

UNIVERSIDADE ESTADUAL DE CAMPINAS

**Conhecimento sintático-semântico e processamento
de sentenças em rede neural recorrente simples**

RICARDO BASSO GARCIA

Dissertação de mestrado apresentada
ao Programa de Pós-Graduação do
Departamento de Linguística do Ins-
tituto de Estudos da Linguagem,
como requisito parcial para a
obtenção do título de mestre em
Linguística.

INSTITUTO DE ESTUDOS DA LINGUAGEM

Campinas

Julho de 2006

UNIVERSIDADE ESTADUAL DE CAMPINAS

INSTITUTO DE ESTUDOS DA LINGUAGEM

PROGRAMA DE PÓS-GRADUAÇÃO EM LINGÜÍSTICA

(NÍVEL DE MESTRADO)

**Conhecimento sintático-semântico e processamento
de sentenças em rede neural recorrente simples**

RICARDO BASSO GARCIA

Banca Examinadora:

Prof. Dr. Edson Françoze (IEL-UNICAMP) – orientador

Prof. Dr. Fernando José Von Zuben (FEEC-UNICAMP)

Prof. Dr. Plínio Almeida Barbosa (IEL-UNICAMP)

Ficha catalográfica elaborada pela Biblioteca do IEL - Unicamp

G165c Garcia, Ricardo Basso.
Conhecimento sintático-semântico e processamento de sentenças em rede neural recorrente simples / Ricardo Basso Garcia. -- Campinas, SP : [s.n.], 2006.

Orientador : Edson Françaço.
Dissertação (mestrado) - Universidade Estadual de Campinas, Instituto de Estudos da Linguagem.

1. Psicolinguística. 2. Conexionismo. 3. Cognição. I. Françaço, Edson. II. Universidade Estadual de Campinas. Instituto de Estudos da Linguagem. III. Título.

Título em inglês: Syntactic-semantic knowledge and sentence processing in a simple recurrent neural network.

Palavras-chaves em inglês (Keywords): Psycholinguistics; Connectionism; Cognition.

Área de concentração: Lingüística.

Titulação: Mestre em Lingüística.

Banca examinadora: Prof. Dr. Edson Françaço (orientador), Prof. Dr. Fernando José Von Zuben, Prof. Dr. Plínio Almeida Barbosa.

Data da defesa: 11/07/2006.

Programa de Pós-Graduação: Lingüística.

Ao meu pai (*in memoriam*)
À minha mãe
Aos meus irmãos

AGRADECIMENTOS

Gostaria de agradecer, primeiramente, ao professor Edson Françoze, pela orientação sempre atenciosa, inteligente, minuciosa.

Durante os anos de mestrado, contei com o suporte e auxílio dos funcionários do IEL. Agradeço ao Wilson Kawai, do setor de informática. Sou especialmente grato aos funcionários da secretaria de pós-graduação, Rosemeire Aparecida de Almeida Marcelino e Cláudio Pereira Platero. Agradeço aos colegas e amigos do IEL, companheiros de trabalho, conversas, congressos e viagens: Ângela Kajita, Antonio Barros, Denílson Amade Sousa, Denise Seabra, Juliano Bellinazzi Nequirito, Karen Alves, Laudino Roces, Leandro Silveira, Luciana Lucente, Maria Luiza Cunha Lima, Pablo Arantes, Renato Basso e Valderes Rinaldi.

Sou grato à Gabriela Helena Bauab, não só pela amizade, como também por me incentivar a acompanhar o curso sobre redes neurais, ministrado pelo professor Fernando Von Zuben na FEEC/Unicamp.

Agradeço aos professores Fernando e Plínio pela leitura cuidadosa do texto de qualificação, pelos comentários e sugestões que contribuíram para o enriquecimento deste trabalho.

Agradeço o apoio e o incentivo da minha família – minha mãe e meus irmãos – e dos amigos Adélia Rocha e Fernando Stefani. Sou grato a Luiza Carnicero de Castro, não somente pela amizade e companheirismo, como também pelo incentivo, apoio e paciência em momentos decisivos.

Agradeço a CNPq pela bolsa de mestrado que possibilitou minha dedicação à pesquisa.

Seria injusto de minha parte negligenciar a importância das pessoas que me inspiraram e incentivaram quando eu era apenas um aluno do bacharelado em lingüística no IEL. Por esse motivo, na etapa encerrada por esta dissertação, reconheço e destaco a importância que algumas pessoas tiveram na minha formação. Durante os anos de graduação, contei com a inspiração, incentivo e apoio dos professores Ana Luiza Navas,

Edson Françaço, Edwiges Morato, Eleonora Albano e Plínio Almeida Barbosa, bem como dos colegas Antonio Barros, Laudino Rocas, Pablo Arantes e Renato Basso. Agradeço, especialmente, a Martina Marana, pela amizade, paciência e companheirismo nos períodos de calma e de tormenta. Sou grato, também, aos professores Angel Humberto Corbera Mori, Jairo Moraes Nunes, Kanavillil Rajagopalan, Rodolfo Ilari, Sírio Possenti e Vandersí Sant'Ana Castro.

RESUMO

Esta dissertação começa com uma reflexão sobre o caráter interdisciplinar da psicolinguística, que é tomada como o campo de pesquisa que investiga, entre outros, os processos cognitivos subjacentes ao comportamento linguístico. Essa investigação emprega diferentes métodos e técnicas, dada a complexidade do estudo de processos que são internos ao organismo e, portanto, não diretamente observáveis. Nesse quadro, introduzimos a metodologia adotada nesta dissertação – a *modelagem*. O uso de modelos é tratado em termos gerais, para então ser apresentado no contexto do estudo dos processos cognitivos.

Neste ponto, é crucial notar que modelar esses processos implica em fazer suposições sobre *como* eles são. A hipótese geral que fundamenta os modelos de nosso interesse é a de que *cognição é computação*, ou seja, processos cognitivos são computacionais. Essa hipótese deu origem a dois paradigmas de modelagem distintos – um baseado nos fundamentos da *computação digital* e outro nos princípios da *neurocomputação*. Este segundo é apresentado em detalhes, depois de uma exposição das principais diferenças entre esses paradigmas.

Estabelecidas as bases teórico-metodológicas dessa abordagem, replicamos o experimento de Elman (1990) em que uma rede neural recorrente simples é treinada para processar palavras dispostas em seqüências que compõem sentenças simples. É interessante notar que essa disposição reflete informações gramaticais, que podem ser descritas em termos sintáticos (como sujeito, verbo e objeto) e semânticos (como agente, paciente e tema). Nosso intuito foi investigar *se* e *como* informações desse tipo são aprendidas pelo sistema. A análise dessa rede mostrou que seu processamento opera principalmente por distinções lexicais, isto é, não há o uso efetivo de informações linguísticas do nível de sentenças.

O passo seguinte foi modificar esse experimento – treinamos a rede com o mesmo conjunto de sentenças, desta vez marcando a fronteira entre elas. A análise dessa rede mostrou que, nesse novo experimento, o processamento opera principalmente por

distinções sintáticas, havendo também a presença de distinções semânticas de papel temático.

A partir desses resultados, elaboramos um novo experimento – treinamos uma rede com um conjunto de sentenças nas quais a disposição das palavras reflete principalmente informações semânticas de papel temático, havendo também a presença de informações sintáticas e da fronteira entre as sentenças. O objetivo foi controlar a informação linguística presente nas sentenças, de modo a permitir maior confiabilidade nos resultados. A análise dessa rede mostrou-se coerente com a do experimento anterior, ou seja, o processamento está, de fato, operando segundo conhecimentos sintático-semânticos do nível de sentenças.

A consistência dos resultados obtidos permite afirmar, com um grau de confiabilidade satisfatório, que há a presença de conhecimento sintático-semântico, durante o processamento de sentenças, ao mesmo tempo em que mostra *como* essas informações estão presentes. Esses resultados são importantes porque estendem a já conhecida capacidade da rede recorrente simples em codificar informação sintática, mostrando também sua capacidade em captar informação semântica.

ABSTRACT

This dissertation begins with a reflection about the interdisciplinary roots of psycholinguistics, which is taken as the research field concerned with the investigation of the cognitive processes subjacent to the verbal behavior. This investigation employs different methods and techniques, given the complexity inherent to the study of processes that take place inside the organism, and therefore are not directly observable. In view of this, we introduce the methodology adopted in this dissertation – *modeling*. The usage of models is initially dealt with in general terms, and then is analysed in the context of the study of cognitive processes.

At this point, it is crucial to notice that modeling such processes implies making assumptions about *how* they are. The general hypothesis underlying the models we are interested in is that *cognition is computation*, that is, cognitive processes are computational. This hypothesis has given birth to two different paradigms of modeling – one based on the fundamentals of *digital computing*, and other based on the principles of *neurocomputing*. After a brief discussion of their main differences, the later is presented in details.

Once the theoretical-methodological basis of the neurocomputing framework is established, we replicate Elman's (1990) experiment in which a simple recurrent neural network is trained to process words in sequences, forming simple sentences. It should be noted that the sequences reflect grammatical information that can be described in syntactic (subjects, verbs and objects) and semantic (agents, patients and themes) terms. Our intent was to investigate *whether* and *how* such information is learned by the system. Analyses of this experiment have shown that the processing is guided by lexical information, that is, linguistic information at the sentence level is not used.

In the next step we have modified this experiment by training the network with the same set of sentences, just introducing an end-of-period mark between them. The analyses have shown that, in this new experiment, the processing is guided by both syntactic and semantic information, such as thematic roles.

Taking these results as starting point, we devised a new experiment – a similar network was trained with a set of sentences in which word sequences were mainly led by semantic information like thematic role; the end-of-period mark between sentences was used. The intent was to control linguistic information so as to increase our confidence in the results. Results consistent with the former experiments were obtained, that is, the processing is indeed guided by both syntactic and semantic information at the sentence level.

The consistency of the results achieved allows us to state that there are both syntactic and semantic knowledge operating during sentence processing. This result is important because it extends the already known simple recurrent network ability in coding syntactic information, showing that SRNs also are capable of handling semantic information like thematic roles.

SUMÁRIO

Capítulo I – Psicolingüística e Conexionismo

1. A psicolingüística.....	1
2. Aspectos metodológicos.....	7
3. Modelos nos estudos da cognição.....	13
3.1 A computação digital.....	14
3.2 A neurocomputação.....	18
4. A abordagem conexionista.....	22
4.1 Modelos de neurônios artificiais.....	23
4.2 Redes de neurônios – arquiteturas.....	25
4.3 Aprendizagem.....	28
4.4 Um exemplo prático.....	30

Capítulo II – Conhecimento sintático-semântico e processamento de sentenças em redes neurais

1. Treinando redes neurais com sentenças.....	33
2. Questões de interesse.....	35
3. Replicando Elman (1990)	39
3.1 Representação lexical e análise de agrupamentos.....	42
3.2 Representação gramatical e análise de componentes principais.....	45
4. Em busca da estrutura sintática.....	52
4.1 Análise de agrupamentos.....	52
4.2 Análise de componentes principais.....	53
5. Um novo experimento.....	58
5.1 A simulação.....	60
5.2 Análise de agrupamentos.....	61
5.3 Análise de componentes principais.....	62

Capítulo III – Considerações Finais.....

Referências Bibliográficas.....

Anexo A.....

Anexo B.....

CAPÍTULO I

Psicolingüística e Conexionismo

1. A psicolingüística

O que é psicolingüística? O que ela estuda? Ela é uma ciência? Estas e outras questões estarão direta ou indiretamente relacionadas com o exercício que faremos ao caracterizar a psicolingüística como um campo de estudo científico. O objetivo desta primeira seção, portanto, consiste em refletir sobre tópicos como o lugar da psicolingüística na ciência, quais são as questões de seu domínio de investigação, suas práticas metodológicas e a sua relevância teórica. Acreditamos que esta reflexão, somada à contextualização teórico-metodológica que complementa o capítulo, é de fundamental importância para uma adequada compreensão dos estudos que realizamos e que serão discutidos nos capítulos 2 e 3.

Começemos nossa reflexão pelo próprio termo *psicolingüística*. Pela sua composição, pode-se depreender que o tipo de estudo que recebe esse nome diz respeito tanto à psicologia quanto à lingüística. Estamos, portanto, diante de um domínio interdisciplinar que partilha, com duas disciplinas reconhecidas, questões e conhecimentos sobre um mesmo assunto. Feita essa consideração, uma questão pertinente seria averiguar como cada disciplina reconhece esse domínio em comum. Para tanto, consultamos um dicionário técnico de lingüística e outro de psicologia. De acordo com o primeiro, a psicolingüística é o “ramo da lingüística que estuda a correlação entre o comportamento lingüístico e os processos psicológicos que se encontram, supostamente, por trás deste comportamento” Crystal (1988:214). O segundo a caracteriza como a subdisciplina da psicologia cujo objeto de estudo é o desenvolvimento e o funcionamento dos comportamentos verbais, sendo que esse estudo emprega modelos e técnicas da lingüística na formulação de hipóteses e na interpretação dos dados (Doron & Parot, 1998:627).

Conforme seria o esperado, tanto a psicologia quanto a lingüística reconhecem a psicolingüística como uma subdisciplina ou um ramo de seus estudos. Não achamos, contudo, que essas definições caracterizam adequadamente esse domínio como

interdisciplinar, principalmente porque as relações ‘ser subdisciplina de’ ou ‘ser um ramo de’ impõem uma classificação hierárquica que subestima o papel da outra disciplina. Em outras palavras, não existe um ramo de dois galhos distintos.

Não pretendemos, entretanto, travar aqui uma luta por demarcação de territórios, mesmo porque essa luta não existe e o domínio da psicolinguística claramente reside na região fronteira entre a psicologia e a linguística. Não se trata, como podemos ilustrar, de uma cerca que delimita precisamente duas regiões, mas de um rio cuja água abastece ambos os territórios sem que nenhum reivindique para si sua propriedade.

Continuando a busca por uma definição mais adequada, consultamos um dicionário de filosofia, o que se justifica tanto por este ser o campo do qual a psicologia e a linguística descendem quanto por ser o campo de reflexão sobre as ciências em geral. Nessa referência, a psicolinguística está definida como uma pesquisa interdisciplinar que usa descrições teóricas da linguagem provenientes da linguística na investigação de processos psicológicos que estão por trás da produção, percepção e aprendizagem da linguagem. Menciona-se, também, que não há um consenso sobre uma caracterização apropriada do campo e suas principais questões (Audi, 1999:757).

Conforme podemos notar, comparando essas posturas, tanto Crystal quanto Doron & Parot situam a psicolinguística parcialmente, uma vez que não definem adequadamente a identidade interdisciplinar desse campo. Enquanto o primeiro não faz menção alguma à contribuição da psicologia, as segundas atribuem à linguística o papel da formulação de hipóteses e a interpretação dos dados. Trata-se de visões limitadas, como veremos no decorrer desta seção. A mesma parcialidade não é observada em Audi, onde a definição reconhece a interdisciplinaridade ao mesmo tempo em que menciona uma falta de consenso, que está refletida nas questões de interesse: onde Crystal diz *estudo de correlações*, Doron & Parot falam em *objeto de estudo*. Essa discrepância revela, conforme mostramos abaixo, a coerência dos autores em relação aos preceitos fundamentais das disciplinas. Neste ponto, pelas considerações feitas, temos que as definições são divergentes em relação ao reconhecimento desse domínio interdisciplinar e à especificação de um objeto de estudo. Para abordar esse assunto, apresentamos brevemente as questões

que são do domínio da psicologia e da lingüística, para então buscarmos depreender aquelas que são de comum interesse.

Quando tomadas separadamente, cada qual circunscreve seu próprio domínio de investigação, possui seus métodos e teorias. Enquanto a psicologia (ao menos a psicologia cognitiva) caracteriza-se pelo estudo dos diversos aspectos relacionados direta ou indiretamente ao *comportamento* dos organismos, a lingüística caracteriza-se pelo estudo de aspectos essencialmente estruturais da linguagem, como a fonologia, a morfologia, o léxico, a sintaxe e a semântica, por exemplo. Apresentado dessa maneira, não é direta a identificação de uma intersecção desses estudos. Conforme veremos, o diálogo entre essas duas disciplinas é possível quando há o interesse na linguagem como uma forma de comportamento, o que tradicionalmente é uma questão secundária nos estudos da linguagem. De fato, a lingüística moderna pouco se interessou por esse aspecto, assentando suas investigações no pressuposto de que a linguagem é, ou pelo menos assim pode ser submetida à investigação científica, uma estrutura encerrada em si mesma e como tal constitui um objeto de estudo. Mesmo quando se trata de uma teoria de apelo biológico e psicológico, como é o caso da gramática gerativa (por exemplo, veja o apelo ao inatismo e ao mentalismo de Chomsky), o fato de a linguagem ser um comportamento é de menor importância, cabendo à lingüística o estudo da competência que um falante ideal possui (um conhecimento abstrato e interno, portanto).

Em suma, embora a linguagem possa ser vista como um comportamento, o estudo lingüístico da mesma privilegiou questões tipicamente estruturais e abstratas em detrimento de aspectos psicológicos e sociais, inerentes ao comportamento humano. Essa postura está refletida na visão de psicolingüística presente em Crystal. Por outro lado, um recorte semelhante não pode ser feito pela psicologia – a explicação do comportamento inteligente não pode excluir a linguagem. É por este motivo que Doron & Parot definem um objeto de estudo para a psicolingüística, alçando-a ao nível de uma disciplina autônoma.

Cumpramos alertar para o corte epistemológico, na lingüística: não há a negação dos aspectos sociais e psicológicos da linguagem; estes são, outrossim, reconhecidos como sendo externos ao objeto propriamente dito. Isso tanto é verdade que existem estudos da linguagem situados em regiões fronteiriças com os estudos do social e do psicológico,

embora cada qual se desenvolva de forma independente. Podemos visualizar esse argumento através de um diagrama (figura 1.1), no qual representamos as intersecções desses três tipos de estudo. Podemos situar temas e pesquisas em cada uma das regiões comuns. Por exemplo: na região ‘a’ teríamos a psicologia social; na região ‘b’ a sociolingüística, a pragmática e a análise do discurso; e na região ‘c’ encontramos as questões tipicamente de interesse da psicolingüística. O diagrama não esgota as possibilidades: também poderíamos incluir aí uma intersecção com os estudos biológicos. É interessante mencionar que é possível elaborar teorias na intersecção desses quatro campos; tal é o caso, por exemplo, de estudos que vêm a linguagem e a cognição como um fenômeno situado e distribuído, reconhecendo a importância do organismo, seu corpo, seu ambiente e suas diversas interações.

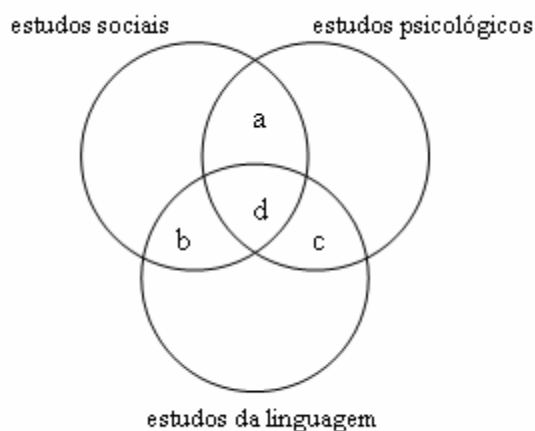


Figura 1.1 A psicolingüística vista como uma intersecção de campos de estudo, na região ‘c’.

Voltando à questão de nosso interesse, resta explicitar como se faz uma ponte entre a psicologia e a lingüística. Isso se dá pela possibilidade de se questionar, a partir da lingüística, quais são os processos psicológicos ou cognitivos que estão na base do comportamento lingüístico, assim como é possível indagar, a partir da psicologia, sobre as relações entre o comportamento verbal e os processos psicológicos. Essas questões remetem, em essência, à mesma temática sobre como o conhecimento lingüístico está incorporado no organismo, como ele é aprendido e quais são e como funcionam os processos que o tornam possível. É notório que não se pode responder a esse tipo de indagação a partir de uma abordagem unilateral – tanto a psicologia quanto a lingüística necessitam unir conhecimentos. Por um lado, a psicologia (em especial a psicologia

cognitiva) é uma forma de estudar aspectos internos ao organismo; por outro lado, a lingüística fornece conhecimentos sobre a linguagem. É neste ponto que conseguimos captar a natureza da psicolingüística: as relações entre linguagem e processos cognitivos não podem prescindir de uma abordagem conjunta. Não se trata, portanto, de trocar ou importar métodos e teorias conforme a visão de Doron & Parot, mas sim de abordar questões de natureza empírica de modo a aliar conhecimentos e técnicas que cada disciplina em si não possui.

A psicolingüística, portanto, pode ser caracterizada como a confluência de duas ciências no estudo dos processos cognitivos subjacentes à aquisição, compreensão e produção da linguagem. Se, por um lado, usar uma língua requer habilidades lingüísticas como codificar e decodificar sons, gestos e letras, reconhecer e acessar o significado das palavras, analisar estruturas sintáticas e construir unidades de sentido mais amplas, por outro, envolve processos psicológicos gerais como ação, percepção, atenção, antecipação e memória (acesso e organização). A psicolingüística, então, procura entender e explicar como o uso da linguagem envolve processos psicológicos gerais e específicos. Não há, portanto, um objeto de estudo definido, mas um conjunto de fenômenos que são abordados de modo interdisciplinar. Dentre os temas de investigação, podemos citar a aquisição da linguagem, a percepção e a produção da fala, a organização e o acesso do conhecimento lexical, o processamento de sentenças e a leitura¹.

