênticos são aqueles que se sobrepõem ortograficamente e semanticamente, como *chocolate*, em português, que em inglês é *chocolate*. Os cognatos muito semelhantes são diferentes entre si apenas em uma ou duas letras, como em *museu* em português, que em inglês é *museum*. Já os cognatos não idênticos possuem mais que duas letras diversas entre si, como em *identificação* em L1 e *identification* em L2.

Para explicar os efeitos de facilitação cognata dentro da estrutura e da arquitetura do léxico mental bilíngue e para contemplar a lacuna deixada por Kroll e Stewart (1994) e seu Modelo Hierárquico Revisado, surgiu o modelo *Bilingual Interactive Activation* (BIA, daqui por diante) de Dijkstra e Heuven (1998) e seus aprimoramentos, BIA+, de Dijkstra e Heuven (2002), BIA –d, de Grainger *et al.* (2010), e BIA+S, de Casaponsa *et al.* (2020). Esses modelos explicam que quando uma palavra é ativada mentalmente em uma língua, sua tradução correspondente no outro idioma também é ativada. Vejamos a figura 8, que apresenta o esquema proposto pelo modelo BIA+ (Dijkstra e Heuven, 2002):

Figura 8 - *Bilingual Interactive Activation Plus*



Fonte: Dijkstra e Heuven (2002)

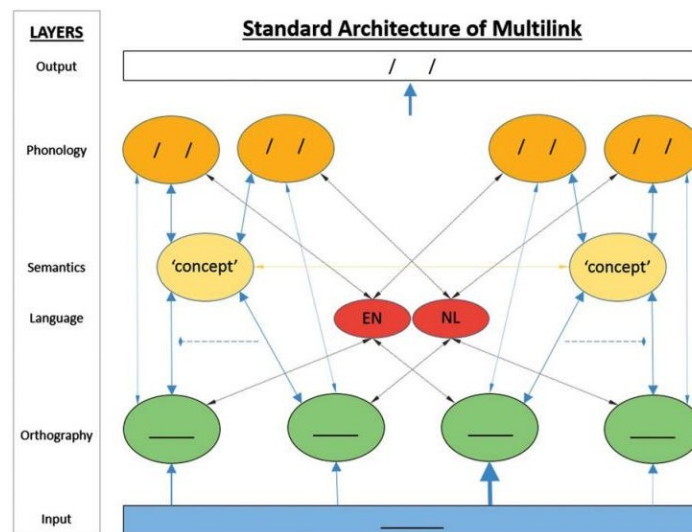
Através do modelo BIA+, Dijkstra e Heuven (2002) explicam que quando um bilíngue tem contato com um cognato, representações ortográficas semelhantes são ativadas nas duas línguas, alcançando a representação semântica comum, o que torna o reconhecimento da palavra cognata mais rápido, em comparação com palavras não cognatas. Casaponsa e Duñabeitia (2016) defendem que essa ativação simultânea ocorre de modo especial quando as palavras possuem a mesma estrutura ortográfica.

De acordo com Hall (2002), o efeito de facilitação cognata seria explicado porque características semelhantes na forma das palavras, tanto na L1 quanto na L2, são detectadas de forma automática e são retomadas para o estabelecimento de traços de memória para a aprendizagem de novas palavras na L2. Para Dijkstra *et al.* (1999), esse compartilhamento entre idiomas diferentes na forma das palavras pode ocorrer também com homógrafos interlinguísticos, que são palavras com a mesma grafia nas duas línguas, mas com significados diferentes. Com o avanço na aprendizagem de L2, o bilíngue vai tendo acesso a mais palavras do novo léxico e se tornando especialista. Casaponsa e Duñabeitia (2016) explicam que quando uma nova palavra é aprendida e armazenada no léxico mental bilíngue, esse novo item se fortalece e desenvolve conexões nos níveis semântico, lexical e sublexical, dentro da língua e entre línguas. E, curiosamente, as representações já existentes podem determinar alguma facilidade na aprendizagem e na incorporação do novo item lexical.

Outro modelo que busca a compreensão das atividades cognitivas desenvolvidas pela mente bilíngue e contempla os cognatos é o Multilink, de Dijkstra *et al.* (2018). É um modelo de acesso lexical não seletivo e busca cobrir processos de reconhecimento, a

recuperação e a produção de itens lexicais, mas incorpora elementos do Modelo Hierárquico Revisado e do BIA. O modelo Multilink permite a investigação de processamento de palavras em mentes bilíngues à medida em que ocorre, sendo ajustado aos níveis de proficiência em L2, conforme a necessidade de cada estudo. Vejamos o modelo computacional de Dijkstra *et al.* (2018):

Figura 9 - *Multilink*



Fonte: Dijkstra *et al.* (2018)

Os dois idiomas investigados estão no centro do modelo, inglês (EN) e holandês (NL). A entrada de informação (*input*) encontra-se na base do modelo, indicando o início do processamento linguístico. A ortografia está logo acima da entrada de informação, seguida dos idiomas do bilíngue. Logo após, a semântica é ativada e somada à metade da entrada de informação, para que a fonologia seja ativada e haja a saída de informação (*output*). A tarefa solicitada vai indicar como a informação vai sair fonologicamente, podendo ser no mesmo idioma de entrada ou em um idioma diferente, como no caso de tarefas de tradução. Ressalte-se que, apesar de haver um caminho a ser percorrido entre a entrada e a saída de informação, todas as etapas do processamento linguístico estão interligadas (o que pode ser visto através das setas). Essas interligações podem ser mais fortes ou mais fracas e a maioria é bidirecional, indicando que os processos são retomados o tempo inteiro.

De Groot (2011) credita a três motivos os efeitos de facilitação cognata: os cognatos são coativados através da sobreposição ortográfica no momento da entrada da informação, a partilha da (morfo)semântica entre línguas e coativação das representações fonológicas durante a produção linguística. Dijkstra *et al.* (2018) acrescenta um quarto motivo da facilitação cognata

na produção linguística: a ativação lexical que se expande desde as representações ortográficas até as representações fonológicas.

Devido à forma compartilhada com os itens da L1 durante a aquisição da L2, presume-se que os cognatos poderiam ser o primeiro caminho do aprendiz para o novo léxico, o que resultaria em padrão diferente de processamento para cognatos e não-cognatos. Dessa maneira, como consequência dos efeitos de facilitação cognata, várias pesquisas demonstram que cognatos são recuperados mais rapidamente da memória e são menos propensos a serem esquecidos do que não cognatos. Então, cognatos seriam mais fáceis de aprender do que não cognatos, como afirmam De Groot e Van Hell (2005).

Evidentemente, a constatação de que cognatos são mais fáceis de aprender do que não cognatos traz algumas implicações na aprendizagem de L2, como a inferência bem-sucedida do significado de palavras desconhecidas que são cognatas, gerando confiança no desenvolvimento do vocabulário da nova língua. Contudo, Casaponsa *et al.* (2015) argumentam que o efeito de facilitação cognata se torna mais relevante em estágios iniciais de aprendizagem de L2, uma vez que os cognatos podem ter uma função de guia para o desenvolvimento da competência linguística nos primeiros anos de contato com a nova língua.

Dentro dos estudos da Psicolinguística do Bilinguismo, que estuda a representação mental e os processos cognitivos que nos permitem adquirir, compreender e falar mais de uma língua, os cognatos têm um lugar especial. Nestes estudos e nas pesquisas de aquisição de L2, geralmente se chega à conclusão de que as palavras cognatas são mais facilmente processadas do que as palavras não cognatas. O fenômeno chamado de efeito de facilitação cognata já foi verificado tanto em tarefas de compreensão quanto em tarefas de produção, sob diferentes paradigmas experimentais.

Por exemplo, em Dijkstra *et al.* (1999), foram pesquisados monolíngues e bilíngues, a L1 era holandês e a L2 era inglês. Foram realizados três experimentos, nos quais os bilíngues foram participantes nos dois primeiros e os monolíngues no último. Os dois primeiros experimentos mostraram resultados bem parecidos, indicando efeitos facilitadores de processamento, por meio da semelhança ortográfica e semântica. No entanto, a sobreposição fonológica não foi identificada. No terceiro experimento, que envolvia apenas monolíngues, os resultados da mesma tarefa de decisão lexical realizada pelos bilíngues como segunda tarefa se mostraram não relacionados ao material de estímulo.

Na pesquisa de Schwartz e Kroll (2006), a L1 era espanhol e a L2 era inglês. Em tarefas de nomeação de figuras, o efeito de facilitação cognata foi significativo quando foram utilizadas sentenças mais comuns, o que sugere ativação das duas línguas. No entanto, não

houve efeito de facilitação cognata em sentenças consideradas raras.

Casaponsa *et al.* (2015) investigaram como os cognatos podem contribuir para o processo de compreensão da escrita, que é o processamento da leitura. As línguas pesquisadas foram espanhol como L1 e inglês como L2 e os participantes possuíam nível básico ou intermediário em L2. Foi detectado que em níveis mais baixos de proficiência, o efeito de facilitação cognata foi positivamente relacionado com o desempenho em uma tarefa de leitura. Essa relação tornou-se negativa em níveis intermediários de aprendizagem de L2.

Através de dois experimentos de decisão lexical, Dijkstra *et al.* (2015) examinaram como a linguagem de uma sentença pode afetar o reconhecimento de itens lexicais em L1 (holandês) ou L2 (inglês). Os participantes viam as palavras em L1 ou L2 isoladamente ou após a leitura de sentenças em L1 ou L2 apresentadas palavra por palavra. Os resultados mostraram que os cognatos em L2 foram mais facilmente lidos tanto quando foram apresentados isolados, quanto quando foram exibidos após sentenças.

O estudo de Sheng *et al.* (2016) examinou o desempenho em tarefas de nomeação de figuras em três grupos diferentes de bilíngues com inglês como L2. Foram investigadas 34 crianças de 4 a 7 anos bilíngues com L1 espanhol e 52 crianças com L1 mandarim e 37 crianças monolíngues falantes de inglês. Os bilíngues com espanhol como L1 apresentaram resultados robustos para o efeito de facilitação cognata, em detrimento do mandarim como L1 e dos monolíngues falantes de inglês. Estes resultados podem indicar que a semelhança interlinguística pode interferir na aprendizagem de vocabulário, diminuindo o custo de processamento.

Poort e Rodd (2017) realizaram dois experimentos de decisão lexical, utilizando os idiomas holandês e inglês, como L1 e L2, respectivamente. Os resultados mostraram significativo efeito de facilitação cognata no experimento com cognatos, palavras controle e não palavras. No experimento que incluiu homógrafos interlinguais, pseudohomófonos e palavras em L1, o efeito de facilitação cognata foi significativamente reduzido.

Em Soares *et al.* (2018), os participantes tinham português europeu como L1 e inglês como L2 com diferentes níveis de proficiência. A tarefa experimental era a conclusão de uma sentença, com cognatos e não cognatos encaixados em um sintagma nominal complexo dentro de orações relativas, a fim de resolver ambiguidades sintáticas. Os cognatos nos sintagmas nominais complexos modularam os resultados, independentemente da proficiência, mas os pesquisadores esperavam maior interferência dos cognatos na sintaxe da L1, em comparação com os não cognatos.

Toassi e Pereira (2019) realizaram um experimento com monolíngues que nunca

tinham tido contato com inglês como L2. Em uma tarefa de tradução com palavras da língua inglesa, mais de metade das palavras apresentadas foram reconhecidas pelos participantes, o que indica que as palavras cognatas são facilmente reconhecidas.

Antón e Duñabeitia (2020) analisaram os efeitos de facilitação cognata utilizando uma língua fictícia como L2. Os participantes foram divididos em dois grupos: o grupo A aprendeu itens lexicais cognatos e não cognatos e o grupo B aprendeu apenas itens lexicais não cognatos. Logo após, a memória dos participantes foi testada em uma tarefa de correspondência de palavras e imagens e em uma tarefa de reconhecimento de tradução. Os resultados indicaram que os cognatos foram muito mais lembrados do que os não cognatos. Entretanto, os itens lexicais não cognatos foram mais lembrados no grupo B, quando não estavam acompanhados dos cognatos.

Os efeitos de facilitação cognata podem ser evidenciados também em multilíngues, como em Toassi *et al.* (2020). Nesta pesquisa, os participantes eram brasileiros com português como L1, alemão como L2 e inglês como L3. Foi realizada uma tarefa de rastreamento ocular com 60 sentenças, que possuíam cognatos duplos entre a L1 e a L3 e a L2 e a L3 e cognatos triplos. Os resultados mostraram que os cognatos triplos foram processados mais rapidamente, evidenciando um acesso lexical não seletivo e um léxico integrado para todas as línguas dos multilíngues.

O efeito de facilitação cognata pressupõe coativação interlinguística e dá suporte à hipótese de não seletividade das línguas em mentes bilíngues, o que causa menor custo linguístico. Isto quer dizer que quando um bilíngue está usando uma de suas línguas, a outra também está ativa, como Traxler (2012) defende.

Neste estudo, os efeitos de facilitação cognata foram examinados na estrutura alvo: os sintagmas nominais preposicionados com dois substantivos. Conforme detalhamento que está apresentado na metodologia, investigamos estes efeitos, considerando três preditoras: (1) a presença de itens lexicais cognatos nos dois substantivos que formam o sintagma nominal; (2) a presença de palavras transparentes apenas no núcleo do sintagma nominal; e (3) sem a presença de itens cognatos na construção sintática.

Dessa forma, buscamos entender se os efeitos de facilitação cognata podem surgir no processamento de sintagmas nominais, analisando a acurácia e o tempo de resposta na tarefa de reconhecimento de tradução e na tarefa de julgamento de aceitabilidade, e a preferência sintática e o tempo de resposta na tarefa de seleção sintática. Neste caso, haveria pressuposição de ativação não seletiva das línguas, o que costuma ser constatado em estudos com itens lexicais isolados.

2.5 Nível de proficiência em L2

Cardoso (2016) afirma que o conceito de proficiência vem sendo debatido nas últimas décadas de forma intensa em fóruns, congressos, conferências e simpósios. Um dos principais pontos de discussão é de que forma o conceito de proficiência deve ser considerado. Vejamos, portanto, algumas considerações sobre proficiência.

Nos estudos sobre aquisição e processamento de L2, a mensuração do nível de habilidade no uso dos idiomas que um bilíngue domina é um dispositivo metodológico bastante comum, conforme Souza (2019). Dentro da Psicolinguística do Bilinguismo, a proficiência em L2 pode ser compreendida como uma dimensão cognitiva capaz de moldar diferencialmente a arquitetura do processamento da linguagem. Em outras palavras, a proficiência em L2 pode estar relacionada à disponibilidade de representações gramaticais da L2. Além disso, pode também estar ligada à fluidez de acesso a essas representações nos eventos de processamento linguístico em tempo real.

Souza (2019) explica que a noção de proficiência em L2 é remissiva ao campo da testagem de habilidades linguísticas. Dessa maneira, trata-se de um conceito que indica clara conotação psicométrica, uma vez que se trata de um construto ou um traço latente.

Qian e Lin (2020) fazem alguns apontamentos sobre a proficiência em L2 e acreditam que o desenvolvimento lexical é mais oneroso para a L2 do que para a L1 porque aprendizes de L2 enfrentam duas restrições práticas. A primeira restrição se refere à falta de oportunidades para aplicar seus insumos lexicais com quantidade e qualidade satisfatórias. Ou seja, há falta de oportunidades suficientes para que os aprendizes sejam expostos à L2 fora da sala de aula, o que pode gerar dificuldades significativas para que os conhecimentos semânticos, sintáticos e morfológicos sejam extraídos de uma palavra. A outra restrição é a existência de um sistema conceitual da L1 já estabelecido. Isso pode gerar uma dependência excessiva da L1 para aprender novas palavras, especialmente quando os aprendizes de L2 são bilíngues tardios com a L1 muito bem estabelecida.

O conhecimento de um vocabulário pode ser observado sob a perspectiva quantitativa e sob a perspectiva qualitativa, como apresentam Qian e Lin (2020). O aspecto quantitativo do conhecimento lexical refere-se ao número de palavras existentes no léxico do aprendiz, o que inclui o conhecimento do aprendiz sobre a forma e o significado de uma palavra que pode ser simplesmente lembrada. Já o aspecto qualitativo do conhecimento do vocabulário, por outro lado, refere-se ao conhecimento lexical profundo e extenso do aluno e à capacidade de usar a palavra de maneira adequada e eficiente. Ter aprendido uma palavra não significa

apenas que o aluno pode recordar a forma e recuperar o significado da palavra em um teste de vocabulário, mas também implica o uso efetivo dessa palavra para fins de comunicação autêntica, como ler um livro ou escrever uma mensagem de e-mail.

O conhecimento de palavras também pode ser dividido em receptivo e produtivo, para Nation (2001). Para o autor, o conhecimento receptivo se dá quando uma palavra é reconhecida na leitura ou na audição, enquanto o conhecimento produtivo ocorre quando a palavra é usada na fala ou na escrita. Além disso, o vocabulário receptivo é formado pelo tamanho (amplitude) do vocabulário no léxico mental de um aprendiz de L2. Esta parte do conhecimento do vocabulário precisa ser transformada em conhecimento produtivo para ser usada em um contexto comunicativo. A transformação do vocabulário receptivo para o produtivo pode ser vista em um continuum, que começa com a familiaridade superficial com a palavra e termina com a habilidade de usar a palavra corretamente na produção livre.

As investigações sobre aquisição de vocabulário de L2 normalmente se concentram no tamanho do conhecimento do vocabulário, e não em sua profundidade, como também ponderam Qian e Lin (2020). Isso ocorre, provavelmente, devido à dificuldade em desenvolver medidas de qualidade para avaliar a dimensão da profundidade. Por outro lado, a construção direta de testes de tamanho de vocabulário, que normalmente medem apenas uma dimensão superficial do conhecimento lexical de um aluno, resulta em uma variedade de medidas para avaliar a amplitude do conhecimento de vocabulário.

Nas pesquisas científicas, os testes de tamanho de vocabulário podem desempenhar um papel importante na previsão do sucesso na proficiência dos alunos, como Stæhr (2009) conseguiu demonstrar. Em seu estudo, foi possível encontrar evidência de uma forte associação positiva do tamanho de vocabulário com habilidades de leitura e escrita e uma associação positiva moderada com sua compreensão auditiva.

Como se trata de um construto ou traço latente, é possível investigar a proficiência em um idioma e relacioná-lo a qualquer fenômeno linguístico que se queira examinar. Na dissertação de Jesus (2012), foram investigados 8 monolíngues e 32 bilíngues, com português como L1 e inglês como L2, a fim de examinar se os sistemas de memória seriam afetados pelo bilinguismo. Os bilíngues foram divididos em grupos, de acordo com sua proficiência em L2. Os resultados gerais mostraram que a maioria das comparações entre bilíngues e monolíngues favoreceu os bilíngues no desempenho de tarefas de memória. Ademais, houve uma diferença de desempenho muito significativa favorecendo o grupo de alta proficiência em relação aos participantes com baixa proficiência e aos monolíngues, sugerindo um efeito positivo de proficiência em L2 nessas tarefas.

Em uma pesquisa com dois experimentos, Bultena *et al.* (2014) verificaram se a troca de direção dos idiomas em traduções pode ser modulada pela proficiência em L2. Neste estudo, a L1 foi holandês e a L2 foi inglês e o paradigma utilizado foi a leitura automonitorada. Foi observado custo de processamento linguístico da L1 para a L2, mas não da L2 para a L1, o que indica relação com a proficiência em L2.

Casaponsa *et al.* (2015) pesquisaram o efeito de facilitação cognata, considerando o nível de proficiência em L2. Nesta pesquisa, o espanhol era a L1 e o inglês era a L2. Os resultados indicaram que em níveis mais baixos de proficiência, o desempenho em uma tarefa de leitura foi significativamente melhor. No entanto, em níveis intermediários de aprendizagem de L2, o desempenho foi negativo.

Lan *et al.* (2019) investigaram a relação entre o uso de sintagmas nominais complexos e a proficiência em L2. Neste estudo, a L1 foi japonês e a L2 foi inglês. Através de análises estatísticas, foi comprovado que quanto maior a proficiência em L2 do bilingue, mais sintagmas nominais complexos são construídos, através do uso de adjetivos atributivos e orações relativas. Em contrapartida, participantes com baixa ou média proficiência em L2 produziram sintagmas nominais menos complexos, com menos recursos relacionados a seu núcleo.

O estudo de Liu *et al.* (2022) examinou se a exposição ao priming sintático em línguas semelhantes (mandarim como L1 e cantonês como L2) seria afetado pela proficiência em L2. Os resultados de dois experimentos mostraram que o efeito de priming ocorreu independentemente da proficiência em L2.

O exame do nível de proficiência em L2 dos bilíngues e sua relação com a construção sintática utilizada para representar sintagmas nominais também em L2 foi uma das limitações identificadas em minha dissertação (Miranda, 2021). Em meu estudo anterior, restringimos o alcance a pesquisar bilíngues que cursavam inglês como L2 nos semestres 3 e 4, que apresentaram média de 57,5% de proficiência, com desvio padrão de 25,39%. O nível de proficiência desta pesquisa foi medido através do sítio gratuito <https://itt-leipzig.de/wortschatztests>, que pertence ao Instituto de Pesquisa e Desenvolvimento de Testes da Universidade de *Leipzig*, na Alemanha, e é capaz de testar o conhecimento lexical em quinze idiomas diferentes, no formato produtivo e no formato receptivo. Este escopo foi ampliado na pesquisa atual, incluindo bilíngues que estão estudando L2 em todos os estágios de aprendizagem.

No caso dos bilíngues que investigamos, eles podem cursar até 6 semestres de L2, cada um com 60 horas/aula. A proposta pedagógica da escola de idiomas que os bilíngues

cursam explica que os níveis de habilidades linguísticas que são almejados ao longo dos semestres são distribuídos da seguinte maneira, considerando o Quadro Comum Europeu (2001) como referência: semestres 1 e 2 – nível A1, semestres 3 e 4 – nível A2 e semestres 5 e 6 – nível B1.

Sabemos que o nível de proficiência em L2 costuma levar em consideração compreensão e produção oral e escrita. Mas, por uma questão prática, em meu estudo utilizaremos um teste de conhecimento de vocabulário em L2 para auxiliar a traçar o perfil dos participantes, levando em consideração seu nível de domínio de L2, ao qual consideraremos como nível de proficiência em L2.

Assim, buscamos compreender se podemos relacionar o processamento de sintagmas nominais preposicionados com dois substantivos ao nível de proficiência em L2, verificando se estas estruturas em L2 são processadas sintaticamente por meio da L1 dos bilíngues e se estratégias de processamento sintático exclusivas podem ser desenvolvidas, buscando uma mistura de construções das duas línguas, que vão se transformando conforme a proficiência na L2 vai se tornando mais alta, como sugerem Bernolet e Hartsuiker (2018).

A seguir, apresentamos a metodologia desta tese, com informações acerca dos participantes, dos materiais e instrumentos, do desenho do estudo, da coleta e da análise de dados e dos métodos experimentais que foram utilizados nesta pesquisa.

3 METODOLOGIA

Para Warren (2013), a Psicolinguística do Bilinguismo é uma ciência que busca experimentalmente entender como a L2 é aprendida, processada e produzida. Nesta pesquisa, investigamos o processamento de sintagmas nominais preposicionados com dois substantivos, explorando as três construções sintáticas possíveis: preposição *of*, genitivo e inversão de substantivos, a fim de estabelecer possível relação com o efeito cognato e com o nível de proficiência em L2. Para tanto, este estudo, com viés psicolinguístico, foi desenvolvido por meio de três experimentos, que estão detalhados nesse capítulo. Retomamos, então, o objetivo geral e os objetivos específicos apresentados nas considerações iniciais, bem como as perguntas e hipóteses que norteiam este estudo.

Também, apresentamos os participantes, o desenho do estudo e como foi realizada a coleta de dados, os materiais utilizados, que nos auxiliaram a traçar o perfil dos participantes da pesquisa e a controlar as variáveis deste estudo, bem como os procedimentos adotados para a análise dos dados coletados.

Igualmente, abordamos os métodos experimentais selecionados para analisar o processamento de sintagmas nominais preposicionados com dois substantivos por bilíngues tardios brasileiros com inglês como L2, assim como descrevemos cada um dos três experimentos.

Salientamos que a coleta de dados foi iniciada após submissão e aprovação do projeto de pesquisa pelo Comitê de Ética em Pesquisa com Seres Humanos da Universidade Federal do Ceará, com CAEE 79724124.9.0000.5054.

3.1 Objetivos

Este estudo apresenta um objetivo geral e três objetivos específicos, como descritos abaixo:

3.1.1 Objetivo geral

O objetivo geral desta tese foi analisar o processamento de sintagmas nominais preposicionados com dois substantivos por bilíngues tardios, que têm português como L1 e inglês como L2.

3.1.2 Objetivos específicos

O objetivo geral desta tese deu origem a três objetivos específicos, que foram: (1) identificar qual construção sintática em L2 tem menor custo de processamento, dentre as três estruturas possíveis: preposição *of*, caso genitivo e inversão dos substantivos; (2) examinar como o efeito cognato pode interferir no processamento destes sintagmas nominais; e (3) investigar qual a relação entre o custo de processamento de sintagmas nominais e o nível de proficiência em L2.

3.2 Questões e hipóteses da pesquisa

As perguntas e hipóteses deste estudo estão relacionadas com o processamento de construções sintáticas em L2 na mente bilíngue e sua possível relação com o efeito cognato e com o nível de proficiência em L2. Portanto, foram investigadas as seguintes questões e hipóteses:

Questão 1: Qual construção sintática tem o menor custo de processamento, considerando as três estruturas possíveis: preposição *of*, caso genitivo e inversão de substantivos?

Hipótese 1: O custo de processamento da estrutura com a preposição *of* será menor, porque se assemelha mais com a L1. Esta hipótese se baseia em Bernolet e Hartsuiker (2018), que apontam que a compreensão da L2 pode ser impulsionada pela sintaxe da L1, especialmente em estágios iniciais de aquisição da L2.

Questão 2: Como os cognatos interferem no processamento de sintagmas nominais em L2?

Hipótese 2: Os cognatos irão interferir no processamento de L2, aumentando a acurácia e diminuindo o tempo de resposta dos itens. Para esta hipótese, tivemos vários estudos como referência, tais quais Dijkstra *et al.* (2015), Poort e Rodd (2017), Soares *et al.* (2018), Toassi e Pereira (2019) e Antón e Duñabeitia (2020), que detectaram que os cognatos podem facilitar o processamento de L2.

Questão 3: Como o nível de proficiência em L2 pode afetar o processamento de sintagmas nominais?

Hipótese 3: Os bilíngues com menor nível de proficiência em inglês como L2 apresentarão menos acurácia e precisarão de mais tempo para processar os sintagmas nominais. Os estudos de Bernolet e Hartsuiker (2018) e Lan *et al.* (2019) nos embasam para esta hipótese.

Bernolet e Hartsuiker (2018) propuseram um modelo de aquisição de L2 para bilíngues tardios, que indica diferentes processamentos de L2 para diferentes estágios de proficiência. E Lan *et al.* (2019) identificaram que o processamento de sintagmas nominais é dependente do nível de proficiência dos bilíngues.

3.3 Participantes

Este estudo experimental foi aplicado com bilíngues tardios, conforme classificação proposta por Wei (2000), que são brasileiros com português como L1 e que estudam inglês como L2. Além disso, de acordo com a descrição apresentada pelas mesmas autoras, podemos classificar o contexto dos participantes desta pesquisa como coordenado, pois ocorre em contextos separados, de formas diferentes. Isso ocorre porque os participantes usam a L1 em seu cotidiano e a L2 é utilizada de forma preponderante na sala de aula de inglês. Pelo mesmo motivo o grau de uso das línguas é recessivo. Já o bilinguismo é incipiente, pois a língua não nativa ainda está em desenvolvimento.

Limberger *et al.* (2023) afirmam que o bilinguismo crescente nos últimos anos fomenta o interesse de pesquisadores sobre fenômenos complexos, multifacetados e diversificados, que podem ser observados na aprendizagem de línguas em cursos livres, através de recursos tecnológicos ou em aulas particulares. Assim, o convite para a participação nesta pesquisa foi feito a aprendizes de inglês como L2 do Centro Cearense de Idiomas (doravante, CCI). Trata-se de um curso livre de idiomas, cujo público-alvo são alunos que estão regularmente matriculados no Ensino Médio da rede pública estadual do Ceará. O CCI também atende professores da rede pública estadual, que desejam estudar uma língua estrangeira. Tendo em vista a amplitude dos CCI e a viabilidade da pesquisa, o convite foi feito a discentes da unidade Jóquei.

A divulgação da pesquisa foi feita em 11 turmas de diferentes semestres de inglês como L2 no CCI. Das 11 turmas convidadas, 10 turmas tinham como alunos adolescentes do Ensino Médio e 1 turma tinha como alunos adultos, que são professores da rede estadual de ensino, e são aprendizes de inglês como L2 no CCI.

Após o convite feito pela pesquisadora nas 11 turmas, 66 bilíngues se mostraram interessados em participar do estudo, sendo 25 maiores de idade e 41 menores de idade. Estes últimos levaram o TCLE para seus responsáveis autorizarem suas participações, o que era condição para iniciar a participação na pesquisa. Dos 66 interessados, 64 efetivaram a participação na pesquisa: 23 maiores de idade e 41 menores de idade.

Após a análise prévia dos dados coletados, detectamos que os dados relativos ao tempo de resposta de 7 participantes não estavam completos (4 menores de idade e 3 maiores de idade). Esses dados que estavam incompletos não registraram o tempo de resposta de alguns itens desses 7 participantes, aparecendo 0 na coluna do tempo. Acreditamos que os problemas técnicos que tenham ocasionado essa coleta parcial do tempo de resposta dos itens estejam ligados à oscilação de energia e, consequentemente, de internet, uma vez que os dados foram coletados *online*, em tempo real. Dessa maneira, esta pesquisa contou com 57 participantes, de forma efetiva, com todos os dados registrados de forma completa.

3.4 Desenho do estudo e coleta de dados

Esta pesquisa foi realizada através do sítio gratuito PsyToolkit (Stoet, 2010; 2017) versão 3.4.6. Trata-se de um sítio voltado a pesquisas e experimentos psicológicos programáveis, utilizado por pesquisadores do mundo todo, capaz de suportar idiomas diferentes. O PsyToolkit (Stoet, 2010; 2017) está disponível no endereço <https://www.psychtoolkit.org/>.

Kim *et al.* (2019) realizaram comparações entre o PsyToolkit e o E-prime 3.0, que é um software bastante utilizado em pesquisas de Psicolinguística e concluíram que o PsyToolkit é válido metodologicamente como meio para coletar dados, como a acurácia e o tempo de resposta de forma satisfatória. Outrossim, este aparato é bastante utilizado pelo nosso grupo de pesquisa, o Laboratório de Processamento da Linguagem de Bilíngues e Multilíngues (PLIBIMULT, daqui por diante). Gadelha (2021), Miranda (2021), Freitas (2023), Borém (2023) e Sousa (2024) são algumas das pesquisas psicolinguísticas desenvolvidas no PLIBIMULT, que utilizaram o PsyToolkit como meio para a coleta de dados de maneira bem-sucedida.

Como tivemos participantes menores e maiores de idade nesta pesquisa, disponibilizamos dois endereços diferentes para acesso aos experimentos. Isso ocorreu porque os participantes menores de idade precisavam iniciar a pesquisa lendo e aceitando o Termo de Assentimento Livre e Esclarecido (doravante, TALE), enquanto os participantes maiores de idade precisavam ler e aceitar o Termo de Consentimento Livre e Esclarecido (a partir de agora, TCLE). Dessa maneira, o link para menores de idade foi <https://www.psychtoolkit.org/c/3.4.4/survey?s=JDyHs> e para os participantes maiores de idade o link foi <https://www.psychtoolkit.org/c/3.4.6/survey?s=zThgR>.

Para validar o desenho dos três experimentos que fazem parte dessa pesquisa, foi realizado um estudo piloto com os membros do grupo de pesquisa PLIBIMULT, que não

sabiam do que se tratava o estudo. Alguns ajustes foram realizados com base nas sugestões dadas, a fim de deixar as instruções claras.

3.5 Análise e processamento de dados

Para as análises estatísticas descritiva e inferencial dos dados coletados, optamos por utilizar o programa RStudio (*RCore Development Team*, 2024) versão 2024.12.1. Lima Júnior (2021) apresenta alguns pontos positivos na utilização do RStudio. Dentre os benefícios de uso do RStudio, Lima Júnior (2021) cita: R é uma linguagem e um programa que dá poder e flexibilidade ao pesquisador; é mais rápido do que outros programas comerciais; consegue manipular conjunto de dados enormes; é capaz de produzir gráficos esteticamente agradáveis com muitas informações, podendo até mesmo ser interativos; e é provavelmente a ferramenta de análise de dados mais utilizada na academia atualmente.

3.6 Documentos, instrumentos e materiais

Apresentamos a seguir os documentos, instrumentos e materiais que foram utilizados neste estudo em detalhes.

3.6.1 TCLE e TALE

Nesta pesquisa, participaram menores e maiores de idade. Ao mostrar interesse em participar desta pesquisa, o pretenso participante menor de idade foi orientado a ler e a levar o TCLE para que seu responsável o preenchesse, o que ocorreu cerca de uma semana antes da sessão experimental. Nesta oportunidade, foi solicitada a leitura atenta do TCLE, a fim de evitar qualquer dúvida. Ao final, com a concordância do responsável, o documento foi preenchido e devolvido para a pesquisadora. Em seguida, foi agendado um momento individual com o participante para a realização da pesquisa, que iniciava com a leitura do TALE. Ao final da leitura, o participante fazia o aceite de participação na pesquisa, clicando em “Aceito”. Somente após leitura e aceite do TALE, o participante teve acesso a todas as etapas da pesquisa: questionário demográfico e linguístico, tarefa 1, tarefa 2, tarefa 3 e teste de conhecimento de vocabulário, nesta ordem. Tanto o TCLE quanto o TALE encontram-se nos apêndices A e B, respectivamente.

Já para o pretenso participante maior de idade foi agendado um momento individual para a realização da pesquisa, que iniciava com a leitura do TCLE. Em seguida, caso realmente

concordasse com sua participação, o aceite na pesquisa era feito, clicando em “Aceito”. Somente após leitura e aceite do TCLE, o participante teve acesso às fases seguintes da pesquisa: questionário demográfico e linguístico e linguístico, tarefa 1, tarefa 2, tarefa 3 e teste de conhecimento de vocabulário, nesta ordem.

3.6.2 Questionário demográfico e linguístico

Um desafio metodológico muito comum em estudos da Psicolinguística se refere à seleção de participantes e ao uso de ferramentas que possam auxiliar em seus perfilamentos, segundo Silva *et al.* (2022). Considerando o controle das variáveis que podem interferir no processamento da linguagem, Grosjean (1998) sugere que questionários sejam utilizados na seleção dos participantes de pesquisas linguísticas, com vistas a coletar e sintetizar informações que sejam relevantes. Dörnyei (2003) esclarece que os questionários podem gerar três tipos de dados: (1) factuais, como idade, gênero e escolaridade; (2) comportamentais, como hábitos de uso e estudo da L1 e da L2; e (3) questões atitudinais, como afiliação étnica ou cultural.

Decidimos utilizar questões factuais e questões comportamentais no questionário deste estudo, que nos auxiliará a perfilar os participantes e controlar algumas variáveis importantes, tais quais idade, gênero e tempo de experiência com a L2. O questionário demográfico e linguístico e linguístico que foi utilizado nesta pesquisa encontra-se no apêndice C.

3.6.3 Teste de conhecimento de vocabulário

Souza (2019) não julga suficientemente confiável a classificação de proficiência pautada apenas pela informação do tempo total de estudos de uma L2 ou pela exposição sistemática à L2, como o tempo de residência em uma comunidade em que a L2 é sociolinguisticamente dominante, por exemplo. Igualmente, o autor chama de questionável a autoavaliação do nível de proficiência ou o relato de autossatisfação no processo de aprendizagem da L2. No mesmo estudo, por exemplo, a correlação entre autoavaliação e proficiência medida por instrumento objetivo não alcançou o resultado desejável.

Concordamos com as colocações de Souza (2019) e, portanto, a fim de medir com precisão e confiabilidade a proficiência dos participantes desta pesquisa, utilizamos um teste de conhecimento de vocabulário, que foi adaptado pela professora Pâmela Toassi (UnB), em parceria com os professores Justin Lauro (Barrys University) e Bernhard Angele (Universidad de Madrid), a partir do teste de vocabulário do *Shipley Institute of Living Scale* (SILS), que é

uma avaliação que compara as pontuações de vocabulário e o pensamento abstrato. O teste de vocabulário do SILS é composto por 40 itens, em que o participante deve sublinhar a palavra que mais se assemelha à primeira palavra de cada linha.

No teste de vocabulário adaptado pelos professores Pâmela Toassi, Justin Lauro e Bernhard Angele, os itens para tradução foram os mesmos do teste SILS. Neste teste adaptado com 40 itens, os bilíngues devem selecionar um item lexical que tenha o mesmo significado ou o significado mais próximo do item a ser traduzido intralinguisticamente, em língua inglesa. Para tanto, os bilíngues têm até 3 segundos para clicar na opção que julgarem mais adequada.

Antes do desenvolvimento deste teste de vocabulário pelo grupo de pesquisa, os trabalhos desenvolvidos pelos pesquisadores deste grupo costumavam utilizar em seus estudos um teste de conhecimento de vocabulário, que pode ser acessado no sítio <https://itt-leipzig.de/wortschatztests>. Esta plataforma pertence ao Instituto de Pesquisa e Desenvolvimento de Testes da Universidade de *Leipzig* na Alemanha, que é um órgão de referência em testes psicométricos e avaliações, tendo desenvolvido uma ampla gama de testes confiáveis e validados, como teste de habilidades cognitivas, teste de habilidades interpessoais e teste de habilidades relacionadas ao trabalho. A plataforma é gratuita e pode testar o conhecimento lexical em quinze idiomas diferentes, no formato produtivo e no formato receptivo. O teste consiste em cinco subtestes, que registram níveis de conhecimento linguístico. O teste produtivo possui 90 itens *cloze* distribuídos em 5 níveis, que consistem na escrita (digitação) de palavras-chave para preencher as sentenças. Já o teste receptivo possui 150 itens, distribuídos em 5 níveis, que consistem na seleção de palavras que sejam a definição ou sinônimo dos termos apresentados. Neste formato, o participante tem acesso a um considerável número de palavras, porque para cada item há 6 opções de escolha. Em ambos os formatos, o teste de vocabulário dura no máximo 30 minutos. O tempo utilizado pelo participante é cronometrado pelo próprio sítio.

Os resultados alcançados no teste de vocabulário desta plataforma podem indicar que nível de proficiência o participante possui na habilidade escrita em L2, de acordo com o Quadro Comum Europeu (*Council of Europe*, 2020). A conclusão com sucesso dos dois primeiros níveis do teste, que explora as 2.000 palavras mais comuns do inglês, indica que o nível de proficiência corresponde ao A2. O terceiro nível do teste utiliza as 3.000 palavras mais comuns do idioma e o sucesso nessa etapa indica nível B1. Para ser considerado bilíngue com nível B2, é necessário que o participante seja bem-sucedido nos níveis 4 e 5 do teste, que avalia o conhecimento das 5.000 palavras mais comuns do inglês.

O participante deve ter no mínimo 80% de acerto em cada etapa, o que faz com que ele seja considerado aprovado naquele nível.

No próximo capítulo, apresentamos os experimentos que foram realizados em nesta pesquisa. Detalhamos os três experimentos que foram aplicados, os paradigmas escolhidos, os desenhos e os *corpora* da pesquisa.

4 OS EXPERIMENTOS

Apresentamos neste capítulo os três experimentos que foram aplicados nesta pesquisa, os paradigmas escolhidos, os desenhos e os *corpora* da pesquisa. Esclareço que todos os sintagmas nominais preposicionados com dois substantivos utilizados como itens nos três experimentos foram construídos a partir da observação das salas de aula de inglês como L2 e dos ambientes escolares na instituição em que trabalho. Por mais de um mês, fiz anotações de sintagmas nominais utilizados por alunos, professores e funcionários, tanto em L1 como em L2. Após este registro inicial, selecionei os itens que poderiam ser utilizados, conforme os controles metodológicos propostos em meu estudo, que serão detalhados nas sessões que tratam do desenho e do *corpus* de cada experimento.

4.1 Experimento 1

Optamos por iniciar a pesquisa com uma tarefa de reconhecimento de tradução, que foi aplicada a bilíngues com diferentes níveis de proficiência em L2. O objetivo foi identificar qual construção sintática é menos custosa cognitivamente para a representação em L2 de sintagmas nominais preposicionados com dois substantivos dentre as três construções possíveis: preposição *of*, caso genitivo e inversão de substantivos. Outrossim, buscamos relação entre os resultados obtidos com o efeito cognato e com o nível de proficiência em L2.

Dessa maneira, a primeira variável preditora foi as diferentes estruturas sintáticas possíveis para representar os sintagmas nominais preposicionados com dois substantivos. E a segunda variável preditora foi a ausência ou a presença de cognatos nestes sintagmas.

A primeira variável de resposta deste experimento foi a convergência dos julgamentos realizados pelos participantes, mediante a previsão de equivalência ou não dos itens. A segunda variável de resposta foi a latência temporal para a reação a cada um dos estímulos, ou seja, o tempo de reação para realizar ou não o reconhecimento da equivalência de traduções. Nas próximas subseções, discorreremos sobre o paradigma escolhido para este experimento, explicamos o desenho do experimento e apresentamos o *corpus* utilizado.

4.1.1 Tarefa de reconhecimento de tradução

Sunderman (2014) considera a tarefa de reconhecimento de tradução como uma boa estratégia para investigar o léxico bilíngue e explica que o processamento das palavras ocorre de forma mais natural e significativa do que, por exemplo, em uma tarefa de decisão lexical, que

é uma estratégia mais comumente utilizada em pesquisas bilíngues. Reconhecimento de tradução é uma tarefa direta, em que dois itens são apresentados para que o participante decida se os dois itens significam a mesma coisa ou se são equivalentes na tradução.

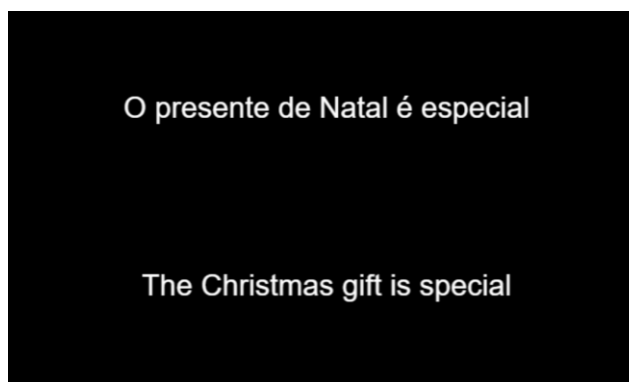
De Groot (1992) lista alguns fatores que justificam a utilização da tarefa de reconhecimento de tradução: (1) por não ser uma produção aberta como uma tarefa de tradução, é possível anular efeitos que antecedem a produção; (2) é possível vislumbrar o que os participantes estariam ativando em suas mentes e pretendendo produzir; e (3) para bilíngues com baixo nível de proficiência em L2 é uma tarefa mais adequada por ser menos custosa linguisticamente, o que pode auxiliar na análise do processamento bilíngue em L2.

Em tarefas de reconhecimento de tradução são utilizados as medidas de tempo de reação em milissegundos e a precisão das respostas, conforme Sunderman (2014). O mecanismo desta tarefa é muito simples, pois os participantes apenas precisam pressionar alguma tecla do computador para indicar sim ou não para a equivalência das traduções. Por um lado, se o participante responde em menor tempo e com maior precisão, isto significa que há alguma facilitação entre os fatores para o processamento da L2. Por outro lado, se o participante leva mais tempo e comete mais erros, há alguma interferência neste processamento.

4.1.2 Desenho e *corpus* do experimento 1

Os sintagmas nominais apareceram simultaneamente em L1 e em L2 no centro da tela do computador para que o participante realizasse seu julgamento, pressionando a tecla 1 para o reconhecimento de equivalência das traduções ou a tecla 2, para o não reconhecimento das traduções como equivalentes. O tempo máximo para que o participante realizasse o reconhecimento de tradução foi de 10.000 milissegundos (10 segundos). Quando o participante não fez o reconhecimento até o limite máximo de tempo, as sentenças desapareceram e surgiram os próximos itens a serem julgados. O processo foi repetido até que todas as traduções fossem reconhecidas como equivalentes ou não. Registramos as respostas dos participantes, assim como o tempo de resposta para cada item. Os participantes viram a tela, como mostra a figura 10:

Figura 10 - Tela com exemplo de item do experimento 1



Fonte: autoria própria

Utilizamos as mesmas sentenças em L1 para todos os participantes, mas distribuimos as sentenças em L2 em três listas de estímulos diferentes para que não houvesse efeito de sentenças nos resultados. Desta maneira, todos os participantes viram as mesmas sentenças em L1 mas cada lista de estímulo tinha uma construção sintática possível. As três listas foram distribuídas aos participantes de forma aleatória. Cada participante teve acesso a 60 itens experimentais e 45 itens distratores, totalizando 105 itens. Este cálculo da quantidade de itens experimentais e de itens distratores foi feito conforme Keating e Jegerski (2015), que propõem que a quantidade adequada dos itens distratores seja 75% dos itens experimentais. A ordem de apresentação dos itens também foi aleatorizada. Para os itens distratores, utilizamos sentenças no presente contínuo, conforme apresentado no apêndice D.