A psicolingüística, dada sua natureza híbrida, não se constitui como uma ciência propriamente dita, ou seja, não produz teorias nem possui métodos próprios. Mas os resultados de suas pesquisas fornecem evidências sobre o processamento lingüístico, possuindo relevância teórica na psicologia e na lingüística. Por exemplo, fenômenos do comportamento lingüístico (como o conjunto de fatores que operam no processo de reconhecimento lexical) devem ser explicados por uma teoria sobre a cognição. Da mesma maneira, evidências sobre o processamento da linguagem podem contribuir, embora em menor grau, para teorias cognitivistas da linguagem, como é o caso da gramática gerativa que, embora procure explicar o conhecimento lingüístico abstrato, espera que este seja, em última instância, adequadamente integrado por sistemas de produção (Corrêa, 2002).

¹ Ver Gernsbacher (1994) e Traxler & Gernsbacher (2006) para um panorama da área.

A contribuição da psicolingüística para suas “ciências-mãe”, no entanto, é assimétrica, i.e., não é o caso que toda contribuição tenha o mesmo valor para ambas as ciências. A psicolingüística tem se mostrado mais frutífera no campo dos estudos da cognição do que no dos estudos da linguagem. Como foi anteriormente discutido, isso advém do corte epistemológico feito pelas teorias da lingüística, que se desenvolvem sem a necessidade de explicar fenômenos do processamento lingüístico.

Retomando os objetivos desta seção, acreditamos que as reflexões aqui expostas levam ao reconhecimento de que a psicolingüística, embora não seja uma ciência no sentido estrito do termo, constitui-se como um campo interdisciplinar de estudo científico, no qual lingüistas e psicólogos procuram entender e explicar um comportamento que, apesar de ser publicamente observável, advém de processos internos ao organismo.

Discutida a relevância teórica da psicolingüística, ainda resta mencionar que caracterizá-la como estudo científico também se deve à legitimidade de suas práticas metodológicas. Durante o predomínio do behaviorismo, principalmente na primeira metade do século XX, houve uma baixa produtividade desse campo de pesquisa. A postura científica adotada nessa época determinava que todo comportamento do organismo deveria ser explicado através da pura observação, não sendo possível fazer ciência de fenômenos introspectivos como consciência, processos mentais etc. Por esse motivo, os behavioristas limitavam-se a observar e catalogar o comportamento dos organismos em termos dos estímulos que esses recebiam do ambiente e das respostas produzidas na presença desses estímulos. Dessa forma, procurava-se descobrir leis que pudessem descrever, prever e explicar o comportamento (Stillings et al., 1987:310). As dificuldades do behaviorismo em explicar o comportamento inteligente em termos de estímulos e respostas proporcionaram um ambiente favorável a teorias mentalistas. Uma compreensão adequada da cognição passaria então, necessariamente, pelo estudo dos processos internos que fazem a mediação entre percepção e ação, precisando, portanto, de um olhar voltado para dentro do organismo (Stillings et al., 1987:311).

A consolidação da psicolingüística como uma disciplina científica remonta aos anos de 1950, época marcada pelo reconhecimento das limitações do método observacional e pela aceitação da metodologia experimental como um método legítimo no estudo dos

processos intra-organismos. Nas décadas posteriores, avanços tecnológicos proporcionaram melhores técnicas de medida, a coleta de dados mais informativos e um melhor entendimento do processamento da linguagem (Dijkstra & Smedt, 1996b:3). Atualmente, a metodologia experimental tornou-se mais produtiva com a aplicação de técnicas provenientes de outras áreas, como o uso da modelagem e simulação em computadores, bem como das técnicas de neuroimageamento como a tomografia computadorizada (*PET, Positron Emission Tomography*) e a ressonância magnética funcional (*fMRI, functional Magnetic Ressonance Image*).

2. Aspectos metodológicos

Esta dissertação enquadra-se na psicolingüística que estuda o processamento lingüístico através de recursos como modelos e simulações, para a qual se cunhou o termo *psicolingüística computacional* (Crocker, 1996; Dijkstra & Smedt, 1996a). Nesta seção, justificamos as motivações para o uso dessa metodologia e apresentamos noções fundamentais como *sistema, modelo, modelagem e simulação*.

Os termos *sistema* e *modelo* são amplamente usados nas ciências em geral, assim como na psicologia e na lingüística. Cumpre alertar, porém, que esses termos são polissêmicos e não possuem um uso único, mesmo na esfera científica, na qual se pretende que haja rigor terminológico. Por este motivo, delimitaremos o uso que faremos desses termos, atrelando-os ao contexto de uma prática metodológica comum a diversas ciências. Essa prática, que não é específica nem tampouco restrita ao domínio científico, denomina-se *modelagem*, e é de tal importância e vasta aplicação que há uma extensa bibliografia dedicada exclusivamente à teorização e à sistematização das diversas técnicas disponíveis (ver, e.g., Kheir, 1988; Shannon, 1975 e Zeigler, Praehofer & Kim, 2000).

Conforme a abordagem delineada em Zeigler, Praehofer & Kim (2000), a base formal e conceitual dessa metodologia é a da teoria dos sistemas matemáticos, que fornecem um formalismo rigoroso e fundamental para a representação de sistemas dinâmicos. Essa teoria subdivide-se em dois ramos principais – por um lado, explora *os níveis da especificação de sistemas*, i.e., níveis descritivos do comportamento de sistemas e dos elementos que os compõem; por outro, explora *os formalismos da especificação de sistemas*, i.e., os estilos de modelagem, como discreta e contínua, que podem ser usados na construção de modelos

de sistemas. Para os propósitos desta dissertação, é necessário tratar de aspectos relacionados aos níveis de especificação de sistemas. Apresentamos primeiro o aparato conceitual básico, para posteriormente, na seção 4, introduzirmos o formalismo com o qual trabalhamos.

Em primeiro lugar, precisamos especificar o que se entende por sistema no âmbito dessa metodologia. Em outras palavras, é necessário especificar a que tipo de ‘coisa’ atribuímos o nome de *sistema*. As respostas que procuramos podem ser encontradas na *sistêmica* ou *teoria dos sistemas gerais*, que procura organizar, a partir dos mesmos princípios e conceitos, o estudo de sistemas complexos que são do interesse de diversas disciplinas científicas.

A *teoria dos sistemas gerais* é o estudo da organização abstrata dos fenômenos, independente de sua substância, tipo, ou escala espacial ou temporal de existência. Ela investiga tanto os princípios comuns a todas as entidades complexas quanto os modelos (geralmente matemáticos) que podem ser usados para descrevê-las (Audi, 1999:898). A consequência direta desse compromisso com a generalidade é o alto nível de abstração conceitual que a *sistêmica* fornece. Uma apresentação sintetizada pode ser encontrada em Bresciani & D’Ottaviano (2000). Um tratamento aprofundado pode ser encontrado em Bertalanffy (1968), Klir (1969, 1991) e Weinberg (1975).

Em um sentido amplo, um sistema é um conjunto de aspectos do mundo que, de alguma maneira, estão relacionados e interagem entre si (van Gelder & Port, 1995:5). Cabe ao sujeito delimitar, dentre esses aspectos, o que observar, estudar, abstrair, conceituar, analisar, simular, modelar ou representar, na busca pela compreensão e explicação de um determinado objeto de estudo (Bresciani & D’Ottaviano, 2000:287). Sistemas, então, podem ser estudados pela observação, pela experimentação, pela construção de protótipos ou pela elaboração de modelos lógico-matemáticos (Neelankavil, 1987:1), com o objetivo de investigar elementos de sua organização que caracterizam sua existência, estrutura e funcionalidade (Bresciani & D’Ottaviano, 2000:287).

Em um sentido mais específico, seguindo a síntese de Bresciani & D’Ottaviano (2000), um sistema é um conjunto de elementos interdependentes e inter-relacionados que formam uma estrutura, a qual possui uma funcionalidade. Esses elementos conduzem

processos e operações, produzem fenômenos e são responsáveis por transformações, conversões e eventos que caracterizam os seus comportamentos. As características e as propriedades dos elementos podem ser expressas através de parâmetros variáveis ou constantes, que podem assumir valores que descrevem o estado de cada elemento. Tais valores são estabelecidos pelas características do elemento, pelas relações deste com outros e por fatores externos a ele. Os elementos, portanto, possuem características individuais e relacionais, e podem ser classificados em três grupos: elementos de entrada (ou importação), elementos dos processos de transformação interna e elementos de saída (ou de exportação) do sistema. Convém mencionar que um sistema pode ser composto por subsistemas, para os quais valem as mesmas características e propriedades supracitadas (figura 1.2).

Logo, um sistema pode ser descrito em termos de:

- seu *comportamento externo*: a relação existente entre o histórico temporal de suas entradas e de suas saídas (Zeigler, Praehofer & Kim, 2000:3-4). Nesse caso, o sistema é descrito como um modelo de entrada-saída, ou seja, é visto externamente como uma “caixa-preta”, à qual não está associado um conceito bem definido de estado ou situação interna (Delchamps, 1988:4). O enfoque é mais o comportamento do sistema e menos seus aspectos internos.

- sua *constituição interna*: seu estado interno e seu mecanismo de transição de estados, que especifica como as entradas transformam estados atuais em sucessores, bem como é responsável pelo mapeamento estado-saída (Zeigler, Praehofer & Kim, 2000:4). Nesse caso, o sistema é descrito como um modelo de espaço de estados, ou seja, é visto internamente a partir de seus elementos ou componentes constituintes. O enfoque é tanto o comportamento do sistema quanto seus componentes internos.

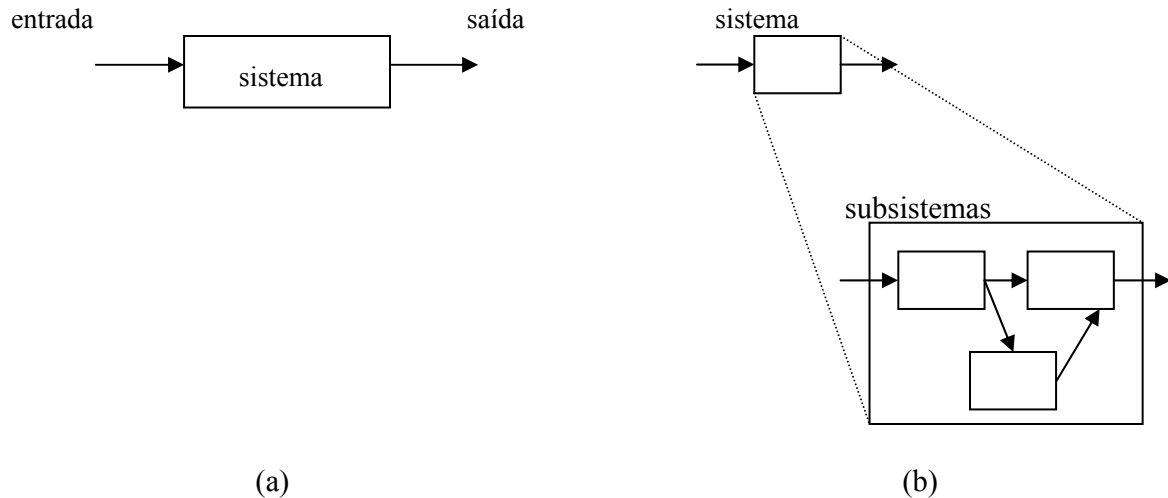


Figura 1.2 Em (a) temos um esquema geral de um sistema. Em (b) vemos que um sistema pode ser decomposto em subsistemas. Figura adaptada de Zeigler, Praehofer & Kim (2000:5).

Temos, agora, uma idéia mais precisa do que é um sistema. Devemos, pois, fazer uma escolha entre a descrição em termos de entrada-saída ou de espaço de estados, o que está intimamente relacionado com a complexidade do objeto de estudo e do enfoque que queremos dar. Para a psicologia que embasa nosso trabalho, não basta o mero estudo do comportamento – é preciso compreender os processos que o subjazem. Faz-se necessário, então, usar um modelo de espaço de estados. Veremos, na seção 4, como é esse modelo.

Antes de prosseguirmos, convém mencionar que consideramos um sistema cognitivo qualquer sistema que exiba comportamento cognitivo ou tenha propriedades cognitivas. Por exemplo, sistemas que possuem uma ou mais destas características: percepção, ação, decisão, memória, associação, linguagem etc. O sistema cognitivo global pode ser estudado em termos de seus subsistemas. Dada essa possibilidade da decomposição, podemos limitar precisamente o que nos é de interesse. Porém, como o sistema cognitivo é, em grande medida, interno ao organismo e escapa à pura observação, faz-se necessário o uso de métodos e técnicas que nos permitam investigar quais são os elementos desse sistema e como ele funciona.

Nesta dissertação, faremos uso da modelagem, que é o processo pelo qual se elabora um modelo, entendendo este como uma representação física, matemática ou lógica de um sistema, entidade, fenômeno ou processo (Zeigler, Praehofer & Kim, 2000:30). Toda modelagem é um processo de simplificação, pois não há a pretensão de que um modelo seja

a reprodução completa de algo, mas sim de que ele seja análogo em algum nível ao que se deseja modelar. É por esse motivo que se pode dizer que modelos são metáforas ou analogias, uma vez que a modelagem necessariamente envolve a exclusão de aspectos contingentes da realidade e o destaque daqueles suficientes ao entendimento do que está sob estudo (Dijkstra & Smedt, 1996b:4).

Cumprе estabelecer a distinção entre teoria e modelo; aqui, esses termos não são sinônimos e não são intercambiáveis, como algumas vezes o são na psicologia (Lachman, 1960) e na lingüística (Crystal, 1988:173). Modelos são analogias, abstrações ou formalizações de alguma entidade ou fenômeno, elaboradas de acordo com finalidades pré-definidas (investigação científica, previsão, controle de processos etc.). Nessa perspectiva, um modelo pode ser visto como um dispositivo de heurística, i.e., um meio através do qual teorias são testadas e aprimoradas.

Diversos interesses motivam a modelagem: destacamos aqui que um modelo pode ser submetido a experimentos, testes e análises que são difíceis ou mesmo impossíveis de se obter a partir do mundo real. Desse modo, trabalhar com uma simplificação de uma entidade ou de um processo tem vantagens como maior facilidade de tratamento e controle, mas possui a desvantagem de requerer um processo de indução, através do qual fazemos inferências (tiramos conseqüências) acerca de um objeto (de estudo) a partir de seu modelo – e somente um modelo adequado permitirá maior segurança quanto às inferências possíveis.

Mas como sabermos se estamos diante de um bom modelo e, portanto, se temos uma boa compreensão do fenômeno em estudo? Essa pergunta só pode ser respondida definitivamente em nível empírico, ou seja, devemos observar se o modelo funciona satisfatoriamente. Neste sentido, um dos recursos que permite instanciar modelos é o da simulação, sendo o computador digital uma ferramenta poderosa e adequada para tal implementação. A simulação não precisa necessariamente ser feita nesse tipo de computador, mas o fato de ele ser *uma máquina de propósitos gerais* o torna um *simulador universal*, ou seja, um sistema capaz de realizar os procedimentos essenciais de qualquer outro sistema ou computador. Simular um modelo, então, é executar um conjunto de instruções, regras, equações ou restrições que geram um comportamento de entrada-saída,

entendendo por modelagem e simulação a construção de um sistema dotado de mecanismos de transição de estados e geração de saídas a partir de entradas recebidas (Zeigler, Praehofer & Kim, 2000:29-30). De fato, o uso da simulação como um instrumento científico traz diversos benefícios – além do automatismo, da velocidade e precisão do computador, ela possibilita verificar se o modelo é adequado e se suas partes são internamente consistentes (Dijkstra & Smedt, 1996b:6). Ademais, através de simulações podemos controlar variáveis e parâmetros do modelo de modo a obtermos com maior facilidade as conseqüências empíricas daquilo que alteramos. A simulação é, portanto, o processo através do qual implementamos um modelo de modo a submetê-lo a experimentos, testes e análises com o propósito de melhor entender um dado comportamento.

Antes de prosseguirmos, uma última consideração faz-se necessária – existe uma distinção, nem sempre feita, entre modelo *em computador* e modelo *computacional* (Dijkstra & Smedt, 1996b:4). Enquanto o primeiro diz respeito à mera implementação ou simulação, o segundo implica que o funcionamento do modelo é regido por procedimentos computacionais, ou seja, há uma relação estrita, uma equivalência forte, entre modelo e computador ou computação. Esta distinção é de especial importância, e será retomada na próxima seção.

Até o presente momento, introduzimos conceitos metodológicos fundamentais, mas ainda resta tratar das motivações para adotá-los na psicolingüística. Em Plunkett & Elman (1997:19-20), encontramos alguns argumentos em favor do uso de modelos e simulações nos estudos da cognição:

- Clareza: a elaboração de um modelo e sua implementação em computador demanda um maior detalhamento do que uma mera descrição verbal;
- Pode-se estudar o desempenho do modelo através da simulação, o que é essencial em casos de sistemas complexos, dinâmicos e não-lineares;
- Os resultados podem sugerir novos experimentos com humanos;
- As dificuldades de se testar teorias e hipóteses no mundo real;

- As simulações ajudam a entender os processos que estão por trás de um comportamento.

Na psicolingüística, por exemplo, o objetivo é entender e explicar como funcionam partes específicas do complexo fenômeno que é a linguagem. Por esse motivo, a idéia de simplificação através de modelos torna-se atraente. Assim, ao investigar os processos envolvidos na leitura, o psicolingüista pode limitar-se a estudar como reconhecemos e compreendemos palavras escritas, dispensando-se de pensar nos processos envolvidos no reconhecimento da fala – trata-se de processos cognitivos distintos, a leitura implica percepção visual mas não a percepção auditiva. O conhecimento acumulado sobre determinado subsistema do processamento lingüístico, então, pode ser representado através de um modelo contendo uma descrição de seu funcionamento. Através de uma simulação, podemos testá-lo e analisá-lo, bem como verificar o seu nível de consistência e adequação.

3. Modelos nos estudos da cognição

Conforme apresentado até agora, o trabalho com modelos e simulações, na psicolingüística, advém da adoção de uma prática metodológica. Nesta seção, o objetivo é explicitar que não se trata apenas disso – o uso de modelos, nas ciências cognitivas em geral, reflete pressupostos teóricos importantes em relação ao funcionamento da cognição. De fato, retomando a distinção feita na seção anterior, uma parte significativa dos modelos de fenômenos cognitivos são computacionais.

A hipótese geral de que *cognição é computação* é amplamente explorada no estudo dos processos cognitivos. Segundo ela, o comportamento inteligente, inclusive o lingüístico, é produto de processos computacionais. Na raiz dessa hipótese temos a analogia entre mente/cérebro e computador – ambos são vistos como sistemas capazes de processar *estímulos* ou *entradas* e emitir *respostas* ou *saídas*, ou seja, são capazes de receber, armazenar, recuperar, transformar e transmitir informações.

A melhor maneira de testar essa hipótese é elaborar um modelo: delimitar um sistema de interesse, concebê-lo como um sistema computacional e simulá-lo em um computador. A existência de dois tipos de computação, no entanto, implica em duas maneiras teórica e metodologicamente distintas de explicar e modelar fenômenos cognitivos.

3.1 A computação digital

Uma possibilidade é ver a cognição em termos de um computador digital², i.e., em termos da chamada arquitetura von Neumann (figura 1.3). Um computador como esse é composto por três blocos principais– um processador, uma memória e dispositivos de entrada e saída, interligados entre si por linhas de comunicação ou *barramentos*:

- Processador: também chamado de unidade de processamento central, é composto por dois dispositivos, a saber: uma *unidade aritmética e lógica*, responsável pela realização sequencial de operações elementares (como adição, subtração etc.); e uma *unidade de controle*, responsável por buscar instruções e dados da memória e controlar o funcionamento de todos os dispositivos.

- Memória: armazena dados a serem processados e programas que contêm as instruções para processá-los. A memória é dividida em espaços com endereços.

- Dispositivos de entrada e de saída: são responsáveis pela comunicação do computador com o ambiente. Por exemplo, receber letras e números por meio de um teclado e projetá-los em um monitor de vídeo.

Em computadores que antecederam essa arquitetura, a memória continha apenas dados e o processador era programado para realizar tarefas específicas. Caso uma outra tarefa precisasse ser realizada, o processador tinha de ser reprogramado para executá-la. Na arquitetura proposta por von Neumann, o processador é um dispositivo geral que busca dados e instruções na memória e realiza as operações aritméticas e lógicas elementares determinadas por um programa contido em sua memória. Esta é a forma de computação mais difundida e utilizada, uma vez que a separação entre memória e processamento, bem como o armazenamento de dados e programas no mesmo dispositivo, torna esse computador uma *máquina de propósitos gerais*.

² “Um computador digital opera sobre unidades distintas quase sempre representadas como ‘presença-ausência’ (*on-off*), que indicam se o valor de uma variável binária é 0 ou 1” (Blackburn, 1997:65).

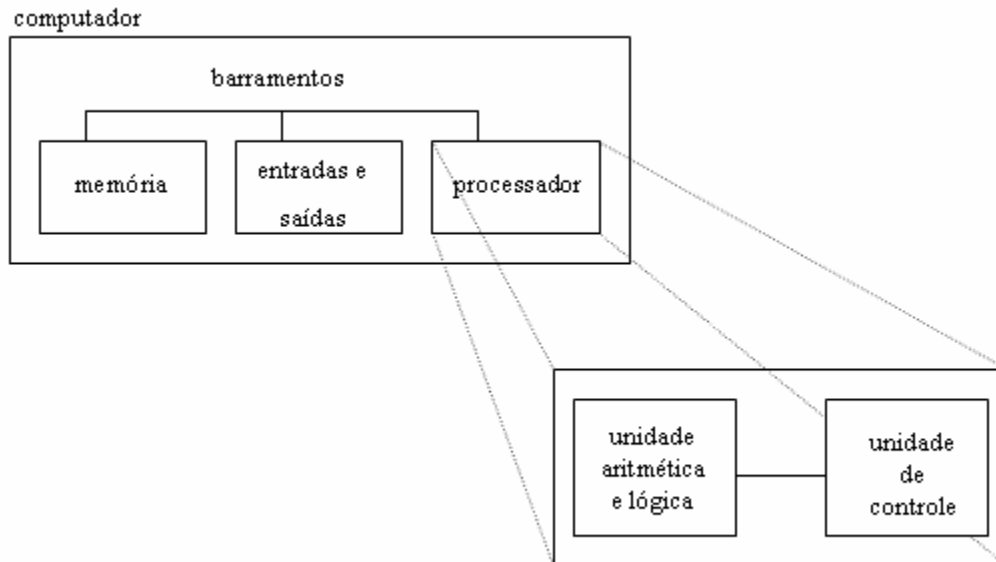


Figura 1.3 Esquema da arquitetura von Neumann: processamento e memória em dispositivos distintos. A memória armazena dados e programas, o processador executa as instruções especificadas por programas.

Resta-nos, ainda, falar sobre a capacidade computacional desse tipo de computador. Supondo que pudéssemos dotá-lo de memória infinita e programá-lo adequadamente, ele seria uma *máquina de Turing universal* (Bechtel, Abrahamsen & Graham, 1998:10).

A máquina de Turing é um dispositivo conceitual (um experimento mental) que foi concebido por Alan Turing (1936) para formalizar a noção de algoritmo ou procedimento efetivo, que está na base da definição de computabilidade. Em termos gerais, uma expressão matemática é computável se existe um procedimento efetivo, i.e., mecânico, capaz de executar um número finito de passos e parar na solução dessa expressão.

Uma máquina de Turing (figura 1.4) é composta por uma fita unidimensional infinita subdividida em células e por uma cabeça de leitura e gravação (um *scanner*) capaz de se mover à direita ou à esquerda da fita e de ler, escrever e apagar símbolos em suas células. Esses símbolos pertencem a um alfabeto finito (por exemplo, 0 e 1), cada célula pode abrigar somente um único símbolo e existe um conjunto finito de instruções (também chamado de programa) que controla o *scanner*, capaz de executar uma operação por vez. O *scanner* pode estar em qualquer um dos estados finitos da máquina, cada qual com a especificação da operação a ser executada de acordo com o símbolo encontrado na fita (Bechtel, Abrahamsen & Graham, 1998:10). Em termos mais precisos, uma função

matemática $f(x)$ é computável se, estando esse número x representado na fita, existe um número finito de operações que leva a máquina a parar na representação do número $f(x)$ (Blackburn, 1997:237).

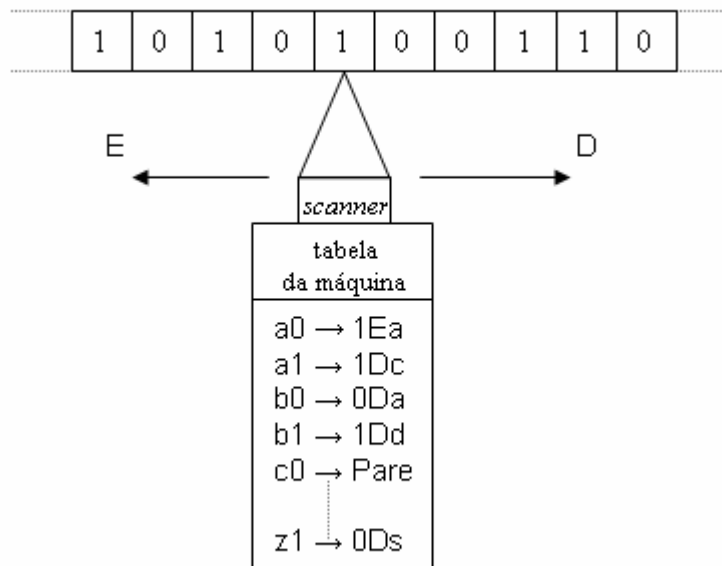


Figura 1.4 Esquema de uma máquina de Turing. A tabela da máquina especifica a operação a ser realizada de acordo com o estado (a, b, c, ...) em que ela se encontra e do símbolo lido na célula. Por exemplo: se a máquina está no estado *a* e na célula está o símbolo *1*, então o *scanner* deve mover uma célula à direita e executar a instrução do estado *c*. Figura adaptada de Bechtel, Abrahamsen & Graham (1998:10).

Conforme descrito acima, trata-se de uma máquina de Turing específica – repare que o *scanner* é guiado por um programa, i.e., um conjunto de instruções especificado na tabela da máquina. Mas isso poderia ser diferente: suponha que tanto os dados de um problema quanto o algoritmo para resolvê-lo sejam colocados na fita. Nesse caso, a máquina pode executar a ação de qualquer programa que seja dado como entrada. Trata-se, portanto, de uma *máquina de Turing universal*, capaz de simular qualquer máquina de Turing específica. Conclui-se, assim, que o poder computacional está limitado, por princípio, a processar qualquer expressão que seja resolvida por um número finito de passos.

Um computador digital, portanto, é um sistema que realiza operações sequenciadas sobre dados representados em sua memória. Essas operações podem ser aritméticas e lógicas e tanto os dados quanto os programas são representados na memória por valores

binários (0 ou 1). Outra característica desse tipo de computação é seu comportamento determinístico, i.e., os mesmos dados de entrada produzem sempre a mesma saída.

Mas o que significa entender, explicar e modelar a cognição segundo os princípios dessa computação? Dado que computação é um conjunto de procedimentos efetivos (mecânicos) que operam sobre símbolos, o estudo da cognição deve contemplar dois aspectos essenciais: regras e representações. De acordo com essa perspectiva, os processos cognitivos são como procedimentos formais (i.e., regras) que operam sobre representações mentais. Defende-se a hipótese de que essa computação é robusta o suficiente para lidar com a complexidade do comportamento e é conhecida como a hipótese de que um sistema físico simbólico é necessário e suficiente para o comportamento inteligente (Newell & Simon, 1976).

Essa é a hipótese nuclear de um programa de pesquisa interdisciplinar conhecido por cognitivismo, de ampla influência nas ciências cognitivas (Varela, 1994; Françoze & Albano, 2003). Na lingüística, por exemplo, uma teoria de cunho cognitivista é a gramática gerativa, que procura investigar as regras e as representações necessárias para gerar o conjunto de todas as sentenças possíveis de uma língua. Nesse caso, a linguagem é uma capacidade cognitiva específica – é o conhecimento gramatical.

O programa cognitivista mostrou-se fecundo, produtivo e hegemônico durante algumas décadas, principalmente depois dos anos de 1950. A partir dos anos de 1980, no entanto, esse programa contabilizou alguns insucessos ao propor-se explicar a totalidade do comportamento em termos de processos computacionais, principalmente porque a computação simbólica não mostra a mesma robustez computacional ao lidar com tarefas cognitivas. O cérebro realiza com extrema rapidez tarefas como o reconhecimento de um rosto familiar em contexto não familiar, enquanto um computador convencional poderia demorar dias para realizar tarefas de menor complexidade (Haykin, 2001:27).

Os processos cognitivos lidam com diversas informações ao mesmo tempo, sendo que estas são comparadas, contrapostas e ponderadas entre si com a finalidade de atingir determinados objetivos de uma maneira econômica e eficiente. Um sistema robusto e flexível como a cognição humana funciona bem na presença de ambigüidades e de informações incompletas e falsas, sendo tais características essenciais para o entendimento

dos mecanismos de processamento (Norman, 1986). Muito já foi discutido em relação às insuficiências dessa abordagem computacional (ver, e.g., Bechtel & Abrahamsen, 2002; Port & van Gelder, 1995), havendo um consenso de que as limitações da computação simbólica, no que tange ao tratamento de fenômenos cognitivos, advêm do caráter serial e simbólico de seu funcionamento. Sistemas como esses são pouco flexíveis, intolerantes a informações ruidosas ou ausentes.

As limitações desse tipo de modelo podem ser tomadas como indício de que os processos cognitivos são diferentes dos processos computacionais do tipo regras e representações. Mas existe uma outra maneira de formalizar uma computação, o que implica em uma outra visão computacional da cognição.

3.2 A neurocomputação

Outra possibilidade é ver a cognição a partir da realidade biológica que a embasa. Na abordagem conexionista (também conhecida por neurocomputação, redes neurais artificiais ou processamento paralelo e distribuído), a computação é inspirada na estrutura e nos processos cerebrais. A idéia central é ver o cérebro como um computador complexo, não-linear e paralelo, dotado de uma capacidade natural de aprendizagem e de adaptação ao ambiente (Haykin, 2001:27). Um neurocomputador (figura 1.5) é um modelo que capta aspectos essenciais do funcionamento do cérebro, sendo formado por dois componentes básicos – neurônios artificiais e conexões:

- neurônios artificiais: são unidades simples de processamento que, assim como neurônios naturais, podem receber estímulos do ambiente ou de outras unidades, bem como propagar uma ativação para o ambiente ou para outras unidades. Cada unidade, portanto, recebe e propaga estímulos: sua capacidade de processamento consiste em somar todos os sinais de entrada recebidos, aplicar o resultado obtido em uma função de ativação e propagar sua resposta. Como os neurônios artificiais compõem um sistema complexo, podemos classificá-los em três tipos: *unidades de entrada*, que recebem estímulos do ambiente; *unidades de saída*, que emitem uma resposta ao ambiente; e *unidades ocultas*, que são internas ao sistema e não interagem com seu exterior.

- conexões: uma conexão liga dois neurônios artificiais, de modo a multiplicar o valor de saída de uma unidade por um *peso de conexão* e remeter o valor obtido para outra

unidade. Desse modo, os pesos podem fortalecer certas conexões enquanto outras são enfraquecidas. As conexões são ajustadas através de um período de treinamento, durante o qual a rede de neurônios é ensinada a responder adequadamente a um ambiente de interesse.

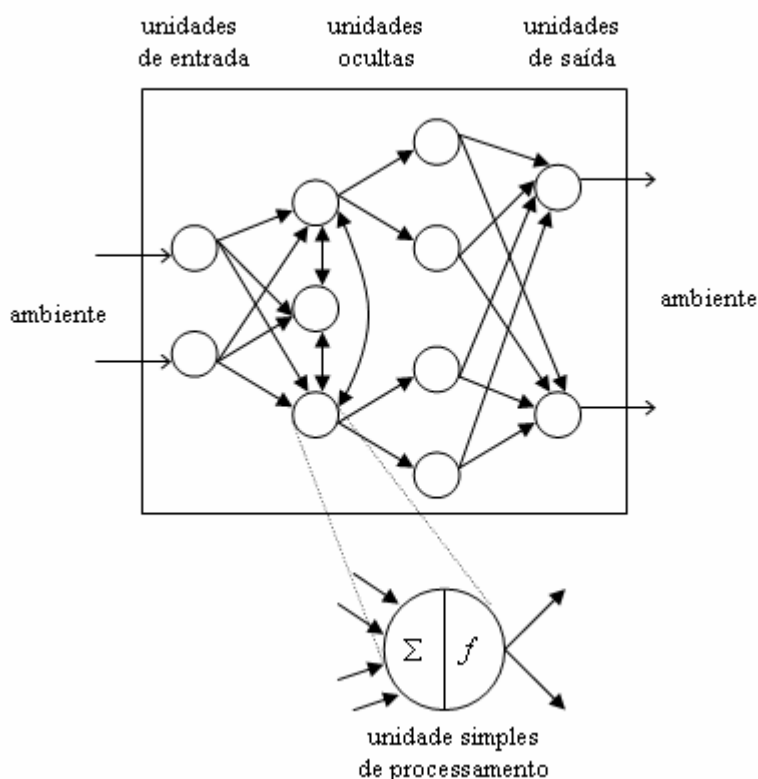


Figura 1.5 Esquema geral de um neurocomputador: unidades simples de processamento interconectadas entre si. As setas cheias indicam o sentido do fluxo de ativação. As setas magras indicam os sinais que entram e que saem do sistema. Neurônios artificiais (no detalhe) somam os sinais de entrada recebidos e uma função dá o valor de sua ativação, que então é propagada.

Repare que, na neurocomputação, não há uma unidade de processamento central, não há um dispositivo de memória e não há a manipulação algorítmica de símbolos como há na computação simbólica. O processamento e a memória estão distribuídos nos componentes da rede neural – seus neurônios artificiais e suas conexões. Esse sistema processa informações localmente (nas unidades) e produz um comportamento global que resulta das diversas computações locais e da interação entre estas. Além disso, o processamento se dá de forma paralela, ou seja, a computação feita por um conjunto de unidades ocorre simultaneamente. Outro aspecto essencial de um sistema conexionista diz respeito à sua

memória: ela não é subdividida em endereços – o conhecimento, nesse sistema, está distribuído nos pesos de conexão. É o padrão de conexões que permite ao sistema responder adequadamente ao ambiente em que foi treinado.

Na próxima seção, apresentamos a abordagem conexionista em detalhes. Por hora, é importante salientar suas diferenças em relação à computação simbólica: processamento paralelo, distribuído e capacidade de aprendizagem. Há diversas motivações para o uso desse tipo de computação: a inspiração biológica, de fato, produz uma computação robusta que exhibe comportamentos globais semelhantes aos do cérebro. Por exemplo, a *tolerância a falhas* e a *degradação suave* da performance. O cérebro não pára de funcionar quando recebe muitas informações e maior demanda de processamento, mas funciona parcialmente, ignorando parte da demanda ou das informações. Mesmo em casos de lesão, o cérebro continua a funcionar, embora possa apresentar ineficiências (de acordo com o tipo e o grau da lesão). Comportamentos semelhantes podem ser encontrados em redes neurais artificiais, cujo processamento paralelo e distribuído não permite que o sistema entre em colapso, mesmo que exposto à sobrecarga de demanda, a informações em excesso e mesmo que parte de suas conexões e unidades tenha sido danificada de alguma maneira. Isso se deve, em grande parte, a sua capacidade de processar as informações de forma paralela e distribuída, de modo que a rede exiba um comportamento global mesmo com informações imprecisas ou perda de algumas unidades e conexões. Computadores simbólicos, ao contrário, teriam uma perda considerável na qualidade do processamento, caso houvesse danos em símbolos ou regras de um algoritmo. Assim, as redes neurais constituem um modelo da cognição com a robustez e a flexibilidade difíceis de obter em modelos computacionais simbólicos.

Entender, explicar e modelar a cognição segundo os princípios da neurocomputação, portanto, é reconhecer que os processos cognitivos são computações locais simples (integração/ativação) que interagem para produzir um determinado comportamento, ou seja, a cognição é a emergência de estados globais em uma rede de unidades simples de processamento (Varela, 1994:62).

Já vimos a importância dos trabalhos de Turing e de von Neumann para o desenvolvimento da computação digital. Aqui, veremos que os princípios da

neurocomputação remontam ao trabalho de McCulloch & Pitts (1943), o primeiro a propor um modelo formal do neurônio biológico. Tal como proposto inicialmente, trata-se de um modelo simples: o neurônio recebe entradas e produz uma saída, sendo um dispositivo binário cuja saída pode ser *1* ou *0*, *ativação* ou *não-ativação*, a depender das *entradas recebidas* e de um *limiar de ativação*. McCulloch & Pitts mostraram como combinações desses neurônios poderiam realizar as operações lógicas *e*, *ou* e *não*. Demonstraram, em seguida, que qualquer processo passível de ser realizado por um número finito dessas operações poderia ser implementado por redes desses neurônios formais e estas, desde que dotadas de grande capacidade de memória, poderiam ter o mesmo poder de uma *máquina de Turing universal* (Bechtel & Abrahamsen, 2002:3). Esse último argumento é de fundamental importância, uma vez que equipara matematicamente o poder computacional da neurocomputação e da computação digital.

Cinco anos após a publicação desse artigo, McCulloch apresenta na conferência “Cerebral Mechanisms in Behavior” uma palestra intitulada “Why the Mind is in the Head”, na qual traça alguns paralelos entre o sistema nervoso e dispositivos lógicos na tentativa de explicar como o cérebro processa informação (Gardner, 1996:25). A ideia essencial, segundo Kovács (1996:29) era:

A inteligência é equivalente ao cálculo de predicados que por sua vez pode ser implementado por funções booleanas. Por outro lado, o sistema nervoso é composto de redes de neurônios, que com as devidas simplificações, tem a capacidade básica de implementar essas funções booleanas. Conclusão: a ligação entre inteligência e atividade nervosa fica estabelecida de forma científica.

Após o artigo seminal de McCulloch & Pitts (1943), trabalhos subsequentes trataram de investigar a capacidade de redes neurais em desempenhar processos cognitivos. Em específico, explorou-se a capacidade dessas redes em reconhecer padrões (Pitts & McCulloch, 1947; Rosenblatt, 1958; Selfridge, 1959) e construir memórias associativas (Taylor, 1956; Anderson, 1972; Kohonen, 1972; Nakano, 1972).

O conexionismo desenvolve-se com vigor no âmbito das ciências cognitivas desde a década de 1980, seja por razões internas como o aprimoramento metodológico, seja por razões externas como o enfraquecimento do cognitivismo clássico (ver Françoze & Albano, 2004). Embora a modelagem de fenômenos lingüísticos tenha sempre acompanhado o desenvolvimento do conexionismo, somente nos últimos anos observamos a aproximação

da psicolingüística com a abordagem conexionista, evidenciada pela publicação internacional *Connectionist Psycholinguistics* (Christiansen & Chater, 2001) e nacional *Rumo à Psicolingüística Conexionista* (Rossa & Rossa, 2004). No Brasil, o interesse da psicolingüística nessa abordagem é ainda recente, o que explica os esparsos trabalhos e publicações na área. O trabalho desenvolvido nesta dissertação insere-se no contexto da *Psicolingüística Conexionista* e procura investigar como a linguagem e processos cognitivos podem decorrer do processamento paralelo e distribuído por unidades maciçamente conectadas.

4. A abordagem conexionista

Modelos conexionistas ou redes neurais artificiais são modelos computacionais cujo processamento é inspirado nas estruturas e processos cerebrais. Unidades simples de processamento, ou neurônios formais, recebem e transmitem ativações (valores numéricos) através de camadas de conexões (ponderações numéricas), assim como neurônios transmitem informações através de sinapses. Temos, portanto, um modelo que simplifica e mimetiza aspectos fundamentais do funcionamento do cérebro. Uma rede de neurônios pode passar por um *período de aprendizagem*, durante o qual as conexões são ajustadas de modo a mapear dados de entrada a saídas com o mínimo de erro. Após o final desse processo, análises estatísticas permitem avaliar a performance do sistema, bem como a qualidade do conhecimento adquirido. Temos, então, um modelo do fenômeno que estamos investigando.

A computação paralela e distribuída por unidades simples de processamento permite replicar características importantes da cognição, como aprendizagem, generalização, reconhecimento de padrões, associação e tolerância a falhas. A representação e a memória desenvolvidas por esses sistemas através da aprendizagem têm um caráter altamente distribuído e atrelado ao processamento – não há como localizar informações específicas como símbolos no interior da rede, mas há como analisar seu comportamento em termos do padrão de suas ativações aos estímulos recebidos. Por esse motivo a abordagem conexionista também é chamada de subsimbólica.

Esse tipo de sistema é interessante não apenas por imitar a estrutura e funcionamento do cérebro, mas por possibilitar a emergência de fenômenos de tipo cognitivo a partir de

um sistema computacional simples em seus fundamentos. São modelos cuja inspiração neural visa dar uma explicação psicológica e computacional para a cognição, i.e.,

4.1 Modelos de neurônios artificiais

Um neurônio é uma unidade simples de processamento composta por três elementos básicos (ver figura 1.6):

- um conjunto de conexões, cada qual caracterizada por um *peso* ou *força de conexão* w_{ij} entre um neurônio i e um sinal de entrada j .
- um somador que soma os sinais de entrada recebidos, devidamente ponderados pelas conexões do neurônio.
- uma função de ativação que produz a saída do neurônio.