A respeito da frequência dos *corpora* dos experimentos que foram aplicados, optamos por pesquisar os vocábulos em dois bancos bastante utilizados em pesquisas de língua portuguesa e de língua inglesa: o LexPorBR e o SubtlexUS, respectivamente. O LexPorBR é um sítio que se encontra no endereço <https://www.lexicodoportugues.com/> (Estivalet, 2019) e possui mais de 200.000 palavras em seu banco de dados, com informações como significados, categoria semântica e a frequência com que aparece dentro do idioma. Já o SubtlexUS é um banco de palavras que leva em consideração o inglês americano, com mais de 51 milhões de itens e que tem por base legendas de filmes e séries de TV. O SubtlexUS pode ser acessado no sítio <http://www.lexique.org/> (Meta, 2023). Todas as frequências em escala Zipf dos itens lexicais que formam os sintagmas nominais das sentenças experimentais estão listados nos apêndices E e H. Nesta unidade de medida, a frequência das palavras varia de 1 a 7, em que palavras que têm frequências entre 1 e 2 são palavras mais raras, entre 3 e 4 são palavras que aparecem com certa frequência no idioma e com frequência 5, 6 ou 7 são vocábulos muito frequentes.

Alguns itens foram excluídos, conforme será explicado na seção dos resultados. Todos os itens excluídos estão indicados com asterisco (*) na acurácia. Apresentamos na tabela 1 as listas de estímulos que foram utilizados no experimento 1:

Tabela 1 – Lista de estímulos do experimento 1

ITENS SEM SUBSTANTIVOS COGNATOS							
ITEM	SENTENÇAS EM L1	LISTA 1	ACURÁCIA	LISTA 2	ACURÁCIA	LISTA 3	ACURÁCIA
1	O tamanho da blusa é médio	The size of shirt is medium	*	The shirt's size is medium	2	The shirt size is medium	1
2	A caixa de sapato é redonda	The shoe's box is round	2	The shoe box is round	1	The box of shoe is round	1
3	A vida de casado é ocupada	The married life is busy	*	The life of married is busy	*	The married's life is busy	*
4	A cadeira de praia é colorida	The beach of chair is colorful	2	The chair's beach is colorful	2	The chair beach is colorful	2
5	O sino da igreja é grande	The bell's church is big	2	The church bell is big	1	The church of bell is big	2
6	O pano de chão está sujo	The cloth floor is dirty	2	The floor of cloth is dirty	2	The cloth's floor is dirty	2
7	A moldura do espelho é oval	The frame of mirror is oval	*	The mirror's frame is oval	2	The mirror frame is oval	1
8	A capa do livro é branca	The book's cover is white	2	The book cover is white	1	The cover of book is white	*
9	O portão da casa é roxa	The house gate is purple	1	The gate of house is purple	*	The house's gate is purple	2
10	O saco de pão está cheio	The bag of bread is full	1	The bag's bread is full	2	The bag bread is full	2
11	Os olhos da garota são verdes	The girl's eyes are green	1	The eyes girl are green	2	The eyes of girl are green	*
12	O rabo do gato é longo	The tail cat is long	2	The tail of cat is long	*	The cat's tail is long	1
13	A tinta da caneta é azul	The ink of pen is blue	*	The pen's ink is blue	2	The pen ink is blue	1
14	O anel de prata é seu	The silver's ring is yours	2	The silver ring is yours	1	The ring of silver is yours	1
15	O lençol de algodão é macio	The cotton sheet is soft	1	The sheet of cotton is soft	1	The cotton's sheet is soft	2
16	O jogo de futebol é hoje	The soccer of match is today	2	The match's soccer is today	2	The match soccer is today	2
17	O vestido de festa é dourado	The dress' party is golden	2	The dress party is golden	2	The party of dress is golden	2
18	A viagem de navio é amanhã	The trip ship is tomorrow	2	The ship of trip is tomorrow	2	The trip's ship is tomorrow	2
19	A borracha do lápis é preta	The eraser of pencil is black	*	The pencil's eraser is black	2	The pencil eraser is black	1
20	O nível do mar	The sea's level	2	The sea level is	1	The level of	*

	está alto	is high		high		sea is high	
21	A tampa da panela é limpa	The pot lid is clean	1	The lid of pot is clean	*	The pot's lid is clean	2
22	O caminhão de lixo está lá	The truck of garbage is there	1	The truck's garbage is there	2	The truck garbage is there	2
23	O molho de carne é saboroso	The sauce's meat is tasty	2	The sauce meat is tasty	2	The meat of sauce is tasty	2
24	A loja de brinquedo é nova	The store toys is new	2	The store of toys is new	1	The store's toys is new	2
25	A água da piscina é aquecida	The water of pool is warmed	*	The pool's water is warmed	2	The pool water is warmed	1
26	O nariz de palhaço é vermelho	The clown's nose is red	1	The clown nose is red	1	The nose of clown is red	1
27	A lata de milho é amarela	The corn can is yellow	1	The can of corn is yellow	1	The corn's can is yellow	2
28	O quadro da sala é pequeno	The picture of room is small	*	The picture's room is small	2	The room picture is small	1
29	A torta de maçã é deliciosa	The pie's apple is delicious	2	The apple pie is delicious	1	The pie of apple is delicious	1
30	O suco de uva é doce	The grape juice is sweet	1	The grape of juice is sweet	2	The juice's grape is sweet	2
ITENS COM DOIS SUBSTANTIVOS COGNATOS							
ITEM	SENTENÇAS EM L1	LISTA 1	ACU RÁCIA	LISTA 2	ACU RÁCIA	LISTA 3	ACU RÁCIA
31	O museu de arte é perto	The museum of art is near	1	The art's museum is near	2	The art museum is near	1
32	A salada de fruta é boa	The fruit's salad is good	2	The fruit salad is good	1	The salad of fruit is good	1
33	O clipe de papel está aqui	The paper clip is here	1	The clip of paper is here	1	The paper's clip is here	2
34	O exame de rotina está ok	The routine of exam is ok	2	The exam's routine is ok	2	The routine exam is ok	1
35	O nome de família é curto	The name's family is short	2	The family name is short	1	The name of family is short	1
36	O carro de polícia é cinza	The car police is gray	2	The car of police is gray	1	The car's police is gray	2
37	O cartão de crédito é laranja	The card of credit is orange	1	The credit's card is orange	2	The credit card is orange	1
38	O fator de risco é alto	The risk's factor is high	2	The risk factor is high	1	The factor of risk is high	1
39	O teste de história é fácil	The History test is easy	1	The test of History is easy	1	The History's test is easy	2
40	A barra de chocolate é marrom	The chocolate of bar is brown	2	The bar's chocolate is brown	2	The chocolate bar is brown	1
41	A medalha de bronze é minha	The medal's bronze is mine	2	The medal bronze is mine	2	The medal of bronze is mine	1
42	A sopa do bebê está fria	The soup baby is cold	2	The baby of soup is cold	2	The soup's baby is cold	2
43	O filme de comédia é chato	The film of comedy is	1	The comedy's film is boring	2	The comedy film is boring	1

		boring					
44	O controle da TV é dele	The TV's control is his	2	The TV control is his	1	The control of TV is his	*
45	A estação de trem é longe	The train station is far	1	The station of train is far	1	The train's station is far	2
ITENS COM UM SUBSTANTIVO COGNATO							
ITEM	SENTENÇAS EM L1	LISTA 1	ACURÁCIA	LISTA 2	ACURÁCIA	LISTA 3	ACURÁCIA
46	A tecla do piano está muda	The piano of key is mute	2	The key's piano is mute	2	The piano key is mute	1
47	A fatia de tomate é fina	The slice's tomato is thin	2	The slice tomato is thin	2	The slice of tomato is thin	1
48	A comida do hospital é ruim	The food's hospital is bad	2	The hospital of food is bad	2	The food's hospital is bad	2
49	O banco do parque é velho	The bench of park is old	*	The park's bench is old	2	The park bench is old	1
50	A casca da manga é grossa	The mango's peel is thick	2	The mango peel is thick	1	The peel of mango is thick	*
51	O doce de banana é gostoso	The banana candy is tasty	1	The candy of banana is tasty	1	The banana's candy is tasty	2
52	O bolo de limão está pronto	The cake of lemon is ready	1	The cake's lemon is ready	2	The cake lemon is ready	2
53	O fio do fone é amarelo	The cord's phone is yellow	2	The phone cord is yellow	1	The phone of cord is yellow	2
54	A mesa de plástico é rosa	The plastic table is pink	1	The plastic of table is pink	2	The table's plastic is pink	2
55	O cabelo do estudante é loiro	The hair of student is blond	*	The student's hair is blond	1	The student hair is blond	1
56	A conta do dentista é barata	The dentist's bill is cheap	1	The dentist's bill is cheap	1	The bill of dentist is cheap	*
57	A dor de estômago é forte	The stomach ache is hard	1	The ache of stomach is hard	1	The stomach's ache is hard	2
58	A garrafa de café está cheia	The bottle of coffee is full	1	The bottle's coffee is full	2	The bottle coffee is full	2
59	O professor de inglês é legal	The teacher's English is kind	2	The English teacher is kind	1	The English of teacher is kind	2
60	A porta do hotel está aberta	The door hotel is open	2	The door of hotel is open	*	The door's hotel is open	2

Fonte: autoria própria

Como mostra a tabela 1, os 60 estímulos foram divididos em 30 sintagmas nominais sem cognatos (itens 1 a 30) e 30 sintagmas com cognatos (itens 31 a 60). Dentre os itens com cognatos, 15 tem os dois substantivos cognatos e 15 tem apenas o substantivo modificador cognato. Esta composição de itens se justifica porque decidimos incluir a análise dos cognatos.

Como tarefa de reconhecimento de tradução, precisamos ter itens com traduções corretas e incorretas. Na tabela 1, a acurácia 1 indica que a tradução está correta e acurácia 2 indica que a tradução está errada. Assim como em Freitas (2023), equilibramos igualmente a

acurácia das traduções. Desta forma, para cada lista com 60 itens experimentais, há 30 traduções que são equivalentes ou corretas e há 30 que não são equivalentes ou incorretas. Como este estudo investiga diferentes estruturas sintáticas da L2 e como em Miranda (2021) foi detectado que há dificuldade no processamento em L2 da ordem dos itens lexicais para que a sintaxe seja estruturada, especialmente com relação ao caso genitivo e à inversão dos substantivos, pois estas duas construções não se assemelham à sintaxe da L1, tornamos alguns itens incorretos através da ordem lexical equivocada dos substantivos que são parte dos sintagmas nominais. Para o caso genitivo, além de alterar a ordem das palavras, também consideramos a regra de uso dessa construção para que um item se tornasse com ou sem acurácia. Ou seja, não consideramos a naturalidade no uso do sintagma nominal, levamos em consideração a ordem lexical e as regras específicas de uso do caso genitivo, que já abordamos na seção 2.3 no capítulo dos referenciais teóricos.

Para o trabalho com os cognatos, controlamos os graus de similaridade lexical. Para isto, utilizamos o sítio <https://psico.fcep.urv.cat/utilitats/nim> (Guasch *et al.*, 2013). Trata-se de um browser com acesso gratuito baseado na web, que se chama NIM. Ele foi desenvolvido para auxiliar pesquisadores em estudos psicolinguísticos, conforme Guasch *et al.* (2013). O NIM possui diversas funções como calcular a similaridade ortográfica entre palavras de qualquer idioma, a frequência e o comprimento de palavras e encontrar vizinhos lexicais intra e interlinguísticos. Na ferramenta de calculadora de similaridade ortográfica do NIM, utilizamos a distância normalizada de Levenshtein ou NLD (*normalized Levenshtein distance*), que foi proposta por Schepens *et al.* (2012). Nesta medida, o valor do grau de similaridade varia entre 0 e 1, em que as palavras que são totalmente idênticas apresentam valor 1 de similaridade. O grau de similaridade dos itens lexicais cognatos que foram utilizados neste estudo variam entre 0,333 e 1 na distância normalizada de Levenshtein.

Vejamos na tabela 2 os substantivos cognatos das listas de estímulos e seus graus de similaridade:

Tabela 2 – Grau de similaridade dos cognatos

ITENS EXPERIMENTAIS	COGNATO EM L1/L2	GRAU DE SIMILARIDADE NA DISTÂNCIA NORMALIZADA DE LEVENSHTein (NLD)
Museu de arte/ Art museum	Museu/Museum	0,83
	Arte/Art	0,75

Salada de fruta/ Fruit salad	Salada/Salad	0,83
	Fruta/Fruit	0,60
Clipe de papel/ Paper clip	Clipe/Clip	0,80
	Papel/Paper	0,80
Exame de rotina/ Routine exam	Exame/Exam	0,80
	Rotina/Routine	0,71
Nome de família/ Family name	Nome/Name	0,75
	Família/Family	0,57
Carro de polícia/ Police car	Carro/Car	0,60
	Polícia/Police	0,57
Cartão de crédito/ Credit card	Cartão/Card	0,50
	Crédito/Credit	0,71
Fator de risco/ Risk factor	Fator/Factor	0,83
	Risco/Risk	0,60
Teste de História/ History test	Teste/Test	0,80
	História/History	0,75
Barra de chocolate/ Chocolate bar	Barra/Bar	0,60
	Chocolate/ Chocolate	1
Medalha de bronze/ Bronze medal	Medalha/Medal	0,71
	Bronze/Bronze	1
Sopa do bebê/ Baby soup	Sopa/Soup	0,50
	Bebê/Baby	0,50
Filme de comédia/ Comedy film	Filme/film	0,80
	Comédia/Comedy	0,71
Controle da TV/ TV control	Controle/Control	0,87
	TV/TV	1
Estação de trem/ Train station	Estação/Station	0,42
	Trem/Train	0,40
Tecla de piano/ Piano key	Piano/Piano	1
Fatia de tomate/ Tomato slice	Tomate/Tomato	0,83
Comida de hospital/ Hospital food	Hospital/Hospital	1

Banco do parque/ Park bench	Parque/Park	0,50
Casca da manga/ Mango peel	Manga/Mango	0,80
Doce de banana/ Banana candy	Banana/Banana	1
Bolo de limão/ Lemon cake	Limão/Lemon	0,40
Fio do fone/ Phone cord	Fone/Phone	0,60
Mesa de plástico/ Plastic table	Plástico/Plastic	0,75
Cabelo do estudante/ Student hair	Estudante/Student	0,66
Conta do dentista/ Dentist bill	Dentista/Dentist	0,87
Dor de estômago/ Stomach ache	Estômago/Stomach	0,50
Garrafa de café/ Coffee bottle	Café/Coffee	0,33
Professor de inglês/ English teacher	Inglês/English	0,57
Porta do hotel/ Hotel door	Hotel/Hotel	1
Mínimo		0,33
Máximo		1
Média		0,71
Mediana		0,75
Desvio padrão		0,18

Fonte: autoria própria

A tabela 2 mostra o grau de similaridade entre os cognatos que foram utilizados no experimento 1, que varia entre 0,33 (mínimo) e 1 (máximo), conforme cálculo no NIM (Guasch *et al.* 2013). Este cuidado metodológico é necessário para que haja o controle, evitando muitos itens idênticos (dentre os 45 cognatos, apenas 7 são idênticos) e itens cujos graus de similaridade sejam muito baixos. Outrossim, a média do grau de similaridade entre os 45 cognatos que é 0,71, a mediana é 0,75 e o desvio padrão é 0,18.

Para evitar itens raros nos experimentos 1 e 2, fizemos o controle da frequência dos substantivos, como está apresentado no apêndice E. Schmitt *et al.* (2001) fazem um alerta com relação ao cuidado com a seleção de palavras para a realização de um experimento, dando ênfase à necessidade de distribuição de frequência dos vocábulos, a fim de garantir que efeitos não intencionais surjam nos resultados. Para os experimentos, utilizamos a unidade de medida de frequência Zipf, que Van Heuven *et al.* (2014) propõem como melhor unidade para compararmos os *corpora* de línguas diferentes, com tamanhos diferentes. No caso do par de

línguas deste estudo, o *corpus* de inglês é bem mais extenso que o *corpus* de português, com 51 milhões de itens que constam no banco de SubtlexUS, comparando com 200.000 itens que estão no banco LexPorBR. Como já afirmado anteriormente, na escala Zipf, a frequência das palavras varia de 1 a 7, em que palavras que têm frequências entre 1 e 2 são palavras mais raras, entre 3 e 4 são palavras que aparecem com certa frequência no idioma e com frequência 5, 6 ou 7 são vocábulos muito frequentes.

Neste estudo, pesquisamos as frequências de todos os substantivos que formam os sintagmas nominais que foram utilizados nos experimentos e decidimos por considerar aqueles que possuem frequência entre 3 e 6 na escala Zipf. Assim, controlamos a distribuição de frequência das palavras e eliminamos os itens que são raros, com frequência menor que 3.

Dentre os substantivos em L1 da pesquisa, o substantivo com menor frequência é *lençol* (do sintagma *lençol de algodão*), que é o mínimo, com valor de 3,24 e o substantivo com maior frequência é *casa* (do sintagma *portão de casa*), que é o máximo, com valor 5,66. A média de frequência dos substantivos em L1 é 4,44, a mediana é 4,34 e o desvio padrão é 0,61.

Já dentre os substantivos em L2 deste estudo, o substantivo com menor frequência é *eraser* (do sintagma *pencil eraser*), com valor de 3,00, que é o mínimo, e o substantivo com maior frequência é *can* (do sintagma *corn can*), com valor 6,71, que é o máximo. A média de frequência dos substantivos em L2 é 4,64, a mediana é 4,68 e o desvio padrão é 0,60. Todos os dados referentes à frequência dos substantivos que formam os sintagmas nominais dos experimentos 1 e 2 estão no apêndice E.

Outro controle metodológico deste estudo foi o de número de caracteres por sentença, para que o tempo de leitura não varie de forma considerável de uma sentença para outra. O número de caracteres das sentenças usadas neste estudo variou de 20 (na sentença *O suco de uva é doce*), que é o mínimo, a 29 (nas sentenças *Os olhos da garota são verdes* e *O cabelo do estudante é loiro*), que é o máximo nas sentenças em L1. A média de número de caracteres é 25,61 caracteres, a mediana é 26 e o desvio padrão é 2,10. Nas sentenças em L2, as sentenças com uso da preposição *of* apresentou variação de números de caracteres de 22 (na sentença *The ink of pen is blue*), que é o mínimo, a 30 (na sentença *The teacher of English is kind*), que é o máximo. A média do número de caracteres é 26,35, a mediana é 27 e o desvio padrão é 1,82. Nas sentenças com o caso genitivo, a variação do número de caracteres é de 21 (na sentença *The pen's ink is blue*), que é mínimo, a 29 (na sentença *The English's teacher is kind*), que é o máximo. A média é 25,35, a mediana é 26 e o desvio padrão é 1,82. Nas sentenças com os substantivos invertidos, os caracteres variaram de 19 (na sentença *The pen ink is blue*), que é o mínimo, a 27 (na sentença *The English teacher is kind*), que é o máximo. A média

23,36, a mediana é 24 e o desvio padrão é 1,84. Todos os dados referentes aos números de caracteres das sentenças dos experimentos 1 e 2 estão no apêndice F.

Um dos controles possíveis desta pesquisa seria o dos variados sentidos que os sintagmas nominais preposicionados com dois substantivos podem estabelecer, como: posse, lugar-fim, material e especificação. Como já descrito na sessão 2.3, que trata da estrutura alvo deste estudo, este não é o foco da minha investigação. Fiz uma pesquisa sobre estas relações de sentido com os dados coletados em minha dissertação em Miranda *et al.* (2023), que não mostrou interferência dos sentidos variados estabelecidos no processamento dos sintagmas nominais.

4.2 Experimento 2

O segundo experimento que aplicamos foi uma tarefa de julgamento de aceitabilidade, que teve como objetivo a estimativa da janela temporal para a formação de julgamentos de diferentes construções sintáticas que possam representar sintagmas nominais preposicionados com dois substantivos em L2. Essa estimativa pode identificar qual construção sintática em L2 é menos custosa cognitivamente para os bilíngues brasileiros.

Assim, a primeira variável preditora foi as diferentes estruturas sintáticas possíveis para representar os sintagmas nominais preposicionados com dois substantivos. E a segunda variável preditora foi a presença de cognatos nestes sintagmas.

A primeira variável de resposta deste experimento foi a convergência dos julgamentos realizados pelos participantes, mediante a previsão de aceitabilidade ou não aceitabilidade dos itens. A segunda variável de resposta foi a latência temporal para a reação a cada um dos estímulos, ou seja, o tempo de reação para realizar o julgamento de cada construção sintática.

Assim como em Souza *et al.* (2015), optamos por um julgamento binário, que aponta aceitação ou rejeição dos itens. Decidimos por esse tipo de julgamento porque queremos ter com maior clareza o entendimento dos bilíngues brasileiros acerca do processamento de diferentes construções sintáticas para representar sintagmas nominais em L2. Outrossim, objetivamos analisar prioritariamente o tempo de reação mínimo médio para a emissão de um julgamento, com o objetivo de identificar a construção sintática que é linguisticamente menos custosa para os bilíngues brasileiros. Nas próximas subseções, discorreremos sobre o paradigma escolhido para este experimento, apresentamos o desenho do experimento e explicamos sobre o *corpus* utilizado.

4.2.1 Tarefa de julgamento de aceitabilidade

Para Sá *et al.* (2022), o teste de julgamento de aceitabilidade é um método experimental utilizado amplamente em estudos psicolinguísticos, que pode ser temporalizado ou não. Esse método consiste na avaliação de sentenças por participantes que costumam não ser tão familiarizados com o fenômeno investigado, para que haja o processamento e a compreensão da estrutura. Ao final, o participante deve atribuir uma nota à sentença que se caracteriza como um julgamento.

Há alguns cuidados metodológicos que devem ser adotados na preparação da tarefa de julgamento de aceitabilidade. O primeiro cuidado é a aleatorização dos estímulos, para que a ordem de apresentação das sentenças não afete o resultado da pesquisa. Além disso, deve haver a inclusão de sentenças experimentais e sentenças distratoras, por três motivos: (a) comparar notas atribuídas na tarefa; (b) avaliar se a sentença experimental é aceita pelos participantes ou não; (c) não evidenciar a construção ou propriedade linguística que está sendo pesquisada. A não evidência do que está sendo pesquisado pode evitar que os participantes fiquem menos sensíveis ao fenômeno, em caso de saturação causada por uma exposição extensiva a uma mesma construção, como afirmam Souza *et al.* (2015).

A escolha da escala é um dos pontos mais importantes do julgamento de aceitabilidade, como pontuam Sá *et al.* (2022). As escalas têm o objetivo de transformar observações individuais subjetivas em um parâmetro objetivo. Isto significa que as percepções qualitativas dos participantes passarão a ser entendidas de forma quantificada. As escalas construídas para tarefas de julgamento de aceitabilidade podem ter diferentes formatos, podendo ser binárias ou com mais opções para julgamento.

A escala binária utiliza apenas duas opções para julgamento, indicando aceitação ou rejeição dos itens. Este tipo de escala aponta com maior clareza o entendimento do item a ser julgado e analisa prioritariamente o tempo de reação para a emissão de um julgamento. As duas escalas com mais opções de julgamento que são as mais comuns nessas tarefas são a escala Likert e a escala com estimativa de magnitude.

A escala Likert pode ser simétrica ou assimétrica. São chamadas de simétricas as escalas que trazem uma avaliação neutra no meio da escala, com a mesma quantidade de valores entre esse ponto médio e os dois extremos. Aquelas que apresentam uma quantidade diferente de opções entre a média e os extremos são consideradas assimétricas. O número de pontos na escala também pode variar, mas a mais comum é a de cinco pontos.

Já em uma escala com estimativa de magnitude, o participante é exposto a um

estímulo padrão, mas outros estímulos são apresentados em seguida. Dessa maneira, a avaliação é feita baseada na comparação, relação, proporção ou diferença entre estímulos. Schültze (1996) explica que é possível explorar em tarefas de julgamento de aceitabilidade construções que não existem na língua ou que são raras de serem encontradas em contexto real, ou ainda que não sejam aceitáveis de um modo geral. Outrossim, Sá *et al.* (2022) apontam como vantagens deste tipo de tarefa a facilidade de acesso a dados quantitativos, a facilidade de planejamento e de programação.

4.2.2 Desenho e *corpus* do experimento 2

As construções sintáticas foram julgadas como bem-formadas ou malformadas no segundo experimento. Por um lado, se o participante aceitou o estímulo apresentado, considerando-o bem-formado, ele clicou no número 1 do teclado de seu computador. Por outro lado, se o participante julgou o estímulo apresentado malformado, ele clicou no número 2, rejeitando sua estrutura. Cada estímulo permaneceu na tela por até 10.000 milissegundos ou até que o participante clicasse em um dos números para realizar seu julgamento. Os participantes viram a tela, como mostra a figura 11:

Figura 11 - Tela com exemplo de item do experimento 2



Fonte: autoria própria

Utilizamos as mesmas sentenças em L2 do experimento 1, com a mesma distribuição em três listas de estímulos, para que não houvesse efeito de sentenças nos resultados. A ordem de apresentação dos itens se deu de forma aleatória. Cada participante teve acesso a 60 itens experimentais e 45 itens distratores, o que totalizou 105 itens. Utilizamos as mesmas sentenças distratoras do experimento 1, conforme apêndice D.

Alguns itens foram excluídos, durante a análise dos dados, conforme será explicado

na seção dos resultados, assim como ocorreu no experimento 1. Também nesta lista, os itens excluídos estão indicados com asterisco (*) na acurácia. Vejamos mais uma vez as sentenças separadas nas três listas de estímulos:

Tabela 3 – Lista de estímulos do experimento 2

ITENS SEM SUBSTANTIVOS COGNATOS						
ITEM	LISTA 1	ACURÁCIA	LISTA 2	ACURÁCIA	LISTA 3	ACURÁCIA
1	The size of shirt is medium	*	The shirt's size is medium	2	The shirt size is medium	1
2	The shoe's box is round	2	The shoe box is round	1	The box of shoe is round	1
3	The married life is busy	*	The life of married is busy	*	The married's life is busy	*
4	The beach of chair is colorful	2	The chair's beach is colorful	2	The chair beach is colorful	2
5	The bell's church is big	2	The church bell is big	1	The church of bell is big	2
6	The cloth floor is dirty	2	The floor of cloth is dirty	2	The cloth's floor is dirty	2
7	The frame of mirror is oval	*	The mirror's frame is oval	2	The mirror frame is oval	1
8	The book's cover is white	2	The book cover is white	1	The cover of book is white	*
9	The house gate is purple	1	The gate of house is purple	*	The house's gate is purple	2
10	The bag of bread is full	1	The bag's bread is full	2	The bag bread is full	2
11	The girl's eyes are green	1	The eyes girl are green	2	The eyes of girl are green	*
12	The tail cat is long	2	The tail of cat is long	*	The cat's tail is long	1
13	The ink of pen is blue	*	The pen's ink is blue	2	The pen ink is blue	1
14	The silver's ring is yours	2	The silver ring is yours	1	The ring of silver is yours	1
15	The cotton sheet is soft	1	The sheet of cotton is soft	1	The cotton's sheet is soft	2
16	The soccer of match is today	2	The match's soccer is today	2	The match soccer is today	2
17	The dress' party is golden	2	The dress party is golden	2	The party of dress is golden	2
18	The trip ship is tomorrow	2	The ship of trip is	2	The trip's ship is	2

			tomorrow		tomorrow	
19	The eraser of pencil is black	*	The pencil's eraser is black	2	The pencil eraser is black	1
20	The sea's level is high	2	The sea level is high	1	The level of sea is high	*
21	The pot lid is clean	1	The lid of pot is clean	*	The pot's lid is clean	2
22	The truck of garbage is there	1	The truck's garbage is there	2	The truck garbage is there	2
23	The sauce's meat is tasty	2	The sauce meat is tasty	2	The meat of sauce is tasty	2
24	The store toys is new	2	The store of toys is new	2	The store's toys is new	2
25	The water of pool is warmed	*	The pool's water is warmed	2	The pool water is warmed	1
26	The clown's nose is red	1	The clown nose is red	1	The nose of clown is red	1
27	The corn can is yellow	1	The can of corn is yellow	1	The corn's can is yellow	2
28	The picture of room is small	*	The picture's room is small	2	The room picture is small	1
29	The pie's apple is delicious	2	The apple pie is delicious	1	The pie of apple is delicious	1
30	The grape juice is sweet	1	The grape of juice is sweet	2	The juice's grape is sweet	2
ITENS COM DOIS SUBSTANTIVOS COGNATOS						
ITEM	LISTA 1	ACURÁCIA	LISTA 2	ACURÁCIA	LISTA 3	ACURÁCIA
31	The museum of art is near	1	The art's museum is near	2	The art museum is near	1
32	The fruit's salad is good	2	The fruit salad is good	1	The salad of fruit is good	1
33	The paper clip is here	1	The clip of paper is here	1	The paper's clip is here	2
34	The routine of exam is ok	1	The exam's routine is ok	2	The routine exam is ok	1
35	The name's family is short	2	The family name is short	1	The name of family is short	1
36	The car police is gray	2	The car of police is gray	1	The car's police is gray	2
37	The card of credit is orange	1	The credit's card is orange	2	The credit card is orange	1
38	The risk's factor is high	2	The risk factor is high	1	The factor of risk is high	1
39	The History test is easy	1	The test of History is	1	The History's	2

			easy		test is easy	
40	The chocolate of bar is brown	2	The bar's chocolate is brown	2	The chocolate bar is brown	1
41	The medal's bronze is mine	2	The medal bronze is mine	2	The medal of bronze is mine	1
42	The soup baby is cold	2	The baby of soup is cold	2	The soup's baby is cold	2
43	The film of comedy is boring	1	The comedy's film is boring	2	The comedy film is boring	1
44	The TV's control is his	2	The TV control is his	1	The control of TV is his	*
45	The train station is far	1	The station of train is far	1	The train's station is far	2
ITENS COM UM SUBSTANTIVO COGNATO						
ITEM	LISTA 1	ACURÁCIA	LISTA 2	ACURÁCIA	LISTA 3	ACURÁCIA
46	The piano of key is mute	2	The key's piano is mute	2	The piano key is mute	1
47	The slice's tomato is thin	2	The slice tomato is thin	2	The slice of tomato is thin	1
48	The food's hospital is bad	2	The hospital of food is bad	2	The food's hospital is bad	2
49	The bench of park is old	*	The park's bench is old	2	The park bench is old	1
50	The mango's peel is thick	2	The mango peel is thick	1	The peel of mango is thick	*
51	The banana candy is tasty	1	The candy of banana is tasty	1	The banana's candy is tasty	2
52	The cake of lemon is ready	1	The cake's lemon is ready	2	The cake lemon is ready	2
53	The cord's phone is yellow	2	The phone cord is yellow	1	The phone of cord is yellow	2
54	The plastic table is pink	1	The plastic of table is pink	2	The table's plastic is pink	2
55	The hair of student is blond	*	The student's hair is blond	1	The student hair is blond	1
56	The dentist's bill is cheap	1	The dentist bill is cheap	1	The bill of dentist is cheap	*
57	The stomach ache is hard	1	The ache of stomach is hard	2	The stomach's ache is hard	2
58	The bottle of coffee is full	1	The bottle's coffee is full	2	The bottle coffee is full	2
59	The	2	The English	1	The English	2

	teacher's English is kind		teacher is kind		of teacher is kind	
60	The door hotel is open	2	The door of hotel is open	*	The door's hotel is open	2

Fonte: autoria própria

Diferentemente do que propusemos no experimento 1, a tabela 3 mostra que os participantes tiveram acesso somente a sentenças em L2, conforme o próprio desenho da tarefa selecionada. Acreditamos que este fator elevou o grau de dificuldade da tarefa 2, se comparada à tarefa 1, uma vez que os bilíngues não terão acesso à sentença em L1 para realizar comparações interlinguísticas.

Como já explicado na subseção 4.1.2, metade dos itens de cada lista de estímulos está correta (acurácia 1) e metade está incorreta (acurácia 2), como dividido em Freitas (2023). Os critérios de acurácia também já foram esclarecidos na mesma subseção. Relembramos que fizemos o controle da frequência dos substantivos, como apresentado no apêndice E, do número de caracteres, conforme apêndice F, e do grau de similaridade dos cognatos, como mostra a tabela 2.

4.3 Experimento 3

O terceiro experimento que aplicamos foi uma tarefa de seleção sintática, que teve como objetivo a estimativa da janela temporal média mínima para a tomada de decisão, tendo como possibilidade diferentes construções sintáticas para representar uma estrutura em L1. Essa estimativa foi capaz de identificar qual construção sintática em L2 foi menos custosa cognitivamente para os bilíngues brasileiros.

Dessa forma, a única variável preditora foi a presença ou a ausência de cognatos nos sintagmas preposicionados com dois substantivos.

A primeira variável de resposta deste experimento foi as seleções sintáticas realizadas pelos participantes, mediante a previsão de seleção dos itens, levando em consideração o nível de proficiência em L2. A segunda variável de resposta foi a latência temporal para a reação a cada um dos estímulos, ou seja, o tempo de reação para tomar a decisão de seleção de uma das construções sintáticas.

4.3.1 Seleção sintática

Chamamos de seleção sintática a tarefa em que se deve escolher uma construção da sintaxe para representar uma determinada ideia. Não encontramos na literatura psicolinguística nenhum estudo que utilize essa nomenclatura como método de pesquisa, sendo, portanto, uma definição minha.

Como esta pesquisa envolve duas línguas, o participante deveria ler a sentença em L1 e decidir qual construção sintática em L2 deveria representá-la. Como a estrutura alvo tem apenas uma possibilidade de representação sintática em L1 (o uso da preposição *de* e suas variações *da* ou *do*), mas possui três possibilidades sintáticas diferentes em L2 (o uso da preposição *of*, o caso genitivo e a inversão de substantivos), acreditamos se tratar de um bom paradigma de pesquisa, uma vez que pode estabelecer relação com várias interfaces, entre elas o nível de proficiência em L2.

4.3.2 Desenho e *corpus* do experimento 3

A partir de três possibilidades sintáticas (preposição *of*, caso genitivo e inversão dos substantivos), os participantes deveriam selecionar uma das três construções possíveis para representar em L2 um sintagma nominal preposicionado com dois substantivos. Analisamos se houve relação entre as escolhas sintáticas com a presença de cognatos e com o nível de proficiência em L2 dos bilíngues. Cada estímulo permaneceu na tela por até 30.000 milissegundos (30 segundos) ou até que o participante clicasse em uma das estruturas sintáticas para realizar sua escolha. Resolvemos ampliar o tempo para cada item neste experimento, em comparação com as tarefas 1 e 2, em que os participantes tiveram 10.000 milissegundos para responder, porque nesta tarefa os bilíngues tinham acesso a quatro sentenças de uma só vez. Os participantes viram a tela, como mostra a figura 11:

Figura 12 - Tela com exemplo de item do experimento 3



Fonte: autoria própria

Neste experimento não houve divisão de lista de estímulos, como ocorreu nos dois primeiros experimentos, porque no mesmo item o participante já tinha acesso às três estruturas sintáticas de uma só vez. No entanto, tanto a ordem de apresentação dos itens quanto as opções de resposta foram aleatorizadas, para evitar efeitos de ordem.

Cada participante teve acesso a 45 itens experimentais e 35 itens distratores, totalizando 80 itens. Assim como nos experimentos anteriores, este cálculo da quantidade de itens experimentais se baseia em Keating e Jegerski (2015). Foram utilizadas 15 sentenças iguais às sentenças dos experimentos 1 e 2 (itens 1 a 15), 15 sentenças com apenas o núcleo dos sintagmas igual aos dos experimentos 1 e 2 (itens 16 a 30) e 15 sentenças inéditas (31 a 45). Para os itens distratores, também utilizamos sentenças no presente contínuo, conforme apresentado no apêndice G. Vejamos a lista de estímulos do experimento 3:

Tabela 4 – Lista de estímulos do experimento 3

ITEM	CARACTERIZAÇÃO DOS SINTAGMAS NOMINAIS	SENTENÇAS EM L1	OPÇÃO SINTÁTICA 1 EM L2	OPÇÃO SINTÁTICA 2 EM L2	OPÇÃO SINTÁTICA 3 EM L2
ITENS SEM SUBSTANTIVOS COGNATOS					
1	Sintagmas repetidos dos experimentos 1 e 2	O jogo de futebol é hoje	The match of soccer is today	The soccer's match is today	The soccer match is today
2	Sintagmas repetidos dos experimentos 1 e 2	O vestido de festa é dourado	The dress of party is golden	The party's dress is golden	The party dress is golden
3	Sintagmas repetidos dos experimentos 1 e 2	A borracha do lápis é preta	The eraser of pencil is black	The pencil's eraser is black	The pencil eraser is black
4	Sintagmas repetidos dos experimentos 1 e 2	O caminhão de lixo está lá	The truck of garbage is there	The garbage's truck is there	The garbage truck is there
5	Sintagmas repetidos dos experimentos 1 e 2	O molho de carne é saboroso	The sauce of meat is tasty	The meat's sauce is tasty	The meat sauce is tasty

ITENS COM UM SUBSTANTIVO COGNATO					
6	Sintagmas repetidos dos experimentos 1 e 2	A tecla do piano está muda	The key of piano is mute	The piano's key is mute	The piano key is mute
7	Sintagmas repetidos dos experimentos 1 e 2	A fatia de tomate é fina	The slice of tomato is thin	The tomato's slice is thin	The tomato slice is thin
8	Sintagmas repetidos dos experimentos 1 e 2	A comida do hospital é ruim	The food of hospital is bad	The hospital's food is bad	The hospital food is bad
9	Sintagmas repetidos dos experimentos 1 e 2	O banco do parque é velho	The bench of park is old	The parks's bench is old	The park bench is old
10	Sintagmas repetidos dos experimentos 1 e 2	A casca da manga é grossa	The peel of mango is thick	The mango's peel is thick	The mango peel is thick
ITENS COM DOIS SUBSTANTIVOS COGNATOS					
11	Sintagmas repetidos dos experimentos 1 e 2	A salada de fruta é boa	The salad of fruit is good	The fruit's salad is good	The fruit salad is good
12	Sintagmas repetidos dos experimentos 1 e 2	O clipe de papel está aqui	The clip of paper is here	The paper's clip is here	The paper clip is here
13	Sintagmas repetidos dos experimentos 1 e 2	O exame de rotina está ok	The exam of routine is ok	The routine's exam is ok	The routine exam is ok
14	Sintagmas repetidos dos experimentos 1 e 2	O nome de família é curto	The name of family is short	The family's name is short	The family name is short
15	Sintagmas repetidos dos experimentos 1 e 2	O carro de polícia é cinza	The car of police is gray	The police's car is gray	The police car is gray
ITENS SEM SUBSTANTIVOS COGNATOS					
16	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	O tamanho da saia é médio	The size of skirt is medium	The skirt's size is medium	The skirt size is medium
17	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	A caixa de joias é redonda	The box of jewel is round	The jewel's box is round	The jewel box is round
18	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	A vida de solteiro é ocupada	The life of single is busy	The single's life is busy	The single life is busy
19	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	O pano de limpeza está sujo	The cloth of cleaning is dirty	The cleaning's cloth is dirty	The cleaning cloth is dirty
20	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	A capa do caderno é branca	The cover of notebook is white	The notebook's cover is white	The notebook cover is white
ITENS COM UM SUBSTANTIVO COGNATO					
21	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	A mesa de cristal é frágil	The table of crystal is fragile	The crystal's table is fragile	The crystal table is fragile
22	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	A conta do restaurante é cara	The bill of restaurant is expensive	The restaurant's bill is expensive	The restaurant bill is expensive
23	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	O cabelo da secretária é ruivo	The hair of secretary is red	The secretary's hair is red	The secretary hair is red
24	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	A porta do aeroporto está fechada	The door of airport is closed	The airport's door is closed	The airport door is closed
25	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	O professor de português é novato	The teacher of Portuguese is newbie	The Portuguese's teacher is newbie	The Portuguese teacher is newbie

ITENS COM DOIS SUBSTANTIVOS COGNATOS					
26	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	O museu de dinossauro é maravilhoso	The museum of dinosaur is wonderful	The dinosaur's museum is wonderful	The dinosaur museum is wonderful
27	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	O cartão de vacina é preto	The card of vaccine is black	The vaccine's card is black	The vaccine card is black
28	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	O teste de geografia é difícil	The test of geography is easy	The geography's test is easy	The geography test is easy
29	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	A barra de proteína é saudável	The bar of protein is healthy	The protein's bar is healthy	The protein bar is healthy
30	Apenas o núcleo do sintagma é repetido dos experimentos 1 e 2	O filme de suspense é surpreendente	The film of suspense is surprising	The suspense's film is surprising	The suspense film is surprising
ITENS SEM SUBSTANTIVOS COGNATOS					
31	Sintagmas inéditos na pesquisa	O relógio de parede está atrasado	The clock of wall is late	The wall's clock is late	The wall clock is late
32	Sintagmas inéditos na pesquisa	A coroa do rei é valiosa	The crown of king is valuable	The king's crown is valuable	The king crown is valuable
33	Sintagmas inéditos na pesquisa	A gaveta do armário está aberta	The drawer of closet is open	The closet's drawer is open	The closet drawer is open
34	Sintagmas inéditos na pesquisa	A bandeja de ovo é amarela	The tray of egg is yellow	The egg's tray is yellow	The egg tray is yellow
35	Sintagmas inéditos na pesquisa	A asa do frango está cozida	The wing of chicken is cooked	The chicken's wing is cooked	The chicken wing is cooked
ITENS COM UM SUBSTANTIVO COGNATO					
36	Sintagmas inéditos na pesquisa	A via de acesso está bloqueada	The way of access is blocked	The access' way is blocked	The access way is blocked
37	Sintagmas inéditos na pesquisa	O passeio de balão é emocionante	The ride of balloon is exciting	The balloon's ride is exciting	The balloon ride is exciting
38	Sintagmas inéditos na pesquisa	A rede de tênis é branca	The net of tennis is white	The tennis' net is white	The tennis net is white
39	Sintagmas inéditos na pesquisa	A placa do computador é nova	The board of computer is new	The computer's board is new	The computer board is new
40	Sintagmas inéditos na pesquisa	O cenário da foto é legal	The setting of photo is cool	The photo's setting is cool	The photo setting is cool
ITENS COM DOIS SUBSTANTIVOS COGNATOS					
41	Sintagmas inéditos na pesquisa	O artigo de opinião está perfeito	The article of opinion is perfect	The opinion's article is perfect	The opinion article is perfect
42	Sintagmas inéditos na pesquisa	O clube de golfe é amplo	The club of golf is wide	The golf's club is wide	The golf club is wide
43	Sintagmas inéditos na pesquisa	O botão do elevador está desligado	The button of elevator is off	The elevator's button is off	The elevator button is off
44	Sintagmas inéditos na pesquisa	O grupo de mensagem está calmo	The group of message is calm	The message's group is calm	The message group is calm
45	Sintagmas inéditos na pesquisa	A poluição do ar é preocupante	The pollution of air is worrying	The air's pollution is worrying	The air pollution is worrying

Fonte: autoria própria

A tabela 4 mostra que as opções de resposta dos 60 itens experimentais se dividem entre as três estruturas sintáticas em L2 pesquisadas e que, diferentemente do que ocorreu nos experimentos 1 e 2, poderia haver mais de um item adequado para a tradução. Isso ocorreu porque o objetivo nesse experimento foi verificar o processamento por meio da escolha ou seleção das três construções sintáticas, sem dar destaque para a acurácia. Entretanto, ressaltamos que para o caso genitivo, nem todas as traduções seriam adequadas.

Também como mostra a tabela 4, os itens incluem substantivos cognatos e não cognatos. De forma detalhada, temos 15 substantivos sem substantivos cognatos, que são os itens 1 a 5, 16 a 20 e 31 a 35; 15 sintagmas nominais com apenas o substantivo modificador cognato, que são os itens 6 a 10, 21 a 25 e 36 a 40; 15 sintagmas nominais com os dois substantivos cognatos, que são os itens 11 a 15, 26 a 30 e 41 a 45. Esta seleção de itens se justifica porque damos continuidade à investigação do efeito cognato.