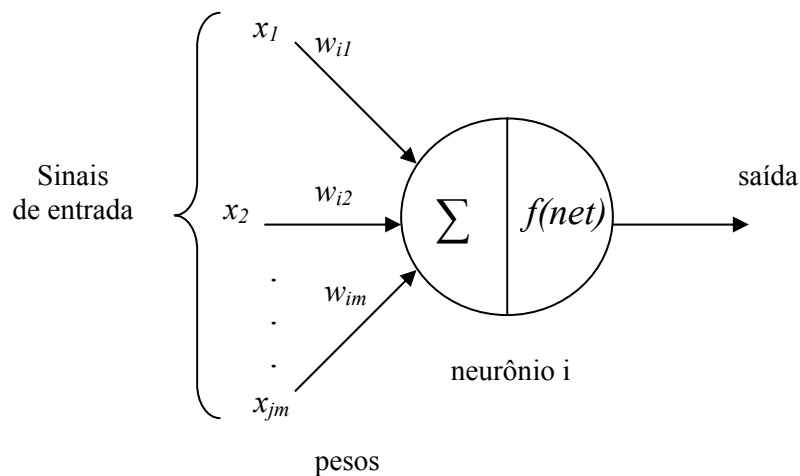


Figura 1.6 Esquema de um neurônio i : as entradas são somadas e aplicadas em uma função de ativação que produz a saída.

Formalmente, temos que um neurônio i recebe m sinais, cada qual ponderado pela respectiva conexão w_{ij} . A primeira operação feita é a soma dos produtos, dada pela equação

$$net_i = \sum_{j=1}^m w_{ij}x_j$$

na qual i é o neurônio receptor, net_i é o somatório de todos os sinais x_j ponderados pela conexão w_{ij} . A segunda operação é uma função de ativação f aplicada sobre o valor obtido

em net_i . Existem diferentes tipos de função de ativação que podem ser usados, a depender do problema a ser modelado.

Conforme já mencionamos, o neurônio proposto por McCulloch & Pitts (1943) tem um comportamento binário. Nesse modelo, a saída de um neurônio assume o valor 1 se net_i for não-negativo, caso contrário sua saída é 0. Esta é a chamada função degrau (figura 1.7), dada pela equação:

$$f(net_i) = \begin{cases} 1 & \text{se } net_i \geq 0 \\ 0 & \text{se } net_i < 0 \end{cases}$$

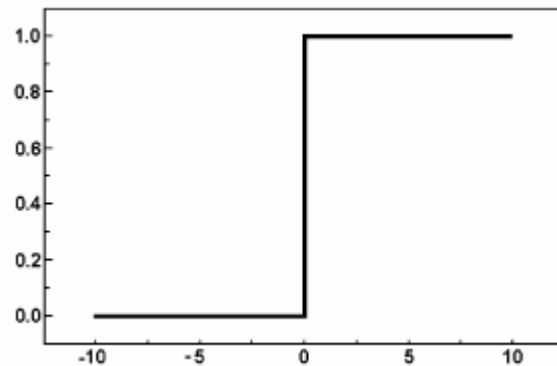


Figura 1.7 Gráfico da função degrau: a ativação (eixo y) da unidade diante das entradas (eixo x).

Outro tipo de função que pode ser usada é a sigmóide, que tem a forma de um ‘s’. Trata-se de uma função crescente que exibe um comportamento balanceado entre linearidade e não-linearidade, e pode ser dada pela equação e seu respectivo gráfico (figura 1.8):

$$f(net_i) = \frac{1}{1 + e^{-net_i}}$$

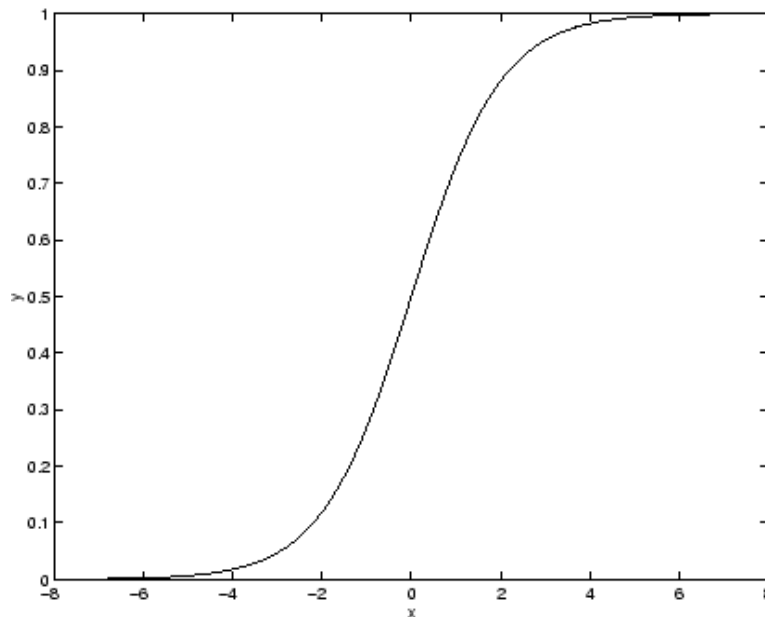


Figura 1.8 Gráfico da função sigmóide: ativação não-linear (eixo y) da unidade diante das entradas (eixo x).

Observamos que a ativação dada pela função sigmóide pode ser “tudo ou nada” para entradas que *não* estão no intervalo $[-5,5]$, dentro do qual a função assume um caráter não-linear. A ativação “tudo ou nada” permite que a unidade responda categorialmente às entradas, enquanto o outro tipo permite uma resposta gradual e reflete a sensibilidade da unidade diante de entradas que estão dentro desse intervalo (Elman et al., 1996:53).

O modelo de neurônio artificial descrito acima é chamado de *neurônio computacional*, pois processa sinais de entrada e produz uma saída. São os neurônios ocultos e de saída do sistema. Geralmente, os neurônios da mesma camada possuem a mesma função de ativação. É comum o uso de função sigmóide na camada oculta e de função linear na camada de saída. Os neurônios da camada de entrada exercem um papel “sensorial” que consiste em receber, em um dado instante, um sinal do ambiente, que então é transmitido para a próxima camada de neurônios, sem qualquer modificação. Ou seja, não há computação: o valor de sua ativação é simplesmente o sinal recebido do ambiente.

4.2 Redes de neurônios – arquiteturas

O padrão de conectividade das unidades constitui a arquitetura da rede. Observe na figura 1.9 que as unidades estão organizadas em três camadas: uma de entrada e outra de saída, entre as quais existe uma camada com unidades ocultas, internas em relação ao ambiente da rede. Unidades da mesma camada não se conectam entre si e as ativações

seguem o fluxo indicado pelas setas. Trata-se da arquitetura conhecida como *perceptron de múltiplas camadas* que, aliada ao algoritmo de retropropagação do erro, é amplamente utilizada em projetos de redes neurais. Considerando que suas unidades ocultas tenham ativação não-linear, uma rede como essa tem a capacidade de *aproximação universal*, i.e., pode realizar um mapeamento não-linear de entrada-saída de natureza geral (Haykin, 2001:234). A camada de unidades ocultas é capaz de extrair estatísticas de ordem elevada e características importantes do ambiente (Haykin, 2001:47).

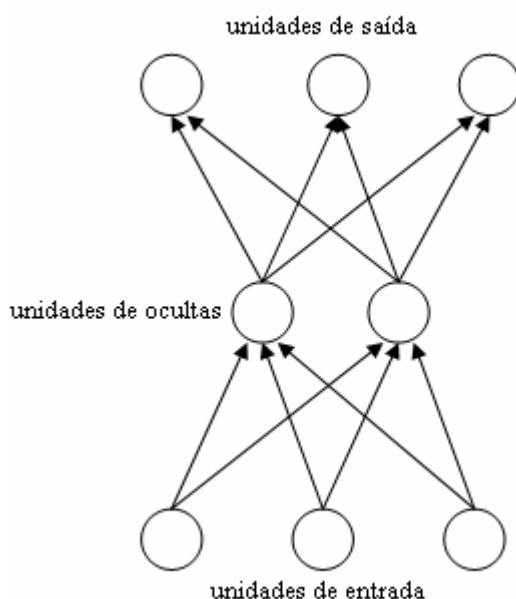


Figura 1.9 Esquema de um *perceptron de múltiplas camadas*: repare que as unidades da mesma camada não se conectam entre si. O fluxo do sinal propaga no sentido indicado pelas setas.

Retomando os conceitos da sistêmica, apresentados na seção 2, podemos classificar uma rede neural com unidades ocultas como um modelo de espaço de estados. Quando do processamento de um sinal de entrada na rede, cada unidade oculta produz uma ativação, que caracteriza seu estado. Os neurônios ocultos, portanto, definem o espaço de estados da rede³.

Existem diversos tipos de arquitetura de redes neurais, cada qual com sua funcionalidade específica. A rede da figura 1.9, por exemplo, realiza computações

³ Formalmente, o espaço de estados é um espaço n -dimensional no qual cada dimensão corresponde a uma das variáveis do sistema. Cada possível estado do sistema tem um ponto em cada dimensão desse espaço (Bechtel & Abrahamsen, 2002:359).

unicamente a partir dos sinais de entrada recebidos, ou seja, o mapeamento realizado é estático e a rede não reutiliza a informação que ela mesma computou. Em *redes recorrentes* existe pelo menos um laço de realimentação, i.e., informações (ativações) produzidas pela rede realimentam alguma parte do próprio sistema. Há diversas possibilidades de implementar recorrência: por exemplo, podemos ter realimentação dos neurônios de saída para a camada de entrada ou para a camada oculta, ou realimentação da camada oculta para a entrada ou para outra camada oculta. Pode haver, também, auto-realimentação, no caso da saída de um neurônio ser realimentada para sua própria entrada. Nesta dissertação, utilizamos uma *Rede Recorrente Simples* (RRS; Elman, 1990), mostrada na figura 1.10. A RRS tem uma camada de *unidades de contexto*, que armazena as ativações da camada oculta para realimentá-la juntamente com a próxima entrada. Ou seja, o espaço de estados de uma RRS é afetado tanto pela entrada atual da rede quanto pelo seu estado anterior, o que a torna capaz de realizar mapeamentos que necessitam de informações temporais.

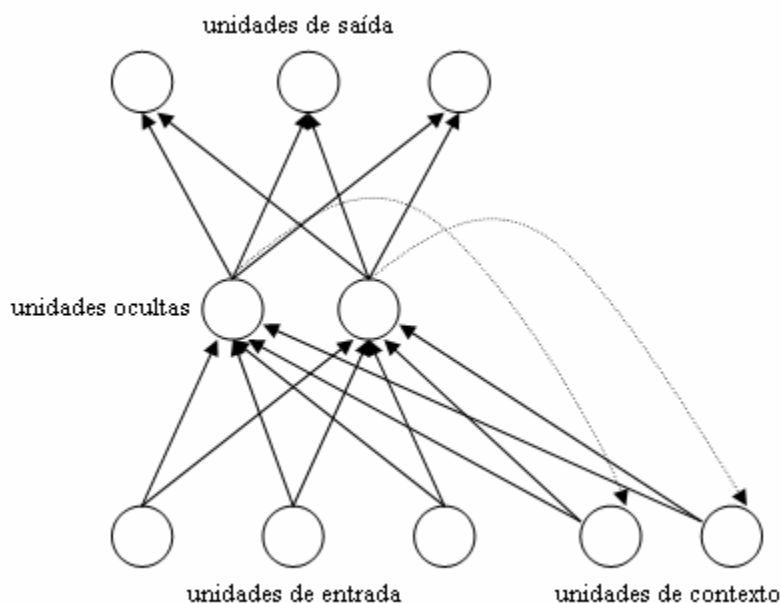


Figura 1.10 Esquema da rede recorrente simples: cada unidade oculta se conecta com uma unidade de contexto, ponderada por um peso de conexão fixo em 1. Repare que as unidades de contexto propagam suas ativações para todas unidades ocultas.

4.3 Aprendizagem

Um dos principais atrativos dos sistemas conexionistas é sua capacidade de aprender a partir da experiência. Chamamos de experiência a apresentação de dados de entrada que são processados pelo sistema. Chamamos de aprendizagem ou treinamento o processo pelo qual os pesos das conexões entre as unidades são ajustados, de modo a produzir saídas esperadas em relação às entradas recebidas.

O objetivo do processo de treinamento é obter uma rede neural que faça um mapeamento de entrada-saída, característico de algum fenômeno bem delimitado (por exemplo, mapear letras a palavras). A rede recebe informações de entrada, processa-as internamente e produz outras informações como saída. Como não é possível prever o conjunto de pesos adequados para fazer um determinado mapeamento, no estado inicial da rede os pesos são determinados aleatoriamente, de preferência números reais próximos, mas diferentes, de zero. Dessa forma, qualquer sinal que entre no sistema resulta em saídas relativamente parecidas. A finalidade do treinamento, então, é obter um conjunto de pesos de conexão tal que, para todos os pares de entrada-saída que caracterizam um fenômeno, o sistema produza as saídas adequadamente.

Esse processo de treinamento é um tipo de *aprendizagem supervisionada*, através do qual o sistema é “ensinado”, “corrigido”, de maneira a emitir uma saída com o mínimo de erro. Ou seja, uma rede neural é treinada para realizar uma tarefa – dada uma entrada, pretende-se que suas conexões sejam tais que ela produza uma saída esperada. O treinamento chega ao final quando o ajuste dos pesos não leva mais a uma redução significativa do erro.

A arquitetura do sistema somada à eficácia do algoritmo de treinamento permite que a camada oculta processe entradas e produza as saídas esperadas, de maneira que podemos dizer que a rede neural *aprendeu* a tarefa ou *internalizou* (construiu uma *representação interna*) o conhecimento acerca do ambiente no qual foi treinada.

O procedimento usual para modificar os pesos de redes como o *perceptron de múltiplas camadas* e a rede recorrente simples é a *regra delta generalizada* aliada ao algoritmo de *retropropagação do erro* (*error backpropagation*; Rumelhart, Hinton & Willians, 1986). O algoritmo recebe esse nome porque o erro produzido na saída é usado

para estimar o ajuste de todos os pesos da rede. Para os detalhes matemáticos e computacionais envolvidos nesse procedimento, remetemos a Rumelhart, Hinton & Williams (1986) ou ao capítulo 4 de Haykin (2001). Aqui, vemos que o ajuste $\Delta w_{ij}(n)$ aplicado ao peso w_{ij} em uma iteração n (passo de tempo no qual o n -ésimo exemplo de treinamento é apresentado à rede) é dado pela *regra delta*

$$\Delta w_{ij}(n) = -\eta \frac{\partial E_i(n)}{\partial w_{ij}(n)}$$

na qual n é a iteração, Δw_{ij} é o ajuste do peso, η é o parâmetro da taxa de aprendizagem, E_i é o erro da saída do neurônio i e w_{ij} é o peso que liga os neurônios j e i . Não entraremos em detalhes, o importante é entender que a fração mede como o erro se comporta quando alteramos o peso. Por isso, o ajuste de peso Δw_{ij} (em uma iteração n) é uma proporção da taxa de variação do erro (dada pela derivada parcial $\partial E_i(n)$) em relação à taxa de variação do peso (dada pela derivada parcial $\partial w_{ij}(n)$), dimensionada pelo parâmetro η . O sinal negativo de η indica a busca de uma mudança de peso que reduza o valor de $E_i(n)$. Essa equação expressa algo que pode ser entendido intuitivamente: o erro da rede varia de acordo com os pesos. Caso os pesos não sejam adequados, a regra permite mensurar o ajuste necessário para minimizar o erro em cada conexão. Quanto menor for η , menor será o ajuste de pesos entre as iterações (considerando uma correção de pesos feita a cada iteração). O uso de um η próximo de 1 indica um ajuste de alta proporção, o que não é interessante: a cada iteração o sinal de entrada do sistema se altera. Um ajuste alto dos pesos a cada iteração pode aumentar a taxa de erro, tornando a rede instável. O uso de um termo de *momento* permite um aumento da taxa de aprendizagem sem que haja perda da estabilidade no ajuste dos pesos (Haykin, 2001:196-7). Esse termo determina qual proporção da mudança de peso na iteração anterior ($n-1$) será usada na iteração atual (n), veja na equação

$$\Delta w_{ij}(n) = \alpha \Delta w_{ij}(n-1) - \eta \frac{\partial E_i(n)}{\partial w_{ij}(n)}$$

na qual α é uma constante de *momento*.

4.4 Um exemplo prático

Para exemplificar, veremos uma implementação da função booleana *ou-exclusivo*, simbolizada por ' $\underline{\vee}$ ' e dada pela tabela 1.1. Esse operador realiza uma *disjunção exclusiva*, de modo a combinar duas variáveis binárias p e q na expressão ' $p \underline{\vee} q$ '. As variáveis podem assumir o valor 1 ou 0 e a expressão ' $p \underline{\vee} q$ ' assume o valor 1 somente se uma das variáveis é 1. Trata-se de um exemplo clássico no campo de redes neurais, pois é o problema mais simples que necessita de uma rede com unidades ocultas.

Tabela 1.1

p	q	$p \underline{\vee} q$
0	0	0
0	1	1
1	0	1
1	1	0

A rede utilizada (figura 1.11) tem duas unidades de entrada (i_1 e i_2), duas unidades ocultas (h_1 e h_2) e uma unidade de saída (o_1). Cada unidade de entrada corresponde a uma variável do problema: digamos que i_1 corresponda a p e i_2 a q . Na saída, queremos que o sistema produza corretamente o valor da expressão ' $p \underline{\vee} q$ ' para todas as combinações dos valores de p e q .

O treinamento dessa rede consiste em receber os vetores de entrada 00, 10, 01 e 00 aleatoriamente por diversas vezes, tantas quantas forem necessárias para o sistema aprender a tarefa. Ao receber um vetor, as ativações se propagam e a rede produz uma saída. A saída obtida é comparada à saída-alvo pelo algoritmo de retropropagação do erro, que utiliza a possível discrepância como base para a alteração dos pesos, visando minimizar o erro. O objetivo, portanto, é obter seis pesos de conexão de modo que, recebido qualquer um dos quatro vetores de entrada, o sistema produza a saída próxima ao desejado.

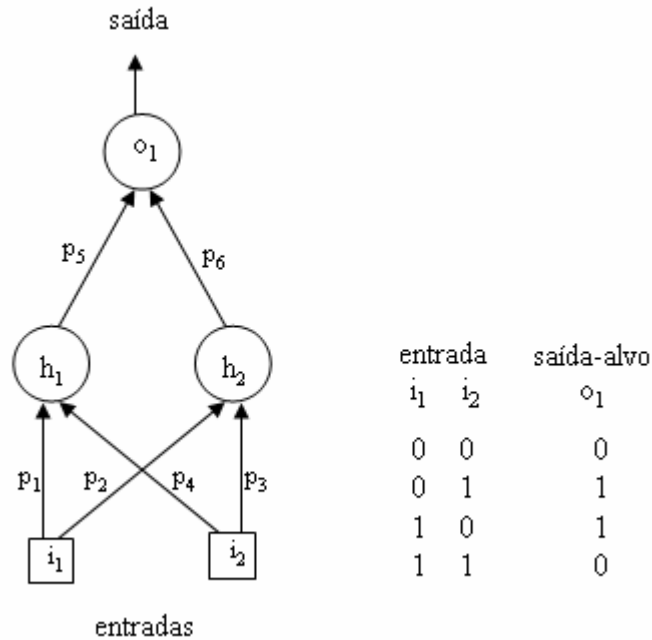


Figura 1.11 Rede neural para a implementação da função *ou-exclusivo*.

Antes de começar o treinamento, é preciso iniciar o sistema com pesos aleatórios. Sorteamos números no intervalo $[-0,2, 0,2]$. Também é preciso determinar os parâmetros do algoritmo de retropropagação do erro: ajustamos a taxa de aprendizagem para 0,25 e a taxa de momento para 0,9.

No estado inicial da rede, os pesos são aleatórios (tabela 1.2) e não há conhecimento algum sobre o mapeamento entrada-saída: repare nas ativações do neurônio o_1 (tabela 1.3) para cada evento. Conforme o treinamento progride, o ajuste dos pesos resulta em melhores ativações do neurônio o_1 para cada evento. Paramos o treinamento após 700 épocas, i.e., 700 apresentações do conjunto de eventos, mas poderíamos ter prosseguido até obter saídas ainda mais próximas do alvo.

Tabela 1.2
Pesos de conexão da rede neural

Pesos	Valor dos pesos de conexão / época		
	início	600	700
p1	-0,070	-5,015	-5,698
p2	0,120	-1,456	-3,356
p3	0,038	-1,483	-3,357
p4	0,116	-5,039	-5,711
p5	-0,049	-5,212	-7,079
p6	-0,004	3,109	6,593

Tabela 1.3
Ativação do neurônio de saída para cada evento

Evento	Ativação do neurônio o_1 / época		
	início	600	700
00→0	0,4599	0,1392	0,0722
01→1	0,4595	0,7057	0,9117
10→1	0,4601	0,7098	0,9116
11→0	0,4597	0,4957	0,1190

Através deste exemplo, mostramos que a mudança de pesos, durante o treinamento, resulta em melhores ativações do neurônio de saída. Dessa forma, podemos dizer que a rede aprendeu a função *ou-exclusivo* pela sua experiência com os eventos que caracterizam essa função, ou seja, o sistema aprende uma regra lógica (nesse caso), sem que houvesse o conhecimento prévio desta.

Terminamos, assim, este capítulo, cujo objetivo principal foi fornecer os fundamentos necessários ao entendimento do trabalho relatado e discutido no restante desta dissertação.