Para o trabalho com os cognatos, controlamos seus graus de similaridade, assim como fizemos nos experimentos 1 e 2. Relembramos que o grau de similaridade dos itens lexicais cognatos que foram utilizados neste estudo variam entre 0,33 e 1 na distância normalizada de Levenshtein. Vejamos na tabela 5 os substantivos cognatos da lista de estímulos do experimento 3 e seus graus de similaridade:

Tabela 5 – Grau de similaridade dos cognatos do experimento 3

ITENS EXPERIMENTAIS	COGNATO EM L1/L2	GRAU DE SIMILARIDADE NA DISTÂNCIA NORMALIZADA DE LEVENSHTTEIN (NLD)
Salada de fruta/ Fruit salad	Salada/Salad	0,83
	Fruta/Fruit	0,60
Clipe de papel/ Paper clip	Clipe/Clip	0,80
	Papel/Paper	0,80
Exame de rotina/ Routine exam	Exame/Exam	0,80
	Rotina/Routine	0,71
Nome de família/ Family name	Nome/Name	0,75
	Família/Family	0,57
Carro de polícia/ Police car	Carro/Car	0,60
	Polícia/Police	0,57
Museu de dinossauro/ Dinosaur museum	Museu/Museum	0,83
	Dinossauro/Dinosaur	0,80

Cartão de vacina/ Vaccine card	Cartão/Card	0,50
	Vacina/Vaccine	0,71
Teste de geografia/ Geography test	Teste/Test	0,80
	Geografia/Geography	0,66
Barra de proteína/ Protein bar	Barra/Bar	0,60
	Proteína/Protein	0,75
Filme de suspense	Filme/Film	0,80
	Suspense/Suspense	1
Artigo de opinião/ Opinion article	Artigo/Article	0,57
	Opinião/Opinion	0,71
Clube de golfe/ Golf club	Clube/Club	0,80
	Golfe/Golf	0,80
Botão do elevador/ Elevator button	Botão/Button	0,50
	Elevador/Elevator	0,87
Grupo de mensagem/ Message group	Grupo/Group	0,60
	Mensagem/Message	0,87
Poluição do ar/ Air pollution	Poluição/Pollution	0,44
	Ar/Air	0,66
Tecla do piano/ Piano key	Piano/Piano	1
Fatia de tomate/ Tomato slice	Tomate/Tomato	0,83
Comida do hospital/ Hospital food	Hospital/Hospital	1
Banco do parque/ Park bench	Parque/Park	0,50
Casca da manga/ Mango peel	Manga/Mango	0,80
Mesa de cristal/ Crystal table	Cristal/Crystal	0,85
Conta do restaurante/ Restaurant bill	Restaurante/Restaurant	0,90
Cabelo da secretária/ Secretary hair	Secretária/Secretary	0,80
Porta do aeroporto/ Airport door	Aeroporto/Airport	0,66
Professor de português/ Portuguese teacher	Português/Portuguese	0,80
Via de acesso/ Access way	Acesso/Access	0,66

Passeio de balão/ Balloon ride	Balão/Balloon	0,57
Rede de tênis/ Tennis net	Tênis/Tennis	0,67
Carregador do computador/ Computer charger	Computador/Computer	0,70
Cenário da foto/ Setting photo	Foto/Photo	0,60
Mínimo do grau de similaridade dos 45 cognatos		0,44
Máximo do grau de similaridade dos 45 cognatos		1
Média do grau de similaridade dos 45 cognatos		0,73
Mediana do grau de similaridade dos 45 cognatos		0,75
Desvio padrão		0,20

Fonte: autoria própria

A tabela 5 mostra o grau de similaridade entre os cognatos que foram utilizados no experimento 3, que varia entre 0,44 (mínimo) e 1 (máximo), conforme cálculo no NIM (Guasch *et al.* 2013). Este cuidado metodológico é necessário para que haja o controle, evitando muitos itens idênticos (dentre os 45 cognatos, apenas 5 são idênticos) e itens cujos graus de similaridade sejam muito baixos. Outrossim, a média do grau de similaridade entre os 45 cognatos é 0,73, a mediana é 0,75 e o desvio padrão é 0,20.

Assim como fizemos com os experimentos 1 e 2, também pesquisamos as frequências de todos os substantivos que formam os sintagmas nominais que foram utilizados no experimento 3. Relembramos que decidimos por considerar aqueles que possuem frequência entre 3 e 6 na escala Zipf, a fim de controlar a distribuição de frequência das palavras e eliminamos os itens que são raros, com frequência menor que 3.

Dentre os substantivos em L1 do experimento 3, o substantivo com menor frequência é *bandeja* (do sintagma *bandeja de ovo*), com valor de 3,15 e o substantivo com maior frequência é *jogo* (do sintagma *jogo de futebol*), com valor 5,66. A média de frequência dos substantivos em L1 é 4,48, a mediana é 4,47 e o desvio padrão é 0,60.

Já dentre os substantivos em L2 do experimento 3, o substantivo com menor frequência é *eraser* (do sintagma *pencil eraser*), com valor de 3,00 e o substantivo com maior frequência é *way* (do sintagma *access way*), com valor 6,15. A média de frequência dos substantivos em L2 é 4,50, a mediana é 4,75 e o desvio padrão de 0,78. Todos os dados referentes à frequência dos substantivos que formam os sintagmas nominais do experimento 3 estão no apêndice H.

Controlamos metodologicamente também o de número de caracteres por sentença.

Este cuidado foi tomado para que o tempo de leitura não apresente variação considerável de uma sentença para outra. O número de caracteres das sentenças usadas no estudo variou de 23 (mínimo) a 35 (máximo) nas sentenças em L1, com média de 27,82 caracteres, mediana 27 e desvio padrão 3,36. Nas sentenças em L2, as sentenças com uso da preposição *of* apresentou variação de números de caracteres de 24 (mínimo) a 35 (máximo), a média de caracteres é 28,15, a mediana é 28 e o desvio padrão 2,99; nas sentenças com o caso genitivo, a variação do número de caracteres é 23 (mínimo) a 34 (máximo), a média é 27,13, a mediana é 27 e o desvio padrão é 3,00; nas sentenças com os substantivos invertidos, os caracteres variaram de 21 (mínimo) a 32 (máximo), com média 25,13, a mediana 25 e o desvio padrão é 2,99. Todos os dados referentes aos números de caracteres do experimento 3 estão no apêndice I.

No próximo capítulo, apresentamos a análise dos dados coletados nas três tarefas realizadas pelos participantes. Analisamos cada tarefa por vez, descrevendo os dados que se referem à acurácia e ao tempo de resposta para os experimentos 1 e 2. Para o experimento 3, apresentamos os dados de escolha/seleção de respostas e de tempo de resposta. Após a descrição dos resultados, apresentamos as considerações finais, que trazem algumas reflexões sobre os resultados obtidos nesta pesquisa e como os dados obtidos nas três tarefas podem estar interligados.

5 RESULTADOS

Nas seções deste capítulo apresentamos os resultados obtidos por meio da análise dos dados coletados no programa RStudio. Primeiramente, fizemos análises estatísticas descritivas. Em seguida, fizemos regressões logísticas, utilizando Modelo Misto Linear Generalizado (Bates *et al.* 2015), Modelo Misto Linear (Bates *et al.* 2015) e Modelo Logístico Multinomial (Agresti, 2002), com o objetivo de explorarmos todas as variáveis possíveis na pesquisa e analisar se as diferenças são significativas ou não. O Modelo Misto Linear Generalizado e o Modelo Misto Linear são modelos de efeitos mistos. Lima Júnior (2022) explica que modelos de efeitos mistos combinam de forma equilibrada efeitos fixos e efeitos aleatórios, com o objetivo de modelar para generalizar para novos dados, indo além de participantes ou itens específicos, mas informando que há dados dependentes e dados aleatórios. Por um lado, o Modelo Misto Linear Generalizado é o mais adequado para analisar variáveis preditoras binominais, como o uso da preposição *of* para representar o sintagma nominal preposicionado com dois substantivos, por exemplo. Por outro lado, o Modelo Misto Linear é o mais adequado para a análise de variáveis contínuas, como, por exemplo, o tempo de resposta. Já o Modelo Logístico Multinomial é utilizado para variáveis de resposta com mais de duas categorias.

Os resultados do questionário demográfico e linguístico e do teste de vocabulário nos auxiliaram a traçar o perfil dos participantes desta pesquisa e encontram-se na seção 5.1. Já os resultados dos três experimentos são apresentados nas seguintes seções: na seção 5.2, encontram-se os resultados da tarefa de reconhecimento de tradução; na seção 5.3, encontram-se os resultados da tarefa de julgamento de aceitabilidade; e na seção 5.4, encontram-se os resultados da tarefa de seleção sintática.

5.1 Perfil dos participantes

Apresentamos nesta seção, os resultados obtidos por meio da análise de dados do questionário demográfico e linguístico e do teste de vocabulário. Estes dois instrumentos foram utilizados a fim de traçar o perfil dos participantes desse estudo com mais detalhes.

5.1.1 Questionário demográfico e linguístico

A partir do questionário demográfico e linguístico contendo sete perguntas, apresentamos o perfil demográfico dos participantes da pesquisa de forma detalhada. A

primeira pergunta do questionário demográfico e linguístico foi: qual seu sexo biológico? Vejamos os dados referentes na tabela 6:

Tabela 6 - Gênero dos participantes

GÊNERO	PARTICIPANTES
FEMININO	25 (44%)
MASCULINO	32 (56%)

Fonte: autoria própria

P = 57

P = número de participantes

Dos 57 participantes, 25 (44%) são do gênero feminino e 32 (56%) são do gênero masculino. do total de participantes. Assim, tivemos a maioria dos participantes do gênero masculino.

A segunda pergunta do questionário foi: qual sua idade? A tabela 7 mostra as medidas padrões dos dados a respeito da idade:

Tabela 7 – Idade dos participantes

MEDIDAS	IDADE
MÉDIA (DESVIO PADRÃO)	20 (9,35)
MÍNIMO	15
MÁXIMO	59

Fonte: autoria própria

P = 57

P = número de participantes

A tabela 7 mostra que a idade média dos participantes foi 20 anos, com desvio padrão de 9,35. Os participantes mais novos a participarem desta pesquisa tinham 15 anos e o participante mais velho tinha 59 anos.

A terceira pergunta do questionário demográfico e linguístico foi: qual semestre você cursa no CCI? Vejamos as respostas coletadas na tabela 8:

Tabela 8 – Semestre dos participantes

SEMESTRE	PARTICIPANTES
2	25 (44%)

4	11 (19%)
5	08 (14%)
6	13 (23%)

Fonte: autoria própria

P = 57

P = número de participantes

Os alunos do CCI podem cursar até seis semestres de língua estrangeira. Nos resultados, como é possível observar na tabela 8, nenhum participante indicou cursar os semestres 1 e 3. Este fato ocorreu porque durante o período em que os dados foram coletados, esses semestres não foram ofertados na unidade em que o convite à participação nesta pesquisa foi feito. Dessa maneira, dos 57 participantes, 25 (44%) cursavam o semestre 2, 11 (19%) cursavam o semestre 4, 8 cursavam o semestre 5 (14%) e 13 (23%) cursavam o semestre 6. Dessa forma, a maioria dos participantes estavam no semestre 2 de estudo no CCI.

A quarta pergunta do questionário demográfico e linguístico foi: com qual das mãos você escreve? Vejamos as respostas coletadas na tabela 9:

Tabela 9 – Mão utilizada para escrever

MÃO	PARTICIPANTES
DIREITA	53 (93%)
ESQUERDA	04 (7%)
DUAS MÃOS	-

Fonte: autoria própria

P = 57

P = número de participantes

Dos 57 participantes, 53 (93%) são destros e 4 (7%) são canhotos. Nenhum dos participantes afirmou ser ambidestro.

A quinta pergunta do questionário demográfico e linguístico foi: desde que nível de ensino você estuda inglês na escola? Vejamos as respostas coletadas na tabela 10:

Tabela 10 – Nível inicial de estudo de inglês na escola

NÍVEL DE ENSINO	PARTICIPANTES
Educação Infantil	06 (10%)
Ensino Fundamental I (1º ao 5º ano)	21 (37%)

Ensino Fundamental II (6º ao 9º ano)	21 (37%)
Ensino Médio	09 (16%)

Fonte: autoria própria

P = 57

P = número de participantes

Dos 57 participantes, 6 (10%) estudam inglês desde a Educação Infantil, 21 (37%) estudam inglês desde o Ensino Fundamental I, 21 (37%) estudam inglês desde o Ensino Fundamental II e 9 (16%) estudam inglês desde o Ensino Médio.

A sexta pergunta do questionário demográfico e linguístico foi: por qual motivo você estuda inglês? Vejamos as respostas coletadas na tabela 11:

Tabela 11 – Motivos para estudar inglês

MOTIVO PARA ESTUDAR INGLÊS	PARTICIPANTES
Eu me identifico com a língua e gosto de estudá-la	37 (65%)
É importante para o futuro profissional	43 (75%)
Eu pretendo viajar para países que falam inglês	42 (74%)
Meus pais ou responsáveis me obrigam	02 (3,50%)

Fonte: autoria própria

P = 57

P = número de participantes

Nesta pergunta do questionário demográfico e linguístico, os participantes tinham a possibilidade de assinalar mais de um motivo pelo qual estudam inglês. Dos 57 participantes, 37 (65%) indicaram que se identificavam com a língua inglesa e que gostavam de estudá-la, 43 (75%) indicaram que estudar inglês é importante para o futuro profissional, 42 (74%) indicaram que pretendem viajar para países que falam inglês e apenas 2 (3,50%) indicaram que seus pais ou responsáveis o (a) obrigavam a estudar inglês. Esses dados indicam que a maioria dos participantes têm a compreensão de que estudar inglês como segunda língua vai contribuir de alguma forma para sua vida profissional (75%). Outrossim, um número considerável pretende viajar para países que falam inglês (74%) e se identificam com a língua e gostam de estudá-la (65%).

A sétima pergunta do questionário demográfico e linguístico foi: sem contar os dias que você estuda no CCI, com qual frequência você lê em inglês? Vejamos as respostas coletadas

na tabela 12:

Tabela 12 – Frequência de leitura em inglês

FREQUÊNCIA DE LEITURA	PARTICIPANTES
Todos os dias da semana	14 (25%)
De 3 a 5 vezes por semana	15 (26%)
2 ou 3 vezes por semana	19 (33%)
1 vez por semana	05 (9%)
Passo mais de uma semana sem ler em inglês	04 (7%)

Fonte: autoria própria

P = 57

P = número de participantes

Dos 57 participantes, 14 (25%) indicaram que lêem em inglês todos os dias, 15 (26%) indicaram que lêem em inglês entre 3 e 5 vezes por semana, 19 (33%) indicaram que lêem em inglês duas ou três vezes por semana, 5 (9%) indicaram que costumam ler em inglês uma vez por semana e 4 (7%) indicaram que passam mais de uma semana sem ler em inglês. Estes dados indicam que 48 participantes (84%) lêem em inglês pelo menos duas ou três vezes por semana, sem contar os dias que estudam inglês no CCI.

5.1.2 Teste de conhecimento de vocabulário

Também para auxiliar a compreender o perfil dos participantes, aplicamos o teste de conhecimento de vocabulário em inglês, que continha 40 itens. Vejamos as medidas padrões relativas ao teste na tabela 13:

Tabela 13 – Proficiência dos participantes

MEDIDAS	PROFICIÊNCIA (%)
MÉDIA (DESVIO PADRÃO)	36,15 (10,89)
MEDIANA	35
MODA	30
MÍNIMO	18
MÁXIMO	63

Fonte: autoria própria
 P = 57
 P = número de participantes

A tabela 13 mostra que a nota média de proficiência dos participantes foi 36,15%, com desvio padrão 10,89%. A mediana foi 35% e a moda foi 30%. A menor proficiência entre os participantes foi 18% a maior proficiência foi 63%.

Os dados apresentados na seção 5.1 indicam que, em resumo, a maioria dos participantes dessa pesquisa são do gênero masculino (56%) e têm idade média de 20 anos. A maioria dos participantes cursa o semestre 2 (44%), o que indica que quase a metade dos participantes está em fase inicial de aprendizagem de inglês como língua não nativa. Além disso, a maioria dos participantes é destro (93%), iniciou os estudos de inglês na escola no ensino fundamental I (37%) ou no ensino fundamental II (37%), estuda inglês porque compreende que a língua é importante para seu futuro profissional (75%) ou porque pretende viajar (74%) e leem pelo menos duas ou três por semana em inglês (84%). A média de proficiência dos participantes foi 36,15%.

5.2 Experimento 1

O primeiro experimento deste estudo consistiu em uma tarefa de reconhecimento de tradução da L1 para a L2, com 105 itens, dentre os quais 60 itens eram experimentais e 45 itens eram distratores. Cada item continha uma sentença em português e uma sentença em inglês, para que fosse julgado se elas eram equivalentes em tradução (tecla 1) ou se as traduções não eram equivalentes (tecla 2). Ressalte-se que a ordem dos 105 itens foi aleatorizada pelo Psytoolkit (Stoet, 2010; 2017) e que foram elaboradas três listas de estímulos diferentes. Para fazer o reconhecimento de cada tradução, o participante tinha até 10 segundos.

O objetivo foi identificar a construção sintática menos custosa cognitivamente para a representação em L2 de sintagmas nominais preposicionados com dois substantivos, dentre as três construções possíveis: preposição *of*, caso genitivo e inversão de substantivos. Outrossim, buscamos relação entre os resultados obtidos com o efeito cognato e com o nível de proficiência em L2.

Apresentaremos nesta seção a análise dos itens experimentais. Assim, salientamos que coletamos 5.985 dados, dos quais nos interessam os 3.420 dados que se referem aos itens experimentais.

Após uma análise inicial, decidimos excluir 21 itens de um total de 180 experimentais. Minha decisão se deu porque detectamos que não inserimos o artigo *the* em

algumas traduções que precisavam desse item lexical, como em *The tail of cat is long* para a tradução de *O rabo do gato é longo*. Como forma correta, a frase em inglês deveria ter o artigo *the* antes do segundo substantivo e ficaria *The tail of the cat is long*. Além dessa situação, percebemos que a sentença *A vida de casado é ocupada* deveria ser excluída porque suas possíveis traduções em L2 não configuram sintagmas nominais formados por dois substantivos. Isso ocorre porque a palavra *casado* em L1 pode ser substantivo ou adjetivo, mas em L2 a palavra *married* não pode ser considerada substantivo, apenas adjetivo. Dessa maneira, as três listas de estímulos do experimento 1 ficaram com as seguintes quantidades de itens para análise de resultados: lista 1 → 51 itens, lista 2 → 55 itens; e lista 3 → 53 itens. Após a exclusão desses itens, havia 3.016 itens experimentais.

Em seguida, analisamos os dados do experimento 1, no que diz respeito ao tempo de resposta. Dos 3.016 itens experimentais, 58 não foram respondidos dentro do tempo previsto de 10.000 milissegundos (10 segundos). Assim, 2.958 itens foram respondidos pelos participantes.

Também, identificamos 2 respostas com registros abaixo de 250 milissegundos. Estudos psicolinguísticos como os de Schoonbaert *et al.* (2009) e Duñabeitia *et al.* (2010) defendem um padrão de tempo mínimo para observação de fenômenos da Psicolinguística do Bilinguismo: 250 milissegundos. Abaixo desse registro, os autores acreditam que pode ter havido resposta antecipada, erro de ativação motora ou lapso de atenção e ponderam ser melhor não incluir esses itens na análise. Dessa maneira, excluímos as 2 respostas que registraram dados abaixo de 250 milissegundos e a análise feita levou em consideração 2.956 itens.

Nas subseções a seguir apresentamos o resultado da análise dos dados em seis perspectivas: análise geral, análise por sintaxe, análise por condição cognata, análise por quantidade de cognatos, análise por sintaxe e quantidade de cognatos e análise inferencial. Neste estudo, temos duas variáveis de resposta: a acurácia e o tempo de resposta. Ressaltamos que consideramos todos os itens respondidos para a análise da acurácia. Para a análise do tempo de resposta, levamos em consideração os itens respondidos com acurácia, que é uma prática padronizada nos estudos psicolinguísticos.

5.2.1 Análise geral

Considerando a primeira variável de resposta, analisamos o *status* geral dos itens experimentais respondidos:

Tabela 14 – Status geral do experimento 1

STATUS GERAL	ACURÁCIA	NÃO ACURÁCIA	TOTAL
QUANTIDADE	1.506 (50,95%)	1.450 (49,05%)	2.956 (100%)

Fonte: autoria própria

P = 57 participantes

P = número de participantes

VER = 2.956 variáveis experimentais respondidas

VER = número total de variáveis experimentais respondidas

Considerando os 2.956 itens experimentais respondidos pelos participantes, os dados da tabela apontam que a acurácia foi um pouco maior do que a não acurácia. Tivemos 1.506 (50,95%) itens com acurácia e 1.450 (49,05%) itens sem acurácia. Vejamos a acurácia geral em um gráfico de barras:

Gráfico 1 – Acurácia do experimento 1



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.956 itens respondidos

IR = número total de itens respondidos

O gráfico mostra visualmente os dados de acurácia geral dos itens respondidos. Demos destaque aos itens com acurácia, que representaram 50,95% dos itens respondidos.

Com relação à segunda variável de resposta, analisamos o tempo de resposta dos 1.506 itens respondidos com acurácia:

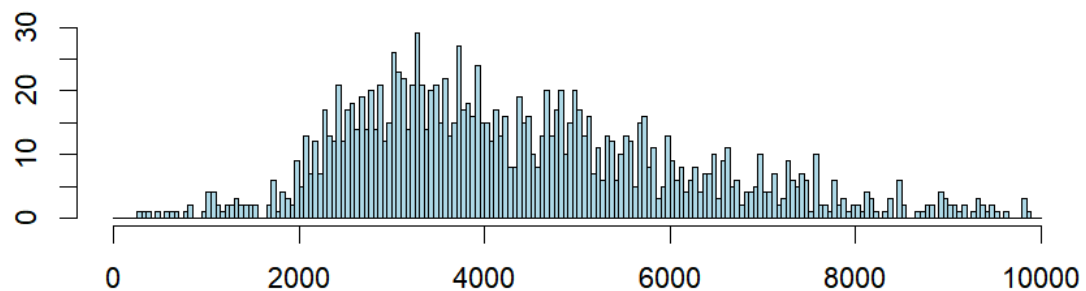
Tabela 15 – Tempo de resposta do experimento 1

MEDIDAS PADRÕES	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
TEMPO DE RESPOSTA (EM MS)	4.372 (1.772)	4.036	264	9.877

Fonte: autoria própria
 P = 57 participantes
 P = número de participantes
 IRC = 1.506 itens respondidos com acurácia
 IRC = número total de itens respondidos com acurácia

A tabela mostra que a média do tempo de resposta dos itens experimentais com acurácia foi 4.372 milissegundos, com desvio padrão 1.772 milissegundos. A mediana foi 4.036 milissegundos. O tempo mínimo utilizado para responder um item corretamente foi 264 milissegundos e o tempo máximo foi 9.877 milissegundos. Vejamos em um histograma a distribuição da média de tempo de resposta:

Gráfico 2 – Histograma do tempo de resposta (ms) do experimento 1



Fonte: autoria própria
 P = 57 participantes
 P = número de participantes
 IRC = 1.506 itens respondidos com acurácia
 IRC = número total de itens respondidos com acurácia

O histograma aponta que o tempo de resposta geral ficou bem distribuído. A maior parte das respostas foi dada entre 2.000 e 8.000 milissegundos, com maior concentração de respostas entre 2.000 e 6.000 milissegundos. Há alguns registros abaixo de 2.000 milissegundos e também acima de 8.000 milissegundos.

5.2.2 Análise por construção sintática

Analisamos, também, o *status* dos itens experimentais respondidos por construção sintática:

Tabela 16 – *Status* por construção sintática do experimento 1

STATUS POR CONSTRUÇÃO SINTÁTICA	ACURÁCIA	NÃO ACURÁCIA	TOTAL
OF	510 (64,64%)	279 (35,36%)	789 (100%)
GENITIVO	279 (26,25%)	784 (73,75%)	1.063 (100%)
INVERSÃO DE SUBSTANTIVOS	717 (64,95%)	387 (35,05%)	1.104 (100%)

Fonte: autoria própria

P = 57 participantes

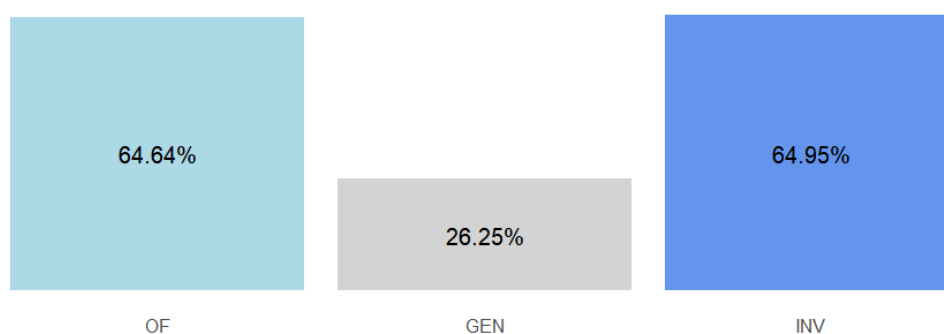
P = número de participantes

IR = 2.956 itens respondidos

IR = número total de itens respondidos

Os dados da tabela apontam que analisamos 789 sintagmas nominais com preposição *of*, 1.063 sintagmas nominais com genitivo e 1.104 sintagmas nominais com inversão de substantivos. Os dados também mostram que a acurácia foi maior do que a não acurácia para duas das três construções sintáticas possíveis para representar os sintagmas nominais preposicionados com dois substantivos: *of* e inversão de substantivos. A acurácia nos sintagmas nominais com preposição *of* foi 510 (64,64%) e nos sintagmas com inversão dos substantivos foi 717 (64,95%). A não acurácia se mostrou maior do que a acurácia apenas nos sintagmas nominais com genitivo, com índice de acurácia de 279 (26,25%). Vejamos a acurácia por construção sintática em um gráfico de barras:

Gráfico 3 - Acurácia por construção sintática do experimento 1



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.956 itens respondidos

IR = número total de itens respondidos

O gráfico mostra visualmente os dados de acurácia por construção sintática.

Destacamos os dois maiores índices de acurácia, que ficaram muito próximos: para sintagmas com inversão de substantivos e para sintagmas com preposição *of*: 64,95% e 64,64%, respectivamente.

Também, analisamos o tempo de resposta por construção sintática:

Tabela 17 – Tempo de resposta (ms) por construção sintática do experimento 1

TEMPO DE RESPOSTA (EM MS) POR CONSTRUÇÃO SINTÁTICA	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
OF	4.367 (1.738)	4.028	616	9.846
GENITIVO	4.727 (1.959)	4.529	264	9.877
INVERSÃO DE SUBSTANTIVOS	4.236 (1.701)	3.915	303	9.850

Fonte: autoria própria

P = 57 participantes

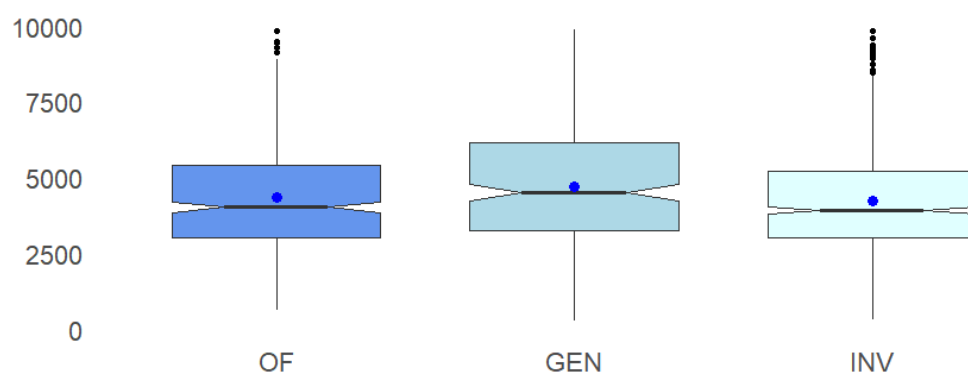
P = número de participantes

IRC = 1.506 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

A tabela mostra que as medidas padrões ficaram próximas, independentemente da construção sintática. Mesmo com dados muito próximos, a média de tempo se mostrou menor para a inversão dos substantivos: 4.236 milissegundos. O mesmo ocorreu com o desvio padrão e a mediana: 1.701 e 3.915 milissegundos, respectivamente. O menor tempo para responder um item com acurácia foi para o caso genitivo: 264 milissegundos. Já o maior tempo foi 9.877 milissegundos, também para um sintagma com genitivo. Vejamos em um gráfico de caixas o tempo de resposta por construção sintática:

Gráfico 4 - Tempo de resposta (ms) por construção sintática do experimento 1



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IRC = 1.506 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

Os gráficos de caixas mostram visualmente que tanto a média quanto a mediana ficaram próximas nas três construções sintáticas. No entanto, a média e a mediana foram maiores quando a construção sintática utilizou o caso genitivo. O mesmo aconteceu com o primeiro e o segundo quartis. Além disso, não apareceram *outliers* para os sintagmas com o caso genitivo.

5.2.3 Análise por condição cognata

Dando continuidade à análise detalhada dos 2.958 itens experimentais respondidos, investigamos o *status* dos itens experimentais respondidos por condição cognata:

Tabela 18 – *Status* por condição cognata do experimento 1

ACURÁCIA POR CONDIÇÃO COGNATA	ACURÁCIA	NÃO ACURÁCIA	TOTAL
SEM COGNATOS	651 (46,67%)	744 (53,33%)	1.395 (100%)
COM COGNATO	855 (54,77%)	706 (45,23%)	1.561 (100%)

Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.956 variáveis itens respondidos

IR = número total de itens respondidos

Os dados da tabela apontam que 1.395 itens não continham cognatos nos sintagmas nominais e 1.561 itens continham cognatos nos sintagmas nominais. Os dados também mostram que a acurácia foi maior do que a não acurácia apenas quando os itens continham cognatos (54,77%). Em itens sem cognatos, a acurácia foi menor do que a não acurácia (46,67%). Vejamos a acurácia por condição cognata em um gráfico de barras:

Gráfico 5 – Acurácia por condição cognata do experimento 1



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.956 itens respondidos

IR = número total de itens respondidos

O gráfico mostra a acurácia das duas condições cognatas. Demos destaque à acurácia dos itens com cognatos (54,77%), que foi maior do que a não acurácia.

Também fizemos a análise do tempo de resposta, levando em consideração a condição cognata:

Tabela 19 – Tempo de resposta (ms) por condição cognata do experimento 1

TEMPO DE RESPOSTA (EM MS) POR CONDIÇÃO COGNATA	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
SEM COGNATO	4.589 (1.841)	4.430	368	9.850
COM COGNATO	4.206 (1.700)	3.894	264	9.877

Fonte: autoria própria

P = 57 participantes

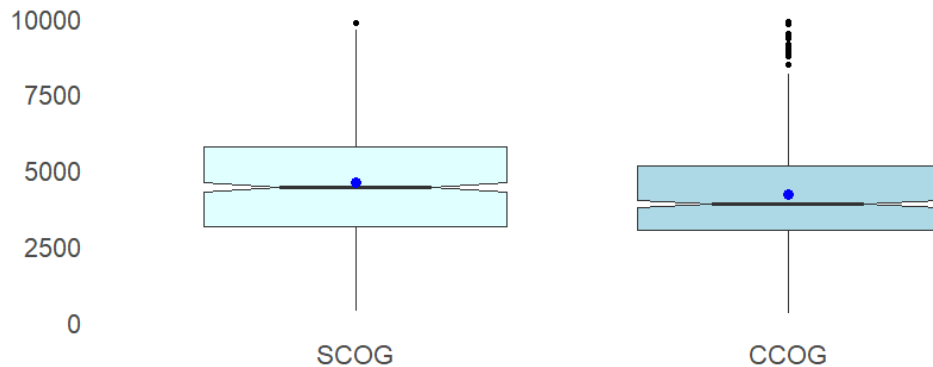
P = número de participantes

IRC = 1.506 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

A tabela mostra que as medidas padrões ficaram próximas, independentemente da condição cognata. A média, a mediana e o desvio padrão foram menores quando havia cognatos nos sintagmas, com registros de 4.206, 1.700 e 3.894 milissegundos, respectivamente. O menor tempo para responder um item foi quando havia cognato no sintagma: 264 milissegundos. O maior tempo foi para itens com cognatos, também: 9.877 milissegundos. Vejamos em um gráfico de caixas o tempo de resposta por condição cognata:

Gráfico 6 – Tempo de resposta (ms) por condição cognata do experimento 1



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IRC = 1.506 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

Os gráficos de caixas mostram visualmente que as medidas padrões ficaram próximas nas duas condições cognatas. Mas é possível perceber que a média e mediana estão menores para os itens com cognatos. O mesmo ocorreu com o primeiro e o segundo quartis. Os *outliers* apareceram em menor quantidade quando os sintagmas não tinham cognatos.

5.2.4 Análise por quantidade de cognatos

Nos itens com cognatos, metade continha um substantivo cognato e metade continha os dois substantivos cognatos. Essa divisão nos permitiu analisar se a quantidade de cognatos nos sintagmas interferiu nos resultados:

Tabela 20 – *Status* por quantidade de cognatos do experimento 1

ACURÁCIA POR POR QUANTIDADE DE COGNATOS	ACURÁCIA	NÃO ACURÁCIA	TOTAL
0 COGNATO	651 (46,67%)	744 (53,33%)	1.395 (100%)
1 COGNATO	387 (52,16%)	355 (47,84%)	742 (100%)
2 COGNATOS	468 (57,14%)	351 (42,86%)	819 (100%)

Fonte: autoria própria

P = 57 participantes

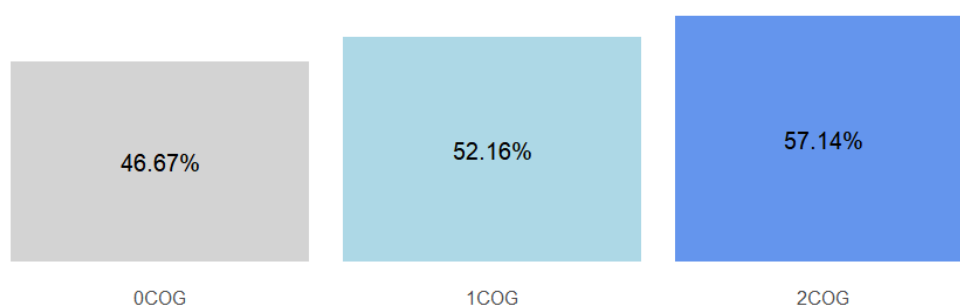
P = número de participantes

IRC = 2.956 itens respondidos

IRC = número total de itens respondidos com cognatos

A tabela mostra que 1.395 itens sem cognatos foram respondidos, 742 itens com um cognato foram respondidos e 819 itens com dois cognatos foram respondidos. A acurácia foi maior do que a não acurácia quando havia cognatos nos itens. Comparando as quantidades de cognatos nas três possibilidades, a acurácia se mostrou crescente, considerando itens sem cognatos, com um cognato e com dois cognatos, nesta ordem. Vejamos os dados da acurácia por quantidade de cognatos em um gráfico de barras:

Gráfico 7 – Acurácia por quantidade de cognatos do experimento 1



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.956 itens respondidos

IR = número total de itens respondidos

O gráfico mostra que a acurácia foi maior do que a não acurácia quando havia cognatos nos itens. Destacamos a acurácia dos sintagmas com dois cognatos (57,14%), que foi a maior entre as três possibilidades.

Analisamos, igualmente, o tempo de resposta dos itens por quantidade de cognatos:

Tabela 21 – Tempo de resposta (ms) por quantidade de cognatos do experimento 1

TEMPO DE RESPOSTA (EM MS) POR QUANTIDADE DE COGNATOS	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
0 COGNATO	4.589 (1.832)	4.430	368	9.850
1 COGNATO	4.499 (1.743)	4.274	264	9.877
2 COGNATOS	3.964 (1.625)	3.589	303	9.846

Fonte: autoria própria

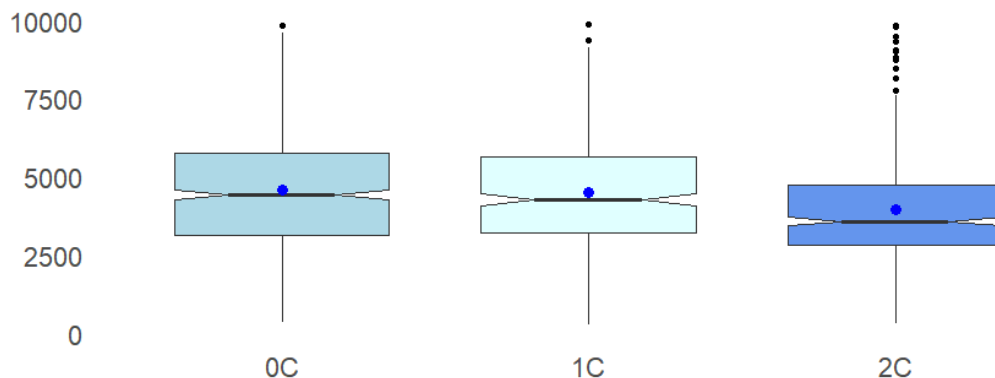
P = 57 participantes

P = número de participantes

IRC = 1.506 itens respondidos com acurácia
 IRC = número total de itens respondidos com acurácia

A tabela mostra que as medidas padrões ficaram próximas, independentemente da quantidade de cognatos. No entanto, todos os dados se mostraram menores para itens com dois cognatos. Além disso, a média, o desvio padrão e a mediana se mostraram decrescentes, considerando itens sem cognatos, com um cognato e com dois cognatos, nesta ordem. O menor tempo para responder um item foi para um sintagma com um cognato (264 milissegundos) e o maior tempo foi para um sintagma com um cognato (9.877 milissegundos). Vejamos o tempo de resposta por quantidade de cognatos em um gráfico de caixas:

Gráfico 8 – Tempo de resposta (ms) por quantidade de cognatos do experimento 1



Fonte: autoria própria
 P = 57 participantes
 P = número de participantes
 IRC = 1.506 itens respondidos com acurácia
 IRC = número total de itens respondidos com acurácia

Os gráficos de caixas mostram visualmente que as medidas de tempo se mostraram decrescentes, considerando itens sem cognatos, com um cognato e com dois cognatos, nesta ordem. Nesta mesma ordem, os *outliers* se mostraram crescentes. Isso significa que apareceram mais *outliers* quando os sintagmas tinham dois cognatos.

5.2.5 Análise por sintaxe e quantidade de cognatos

Após analisar as condições sintáticas e cognatas de forma isolada, resolvemos correlacionar as duas condições, uma vez que para cada construção sintática diferente, havia itens sem cognatos, com um cognato e com dois cognatos em suas estruturas:

Tabela 22 – *Status* por condição sintática e cognata do experimento 1

ACURÁCIA POR CONSTRUÇÃO SINTÁTICA E QUANTIDADE DE COGNATOS	ACURÁCIA	NÃO ACURÁCIA	TOTAL
OF SEM COGNATOS	201 (57,43%)	149 (42,57%)	350 (100%)
OF COM 1 COGNATO	114 (61,96%)	70 (38,04%)	184 (100%)
OF COM 2 COGNATOS	195 (76,47%)	60 (23,53%)	255 (100%)
GENITIVO SEM COGNATOS	135 (26,79%)	369 (73,21%)	504 (100%)
GENITIVO COM 1 COGNATO	92 (32,97%)	187 (67,03%)	279 (100%)
GENITIVO COM 2 COGNATOS	52 (18,57%)	228 (81,43%)	280 (100%)
INVERSÃO DE SUBSTANTIVOS SEM COGNATOS	315 (58,23%)	226 (41,77%)	541 (100%)
INVERSÃO DE SUBSTANTIVOS COM 1 COGNATO	181 (64,87%)	98 (35,13%)	279 (100%)
INVERSÃO DE SUBSTANTIVOS COM 2 COGNATOS	221 (77,82%)	63 (22,18%)	284 (100%)

Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.956 itens respondidos

IR = número total de itens respondidos

A tabela mostra que nos sintagmas com preposição *of* e sem cognatos a acurácia foi de 57,43%. A acurácia para esta construção sintática aumentou quando havia um substantivo cognato (61,96%) e aumentou mais ainda quando havia dois substantivos cognatos (76,47%). Nos sintagmas com uso de genitivo e sem cognatos, a acurácia foi de 26,79%. Nesta estrutura sintática, a acurácia aumentou quando havia um substantivo cognato (32,97%). No entanto, de forma surpreendente, a acurácia em sentenças com genitivo e dois substantivos cognatos se mostrou muito baixa (18,57%), o que indica que nesta condição a acurácia foi a menor de todas as condições cognatas. Nos sintagmas com inversão de substantivos e sem cognatos a acurácia foi de 58,23%. A acurácia para esta construção sintática aumentou quando havia um substantivo cognato (64,87%) e aumentou mais ainda quando havia dois substantivos cognatos (77,82%).

Dando destaque aos resultados relativos às construções sintáticas com caso genitivo, fica claro que os bilíngues apresentam dificuldade de domínio sintático deste tipo de estrutura. Primeiramente, podemos refletir sobre a semelhança visual entre o caso genitivo e a

inversão de substantivos, que se diferenciam por um sinal gráfico, o apóstrofo ('), e uma letra (s), tão somente. Apesar de muito semelhantes estruturalmente, o caso genitivo é utilizado para indicar relações específicas: de posse, familiares, de propósito ou de origem, como explica Murphy (2015). Já a inversão de substantivos pode ser utilizada em relações indicativas de posse e de tipificação, como Santos e Alonso (2022) esclarecem, tornando sua aplicação mais abrangente. Isto significa que o uso do caso genitivo é bem mais restrito do que o da inversão de substantivos, o que pode levar a uma dificuldade de processamento maior, em comparação com as outras duas construções sintáticas investigadas. Esta dificuldade pode levar a uma sobrecarga de processamento para estes itens. Há também a possibilidade de confusão com o plural pela semelhança visual, como a compreensão de *The pot's lid* como *A tampa das panelas*, por exemplo.

No que diz respeito ao efeito cognato, os resultados indicaram que a acurácia dos itens com caso genitivo aumentou quando os sintagmas nominais apresentaram um cognato, em comparação com itens sem cognatos. Esses resultados apontam para o efeito de facilitação cognata, assim como ocorreu nos itens com preposição *of* e com inversão de substantivos. No entanto, a acurácia atingiu os níveis mais baixos, quando os sintagmas tinham dois substantivos cognatos, o que parece um efeito de dificuldade exclusivamente para esses itens. Ou seja, é como se houvesse uma sobrecarga de processamento para os sintagmas com caso genitivo e dois cognatos. Isso pode ter ocorrido porque uma maior quantidade de substantivos cognatos pode criar a sensação de facilidade de processamento de L2, levando o bilíngue a negligenciar a marcação 's. Isto significa que pode ter havido a sobrecarga atencional, como explica Costa *et al.* (2006).

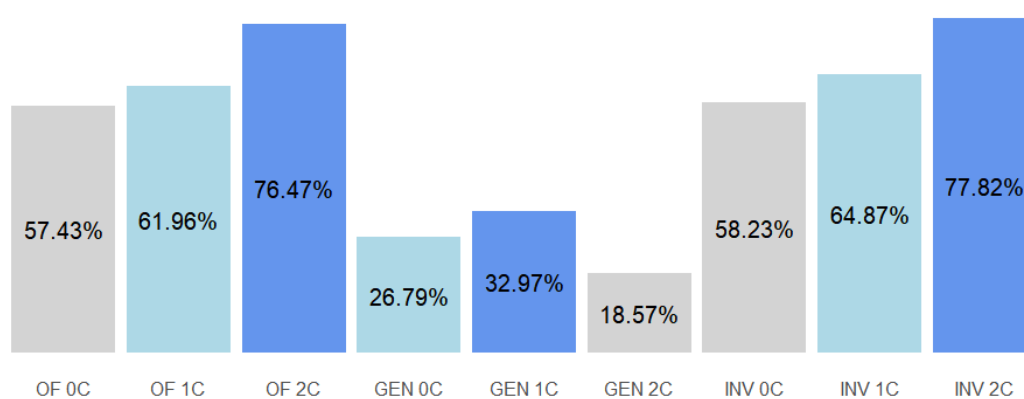
Além disso, ponderamos que o não domínio da sintaxe em L2 se sobrepõe ao domínio lexical, pois segundo Kenedy (2013), a sintaxe é o centro da cognição humana. Também, conforme Bernolet e Hartsuiker (2018), é necessário compreender o significado dos itens lexicais, formar categorias sintáticas, aprender e aplicar o conjunto de regras sintáticas, estabelecendo a correta relação que deve ser estabelecida entre os elementos que formam uma sentença. Todos estes processos cognitivos tornam a aquisição da sintaxe de uma língua um processo complexo e multifacetado. Isto pode significar que mesmo a presença de cognatos, que costumam estar ligados à facilitação de processamento em L2, pode não ser suficiente se a sintaxe é mais difícil de ser processada. Ao processar um item com dois substantivos que são cognatos, o bilíngue pode priorizar a semelhança lexical, em detrimento da sintaxe mais complexa, tornando o processamento mais difícil e incorrendo em mais erros.

Em resumo, nos itens com estruturas sintáticas em que a acurácia é maior, os

cognatos auxiliam na redução do custo de processamento. Isto pode ser comprovado porque nos itens com preposição *of* e com inversão de substantivos, a maior acurácia é para os itens com dois cognatos, que é maior do que para os itens com um cognato, que é maior do que para os itens sem cognatos.

Vejamos em um gráfico de barras a acurácia, correlacionando a sintaxe e a quantidade de cognatos:

Gráfico 9 – Acurácia por condição sintática e cognata do experimento 1



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.956 itens respondidos

IR = número total de itens respondidos

O gráfico de barras mostra que a acurácia se mostrou crescente, na medida em que os cognatos foram surgindo nos sintagmas que utilizaram a preposição *of* e nos sintagmas com inversão de substantivos. Isso indica que a acurácia foi maior em sintagmas nominais com dois cognatos em duas das três construções sintáticas investigadas. Entretanto, os sintagmas com uso de genitivo apresentaram acurácia crescente apenas da condição sem cognatos para a condição com um cognato. A acurácia decresceu em sintagmas com genitivo e dois substantivos cognatos a índices menores do que quando não havia cognatos.

Igualmente, fizemos a análise do tempo de resposta dos itens, considerando a correlação entre construção sintática e quantidade de cognatos:

Tabela 23 – Tempo de resposta (ms) por condição sintática e cognata do experimento 1

TEMPO DE RESPOSTA (MS) POR CONSTRUÇÃO SINTÁTICA E QUANTIDADE DE COGNATOS	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
OF SEM COGNATOS	4.586 (1.866)	4.320	998	9.476
OF COM 1 COGNATO	4.648 (1.656)	4.660	848	9.124
OF COM 2 COGNATOS	3.977 (1.577)	3.614	616	9.846
GENITIVO SEM COGNATOS	4.730 (1.910)	4.695	1.075	9.539
GENITIVO COM 1 COGNATO	4.748 (1.961)	4.594	264	9.877
GENITIVO COM 2 COGNATOS	4.683 (2.114)	4.036	1.207	9.819
INVERSÃO DE SUBSTANTIVOS SEM COGNATOS	4.530 (1.796)	4.375	368	9.850
INVERSÃO DE SUBSTANTIVOS COM 1 COGNATO	4.278 (1.661)	4.116	490	9.389
INVERSÃO DE SUBSTANTIVOS COM 2 COGNATOS	3.784 (1.491)	3.459	303	9.333

Fonte: autoria própria

P = 57 participantes

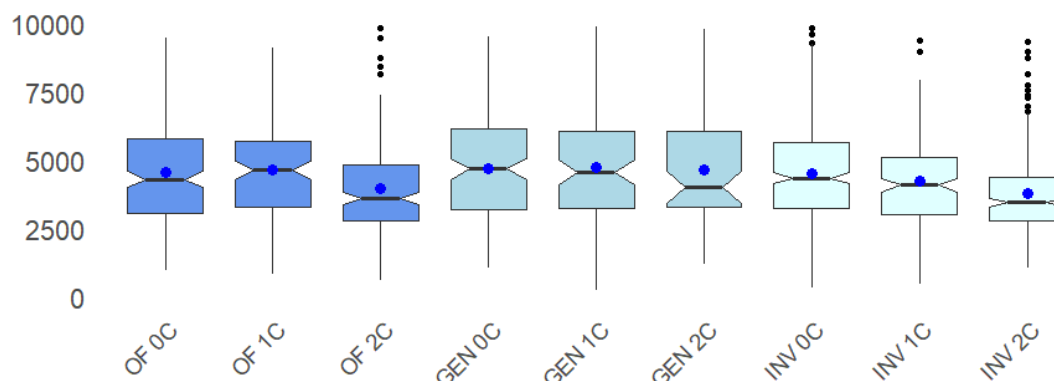
P = número de participantes

IRC = 1.506 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

A tabela mostra que para as construções sintáticas com inversão de substantivos, a presença de cognatos reduziu a média do tempo de resposta de forma gradativa. Nos sintagmas com preposição *of* e uso de genitivo, a presença de cognatos reduziu a média do tempo de resposta apenas quando os itens continham dois cognatos, em comparação com itens sem cognatos. Nestas duas construções sintáticas, a maior média do tempo de resposta foi para itens que tinham um cognato. Vejamos em um gráfico de caixas o tempo de resposta, correlacionando a sintaxe e a quantidade de cognatos:

Gráfico 10 – Tempo de resposta (ms) por condição sintática e cognata do experimento 1



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IRC = 1.506 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

Os gráficos de caixas mostram visualmente que a média do tempo de resposta se mostrou menor quando os sintagmas tinham dois cognatos para itens com uso da preposição *of* e com inversão de substantivos. No caso dos sintagmas com inversão de substantivos, o tempo foi se tornando decrescente, se considerarmos itens sem cognatos, com um cognato e com dois cognatos, nesta ordem. Por outro lado, as duas maiores médias de tempo de resposta foram para os sintagmas com genitivo e um cognato e para os sintagmas com genitivo sem cognatos. Os *outliers* apareceram em maior quantidade para sintagmas com inversão de substantivos e dois cognatos.

Para fazer um resumo, nos itens com estruturas sintáticas em que o tempo de resposta foi menor para itens sem cognatos, essas palavras transparentes auxiliam na redução do custo de processamento. Isto pode ser comprovado porque nos itens com preposição *of* e com inversão de substantivos, o menor tempo de resposta é para os itens com dois cognatos, que é menor do que para os itens com um cognato, que é menor do que para os itens sem cognatos.

5.2.6 Análise inferencial do experimento 1

Após fazer análise descritiva dos dados do experimento 1, fizemos também a análise inferencial, a fim de testar as hipóteses de pesquisa. Primeiramente, fizemos uma regressão logística, a fim de estimar a probabilidade de acerto nos reconhecimentos das traduções, por meio de um Modelo Misto Linear Generalizado (glmer). Tivemos como efeitos

fixos as estruturas sintáticas (SINTAXE), a quantidade de cognatos nos sintagmas (NCOG) e o nível de proficiência em L2 (PROFICIENCIA) dos participantes. Como efeitos aleatórios incluímos os itens (IDITEM) e os participantes (PARTICIPANTE). Utilizamos o seguinte modelo: $glmer(STATUS \sim SINTAXE + NCOG + PROFICIENCIA + (1 | PARTICIPANTE) + (1 | IDITEM), data = T1, family = binomial)$ e os resultados são apresentados na tabela:

Tabela 24 – Modelo Misto Linear Generalizado (glmer) do experimento 1

PREDITORES	ACURÁCIA		
	LOG-ODDS	INTERVALO DE CONFIANÇA	VALOR DE P
INTERCEPTO	-1,88	-2,55 – -1,22	<0,001
SINTAXE INV	2,14	1,60 – 2,67	<0,001
SINTAXE OF	2,02	1,44 – 2,59	<0,001
1 COGNATO	0,31	-0,25 – 0,86	0,280
2 COGNATOS	0,49	-0,06 – 1,04	0,078
PROFICIÊNCIA	0,01	0,00 – 0,02	0,082
Efeitos aleatórios			
σ^2	3,29		
Desvio padrão $\tau_{00 \text{ IDITEM}}$	1,74		
Desvio padrão $\tau_{00 \text{ PARTICIPANTE}}$	0,18		
ICC	0,37		
$N_{\text{PARTICIPANTE}}$	57		
N_{IDITEM}	159		
Observações	2.956		
R^2 Marginal/ R^2 Condicional	0,171/0,476		

Fonte: autoria própria

Nesta tabela, os coeficientes apresentados estão em escala *log-odds*, que são as chances de acurácia em logaritmos. Para uma interpretação mais intuitiva dos dados, utilizamos a função *invlogit*, para transformar *log-odds* em probabilidade. Para isso, consideramos a proficiência média em L2 dos participantes, que foi 36,15 e como intercepto os sintagmas nominais com caso genitivo sem cognatos.

O Modelo Misto Linear Generalizado (glmer) indica que as probabilidades de acerto no experimento 1 foram as seguintes: 18,82% para sintagmas com genitivo sem cognatos; 23,94% para sintagmas com genitivo com um cognato; 27,52% para sintagmas com genitivo e dois cognatos; 66,26% para sintagmas com inversão de substantivos sem cognatos; 72,72% para sintagmas com inversão de substantivos com um cognato; 76,28% para sintagmas

com inversão de substantivos e dois cognatos; 63,53% para sintagmas com preposição *of* sem cognatos; 70,28% para sintagmas com preposição *of* com um cognato; e 74,05% para sintagmas com preposição *of* e dois cognatos.

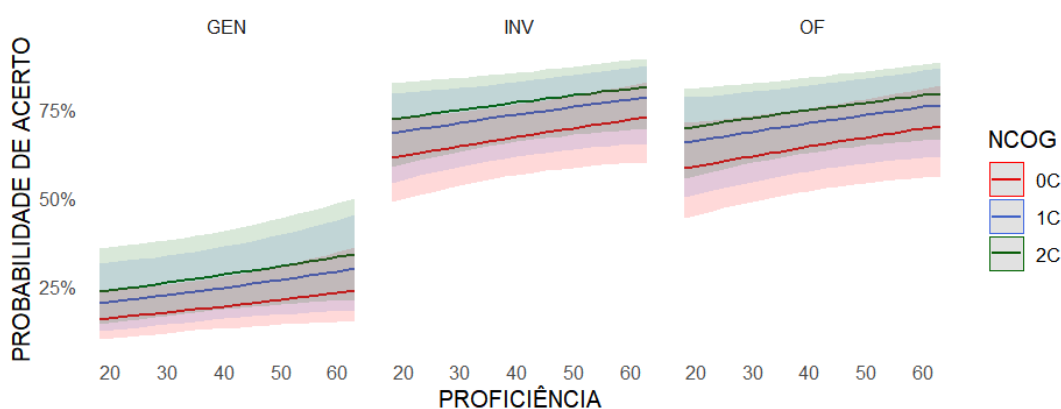
O Modelo Misto Linear Generalizado (glmer) também mostra que a variabilidade do desvio padrão do intercepto por item é 1,74 e que a variabilidade do desvio padrão por participante é 0,18. Utilizando a função *invlogit* mais uma vez, calculamos que o intercepto por item varia entre 36,10% e 39,07% e por participante varia entre 9,04% e 18,75%. Outrossim, o R^2 condicional indica quase 48% da variância em acurácia pode ser explicada pela variância dos efeitos fixos e dos efeitos aleatórios.

Também, o Modelo Misto Linear Generalizado (glmer) sugere diferença significativa, com valor de $p < 0,05$, para a construção sintática com preposição *of* e com inversão de substantivos. Isso significa que uma das variáveis preditoras, que é a sintaxe, exerce influência na acurácia.

É válido destacar que consideramos incluir a idade dos participantes como variável preditora, mesmo não tendo pensado neste fator desde o início da pesquisa, porque tivemos participantes adolescentes (com idade entre 15 e 19 anos) e participantes adultos (com idades que variaram entre 30 e 59 anos). No entanto, decidimos por não levar em consideração este fator, após uma análise prévia, que indicou que a idade não se mostrou relevante nos resultados.

Apresentamos em um gráfico a probabilidade estimada de acerto no experimento 1, que foi uma tarefa de reconhecimento de tradução:

Gráfico 11 - Probabilidade de acerto no experimento 1



Fonte: autoria própria

O gráfico de probabilidade de acerto nos traz muitas informações. Com relação à

sintaxe, indica que a acurácia dos sintagmas com caso genitivo está bem abaixo da acurácia dos sintagmas com preposição *of* e com inversão de substantivos. Outra informação é de que à medida em que a proficiência em L2 vai aumentando, a probabilidade de acurácia também vai crescendo para as três construções sintáticas, porém sem significância estatística. E com relação ao número de cognatos presentes nos sintagmas nominais, é possível perceber que a acurácia dos sintagmas nominais com dois cognatos é maior do que a acurácia dos sintagmas com um cognato, que é maior do que a acurácia dos sintagmas sem cognatos. Porém, as três condições cognatas estão tão próximas que acabam por se sobrepor, o que pode reforçar que a diferença entre as quantidades de cognatos não pode ser considerada significativa.

Igualmente, fizemos uma regressão linear, a fim de prever a probabilidade do tempo de resposta para realizar os reconhecimentos das traduções, por meio de um Modelo Misto Linear (lmer). Tivemos como efeitos fixos as estruturas sintáticas (SINTAXE), as quantidades de cognatos nos sintagmas (NCOG) e o nível de proficiência em L2 (PROFICIENCIA) dos participantes. Como efeitos aleatórios incluímos os itens (IDITEM) e os participantes (PARTICIPANTE). Utilizamos a fórmula: $\text{lmer}(\text{TEMPO} \sim \text{SINTAXE} + \text{NCOG} + \text{PROFICIENCIA} + (1 | \text{PARTICIPANTE}) + (1 | \text{IDITEM}), \text{data} = \text{T1})$ e os resultados são apresentados na tabela:

Tabela 25 – Modelo Misto Linear (lmer) do experimento 1

PREDITORES	TEMPO DE RESPOSTA		
	ESTIMATIVA	INTERVALO DE CONFIANÇA	VALOR DE P
INTERCEPTO	5.514,85	4.613,79 – 6.415,92	<0,001
SINTAXE INV	- 497,64	-749,37 – -245,91	<0,001
SINTAXE OF	- 345,78	-612,65 – -78,91	0,011
1 COGNATO	- 90,35	-323,85 – 143,14	0,448
2 COGNATOS	- 560,57	-788,83 – -332,31	<0,001
PROFICIÊNCIA	- 15,25	-38,49 – 7,98	0,198
Efeitos aleatórios			
σ^2	2.062.156,14		
Desvio padrão $\tau_{00 \text{ IDITEM}}$	340,60		
Desvio padrão $\tau_{00 \text{ PARTICIPANTE}}$	922,30		
ICC	0,32		
$N_{\text{PARTICIPANTE}}$	57		
N_{IDITEM}	156		
Observações	1.506		
$R^2 \text{ Marginal}/R^2 \text{ Condicional}$	0,041/0,347		

Fonte: autoria própria

Nesta análise, os coeficientes apresentados estão em milissegundos. Para calcular a probabilidade do tempo de resposta de todas as condições, consideramos a proficiência média em L2 dos participantes, que foi 36,15 e como intercepto os sintagmas nominais com caso genitivo sem cognatos.

O Modelo Misto Linear (lmer) indica que as probabilidades de tempo de resposta no experimento 1 foram as seguintes: 4.963 milissegundos para sintagmas com genitivo sem cognatos; 4.873 milissegundos para sintagmas com genitivo com um cognato; 4.402 milissegundos para sintagmas com genitivo e dois cognatos; 4.465 milissegundos para sintagmas com inversão de substantivos sem cognatos; 4.375 milissegundos para sintagmas com inversão de substantivos com um cognato; 3.905 milissegundos para sintagmas com inversão de substantivos e dois cognatos; 4.617 milissegundos para sintagmas com preposição *of* sem cognatos; 4.527 milissegundos para sintagmas com preposição *of* com um cognato; e 4.057 milissegundos para sintagmas com preposição *of* e dois cognatos.

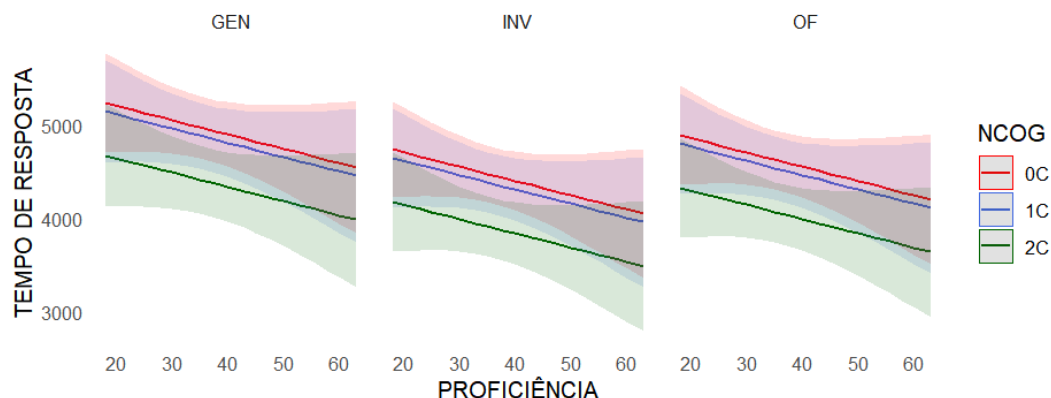
O Modelo Misto Linear (lmer) também mostra que a variabilidade do desvio padrão do intercepto por item é de 340,60 milissegundos e que a variabilidade do desvio padrão por participante é de 922,30 milissegundos. Dessa maneira, calculamos que o intercepto por item varia entre 5.169 e 5.845 milissegundos e por participante varia entre 4.589 e 6.425 milissegundos. Outrossim, o R^2 condicional indica que 34% da variância do tempo de resposta pode ser explicada pela variância dos efeitos fixos e dos efeitos aleatórios.

Também, o Modelo Misto Linear (lmer) sugere diferença significativa, com valor de $p < 0,05$, para a construção sintática com preposição *of* e com inversão de substantivos e para sintagmas com dois cognatos. Isso significa que duas das variáveis preditoras exercem influência no tempo de resposta: a sintaxe e a quantidade de cognatos (quando havia dois substantivos cognatos).

Assim como ocorreu com a acurácia, decidimos por não levar em consideração a idade como uma variável preditora, após uma análise prévia, que indicou que a idade não se mostrou relevante nos resultados.

Apresentamos em um gráfico a probabilidade de tempo de resposta prevista no experimento 1:

Gráfico 12 - Previsão do tempo de resposta do experimento 1



Fonte: autoria própria

O gráfico de previsão de tempo de resposta nos informa que, com relação à sintaxe, o tempo de resposta dos sintagmas com caso genitivo está maior do que o tempo de resposta dos sintagmas com preposição *of* e com inversão de substantivos. Outra informação é de que à medida em que a proficiência em L2 vai aumentando, a previsão do tempo de resposta vai diminuindo para as três construções sintáticas, porém sem significância estatística. E com relação ao número de cognatos presentes nos sintagmas nominais, é possível perceber que o tempo de resposta dos sintagmas nominais com dois cognatos é menor do que o tempo de resposta dos sintagmas com um cognato, que é menor do que o tempo de resposta dos sintagmas sem cognatos. Outrossim, o tempo de resposta para itens com dois cognatos se distancia das outras duas condições cognatas nas três construções sintáticas, o que pode reforçar que a diferença de tempo de resposta para itens com dois cognatos pode ser considerada significativa.

5.3 Experimento 2

O segundo experimento realizado pelos participantes consistia em uma tarefa de julgamento de aceitabilidade, com 105 itens, dentre os quais 60 itens eram experimentais e 45 itens eram distratores. Isto significa que cada um dos 57 participantes tiveram acesso a 105 itens que continham uma sentença em inglês e uma sentença em português, para que julgassem se essas sentenças estavam bem construídas (tecla 1) ou não estavam bem construídas (tecla 2). Ressalte-se que a ordem dos 105 itens foi aleatorizada pelo Psytoolkit (Stoet, 2010; 2017).

O objetivo foi investigar a estimativa da janela temporal para a formação de julgamentos para as três construções sintáticas que possam representar sintagmas nominais

preposicionados com dois substantivos em L2: preposição *of*, caso genitivo e inversão de substantivos. Outrossim, buscamos relação entre os resultados obtidos com o efeito cognato e com o nível de proficiência em L2.

Apresentaremos nesta seção a análise dos itens experimentais. Dessa forma, salientamos que coletamos 5.985 dados, dos quais somente 3.420 dados nos interessam, porque se referem aos itens experimentais.

Após uma análise inicial, decidimos excluir 21 itens de um total de 180 experimentais, assim como ocorreu no experimento 1. Excluimos as sentenças *The life of married is busy*, *The married's life is busy* e *The married life is busy* porque a palavra *married* não pode ser considerada substantivo, apenas adjetivo em L2 e, dessa maneira, não configuram sintagmas nominais formados por dois substantivos. Também, excluimos mais 18 itens, porque não inserimos o artigo *the* em algumas traduções que precisavam desse item lexical, como em *The size of shirt is medium*. Como não há sentença em L1 neste experimento para assegurar a acurácia e seria possível apontar acurácia ou não acurácia nesses itens, achamos mais adequado excluir os itens da análise. Dessa maneira, as três listas de estímulos do experimento 1 ficaram com as seguintes quantidades de itens para análise de resultados: lista 1 → 51 itens, lista 2 → 55 itens; e lista 3 → 53 itens.

Em seguida, fizemos a análise dos registros de tempo de resposta do experimento 2. Dos 3.023 itens experimentais, 18 não foram respondidos dentro do tempo estabelecido de 10.000 milissegundos (10 segundos). Assim, 3.005 itens experimentais foram respondidos pelos participantes no experimento 2.

Além disso, verificamos 12 respostas com registros abaixo de 250 milissegundos, que estão fora dos padrões de tempo mínimo para observação de fenômenos da Psicolinguística do Bilinguismo, segundo Schoonbaert *et al.* (2009) e Duñabeitia *et al.* (2010). Dessa maneira, excluimos as 12 respostas que registraram dados abaixo de 250 milissegundos e a análise levou em consideração 2.993 itens.

Nas subseções a seguir apresentamos o resultado da análise dos dados em seis perspectivas: análise geral, análise por sintaxe, análise por condição cognata, análise por quantidade de cognatos, análise por sintaxe e quantidade de cognatos e análise inferencial. Neste estudo, temos duas variáveis de resposta: a acurácia e o tempo de resposta. Ressaltamos que consideramos todos os itens respondidos para a análise da acurácia. Para a análise do tempo de resposta, levamos em consideração os itens respondidos com acurácia.

5.3.1 Análise geral

Considerando a primeira variável de resposta, analisamos o *status* geral dos itens experimentais respondidos:

Tabela 26 – *Status* geral do experimento 2

STATUS GERAL	ACURÁCIA	NÃO ACURÁCIA	TOTAL
QUANTIDADE	1.552 (51,85%)	1.441 (48,15%)	2.993 (100%)

Fonte: autoria própria

P = 57 participantes

P = número de participantes

VER = 2.993 variáveis experimentais respondidas

VER = número total de variáveis experimentais respondidas

Considerando os 2.993 itens experimentais respondidos pelos participantes, os dados da tabela apontam que a acurácia foi um pouco maior do que a não acurácia. Tivemos 1.552 (51,85%) itens com acurácia e 1.441 (48,15%) itens sem acurácia. Vejamos a acurácia geral em um gráfico de barras:

Gráfico 13 – Acurácia do experimento 2



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.993 itens respondidos

IR = número total de itens respondidos

O gráfico mostra visualmente os dados de acurácia geral dos itens respondidos. Demos destaque aos itens com acurácia, que representaram 51,85% dos itens respondidos.

Com relação à segunda variável de resposta, analisamos o tempo de resposta dos itens respondidos com acurácia:

Tabela 27 – Tempo de resposta do experimento 2

MEDIDAS PADRÕES	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
TEMPO DE RESPOSTA (EM MS)	3.029 (1.692)	2.672	253	9.792

Fonte: autoria própria

P = 57 participantes

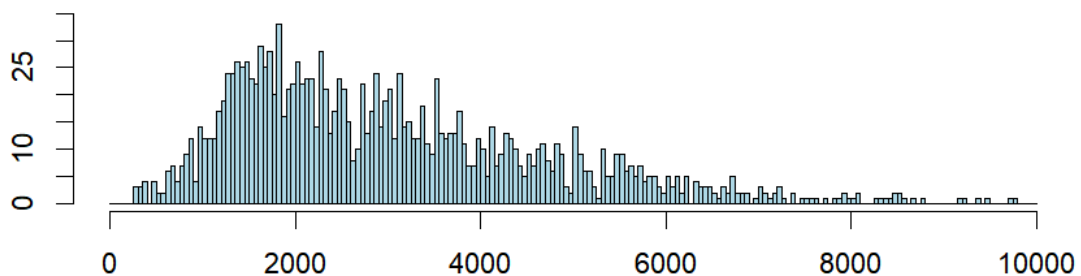
P = número de participantes

IRC = 1.552 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

A tabela mostra que a média do tempo de resposta dos itens experimentais com acurácia foi 3.029 milissegundos, com desvio padrão 1.692 milissegundos. A mediana foi 2.672 milissegundos. O tempo mínimo utilizado para responder um item corretamente foi 253 milissegundos e o tempo máximo foi 9.792 milissegundos. Vejamos em um histograma a distribuição da média de tempo de resposta:

Gráfico 14 – Tempo de resposta (ms) do experimento 2



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IRC = 1.552 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

O histograma aponta que o tempo de resposta geral ficou bem distribuído. A maior parte das respostas foi dada em até 6.000 milissegundos. Há alguns registros acima de 6.000 milissegundos.

5.3.2 Análise por construção sintática

Analisamos, também, o *status* dos itens experimentais respondidos por construção sintática:

Tabela 28 – *Status* por construção sintática do experimento 2

STATUS POR CONSTRUÇÃO SINTÁTICA	ACURÁCIA	NÃO ACURÁCIA	TOTAL
OF	473 (61,03%)	302 (38,97%)	775 (100%)
GENITIVO	346 (31,26%)	761 (68,74%)	1.107 (100%)
INVERSÃO DE SUBSTANTIVOS	733 (65,98%)	378 (34,02%)	1.111 (100%)

Fonte: autoria própria

P = 57 participantes

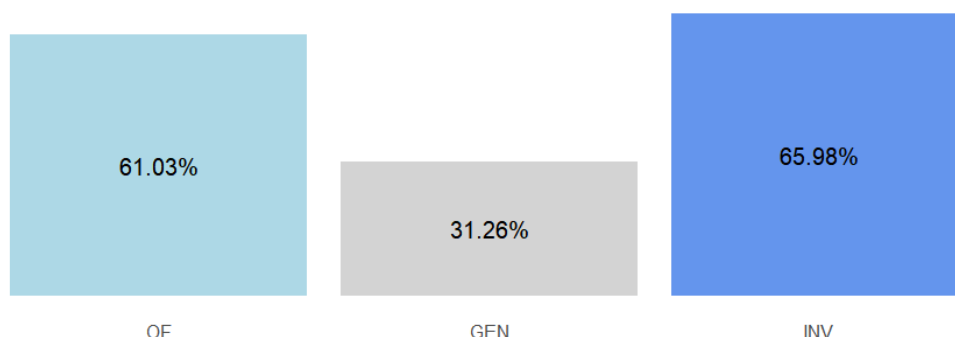
P = número de participantes

IR = 2.993 itens respondidos

IR = número total de itens respondidos

Os dados da tabela apontam que analisamos 775 sintagmas nominais com preposição *of*, 1.107 sintagmas nominais com genitivo e 1.111 sintagmas nominais com inversão de substantivos. Os dados também mostram que a acurácia foi maior do que a não acurácia para duas das três construções sintáticas possíveis para representar os sintagmas nominais preposicionados com dois substantivos: *of* e inversão de substantivos. A acurácia nos sintagmas nominais com preposição *of* foi 473 (61,03%) e nos sintagmas com inversão dos substantivos foi 733 (65,98%). A não acurácia se mostrou maior do que a acurácia apenas nos sintagmas nominais com genitivo, com índice de acurácia de 346 (31,26%). Vejamos a acurácia por construção sintática em um gráfico de barras:

Gráfico 15 – Acurácia por construção sintática do experimento 2



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.993 itens respondidos

IR = número total de itens respondidos

O gráfico mostra visualmente os dados de acurácia por construção sintática. A acurácia foi maior nos sintagmas com inversão de substantivos (65,98%). Mas a acurácia nos sintagmas com preposição *of* (61,03%) ficou muito próxima da maior acurácia. Assim, destacamos a acurácia das duas construções sintáticas.

Vejamos o tempo de resposta dos itens experimentais por construção sintática:

Tabela 29 – Tempo de resposta (ms) por construção sintática do experimento 2

TEMPO DE RESPOSTA (EM MS) POR CONSTRUÇÃO SINTÁTICA	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
OF	2.971 (1.646)	2.525	290	9.151
GENITIVO	3.338 (1.799)	3.019	277	8.667
INVERSÃO DE SUBSTANTIVOS	2.921 (1.655)	2.545	253	9.792

Fonte: autoria própria

P = 57 participantes

P = número de participantes

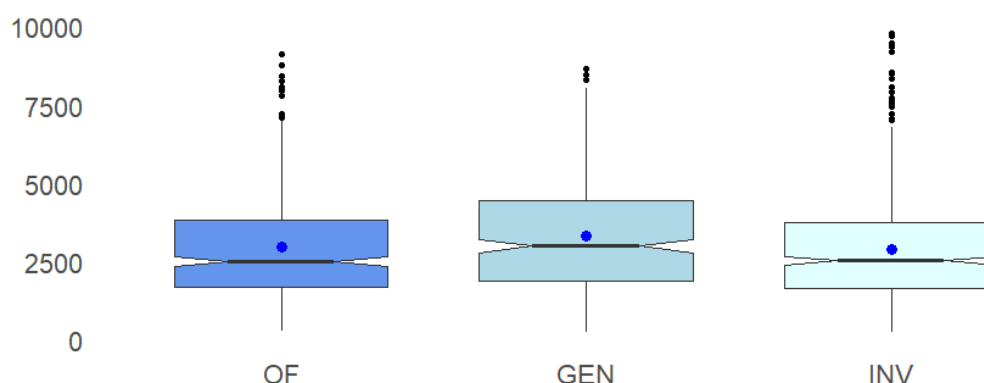
IRC = 1.552 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

A tabela mostra que as medidas padrões ficaram próximas, independentemente da construção sintática. Mesmo com dados muito próximos, a média de tempo se mostrou menor para a inversão dos substantivos: 2.921 milissegundos. O mesmo ocorreu com a mediana: 2.545 milissegundos. O desvio padrão foi menor para sintagmas com preposição *of*: 1.646

milissegundos. Tanto o menor tempo quanto o maior tempo para responder um item com acurácia foi para a inversão dos substantivos: 253 e 9.792 milissegundos, respectivamente. Vejamos em um gráfico de caixas o tempo de resposta por construção sintática:

Gráfico 16 – Tempo de resposta (ms) por construção sintática do experimento 2



Fonte: autoria própria
 P = 57 participantes
 P = número de participantes
 IRC = 1.552 itens respondidos com acurácia
 IRC = número total de itens respondidos com acurácia

Os gráficos de caixas mostram visualmente que tanto a média quanto a mediana ficaram próximas nas três construções sintáticas. No entanto, a média e a mediana foram maiores quando a construção sintática utilizou o caso genitivo. O mesmo aconteceu com o primeiro e o segundo quartis. Além disso, quase não apareceram *outliers* para os sintagmas com o caso genitivo.

5.3.3 Análise por condição cognata

Dando continuidade à análise detalhada dos 2.993 itens experimentais respondidos, investigamos o *status* dos itens experimentais respondidos por condição cognata:

Tabela 30 – *Status* por condição cognata do experimento 2

STATUS POR CONDIÇÃO COGNATA	ACURÁCIA	NÃO ACURÁCIA	TOTAL
SEM COGNATOS	688 (48,66%)	726 (51,34%)	1.414 (100%)
COM COGNATO	864 (54,72%)	715 (45,28%)	1.579 (100%)

Fonte: autoria própria
 P = 57 participantes

P = número de participantes
 IR = 2.993 itens respondidos
 IR = número total de itens respondidos

Os dados da tabela apontam que 1.414 itens não continham cognatos nos sintagmas nominais e 1.579 itens continham cognatos nos sintagmas nominais. Os dados também mostram que a acurácia foi maior do que a não acurácia nos itens com cognatos. A acurácia em itens com cognatos foi 54,72%, enquanto para itens sem cognatos foi 48,66%. Vejamos a acurácia por condição cognata em um gráfico de barras:

Gráfico 17 – Acurácia por condição cognata do experimento 2



Fonte: autoria própria
 P = 57 participantes
 P = número de participantes
 IR = 2.993 itens respondidos
 IR = número total de itens respondidos

O gráfico mostra a acurácia das duas condições cognatas. Damos destaque à acurácia dos itens com cognatos (54,72%), que foi maior do que a acurácia dos itens sem cognatos.

Também fizemos a análise do tempo de resposta, levando em consideração a condição cognata:

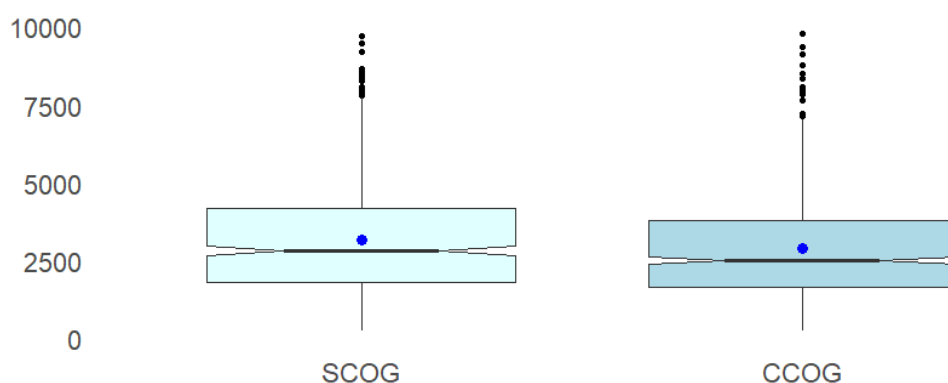
Tabela 31 – Tempo de resposta (ms) por condição cognata do experimento 2

TEMPO DE RESPOSTA (EM MS) POR CONDIÇÃO COGNATA	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
SEM COGNATO	3.181 (1.790)	2.839	253	9.705
COM COGNATO	2.909 (1.601)	2.517	277	9.792

Fonte: autoria própria
 P = 57 participantes
 P = número de participantes
 IRC = 1.552 itens respondidos com acurácia
 IRC = número total de itens respondidos com acurácia

A tabela mostra que as medidas padrões ficaram próximas, independentemente da condição cognata. A média, a mediana e o desvio padrão foram menores quando havia cognatos nos sintagmas, com registros de 2.909, 2.517 e 1.601 milissegundos, respectivamente. O menor tempo para responder um item foi quando não havia cognato no sintagma: 253 milissegundos. O maior tempo foi para itens com cognatos: 9.792 milissegundos. Vejamos em um gráfico de caixas o tempo de resposta por condição cognata:

Gráfico 18 – Tempo de resposta (ms) por condição cognata do experimento 2



Fonte: autoria própria
 P = 57 participantes
 P = número de participantes
 IRC = 1.552 itens respondidos com acurácia
 IRC = número total de itens respondidos com acurácia

Os gráficos de caixas mostram visualmente que as medidas padrões ficaram próximas nas duas condições cognatas. Mas é possível perceber que a média e mediana estão menores para os itens com cognatos. O mesmo ocorreu com o primeiro e o segundo quartis. Os *outliers* apareceram nas duas condições cognatas.

5.3.4 Análise por quantidade de cognatos

Nos itens com cognatos, metade continha um substantivo cognato e metade continha os dois substantivos cognatos. Essa divisão nos permitiu analisar se a quantidade de cognatos nos sintagmas interferiu nos resultados:

Tabela 32 – *Status* por quantidade de cognatos do experimento 2

STATUS POR POR QUANTIDADE DE COGNATOS	ACURÁCIA	NÃO ACURÁCIA	TOTAL
0 COGNATO	688 (48,66%)	726 (51,34%)	1.414 (100%)
1 COGNATO	405 (53,78%)	348 (46,22%)	753 (100%)
2 COGNATOS	459 (55,57%)	367 (44,43%)	826 (100%)

Fonte: autoria própria

P = 57 participantes

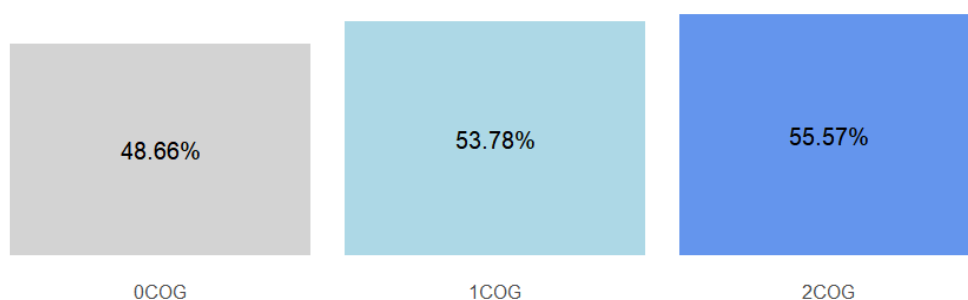
P = número de participantes

IRC = 2.993 itens respondidos

IRC = número total de itens respondidos

A tabela mostra que 1.414 itens sem cognatos foram respondidos, 753 itens com um cognato foram respondidos e 826 itens com dois cognatos foram respondidos. A acurácia foi maior do que a não acurácia em duas das três condições. Comparando as quantidades de cognatos, a acurácia se mostrou crescente, considerando itens sem cognatos, com um cognato e com dois cognatos, nesta ordem. Vejamos os dados da acurácia por quantidade de cognatos em um gráfico de barras:

Gráfico 19 – Acurácia por quantidade de cognatos do experimento 2



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.993 itens respondidos

IR = número total de itens respondidos

O gráfico mostra visualmente a acurácia crescente, à medida em que os cognatos iam surgindo. Destacamos a acurácia dos sintagmas com um cognato (53,78%) e com dois cognatos (55,57%), que foram maiores entre as três possibilidades.

Analisamos, igualmente, o tempo de resposta dos itens por quantidade de cognatos:

Tabela 33 – Tempo de resposta (ms) por quantidade de cognatos do experimento 2

TEMPO DE RESPOSTA (EM MS) POR QUANTIDADE DE COGNATOS	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
0 COGNATO	3.181 (1.790)	2.839	253	9.705
1 COGNATO	2.993 (1.610)	2.626	277	9.151
2 COGNATOS	2.835 (1.591)	2.455	328	9.792

Fonte: autoria própria

P = 57 participantes

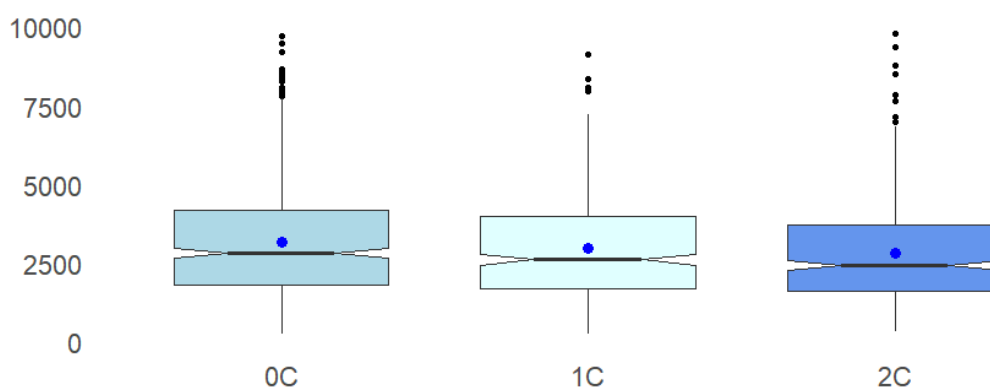
P = número de participantes

IRC = 1.552 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

A tabela mostra que as medidas padrões ficaram próximas, independentemente da quantidade de cognatos. A média, o desvio padrão e a mediana se mostraram decrescentes, considerando itens sem cognatos, com um cognato e com dois cognatos, nesta ordem. O menor tempo para responder um item foi para um sintagma sem cognato (253 milissegundos) e o maior tempo para responder um item foi em um item com dois cognatos (9.792 milissegundos). Vejamos o tempo de resposta por quantidade de cognatos em um gráfico de caixas:

Gráfico 20 – Tempo de resposta por quantidade de cognatos do experimento 2



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IRC = 1.552 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

Os gráficos de caixas mostram visualmente que as medidas de tempo se mostraram decrescentes, considerando itens sem cognatos, com um cognato e com dois cognatos, nesta ordem. Além disso, apareceram menos *outliers* quando os sintagmas tinham um cognato.

5.3.5 Análise por sintaxe e quantidade de cognatos

Após analisar as condições sintáticas e cognatas de forma isolada, resolvemos correlacionar as duas condições, uma vez que para cada construção sintática diferente, havia itens sem cognatos, com um cognato e com dois cognatos em suas estruturas:

Tabela 34 – *Status* por condição sintática e cognata do experimento 2

ACURÁCIA POR CONSTRUÇÃO SINTÁTICA E QUANTIDADE DE COGNATOS	ACURÁCIA	NÃO ACURÁCIA	TOTAL
OF SEM COGNATOS	177 (54,80%)	146 (45,20%)	323 (100%)
OF COM 1 COGNATO	113 (59,79%)	76 (40,21%)	189 (100%)
OF COM 2 COGNATOS	183 (69,58%)	80 (30,42%)	263 (100%)
GENITIVO SEM COGNATOS	177 (32,48%)	368 (67,52%)	545 (100%)
GENITIVO COM 1 COGNATO	116 (41,28%)	165 (58,72%)	281 (100%)
GENITIVO COM 2 COGNATOS	53 (18,86%)	228 (81,14%)	281 (100%)
INVERSÃO DE SUBSTANTIVOS SEM COGNATOS	334 (61,17%)	212 (38,63%)	546 (100%)
INVERSÃO DE SUBSTANTIVOS COM 1 COGNATO	176 (62,19%)	107 (37,81%)	283 (100%)
INVERSÃO DE SUBSTANTIVOS COM 2 COGNATOS	223 (79,08%)	59 (20,92%)	282 (100%)

Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.993 itens respondidos

IR = número total de itens respondidos com cognatos

A tabela mostra que nos sintagmas com preposição *of* e sem cognatos a acurácia foi de 54,80%. A acurácia para esta construção sintática aumentou quando havia um substantivo cognato (59,79%) e aumentou mais ainda quando havia dois substantivos cognatos (69,58%). Nos sintagmas com uso de genitivo e sem cognatos, a acurácia foi de 32,48%. Nesta estrutura sintática, a acurácia aumentou quando havia um substantivo cognato (41,28%). No entanto, de forma surpreendente, a acurácia em sentenças com genitivo e dois substantivos cognatos se mostrou muito baixa (18,86%), o que indica que nesta condição a acurácia foi a menor de todas as condições cognatas. Nos sintagmas com inversão de substantivos e sem cognatos a acurácia

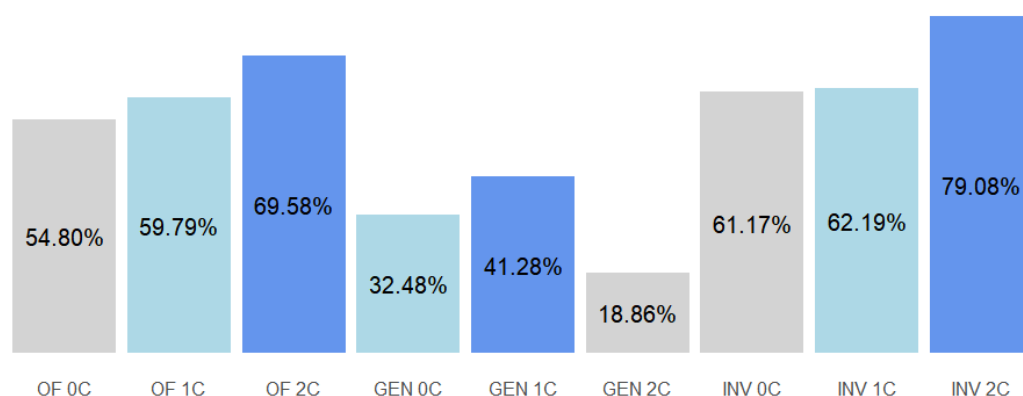
foi de 61,17%. A acurácia para esta construção sintática aumentou quando havia um substantivo cognato (62,19%) e aumentou mais ainda quando havia dois substantivos cognatos (79,08%).

Com relação aos resultados relativos às construções sintáticas com caso genitivo, as considerações são semelhantes às aquelas feitas no experimento 1, tendo em vista que os resultados deste experimento se mostraram muito próximos aos da tarefa de reconhecimento de tradução. Assim, mais uma vez acreditamos que ficou evidente a dificuldade de domínio sintático do caso genitivo. No que diz respeito aos menores índices de acurácia dessa estrutura sintática terem sido identificados em itens com dois cognatos, consideramos que pode ter ocorrido sobrecarga atencional, como postula Costa *et al.* (2006).