CAPÍTULO II

Conhecimento sintático-semântico e processamento de sentenças em redes neurais

1. Treinando redes neurais com sentenças

No capítulo anterior, apresentamos a abordagem conexionista, bem como alguns aspectos que consideramos necessários ao entendimento das simulações realizadas. Não mostramos, contudo, como aqueles conceitos se aplicam ao estudo do processamento lingüístico. Essa lacuna é preenchida neste capítulo, no qual mostramos uma abordagem conexionista da linguagem, tal como desenvolvida nos trabalhos de Elman (1990, 1991, 1995).

Vimos anteriormente que a modelagem através de redes neurais desenvolveu-se principalmente a partir dos anos de 1980, em parte impulsionada pelo advento do algoritmo de retropropagação do erro como uma maneira de ajustar os pesos de conexão. Esse algoritmo permite que o *perceptron de múltiplas camadas* construa uma rica representação interna do ambiente (Rumelhart, Hinton & Williams, 1986). A questão de como modelar fenômenos temporais nesse tipo de rede mostrou-se de especial interesse, uma vez que diversos fenômenos físicos, comportamentais e cognitivos ocorrem no tempo. Em Elman (1990), encontramos uma revisão dessa questão, bem como a proposta de uma modificação que resulta em uma rede capaz de reconhecer padrões temporais.

A rede recorrente simples (RRS; Elman, 1990) é considerada um dos modelos conexionistas mais bem sucedidos nos estudos da cognição (Landy, 2004). Sua arquitetura é uma resposta ao problema do processamento temporal. A solução vem pela introdução de uma camada de contexto, responsável por executar uma recorrência simples – copiar o estado interno do sistema para torná-lo disponível no passo seguinte. Isso permite que a camada oculta processe uma entrada juntamente com seu estado interno anterior, o que significa que o tempo está representado intrinsecamente por seus efeitos no processamento (Elman, 1990:182). Além disso, o mapeamento de entrada-saída que a RRS é treinada a fazer necessita de informações temporais – o objetivo é emitir como saída uma previsão de sua próxima entrada. Dessa forma, a RRS desenvolve representações que refletem a

demanda por informações contidas no contexto de seus estados internos, o que torna a memória desse sistema fortemente atrelada ao processamento da tarefa e as representações altamente dependentes do contexto (Elman, 1990:194).

Dentre os experimentos relatados em Elman (1990), o principal mostra os resultados de uma RRS treinada com sentenças simples de 2 e 3 palavras, geradas a partir de um léxico composto por 29 itens (tabela 2.1) e 16 padrões de combinação entre eles (tabela 2.2). Embora simples, esse experimento mostrou-se promissor: a RRS foi capaz de construir representações internas que refletem conhecimento lingüístico como classes de palavras (verbos e nomes) e tipos semânticos (animados e inanimados), ou seja, o processo de aprendizagem permitiu que a rede captasse informação lingüística presente nas sentenças⁴.

Tabela 2.1
categorias de itens lexicais

Categorias	Membros
Nomes	
Humanos	woman, girl, man, boy
Animais	cat, dog, mouse, lion, monster, dragon
Agressivos	lion, monster, dragon
Inanimados	rock, car, book, plate, glass, cookie, sandwich, bread
Quebráveis	plate, glass
Alimentos	cookie, sandwich, bread
Verbos	
Transitivos	chase, like, eat, break, smash, move, see, smell
Intransitivos	sleep, think, exist
Ag/pac	move, break, smash
Destrutivos	break, smash
Perceptuais	see, smell
Comer	Eat

Tabela 2.2
padrões para o programa gerador de sentenças

Palavra 1 Nome	Palavra 2 Verbo	Palavra 3 Nome
Humano	Comer	Alimento
Humano	Perceptuais	Inanimado
Humano	Destrutivos	Quebrável
Humano	Intransitivos	
Humano	Transitivos	Humano
Humano	Ag/pac	Inanimado
Humano	Ag/pac	
Animal	Comer	Alimento
Animal	Transitivos	Animal
Animal	Ag/pac	Inanimado
Animal	Ag/pac	
Inanimado	Ag/pac	
Agressivo	Destrutivos	Quebrável
Agressivo	Comer	Humano
Agressivo	Comer	Animal
Agressivo	Comer	Alimento

Com o intuito de investigar como a RRS reagiria a padrões mais complexos, o trabalho seguinte de Elman (1991) apresenta experimentos realizados com sentenças de maior complexidade sintática (3 a 16 palavras), contendo concordância, subordinação (recursividade) e dependências de longa distância. O objetivo desses experimentos era testar a capacidade da RRS de aprender, apenas pela co-ocorrência lexical, conhecimento

⁴ Esse experimento é mostrado em detalhes na seção 3 deste capítulo.

sintático tradicionalmente considerado abstrato e impossível de ser aprendido apenas pela informação superficial das sentenças. A aprendizagem dessa sintaxe não foi simples, havendo a necessidade de garantir (seja pelo aumento gradativo da complexidade do conjunto de treinamento, seja pelo aumento gradual da ‘janela’ de memória do contexto) que as sentenças simples fossem aprendidas primeiro, para então a complexidade sintática ser, de fato, assimilada pela rede⁵. O processamento de sentenças e as representações internas foram avaliados em termos de trajetórias pelo espaço de estados, i.e., padrões de ativação das unidades ocultas no tempo. Uma explicação mais teórica, no entanto, aparece posteriormente em Elman (1995), quando a RRS é vista sob a ótica da teoria dos sistemas dinâmicos. Nessa perspectiva, o processamento de sentenças na RRS é visto como o comportamento de um sistema dinâmico: as representações são regiões do espaço de estados e as regras gramaticais estão incorporadas na dinâmica do sistema, a qual facilita ou dificulta certas transições (Elman, 1995). Desse modo, explicou-se o conhecimento gramatical complexo pelas trajetórias das sentenças processadas no espaço de estados da rede, observando que o comportamento do sistema refletia a presença de atratores sintáticos.

2. Questões de interesse

Trataremos aqui sobre o conhecimento gramatical, entendendo este como as informações sintáticas e semânticas que permite a produção e compreensão das sentenças de uma língua. Especificamente, nosso interesse está em relações gramaticais que ocorrem no nível de sentenças simples e, por este motivo, nosso enfoque é o experimento de Elman (1990).

Um tipo de *análise em constituintes* considera a sentença como o maior constituinte estrutural de uma língua, que pode ser decomposta primeiramente em dois constituintes – sujeito e predicado – e, em seguida, este último pode ser decomposto em verbo e objeto. Veja o caso da sentença *o menino empurrou a menina*: podemos decompô-la em sujeito (*o*

⁵ Elman (1993) explica em detalhes a aprendizagem da complexidade sintática em RRS, correlacionando os resultados obtidos com evidências de uma relação entre aprendizagem linguística e restrições maturacionais (capacidade de memória) durante o desenvolvimento da criança, conhecido na literatura como a hipótese do “less is more” (Newport, 1988; 1990).

menino) e predicado (*empurrou a menina*), que ainda pode ser decomposto em verbo e objeto. Não é necessário, entretanto, que todos os constituintes sejam expressos – a elipse é permitida pela gramática quando o contexto está claro, tendo função estilística, enfática ou de economia.

Os constituintes de uma sentença têm relações estruturais ou gramaticais entre si, refletidas pela sua posição na sentença em línguas nas quais a ordem de palavras é fixa. Compare as sentenças *o menino empurrou a menina* e *a menina empurrou o menino*: a diferença de sentido é explicada pela relação gramatical entre os constituintes que antecedem e sucedem ao verbo, a saber sujeito e objeto.

Uma das maneiras de descrever relações gramaticais desse tipo é pensar em estrutura argumental e grade temática. A estrutura argumental é uma especificação do número de argumentos (posições de variáveis) que um verbo requer (Haegeman, 1991) e pode ser visualizada como:

Verbo: $\text{arg}_1 (\text{arg}_2 \text{arg}_3 \dots)$

na qual $\text{arg}_1 \dots \text{arg}_n$ representam os argumentos de um verbo. O primeiro argumento, que é obrigatório, é o sujeito do verbo e os outros representam os demais termos de uma sentença. Assim, um verbo intransitivo como *sleep* pode ter sua estrutura argumental representada como⁶:

Sleep: verbo; 1
 nome

No caso de verbos transitivos:

break: verbo; 1 (2)
 nome nome

chase: verbo; 1 2
 nome nome

A depender do verbo, permite-se a elipse do objeto, como ocorre acima com o verbo *break* – e isso é representado pela inclusão de arg_2 entre parênteses. O conjunto de treinamento de Elman (1990), portanto, contém os seguintes tipos de estrutura argumental:

⁶ Exemplificamos segundo as estruturas argumentais das sentenças usadas em Elman (1990).

Tabela 2.3

<i>Intransitivo</i>		
verbo;	1	
	nome	
<i>Transitivos Obrigatórios</i>		
verbo;	1	2
	nome	nome
<i>Transitivos Opcionais</i>		
verbo;	1	(2)
	nome	nome

A grade temática dos verbos, por sua vez, representa as relações semânticas entre os verbos e seus argumentos (Haegeman, 1991). Uma lista não exaustiva de papéis temáticos é a seguinte:

agente	aquele que intencionalmente inicia a ação expressa pelo verbo
paciente	pessoa ou coisa que sofre a ação expressa pelo verbo
tema	pessoa ou coisa movida pela ação expressa pelo verbo
experienciador	a entidade que experiencia algum estado (psicológico) expresso pelo verbo
beneficiário	a entidade que se beneficia da ação expressa pelo verbo
alvo	a entidade em direção a qual vai a ação expressa pelo verbo
fonte	a entidade a partir da qual algo é movido como resultado da ação expressa pelo verbo
local	local onde se situa a ação ou estado expresso pelo verbo

Um exame da gramática artificial usada em Elman (1990) permite depreender que agente, experienciador, paciente e tema são os papéis temáticos utilizados. Esses papéis temáticos distribuem-se nas duas posições argumentais (arg_1 e arg_2) dos verbos usados.

Portanto, os padrões de co-ocorrência de palavras em sentenças, no experimento de Elman, além de refletirem classes semânticas (como inanimados e animados) e transitividade (como sujeito, objeto obrigatório ou opcional), também refletem distribuições de papéis temáticos. Note que, diferentemente de traços semânticos como animado e inanimado, que podem ser vistos como propriedades das palavras, os papéis temáticos são propriedades da relação de certas palavras com certos verbos, ou seja, uma relação que é condicionada pela posição argumental da palavra em questão. Some a isso as evidências que indicam a existência de uma relação estreita entre o sentido de um verbo e os perfis de

sua subcategorização (Hare, McRae & Elman, 2003), sendo que o verbo requer argumentos que assumem determinados papéis temáticos e isso, de certa forma, se manifesta em perfis distribucionais.

O objetivo de nosso trabalho, então, consiste em explorar a capacidade da RRS em representar conhecimento gramatical também de nível semântico, como papéis temáticos, apenas pelos padrões de co-ocorrência. A resposta poderá ser afirmativa caso os padrões distribucionais revelem informações de papel temático, e este é, em parte, o caso. Repare, por exemplo, que um verbo perceptual como *see* seleciona como arg_1 apenas seres animados e, como arg_2 , tanto seres animados quanto inanimados. Um verbo como *chase*, por sua vez, seleciona somente seres animados para ambas as posições argumentais. Isso reflete, ao menos parcialmente, o sentido do verbo e fornece pistas para uma distinção temática: verbos como *see*, cuja grade temática é [experienciador, tema], têm menos restrições de seleção do que verbos como *chase*, cuja grade temática é [agente, paciente]. A semântica do item lexical que ocupa arg_2 é condicionada pelo sentido do verbo, que pode restringir mais ou menos seus tipos semânticos e temáticos.

Os padrões distribucionais, no entanto, diferencia somente grandes classes de verbos (por exemplo, os transitivos *chase* e *like* dos transitivos perceptuais *see* e *smell*), mas não faz distinções temáticas mais específicas. Obviamente, o conhecimento de papel temático necessita de informação semântica mais rica. Feita essa ressalva, nosso objetivo consiste em investigar se emerge algum atrator de tipo semântico apenas pela presença de vieses estatísticos decorrentes das diferentes seleções que os verbos fazem.

Com essa problemática em vista, decidimos replicar o experimento de Elman (1990) a fim de realizar análises adicionais, testar hipóteses e discutir tópicos não tratados no trabalho original, no qual os objetivos eram outros. Repare, pelos excertos arrolados abaixo, que Elman menciona o caráter preliminar de seus experimentos e antecipa a possibilidade de análises complementares.

The results described here are preliminary in nature. They are highly suggestive, and often raise more question than they answer. These networks are properly thought of as dynamical systems, and one would like to know more about their property as such. (...) One would like to know what the trajectories between states (i.e., the vector field) look like. What sort of attractors develop in these systems? It is a problem, of course, that the networks studied here are high-dimensional systems and consequently difficult to study

using traditional techniques. One promising approach which is currently being studied is to carry out a principal components analysis of the hidden unit activation pattern time series, and then to construct phase state portraits of the most significant principal components. (Elman, 1990:208)

No excerto acima, podemos perceber que os resultados apresentados no artigo de 1990 são parciais, de modo que existem questões que permanecem sem respostas. Além disso, a RRS pode ser vista como um sistema dinâmico, o que requer um aparato de análise adicional ao usado nesse trabalho. Em específico, ele menciona a análise de componentes principais como uma maneira promissora de examinar as ativações das unidades ocultas (o espaço de estados do sistema) durante o processamento. De fato, seu trabalho de 1991 mostra que essa análise permite investigar características importantes do processamento de uma RRS.

No excerto abaixo fica claro, mais uma vez, que os resultados apresentados são parciais, bem como menciona a possibilidade de testes adicionais para testar a capacidade representacional da RRS.

The results are preliminary in a number of ways. First, one can imagine a number of additional tests that could be performed to test the representational capacity of the simple recurrent network. (Elman, 1991:19-20)

Neste outro excerto, abaixo, Elman menciona que a sensibilidade ao contexto, inerente à RRS, não a impede de captar generalizações que estão em um alto nível de abstração.

The important result of the current work is to suggest that sensitivity to context which is characteristic of many connectionist models, and which is built-in to the architecture of the networks used here, does not preclude the ability to capture generalizations which are at a high level of abstraction. (Elman, 1991:20)

Através desses excertos, portanto, fundamentamos a necessidade da replicação do experimento com sentenças simples como uma forma de (a) realizar análises adicionais; (b) investigar a capacidade representacional da RRS; e (c) investigar se a RRS pode representar conhecimento gramatical mesmo com sua alta sensibilidade ao contexto.

3. Replicando Elman (1990)

Seguindo o procedimento relatado em Elman (1990), geramos⁷ um conjunto aleatório de 10.000 sentenças simples, utilizando os padrões da tabela 2.2 e as 29 palavras

⁷ O código em C do programa gerador de sentenças pode ser conferido no anexo A.

pertencentes às 12 categorias da tabela 2.1. Veja algumas sentenças produzidas: *woman smash cookie, man sleep, lion chase dog, car exist*.

Em primeiro lugar, é preciso atribuir às palavras um padrão de números que será processado pela rede. Cada palavra é representada por um vetor de 29 *bits* (tabela 2.4), no qual um único *bit* é 1, o restante é 0 e não há dois vetores iguais⁸. Isso significa que uma palavra ativa um único neurônio de entrada⁹. Como foram usadas 29 palavras, temos um vetor de 29 *bits* e uma rede com 29 unidades de entrada e 29 de saída.

Tabela 2.4

Representação das palavras por vetores

Palavras	Vetores
LION	1 0
CAR	0 1 0
SMASH	0 0 1 0
LIKE	0 0 0 1 0
ROCK	0 0 0 0 1 0
SMELL	0 0 0 0 0 1 0
SLEEP	0 0 0 0 0 0 1 0
EAT	0 0 0 0 0 0 0 1 0
COOKIE	0 0 0 0 0 0 0 0 1 0
SEE	0 0 0 0 0 0 0 0 0 1 0
BOY	0 0 0 0 0 0 0 0 0 0 1 0
SANDWICH	0 0 0 0 0 0 0 0 0 0 0 1 0
DOG	0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0
GIRL	0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0
MONSTER	0 0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0
MOVE	0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0
WOMAN	0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0
CHASE	0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0 0
PLATE	0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0
MOUSE	0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0
THINK	0 1 0 0 0 0 0 0 0 0 0 0 0
GLASS	0 1 0 0 0 0 0 0 0 0 0 0
BREAD	0 1 0 0 0 0 0 0 0 0 0
BREAK	0 1 0 0 0 0 0 0 0 0
BOOK	0 1 0 0 0 0 0 0 0
MAN	0 1 0 0 0 0 0 0
EXIST	0 1 0 0 0 0
DRAGON	0 1 0 0
CAT	0 1

As palavras são apresentadas continuamente e a tarefa da rede é emitir uma saída que corresponde à previsão da próxima palavra. Por exemplo, no caso de *dog chase cat*, quando da entrada do vetor correspondente a *dog*, espera-se que a rede produza em sua camada de

⁸ Em termos geométricos, esses vetores são ortogonais entre si.

⁹ Denomina-se de representação localista esse tipo de codificação “um conceito = uma unidade”.

saída o vetor correspondente a *chase* e assim sucessivamente. Depois que a rede emite a saída, o algoritmo de treinamento ajusta os pesos de conexão e, em seguida, a próxima palavra é apresentada. Veja na tabela 2.5 um fragmento do conjunto de treinamento.

Tabela 2.5
Exemplo do mapeamento entrada-saída

ENTRADA	SAÍDA
000000000000100000000000000000 (dog)	000000000000000000000000000000 (chase)
000000000000000000000000000000 (chase)	00000000000000000000000000000001 (cat)
00000000000000000000000000000001 (cat)	00000000000000000000000000000000 (woman)
00000000000000000000000000000000 (woman)	00000001000000000000000000000000 (eat)
00000001000000000000000000000000 (eat)	00000000000010000000000000000000 (sandwich)
00000000000010000000000000000000 (sandwich)	10000000000000000000000000000000 (lion)

É importante notar que o mapeamento *não é determinístico*, i.e., não existe uma única saída-alvo possível para qualquer uma das entradas. Diversas palavras podem ocorrer depois de *boy*, por exemplo: *boy chase*, *boy see*, *boy eat*. Ou seja, a saída da rede nunca será o vetor exato de uma palavra, mas refletirá um conjunto de possibilidades.

Convém salientar a simplificação feita: o uso de vetores ortogonais (uma representação localista) significa que, sob o ponto de vista da entrada da rede, não há nenhuma informação relativa a classes de palavras (nome ou verbo) e função sintática (sujeito e objeto). A entrada (o estímulo) é empobrecida dessa forma para permitir um enfoque nas relações de co-ocorrência de palavras, que refletem restrições sintático-semânticas. Vejamos, agora, detalhes mais técnicos do modelo.

O conjunto de treinamento é composto por uma seqüência ininterrupta de palavras, totalizando uma cadeia de 27.360 vetores de entrada. Note que não há nenhuma informação relativa à fronteira entre as sentenças. A RRS utilizada tem 29 unidades na camada de entrada e na de saída, 150 na camada oculta e na de contexto¹⁰. Os pesos iniciais são dados por uma função que gera números uniformemente distribuídos entre -0,5 e 0,5. Todas as unidades de contexto possuem a ativação inicial de 0,5. Diferentes simulações foram

¹⁰ Essa rede é semelhante à usada em Elman (1990), mas na original Elman usou 31 unidades de entrada e de saída para realização de testes.

realizadas, ora mantendo, ora modificando parâmetros do algoritmo de aprendizagem. Os resultados descritos a seguir foram obtidos após um treinamento de 6 épocas, i.e., 6 passagens completas do conjunto de 10.000 sentenças, utilizando uma taxa de aprendizagem¹¹ de 0,1 e um momento de 0,9.

3.1 Representação lexical e análise de agrupamentos

Quando o treinamento da rede termina, podemos submetê-la a testes e análises que nos permitem investigar o seu processamento. O teste feito por Elman (e que repetimos nesta seção) consistiu em processar o mesmo conjunto de sentenças do treinamento sem que houvesse ajuste dos pesos. Durante o processamento dessas sentenças, as ativações da camada oculta são gravadas. É importante lembrar que a camada oculta é ativada tanto pelas unidades de entrada quanto pelas de contexto, ou seja, não há processamento de palavras isoladas. Feita essa ressalva, apresentamos uma análise que permite observarmos o efeito que cada palavra (vetor) causa na camada oculta.

O objetivo desse teste é analisar as ativações da camada oculta quando as palavras são processadas, mas cada palavra ocorre muitas vezes e em diversos contextos no conjunto de 10.000 sentenças. Uma maneira de minimizar os efeitos da ocorrência e do contexto é calcular a média das ativações da camada oculta para cada item, o que resulta em 29 vetores de 150 elementos cada um (são 150 unidades ocultas). Pretende-se, com isso, obter uma aproximação da representação *prototípica* de cada palavra. O passo seguinte é submeter a matriz obtida a uma *análise de agrupamentos*. Repare que se trata de uma matriz 29x150 de dados multivariados, na qual 29 objetos são descritos por 150 variáveis. Essa alta dimensionalidade impossibilita a comparação direta entre os vetores, sendo indispensável o uso de uma análise multivariada.