Resumindo, nos itens com estruturas sintáticas em que a acurácia é maior, os cognatos auxiliam na redução do custo de processamento. Isto pode ser comprovado porque nos itens com preposição *of* e com inversão de substantivos, a maior acurácia é para os itens com dois cognatos, que é maior do que para os itens com um cognato, que é maior do que para os itens sem cognatos.

Vejam os em um gráfico de barras a acurácia, correlacionando a sintaxe e a quantidade de cognatos:

Gráfico 21 – Acurácia por condição sintática e cognata do experimento 2



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.993 itens respondidos

IR = número total de itens respondidos com cognatos

O gráfico de barras mostra que a acurácia se mostrou crescente, na medida em que os cognatos foram surgindo nos sintagmas que utilizaram a preposição *of* e nos sintagmas com

inversão de substantivos. Isso indica que a acurácia foi maior em sintagmas nominais com dois cognatos em duas das três construções sintáticas investigadas. Entretanto, os sintagmas com uso de genitivo apresentaram acurácia crescente apenas da condição sem cognatos para a condição com um cognato. A acurácia decresceu em sintagmas com genitivo e dois substantivos cognatos a índices menores do que quando não havia cognatos.

Igualmente, fizemos a análise do tempo de resposta dos itens, considerando a correlação entre construção sintática e quantidade de cognatos:

Tabela 35 – Tempo de resposta (ms) por condição sintática e cognata do experimento 2

TEMPO POR CONSTRUÇÃO SINTÁTICA E QUANTIDADE DE COGNATOS	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
OF SEM COGNATOS	3.188 (1.710)	2.817	395	8.444
OF COM 1 COGNATO	3.029 (1.734)	2.607	290	9.151
OF COM 2 COGNATOS	2.720 (1.494)	2.343	328	8.774
GENITIVO SEM COGNATOS	3.512 (1.951)	3.190	316	8.667
GENITIVO COM 1 COGNATO	2.996 (1.533)	2.796	277	6.725
GENITIVO COM 2 COGNATOS	3.507 (1.730)	3.125	830	7.856
INVERSÃO DE SUBSTANTIVOS SEM COGNATOS	3.001 (1.722)	2.731	253	9.705
INVERSÃO DE SUBSTANTIVOS COM 1 COGNATO	2.961 (1.586)	2.515	490	8.357
INVERSÃO DE SUBSTANTIVOS COM 2 COGNATOS	2.769 (1.602)	2.426	342	9.792

Fonte: autoria própria

P = 57 participantes

P = número de participantes

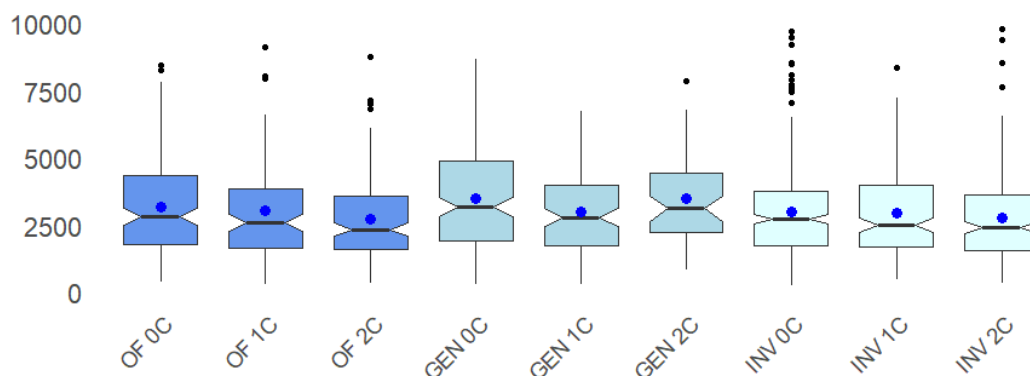
IRC = 1.552 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

A tabela mostra que para as construções sintáticas com preposição *of* e com inversão de substantivos, a presença de cognatos reduziu a média do tempo de resposta de forma gradativa. Isso indica, que o tempo de resposta foi menor quando os sintagmas tinham dois cognatos para estas duas construções sintáticas. Nos sintagmas com caso genitivo, a presença de cognatos reduziu a média do tempo de resposta, mas não de forma gradativa. Nesta construção sintática, a maior média do tempo de resposta foi para itens sem cognatos, mas a

menor média foi para itens com um cognato. Vejamos em um gráfico de caixas o tempo de resposta, correlacionando a sintaxe e a quantidade de cognatos:

Gráfico 22 – Tempo de resposta (ms) por condição sintática e cognata do experimento 2



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IRC = 1.552 itens respondidos com acurácia

IRC = número total de itens respondidos com acurácia

Os gráficos de caixas mostram visualmente que a média do tempo de resposta se mostrou menor quando os sintagmas tinham dois cognatos para itens com uso da preposição *of* e com inversão de substantivos. Nestas duas construções sintáticas, o tempo foi se tornando decrescente, se considerarmos itens sem cognatos, com um cognato e com dois cognatos, nesta ordem. Por outro lado, as duas maiores médias de tempo de resposta foram para os sintagmas com genitivo sem cognatos e para os sintagmas com genitivo e dois cognatos. Os *outliers* apareceram em maior quantidade para sintagmas com inversão de substantivos sem cognatos.

Para resumir, nos itens com estruturas sintáticas em que o tempo de resposta foi menor para itens sem cognatos, essas palavras transparentes auxiliam na redução do custo de processamento. Isto pode ser comprovado porque nos itens com preposição *of* e com inversão de substantivos, o menor tempo de resposta é para os itens com dois cognatos, que é menor do que para os itens com um cognato, que é menor do que para os itens sem cognatos.

5.3.6 Análise inferencial do experimento 2

Após fazer análise descritiva dos dados do experimento 2, fizemos também a análise inferencial, a fim de testar as hipóteses de pesquisa mais uma vez. Primeiramente, fizemos uma regressão logística, a fim de estimar a probabilidade de acerto nos

reconhecimentos das traduções, por meio de um Modelo Misto Linear Generalizado (glmer). Tivemos como efeitos fixos as estruturas sintáticas (SINTAXE), a quantidade de cognatos nos sintagmas (NCOG) e o nível de proficiência em L2 (PROFICIENCIA) dos participantes. Como efeitos aleatórios incluímos os itens (IDITEM) e os participantes (PARTICIPANTE). Utilizamos o seguinte modelo: $glmer(STATUS \sim SINTAXE + NCOG + PROFICIENCIA + (1 | PARTICIPANTE) + (1 | IDITEM), data = T2, family = binomial)$ e os resultados são apresentados na tabela:

Tabela 36 - Modelo Misto Linear Generalizado (glmer) do experimento 2

ACURÁCIA			
PREDITORES	LOG-ODDS	INTERVALO DE CONFIANÇA	VALOR DE P
INTERCEPTO	-1,58	-2,19 – -0,97	<0,001
SINTAXE INV	1,80	1,38 – 2,22	<0,001
SINTAXE OF	1,43	0,98 – 1,89	<0,001
1 COGNATO	0,26	-0,17 – 0,70	0,240
2 COGNATOS	0,30	-0,14 – 0,73	0,179
PROFICIÊNCIA	0,01	0,00 – 0,03	0,041
Efeitos aleatórios			
σ^2	3,29		
Desvio padrão $\tau_{00 \text{ IDITEM}}$	1,00		
Desvio padrão $\tau_{00 \text{ PARTICIPANTE}}$	0,20		
ICC	0,27		
$N_{\text{PARTICIPANTE}}$	57		
N_{IDITEM}	159		
Observações	2.993		
R^2 Marginal/ R^2 Condicional	0,135/0,367		

Fonte: autoria própria

Nesta tabela, os coeficientes apresentados estão em escala *log-odds*, que são as chances de acurácia em logaritmos. Para uma interpretação mais intuitiva dos dados, utilizamos a função *invlogit*, para transformar *log-odds* em probabilidade. Para isso, consideramos a proficiência média em L2 dos participantes, que foi 36,15 e como intercepto os sintagmas nominais com caso genitivo sem cognatos.

O Modelo Misto Linear Generalizado (glmer) indica que as probabilidades de acerto no experimento 2 foram as seguintes: 25,41% para sintagmas com genitivo sem cognatos; 30,68% para sintagmas com genitivo com um cognato; 31,41% para sintagmas com genitivo e dois cognatos; 67,34% para sintagmas com inversão de substantivos sem cognatos;

72,82% para sintagmas com inversão de substantivos com um cognato; 73,48% para sintagmas com inversão de substantivos e dois cognatos; 58,84% para sintagmas com preposição *of* sem cognatos; 65% para sintagmas com preposição *of* com um cognato; e 65,76% para sintagmas com preposição *of* e dois cognatos.

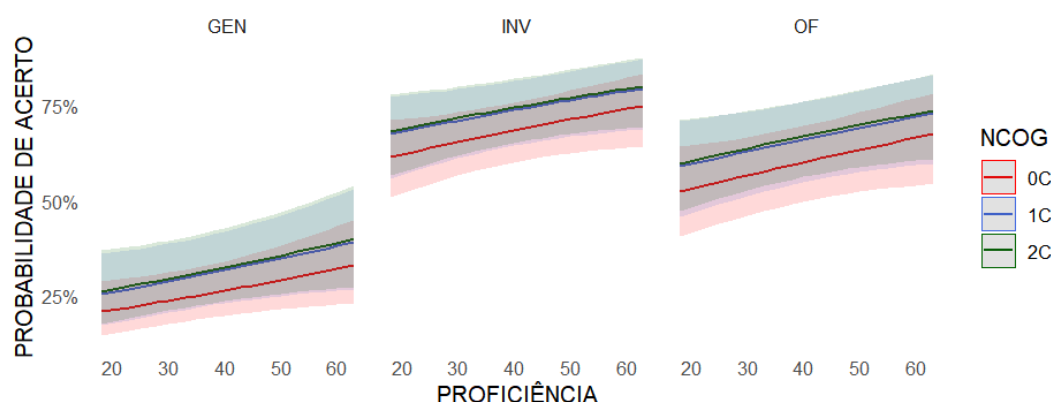
O Modelo Misto Linear Generalizado (glmer) também mostra que a variabilidade do desvio padrão do intercepto por item é 1,00 e que a variabilidade do desvio padrão por participante é 0,20. Utilizando a função *invlogit* mais uma vez, calculamos que o intercepto por item varia entre 6,96% e 35,85% e por participante varia entre 11,54% e 24,26%. Outrossim, o R^2 condicional indica que quase 37% da variância em acurácia pode ser explicada pela variância dos efeitos fixos e dos efeitos aleatórios.

Em resumo, o Modelo Misto Linear Generalizado (glmer) sugere diferença significativa, com valor de $p < 0,05$, para a construção sintática com preposição *of* e com inversão de substantivos e para o nível de proficiência em L2. Isso significa que duas das variáveis preditoras, que são a sintaxe e o nível de proficiência em L2, exercem influência na acurácia.

Assim como ocorreu no experimento 1, decidimos por não levar em consideração a idade como uma variável preditora, após uma análise prévia, que indicou que a idade não se mostrou relevante nos resultados.

Apresentamos em um gráfico a probabilidade de acertos estimada no experimento 2, que foi um julgamento de aceitabilidade:

Gráfico 23 - Probabilidade de acerto no experimento 2



Fonte: autoria própria

O gráfico de probabilidade de acurácia nos traz muitas informações. Com relação à

sintaxe, indica que a acurácia dos sintagmas com caso genitivo está bem abaixo da acurácia dos sintagmas com preposição *of* e com inversão de substantivos. Outra informação é de que à medida em que a proficiência em L2 vai aumentando, a probabilidade de acurácia também vai crescendo para as três construções sintáticas. E com relação ao número de cognatos presentes nos sintagmas nominais, é possível perceber que a acurácia dos sintagmas nominais com dois cognatos é maior do que a acurácia dos sintagmas com um cognato, que é maior do que a acurácia dos sintagmas sem cognatos. Além disso, as duas condições cognatas (sintagmas sem cognatos e sintagmas com um cognato) estão tão próximas que acabam por se sobrepor, o que pode reforçar que a diferença entre os sintagmas com dois cognatos pode ser considerada significativa.

Igualmente, fizemos uma regressão linear, a fim de prever a probabilidade do tempo de resposta para realizar os reconhecimentos das traduções, por meio de um Modelo Misto Linear (lmer). Tivemos como efeitos fixos as estruturas sintáticas (SINTAXE), a quantidade de cognatos nos sintagmas (NCOG) e o nível de proficiência em L2 (PROFICIENCIA) dos participantes. Como efeitos aleatórios incluímos os itens (IDITEM) e os participantes (PARTICIPANTE). Utilizamos a fórmula: $\text{lmer}(\text{TEMPO} \sim \text{SINTAXE} + \text{NCOG} + \text{PROFICIENCIA} + (1 | \text{PARTICIPANTE}) + (1 | \text{IDITEM}), \text{data} = \text{T2})$ e os resultados são apresentados na tabela:

[Tabela 37] - Modelo Misto Linear (lmer) do experimento 2

PREDITORES	TEMPO DE RESPOSTA		
	ESTIMATIVA	INTERVALO DE CONFIANÇA	VALOR DE P
INTERCEPTO	3.542,48	2.625,67 – 4.479,29	<0,001
SINTAXE INV	- 284,15	-505,19 – -63,11	0,012
SINTAXE OF	- 159,39	-400,13 – 81,35	0,194
1 COGNATO	- 245,93	-458,99 – -32,86	0,024
2 COGNATOS	- 354,04	-564,39 – -143,69	0,001
PROFICIÊNCIA	- 5,17	-29,30 – 18,76	0,667
Efeitos aleatórios			
σ^2	1.813.108,61		
Desvio padrão $\tau_{00 \text{ IDITEM}}$	101,76		
Desvio padrão $\tau_{00 \text{ PARTICIPANTE}}$	930,48		
ICC	0,36		
$N_{\text{PARTICIPANTE}}$	57		
N_{IDITEM}	159		
Observações	1.552		

R² Marginal/R² Condicional 0,015/0,372

Fonte: autoria própria

Nesta análise, os coeficientes apresentados estão em milissegundos. Para calcular a probabilidade do tempo de resposta de todas as condições, consideramos a proficiência média em L2 dos participantes, que foi 36,15 e como intercepto os sintagmas nominais com caso genitivo sem cognatos.

O Modelo Misto Linear (lmer) indica que as probabilidades de tempo de resposta no experimento 1 foram as seguintes: 3.361 milissegundos para sintagmas com genitivo sem cognatos; 3.115 milissegundos para sintagmas com genitivo com um cognato; 3.007 milissegundos para sintagmas com genitivo e dois cognatos; 3.077 milissegundos para sintagmas com inversão de substantivos sem cognatos; 2.831 milissegundos para sintagmas com inversão de substantivos com um cognato; 2.723 milissegundos para sintagmas com inversão de substantivos e dois cognatos; 3.202 milissegundos para sintagmas com preposição *of* sem cognatos; 2.956 milissegundos para sintagmas com preposição *of* com um cognato; e 2.848 milissegundos para sintagmas com preposição *of* e dois cognatos.

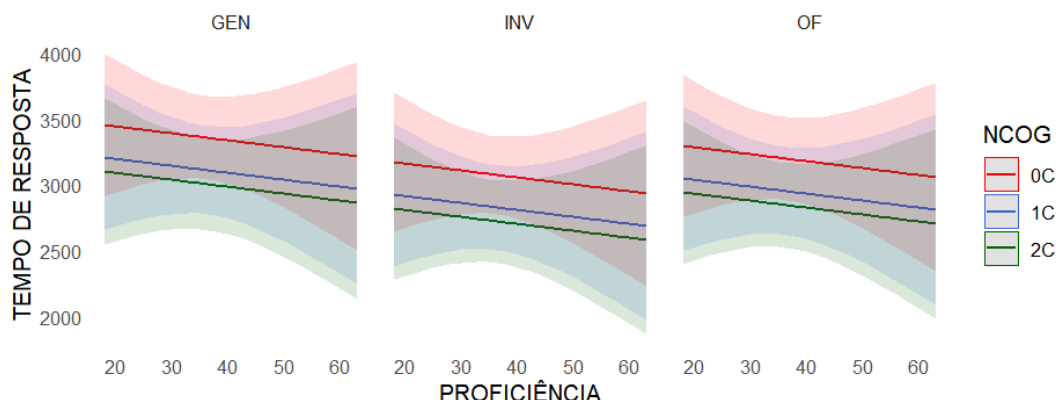
O Modelo Misto Linear (lmer) também mostra que a variabilidade do desvio padrão do intercepto por item é 101 milissegundos e que a variabilidade do desvio padrão por participante é 930 milissegundos. Calculamos que o intercepto por item varia entre 3.226 e 3.862 milissegundos e por participante varia entre 2.573 e 4.515 milissegundos. Outrossim, o R² condicional indica que cerca de 38% da variância em acurácia pode ser explicada pela variância dos efeitos fixos e dos efeitos aleatórios.

Resumidamente, o Modelo Misto Linear (lmer) sugere diferença significativa, com valor de $p < 0,05$, para a construção sintática com inversão de substantivos e para sintagmas com um cognato e com dois cognatos. Isso significa que duas das variáveis preditoras exercem influência no tempo de resposta: na sintaxe com inversão dos substantivos e na condição cognata com um e com dois itens cognatos.

Mais uma vez, decidimos por não levar em consideração a idade como uma variável preditora, após uma análise prévia, que indicou que a idade não se mostrou relevante nos resultados.

Apresentamos em um gráfico a probabilidade de tempo de resposta estimada no experimento 2:

Gráfico 24 - Previsão do tempo de resposta (ms) do experimento 2



Fonte: autoria própria

O gráfico de previsão de tempo de resposta nos informa que, com relação à sintaxe, o tempo de resposta dos sintagmas com inversão de substantivos está menor do que o tempo de resposta dos sintagmas com preposição *of* e com caso genitivo. Outra informação é de que à medida em que a proficiência em L2 vai aumentando, a previsão do tempo de resposta vai diminuindo para as três construções sintáticas. E com relação ao número de cognatos presentes nos sintagmas nominais, é possível perceber que o tempo de resposta dos sintagmas nominais com dois cognatos é menor do que o tempo de resposta dos sintagmas com um cognato, que é menor do que o tempo de resposta dos sintagmas sem cognatos. Outrossim, o tempo de resposta para itens com um cognato e com dois cognatos se distancia da condição sem cognatos nas três construções sintáticas, o que pode reforçar que a diferença entre sintagmas com um cognato e com dois cognatos pode ser considerada significativa.

5.4 Experimento 3

O terceiro experimento realizado pelos participantes consistiu em uma tarefa de tradução, com 80 itens, dentre os quais 45 itens eram experimentais e 35 itens eram distratores. Isto significa que cada um dos 57 participantes tiveram acesso a 80 itens que continham uma sentença em português a ser traduzida e três opções de resposta em inglês. Assim, eles deveriam escolher qual sentença seria a tradução mais adequada para a frase em L1, clicando na tecla 1, 2 ou 3, de acordo com sua opção de resposta. Ressalte-se que a ordem dos 80 itens foi aleatorizada pelo Psytoolkit (Stoet, 2010; 2017).

O objetivo foi investigar a preferência sintática e a estimativa da janela temporal para a seleção da tradução de preferência, levando-se em consideração as três construções sintáticas

que podem representar os sintagmas nominais preposicionados com dois substantivos em L2 (preposição *of*, caso genitivo e inversão de substantivos) que estavam contempladas em cada uma das opções. Outrossim, buscamos relação entre os resultados obtidos com o efeito cognato e com o nível de proficiência em L2.

Apresentaremos nesta seção a análise dos itens experimentais. Dessa forma, salientamos que coletamos 4.560 dados, dos quais nos interessam os 2.565 dados que se referem aos itens experimentais.

Em seguida, analisamos os dados que se referem aos registros do tempo de resposta do experimento 3. Verificamos que apenas 3 itens não foram respondidos dentro do tempo previsto, que era 30.000 milissegundos (30 segundos). Dessa maneira, 2.562 itens foram respondidos pelos participantes no experimento 3.

Do mesmo modo, identificamos 9 respostas com registros abaixo de 250 milissegundos, que estão fora dos padrões de tempo mínimo para observação de fenômenos da Psicolinguística do Bilinguismo, segundo Schoonbaert *et al.* (2009) e Duñabeitia *et al.* (2010). Dessa maneira, excluímos as 9 respostas que registraram dados abaixo de 250 milissegundos e a análise feita levou em consideração 2.553 itens.

Nas subseções a seguir apresentamos o resultado da análise dos dados em quatro perspectivas: análise geral, análise por seleção sintática, análise por seleção sintática e quantidade de cognatos e análise inferencial. Neste experimento, temos duas variáveis de resposta: a seleção sintática e o tempo de resposta. Ressaltamos que consideramos todos os itens respondidos para a análise da seleção sintática. Para a análise do tempo de resposta, também levamos em consideração todos os itens experimentais respondidos, uma vez que neste experimento não analisamos a acurácia, apenas a preferência por cada construção sintática.

5.4.1 Análise geral

Como neste experimento não temos análise de acurácia, iniciamos com a análise do tempo de resposta, de um modo geral. Vejamos os dados na tabela:

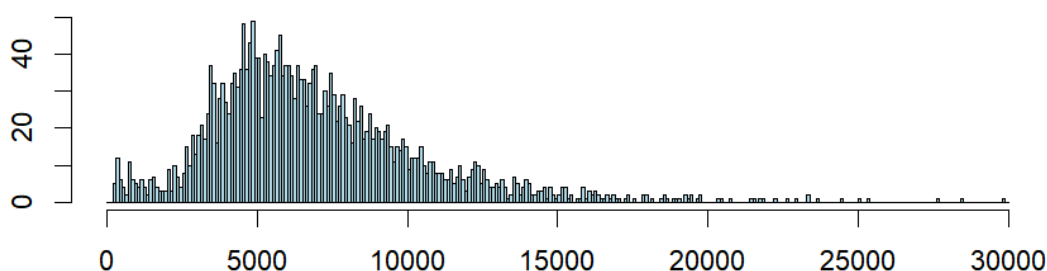
Tabela 38 - Tempo de resposta (ms) do experimento 3

MEDIDAS PADRÕES	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
TEMPO DE RESPOSTA (EM MS)	6.933 (3.586)	6.271	270	29.884

Fonte: autoria própria
 P = 57 participantes
 P = número de participantes
 IR = 2.553 itens respondidos
 IR = número total de itens respondidos

A tabela mostra que a média do tempo de resposta dos itens foi 6.933 milissegundos, com desvio padrão 3.586 milissegundos. A mediana foi 6.271 milissegundos. O tempo mínimo utilizado para fazer a seleção sintática foi 270 milissegundos e o tempo máximo foi 29.884 milissegundos. Vejamos em um histograma a distribuição da média de tempo de resposta:

Gráfico 25 - Tempo de resposta (ms) do experimento 3



Fonte: autoria própria
 P = 57 participantes
 P = número de participantes
 IR = 2.553 itens respondidos
 IR = número total de itens respondidos

O histograma aponta que o tempo de resposta geral se concentrou abaixo de 20.000 milissegundos. A maior parte das respostas foi dada em até 10.000 milissegundos. Mas há alguns registros acima de 20.000 milissegundos.

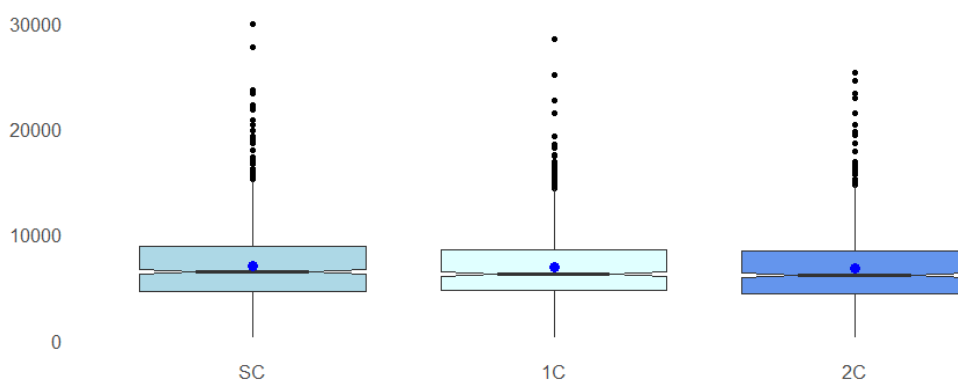
Os itens do experimento 3 tinham três condições cognatas diferentes: itens sem cognatos, itens com um cognato e itens com dois cognatos. Vejamos em um gráfico de caixas o tempo de resposta por condição cognata:

Tabela 39 - Tempo de resposta (ms) por condição cognata do experimento 3

TEMPO DE RESPOSTA (EM MS) POR CONDIÇÃO COGNATA	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
SEM COGNATO	7.088 (3.740)	6.462	285	29.884
1 COGNATO	6.909 (3.419)	6.287	270	28.432
2 COGNATOS	6.800 (3.590)	6.126	272	25.314

A tabela mostra que as medidas padrões ficaram muito próximas, independentemente da condição cognata. No entanto, é possível perceber que a média e a mediana do tempo de resposta foram menores quando havia dois cognatos nos sintagmas. Esses índices se mostraram decrescentes na medida em que os cognatos foram aparecendo. O menor tempo utilizado para fazer a seleção sintática foi para um item com um cognato (270 ms) e o maior tempo foi para um item sem cognatos (29.884). Vejamos em um gráfico de caixas o tempo de resposta por condição cognata:

Gráfico 26 - Tempo de resposta (ms) por condição cognata do experimento 3



Fonte: autoria própria

O gráfico de caixas mostra visualmente que as medidas padrões ficaram muito próximas nas três condições cognatas. Mas é possível perceber que a média e mediana se tornam decrescentes na medida em que os cognatos aparecem. O mesmo ocorreu com o primeiro e o segundo quartis. Os *outliers* apareceram quantidade razoável para todas as condições cognatas dos sintagmas nominais.

5.4.2 Análise por seleção sintática

Analizamos os 2.553 itens experimentais respondidos, considerando a seleção sintática dos itens experimentais respondidos:

Tabela 40 - Seleção sintática do experimento 3

SELEÇÃO SINTÁTICA	OF	GENITIVO	INVERSÃO DE SUBSTANTIVOS	TOTAL
ITENS	525 (20,56%)	975 (38,19%)	1.053 (41,25%)	2.553 (100%)

Fonte: autoria própria

P = 57 participantes

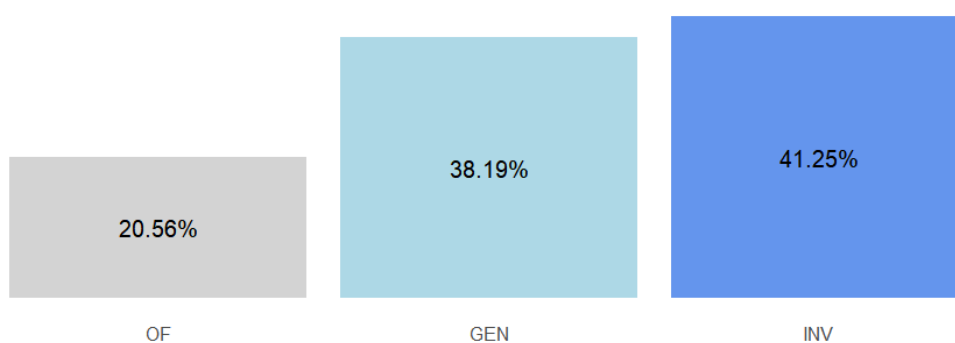
P = número de participantes

IR = 2.553 itens respondidos

IR = número total de itens respondidos

Os dados da tabela apontam que em 525 (20,56%) itens os participantes optaram pelo uso da preposição *of* como estrutura mais adequada para representar a tradução dos sintagmas nominais preposicionados com dois substantivos, em 975 (38,19%) itens a opção foi pelo caso genitivo e em 1.053 (41,25%) itens a escolha foi pela inversão dos substantivos. Vejamos a seleção sintática em um gráfico de barras:

Gráfico 27 - Seleção sintática do experimento 3



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.553 itens respondidos

IR = número total de itens respondidos

O gráfico mostra que a opção sintática mais escolhida para representar os sintagmas nominais em L2 foi a inversão de substantivos (41,25%), seguida do caso genitivo (38,19%). A opção menos escolhida dentre as três possibilidades sintáticas foi o uso da preposição *of*

(20,56%). Destacamos no gráfico as duas construções mais escolhidas pelos participantes.

Também fizemos a análise do tempo de resposta, levando em consideração a seleção sintática dos participantes:

Tabela 41 - Tempo de resposta (ms) por seleção sintática do experimento 3

TEMPO DE RESPOSTA (EM MS) POR SELEÇÃO SINTÁTICA	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
OF	7.396 (3.944)	6.867	270	28.432
GENITIVO	6.415 (3.266)	5.800	272	23.309
INVERSÃO DOS SUBSTANTIVOS	7.181 (3.627)	6.498	280	29.884

Fonte: autoria própria

P = 57 participantes

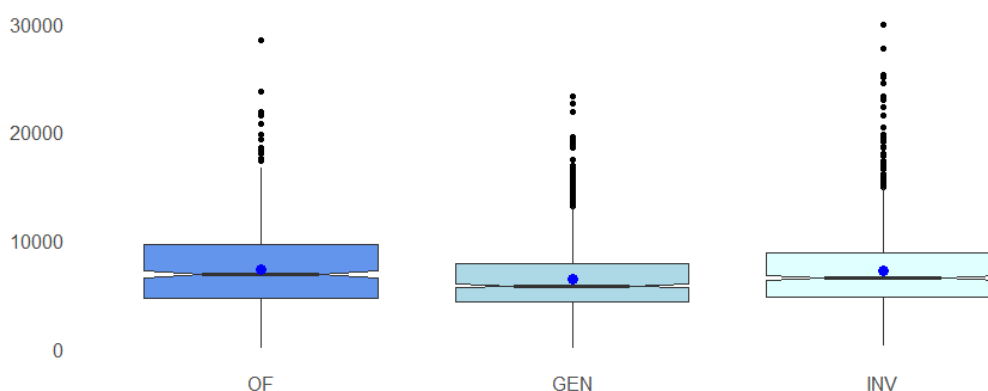
P = número de participantes

IR = 2.553 itens respondidos

IR = número total de itens respondidos

A tabela mostra que medidas padrões ficaram muito próximas, independentemente da seleção sintática. No entanto, é possível perceber que a média, o desvio padrão e a mediana do tempo de resposta foram menores quando a seleção sintática foi o caso genitivo. Vejamos em um gráfico de caixas o tempo de resposta por seleção sintática:

Gráfico 28 - Tempo de resposta (ms) por seleção sintática do experimento 3



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.553 itens respondidos

IR = número total de itens respondidos

O gráfico de caixas mostra visualmente que as medidas padrões ficaram muito

próximas nas três seleções sintáticas. Mas é possível perceber que a média e mediana estão um pouco menores quando os participantes selecionaram o caso genitivo. O mesmo ocorreu com o primeiro e o segundo quartis. Os *outliers* apareceram em menor quantidade quando os participantes selecionaram como resposta itens com a preposição *of*.

5.4.3 Análise por seleção sintática e quantidade de cognatos

Neste experimento, tínhamos 45 itens experimentais, dentre os quais 15 não tinham cognatos, 15 tinham um substantivo cognato e 15 tinham dois substantivos cognatos. Vejamos as seleções sintáticas dos 2.553 itens experimentais respondidos pelos participantes, considerando as condições cognatas:

Tabela 42 - Seleção sintática por condição cognata do experimento 3

SELEÇÃO SINTÁTICA POR CONDIÇÃO COGNATA	OF	GENITIVO	INVERSÃO DE SUBSTANTIVOS	TOTAL
SEM COGNATOS	192 (22,51%)	309 (36,23%)	352 (41,27%)	853 (100%)
1 COGNATO	179 (21,01%)	352 (41,31%)	321 (37,68%)	852 (100%)
2 COGNATOS	154 (18,16%)	314 (37,03%)	380 (44,81%)	848 (100%)

Fonte: autoria própria

P = 57 participantes

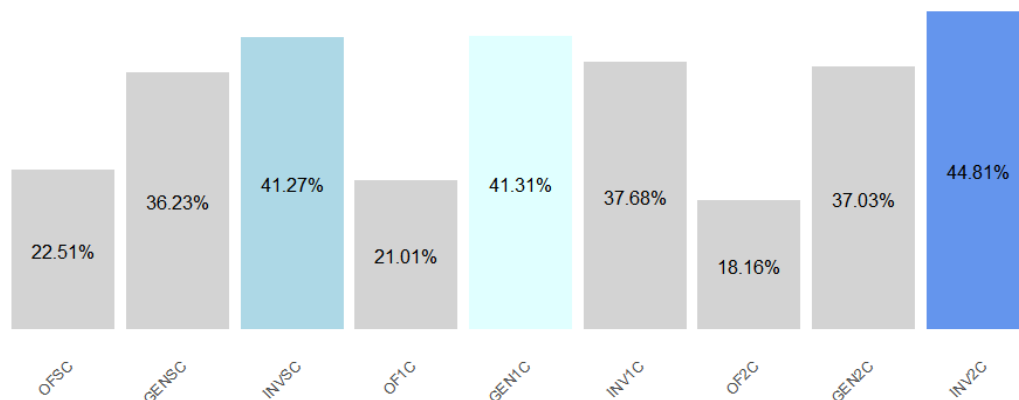
P = número de participantes

IR = 2.553 itens respondidos

IR = número total de itens respondidos

Os dados da tabela apontam que nos itens sem cognatos e com dois cognatos, a inversão de substantivos foi selecionada mais vezes, com 41,27% e 44,81% das seleções, respectivamente. Nos itens com um cognato, a maioria das opções foi pelo caso genitivo, com 41,31% das respostas. Vejamos as seleções sintáticas por condição cognata em um gráfico de barras:

Gráfico 29 - Seleção sintática por condição cognata do experimento 3



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.553 itens respondidos

IR = número total de itens respondidos

O gráfico mostra visualmente que os participantes escolheram mais vezes o caso genitivo quando os itens tinham apenas um cognato. E escolheram mais vezes a inversão dos substantivos quando os itens não tinham cognatos e quando os itens tinham dois cognatos.

Fizemos a análise do tempo de resposta por seleção sintática e quantidade de cognatos nos itens:

Tabela 43 - Tempo de resposta (ms) por seleção sintática e cognata do experimento 3

TEMPO DE RESPOSTA (MS) POR SELEÇÃO SINTÁTICA E CONDIÇÃO COGNATA	MÉDIA (DESVIO PADRÃO)	MEDIANA	MÍNIMO	MÁXIMO
OF SEM COGNATOS	7.712 (4.059)	7.301	377	23.686
GENITIVO SEM COGNATOS	6.326 (3.227)	5.658	365	21.868
INVERSÃO DE SUBSTANTIVOS SEM COGNATOS	7.416 (3.876)	6.736	285	29.884
OF 1 COGNATO	7.325 (3.817)	6.859	270	28.432
GENITIVO 1 COGNATO	6.578 (3.192)	5.940	311	22.631
INVERSÃO DE SUBSTANTIVOS 1 COGNATO	7.091 (3.406)	6.464	342	25.095
OF 2 COGNATOS	7.181 (3.943)	6.289	285	21.511
GENITIVO 2 COGNATOS	6.321 (3.388)	5.793	272	23.309

INVERSÃO DE SUBSTANTIVOS 2 COGNATOS	7.040 (3.569)	6.321	280	25.314
--	------------------	-------	-----	--------

Fonte: autoria própria

P = 57 participantes

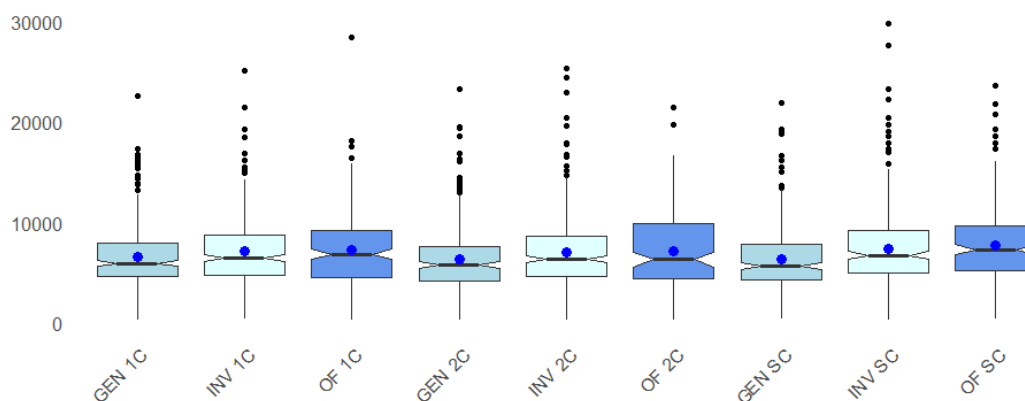
P = número de participantes

IR = 2.553 itens respondidos

IR = número total de itens respondidos

Os dados da tabela apontam que, dentre as nove correlações possíveis, a menor média de tempo de resposta foi quando os participantes selecionaram o caso genitivo e os itens tinham dois cognatos, que registrou 6.321 milissegundos e a maior média de tempo de resposta foi quando os participantes selecionaram o item com preposição *of* e os itens não tinham cognatos, que registrou 7.712 milissegundos. O menor tempo de resposta utilizado foi 270 milissegundos, quando um participante selecionou a preposição *of* para um item que tinha um cognato. O maior tempo de resposta utilizado para fazer a seleção sintática foi 29.884 milissegundos quando um participante selecionou a inversão de substantivos para um item não tinha cognatos. Vejamos o tempo de resposta considerando as seleções sintáticas por condição cognata em um gráfico de barras:

Gráfico 30 - Tempo de resposta (ms) por seleção sintática e cognata do experimento 3



Fonte: autoria própria

P = 57 participantes

P = número de participantes

IR = 2.553 itens respondidos

IR = número total de itens respondidos

O gráfico mostra que as medidas padrões ficaram bem próximas, independentemente da seleção sintática e do número de cognatos. É possível perceber que os registros de tempo foram menores quando os participantes selecionaram o caso genitivo nas três condições cognatas. Os *outliers* apareceram menos vezes quando a seleção sintática foi a

preposição *of* e os itens tinham dois cognatos.

5.4.4 Análise inferencial do experimento 3

Após fazer análise descritiva dos dados do experimento 3, fizemos também a análise inferencial, a fim de testar as hipóteses de pesquisa. Primeiramente, fizemos uma regressão logística, a fim de estimar a probabilidade das seleções sintáticas para a representação dos sintagmas nominais preposicionados com dois substantivos, por meio de um Modelo Logístico Multinomial (*polytomous*), porque tínhamos como opções de tradução três construções sintáticas: preposição *of*, caso genitivo e inversão de substantivos. Levshina (2015) esclarece que o Modelo Logístico Multinomial (*polytomous*) compara cada uma das estruturas sintáticas com a soma das outras duas. Usamos como efeitos fixos a quantidade de cognatos nos sintagmas (CONDIÇÃO) e o nível de proficiência em L2 (PROFICIENCIA) dos participantes. Não incluímos efeitos aleatórios porque este modelo não suporta este tipo de efeito. Utilizamos o seguinte modelo: $\text{polytomous}(\text{ESCOLHA} \sim \text{CONDICAO} + \text{PROFICIENCIA}, \text{data} = \text{T3})$ e os resultados são apresentados:

Tabela 44 - Modelo Logístico Multinomial (*polytomous*) do experimento 3

LOG-ODDS	SELEÇÃO SINTÁTICA		
	OF	GENITIVO	INVERSÃO DE SUBSTANTIVOS
INTERCEPTO	0,4129	-1,135	-0,8324
1 COGNATO	(-0,09132)	0,216	(-0,1512)
2 COGNATOS	-0,2783	(0,03456)	(0,1452)
PROFICIÊNCIA	(-0,04754)	0,01565	0,01321
Observações	2.553		
Probabilidade R ²	0,01956		

Fonte: autoria própria

Nesta tabela, os coeficientes apresentados estão em escala *log-odds*, que são as chances de seleção sintática em logaritmos. Para uma interpretação mais intuitiva dos dados, utilizamos a função *invlogit*, para transformar *log-odds* em probabilidade. Para isso, consideramos a proficiência média em L2 dos participantes, que foi 36,15.

Dessa maneira, o Modelo Logístico Multinomial (*polytomous*) indica que as probabilidades de seleção sintática no experimento 3, de acordo com o número de cognatos

presentes nos sintagmas são as seguintes: em sintagmas sem cognatos, 21,32% de chance de selecionar a opção com preposição *of*, 36,14% de selecionar a opção com genitivo e 41,22% de selecionar a opção com inversão de substantivos; em sintagmas com um cognato, 19,82% de chance de selecionar a opção com preposição *of*, 41,25% de selecionar a opção com genitivo e 37,61% de selecionar a opção com inversão de substantivos; e em sintagmas com dois cognatos, 17,02% de chance de selecionar a opção com preposição *of*, 36,94% de selecionar a opção com genitivo e 44,77% de selecionar a opção com inversão de substantivos.

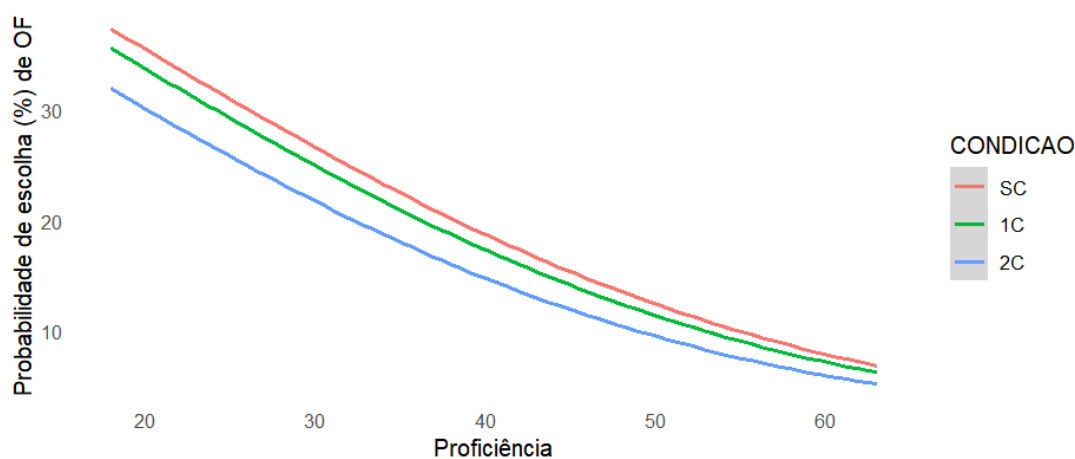
Outrossim, a probabilidade do R^2 aponta que cerca de 2% da variância em seleção sintática pode ser explicada pela variância dos efeitos fixos e dos efeitos aleatórios. Isso significa que a condição cognata e o nível de proficiência em L2 podem explicar 2% das seleções sintáticas.

Também, o Modelo Logístico Multinomial (*polytomous*) sugere diferença significativa, com valor que está fora dos parênteses, para o nível de proficiência em L2 nas três seleções sintáticas e de quantidade de cognatos quando houve seleção da preposição *of* (dois cognatos) e do caso genitivo (um cognato). Isso significa que as duas variáveis preditoras, o nível de proficiência em L2 e a quantidade de cognatos, podem exercer influência na seleção sintática.

Assim como ocorreu nos experimentos 1 e 2, decidimos por não levar em consideração a idade como uma variável preditora, após uma análise prévia, que indicou que a idade não se mostrou relevante nos resultados.

Tendo como base o Modelo Logístico Multinomial (*polytomous*), apresentamos em um gráfico a probabilidade de seleção da estrutura sintática com *of*:

Gráfico 31 - Probabilidade de seleção da preposição *of* no experimento 3

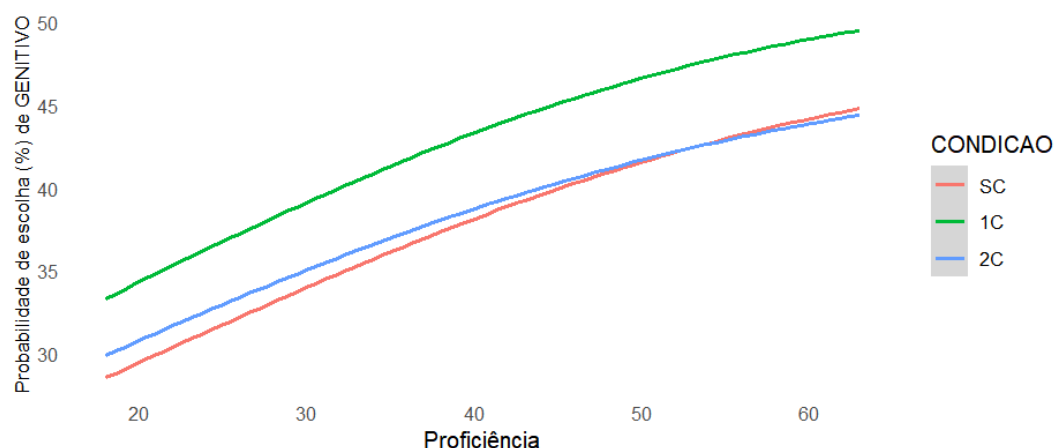


Fonte: autoria própria

O gráfico de probabilidade de seleção sintática da preposição *of*, baseado no Modelo Logístico Multinomial (*polytomous*), nos informa que à medida em que a proficiência em L2 vai aumentando, a probabilidade de seleção por esta construção sintática vai diminuindo. Com relação ao número de cognatos presentes nos sintagmas nominais, é possível perceber que a probabilidade de selecionar os sintagmas com preposição *of* é maior para itens sem cognatos, que é maior para itens com um cognato, que é maior para itens com dois cognatos. Ou seja, quanto maior o número de cognatos, menor a chance da opção com preposição *of* ser selecionada. Além disso, as três condições cognatas estão próximas, mas a condição com dois cognatos está um pouco mais distante das condições com um cognato e sem cognatos, o que pode sugerir a diferença significativa para a condição com dois cognatos, como mostrado na tabela 44.

Apresentamos em um gráfico a probabilidade de seleção da estrutura sintática com caso genitivo, tendo como base o Modelo Logístico Multinomial (*polytomous*):

Gráfico 32 - Probabilidade de seleção do caso genitivo no experimento 3



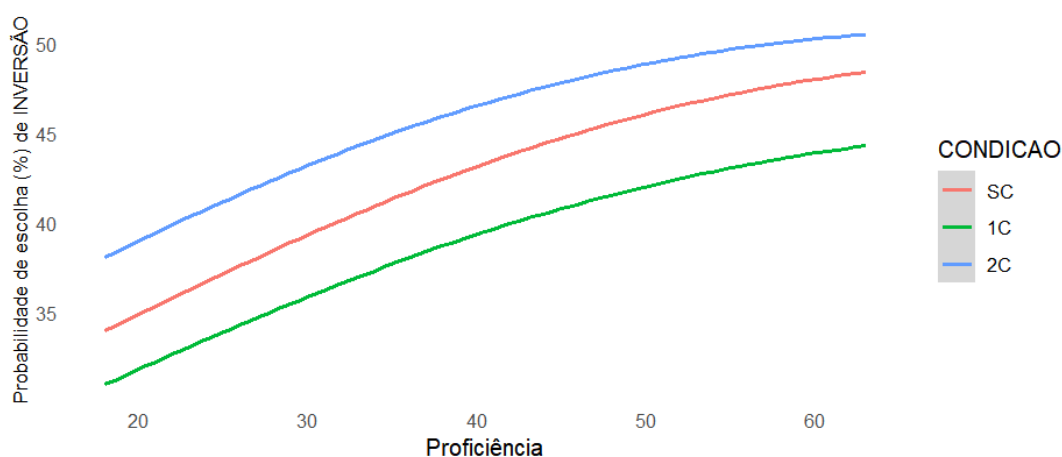
Fonte: autoria própria

O gráfico de probabilidade de seleção sintática do caso genitivo, baseado no Modelo Logístico Multinomial (*polytomous*), nos informa que à medida em que a proficiência em L2 vai aumentando, a probabilidade de seleção por esta construção sintática também cresce. Como a curvatura da proficiência está acentuada, isto pode sugerir a diferença significativa para o nível de proficiência em L2, como mostrado na tabela 44. Com relação ao número de cognatos presentes nos sintagmas nominais, é possível perceber que a probabilidade de selecionar os sintagmas com caso genitivo é maior para itens com um cognato, que é maior para itens com dois cognatos, que é maior para itens sem cognatos, quando o nível de proficiência em L2 é até

cerca de 55. Quando o nível de proficiência em L2 é maior do que aproximadamente 55, a probabilidade de selecionar os sintagmas com caso genitivo é maior para itens com um cognato, que é maior para itens sem cognatos, que é maior para itens com dois cognatos. Ou seja, quanto maior o número de cognatos, menor a chance de a opção com caso genitivo ser selecionada, para bilíngues com nível de proficiência em L2 mais alto. Além disso, as três condições cognatas estão próximas, mas a condição com um cognato está um pouco mais distante das condições com dois cognatos e sem cognatos, o que pode sugerir a diferença significativa para a condição com um cognato, como mostrado na tabela 44.

Igualmente, apresentamos em um gráfico a probabilidade de seleção da estrutura sintática com inversão de substantivos, tendo como base o Modelo Logístico Multinomial (*polytomous*):

Gráfico 33 - Probabilidade de seleção da inversão de substantivos no experimento 3



Fonte: autoria própria

O gráfico de probabilidade de seleção sintática da inversão de substantivos, baseado no Modelo Logístico Multinomial (*polytomous*), nos informa que à medida em que a proficiência em L2 vai aumentando, a probabilidade de seleção por esta construção sintática também cresce. Como a curvatura da proficiência está acentuada, isto pode sugerir a diferença significativa para o nível de proficiência em L2, como mostrado na tabela 44. Com relação ao número de cognatos presentes nos sintagmas nominais, é possível perceber que a probabilidade de selecionar os sintagmas com inversão de substantivos é maior para itens com dois cognatos, que é maior do que para itens sem cognatos, que é maior do que para itens com um cognato. Ou seja, quando os itens possuem dois cognatos, maior a chance da opção com inversão de substantivos ser selecionada. Além disso, as três condições cognatas estão próximas, o que pode

sugerir que não há diferença significativa para a condição cognata, como também foi mostrado na tabela 44.

Também, fizemos uma regressão linear, a fim de prever a probabilidade do tempo de resposta para realizar seleções sintáticas, por meio de um Modelo Misto Linear (lmer). Tivemos como efeitos fixos as condições cognatas (CONDICAO) e o nível de proficiência em L2 (PROFICIENCIA) dos participantes. Como efeitos aleatórios incluímos os itens (IDITEM) e os participantes (PARTICIPANTE). Utilizamos a fórmula: $\text{lmer}(\text{TEMPO} \sim \text{CONDICAO} + \text{PROFICIENCIA} + (1 | \text{PARTICIPANTE}) + (1 | \text{IDITEM}), \text{data} = \text{T3})$ e os resultados são apresentados na tabela:

Tabela 45 - Modelo Misto Linear (lmer) do experimento 3

PREDITORES	TEMPO DE RESPOSTA		
	ESTIMATIVA	INTERVALO DE CONFIANÇA	VALOR DE P
INTERCEPTO	8.870,81	7.000,84 – 10.740,78	<0,001
1 COGNATO	-176,42	-765,71 – 412,86	0,557
2 COGNATOS	-282,55	-871,98 – 306,89	0,347
PROFICIÊNCIA	-49,54	-97,94 – -1,14	0,045
Efeitos aleatórios			
σ^2	8.373.593,62		
Desvio padrão $\tau_{00 \text{ IDITEM}}$	728		
Desvio padrão $\tau_{00 \text{ PARTICIPANTE}}$	1.965		
ICC	0,34		
$N_{\text{PARTICIPANTE}}$	57		
N_{IDITEM}	45		
Observações	2.553		
R^2 Marginal/ R^2 Condicional	0,023 / 0,359		

Fonte: autoria própria

Nesta análise, os coeficientes apresentados estão em milissegundos. Para calcular a probabilidade do tempo de resposta de todas as condições cognatas, consideramos a proficiência média em L2 dos participantes, que foi 36,15 e como intercepto os sintagmas nominais sem cognatos.

O Modelo Misto Linear (lmer) indica que as probabilidades de tempo de resposta no experimento 3 foram as seguintes: 7.079 milissegundos para sintagmas sem cognatos; 6.903 milissegundos para sintagmas com um cognato; e 6.797 milissegundos para sintagmas com dois

cognatos.

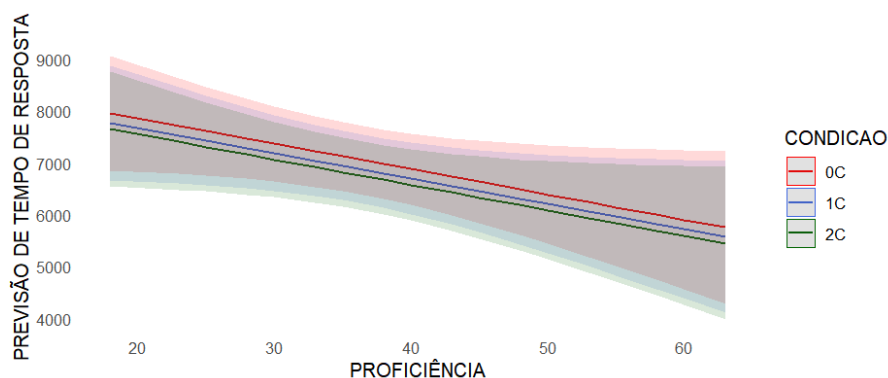
O Modelo Misto Linear também mostra que a variabilidade do desvio padrão do intercepto por item é de 728 milissegundos e que a variabilidade do desvio padrão por participante é de 1.965 milissegundos. Calculamos que o intercepto por item varia entre 8.107 e 9.567 milissegundos e por participante varia entre 6.856 e 10.819 milissegundos. Outrossim, o R^2 condicional indica que 36% da variância em acurácia pode ser explicada pela variância dos efeitos fixos e dos efeitos aleatórios.

Também, o Modelo Misto Linear (lmer) sugere diferença significativa, com valor de $p < 0,05$, para o nível de proficiência em L2 dos participantes. Isso significa que uma das variáveis preditoras exerce influência no tempo de resposta: a proficiência.

Reiteramos que decidimos por não levar em consideração a idade como uma variável preditora, após uma análise prévia, que indicou que a idade não se mostrou relevante nos resultados.

Apresentamos em um gráfico a probabilidade estimada de tempo de resposta no experimento 3:

Gráfico 34 - Previsão do tempo de resposta (ms) do experimento 3



Fonte: autoria própria

O gráfico de previsão de tempo de resposta nos informa que, com relação ao número de cognatos presentes nos sintagmas nominais, é possível perceber que o tempo de resposta dos sintagmas nominais com dois cognatos é menor do que o tempo de resposta dos sintagmas com um cognato, que é menor do que o tempo de resposta dos sintagmas sem cognatos, porém não de forma significativa. Outra informação é que à medida em que a proficiência em L2 vai aumentando, a previsão do tempo de resposta vai diminuindo. Outrossim, a curva do tempo de resposta está decrescendo de forma bastante acentuada, o que pode reforçar que o nível de

proficiência em L2 dos participantes interferiu no tempo de resposta de forma significativa.

Após a descrição dos resultados neste capítulo, apresentamos a discussão dos resultados no próximo capítulo.

6 DISCUSSÃO DOS RESULTADOS

Neste capítulo, apresentamos a discussão dos resultados obtidos na tarefa de reconhecimento de tradução, na tarefa de julgamento de aceitabilidade e na tarefa de seleção sintática, que são os experimentos que compõem esta tese. Também, apresentamos a comparação dos resultados obtidos entre os três experimentos.

Recordamos que minha tese é de que as diferentes estruturas sintáticas em L2 para representar os sintagmas nominais preposicionados com dois substantivos têm custos de processamento diferentes, podendo haver interferência do efeito cognato e do nível de proficiência em L2. Também, relembramos que o objetivo principal desta pesquisa foi analisar o processamento de sintagmas nominais preposicionados com dois substantivos por bilíngues tardios, que têm português como L1 e inglês como L2. Os três objetivos específicos foram: (1) identificar qual construção sintática em L2 tem menor custo de processamento, dentre as três estruturas possíveis: preposição *of*, caso genitivo e inversão dos substantivos; (2) examinar se há efeito cognato no processamento destes sintagmas nominais; e (3) investigar se há relação entre o processamento de sintagmas nominais e o nível de proficiência em L2.

Retomando as hipóteses desta pesquisa, temos a primeira hipótese, que considera que o custo de processamento da estrutura com a preposição *of* deve ser menor, porque se assemelha mais com a L1. Esta hipótese se baseia em Bernolet e Hartsuiker (2018), que apontam que a compreensão da L2 pode ser impulsionada pela sintaxe da L1, especialmente em estágios iniciais de aquisição da L2. A segunda hipótese pressupõe que há interferência de cognatos no processamento de sintagmas nominais em L2, levando ao aumento da acurácia e à diminuição do tempo de resposta. Para esta hipótese, tivemos vários estudos como referência, tais quais Dijkstra *et al.* (2015), Poort e Rodd (2017), Soares *et al.* (2018), Toassi e Pereira (2019) e Antón e Duñabeitia (2020), que detectaram que os cognatos podem facilitar o processamento de L2. Por fim, a terceira hipótese, que acredita que há interferência do nível de proficiência no processamento de sintagmas nominais em L2. Acreditamos que os bilíngues com menor nível de proficiência em inglês como L2 apresentam menos acurácia e precisam de mais tempo para processar os sintagmas nominais. Os estudos de Bernolet e Hartsuiker (2018) e Lan *et al.* (2019) nos embasam para esta hipótese. Bernolet e Hartsuiker (2018) propuseram um modelo de aquisição de L2 para bilíngues tardios, que indica diferentes processamentos de L2 para diferentes estágios de proficiência. E Lan *et al.* (2019) identificaram que o processamento de sintagmas nominais é dependente do nível de proficiência dos bilíngues.

A seguir, descrevemos os resultados obtidos com relação à acurácia e ao tempo de resposta nos dois primeiros experimentos e com relação à seleção sintática e ao tempo de resposta no terceiro experimento, por meio da estatística descritiva e inferencial. Para isso, detalhamos em cada seção um experimento. Ao final, apresentamos a comparação dos resultados dos três experimentos.

6.1 Tarefa de reconhecimento de tradução

Na tarefa de reconhecimento de tradução, os resultados gerais apontaram acurácia de 50,95% e tempo de resposta médio de 4.372 milissegundos por item.

Como um dos objetivos era investigar qual construção sintática seria a menos custosa para os bilíngues, fizemos análise dos dados por construção sintática. Os dados mostraram que o custo de processamento da sintaxe com uso da preposição *of* e da sintaxe com inversão de substantivos foram os menores, se mostrando muito próximos. Com relação à acurácia, tivemos 64,64% de acerto para as sentenças com uso da preposição *of* e 64,95% de acerto para as sentenças com inversão de substantivos. Sobre o tempo de resposta, foram gastos em média 4.367 milissegundos para processar corretamente uma sentença com preposição *of* e 4.236 milissegundos para o processamento de uma sentença com inversão de substantivos.

Fizemos também a análise por condição cognata, pois um dos objetivos era verificar se havia efeito cognato no processamento de sintagmas nominais preposicionados com dois substantivos. Os dados mostraram que a acurácia foi maior em sentenças com cognatos, com registros de 54,77% de acerto, do que em sentenças sem cognatos, que registraram 46,67% de acerto. O tempo de resposta também foi menor para sentenças com cognatos, com média de 4.206 milissegundos, em comparação com sentenças sem cognatos, com média de 4.589 milissegundos.

Analisamos com mais detalhes o efeito cognato. Para verificar se a quantidade de cognatos fazia diferença no processamento dos sintagmas nominais, comparamos as sentenças com um cognato e com dois cognatos. A acurácia se mostrou maior quando havia dois cognatos: 57,14%. Quando havia um cognato, a acurácia foi 52,16% e quando não havia cognatos, a acurácia foi 46,67%. O tempo de resposta também foi menor para sintagmas com dois cognatos: 3.964 milissegundos, em média. Para sintagmas com um cognato, a média do tempo de resposta foi 4.499 milissegundos e para sintagmas sem cognatos, a média do tempo de resposta foi 4.589 milissegundos.

Fizemos comparações entre todas as condições sintáticas e cognatas possíveis neste

estudo. A acurácia foi maior para as sentenças com inversão de substantivos e dois cognatos e para as sentenças com preposição *of* e dois cognatos, com acerto de 77,82% e 76,47%, respectivamente. O tempo de resposta foi menor para as sentenças com inversão de substantivos e dois cognatos e para as sentenças com preposição *of* e dois cognatos, com registros de 3.784 e 3.977 milissegundos, em média, respectivamente.

Os resultados deste experimento apontaram dificuldade de domínio sintático do caso genitivo pelos bilíngues tardios, o que foi evidenciado pelos baixos índices de acurácia e pelos maiores tempos de resposta para reconhecer as traduções. Já na apresentação dos resultados, fizemos algumas ponderações sobre a semelhança estrutural entre o caso genitivo e a inversão de substantivos e sobre o uso mais restrito do caso genitivo, o que pode levar a uma dificuldade de processamento maior, em comparação com as outras duas construções sintáticas investigadas. Um dos resultados mais surpreendentes do meu estudo se refere à acurácia dos itens com caso genitivo e dois cognatos (18,57%), que foi a menor dentre todas as condições sintáticas e cognatas possíveis. Além disso, os resultados se mostraram semelhantes aos encontrados por Ortiz-Preuss (2011), cujos participantes da pesquisa foram mais rápidos para responder itens com cognatos, mas apresentaram baixa acurácia nesses mesmos itens. A autora defende que seus resultados podem estar relacionados à sobrecarga atencional, proposta por Costa *et al.* (2006), que afirma que a similaridade interlinguística nem sempre pode estar associada à facilidade de processamento de L2.

A análise inferencial, por meio do Modelo Linear Misto Generalizado (glmer) revelou que a estrutura sintática pode interferir de forma significativa na acurácia do reconhecimento de tradução, podendo levar a uma probabilidade de acerto maior para itens com preposição *of* e para itens com inversão de substantivos. Isto significa que o custo de processamento foi menor para itens com preposição *of* e para itens com inversão de substantivos do que para itens com caso genitivo.

Através do Modelo Linear Misto (lmer), a análise inferencial também revelou que a estrutura sintática pode interferir de forma significativa no tempo de resposta do reconhecimento de tradução, podendo levar a uma previsão de tempo de resposta menor para itens com preposição *of* e para itens com inversão de substantivos. Em outras palavras, o custo de processamento foi menor para itens com preposição *of* e para itens com inversão de substantivos do que para itens com caso genitivo. Além disso, o tempo de resposta também foi significativamente menor para itens com dois cognatos.

Considerando as variáveis preditoras e as variáveis de resposta, temos os seguintes resultados: (1) a sintaxe exerceu influência na acurácia; (2) e a sintaxe e o efeito cognato

exerceram influência no tempo de resposta. Em resumo, as diferentes estruturas sintáticas levaram a menor custo de processamento dos itens com preposição *of* e nos itens com inversão de substantivos, no que se refere a maior acerto e a menor tempo de resposta. Outrossim, as quantidades de cognatos sugerem menor custo de processamento, no que se refere a menor tempo de resposta nos itens com dois cognatos.

Os resultados obtidos no experimento 1 confirmam minha primeira hipótese, que defendia que o custo de processamento com a preposição *of* deveria ser menor, porque é a estrutura sintática que mais se assemelha à L1 dos bilíngues. Em minha pesquisa, os resultados de acurácia e tempo de resposta ficaram muito próximos para sintagmas com preposição *of* e para sintagmas com inversão de substantivos. Isso significa que o custo de processamento foi menor para estas duas estruturas sintáticas, o que pode revelar que os bilíngues utilizaram a semelhança entre suas línguas para reconhecer as traduções. Este resultado também pode reforçar o que propõem Bernolet e Hartsuiker (2018), em seu modelo de desenvolvimento sintático da L2 para bilíngues tardios. Neste modelo, as representações sintáticas ainda não estão formadas em estágios iniciais de aprendizagem da L2 e os bilíngues tendem a transferir construções sintáticas da L1 para a L2, o que faz com que a compreensão da L2 seja impulsionada pelas preferências sintáticas da L1. Neste mesmo modelo, Bernolet e Hartsuiker (2018) explicam que, à medida em que tem mais contato com a L2, o bilíngue começa a aproximar as conexões entre palavras e representações sintáticas específicas. E em um estágio mais avançado, a informação sintática começa a se tornar mais abstrata e, conseqüentemente, as representações sintáticas da L2 se desenvolvem de forma separada das representações sintáticas da L1. Estes estágios posteriores ao estágio inicial podem explicar o baixo custo cognitivo também para os sintagmas com inversão de substantivos. Como os participantes tinham diferentes níveis de proficiência em L2, ou seja, estão em diferentes estágios de aprendizagem de L2, tanto as transferências das construções sintáticas da L1 para a L2, quanto a abstração das informações sintáticas ocorreram. Já a baixa acurácia e o tempo de resposta maior para itens com caso genitivo pode explicado pela supergeneralização da regra do genitivo, fenômeno observado em salas de aula de inglês como L2, que já havíamos citado anteriormente, e que foi reforçado por Romano (2016), que encontrou no caso genitivo maior grau de dificuldade em falantes não-nativos. Acreditamos que a supergeneralização do caso genitivo possa ocorrer pela semelhança com a estrutura que utiliza a inversão de substantivos, que se diferenciam apenas pela presença do apóstrofo (') e do *s*, que acompanham o primeiro substantivo. Por exemplo, temos o sintagma *chocolate bar*, que pode ser considerado acurado, enquanto *chocolate's bar* não pode ser considerado acurado, conforme Murphy (2015), pois o

caso genitivo deve ser utilizado para indicar relações de posse, familiares, de propósito ou de origem, quando o substantivo modificador é animado, como em *John's pen* (A caneta de João) e em *women's rights* (direitos das mulheres).

A segunda hipótese também foi confirmada no experimento 1. Os resultados apontam o efeito de facilitação cognata no que diz respeito ao tempo de resposta. Ou seja, os bilíngues levaram menos tempo para reconhecer as traduções quando havia dois cognatos nos itens, o que gerou menor custo de processamento. Estes resultados estão em consonância com os resultados encontrados em Dijkstra *et al.* (2015), Poort e Rodd (2017), Soares *et al.* (2018), Toassi e Pereira (2019) e Antón e Duñabeitia (2020), que confirmaram o efeito de facilitação cognata em diferentes tipos de tarefas. Destacamos os resultados encontrados em Dijkstra *et al.* (2015) e em Soares *et al.* (2018), que investigaram o efeito cognato em sentenças, assim como foi feito em meu estudo. Vale relembrar que o efeito de facilitação cognata pressupõe coativação interlinguística e dá suporte à hipótese de não seletividade das línguas em mentes bilíngues, o que causa menor custo de processamento linguístico. Como explicado por Traxler (2012), isto quer dizer que quando um bilíngue está usando uma de suas línguas, a outra também está ativa.

Por fim, a tarefa de reconhecimento de tradução não confirmou a terceira hipótese, que acreditava na interferência do nível de proficiência no processamento de sintagmas nominais em L2. Apesar dos gráficos dos modelos mostrarem uma tendência clara de menor custo de processamento para maiores níveis de proficiência em L2, a diferença não se mostrou significativa. Acreditamos que isto tenha ocorrido pela estrutura da própria tarefa, que apresentava o item em L1 e sua tradução em L2, o que permitiu comparação entre as línguas. De certo modo, a comparação interlinguística pode ter auxiliado participantes bilíngues com níveis de proficiência em L2 mais baixos a reconhecer as traduções com acurácia.

6.2 Tarefa de julgamento de aceitabilidade

Na tarefa de julgamento de aceitabilidade, os resultados gerais apontaram acurácia de 51,85% e tempo de resposta médio de 3.029 milissegundos por item.

Fizemos análise dos dados por construção sintática, com o objetivo de investigar qual construção sintática seria a menos custosa para os bilíngues. Os dados mostraram que o custo linguístico da sintaxe com uso da preposição *of* e da sintaxe com inversão de substantivos foram próximos. Com relação à acurácia, tivemos 61,03% de acerto para as sentenças com uso da preposição *of* e 65,98% de acerto para as sentenças com inversão de substantivos. Sobre o tempo de resposta, foram gastos em média 2.971 milissegundos para processar corretamente

uma sentença com preposição *of* e 2.921 milissegundos para o processamento de uma sentença com inversão de substantivos.

Fizemos também a análise por condição cognata, pois um dos objetivos era verificar se havia efeito cognato no processamento de sintagmas nominais preposicionados com dois substantivos. Os dados mostraram que a acurácia foi maior em sentenças com cognatos, com registros de 54,72% de acerto, do que em sentenças sem cognatos, que registraram 48,66% de acerto. O tempo de resposta também foi menor para sentenças com cognatos, com média de 2.909 milissegundos, em comparação com sentenças sem cognatos, com média de 3.181 milissegundos.

Analisamos com mais detalhes o efeito cognato. Para verificar se a quantidade de cognatos fazia diferença no processamento dos sintagmas nominais, comparamos as sentenças com um cognato e com dois cognatos. A acurácia se mostrou maior quando havia dois cognatos: 55,57%. Quando havia um cognato, a acurácia foi 53,78% e quando não havia cognatos, a acurácia foi 48,66%. O tempo de resposta também foi menor para sintagmas com dois cognatos: 2.835 milissegundos, em média. Para sintagmas com um cognato, a média do tempo de resposta foi 2.993 milissegundos e para sintagmas sem cognatos, a média do tempo de resposta foi 3.181 milissegundos.

Fizemos comparações entre todas as condições sintáticas e cognatas possíveis neste estudo. A acurácia foi maior para as sentenças com inversão de substantivos e dois cognatos e para as sentenças com preposição *of* e dois cognatos, com acerto de 79,08% e 69,58%, respectivamente. O tempo de resposta foi menor para as sentenças com preposição *of* e dois cognatos e para as sentenças com inversão de substantivos e dois cognatos, com registros de 2.720 e 2.769 milissegundos, em média, respectivamente.

Assim como ocorreu com o primeiro experimento, os resultados do segundo experimento também indicaram dificuldade de domínio sintático do caso genitivo pelos bilíngues tardios. Na tarefa de julgamento de aceitabilidade, os itens com caso genitivo apresentaram baixos índices de acurácia e maiores tempos de resposta. Já fizemos algumas reflexões sobre a semelhança estrutural entre o caso genitivo e a inversão de substantivos e sobre o uso mais restrito do caso genitivo na subseção anterior. A baixa acurácia dos itens com caso genitivo e dois cognatos (18,86%), que foi a menor dentre todas as condições sintáticas e cognatas possíveis, também se assemelham aos resultados encontrados por Ortiz-Preuss (2011), que acredita que seus resultados podem estar relacionados à sobrecarga atencional, proposta por Costa *et al.* (2006).

A análise inferencial, por meio do Modelo Linear Misto Generalizado (glmer) revelou que a estrutura sintática e o nível de proficiência em L2 podem interferir de forma significativa na acurácia dos julgamentos de aceitabilidade, podendo levar a uma probabilidade de acerto maior para itens com preposição *of* e para itens com inversão de substantivos. Isto significa que o custo cognitivo foi menor para itens com preposição *of* e para itens com inversão de substantivos do que para itens com caso genitivo. Ademais, o custo cognitivo também foi menor para bilíngues com maior nível de proficiência em L2.

Através do Modelo Linear Misto (lmer), a análise inferencial também revelou que a sintaxe e o efeito cognato podem interferir de forma significativa no tempo de resposta dos julgamentos de aceitabilidade, podendo levar a uma previsão de tempo de resposta menor para itens com inversão de substantivos. Isto significa que o custo cognitivo foi significativamente menor para itens com inversão de substantivos do que para itens com caso genitivo. Além disso, o tempo de resposta também foi significativamente menor para itens com um cognato e para itens com dois cognatos, em comparação com itens sem cognatos.

Considerando as variáveis preditoras e as variáveis de resposta, temos os seguintes resultados: (1) a sintaxe e o nível de proficiência em L2 exerceram influência na acurácia; (2) e a sintaxe e o efeito cognato exerceram influência no tempo de resposta. Em resumo, as diferentes estruturas sintáticas levaram a menor custo de processamento, no que se refere a maior acerto e a menor tempo de resposta, nos itens com inversão de substantivos. E maiores níveis de proficiência em L2 podem levar a menor custo de processamento para fazer um julgamento com acurácia. Outrossim, as quantidades de cognatos sugerem menor custo de processamento, no que se refere a menor tempo de resposta nos itens com um cognato e nos itens com dois cognatos, em comparação a itens sem cognatos.

A primeira hipótese postulava que o custo de processamento com a preposição *of* deveria ser menor, porque é a estrutura sintática que mais se assemelha à L1 dos bilíngues, o que foi confirmado pelos resultados do experimento 2. Nesta pesquisa, os resultados de acurácia ficaram muito próximos para sintagmas com preposição *of* e para sintagmas com inversão de substantivos. Isso significa que o custo de processamento foi menor para estas duas estruturas sintáticas, o que pode revelar que os bilíngues foram capazes tanto de utilizar a semelhança entre suas línguas, quanto de aproximar as conexões entre palavras e representações sintáticas específicas dos sintagmas nominais preposicionados com dois substantivos em L2, afastando-se das representações sintáticas em L1. Dessa forma, o resultado do experimento 2 pode reforçar o que propõem Bernolet e Hartsuiker (2018), em seu modelo de desenvolvimento sintático da L2 para bilíngues tardios. Neste modelo, as representações sintáticas ainda não estão formadas

em estágios iniciais de aprendizagem da L2 e os bilíngues tendem a transferir construções sintáticas da L1 para a L2, o que faz com que a compreensão da L2 seja impulsionada pelas semelhanças sintáticas da L1. Neste mesmo modelo, Bernolet e Hartsuiker (2018) explicam que, à medida em que tem mais contato com a L2, o bilíngue começa a aproximar as conexões entre palavras e representações sintáticas específicas. E em um estágio mais avançado, a informação sintática começa a se tornar mais abstrata e, conseqüentemente, as representações sintáticas da L2 se desenvolvem de forma separada das representações sintáticas da L1, como já mencionado. Estes estágios posteriores ao estágio inicial podem explicar o baixo custo de processamento também para os sintagmas com inversão de substantivos. Acreditamos que tanto a transferência das construções sintáticas da L1 para a L2 quanto as abstrações sintáticas tenham sido feitas porque tivemos bilíngues com diferentes níveis de proficiência em L2, em diferentes estágios de aprendizagem de L2. Já a baixa acurácia e o tempo de resposta maior para itens com caso genitivo pode ser explicado pela supergeneralização da regra do genitivo, fenômeno observado em salas de aula de inglês como L2, que já havíamos listado anteriormente e que foi reforçado por Romano (2016), que encontrou no caso genitivo maior grau de dificuldade em falantes não-nativos. Assim, como já explicamos na subseção anterior, que discute os resultados do experimento 1, acreditamos que a supergeneralização do caso genitivo possa ocorrer pela semelhança visual entre o caso genitivo e a construção que utiliza a inversão de substantivos.

O experimento 2 também confirmou a segunda hipótese, indicando efeito de facilitação cognata no que diz respeito ao tempo de resposta. Ou seja, os bilíngues levaram menos tempo para fazer os julgamentos dos itens quando havia um ou dois cognatos nos itens, o que gerou menor custo de processamento. Estes resultados estão em consonância com os resultados encontrados em Dijkstra *et al.* (2015), Poort e Rodd (2017), Soares *et al.* (2018), Toassi e Pereira (2019) e Antón e Duñabeitia (2020), cujos resultados de pesquisa confirmaram o efeito de facilitação cognata em diferentes tipos de tarefas. Os resultados estão em convergência especialmente com os resultados obtidos em Dijkstra *et al.* (2015) e em Soares *et al.* (2018), que, assim como em minha tese, pesquisaram o efeito cognato em sentenças e não em palavras isoladas. Vale lembrar que o efeito de facilitação cognata pressupõe coativação interlinguística e dá suporte à hipótese de não seletividade das línguas em mentes bilíngues, o que causa menor custo de processamento. Como já mencionado, isto quer dizer que quando um bilíngue está usando uma de suas línguas, a outra também está ativa, conforme Traxler (2012).

Finalmente, a tarefa de julgamento de aceitabilidade confirmou a terceira hipótese, que postulava a possibilidade de interferência do nível de proficiência no processamento de sintagmas nominais em L2. Níveis de proficiência mais altos em L2 colaboraram para maior

acurácia de forma significativa, e menor tempo de resposta, mas não de forma significativa. Acreditamos que isto tenha ocorrido pela estrutura da própria tarefa, que apresentava os itens apenas em L2. Diferentemente do que ocorreu no primeiro experimento, não era possível fazer a comparação entre as línguas. Isto significa que níveis de proficiência em L2 mais altos fizeram com os itens fossem julgados corretamente mais vezes. Este resultado reforça o que propõe Bernolet e Hartsuiker (2018), em seu modelo de desenvolvimento sintático da L2 para bilíngues tardios, e o que foi pesquisado por Lan *et al.* (2019).

6.3 Tarefa de seleção sintática

Na tarefa de seleção sintática, os resultados gerais apontaram preferência pela sintaxe com inversão de substantivos e pelo caso genitivo, com registros de 41,25% e 38,19%, nesta ordem. Com relação ao tempo de resposta para fazer a escolha pela sintaxe adequada, a menor média de tempo de resposta foi para o caso genitivo: 6.415 milissegundos. O segundo menor tempo de resposta foi para a inversão de substantivos, com 7.181 milissegundos gastos, em média, para fazer a seleção.

Também analisamos o efeito cognato, assim como fizemos nos dois primeiros experimentos. Para verificar o efeito cognato no processamento dos sintagmas nominais, comparamos as seleções sintáticas feitas pelos participantes, considerando as quantidades de cognatos nos sintagmas. Quando não havia cognatos, a inversão de substantivos foi selecionada mais vezes como a sintaxe mais adequada: 41,27%. Quando havia um cognato, o caso genitivo foi a construção sintática mais selecionada: 41,31%. E quando havia dois cognatos, a construção sintática mais escolhida foi a inversão de substantivos: 44,81%. Já o tempo de resposta se mostrou menor para sintagmas com dois cognatos: 6.800 milissegundos, em média. Para sintagmas com um cognato, a média do tempo de resposta foi 6.909 milissegundos e para sintagmas sem cognatos, a média do tempo de resposta foi 7.088 milissegundos.

Fizemos comparações entre todas as condições sintáticas e cognatas possíveis neste estudo. O tempo de resposta foi menor para as sentenças com caso genitivo e dois cognatos: 6.321 milissegundos, em média.

A análise inferencial, por meio do Modelo Logístico Multinomial (*polynomous*) revelou que o efeito cognato e o nível de proficiência em L2 podem interferir de forma significativa na seleção sintática para representar os sintagmas nominais preposicionados com dois substantivos em L2. Isto significa que a presença de cognatos fez com que os bilíngues selecionassem menos vezes a opção com preposição *of* (quando os itens tinham dois cognatos)

e selecionassem mais vezes a opção com caso genitivo (quando os itens tinham um cognato). Ademais, bilíngues com níveis maiores de proficiência em L2 selecionaram mais vezes as opções com preposição *of* e com inversão de substantivos.

Através do Modelo Linear Misto (lmer), a análise inferencial também revelou que o nível de proficiência em L2 pode interferir de forma significativa no tempo de resposta para fazer a seleção sintática. Isto significa que o custo cognitivo foi significativamente menor para bilíngues com níveis de proficiência em L2 mais altos.

Considerando as variáveis preditoras e as variáveis de resposta, temos os seguintes resultados: (1) o efeito cognato e o nível de proficiência em L2 exerceram influência na seleção sintática; (2) e o nível de proficiência em L2 interferiu no tempo de resposta. Em resumo, a quantidade de cognatos e o nível de proficiência em L2 influenciaram na estrutura sintática a ser selecionada. O nível de proficiência em L2 também levou a um menor custo cognitivo para selecionar itens com caso genitivo e com inversão de substantivos.

O experimento 3 confirmou a segunda hipótese, indicando efeito cognato no que diz respeito às escolhas sintáticas. Os bilíngues se tornaram propensos a selecionar mais vezes as opções com caso genitivo quando havia um cognato nos sintagmas e a selecionar menos vezes as opções com preposição *of* quando havia dois cognatos nos sintagmas. Isto sugere que os cognatos interferiram nas seleções sintáticas, indicando efeito cognato, mas não necessariamente efeito de facilitação cognata, pois neste experimento a acurácia não foi levada em consideração, apenas a preferência sintática dos bilíngues. O efeito de facilitação cognata foi constatado no tempo de resposta, mas não de forma significativa.

Por fim, a tarefa de seleção sintática confirmou a terceira hipótese, que postulava interferência do nível de proficiência no processamento de sintagmas nominais em L2, indicando que quanto maior era o nível de proficiência em L2 dos bilíngues, maiores eram as chances de escolha das opções com caso genitivo e com inversão de substantivos. Além disso, maiores níveis de proficiência em L2 também levaram a tempos de resposta menores, apontando para menor custo de processamento para fazer as seleções sintáticas. Este resultado reforça o que propõe Bernolet e Hartsuiker (2018), em seu modelo de desenvolvimento sintático da L2 para bilíngues tardios, e o que foi pesquisado por Lan *et al.* (2019). Tanto Bernolet e Hartsuiker (2018), quanto Lan *et al.* (2019), explicam que níveis maiores de proficiência em L2 e processamento mais complexo e abstrato da língua não-nativa estão ligados diretamente.

6.4 Comparação dos três experimentos

Primeiramente, é válido explicar que a comparação entre os dois primeiros experimentos parece ser mais viável, uma vez que nestas duas tarefas os participantes tinham acesso a uma estrutura sintática em L2 por vez: na tarefa de reconhecimento de tradução, eles liam a mesma sentença em L1 e em L2, e na tarefa de julgamento de aceitabilidade, eles liam uma sentença em L2. Por outro lado, no terceiro experimento, os participantes tinham acesso a quatro sentenças de uma só vez: eles liam a frase em L1 e liam as três opções em L2, que traziam as três possibilidades sintáticas. Além disso, a acurácia foi analisada somente nos dois primeiros experimentos, uma vez que decidimos analisar somente a preferência sintática dos bilíngues no terceiro experimento. Estas informações esclarecem que como a análise do terceiro experimento se diferenciou um pouco dos dois primeiros, a comparação entre a tarefa de reconhecimento de tradução e a tarefa de julgamento de aceitabilidade pode ser mais aprofundada. Finalmente, comparando os resultados dos três experimentos, apresentamos a análise das variáveis preditoras sobre as variáveis de resposta.

A primeira variável preditora, que foi a sintaxe, se mostrou influente na acurácia e no tempo de resposta nos dois primeiros experimentos. Tanto no reconhecimento de traduções quanto nos julgamentos de aceitabilidade, a sintaxe exerceu influência para uma maior acurácia em sintagmas com preposição *of* e com inversão de substantivos, em detrimento dos sintagmas com caso genitivo. A sintaxe também interferiu no tempo de resposta, fazendo com que o tempo para reconhecer as traduções dos itens com preposição *of* e com inversão de substantivos fosse menor do que em itens com caso genitivo e fazendo com que o tempo para realizar os julgamentos de aceitabilidade dos itens com inversão de substantivos fosse menor do que em itens com preposição *of* e em itens com caso genitivo.

A segunda variável preditora, que foi o efeito cognato, se mostrou influente na seleção sintática do terceiro experimento e no tempo de resposta dos dois primeiros experimentos realizados. Na tarefa de seleção sintática, o efeito cognato exerceu influência de duas formas: os itens com dois cognatos fizeram com que as opções que tinham preposição *of* fossem escolhidas menos vezes, em comparação com itens sem cognatos e com um cognato e os itens com um cognato fizeram com que as opções que tinham caso genitivo fossem escolhidas mais vezes, em comparação com itens sem cognatos e com dois cognatos. O efeito cognato também interferiu no tempo de resposta. No reconhecimento de traduções, os itens com dois cognatos fizeram com que o tempo de resposta fosse menor, em comparação com os itens sem cognatos e com um cognato. Nos julgamentos de aceitabilidade, tanto os itens com um

cognato quanto os itens com dois cognatos fizeram com que o tempo de resposta fosse menor, em comparação com os itens sem cognatos.

A terceira variável preditora, que foi o nível de proficiência em L2, se mostrou influente na acurácia do segundo experimento e no tempo de resposta do terceiro experimento. Nos julgamentos de aceitabilidade, a acurácia foi maior para participantes com maiores níveis de proficiência em L2. Nas seleções sintáticas, bilíngues com maiores níveis de proficiência em L2 levaram menos tempo para fazer as seleções.

Por um lado, com relação aos dois primeiros experimentos, a tarefa de reconhecimento de tradução e a tarefa de julgamento de aceitabilidade, os resultados se mostraram muito próximos, diferindo em poucos aspectos. Vejamos alguns pontos sobre a acurácia e o tempo de resposta.

A acurácia dos dois experimentos se mostrou maior para itens com preposição *of* e para itens com inversão de substantivos, mas o nível de proficiência em L2 interferiu de forma significativa apenas na tarefa de julgamento de aceitabilidade. Tínhamos previsto anteriormente que esta tarefa poderia se mostrar mais desafiadora do que a tarefa de reconhecimento de tradução. Enquanto no primeiro experimento os participantes liam os itens em L1 e suas traduções em L2, no experimento 2 os bilíngues acessaram somente os itens em L2, o que não permitia uma possível comparação com a língua nativa. Desse modo, o nível de proficiência se mostrou relevante na acurácia dos julgamentos.

No que diz respeito ao tempo de resposta, os resultados também se mostraram próximos nestes dois experimentos, pois tanto a sintaxe quanto os cognatos influenciaram no tempo de resposta dos bilíngues. O tempo gasto para reconhecer itens com preposição *of* e com inversão de substantivos foi significativamente menor do que para reconhecer itens com caso genitivo. Igualmente, o tempo gasto para julgar itens com preposição *of* e com inversão de substantivos foi menor do que para julgar itens com caso genitivo, sendo significativamente menor apenas para itens com inversão de substantivos. Ademais, o tempo de resposta para reconhecer as traduções foi reduzido de forma significativa quando os itens tinham dois cognatos e foi também reduzido significativamente para fazer os julgamentos de aceitabilidade quando havia um ou dois cognatos nos itens.

Como descrevemos na primeira hipótese, acreditávamos que o menor custo de processamento seria para os itens com preposição *of*, porque esta estrutura em L2 é a que mais se assemelha da estrutura em L1, mantendo os substantivos ligados por uma preposição na mesma ordem lexical nas duas línguas, ao contrário do que ocorre com o caso genitivo e a inversão de substantivos. Igualmente, o alto custo para o processamento do caso genitivo

também foi previsto, por causa da supergeneralização de seu uso, pois especialmente bilíngues em níveis iniciais de aprendizagem de L2 costumam aplicar o 's a todos os possuidores, não importando se são seres animados ou não. Já o baixo custo de processamento dos itens com inversão de substantivos, mostrando resultados muito próximos aos dos itens com preposição *of* nos surpreendeu. Acreditamos que o baixo custo de processamento de itens com inversão de substantivos pode ser relacionado à familiaridade de alguns sintagmas que aparecem nas primeiras unidades do livro didático utilizado pelos participantes desta pesquisa, como *phone number*, *family name*, *art museum*, e *credit card*.

Por outro lado, a tarefa de seleção sintática se mostrou diferente nos resultados pela sua própria estrutura, pois neste experimento a acurácia não foi levada em consideração, apenas a preferência sintática dos bilíngues. Como esta tarefa apresentava as três opções sintáticas em L2 para a representação dos sintagmas nominais preposicionados com dois substantivos, a análise do efeito cognato foi feita com dois objetivos: para verificar possível interferência nas seleções dos participantes e no tempo de resposta. A quantidade de cognatos interferiu nas seleções dos participantes, mas isso não sugere necessariamente efeito de facilitação cognata, uma vez que neste experimento não consideramos a acurácia, mas tão somente a preferência sintática dos bilíngues. Além disso, a quantidade de cognatos interferiu no tempo de resposta, mas não de forma significativa. Outrossim, o terceiro experimento se mostrou diferente dos dois primeiros experimentos com relação ao nível de proficiência em L2 dos bilíngues. Na tarefa de seleção de sintática, o nível de proficiência reduziu o tempo de resposta de forma significativa.

Como descrevemos na terceira hipótese, acreditávamos que o nível de proficiência em L2 poderia influenciar nas seleções sintáticas dos bilíngues. Por se assemelhar mais à L1 dos participantes, as opções com preposição *of* seriam escolhidas mais vezes por participantes com baixos níveis de proficiência em L2, assim como propõe o modelo de Bernolet e Hartsuiker (2018), que aponta que a compreensão da L2 pode ser impulsionada pela sintaxe da L1, especialmente em estágios iniciais de aquisição da L2. Já as opções com caso genitivo e com inversão de substantivos seriam escolhidas mais vezes por participantes com níveis de proficiência em L2 mais altos, por serem estruturas sintáticas em L2 que se diferenciam em quantidade e em ordem lexical da única construção sintática possível em L1, o que também pode ser explicado pelo modelo proposto por Bernolet e Hartsuiker (2018), que esclarece que a abstração sintática vai se desenvolvendo à medida em que o nível de proficiência em L2 vai aumentando.