Análise de agrupamentos (cluster analysis) é o nome de um conjunto de técnicas estatísticas que permite agrupar objetos segundo suas características, de modo que haja uma homogeneidade interna a cada agrupamento e uma heterogeneidade externa entre os agrupamentos (Hair et al., 2005:384). Com isso, se os dados são informativos o suficiente para uma classificação bem sucedida, obtém-se uma representação gráfica dos agrupamentos na qual objetos semelhantes estão próximos entre si, havendo uma distância

¹¹ Esses parâmetros foram apresentados na seção 4.3 do capítulo 1.

maior entre agrupamentos menos semelhantes. É importante salientar que se trata de uma estatística descritiva, de cunho explanatório e não-inferencial, que pode produzir diferentes soluções de acordo com o tipo de procedimento adotado (Hair et al., 2005:385).

O primeiro critério a ser estabelecido é uma *medida de similaridade* entre os objetos, de modo que as menores distâncias ou diferenças representam maiores similaridades. O segundo critério é a *formação dos agrupamentos* – independentemente da medida adotada, um procedimento deve agrupar as observações similares. Por último, um critério deve determinar a *quantidade de grupos* formados, o que é feito por uma avaliação da média de similaridade dos agrupamentos, de modo que o aumento dessa média implica que os grupos são menos parecidos (Hair et al. 2005:385). Existem diversos métodos propostos, que variam de acordo com esses três critérios. Para uma apresentação detalhada do conjunto de técnicas de análise de agrupamentos, remetemos ao capítulo 9 de Hair et al. (2005).

Aqui, adotamos como medida de similaridade a *distância euclidiana*, i.e., a distância geométrica entre dois pontos que é dada por uma linha reta. O agrupamento é feito seguindo um tipo de *procedimento hierárquico*, em que *cada objeto* começa como sendo seu próprio agrupamento. Os próximos passos combinam os dois agrupamentos (ou objetos) mais próximos (i.e., semelhantes) em um novo, reduzindo o número de grupos em uma unidade a cada passo (Hair et al. 2005:398). Dada essa característica de formar agrupamentos a partir da combinação de outros já existentes, esse tipo de método também é chamado de *aglomerativo*. Adotamos o *método de Ward*, que procura minimizar a soma dos quadrados entre dois *possíveis* agrupamentos que podem se formar a cada passo.

O resultado de uma análise de agrupamentos, então, pode ser visualizado através de um *dendograma*, i.e., um diagrama em formato de árvore que mostra espacialmente os níveis de similaridade e dissimilaridade entre os objetos agrupados. Através do dendograma, podemos observar as combinações feitas pelo conjunto de passos do procedimento.

Retomando nossa matriz de dados, a finalidade de fazer essa análise é agrupar as palavras pelas similaridades das ativações que elas causam na camada oculta. A figura 2.1 traz o resultado dessa análise. A árvore hierárquica obtida reflete a similaridade das representações internas das palavras na camada escondida, ou seja, obtemos uma

“fotografia” da organização lexical no espaço representacional da rede. Podemos observar três grupos principais, que subdividem os itens lexicais em verbos, nomes animados e nomes inanimados. Cada grupo é subdividido em subgrupos que refletem similaridades mais específicas entre os itens lexicais. Com base nesse dendograma, podemos concluir que a representação dos itens lexicais na camada oculta é estruturada segundo propriedades lingüísticas, ou seja, palavras parecidas (*boy* e *girl*, por exemplo) causam um padrão de ativações semelhante na camada oculta. A análise de agrupamentos que obtivemos é bem semelhante à apresentada em Elman (1990:200), mas existem algumas diferenças. Em Elman, o resultado mostra dois grandes grupos – um contendo os nomes e outro, os verbos.

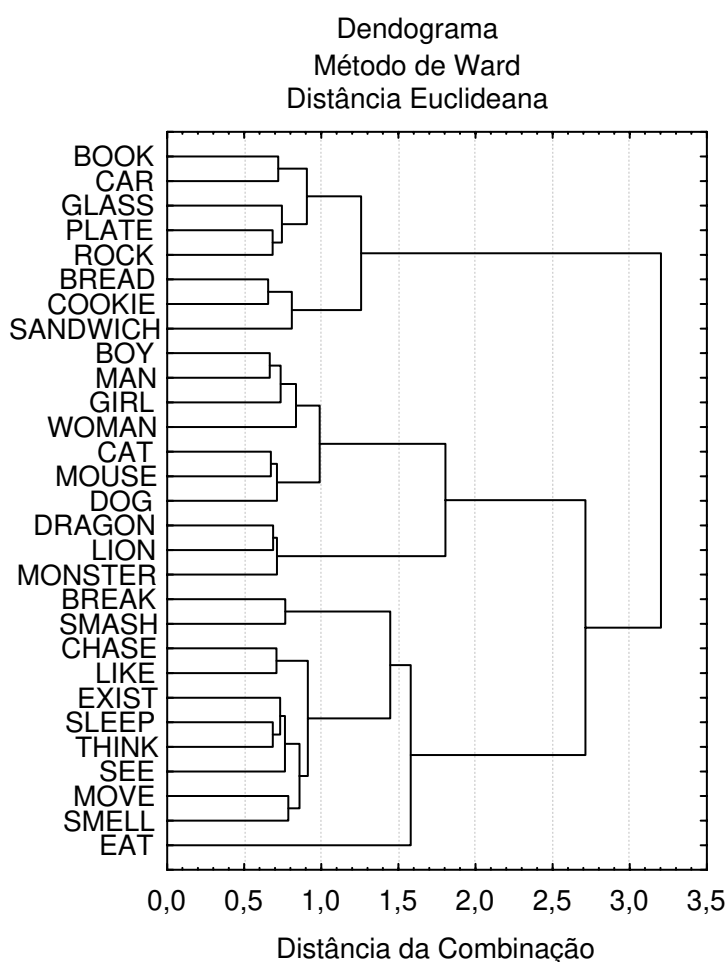


Figura 2.1 A árvore hierárquica mostra a similaridade das representações internas. Da esquerda para a direita, podemos acompanhar o agrupamento realizado e as respectivas distâncias (eixo x) entre as palavras e os grupos. Por exemplo, é de aproximadamente 3,2 a distância entre o grupo dos inanimados (book, car, bread etc.) em relação ao outro grande agrupamento.

A partir do dendograma, podemos observar que o grupo dos animados agrega 3 subgrupos: humanos (*boy, girl, man, woman*), animais pequenos (*cat, dog, mouse*) e animais grandes (*dragon, lion, monster*). O grupo dos inanimados agrega 3 subgrupos: objetos (*book, car*), alimentos (*bread, cookie, sandwich*) e objetos quebráveis (*glass, plate*). O agrupamento dos verbos também reflete categorizações relevantes, observadas em subgrupos de verbos transitivos como *break* e *smash, chase* e *like*, e os intransitivos *exist, sleep* e *think*.

3.2 Representação gramatical e análise de componentes principais

Elman (1990) não prossegue com análises da camada oculta, servindo-se apenas da análise de agrupamentos para discutir tópicos relacionados à representação lexical. Aqui, nosso objetivo é investigar as representações sintático-semânticas que operam no nível das sentenças. Para tanto, é necessário inspecionar a camada oculta da rede, de modo a identificar como o padrão de ativações está se comportando *durante o processamento* das sentenças. Ou seja, é preciso estudar as mudanças que ocorrem no espaço de estados da RRS enquanto sentenças são processadas.

Repare que, novamente, uma inspeção direta da camada oculta é impraticável, devido à grande dimensionalidade do espaço – o processamento e a representação estão distribuídos ao longo de 150 unidades, de modo que nenhuma é significativa por si só. Por esse motivo, podemos supor que há informações codificadas em conjuntos de unidades, o que se reflete na variância dos padrões de ativação da camada oculta durante o processamento das sentenças (Elman, 1991:13). Em outras palavras, não há como analisar as ativações causadas por uma sentença específica e identificar as unidades que respondem a uma certa informação – é preciso uma análise que englobe o processamento de diversas sentenças e identifique um padrão significativo de ativações.

De fato, como apontado por Hair et al. (2005:92), o número elevado de variáveis aumenta as possibilidades de que nem todas as variáveis sejam não-correlacionadas, i.e., pode haver um grupo de variáveis que sejam inter-relacionadas, de modo a representarem um conceito mais geral. Isso dificulta a interpretação dos dados e a identificação de uma estrutura subjacente. É preciso, pois, uma maneira de melhor descrever e visualizar os dados em questão. Conforme proposto em Elman (1991:13-4), podemos submeter os dados

obtidos a partir do processamento do conjunto de sentenças a uma *análise de componentes principais*. Trata-se de um tipo de análise estatística multivariada que se aplica em casos em que diversas variáveis são usadas para caracterizar um dado conjunto de elementos. No nosso caso, o alvo da análise é a matriz 27.360x150, correspondente à ativação das unidades ocultas (descrita por 150 variáveis) para cada palavra processada em seu contexto (27.360 casos).

A análise de componentes principais (ACP) é usada para reduzir o número de variáveis, de forma a reter a maior informação possível das variáveis originais. Essa redução se dá porque a ACP permite identificar as dimensões de maior variação dos vetores, possibilitando a visualização destes em um sistema de coordenadas alinhadas com essa variação. Em outras palavras, a ACP efetua uma rotação de eixos no espaço, de modo que a variabilidade máxima é projetada no eixo do primeiro componente principal. A variabilidade restante é projetada no eixo do segundo CP e assim sucessivamente para quantas forem as dimensões de variação. Logo, esse novo sistema de coordenadas permite uma melhor descrição dos dados. Veja na figura 2.2 um gráfico que mostra a rotação realizada por uma ACP.

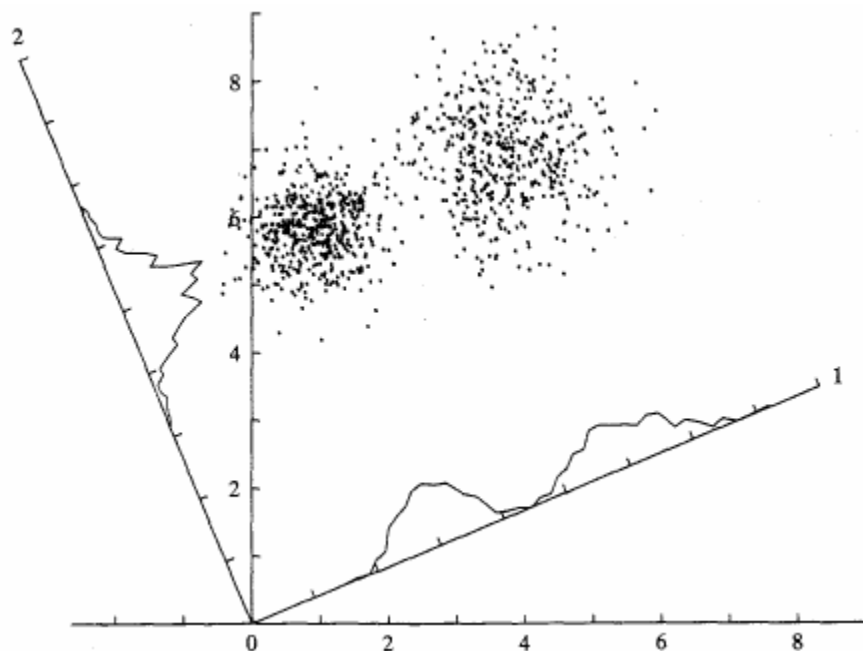


Figura 2.2 Uma nuvem de pontos de dados pode ser projetada nos eixos 1 e 2. A projeção no eixo 1 tem variância máxima. A variância restante é projetada no eixo 2. Figura retirada de Haykin (2001:441).

O resultado de uma ACP, portanto, é a obtenção de um novo conjunto de variáveis (chamadas componentes principais, CPs) não-correlacionadas, de forma que as primeiras retêm grande parte da variabilidade presente no conjunto de dados. Logo, pode-se usar um número relativamente pequeno de variáveis para descrever os dados, com a ressalva de que há perdas de informação. A ACP resulta em eixos ortogonais entre si, e nem sempre uma reta é suficiente para conter a variabilidade em toda sua significância (este é o caso, por exemplo, se tivermos uma distribuição curva). No entanto, as componentes principais fornecem a melhor *aproximação linear* para um conjunto de dados.

A seguir, descrevemos os resultados da ACP da matriz 27.360x150. Veja no gráfico abaixo que os componentes principais 1, 2 e 3 explicam 50% da variabilidade dos dados.

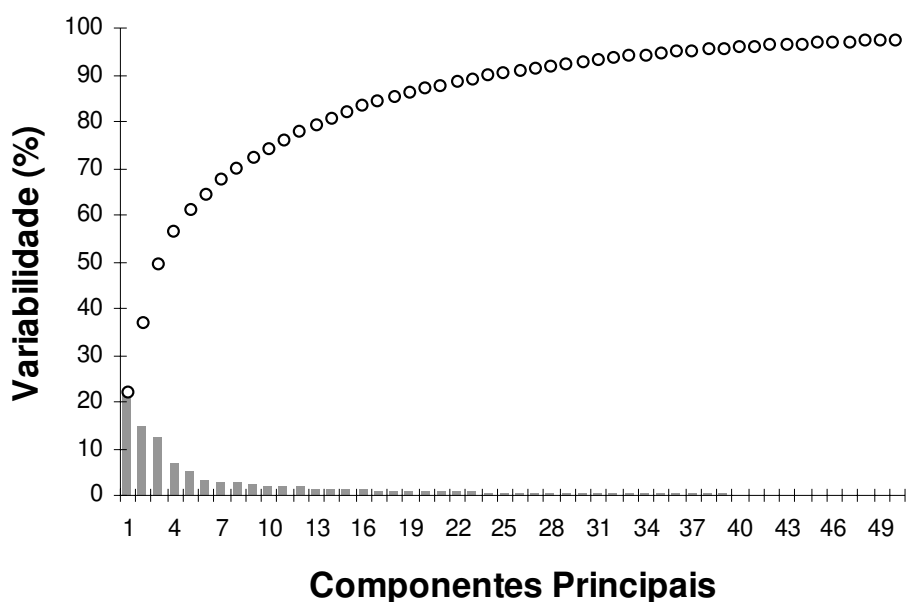


Figura 2.3 As barras verticais indicam a porcentagem da variabilidade contida em cada componente. Os pontos indicam a porcentagem cumulativa ao longo do conjunto de componentes.

Repare que a maior parte da informação está nos primeiros componentes principais, de forma que podemos desconsiderar os CPs depois do 15º, o que nos deixa com 15 variáveis que explicam aproximadamente 80% dos dados. Ou seja: os dados são descritos por 15 variáveis não-correlacionadas, que substituem as 150 originais. Obtido o novo sistema de coordenadas, através da ACP, o passo seguinte consiste em projetar, no eixo das

dimensões mais informativas, o conjunto de pontos que descrevem o processamento das sentenças.

Retomando as questões de interesse, discutidas na seção 2 deste capítulo, selecionamos algumas sentenças-alvo que contêm os tópicos que queremos discutir. Por exemplo, ao processar as sentenças *boy chase boy* e *boy see boy*, queremos verificar se a rede usa informações sintáticas e semânticas. Repare que os itens lexicais que ocupam a posição de arg_1 e arg_2 são os mesmos (sob o ponto de vista da entrada da rede), enquanto os verbos são diferentes e selecionam argumentos com papéis temáticos distintos (paciente e tema, respectivamente). Assim, trabalhamos com duas hipóteses:

(1) se há conhecimento sintático, então algum CP aproxima sujeitos, verbos e objetos entre si. Ou seja, estamos supondo que palavras de mesma função sintática estão na mesma região do espaço de estados, em pelo menos alguma dimensão.

(2) se há conhecimento semântico, então algum CP afasta as ocorrências de *boy* em arg_2 . Nesse caso, estamos supondo que palavras com papéis temáticos distintos estão em regiões diferentes do espaço de estados. O afastamento das instâncias de *boy*, que correspondem aos mesmos vetores de entrada, é tomado como indício da presença de informação relativa a papéis temáticos.

As sentenças-alvo foram, então, submetidas à rede como teste. Antes da apresentação de cada sentença, houve uma “limpeza” da camada de contexto, i.e., todas as ativações foram ajustadas para 0,5. Assim, obtemos as ativações da camada oculta apenas durante o processamento de uma sentença específica.

Veja, na figura 2.4, a trajetória do processamento das sentenças *boy chase boy* e *boy see boy*, ao longo do tempo, no CP 1. A informação incorporada pelo CP 1 mostra a trajetória diferenciada do processamento das duas sentenças. Os verbos são reconhecidos como itens lexicais diferentes, a despeito de exercerem a mesma função sintática. O afastamento das ocorrências de *boy*, em arg_2 , indica que *boy* objeto de *chase* é distinto de *boy* objeto de *see*. Considerando essa dimensão, temos indícios de que a grade temática está representada: os verbos são diferentes e seus objetos também são. As ocorrências de *boy*, em arg_1 , ocupam a mesma posição no espaço de estados (e isso ocorre em todos os CPs),

dado que a rede inicia o processamento das sentenças sem contexto prévio e, portanto, responde da mesma forma diante das mesmas entradas.

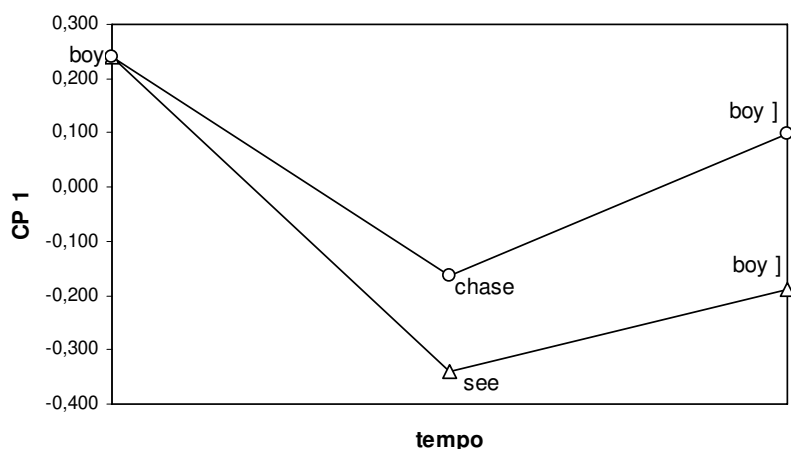


Figura 2.4 Trajetórias no espaço de estados (CP 1) quando a rede processa as sentenças *boy chase boy* e *boy see boy*.

Checando uma outra dimensão, encontramos a situação na qual os objetos estão próximos entre si (figura 2.5). Considerando a trajetória do processamento das sentenças no CP 5, podemos observar, pela proximidade de *boy* em arg_2 , que o fato dos verbos serem diferentes não afeta o processamento do objeto, o que acontece no CP 1. Logo, nessa dimensão, a informação presente é distinção entre os verbos.

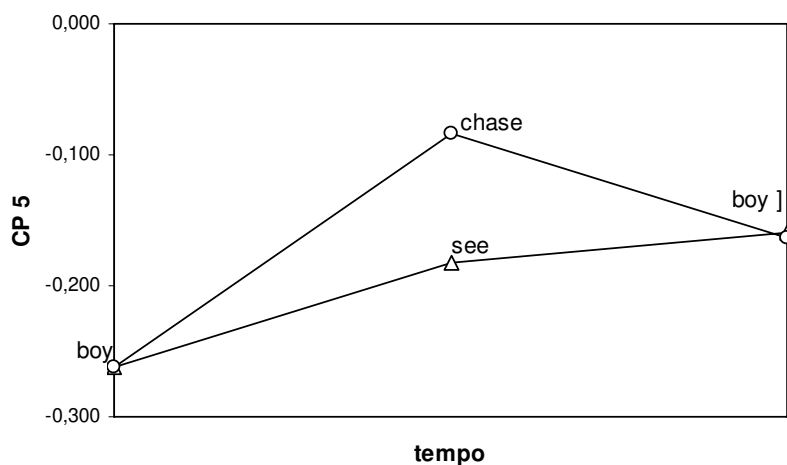


Figura 2.5 Trajetórias no espaço de estados (CP 5) quando a rede processa as sentenças *boy chase boy* e *boy see boy*.

Em outra dimensão, encontramos uma maior proximidade entre as palavras de mesma função sintática (figura 2.6). Ou seja, há um indício de que um CP que carrega pouca informação, como o 7 (cf figura 2.3), está codificando informação de tipo sintática. Ao contrário dos CPs 1 e 5, que mostram uma grande diferença entre os verbos, no CP 7 há uma maior proximidade entre eles.

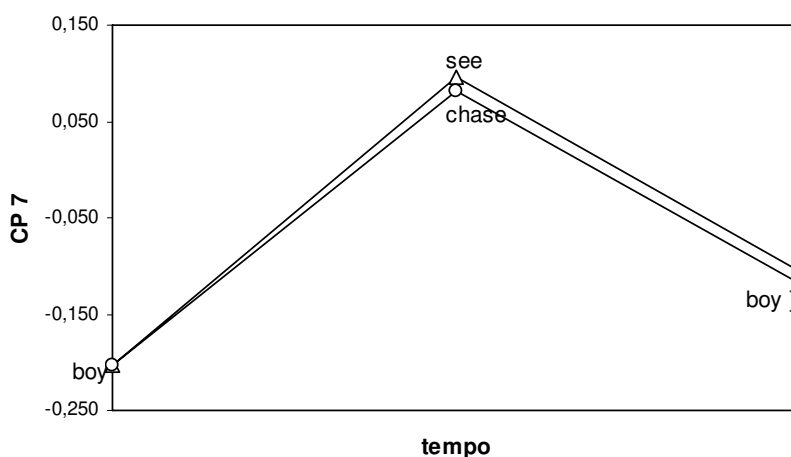


Figura 2.6 Trajetórias no espaço de estados (CP 7) quando a rede processa as sentenças *boy chase boy* e *boy see boy*.