7 CONSIDERAÇÕES FINAIS

Neste capítulo, apresentamos as considerações finais sobre esta pesquisa, destacando os principais achados e possíveis contribuições para a área da Psicolinguística do Bilinguismo. Também, apresentamos possíveis limitações e sugestões para futuros estudos.

Aplicamos três experimentos para analisar o processamento de sintagmas nominais preposicionados com dois substantivos com 57 participantes bilíngues tardios, com português brasileiro como L1 e inglês como L2. Os bilíngues tinham média de 20 anos e média de proficiência em L2 de 36,15%. O nível de proficiência em L2 foi medido por um teste de conhecimento de vocabulário de inglês, baseado no teste de vocabulário do *Shipley Institute of Living Scale* (SILS), e adaptado pela professora Pamela Toassi (UnB), em parceria com os professores Justin Lauro (*Barrys University*) e Bernhard Angele (*Universidad de Madrid*).

O primeiro experimento foi uma tarefa de reconhecimento de tradução de L1 para L2 com 60 itens experimentais. Os bilíngues deveriam clicar na tecla 1, caso achassem que as sentenças eram equivalentes em tradução ou clicar na tecla 2, caso achassem que as sentenças não eram equivalentes em tradução. Para fazer o reconhecimento de tradução, os bilíngues tinham até 10.000 milissegundos (10 segundos). Neste experimento, as variáveis preditoras foram as construções sintáticas, as condições cognatas e o nível de proficiência em L2 dos participantes. As variáveis de respostas foram a acurácia e o tempo de resposta.

Os resultados do experimento 1 mostraram que as diferentes estruturas sintáticas apresentaram diferentes custos de processamento. Os itens com preposição *of* e os itens com inversão de substantivos apresentaram maior acerto e menor tempo de resposta, o que indica menor custo de processamento para estas duas estruturas sintáticas. Ademais, a presença de cognatos nos sintagmas levou os participantes a reconhecerem as traduções de forma mais rápida, o que também gerou menor custo de processamento.

O segundo experimento foi uma tarefa de julgamento de aceitabilidade com 60 itens experimentais em L2, exclusivamente. Os bilíngues deveriam clicar na tecla 1, caso achassem que as sentenças estavam bem construídas ou clicar na tecla 2, caso achassem que as sentenças não estavam bem construídas. Para fazer o julgamento, os bilíngues tinham até 10.000 milissegundos (10 segundos). Assim como ocorreu no experimento 1, as variáveis preditoras foram as construções sintáticas, as condições cognatas e o nível de proficiência em L2 dos participantes. As variáveis de respostas foram a acurácia e o tempo de resposta.

Os resultados do experimento 2 também mostraram custos de processamento

diferentes para as diferentes estruturas sintáticas. Os itens com inversão de substantivos apresentaram maior acerto e menor tempo de resposta de forma significativa e os itens com preposição *of*, exclusivamente, apresentaram maior acurácia significativamente. Outrossim, as quantidades de cognatos levaram os bilíngues a fazerem os julgamentos de aceitabilidade em menores tempos, fazendo com que o custo de processamento fosse menor. Também, maiores níveis de proficiência em L2 levaram os bilíngues a fazerem os julgamentos com maior acurácia, o que também pode gerar menor custo de processamento.

O terceiro experimento foi uma tarefa de seleção sintática com 45 itens experimentais em L1 e três opções de resposta, com diferentes estruturas sintáticas em L2. Os bilíngues deveriam clicar na tecla 1, 2 ou 3, de acordo com sua opção de resposta. Para fazer a seleção, os bilíngues tinham até 30.000 milissegundos (30 segundos). Neste último experimento, as variáveis preditoras foram as condições cognatas e o nível de proficiência em L2 dos participantes. As variáveis de respostas foram as seleções sintáticas e o tempo de resposta.

Os resultados do experimento 3 mostraram que a quantidade de cognatos e o nível de proficiência em L2 influenciaram na estrutura sintática a ser selecionada. Itens que continham um cognato levaram os bilíngues a selecionarem mais vezes as opções com caso genitivo e itens com dois cognatos levaram os bilíngues a selecionarem menos vezes as opções com preposição *of*. Também, participantes com maiores níveis de proficiência em L2 selecionaram mais opções com caso genitivo e com inversão de substantivos, assim como bilíngues com menores índices de proficiência em L2 selecionaram mais vezes como resposta itens com preposição *of*. Ademais, bilíngues com maiores níveis de proficiência em L2 fizeram as seleções sintáticas em menores tempos, o que gerou menor custo cognitivo.

Conforme minha prática em salas de aula de inglês como L2, acreditávamos que o processamento de sintagmas nominais preposicionados com dois substantivos ocorreria de modo diferente entre bilíngues tardios, variando entre as estruturas sintáticas possíveis. Para detalhar ainda mais a investigação dessas construções, resolvemos incluir em minha investigação o efeito cognato e o nível de proficiência em L2, que são aspectos frequentemente investigados em estudos de Psicolinguística do Bilinguismo.

Minha tese de que as três construções sintáticas em L2 são processadas de diferentes formas, no que diz respeito à acurácia e ao tempo de resposta, se confirmou. As três construções sintáticas apresentaram diferentes custos de processamento, indicando que os itens com preposição *of* e com inversão de substantivos tiveram menor custo de processamento, pois apresentaram maior índice de acurácia e menor tempo de resposta, dentre as três construções

pesquisadas, o que pode confirmar a primeira hipótese dessa tese. Também foi constatado que há dificuldade de processamento do caso genitivo, uma vez que os índices de acurácia foram bem menores e de tempo de resposta foram maiores, em comparação com as outras duas estruturas sintáticas.

O efeito cognato foi constatado nos três experimentos, o que pode confirmar a segunda hipótese dessa tese. Houve redução de custo de processamento na tarefa de reconhecimento de tradução, mostrado por meio de menor tempo de resposta para itens com dois cognatos. Também houve redução de custo de processamento na tarefa de julgamento de aceitabilidade, comprovado através de menor tempo de resposta para itens com um ou dois cognatos. Ou seja, o efeito de facilitação cognata foi confirmado nos dois primeiros experimentos. Além disso, a presença de cognatos mostrou influência na seleção sintática no terceiro experimento, mas este resultado não está ligado ao efeito de facilitação cognata, pois neste experimento a acurácia não foi levada em consideração, apenas a preferência sintática dos bilíngues para representar os sintagmas nominais preposicionados com dois substantivos.

O nível de proficiência em L2 interferiu em dois dos três experimentos que aplicamos, o que pode confirmar a terceira hipótese dessa tese. Os bilíngues mais proficientes responderam a tarefa de julgamento de aceitabilidade com maior acurácia e levaram menos tempo para fazer as seleções sintáticas no experimento 3.

Os resultados encontrados nesta tese podem reforçar o modelo de desenvolvimento sintático da L2 para bilíngues tardios, proposto por Bernolet e Hartsuiker (2018). Igualmente, os resultados podem reforçar o efeito de facilitação cognata, fenômeno comumente constatado nas pesquisas de Psicolinguística do Bilinguismo, que pressupõe que a presença de itens lexicais cognatos levam a menores custos de processamento de L2. Consequentemente, o efeito de facilitação cognata pode sustentar que há coativação interlinguística e dá suporte à hipótese de não seletividade das línguas em mentes bilíngues. E, por fim, os resultados também podem corroborar com os estudos que defendem que diferentes níveis de proficiência podem levar a diferentes estágios de processamento de L2.

Dessa maneira, acreditamos esta tese pode contribuir com a discussão acerca do processamento da L2 na mente dos bilíngues, especialmente no que diz respeito ao processamento sintático de sintagmas nominais, ao efeito cognato e ao nível de proficiência em L2. Também acreditamos que podemos contribuir com discussões acerca da interferência sintática da L1 para a L2, pois segundo Glenday *et al.* (2020), ainda há poucos estudos sobre a interferência sintática interlinguística, em comparação com outros aspectos linguísticos.

Este estudo também pode contribuir com as práticas de ensino de L2, uma vez que foi comprovado cientificamente aquilo que já observávamos nestas salas de aula: as três estruturas sintáticas que podem representar os sintagmas preposicionados com dois substantivos em L2 apresentam diferentes custos de processamento, podendo haver interferência do efeito cognato e do nível de proficiência em L2. Com estes resultados, foi possível constatar que a comparação interlinguística pode favorecer o processo de aquisição da sintaxe de L2 em fases iniciais de estudo, especialmente de construções sintáticas semelhantes, como é o caso dos sintagmas nominais preposicionados com dois substantivos em L1 e a construção que utiliza o pronome *of* em L2. Da mesma forma, os bilíngues tardios são capazes de adquirir estruturas sintáticas em L2 que não se assemelham à L1, à medida em que o contato com a L2 vai se intensificando. Isto significa que, conforme o nível de proficiência em L2 vai avançando, o bilíngue também é capaz de incorporar novas estruturas sintáticas em L2. Da mesma maneira, a utilização de itens lexicais que são cognatos pode favorecer o processo de aprendizagem de estruturas sintáticas em L2, levando ao processamento correto de construções sintáticas em menor tempo, sejam elas semelhantes ou não à L1. Além disso, a dificuldade apresentada no processamento dos itens com caso genitivo pode evidenciar a necessidade de instrução explícita, especialmente com substantivos cognatos, nas salas de aula de inglês como L2.

Como limitações, pensamos que em futuras pesquisas o número de participantes pode ser ampliado. Em comparação com minha dissertação, este número já aumentou de 20 para 57 participantes. Mesmo assim, acreditamos que podemos alcançar um número maior de bilíngues e mais participantes por nível de proficiência em L2 em estudos futuros. Outra sugestão para futuras pesquisas seria ampliar as diferentes idades de aquisição da L2 dos bilíngues tardios. Em meu estudo, dos 57 participantes pesquisados, apenas 7 eram adultos que estão cursando inglês como L2. Ou seja, 50 participantes eram adolescentes, com idade de aquisição de L2 entre 15 e 20 anos. E ainda, seria possível incluir no estudo bilíngues precoces e simultâneos. Esses fatores poderiam fornecer uma compreensão mais abrangente sobre os mecanismos subjacentes ao processamento sintático em L2, além de possibilitar a comparação entre bilíngues com diferentes perfis de experiência. Ademais, pensamos na possibilidade de aprofundar a pesquisa sobre os sintagmas preposicionados com dois substantivos utilizando outros métodos experimentais, como o rastreador ocular, por exemplo, que poderia revelar por meio do tempo de fixação e dos padrões de regressão, as palavras ou estruturas com processamento mais complexo em L2. Outra possibilidade é a ampliação da investigação destas estruturas, estudando o processamento oral e a produção destes sintagmas, tanto de forma

escrita quanto oral.

Em conclusão, esta tese pode contribuir para a ampliação da discussão acerca do processamento de L2 em mentes bilíngues, especialmente do processamento sintático de sintagmas nominais por bilíngues tardios e oferece importantes implicações práticas para o ensino de L2. Ao evidenciar a interação entre o processamento de estruturas sintáticas em L2, o efeito cognato e o nível de proficiência em L2, esta tese pode contribuir com o aprimoramento de modelos de desenvolvimento linguístico para bilíngues. Ademais, possibilita o aprofundamento das variáveis aqui exploradas em futuras pesquisas, a fim de colaborar mais ainda com a compreensão dos processos cognitivos na mente bilíngue e da aprendizagem de uma segunda língua.

REFERÊNCIAS

- AGRESTI, Alan. **Categorical data analysis**. 2nd ed. New Jersey: Wiley, 2002, 372 p.
- ANTÓN, Éneko; DUÑABEITIA, Jon Andoni. Better to be alone than in bad company: cognate synonyms impair word learning. **Behavioral Sciences**, Basileia, v. 10, n. 23, 2020.
- BATES, Douglas; MÄCHLER, Martin; BOLKER, Ben; WALKER, Steve. Fitting linear mixed-effects models using lme4. **Journal of Statistical Software**, Los Angeles, v. 67, n. 1, p. 1 - 48, 2015.
- BERNOLET, Sarah; HARTSUIKER, Robert J.; PICKERING, Martin J. From language-specific to shared syntactic representations: the influence of second language proficiency on syntactic sharing bilinguals. **Cognition**, Amsterdam, v. 127, n. 3, p. 287-306, 2013.
- BERNOLET, Sarah; HARTSUIKER, Robert J. Syntactic representations in late learners of a second language: a learning trajectory. In: MILLER, David; BAYRAM, Fatih; ROTHMAN, Jason; SERRATRICE, Ludovica. **Bilingual cognition and language: the state of the science across its subfields**. New York: John Benjamins, p. 205-224, 2018.
- BLOOMFIELD, Leonard. **Language**. London: George Allen and Unwin Ltd., 1933, 566 p. Disponível em: <https://philarchive.org/archive/BLOLAO>. Acesso em: 23 abr. 2023.
- BORÉM, Sandro Almeida. **The translation process of Brazilian Portuguese – English cognate words**. Orientadora: Pamela Freitas Pereira Toassi. 2023. 130 f. Dissertação (Mestrado em Estudos da Tradução) – Faculdade de Letras, Universidade Federal do Ceará, Fortaleza, 2023.
- BULTENA, Sybrine; DIJKSTRA, Ton; VAN HELL, Janet G. Cognate effects insentence context depend on word class, L2 proficiency and task. **The Quarterly Journal of Experimental Psychology**, Londres, n. 67, v. 6, 2014.
- CALABRASE, Rita. "Living on the edge of two languages": Le costruzioni possessive in "In the Second Person" di Smaro Kamboureli. In: PALUSCI, Oriana. **Traduttrici: female voices across languages**. Trento: Tangram Edizioni Scientifiche, p. 177-189, 2011.
- CARDOSO, Lídia Amélia de Barros. **Diversidade lexical e níveis de proficiência (B2 e C1) em português como língua adicional**. Orientadora: Rosemeire Selma Monteiro Plantin. 2016. 712 f. Tese (Doutorado em Linguística) – Faculdade de Letras, Universidade Federal do Ceará, Fortaleza, 2016.
- CASAPONSA, Aina; ANTÓN, Éneko; PÉREZ, Alejandro; DUÑABEITIA, Jon Andoni. Foreign language comprehension achievement: insights from the cognate facilitation effect. **Frontiers in Psychology**, Lausanne, v. 6, 2015.
- CASAPONSA, Aina; DUÑABEITIA, Jon Andoni. Lexical organization of language-ambiguous and language-specific words in bilinguals. **The Quarterly Journal of Experimental Psychology**, Londres, v. 69, n. 3, p. 589-604, 2016.

CASAPONSA, Aina; THIERRY Guillaume; DUÑABEITIA, Jon Andoni. The Role of Orthotactics in Language Switching: An ERP Investigation Using Masked Language Priming. **Brain Sciences**, Basilea, v. 10, n. 1, 2020.

CHANG, Franklin; DELL, Gary S.; BOCK, Kathryn. Becoming syntactic. **Psychological Review**, Whashington, D.C., v. 113, n. 2, p. 234-272, 2006.

CHRISTENSEN, Larry B.; JOHNSON, Burke; TURNER, Lisa A. **Research methods, design and analysis**. 12th ed. South Alabama: Pearson, 2013, 552p.

COSTA, Albert; SANTESTEBAN, Mikel; IVANOVA, Iva. How do highly proficient bilinguals control thier lexicalization process? Inhibitory and language specific selection mechanisms are both functional. **Journal of Experimental Psychology: Learning, Memory and Cogntion**, Whashington, D.C., v. 32, n. 5. P. 1057-1074, 2006.

COUNCIL OF EUROPE. **Common European Framework of Reference for Languages: learning, teaching, assessment - companion volume**. Strasbourg: Council of Europe Publishing, 2020, 278p.

COWLES, H. Wind. **Psycholinguistics 101**. New York: Springer Publishing Company, 2011, 209 p.

CRYSTAL, David. **A Dictionary of Linguistics and Phonetics**. Oxford: Wiley-Blackwell, 2008, 560 p.

DE CAT, Cecile; KLEPOUSNIOTOU, Ekaterini; BAAYEN; R. Harald. Representational deficit or processing effect? An electrophysiological study of noun-noun compound processing by very advanced L2 speakers of English. **Frontiers in Psychology**, Lausanne, v. 6, 2015.

DE GROOT, Annette M. B. Determinants of word translation. **Journal of Experimental Psychology: Learning, Memory and Cognition**, Washington, D.C., v. 18, n. 5, p. 1001 –1018, 1992.

DE GROOT, Annette M. B.; VAN HELL, Janet. G. The learning of foreign language vocabulary. In: KROLL, Judith F.; DE GROOT, Annette M. B. **Handbook of Bilingualism: Psycholinguistics Approaches**. Oxford: Oxford University Press, p. 09-29, 2005.

DE GROOT, Annette M. B. **Language and cognition in bilinguals and multilinguals: an introduction**. New York: Psychology Press, 2011, 529 p.

DIJKSTRA, Ton; HEUVEN, Walter J. B. van. The BIA model and bilingual word recognition. In: GRAINGER, Jonathan; JACOBS, Arthur (Ed.) **Localist Connectionist Approaches to Human Cognition**. 1st edition. New Jersey: Lawrence Erlbaum, p. 189-225, 1998.

DIJKSTRA, Ton; GRAINGER, Jonathan; HEUVEN, Walter J. B. van. Recognition of cognates and interlingual homographs: the neglected role of phonology. **Journal of Memory and Language**, Amsterdam, v. 41, p. 496-518, 1999.

DIJKSTRA, Ton; HEUVEN, Walter J. B. van. The architecture of the bilingual word recognition system: from identification to decision. **Bilingualism: Language and Cognition**, Cambridge, v. 5, n. 3, p. 175-197, 2002.

DIJKSTRA, Ton; MIWA, Koji; BRUMMELHUIS, Bianca; SAPPELLI, Maya; BAAYEN, Harald. How cross-language similarity and task demands affect cognate recognition. **Journal of Memory and Language**, Amsterdam, v. 62, p. 284-301, 2010.

DIJKSTRA, Ton; VAN HELL, Janet G.; BRENDERS, Pascal. Sentence context effects in bilingual word recognition: cognate status, sentence language and semantic constraint. **Bilingualism: Language and Cognition**, Cambridge, v. 18, n. 4, p. 597-613, 2015.

DIJKSTRA, Ton; WAHL, Alexander; BUYTENHUIJS, Franka; VAN HALEM, Nino; AL-JIBOURI, Zina; KORTE, Marcel de; REKKÉ, Steven. Multilink: a computational model for bilingual word recognition and word translation. **Bilingualism: Language and Cognition**, Cambridge, v. 22, n. 4, p. 657-679, 2018.

DÖRNYEI, Zoltán. **Questionnaires in second language research: construction, administration and processing**. London: Lawrence Erlbaum Associates, 2003, 165 p.

DOWNING, Angela. **English grammar: a university course**. 3rd revised edition. New York: Routledge, 2015, 551 p.

DUÑABEITIA, Jon Andoni; PEREA, Manuel; CARREIRAS, Manuel. Masked translation priming effects with highly proficient simultaneous bilinguals. **Experimental Psychology**, Göttingen, v. 57, n. 2, p. 98-107, 2010.

ELGORT, Irina; SIYANOVA-CHANTURIA, Anna; BRYSSBAERT, Marc (Ed.). **Cross-language influences in bilingual processing and second language acquisition**. Amsterdam: John Benjamins Publishing, 2015, 321 p.

ESTIVALET, Gustavo Lopez. **Léxico do Português Brasileiro – LexPorBR**. 2019. Disponível em <https://www.lexicodoportugues.com/index.php>. Acesso em: 03 jun 2023.

FERNÁNDEZ, Eva M.; CAIRNS, Helen Smith. Wind. **Fundamentals of Psycholinguistics**. Oxford: Wiley-Blackwell, 2011, 332 p.

FINGER, Ingrid; ORTIZ-PREUSS Elena. A Psicolinguística do Bilinguismo: estudando o processamento linguístico e cognitivo bilíngue. In: ORTIZ-PREUSS Elena; FINGER, Ingrid (Org.). **A dinâmica do processamento bilíngue**. Campinas: Pontes Editores, p. 31 – 57, 2018.

FIorentino, Robert; POEPEL, David. Compound words and structure in the lexicon. **Language and Cognitive Processes**, Londres, v. 22, n. 7, p. 953 – 1000, 2007.

FREITAS, John Morais de. **O processamento lexical de homógrafos interlinguísticos por bilíngues em tarefas de reconhecimento de tradução**. Orientadora: Pamela Freitas Pereira Toassi. 2023. 129 f. Dissertação (Mestrado em Estudos da Tradução) – Faculdade de Letras, Universidade Federal do Ceará, Fortaleza, 2023.

GADELHA, Liana Maria da Silva. **Efeito de priming no processo tradutório de palavras homógrafas interlinguísticas português brasileiro - inglês**. Orientadora: Pamela Freitas Pereira Toassi. 2021. 214 f. Dissertação (Mestrado em Estudos da Tradução) – Faculdade de Letras, Universidade Federal do Ceará, Fortaleza, 2021.

GLENDAY, Candice Helen; LINS JÚNIOR, José Raymundo Figueiredo; FERRARI NETO, José. Aquisição de língua adicional e bilinguismo na perspectiva gerativista: transferência e influência do fator etário. **Macabéa**, Fortaleza, v. 9, n. 4, p. 126 – 145, 2020.

GODFROID, Aline; HOPP, Holger (Ed.). **The Routledge Handbook of Second Language Acquisition and Psycholinguistics**. New York: Routledge, 2023, 479 p.

GRAINGER, Jonathan; MIDGLEY, Katherine; HOLCOMB, Philip J. Rethinking the bilingual interactive activation model from a developmental perspective (BIA-d). In: KAIL, Michèle; HICKMAN, Maya. **Language acquisition across linguistic and cognitive systems**. New York: John Benjamins, p. 267-284, 2010.

GREEN, Tamara M. **The Greek & Latin Roots of English**. 6th edition. London: Rowman & Littlefield, 2020, 290 p.

GROSJEAN, François. Studying bilinguals: methodological and conceptual issues. In: **Bilingualism: Language and Cognition**, Cambridge, v. 1, n. 2, p. 131 – 149, 1998.

GROSJEAN, François. **Studying bilinguals**. New York: Cambridge University Press, 2008, 323 p.

GROSJEAN, François; LI, Ping. **The Psycholinguistics of Bilingualism**. Oxford: Wiley-Blackwell, 2013, 258 p.

GUASCH, Marc; BOADA, Roger; FERRÉ, Pilar; SÁNCHEZ-CASAS, Rosa. NIM: A web-based Swiss army knife to select stimuli for psycholinguistic studies. **Behavior Research Methods**, Nova Iorque, v. 45, p. 765-771, 2013.

GUIMARÃES, Mara Passos. Mecanismos de aprendizagem distribucional em bilíngues tardios. In: OLIVEIRA, Cândido Samuel Fonseca de; SÁ, Thaís Maíra Machado de (Org.). **Psicolinguística em Minas Gerais**. Contagem: CEFETMG, p. 191-206, 2020.

HALL, Christopher J. The automatic cognate form assumption: evidence for the parasitic model of vocabulary development. **IRAL**, Berlim, v. 40, p. 69-87, 2002.

HALL, Christopher J; NEWBRAND, Denise; ECKE, Peter; SPERR, Ulrike; MARCHAND, Vanessa; HAYES, Lisa. Learners' implicit assumptions about syntactic frame in new L3 words: the role of cognates, typological proximity and L2 status. **Language Learning**, Hoboken, v. 59, p. 153-202, 2009.

HANSEN EDWARDS, Jette G.; ZAMPINI, Mary L. (Eds.) **Phonology in second language acquisition**. Amsterdam: John Benjamin Publishing, 2008, 380p.

HARTSUIKER, Robert J.; PICKERING, Martin J.; VELTKAMP, Eline. Is syntactic separate or shared between languages? **Psychological Science**, Washington, D.C., v. 15, n. 6, p. 409-

414, 2004.

HARTSUIKER, Robert J.; PICKERING, Martin J. Language integration in bilingual sentence production. **Acta Psychologica**, Amsterdam, v. 128, p. 479-489, 2008.

HARTSUIKER, Robert J.; BEERTS, Saskia; LONCKE, Maïke; DESMET, Timothy; BERNOLET, Sarah. Cross-linguistic structural priming in multilinguals: further evidence for shared syntax. **Journal of Memory and Language**, Amsterdam, v. 90, p. 14-30, 2016.

HARTSUIKER, Robert J.; BERNOLET, Sarah. The development of shared syntax in second language learning. **Bilingualism**, Cambridge, v. 20, n. 2, p. 219-234, 2017.

HUDDLESTON, Rodney; PULLUM, Geoffrey K.; REYNOLDS, Brett. **A Student's Introduction to English Grammar**. 2nd edition. Cambridge: Cambridge University Press, 2022, 422 p.

HUNT, Alan; BEGLAR, David. A framework for developing EFL reading vocabulary. **Reading in a Foreign Language**, Honolulu, v. 17, n. 1, p. 23-59, 2005.

HWANG, Heeju; SHIN, Jeong-Ah; HARTSUIKER, Robert J. Late bilinguals share syntax unsparingly between L1 and L2: evidence from crosslinguistically similar and different constructions. **Language Learning**, Hoboken, v. 68, n. 1, p. 177 – 205, 2018.

JESUS, Daniela Brito de. **The Effect of L2 Proficiency on the Declarative and Procedural Memory Systems of Bilinguals**: a psycholinguistic study. Orientadora: Mailce Borges Mota. 2012. 134 f. Dissertação (Mestrado em Letras) – Faculdade de Letras, Universidade Federal de Santa Catarina, Florianópolis, 2012.

KEATING, Gregory D.; JEGERSKI, Jill. Experimental designs in sentence processing research: a methodological review and user's guide. **Studies in Second Language Acquisition**, Cambridge, v. 37, p. 1-32, 2015.

KENEDY, Eduardo. **Curso básico de linguística gerativa**. São Paulo: Contexto, 2013, 299 p.

KIM, Jonathan; GABRIEL, Ute; GYGAX, Pascal. Testing the effectiveness of the Internet-based instrument PsyToolkit: a comparison between web-based (PsyToolkit) and lab-based (E-prime 3.0) measurements of response choice and response time in a complex psycholinguistic task. **Plos One**, São Francisco, n. 14, v. 9, 2019.

KROLL, Judith F.; STEWART, E. Category interference in translation and Picture naming: evidence for asymmetric connections between bilingual memory representations. **Journal of Memory and Language**, Amsterdam, v. 33, n. 2, p. 149-174, 1994. Disponível em: <https://www.proquest.com/openview/23bb358d29db119c242de03cab7e6e90/1?pq-> Acesso em: 24 mai. 2023.

LAN, Ge; LUCAS, Kyle; SUN, Yachao. Does L2 writing proficiency influence noun phrase complexity? A case analysis of argumentative essays written by Chinese students in a first-year composition course. **System**, Amsterdam, v. 85, 2019.

LANGACKER, Ronald W. Strategies of clausal possession. **International Journal of English Studies**, San Diego, v. 3, n. 2, p.1-24, 2003.

LANGACKER, Ronald W. Possession, location and existence. *In*: LANGACKER, Ronald. (Ed.). **Investigations in Cognitive Grammar**. Berlin: Mouton de Gruyter, p.81-108, 2009.

LEVSHINA, Natalia. **How to do Linguistics with R**: data exploration and statistical analysis. Amsterdam: John Benjamin, 2015, 455 p.

LI, Man; JIANG, Nan; GOR, Kira. L1 and L2 processing of compound words: evidence from masked priming experiments in English. **Bilingualism: Language and Cognition**, Cambridge, v. 20, n. 2, p. 384 – 402, 2017.

LIMA JÚNIOR, Ronaldo Manguiera; ALVES, Ubiratã Kickhöfel; KUPSKE, Felipe Flores. Introdução a pesquisas de sons não nativos. *In*: KUPSKE, Felipe Flores; ALVES, Ubiratã Kickhöfel; LIMA JÚNIOR, Ronaldo Manguiera. **Investigando os sons de línguas não nativas**: uma introdução. Campinas: Editora da Abralín, 2021. Disponível em <https://editora.abralin.org/wp-content/uploads/2021/09/Investigando-os-sons-de-linguas-nao-nativas.pdf>. Acesso em: 14 jul 2022.

LIMA JÚNIOR, Ronaldo Manguiera. **Introdução ao R e RStudio** – por que o R? Youtube, 05 de outubro de 2021. 13min49s. Disponível em https://www.youtube.com/watch?v=sjX4rPfA_H0. Acesso em: 12 out 2024.

LIMA JÚNIOR, Ronaldo Manguiera. **Análise quantitativa de dados em Linguística**: modelos de efeitos mistos. Youtube, 18 de janeiro de 2022. 1h01min25s. Disponível em https://www.youtube.com/watch?v=eYRquH4E_1Q. Acesso em: 16 jan 2025.

LIMBERGER, Bernardo Kolling; TOASSI, Pamela Freitas Pereira; KLEIN, Angela Ines (Org.) **Bilinguismo e Leitura**: Contribuições da Psicolinguística. Campinas: Pontes Editores, 2023, 452 p.

LIU, Shun; HONG, Danping; HUANG, Jian; WANG, Suiping; LIU, Xiqin; BRANIGAN, Holly P.; PICKERING, Martin J. Syntactic priming across highly similar languages is not affected by language proficiency. **Language, Cognition and Neuroscience**, Londres, v. 37, n. 4, p. 469-480, 2022.

MAIA, Marcus (Org.). **Psicolinguística e educação**. Campinas: Editora Mercado de Letras, 2018, 260 p.

META. **Lexique**. 2023. Disponível em <http://www.lexique.org/>. Acesso em: 15 ago 2023.

MIRANDA, Daniele Lima. **A tradução de sintagmas nominais por bilíngues com inglês como L2**. Orientadora: Pamela Freitas Pereira Toassi. 2021. 127 f. Dissertação (Mestrado em Estudos da Tradução) – Faculdade de Letras, Universidade Federal do Ceará, Fortaleza, 2021.

MIRANDA, Daniele Lima; TOASSI, Pamela Freitas Pereira. A tradução de sintagmas nominais: um estudo experimental. **Letrônica**, Porto Alegre, v. 16, n. 1, 2023.

MIRANDA, Daniele Lima; TEIXEIRA, Leonardo Antônio Silva; TOASSI, Pamela Freitas

Pereira. A produção de SN1 de SN2 em inglês – L2 por aprendizes brasileiros: um estudo psicolinguístico. **Veredas: Revistas de Estudos Linguísticos**, Juiz de Fora, v. 27, n. 1, 2023.

MURPHY, Raymond. **Essential Grammar in Use**. 4th ed. Cambridge: Cambridge University Press, 2015, 320 p.

NATION, Paul. **Learning vocabulary in another language**. Cambridge: Cambridge University Press, 2001, 640 p.

ORTIZ-PREUSS, Elena. **Acesso lexical e produção de fala em bilíngues português-espanhol e espanhol-português**. Orientadora: Ingrid Finger. 2011. 183 f. Tese (Doutorado em Letras) – Faculdade de Letras, Universidade Federal do Rio Grande do Sul, Porto Alegre, 2011.

PAVLENKO, Aneta. **The bilingual mental lexicon: interdisciplinary approaches**. Bristol: Multilingual Matters, 2009, 272 p.

PICKERING, Martin J.; BRANIGAN, Holly P. The representation of verbs: evidence from syntactic priming in language production. **Journal of Memory and Language**, Amsterdam, v. 39, n. 4, p. 633-651, 1998.

POORT, Eva D.; RODD, Jennifer M. The cognate facilitation effect in bilingual lexical decision is influenced by stimulus list composition. **Acta Psychologica**, Amsterdam, v. 180, p. 52- 63, 2017.

QIAN, David D.; LIN, Linda H. F. The relationship between vocabular knowledge and language proficiency. In: WEBB, Stuart (Ed.). **The Routledge Handbook of Vocabulary Studies**. New York: Routledge, p. 66-80, 2020.

R CORE DEVELOPMENT TEAM. **R: A Language and Environment for Statistical Computing**. Viena, 2024, 3.917 p.

RESENDE, Natália; COWAN, Benjamin; WAY, Andy. MT syntactic priming effects on L2 English speakers. In: **Proceedings of the 22nd Annual Conference of the European Association for Machine Translation**. Lisboa: Creative Commons, p. 245-253, 2020.

ROEMBKE, Tanja; KOCH, Iring; PHILIPP, Andrea. What makes a cognate? Implications for research on bilingualism. **Bilingualism: Language and Cognition**, Cambridge, v. 27, n. 3, p. 1 – 6, 2024. Disponível em <https://www.cambridge.org/core/journals/bilingualism-language-and-cognition/article/what-makes-a-cognate-implications-for-research-on-bilingualism/B925D629EB763E01FB1769AD61E40A88>. Acesso em: 31 ago 2024.

ROMANO, Francesco. Syntactic planning of English genitives in L1 and L2 production: the animacy rule model. **Lingua**, Amsterdam, v. 184, p. 104-121, 2016.

SÁ, Thaís Maíra Machado de; CIRÍACO, Larissa; GODOY, Mahayana. Julgamento de aceitabilidade: um método de fácil acesso a dados quantitativos. In: OLIVEIRA, Cândido Samuel Fonseca de; SÁ, Thaís Maíra Machado de. **Métodos experimentais em Psicolinguística**. 1^a edição. São Paulo: Pá de Palavra, p. 27-39, 2022.

SAER, D. J. The effect of bilingualism on intelligence. **British Journal of Psychology**, Hoboken, v. 14, p. 25-38, 1923. Disponível em: <https://bpspsychub.onlinelibrary.wiley.com/doi/epdf/10.1111/j.2044-8295.1923.tb00110.x>. Acesso em: 23 abr. 2023.

SANTOS, Carolina Piechotta Martins; ALONSO, Karen Braga. A polissemia da construção relacional binominal SN1 de SN2 no Português Brasileiro. **Revista Estudos da Linguagem**, Belo Horizonte, v. 30, n. 2, p. 808-842, 2022.

SCHEPENS, Job; DIJKSTRA, Ton; GROOTJEN, Franc. Distributions of cognates in Europe as based on Levenshtein distance. **Bilingualism: Language and Cognition**, Cambridge, v. 15, n. 1, p. 157-166, 2012.

SCHMITT, Norbert; SCHMITT, Diane; CLAPHAM, Caroline. Developing and exploring the behavior of two new versions of the word frequency profile. **Applied Linguistics**, Oxford, v. 22, n. 1, p. 56 – 57, 2001.

SCHOONBAERT, Sofie; DUYCK, Wolter; BRYSAERT, Marc; HARTSUIKER, Robert J. Semantic and translation priming from a first language to a second and back: making sense of the findings. **Memory and Cognition**, Nova Iorque, v. 37, n. 5, p. 569–586, 2009.

SCHÜLTZE, Carson T. **The empirical base of linguistics: grammaticality judgements and linguistic methodology**. Chicago: University of Chicago Press, 1996, 244 p. Disponível em <https://escholarship.org/content/qt05b2s4wg/qt05b2s4wg.pdf?t=qaylsq>. Acesso em: 19 mai. 2023.

SCHWARTZ, Ana I.; KROLL, Judith F. Bilingual lexical activation in sentence context. **Journal of Memory and Language**, Amsterdam, v. 55, p. 197-212, 2006. Disponível em Bilinguallexicalactivationinsentencecontext–ScienceDirect. Acesso em: 13 mai. 2023.

SHENG, Li; LAM, Boji Pak Wing; CRUZ, Diana; FULTON, Aislynn. A robust demonstration of the cognate facilitation effect in first language and second-language naming. **Journal of Experimental Child Psychology**, Amsterdam, v. 141, p. 229-238, 2016.

SILVA, Giselli Mara da Silva; SOUZA, Ricardo Augusto de; VALADARES, Marcus Guilherme Pinto de Faria. Questionários em pesquisas na área de Bilinguismo e aquisição de segunda língua: elaboração e uso para perfilamento de participantes. In: OLIVEIRA, Cândido Samuel Fonseca de; SÁ, Thaís Maíra Machado de. **Métodos experimentais em Psicolinguística**. 1ª edição. São Paulo: Pá de Palavra, p. 13-26, 2022.

SOARES, Ana Paula; OLIVEIRA, Helena Mendes; COMESAÑA, Montserrat; COSTA, Ana Santos. Lexico-syntactic interactions in the resolution of relative clause ambiguities in a second language (L2): the role of cognate status and L2 proficiency. **Psicológica**, Valência, v. 39, p. 164-197, 2018.

SOUSA, Letícia Rodrigues de. **Portuguese-English interlingual homophones in word reading and translation**. Orientadora: Pamela Freitas Pereira Toassi. 2024. 114 f. Dissertação (Mestrado em Linguística) – Faculdade de Letras, Universidade Federal do Ceará, Fortaleza, 2024.

SOUZA, Ricardo Augusto de; OLIVEIRA Cândido Samuel Fonseca de; SOARES- SILVA, Jesiel; PENZIN, Alberto Gallo Araújo; SANTOS, Alexandre Alves. Estudo sobre um parâmetro de tarefa e um parâmetro amostral para experimentos com julgamentos de aceitabilidade temporalizados. **Revista Estudos da Linguagem**, Belo Horizonte, v. 23, n. 1, p. 211-244, 2015.

SOUZA, Ricardo Augusto de. A proficiência em L2 como objeto da Psicolinguística. *In*: MOTA, Mailce Borges; NAME, Cristina. (Org.). **Interface linguagem e cognição: contribuições da Psicolinguística**. 1ª ed. Tubarão: Copiart, p. 201-218, 2019.

SOUZA, Ricardo Augusto de. Os bilíngues interessam aos linguistas? Uma resposta da Psicolinguística. *In*: OLIVEIRA, Cândido Samuel Fonseca de; SÁ, Thaís Maíra Machado de (Org.). **Psicolinguística em Minas Gerais**. Contagem: CEFETMG, p. 191-206, 2020.

STÆHR, Lars Stenius. Vocabulary size and skills of listening, reading and writing. **The Language Learning Journal**, Londres, v. 36, n. 2, p. 139-152, 2008.

STOET, Gijsbert. PsyToolKit: a software package for programming psychological experiments using Linux. **Behavior Research Methods**, Nova Iorque v. 42, n. 4, p. 1096-1104, 2010.

STOET, Gijsbert. PsyToolKit: a novel web-based method for running online questionnaires and reaction-time experiments. **Teaching of Psychology**, Thousand Oaks, v. 44, n. 1, p.24-31, 2017.

SUNDERMAN, Gretchen. Translation recognition tasks. *In*: JEGERSKI, Jill; VANPATTEN, Bill (Ed.). **Research methods in second language Psycholinguistics**. New York: Taylor & Francis, p. 183 – 211, 2014.

TOASSI, Pamela Freitas Pereira; PEREIRA, Silvia Hedine de Albuquerque. Understanding cognate words in a first contact with English. **Caderno de Letras**, Niterói, n. 35, p. 27-44, 2019. Disponível em: <https://periodicos.ufpel.edu.br/index.php/cadernodeletras/issue/view/887>. Acesso em: 01 mai. 2023.

TOASSI, Pamela Freitas Pereira; MOTA, Mailce Borges; TEIXEIRA, Elisângela Nogueira. The effect of cognate words on lexical access of English as a third language. **Caderno de Traduções de Florianópolis**, Florianópolis, v. 40, n. 2, p. 74-96, 2020. Disponível em <https://periodicos.ufsc.br/index.php/traducao/article/view/78396/44781>. Acesso em: 11 jun. 2023.

TOKOWICZ, Natasha. **Lexical processing and second language acquisition**. New York: Routledge, 2015, 138 p.

TOMASELLO, Michael. First steps toward a usage-based theory of language acquisition. **Cognitive Linguistics**, Berlim, v. 11, n. 1-2, p. 61-82, 2000. Disponível em <https://www.degruyter.com/document/doi/10.1515/cogl.2001.012/html>. Acesso em: 24 mai. 2023.

TRAXLER, Matthew J. **Introduction to Psycholinguistics: Understanding Language**

Science. Oxford: Wiley-Blackwell, 2012, 568 p.

VAN HEUVEN, Walter J. B.; MANDERA, Pavel; KEULEERS, Emmanuel; BRYSSBAERT, Marc. SUBTLEX-UK: A new and improved word frequency database for British English. **The Quarterly Journal of Experimental Psychology**, Londres, v. 67, n. 6, p. 1176 – 1190, 2014.

WARREN, Paul. **Introducing Psycholinguistics**. New York: Cambridge University Press, 2013, 290 p.

WEI, Li. **The bilingualism reader**. Routledge: London, 2000, 592 p.

APÊNDICE A - TERMO DE CONSENTIMENTO LIVRE E ESCLARECIDO (TCLE)

O menor sob sua responsabilidade está sendo convidado por Daniele Lima Miranda, doutoranda do programa de Linguística na Universidade Federal do Estado do Ceará (UFC), orientanda da professora Dra. Pâmela Pereira Freitas Toassi, como participante da pesquisa intitulada “PROCESSAMENTO DE SENTENÇAS PORTUGUÊS - INGLÊS”. Leia atentamente as informações abaixo, para que todos os procedimentos desta pesquisa sejam esclarecidos.

Caro (a) Senhor (a) Responsável,

Gostaria de convidar o menor sob sua responsabilidade a participar do estudo da doutoranda Daniele Lima Miranda, que busca investigar como se dá o processamento de sentenças em português-inglês. Os estudos nessa área visam não só compreender os processos envolvidos no processamento de uma língua estrangeira, mas também desenvolver meios de aperfeiçoar o processo de ensino- aprendizagem dessa língua. Peço que você leia este termo de consentimento e tire todas as dúvidas que possam surgir (através do e-mail danielelimamiranda@gmail.com) antes de concordar em participar do estudo.

Objetivo do estudo

O objetivo geral deste estudo é investigar o processamento de sentenças em português-inglês.

Procedimentos

Se for permitida a participação do menor sob sua responsabilidade neste estudo, ele(a) será solicitado a responder um (01) questionário demográfico e linguístico e linguístico, a participar de três (03) tarefas experimentais e a responder um (01) teste de conhecimento de vocabulário de inglês. Todos estes instrumentos serão disponibilizados no momento oportuno, através de links online. Ele(a) poderá realizar as atividades desta pesquisa no CCI, utilizando um computador de mesa ou um notebook com acesso à internet. Não é possível realizar esta pesquisa através de tablet ou celular. A pesquisa se desdobrará em sete (07) etapas, especificadas abaixo:

1ª etapa: Leitura do TCLE (Termo de Consentimento Livre e Esclarecido)

Se você concordar com a participação do menor sob sua responsabilidade neste estudo, deverá preencher os seus dados e os dados do seu (sua) filho(a), assinalar a opção “Aceito” e assinar. Então, entramos em contato com o menor sob sua responsabilidade para que

ele leia o termo de assentimento. Caso ao final da leitura desse TCLE ainda tiver dúvidas sobre a pesquisa, esclareça-as primeiramente com a pesquisadora principal desta pesquisa, através do e-mail danielelimamiranda@gmail.com.

2ª etapa: Leitura do TALE (Termo de Assentimento Livre e Esclarecido)

Se o menor sob sua responsabilidade concordar com sua participação neste estudo, ele (ela) deverá preencher seus dados, assinalar a opção “Aceito” ao final do formulário e depois pressionar “enter”. Então ele(a) será encaminhado para a próxima etapa da pesquisa. Caso ao final da leitura desse TALE, ele ainda tiver dúvidas sobre a pesquisa, deverá esclarecê-las primeiramente com a pesquisadora principal desta pesquisa, através do e-mail danielelimamiranda@gmail.com.

3ª etapa: Questionário demográfico e linguístico e linguístico (tempo total estimado: 10 minutos)

Esta etapa consistirá no preenchimento de um questionário que contém questões de múltipla escolha e/ou de resposta curta para coletar dados pessoais e para investigar o histórico de aprendizagem da língua estrangeira de seu (sua) filho(a).

4ª etapa: Tarefa 1 (tempo total estimado: 20 minutos)

Nesta etapa, o menor sob sua responsabilidade terá acesso a setenta e cinco (75) pares de sentenças em português e em inglês e deverá decidir se as sentenças são equivalentes ou não.

5ª etapa: Tarefa 2 (tempo total estimado: 20 minutos)

Nesta etapa, o menor sob sua responsabilidade terá acesso a setenta e cinco (75) frases em inglês e deverá decidir se a sentença está bem construída ou não.

6ª etapa: Tarefa 3 (tempo total estimado: 20 minutos)

Nesta etapa, o menor sob sua responsabilidade terá acesso a noventa (90) sentenças em português e deverá decidir entre três (03) opções apresentadas qual é o item que contém a melhor tradução.

7ª etapa: Teste de conhecimento de vocabulário em inglês (tempo total estimado: até 30 minutos)

Para certificar o nível de conhecimento de vocabulário de inglês, o menor sob sua responsabilidade será solicitado a realizar um teste de vocabulário em inglês (através de um link informado ao final das tarefas). Neste teste, ele(a) deve relacionar as palavras às definições apresentadas. O resultado do teste é disponibilizado imediatamente após o seu término.

Estima-se que o tempo total da pesquisa, compreendendo as sete (07) fases, será em torno de duas (02) horas. Ele (ela) poderá fazer intervalos para descansar entre uma etapa e outra. Ele(a) também poderá interromper a sua participação no estudo a qualquer momento. A participação dele(a) nas tarefas desse estudo será voluntária e contribuirá para o entendimento do processamento de sintagmas nominais em inglês. Durante a pesquisa, ele(a) terá a oportunidade de praticar o inglês e, também, terá uma avaliação do seu nível de conhecimento da língua, o que pode certificar sua proficiência em inglês.

Riscos

Toda investigação com a participação de seres humanos, ainda que seja realizada em documentos, é passível de riscos. No caso específico desse estudo, trata-se de um risco mínimo, que poderá ser o cansaço proveniente do preenchimento de um (01) questionário demográfico e linguístico e linguístico, da realização de três (03) tarefas experimentais e de um (01) teste de conhecimento de vocabulário de inglês em formato eletrônico. No entanto, o menor sob sua responsabilidade é livre para interromper o experimento, a qualquer momento, sem que haja nenhum prejuízo a ele. Além disso, todas as etapas da pesquisa são curtas e ele (ela) poderá fazer intervalos para descansar.

Outrossim, há limitação por parte do pesquisador de assegurar total confidencialidade. Isto ocorre por se tratar de uma pesquisa realizada de forma remota em ambientes virtuais, existindo mínimo risco de violação deste estudo.

Benefícios

Um benefício direto da pesquisa será a avaliação do nível de vocabulário em inglês do menor sob sua responsabilidade. Você poderá obter os resultados das três (03) tarefas realizadas, entrando em contato por e-mail, informando seu código de identificação, aquele criado por ele(a), contendo duas letras e dois números.

Direitos dos participantes

O menor sob sua responsabilidade é livre para decidir se deseja participar ou não

desse estudo. Como a participação é voluntária, ele(a) pode desistir a qualquer momento sem nenhum prejuízo para você. A qualquer momento ele(a) poderá recusar a continuar participando da pesquisa, sem que isso lhe traga qualquer prejuízo. Você também poderá a qualquer momento retirar o consentimento de participação do menor sob sua responsabilidade, sem que ele seja prejudicado.

Compensação financeira

Não há despesas pessoais ou compensações financeiras relacionadas à participação no estudo.

Utilização dos dados:

Os dados coletados neste estudo online serão acessados apenas pela responsável desta pesquisa e a divulgação das informações só será feita entre os profissionais estudiosos do assunto. Mesmo após os resultados se tornarem públicos, a identidade do menor sob sua responsabilidade será totalmente preservada. Não haverá nenhuma informação que leve à sua identificação. A qualquer momento você poderá ter acesso a informações referentes à pesquisa, pelo telefone da instituição e endereço de e-mail do grupo de pesquisa, o PLIBIMULT.

Endereço do responsável pela pesquisa:

Nome: Daniele Lima Miranda

Instituição: Universidade Federal do Ceará - UFC

Endereço: XXX

Telefones para contato: XXX

E-mail para contato: danielelimamiranda@gmail.com

ATENÇÃO: Se você tiver alguma consideração ou dúvida, sobre a sua participação na pesquisa, entre em contato com o Comitê de Ética em Pesquisa da UFC/PROPESQ – Rua Coronel Nunes de Melo, 1000 - Rodolfo Teófilo, fone: 3366-8346/44. (Horário: 08:00-12:00 horas de segunda a sexta-feira).

O CEP/UFC/PROPESQ é a instância da Universidade Federal do Ceará responsável pela avaliação e acompanhamento dos aspectos éticos de todas as pesquisas envolvendo seres humanos.

APÊNDICE B - TERMO DE ASSENTIMENTO LIVRE E ESCLARECIDO

Você está sendo convidado(a) como participante da pesquisa: “PROCESSAMENTO DE SENTENÇAS PORTUGUÊS - INGLÊS”.

Nesse estudo pretendemos investigar como se dá o processamento de sentenças em português-inglês. O motivo que nos leva a estudar esse assunto é compreender os processos envolvidos na aprendizagem de uma língua estrangeira, bem como também desenvolver meios de aperfeiçoar o processo de ensino-aprendizagem dessa língua.

Para este estudo adotaremos o(s) seguinte(s) procedimento(s): (1) leitura do termo de consentimento, (2) leitura do termo de assentimento, (3) um questionário demográfico e linguístico e linguístico, (4) tarefa 1, (5) tarefa 2, (6) tarefa 3 e (7) teste de conhecimento de vocabulário de inglês, conforme descrição abaixo:

1ª etapa: Leitura do TCLE (Termo de Consentimento Livre e Esclarecido)

Se seu responsável concordar com sua participação neste estudo, deverá preencher os seus do menor participante e do responsável, assinalar a opção “Aceito” e assinar. Então, entraremos em contato com você, para que ele leia o termo de assentimento. Caso ao final da leitura desse TCLE ainda tiver dúvidas sobre a pesquisa, esclareça-as primeiramente com a pesquisadora principal desta pesquisa, através do e-mail danielelimamiranda@gmail.com.

2ª etapa: Leitura do TALE (Termo de Assentimento Livre e Esclarecido)

Se o menor concordar em participar deste estudo, ele (ela) deverá preencher seus dados, assinalar a opção “Aceito” ao final do formulário e depois pressionar “enter”. Então ele(a) será encaminhado para a próxima etapa da pesquisa. Caso ao final da leitura desse TALE, ele ainda tiver dúvidas sobre a pesquisa, deverá esclarecê-las primeiramente com a pesquisadora principal desta pesquisa, através do e-mail danielelimamiranda@gmail.com.

3ª etapa: Questionário demográfico e linguístico e linguístico (tempo total estimado: 10 minutos)

Esta etapa consistirá no preenchimento de um questionário que contém questões de múltipla escolha e/ou de resposta curta para coletar dados pessoais e para investigar seu histórico de aprendizagem da língua estrangeira.

4ª etapa: Tarefa 1 (tempo total estimado: 20 minutos)

Nesta etapa, você terá acesso a setenta e cinco (75) pares de sentenças em português e em inglês e deverá decidir se as sentenças são equivalentes ou não.

5ª etapa: Tarefa 2 (tempo total estimado: 20 minutos)

Nesta etapa, você terá acesso a setenta e cinco (75) frases em inglês e deverá decidir se a sentença está bem construída ou não.

6ª etapa: Tarefa 3 (tempo total estimado: 20 minutos)

Nesta etapa, você terá acesso a noventa (90) sentenças em português e deverá decidir entre três (03) opções apresentadas qual é o item que contém a melhor tradução para o inglês.

7ª etapa: Teste de conhecimento de vocabulário em inglês (tempo total estimado: até 30 minutos)

Para certificar o nível de conhecimento de vocabulário de inglês você será solicitado a realizar um teste de vocabulário em inglês. Neste teste, você deve relacionar as palavras às definições apresentadas. O resultado do teste é disponibilizado imediatamente após o seu término.

Estima-se que o tempo total da pesquisa, compreendendo as sete (07) fases, será em torno de duas (02) horas. Você poderá fazer intervalos para descansar entre uma etapa e outra. Você também poderá interromper a sua participação no estudo a qualquer momento. A sua participação nas tarefas desse estudo será voluntária e contribuirá para a compreensão de sintagmas nominais em inglês. Durante a pesquisa, você terá a oportunidade de praticar o inglês e, também, terá uma avaliação do seu nível de conhecimento da língua.

Para participar deste estudo, o responsável por você deverá ler e aceitar um termo de consentimento. Você não terá nenhum custo, nem receberá qualquer vantagem financeira. Você será esclarecido(a) em qualquer aspecto que desejar e estará livre para participar ou recusar-se. O responsável por você poderá retirar o consentimento ou interromper a sua participação a qualquer momento. A sua participação é voluntária e a recusa em participar não acarretará qualquer penalidade ou modificação na forma em que é atendido(a) pela pesquisadora, que irá tratar a sua identidade com padrões profissionais de sigilo. Trata-se de uma pesquisa online e você não será identificado em nenhuma publicação.

Toda investigação com a participação de seres humanos, ainda que seja realizada

em documentos, é passível de riscos. No caso específico deste estudo, trata-se de um risco mínimo, que poderá ser o cansaço proveniente do preenchimento de um questionário, da realização de três (03) tarefas e do teste de conhecimento vocabular. Outrossim, há limitação por parte do pesquisador de assegurar total confidencialidade. Isto ocorre por se tratar de uma pesquisa realizada de forma remota em ambientes virtuais, existindo mínimo risco de violação deste estudo. No entanto, você é livre para interromper o experimento, a qualquer momento, sem que haja nenhum prejuízo a você. Além disso, todas as etapas da pesquisa são curtas e você poderá fazer intervalos para descansar.

Os resultados estarão à sua disposição quando forem finalizados. Seu nome ou o material que indique sua participação não será liberado sem a permissão do responsável por você. Os dados e instrumentos utilizados na pesquisa ficarão arquivados com o pesquisador responsável por um período de cinco (05) anos e, após esse tempo, serão destruídos.

Endereço do responsável pela pesquisa:

Nome: Daniele Lima Miranda

Instituição: Universidade Federal do Ceará - UFC

Endereço: XXX

Telefones para contato: XXX

E-mail para contato: danielelimamiranda@gmail.com

ATENÇÃO: Se você tiver alguma consideração ou dúvida, sobre a sua participação na pesquisa, entre em contato com o Comitê de Ética em Pesquisa da UFC/PROPESQ – Rua Coronel Nunes de Melo, 1000 - Rodolfo Teófilo, fone: 3366-8344/46. (Horário: 08:00-12:00 horas de segunda a sexta-feira).

O CEP/UFC/PROPESQ é a instância da Universidade Federal do Ceará responsável pela avaliação e acompanhamento dos aspectos éticos de todas as pesquisas envolvendo seres humanos.

APÊNDICE C - QUESTIONÁRIO DEMOGRÁFICO E LINGUÍSTICO E LINGUÍSTICO

Código com duas letras e dois números	
Nome completo	
Gênero	
Data de nascimento	
Idade	
Cidade/Estado de nascimento	
Bairro em que mora	
Telefone de contato	
Nome da escola em que estuda	
Semestre no CCI	
Turno no CCI	
Mão que utiliza para escrever	
Série em que começou a estudar inglês	
Motivo pelo qual estuda inglês	
Frequência que escreve ou digita em inglês	
Frequência que lê em inglês	

APÊNDICE D - ITENS DISTRATORES DOS EXPERIMENTOS 1 E 2

ITEM	ITENS DISTRATORES EM L1	NÚMERO DE CARACTERES	ITENS DISTRATORES EM L2	NÚMERO DE CARACTERES
61	Ela está assistindo TV	22	She is watching TV	18
62	Ela está assistindo TV	22	She is watched TV	17
63	Ele/a está crescendo	14	It is growing up	16
64	Ele/a está crescendo	14	It is grow up	13
65	Eu estou estudando todo dia	27	I am studying everyday	22
66	Eu estou estudando todo dia	27	I am study everyday	19
67	Está chovendo forte	19	It is raining hard	19
68	Está chovendo forte	19	It is rained hard	18
69	Ele está lendo agora	20	He is reading now	17
70	Ele está lendo agora	20	He is read now	14
71	Eles estão jogando futebol	26	They are playing soccer	23
72	Eles estão jogando futebol	26	They are played soccer	22
73	Nós estamos jantando	20	We are having dinner	20
74	Nós estamos jantando	20	We are had dinner	17
75	Você está escutando música	26	You are listening music	23
76	Você está escutando música	26	You are listened music	22
77	Nós estamos dormindo tarde	26	We are sleeping late	20
78	Nós estamos dormindo tarde	26	We are sleep late	17
79	Ela está escrevendo poemas	26	She is writing poems	20
80	Ela está escrevendo poemas	26	She is write poems	18
81	Ele está falando francês	24	He is speaking French	21
82	Ele está falando francês	24	He is spoke French	18
83	Você está andando rápido	26	You are walking fast	19
84	Você está andando rápido	26	You are walked fast	18
85	Está nevando pouco	21	It is snowing slightly	22
86	Está nevando pouco	21	It is snowed slightly	21
87	Eles estão lanchando	20	They are having a snack	23
88	Eles estão lanchando	20	They are have a snack	21
89	Nós estamos bebendo água	24	We are drinking water	21
90	Nós estamos bebendo água	24	We are drank water	18
91	Paul está nadando bem	21	Paul is swimming well	21
92	Paul está nadando bem	21	Paul is swim well	17
93	Anne está dançando balé	23	Anne is dancing ballet	22
94	Anne está dançando balé	23	Anne is danced ballet	21
95	Peter está comendo arroz	24	Peter is eating rice	20
96	Peter está comendo arroz	24	Peter is eat rice	17
97	Kate está falando devagar	25	Kate is talking quietly	23
98	Kate está falando devagar	25	Kate is talked quietly	22
99	Charles está vestindo jeans	27	Charles is wearing jeans	24
100	Charles está vestindo jeans	27	Charles is wear jeans	21
101	Susan está escovando dentes	27	Susan is brushing teeth	23
102	Susan está escovando dentes	27	Susan is brushed teeth	22
103	Jack está usando short	22	Jack is using shorts	20
104	Jack está usando short	22	Jack is used shorts	19

105	Pam está ensinando inglês	25	Pam is teaching English	23
-----	---------------------------	----	-------------------------	----

APÊNDICE E - FREQUÊNCIA DOS SUBSTANTIVOS DOS EXPERIMENTOS 1 E 2

ITEM	SINTAGMAS NOMINAIS EM L1/L2	ESCALA ZIPF SUBST1 EM L1	ESCALA ZIPF SUBST2 EM L1	ESCALA ZIPF SUBST1 EM L2	ESCALA ZIPF SUBST2 EM L2
1	Tamanho da blusa/ Shirt size	Tamanho 4,64	Blusa 3,49	Shirt 4,66	Size 4,66
2	Caixa de sapato/ Shoe box	Caixa 4,79	Sapato 3,76	Shoe 4,48	Box 4,95
3	Vida de casado/ Married life	Vida 5,66	Casado 4,10	Married 5,37	Life 5,90
4	Cadeira de praia/ Beach chair	Cadeira 4,44	Praia 4,87	Beach 4,75	Chair 4,69
5	Sino da igreja/ Church bell	Sino 3,36	Igreja 4,81	Church 4,84	Bell 4,59
6	Pano de chão/ Floor cloth	Pano 4,04	Chão 4,64	Floor 5,00	Cloth 3,78
7	Moldura do espelho/ Mirror frame	Moldura 3,49	Espelho 3,77	Mirror 4,38	Frame 4,14
8	Capa do livro/ Book cover	Capa 4,47	Livro 5,46	Book 5,24	Cover 4,97
9	Portão da casa/ House gate	Portão 3,98	Casa 5,66	House 5,71	Gate 4,50
10	Saco de pão/ Bread bag	Saco 4,11	Pão 4,31	Bread 4,45	Bag 4,97
11	Olhos da garota/ Girl eyes	Olhos 5,03	Garota 4,24	Girl 5,74	Eyes 5,34
12	Rabo do gato/ Cat tail	Rabo 3,59	Gato 4,04	Cat 4,82	Tail 4,37
13	Tinta da caneta/ Pen ink	Tinta 4,13	Caneta 3,86	Pen 4,39	Ink 3,87
14	Anel de prata/ Silver ring	Anel 3,78	Prata 4,26	Silver 4,50	Ring 4,96
15	Lençol de algodão/ Cotton sheet	Lençol 3,24	Algodão 4,25	Cotton 4,15	Sheet 4,06

16	Jogo de futebol/ Soccer match	Jogo 5,66	Futebol 5,42	Soccer 4,09	Match 4,69
17	Vestido de festa/ Party dress	Vestido 3,99	Festa 5,09	Party 5,36	Dress 4,93
18	Viagem de navio/ Ship trip	Viagem 5,03	Navio 4,45	Ship 4,99	Trip 4,91
19	Borracha do lápis/ Pencil eraser	Borracha 4,02	Lápis 3,76	Pencil 3,99	Eraser 3,00
20	Nível do mar/ Sea level	Nível 5,17	Mar 4,87	Sea 4,77	Level 4,71
21	Tampa da panela/ Pot lid	Tampa 3,59	Panela 3,58	Pot 4,35	Lid 3,69
22	Caminhão do lixo/ Garbage truck	Caminhão 4,56	Lixo 4,62	Garbage 4,41	Truck 4,86
23	Molho de carne/ Meat sauce	Molho 4,07	Carne 4,80	Meat 4,63	Sauce 4,19
24	Loja de brinquedo/ Toy store	Loja 4,96	Brinquedo 3,87	Toy 4,22	Store 4,91
25	Água da piscina/ Pool water	Água 5,29	Piscina 4,20	Pool 4,67	Water 5,35
26	Nariz de palhaço/ Clown nose	Nariz 4,17	Palhaço 3,56	Clown 4,19	Nose 4,84
27	Lata de milho/ Corn can	Lata 3,97	Milho 4,39	Corn 4,15	Can 6,71
28	Quadro da sala/ Room picture	Quadro 5,08	Sala 4,96	Room 5,64	Picture 5,14
29	Torta de maçã/ Apple pie	Torta 3,47	Maçã 3,67	Apple 4,37	Pie 4,45
30	Suco de uva/ Grape juice	Suco 4,16	Uva 3,53	Grape 3,60	Juice 4,42
31	Museu de arte/ Art museum	Museum 4,60	Arte 5,17	Art 4,84	Museum 4,26
32	Salada de fruta/ Fruit salad	Salada 3,84	Fruta 3,88	Fruit 4,33	Salad 4,23
33	Clipe de papel/ Paper clip	Clipe 3,92	Papel 5,27	Paper 5,01	Clip 3,75
34	Exame de rotina/ Routine exam	Exame 4,81	Rotina 4,30	Routine 4,29	Exam 4,12
35	Nome de família/ Family name	Nome 5,55	Família 5,34	Family 5,54	Name 5,80

36	Carro de polícia/ Police car	Carro 5,40	Polícia 5,45	Police 5,37	Car 5,68
37	Cartão de crédito/ Credit card	Cartão 4,80	Crédito 5,10	Credit 4,66	Card 4,93
38	Fator de risco/ Risk factor	Fator 4,63	Risco 5,14	Risk 4,69	Factor 3,86
39	Teste de história/ History test	Teste 4,71	História 5,55	History 4,92	Test 4,92
40	Barra de chocolate/ Chocolate bar	Barra 4,28	Chocolate 4,15	Chocolate 4,46	Bar 4,93
41	Medalha de bronze/ Bronze medal	Medalha 4,25	Bronze 3,94	Bronze 3,45	Medal 4,06
42	Sopa do bebê/ Baby soup	Sopa 3,92	Bebê 4,28	Baby 5,70	Soup 4,40
43	Filme de comédia/ Comedy fim	Filme 5,46	Comédia 4,43	Comedy 4,07	Movie 4,81
44	Controle da TV/ TV control	Controle 5,22	TV 5,26	TV 5,00	Control 5,11
45	Estação de trem/ Train station	Estação 4,30	Trem 4,44	Train 4,97	Station 4,89
46	Tecla do piano/ Piano key	Tecla 3,55	Piano 4,28	Piano 4,39	Key 4,93
47	Fatia de tomate/ Tomato slice	Fatia 4,04	Tomate 4,05	Tomato 3,77	Slice 3,93
48	Comida de hospital/ Hospital food	Comida 4,61	Hospital 4,99	Hospital 5,09	Food 5,18
49	Banco do parque/ Park bench	Banco 5,23	Parque 4,80	Bench 3,98	Park 4,85
50	Casca da manga/ Mango peel	Casca 3,67	Manga 3,49	Mango 3,22	Peel 3,72
51	Doce de banana/ Banana candy	Doce 4,15	Banana 3,92	Candy 4,55	Banana 4,03
52	Bolo de limão/ Lemon cake	Bolo 4,08	Limão 3,67	Lemon 4,08	Cake 4,65
53	Fio do fone/ Phone cord	Fio 4,37	Fone 3,50	Phone 5,43	Cord 6,84
54	Mesa de plástico/ Plastic table	Mesa 4,79	Plástico 4,23	Plastic 4,27	Table 5,02
55	Cabelo do estudante/	Cabelo 4,41	Estudante 4,11	Student 4,63	Hair 5,01

	Student hair				
56	Conta do dentista/ Dentist bill	Conta 5,30	Dentista 3,88	Dentist 4,04	Bill 5,07
57	Dor de estômago/ Stomach ache	Dor 4,55	Estômago 4,52	Stomach 4,52	Ache 3,39
58	Garrafa de café/ Coffee bottle	Garrafa 4,02	Café 4,95	Coffee 5,15	Bottle 4,70
59	Professor de inglês/ English teacher	Professor 5,08	Inglês 4,89	English 4,86	Teacher 4,74
60	Porta do hotel/ Hotel door	Porta 4,85	Hotel 4,92	Hotel 5,01	Door 5,46

**APÊNDICE F - NÚMEROS DE CARACTERES DOS ITENS EXPERIMENTAIS DOS
EXPERIMENTOS 1 E 2**

ITEM	SENTENÇAS EM L1 - Nº DE CARACTERES	OPÇÃO SINTÁTICA 1 EM L2 - Nº DE CARACTERES	OPÇÃO SINTÁTICA 2 EM L2 - Nº DE CARACTERES	OPÇÃO SINTÁTICA 3 EM L2 - Nº DE CARACTERES
1	O tamanho da blusa é médio - 26	The size of shirt is medium - 27	The shirt's size is medium - 26	The shirt size is medium - 24
2	A caixa de sapato é redonda - 27	The box of shoe is round - 24	The shoes' box is round - 23	The shoe box is round - 21
3	A vida de casado é ocupada - 26	The life of married is busy - 27	The married's life is busy - 26	The married life is busy - 24
4	A cadeira de praia é colorida - 29	The chair of beach is colored - 29	The beach's chair is colored - 28	The beach chair is colored - 26
5	O sino da igreja é grande - 25	The bell of church is big - 25	The church's bell is big - 24	The church bell is big - 22
6	O pano de chão está sujo - 24	The cloth of floor is dirty - 27	The floor's cloth is dirty - 26	The floor cloth is dirty - 24
7	A moldura do espelho é oval - 27	The frame of mirror is oval - 27	The mirror's frame is oval - 26	The mirror frame is oval - 24
8	A capa do livro é branca - 24	The cover of book is white - 26	The book's cover is white - 25	The book cover is white - 23
9	O portão da casa é roxo - 23	The gate of house is purple - 27	The house's gate is purple - 26	The house gate is purple - 24
10	O saco de pão está cheio - 24	The bag of bread is full - 24	The bread's bag is full - 23	The bread bag is full - 21
11	Os olhos da garota são verdes - 29	The eyes of girl are green - 26	The girl's eyes are green - 25	The girl eyes are green - 23
12	O rabo do gato é longo - 22	The tail of cat is long - 23	The cat's tail is long - 22	The cat tail is long - 20
13	A tinta da caneta é azul - 24	The ink of pen is blue - 22	The pen's ink is blue - 21	The pen ink is blue - 19
14	O anel de prata é seu - 21	The ring of silver is yours - 27	The silver's ring is yours - 26	The silver ring is yours - 24
15	O lençol de algodão é macio - 27	The sheet of cotton is soft - 28	The cotton's sheet is soft - 27	The cotton sheet is soft - 25
16	O jogo de futebol é hoje - 24	The match of soccer is today	The soccer's match is today -	The soccer match is today -

		- 28	27	25
17	O vestido de festa é dourado - 28	The dress of party is golden - 28	The party's dress is golden - 27	The party dress is golden - 25
18	A viagem de navio é amanhã - 26	The trip of ship is tomorrow - 28	The ship's trip is tomorrow - 27	The ship trip is tomorrow - 25
19	A borracha do lápis é preta - 27	The eraser of pencil is black - 29	The pencil's eraser is black - 28	The pencil eraser is black - 26
20	O nível do mar está alto - 24	The level of sea is high - 24	The sea's level is high - 23	The sea level is high - 21
21	A tampa da panela é prata - 25	The lid of pot is clean - 23	The pot's lid is clean - 22	The pot lid is clean - 20
22	O caminhão de lixo está lá - 26	The truck of garbage is there - 29	The garbage's truck is there - 28	The garbage truck is there - 26
23	O molho de carne é saboroso - 27	The sauce of meat is tasty - 26	The meat's sauce is tasty - 25	The meat sauce is tasty - 23
24	A loja de brinquedo é nova - 26	The store of toy is new - 23	The toy's store is new - 22	The toy store is new - 20
25	A água da piscina é aquecida - 28	The water of pool is warmed - 27	The pool's water is warmed - 26	The pool water is warmed - 24
26	O nariz de palhaço é vermelho - 29	The nose of clown is red - 24	The clown's nose is red - 23	The clown nose is red - 21
27	A lata de milho é amarela - 25	The can of corn is yellow - 25	The corn's can is yellow - 24	The corn can is yellow - 22
28	O quadro da sala é pequeno - 26	The picture of room is small - 28	The room's picture is small - 27	The room picture is small - 25
29	A torta de maçã é deliciosa - 27	The pie of apple is delicious - 29	The apple's pie is delicious - 28	The apple pie is delicious - 26
30	O suco de uva é doce - 20	The juice of grape is sweet - 27	The grape's juice is sweet - 26	The grape juice is sweet - 24
31	O museu de arte é perto - 23	The museum of art is near - 25	The art's museum is near - 24	The art museum is near - 22
32	A salada de fruta é boa - 23	The salad of fruit is good - 26	The fruit's salad is good - 25	The fruit salad is good - 23
33	O clipe de papel está aqui - 26	The clip of paper is here - 25	The paper's clip is here - 24	The paper clip is here - 22
34	O exame de rotina está ok - 25	The exam of routine is ok -	The routine's exam is ok - 24	The routine exam is ok - 22

		25		
35	O nome de família é curto - 25	The name of family is short - 27	The family's name is short - 26	The family name is short - 24
36	O carro de polícia é cinza - 26	The car of police is gray - 25	The police's car is gray - 24	The police car is gray - 22
37	O cartão de crédito é laranja - 29	The card of credit is orange - 28	The credit's card is orange - 27	The credit card is orange - 25
38	O fator de risco é alto - 23	The factor of risk is high - 26	The risk's factor is high - 25	The risk factor is high - 23
39	O teste de história é fácil - 27	The test of History is easy - 27	The History's test is easy - 26	The History test is easy - 24
40	A barra de chocolate é marrom - 29	The bar of chocolate is brown - 29	The chocolate's bar is brown - 28	The chocolate bar is brown - 26
41	A medalha de bronze é minha - 27	The medal of bronze is mine - 27	The bronze's medal is mine - 26	The bronze medal is mine - 24
42	A sopa do bebê está fria - 24	The soup of baby is cold - 24	The baby's soup is cold - 23	The baby soup is cold - 21
43	O filme de comédia é chato - 26	The film of comedy is boring - 28	The comedy's film is boring - 27	The comedy movie is boring - 26
44	O controle da TV é dele - 23	The control of TV is his - 24	The TV's control is his - 23	The TV control is his - 21
45	A estação de trem é longe - 25	The station of train is far - 27	The train's station is far - 26	The train station is far - 24
46	A tecla do piano está muda - 26	The key of piano is mute - 24	The piano's key is mute - 23	The piano key is mute - 21
47	A fatia de tomate é fina - 24	The slice of tomato is thin - 27	The tomato's slice is thin - 26	The tomato slice is thin - 24
48	A comida do hospital é ruim - 27	The food of hospital is bad - 27	The hospital's food is bad - 26	The hospital food is bad - 24
49	O banco do parque é velho - 25	The bench of park is old - 24	The park's bench is old - 23	The park bench is old - 21
50	A casca da manga é doce - 23	The peel of mango is sweet - 26	The mango's juice is sweet - 25	The mango juice is sweet - 23
51	O doce de banana é gostoso - 26	The candy of banana is tasty -	The banana's candy is tasty -	The banana candy is tasty -

		28	27	25
52	O bolo de limão está pronto - 27	The cake of lemon is ready - 26	The lemon's cake is ready - 25	The lemon cake is ready - 23
53	O fio do fone é amarelo - 23	The cord of phone is yellow - 27	The phone's cord is yellow - 26	The phone cord is yellow - 24
54	A mesa de plástico é rosa - 25	The table of plastic is pink - 28	The plastic's table is pink - 27	The plastic table is pink - 25
55	O cabelo do estudante é loiro - 29	The hair of student is blond - 28	The student's hair is blond - 27	The student hair is blond - 25
56	A conta do dentista é barata - 25	The bill of hotel is cheap - 26	The hotel's bill is cheap - 25	The hotel bill is cheap - 23
57	A dor de estômago é forte - 25	The ache of stomach is hard - 27	The stomach's ache is hard - 26	The stomach ache is hard - 24
58	A garrafa de café está cheia - 28	The bottle of coffee is full - 28	The coffee's bottle is full - 27	The coffee bottle is full - 25
59	O professor de inglês é legal - 29	The teacher of English is kind - 30	The English's teacher is kind - 29	The English teacher is kind - 27
60	A porta do hotel está aberta - 28	The door of hotel is open - 25	The hotel's door is open - 24	The hotel door is open - 22

APÊNDICE G - ITENS DISTRATORES DO EXPERIMENTO 3

ITEM	ITENS DISTRATORES EM L1	OPÇÃO SINTÁTICA 1	OPÇÃO SINTÁTICA 2	OPÇÃO SINTÁTICA 3
151	Ela está assistindo TV	She is watching TV	She is watched TV	She is watch TV
152	Ele/a está crescendo	It is growing up	It is grown up	It is grow up
153	Eu estou estudando todo dia	I am studying everyday	I am studied everyday	I am study everyday
154	Está chovendo forte	It is raining hard	It is rained hard	It is rain hard
155	Ele está lendo agora	He is reading now	He is readed now	He is read now
156	Eles estão jogando futebol	They are playing soccer	They are played soccer	They are play soccer
157	Nós estamos jantando	We are having dinner	We are had dinner	We are have dinner
158	Você está escutando música	You are listening music	You are listened music	You are listen music
159	Nós estamos dormindo tarde	We are sleeping late	We are slept late	We are sleep late
160	Ela está escrevendo poemas	She is writing poems	She is wrote poems	She is write poems
161	Ele está falando francês	He is speaking French	He is spoke French	He is speak French
162	Você está andando rápido	You are walking fast	You are walked fast	You are walk fast
163	Está nevando pouco	It is snowing slightly	It is snowed slightly	It is snowing slightly
164	Eles estão lanchando	They are having snacks	They are having snacks	They are have snacks
165	Nós estamos bebendo água	We are drinking water	We are drank water	We are drink water
166	Paul está nadando rápido	Paul is swimming fast	Paul is swam fast	Paul is swim fast
167	Anne está dançando balé	Anne is dancing ballet	Anne is danced ballet	Anne is dance ballet
168	Peter está comendo arroz	Peter is eating rice	Peter is ate rice	Peter is eat rice
169	Kate está falando devagar	Kate is talking quietly	Kate is talked quietly	Kate is talk quietly
170	Charles está vestindo jeans	Charles is wearing jeans	Charles is wore jeans	Charles is wear jeans
171	Susan está escovando dentes	Susan is brushing teeth	Susan is brushed teeth	Susan is brush teeth
172	Jack está usando gravata	Jack is using tie	Jack is used tie	Jack is use tie
173	Pam está ensinando inglês	Pam is teaching English	Pam is taught English	Pam is teach English
174	Eu estou comprando pão	I am buying	I am bought	I am buy bread

		bread	bread	
175	Nós estamos vendendo doces	We are selling candies	We are sold candies	We are sell candies
176	Eles estão chorando alto	They are crying loud	They are cryed loud	They are cry loud
177	Ela está sorrindo feliz	She is smiling happily	She is smiled happily	She is smile happily
178	Ele está pulando corda	He is jumping hope	He is jumped hope	He is jump hope
179	Você está cozinhando feijão	You are cooking beans	You are cooked beans	You are cook beans
180	Vocês estão prestando atenção	You are paying attention	You are payed attention	You are pay attention
181	Estou me sentindo mal	I am feeling bad	I am felt bad	I am feel bad
182	Ele está trabalhando muito	He is working hard	He is worked hard	He is work hard
183	Ela está abrindo portas	She is openning doors	She is opened doors	She is open doors
184	Você está sonhando alto	You are dreaming big	You are dreamed big	You are dream big

OBS: As sentenças da opção sintática 1 são consideradas corretas gramaticalmente.

APÊNDICE H - FREQUÊNCIA DOS SUBSTANTIVOS DO EXPERIMENTO 3

ITEM	SINTAGMAS NOMINAIS EM L1/L2	ESCALA ZIPF SUBST1 EM L1	ESCALA ZIPF SUBST2 EM L1	ESCALA ZIPF SUBST1 EM L2	ESCALA ZIPF SUBST2 EM L2
1	Jogo de futebol/ Soccer match	Jogo 5,66	Futebol 5,42	Soccer 4,09	Match 4,69
2	Vestido de festa/ Party dress	Vestido 3,99	Festa 5,09	Party 5,36	Dress 4,93
3	Borracha do lápis/ Pencil eraser	Borracha 4,02	Lápis 3,76	Pencil 3,99	Eraser 3,00
4	Caminhão de lixo/ Garbage truck	Caminhão 4,56	Lixo 4,62	Garbage 4,41	Truck 4,86
5	Molho de carne/ Meat sauce	Molho 4,07	Carne 4,80	Meat 4,63	Sauce 4,19
6	Tecla do piano/ Piano key	Tecla 3,55	Piano 4,28	Piano 4,39	Key 4,93
7	Fatia de tomate/ Tomato slice	Fatia 4,04	Tomate 4,05	Tomato 3,77	Slice 3,93
8	Comida do hospital/ Hospital food	Comida 4,61	Hospital 4,99	Hospital 5,09	Food 5,18
9	Banco do parque/ Park bench	Banco 5,23	Parque 4,80	Bench 3,98	Park 4,85
10	Casca da manga/ Mango peel	Casca 3,67	Manga 3,49	Mango 3,22	Peel 3,72
11	Salada de fruta/ Fruit salad	Salada 3,84	Fruta 3,88	Fruit 4,33	Salad 4,23
12	Clipe de papel/ Paper clip	Clipe 3,92	Papel 5,27	Paper 5,01	Clip 3,75
13	Exame de rotina/ Routine exam	Exame 4,81	Rotina 4,30	Routine 4,29	Exam 4,12
14	Nome de família/ Family name	Nome 5,55	Família 5,34	Family 5,54	Name 5,80
15	Carro de polícia/ Police car	Carro 5,40	Polícia 5,45	Police 5,37	Car 5,68

16	Tamanho da saia/ Skirt size	Tamanho 4,64	Saia 3,96	Skirt 3,99	Size 4,66
17	Caixa de joias/ Jewel box	Caixa 4,79	Joias 4,07	Jewel 3,86	Box 4,95
18	Vida de solteiro/ Single life	Vida 5,66	Solteiro 3,55	Single 4,85	Life 5,90
19	Pano de limpeza/ Cleaning cloth	Pano 4,04	Limpeza 4,62	Cleaning 4,35	Cloth 3,78
20	Capa do caderno/ Notebook cover	Capa 4,47	Caderno 4,42	Notebook 3,66	Cover 4,97
21	Mesa de cristal/ Crystal table	Mesa 4,79	Cristal 3,88	Crystal 4,20	Table 5,02
22	Conta do restaurante/ Restaurant bill	Conta 5,30	Restaurante 4,77	Restaurant 4,66	Bill 5,07
23	Cabelo da secretária/ Secretary hair	Cabelo 4,41	Secretária 4,57	Secretary 4,52	Hair 5,01
24	Porta do aeroporto/ Airport door	Porta 4,85	Aeroporto 4,71	Airport 4,57	Door 5,46
25	Professor de português/ Portuguese teacher	Professor 5,08	Português 4,66	Portuguese 3,88	Teacher 4,74
26	Museu de dinossauro/ Dinosaur museum	Museum 4,60	Dinossauro 3,34	Dinosaur 3,60	Museum 4,26
27	Cartão de vacina/ Vaccine card	Cartão 4,80	Vacina 4,11	Vaccine 3,28	Card 4,93
28	Teste de geografia/ Geography test	Teste 4,71	Geografia 4,16	Geography 3,36	Test 4,92
29	Barra de Proteína/ Protein bar	Barra 4,28	Proteína 3,92	Protein 3,66	Bar 4,93
30	Filme de suspense/ Suspense film	Filme 5,46	Suspense 3,84	Suspense 3,35	Film 4,81
31	Relógio de parede/ Wall clock	Relógio 4,11	Parede 4,34	Wall 4,84	Clock 4,76
32	Coroa do rei/ King crown	Coroa 4,05	Rei 4,68	King 5,11	Crown 4,13
33	Gaveta do armário/ Closet drawer	Gaveta 3,69	Armário 3,68	Closet 4,43	Drawer 4,11
34	Gema do ovo/ Egg yolk	Gema 3,15	Ovo 3,94	Egg 4,41	Tray 3,90

35	Asa do frango/ Chicken wing	Asa 3,63	Frango 4,21	Chicken 4,79	Wing 4,30
36	Via de acesso/ Access way	Via 4,41	Acesso 5,01	Access 4,50	Way 6,15
37	Passeio de balão/ Balloon ride	Passeio 4,46	Balão 3,71	Balloon 3,93	Ride 5,13
38	Rede de tênis/ Tennis net	Rede 5,22	Tênis 4,43	Tennis 4,13	Net 4,19
39	Placa do computador/ Computer board	Placa 4,54	Computador 4,98	Computer 4,77	Board 4,80
40	Cenário da foto/ Photo setting	Cenário 4,80	Foto 4,95	Photo 4,35	Setting 4,33
41	Artigo de opinião/ Opinion article	Artigo 5,18	Opinião 5,08	Opinion 4,62	Article 4,29
42	Clube de golfe/ Golf club	Clube 5,07	Golfe 3,69	Golf 4,40	Club 4,99
43	Botão do elevador/ Elevator button	Botão 3,87	Elevador 3,80	Elevator 4,38	Button 4,45
44	Grupo de mensagem/ Message group	Grupo 5,66	Mensagem 4,46	Message 4,96	Group 4,86
45	Poluição do ar/ Air pollution	Poluição 4,39	Ar 5,13	Air 5,14	Pollution 3,28

**APÊNDICE I - NÚMEROS DE CARACTERES DOS ITENS EXPERIMENTAIS DO
EXPERIMENTO 3**

ITEM	SENTENÇAS EM L1 - Nº DE CARACTERES	OPÇÃO SINTÁTICA 1 EM L2 - Nº DE CARACTERES	OPÇÃO SINTÁTICA 2 EM L2 - Nº DE CARACTERES	OPÇÃO SINTÁTICA 3 EM L2 - Nº DE CARACTERES
1	O jogo de futebol é hoje - 24	The match of soccer is today - 28	The soccer's match is today - 27	The soccer match is today - 25
2	O vestido de festa é dourado - 28	The dress of party is golden - 28	The party's dress is golden - 27	The party dress is golden - 25
3	O cinto de segurança é preto - 28	The belt of seat is black - 25	The seat's belt is black - 24	The seat belt is black - 22
4	O caminhão de lixo está lá - 26	The truck of garbage is there - 29	The garbage's truck is there - 28	The garbage truck is there - 26
5	O molho de carne está gostoso - 29	The sauce of meat is tasty - 26	The meat's sauce is tasty - 26	The meat sauce is tasty - 23
6	A tecla do piano é muda - 23	The key of piano is mute - 24	The piano's key is mute - 23	The piano key is mute - 21
7	A fatia de tomate é fina - 24	The slice of tomato is thin - 27	The tomato's slice is thin - 26	The tomato slice is thin - 24
8	A comida do hospital é ruim - 27	The food of hospital is bad - 27	The hospital's food is bad - 26	The hospital food is bad - 24
9	O banco do parque é velho - 25	The bench of park is old - 24	The parks's bench is old - 24	The park bench is old - 21
10	A casca da manga é grossa - 25	The peel of mango is thick - 26	The mango's peel is thick - 25	The mango peel is thick - 23
11	A salada de fruta é boa - 23	The salad of fruit is good - 26	The fruit's salad is good - 25	The fruit salad is good - 23
12	O clipe de papel está aqui - 26	The clip of paper is here - 25	The paper's clip is here - 24	The paper clip is here - 22
13	O exame de rotina está ok - 25	The exam of routine is ok - 25	The routine's exam is ok - 24	The routine exam is ok - 22
14	O nome de família é curto -	The name of family is short -	The family's name is short -	The family name is short -

	25	27	26	24
15	O carro de polícia é cinza - 26	The car of police is gray - 25	The police's car is gray - 24	The police car is gray - 22
16	O tamanho da saia é médio - 25	The size of skirt is medium - 27	The skirt's size is medium - 26	The skirt size is medium - 24
17	A caixa de joias é redonda - 26	The box of jewel is round - 25	The jewel's box is round - 23	The jewel box is round - 22
18	A vida de solteiro é ocupada - 28	The life of single is busy - 26	The single's life is busy - 25	The single life is busy - 23
19	O pano de limpeza está sujo - 27	The cloth of cleaning is dirty - 30	The cleaning's cloth is dirty - 29	The cleaning cloth is dirty - 27
20	A capa do caderno é branca - 26	The cover of notebook is white - 30	The notebook's cover is white - 29	The notebook cover is white - 27
21	A mesa de cristal é frágil - 26	The table of crystal is fragile - 31	The crystal's table is fragile - 30	The crystal table is fragile - 28
22	A conta do restaurante é cara - 29	The bill of restaurant is expensive - 35	The restaurant's bill is expensive - 34	The restaurant bill is expensive - 32
23	O cabelo da secretária é ruivo - 30	The hair of secretary is red - 28	The secretary's hair is red - 27	The secretary hair is red - 25
24	A porta do aeroporto está fechada - 33	The door of airport is closed - 29	The airport's door is closed - 28	The airport door is closed - 26
25	O professor de português é novato - 33	The teacher of Portuguese is newbie - 35	The Portuguese's teacher is newbie - 34	The Portuguese teacher is newbie - 32
26	O museu de dinossauro é maravilhoso - 35	The museum of dinosaur is wonderful - 35	The dinosaur's museum is wonderful - 34	The dinosaur museum is wonderful - 32
27	O cartão de vacina é preto - 26	The card of vaccine is black - 28	The vaccine's card is black - 27	The vaccine card is black - 25
28	O teste de geografia é difícil - 30	The test of geography is easy - 29	The geography's test is easy - 28	The geography test is easy - 26
29	A barra de cereal é saudável - 28	The bar of cereal is healthy - 28	The cereal's bar is healthy - 27	The cereal bar is healthy - 25
30	O filme de suspense é	The film of suspense is	The suspense's film is	The suspense film is

	surpreendente - 35	surprising - 34	surprising - 33	surprising - 31
31	O relógio de parede está atrasado - 33	The clock of wall is late - 25	The wall's clock is late - 24	The wall clock is late - 22
32	A coroa do rei é valiosa - 24	The crown of king is valuable - 29	The king's crown is valuable - 28	The king crown is valuable - 26
33	A gaveta do armário está aberta - 31	The drawer of closet is open - 28	The closet's drawer is open - 27	The closet drawer is open - 25
34	A bandeja do ovo é amarela - 26	The tray of egg is yellow - 25	The egg's tray is yellow - 24	The egg tray is yellow - 22
35	A asa do frango está cozida - 27	The wing of chicken is cooked - 29	The chicken's wing is cooked - 28	The chicken wing is cooked - 26
36	A via de acesso está bloqueada - 30	The way of access is blocked - 28	The access' way is blocked - 26	The access way is blocked - 25
37	O passeio de balão é emocionante - 32	The ride of balloon is exciting - 31	The balloon's ride is exciting - 30	The balloon ride is exciting - 28
38	A rede de tênis é branca - 24	The net of tennis is white - 26	The tennis' net is white - 24	The tennis net is white - 23
39	A placa do computador está nova - 28	The board of computer is new - 28	The computer's board is new - 27	The computer board is ok - 24
40	O jogo de batalha é legal - 25	The game of battle is cool - 26	The battle's game is cool - 25	The battle game is cool - 23
41	O artigo de opinião está perfeito - 33	The article of opinion is perfect - 33	The opinion's article is perfect - 32	The opinion article is perfect - 30
42	O clube de golfe é amplo - 24	The club of golf is wide - 24	The golf's club is wide - 23	The golf club is wide - 21
43	O botão do elevador está desligado - 34	The button of elevator is off - 29	The elevator's button is off - 28	The elevator button is off - 26
44	O grupo de mensagem está calmo - 30	The group of message is calm - 28	The message's group is calm - 27	The message group is calm - 25
45	A poluição do ar é preocupante - 30	The pollution of air is worrying - 32	The air's pollution is worrying - 31	The air pollution is worrying - 29