Contudo, em sentenças como *boy break plate* e *plate break*, temos uma outra situação: observe que a posição de arg_1 do verbo *break* é ocupada por itens lexicais distintos, mas sintaticamente ambos são sujeitos. A aproximação de *plate* e *boy* em arg_1 requer conhecimento estrutural, e não pode ser visto no CP 7 (figura 2.7).

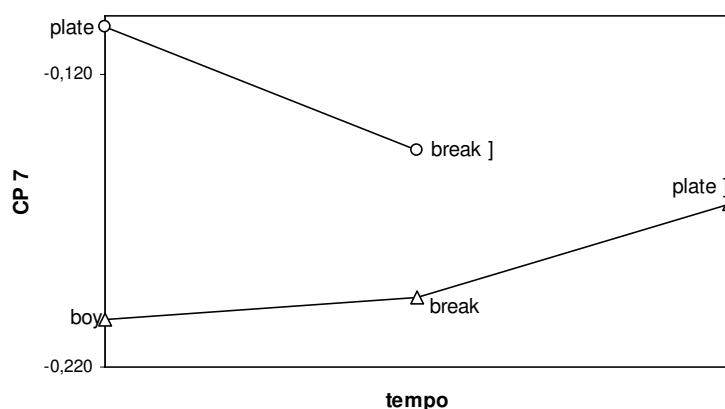


Figura 2.7 Trajetórias no espaço de estados (CP 7) quando a rede processa as sentenças *boy brake plate* e *plate break*.

Portanto, o indício de que o CP 7 codifica função sintática (figura 2.6) é refutado pelo dado apresentado na figura 2.7. Em outras dimensões, quando os verbos estão próximos entre si, no CP 4, os sujeitos não estão, ocorrendo o oposto no CP 9 (figura 2.8). No caso dessas sentenças, o CP 4 é sensível a informações lexicais: o verbo é o mesmo e permanece na mesma região do espaço, mas os nomes são diferentes e isso pode ser visto nessa dimensão. Já no CP 9, há uma maior proximidade entre os nomes, mas as ocorrências de *break* estão separadas. De fato, há uma diferença fundamental nesses dois usos do verbo *break*: quando arg_1 é ocupado por um inanimado o verbo não tem arg_2 , que é opcional quando arg_1 é preenchido por um animado.

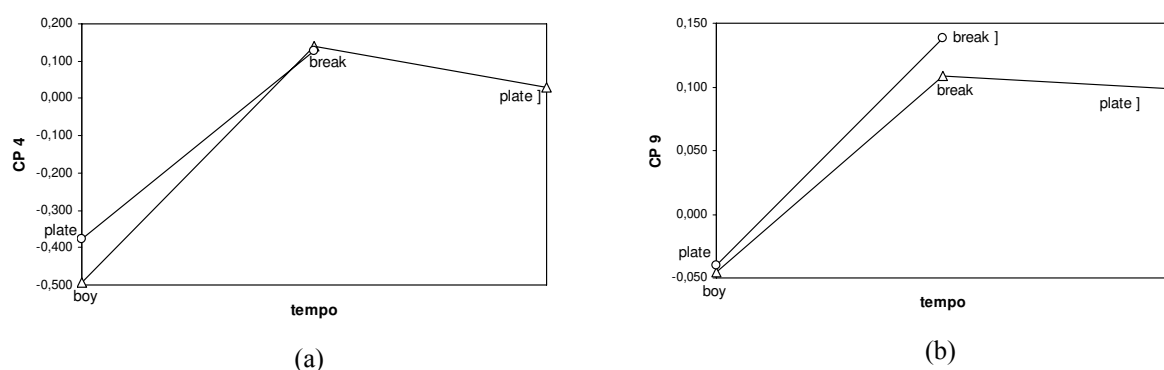


Figura 2.8 Trajetórias no espaço de estados (CPs 4 e 9) quando a rede processa as sentenças *boy brake plate* e *plate break*.

Não encontramos, portanto, nenhum CP que, sozinho, captasse a informação estrutural. Logo, o conhecimento sintático não está satisfatoriamente codificado no espaço de estados, que parece operar segundo informações de nível lexical, apenas. Em outras palavras, não podemos concluir que há conhecimento de constituintes sintáticos como sujeitos, verbos e objetos, o que deveria aparecer em um único CP. Repare que esta é, de fato, uma informação sub-representada no conjunto de treinamento, uma vez que as sentenças são apresentadas sob a forma de uma cadeia ininterrupta de palavras, sem que nenhuma informação estrutural esteja explícita. Não haveria, pois, informações suficientes para a rede induzir que os itens lexicais pertencem ao constituinte maior *sentença* e, conseqüentemente, responder de maneira análoga a itens que exercem a mesma função sintática. Dada a existência de dimensões do espaço de estados que captam as distinções semânticas esperadas, na próxima seção modificamos o experimento em busca da codificação sintática.

4. Em busca da estrutura sintática

Estamos diante de uma questão crucial. As informações presentes no conjunto de treinamento não bastam para levar a rede a aprender por completo as informações lingüísticas do nível da sentença. Em vista disso, devemos buscar uma maneira de treinar a rede de modo a torná-la sensível ao fato das palavras formarem sentenças.

Uma alternativa reconhecida e adotada na literatura (e.g. Elman, 1991) sobre codificações sintáticas em RRS é a utilização da marca de fim de sentença (um ponto final), o que se dá pelo acréscimo de mais um vetor de entrada. Adicionamos à tabela 2.4 (página 40) o seguinte vetor correspondente ao ponto final

[illegible]

no qual apenas o trigésimo elemento é l .

Essa manobra, ao explicitar a fronteira entre sentenças, pode servir como base para a rede construir uma noção de constituinte mais refinada e complexa – como a que precisamos neste experimento. Realizamos, então, novos experimentos usando o mesmo conjunto de sentenças da simulação anterior, apenas introduzindo o ponto final nas sentenças. O conjunto passou de 27.360 para 37.360 vetores (são 10.000 pontos finais) e a rede passou a ter 30 unidades de entrada e de saída (29 palavras mais o ponto final). Os resultados apresentados a seguir foram conseguidos parâmetros similares aos do treinamento anterior. Os pesos iniciais são dados por uma função que gera números uniformemente distribuídos entre -0,5 e 0,5. Todas as unidades de contexto possuem a ativação inicial de 0,5. E usamos uma taxa de aprendizagem de 0,1 e o momento de 0,9. A primeira diferença observada foi no tempo de treinamento – 4 épocas bastaram para o aprendizado da rede. Também realizamos a mesma metodologia de teste: o conjunto de sentenças foi processado sem que houvesse a correção de pesos e as ativações da camada oculta foram gravadas para a realização das análises estatísticas.

4.1 Análise de agrupamentos

Calculamos a média das ativações da camada oculta para cada palavra, o que resulta em 29 vetores de 150 elementos cada um (29 palavras e 150 unidades ocultas). Repare que não incluímos o ponto final na análise, uma vez que o intuito é investigar as representações lexicais. A análise de agrupamentos mostra uma organização do espaço representacional

semelhante à do experimento anterior, embora um pouco mais refinada. Veja na figura 2.9 que há um maior distanciamento entre os três agrupamentos principais (verbos, nomes animados e nomes inanimados). Além disso, comparando à figura 2.1 (página 44), pode-se observar que: (1) os objetos quebráveis *plate* e *glass* agora estão agrupados; (2) o agrupamento dos verbos está mais refinado, ficando mais clara a distinção entre transitivos e intransitivos (a exceção é o verbo *move* que aparece junto aos intransitivos); e (3) os perceptuais *see* e *smell* aparecem juntos, diferentemente do ocorrido no experimento anterior.

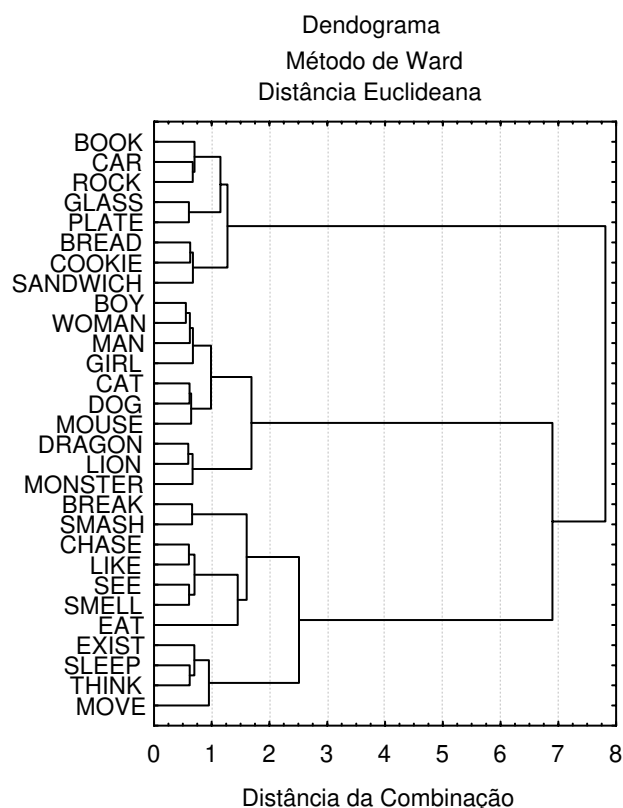


Figura 2.9 A árvore hierárquica mostra a similaridade das representações internas. Da esquerda para a direita, podemos acompanhar o agrupamento realizado e as respectivas distâncias (eixo x) entre as palavras e os grupos. Por exemplo, é de aproximadamente 7,8 a distância entre o grupo dos inanimados (book, car, bread etc.) em relação ao outro grande agrupamento.

4.2 Análise de componentes principais

Submetemos, então, a matriz 37.360x150 a uma ACP. Obtivemos uma qualidade melhor na representação obtida (figura 2.10): dez CPs bastam para explicar aproximadamente 90% dos dados, enquanto no experimento anterior essa mesma

quantidade de componentes explicava menos de 80% dos dados. Repare, ainda, que os CPs 1, 2 e 3 explicam 70% da variabilidade dos dados.

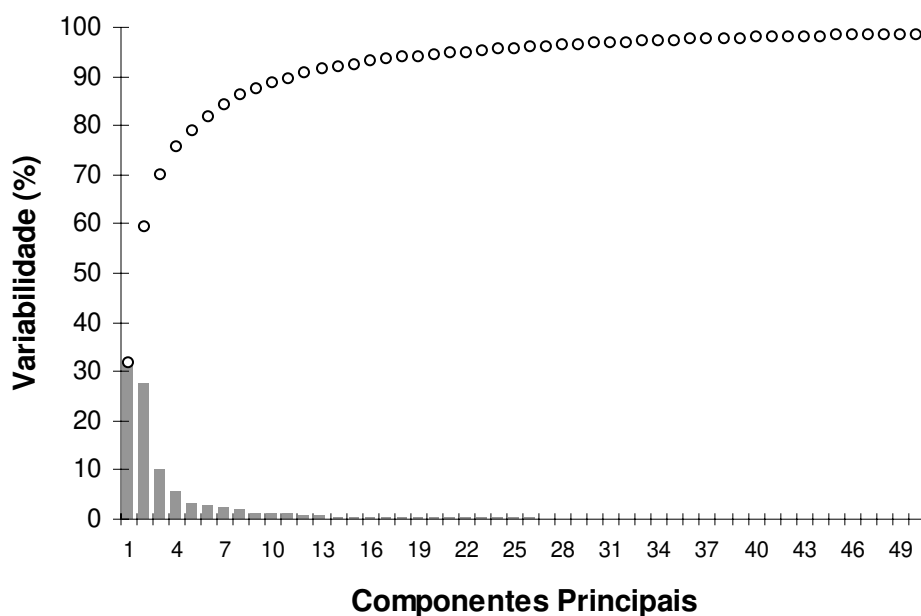


Figura 2.10 As barras verticais indicam a porcentagem da variabilidade contida em cada componente. Os pontos indicam a porcentagem cumulativa ao longo do conjunto de componentes.

Vejamos, pois, como essa “melhor qualidade de representação” está refletida no processamento das sentenças-alvo. Veja na figura 2.11 a trajetória do processamento das sentenças *boy see boy* e *boy chase boy* no CP 1.

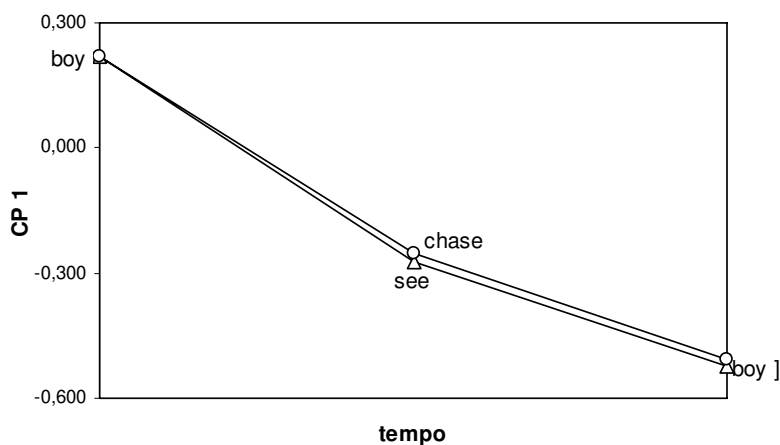


Figura 2.11 Trajetórias no espaço de estados (CP 1) quando a rede processa as sentenças *boy chase boy* e *boy see boy*.

A informação incorporada pelo CP 1 mostra as trajetórias semelhantes e próximas do processamento das sentenças *boy see boy* e *boy chase boy*. Os verbos estão próximos entre si, mesmo sendo itens lexicais distintos. Além disso, a proximidade entre as ocorrências de *boy* em arg_2 mostra que, apesar destes seguirem diferentes tipos de verbos, ambos são tratados da mesma maneira. Interpretamos essas proximidades como um indício de que essa dimensão capta *informação sintática*, ou seja, o CP 1 parece codificar *estrutura argumental* (e isto será confirmado mais adiante). Obtivemos, pois, um primeiro progresso em relação ao experimento anterior.

Devemos verificar, ainda, se há em outros CPs a presença de informações semânticas, como grade temática. No CP 2 (figura 2.12), observamos principalmente a distinção dos verbos *see* e *chase*. A despeito disso, os objetos permanecem próximos, ou seja, o CP 2 está fazendo a categorização dos verbos.

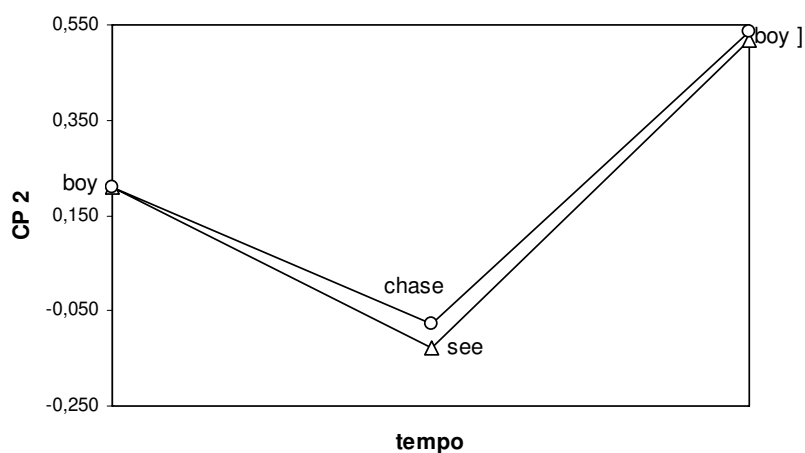


Figura 2.12 Trajetórias no espaço de estados (CP 2) quando a rede processa as sentenças *boy chase boy* e *boy see boy*.

A informação de grade temática, que pode ser vista como um afastamento entre os verbos e entre as ocorrências da palavra *boy* na posição de objeto, pode ser vista no CP 3 (figura 2.13). Note que, nessa dimensão, o processamento das sentenças percorre trajetórias marcadamente diferentes no espaço de estados.

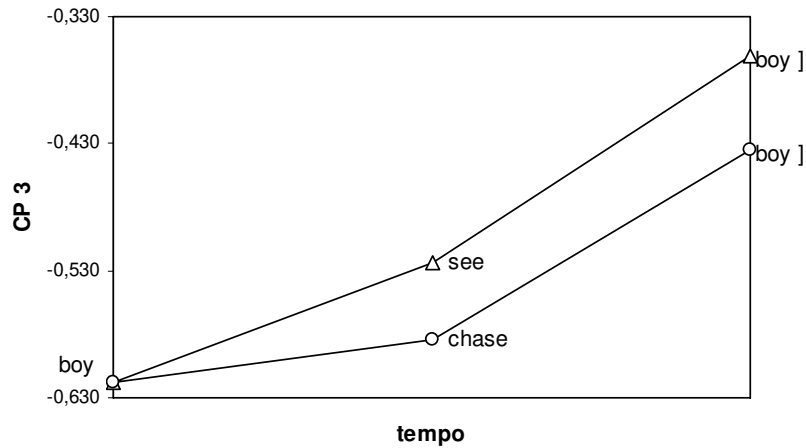


Figura 2.13 Trajetórias no espaço de estados (CP 3) quando a rede processa as sentenças *boy chase boy* e *boy see boy*.

Encontramos, ainda, uma dimensão cujo papel fundamental é diferenciar os objetos entre si (figura 2.14).

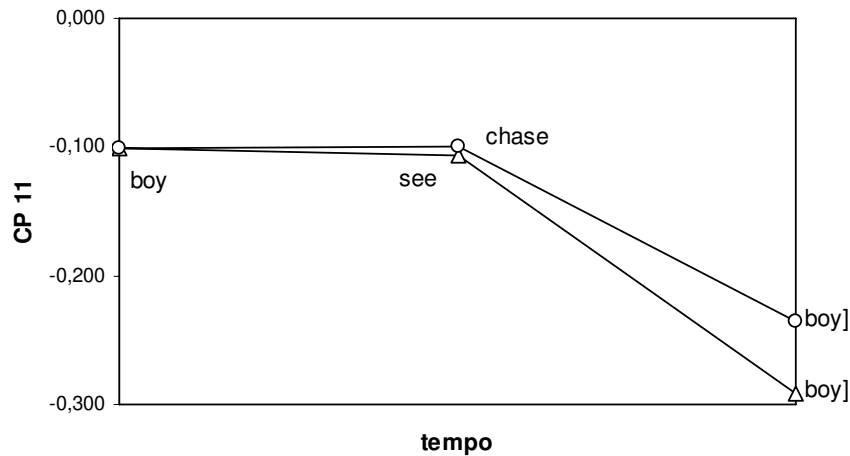


Figura 2.14 Trajetórias no espaço de estados (CP 11) quando a rede processa as sentenças *boy chase boy* e *boy see boy*.

Obtivemos, pois, resultados que confirmam as hipóteses de que, no espaço multidimensional da camada oculta, encontramos dimensões operando distinções sintáticas e outras operando distinções semânticas. Diferentemente do experimento anterior, a camada oculta mostra-se capaz de codificar conhecimento linguístico como estrutura argumental (no CP 1) e grade temática (no CP 3).

Para verificar a consistência do conhecimento de estrutura argumental, vejamos o processamento de algumas sentenças ao longo do CP 1. As sentenças *boy chase woman* e *woman chase boy* têm o mesmo padrão, mas diferem em relação aos nomes que ocupam arg_1 e arg_2 . A despeito disso, as trajetórias são semelhantes, como podemos ver na figura 2.15.

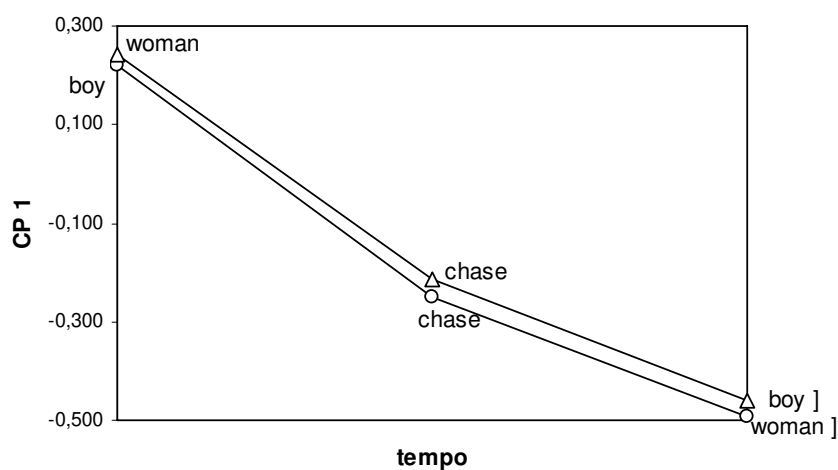


Figura 2.15 Trajetórias no espaço de estados (CP 1) quando a rede processa as sentenças *boy chase woman* e *woman chase boy*.

As sentenças *boy break plate* e *plate break* são de padrões diferentes. O conhecimento de estrutura argumental permite aproximar os sujeitos *boy* e *plate* (figura 2.16). De fato, o CP 1 mostra que a trajetória sujeito-verbo é similar para as duas sentenças.

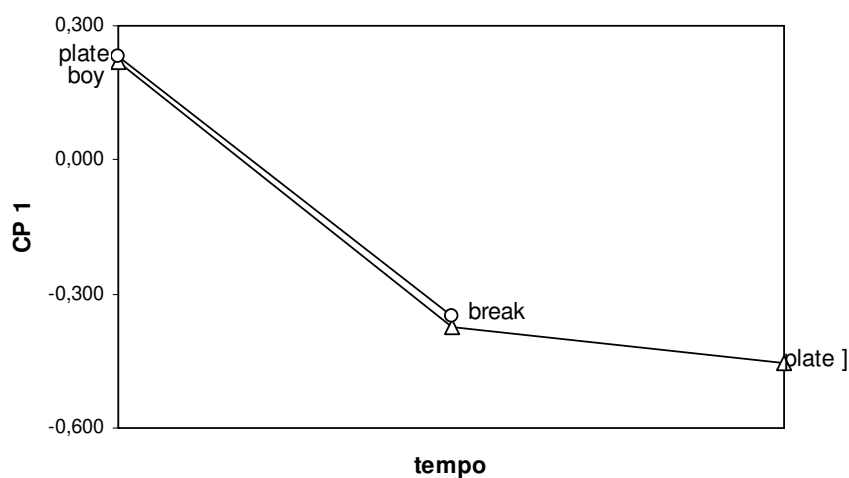


Figura 2.16 Trajetórias no espaço de estados (CP 1) quando a rede processa as sentenças *boy break plate* e *plate break*.

Nessas sentenças, temos duas grades temáticas do verbo *break*: [agente, paciente] e [paciente]. Ou seja, o tipo do verbo muda conforme o item lexical que está na posição de arg_1 . Vemos na figura 2.17 que os CPs 2 e 3 captam essa distinção, corroborando que esses CPs distinguem os verbos e as grades temáticas.

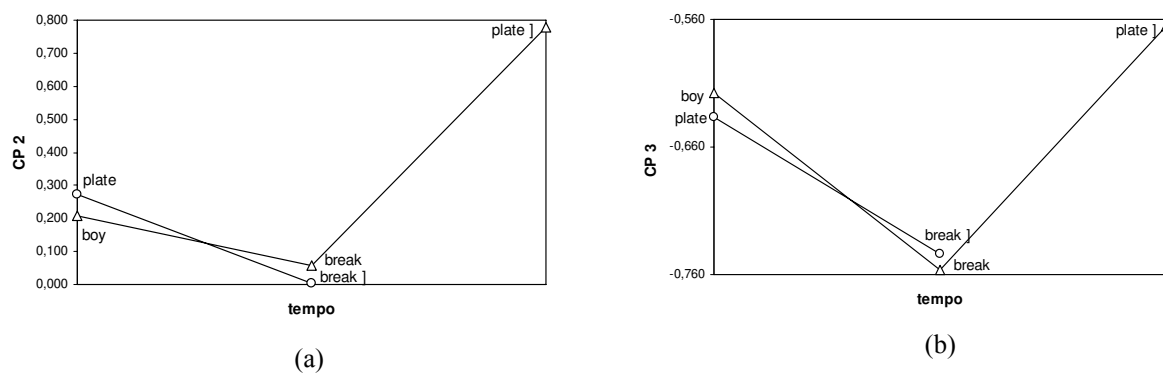


Figura 2.17 Trajetórias no espaço de estados (CPs 2 e 3) quando a rede processa as sentenças *boy brake plate* e *plate break*.

Verificamos, pois, a consistência da presença de informações sintático-semânticas durante o processamento das sentenças. O experimento relatado nesta seção, portanto, mostra que uma RRS é capaz de aprender e de representar conhecimento gramatical que opera durante o processamento de sentenças simples. Conhecimento este que se mostrava incompleto no experimento anterior, no qual havia inconsistências no uso de informações sintáticas. Ademais, podemos assegurar que o segundo experimento foi bem sucedido devido ao uso do ponto final depois de cada sentença, pois se trata da única diferença presente no conjunto de treinamento. Também é importante notar que esse conhecimento é codificado nos primeiros componentes principais, o que nos permite interpretar que os dados são informativos o suficiente e a RRS é capaz de representá-lo. Em vista dos resultados obtidos, realizamos uma nova simulação, desta vez controlando a qualidade de informações semânticas presentes nas sentenças.

5. Um novo experimento

No experimento relatado nesta seção, treinamos uma RRS com sentenças geradas a partir de uma gramática artificial que melhor especifica relações temáticas. Dessa forma, podemos observar com maior segurança as forças que atuam durante o processamento de sentenças em uma RRS. Especificamente, o objetivo é *confirmar* que a RRS é sensível a

informações semânticas presentes sob a forma de vieses distribucionais. Assim, a estrutura argumental permanece a mesma utilizada por Elman (1990), mas alteramos o léxico com o objetivo de minimizar os tipos semânticos dos nomes (tabela 2.6) e aumentar o contraste semântico entre os verbos, enriquecendo dessa forma as relações temáticas (tabela 2.7). Ou seja, os nomes ainda se subdividem em animados e inanimados, mas o contraste entre eles se dá apenas por restrições de seleção relacionadas com as grades temáticas dos verbos (os agentes só podem ser animados, os temas podem ser animados e inanimados). Na gramática de Elman (1990), o contraste entre os nomes se devia principalmente a tipos semânticos como alimentos, objetos quebráveis, animais etc.

Tabela 2.6
Categorias de itens lexicais

Categorias	Membros
Nomes	
Animados	boy, cat, dog, fox, girl, man, woman, wolf
Inanimados	ball, box, case, glass, plate, toy, vase, watch
Verbos	
Transitivos	beat, carry, hurt, move
Transitivos opcionais	break, smash
Intransitivos	fall, run, sink, walk

Tabela 2.7
Verbos e grades temáticas

Verbo	Grade temática
beat, break, hurt, smash	[agente, paciente]
run, walk	[agente]
fall, sink, break, smash	[paciente]
carry, move	[agente, tema]

Os nomes alternam a posição de arg_1 (sujeito) e arg_2 (objeto) e os papéis temáticos de agente, paciente e tema. Na tabela 2.8 temos a gramática artificial usada para gerar as sentenças.

Tabela 2.8
Padrão para o programa gerador de sentenças

arg_1	verbo	arg_2
animados	run, walk, fall, sink	
animados	beat, hurt	animados
animados	break, smash	inanimados
animados	move, carry	animados/inanimados
inanimados	fall, sink, break, smash	

5.1 A simulação

Realizamos procedimentos semelhantes aos relatados anteriormente. Um programa¹² gerou um conjunto de 10.000 sentenças aleatórias segundo a gramática especificada na tabela 2.8. Cada palavra é representada por um vetor de 27 *bits* (26 palavras mais o ponto final, veja a tabela 2.9).

Tabela 2.9
Representação das palavras por vetores

Palavras	Vetores
RUN	1 0
VASE	0 1 0
WATCH	0 0 1 0
WALK	0 0 0 1 0
CAT	0 0 0 0 1 0
FALL	0 0 0 0 0 1 0
PLATE	0 0 0 0 0 0 1 0
DOG	0 0 0 0 0 0 0 1 0
SMASH	0 0 0 0 0 0 0 0 1 0
TOY	0 0 0 0 0 0 0 0 0 1 0
MOVE	0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0
HURT	0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0
MAN	0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0
SINK	0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0
BOY	0 0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0
GLASS	0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0 0
CARRY	0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0 0
WOLF	0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0 0
BREAK	0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0 0
FOX	0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 1 0 0 0 0 0 0 0 0 0 0
GIRL	0 1 0 0 0 0 0 0 0 0 0
WOMAN	0 1 0 0 0 0 0 0 0 0
BOX	0 1 0 0 0 0 0 0 0
CASE	0 1 0 0 0 0 0 0
BEAT	0 1 0 0
BALL	0 1 0
'.'	0 1

Usamos uma RRS com 27 unidades na camada de entrada e na de saída, 150 na camada oculta e na de contexto. Os pesos iniciais são dados por uma função que gera números uniformemente distribuídos entre -0,5 e 0,5. Todas as unidades de contexto possuem a ativação inicial de 0,5. Usamos uma taxa de aprendizagem de 0,1 e um momento de 0,9. Após 5 épocas de treinamento, i.e., 5 passagens completas do conjunto de 10.000 sentenças, usamos o mesmo conjunto de sentenças como teste (sem o ajuste de pesos) e gravamos as ativações da camada oculta para a realização das análises estatísticas.

¹² O código em C do programa gerador de sentenças pode ser conferido no anexo B.

5.2 Análise de agrupamentos

Com o intuito de investigar a organização lexical no espaço representacional da rede, repetimos o procedimento relatado na seção 3.1 deste capítulo. Calculamos a média das ativações da camada oculta para cada item, o que resulta em 26 vetores de 150 elementos cada um (26 palavras e 150 unidades ocultas), lembrando que estamos desconsiderando o ponto final na análise. Submetemos a matriz a uma *análise de agrupamentos*, que pode ser visualizada através de uma árvore hierárquica (figura 2.18).

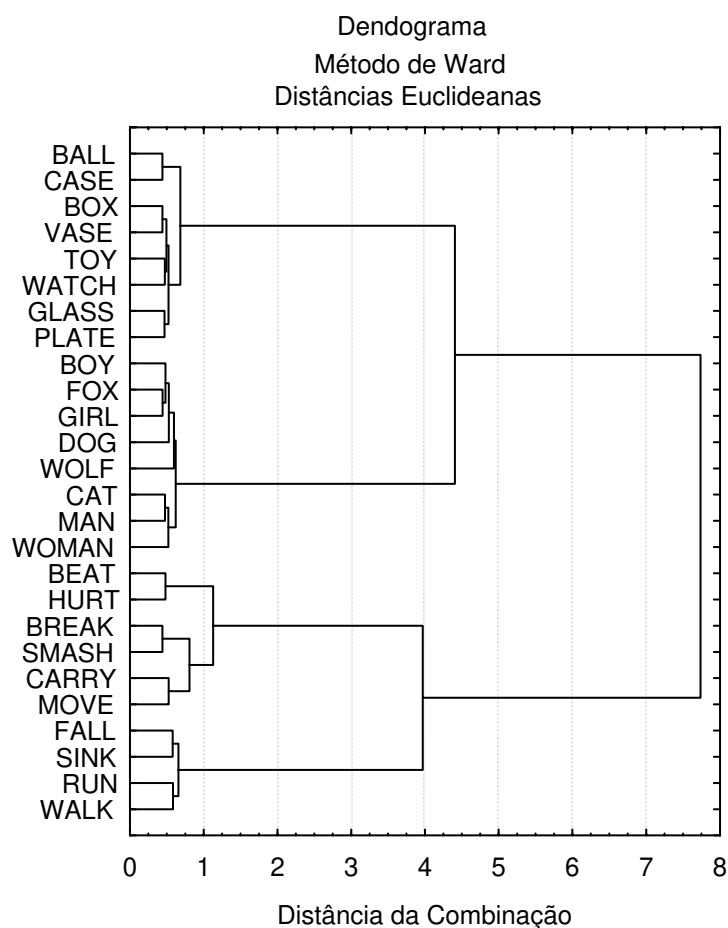


Figura 2.18 A árvore hierárquica mostra a similaridade das representações internas. Da esquerda para a direita, podemos acompanhar o agrupamento realizado e as respectivas distâncias (eixo x) entre as palavras e os grupos. Por exemplo, é de aproximadamente 4 a distância entre o grupo dos verbos transitivos (beat, hurt etc.) em relação ao grupo dos verbos intransitivos (fall, sink etc.).

A partir do dendograma, podemos observar dois grupos principais, que subdividem os itens lexicais em nomes e verbos. O grupo dos nomes está subdividido em animados (*boy*,

fox, girl etc.) e inanimados (*ball, case, box* etc.). O grupo dos verbos está subdividido em transitivos (*beat, hurt, break* etc.) e intransitivos (*fall, sink, run* e *walk*). Com base nesse dendograma, podemos concluir que a representação dos itens lexicais na camada oculta é estruturada segundo informações lingüísticas.

5.3 Análise de componentes principais

Com o intuito de investigar o processamento de sentenças nessa rede, submetemos a matriz obtida a uma ACP. Repare no gráfico (figura 2.19) a alta representatividade dos 4 primeiros CPs, que juntos explicam 80% da variabilidade dos dados. De fato, os dados são mais informativos, e a análise mostra que a rede captou grande parte das regularidades.

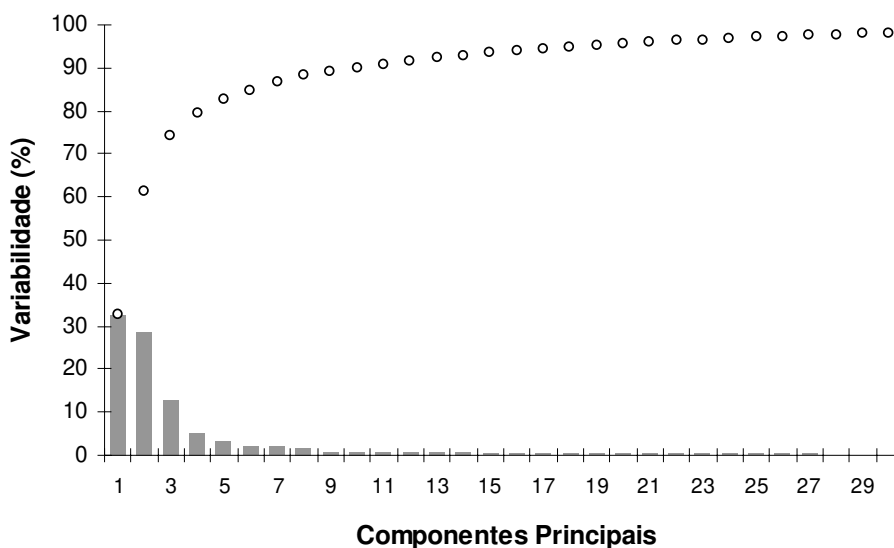


Figura 2.19. As barras verticais indicam a porcentagem da variabilidade contida em cada componente. Os pontos indicam a porcentagem cumulativa ao longo do conjunto de componentes.

Analizamos, então, o processamento das sentenças *boy beat boy* e *boy move boy*, em que os verbos têm diferentes grades temáticas¹³. Veja, na figura 2.20a, a trajetória do processamento dessas sentenças, ao longo do tempo, no CP 1. A informação codificada nesse CP mostra a separação dos verbos. Interpretamos isso como consequência direta do controle da informação contida nas sentenças, ou seja, o fato de termos aumentado o contraste semântico entre os verbos deixou essa informação mais saliente, e aparece

¹³ Usamos o mesmo procedimento de teste e análise relatado na seção 3.2 deste capítulo.

codificada logo no CP mais significativo. Com efeito, a distinção entre os intransitivos *run* e *fall* (o primeiro seleciona arg_1 *agente* enquanto o segundo seleciona *paciente*) também pode ser vista no CP 1 (figura 2.20b).

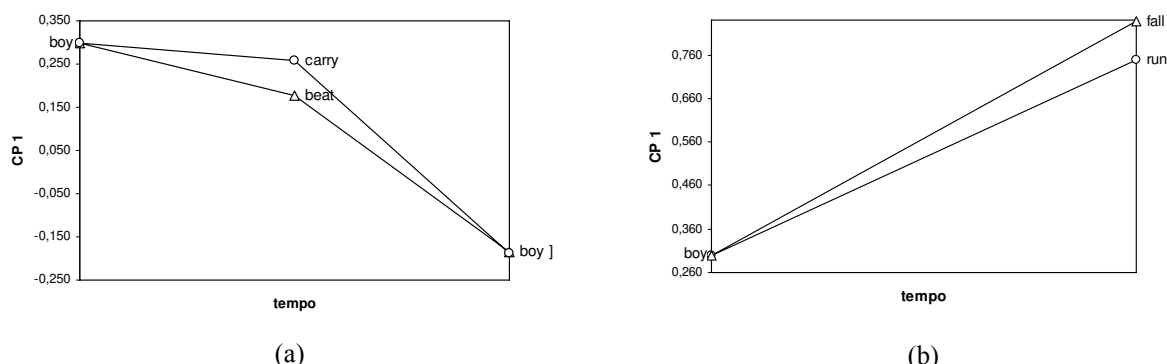


Figura 2.20 Trajetórias no espaço de estados (CP 1) quando a rede processa (a) as sentenças *boy beat boy* e *boy carry boy*; e (b) as sentenças *boy run* e *boy fall*.

Analisando o processamento das sentenças *boy beat boy* e *boy move boy* em outras dimensões, a informação de grade temática pode ser vista no CP 3 (figura 2.21a) e a diferença de papel temático de *boy* em arg_2 está codificada no CP 4 (figura 2.21b).

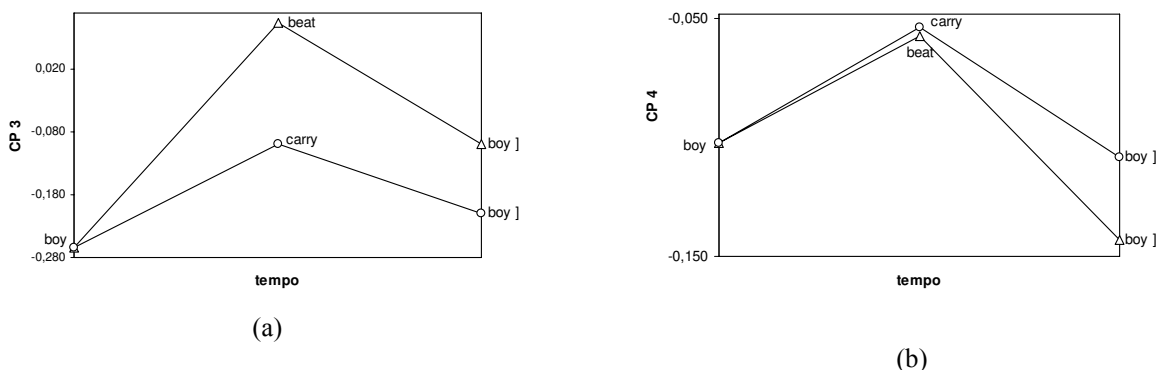


Figura 2.21 Trajetórias no espaço de estados (CPs 3 e 4) quando a rede processa as sentenças *boy beat boy* e *boy carry boy*.

Analisamos um outro par de sentenças, desta vez com a posição de arg_2 preenchida por um inanimado. No caso das sentenças *boy break plate* e *boy carry plate*, observamos que a informação de grade temática pode ser vista no CP 1 (figura 2.22a), enquanto as ocorrências de *plate* estão próximas no CP 7 (figura 2.22b). De fato, a distribuição dos inanimados em arg_1 e arg_2 (tabela 2.8, página 59) é mais restrita do que a dos animados, o que está codificado na dimensão mais significativa.

Em vista dessas análises, é interessante notar que o controle imposto sobre as sentenças geradas (o enfoque sobre as grades temáticas) resultou na perda do conhecimento sintático (conforme obtido no segundo experimento), pois não observamos proximidade entre os verbos. Ou seja, o verbo não é reconhecido como um operador sintático (que tem suas posições argumentais preenchidas por diferentes tipos de nomes), mas sim como um operador semântico (que tem suas posições preenchidas por papéis temáticos).

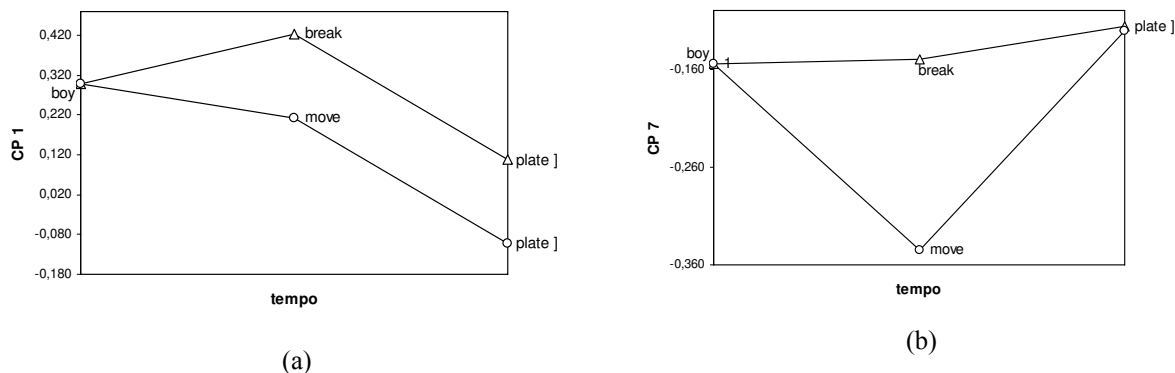


Figura 2.22 Trajetórias no espaço de estados (CPs 1 e 7) quando a rede processa as sentenças *boy break plate* e *boy move plate*.

Notamos, então, que as diferenças nos resultados desta terceira simulação em relação aos obtidos no segundo experimento, nos levam à confirmação de que o modelo que utilizamos (baseado numa RRS) é sensível a informações semânticas (como grade temática) presentes sob a forma de perfis distribucionais.

CAPÍTULO III

Considerações Finais

Retomamos, agora, as questões de interesse apresentadas na seção 2 (capítulo 2). A replicação do experimento relatado em Elman (1990) fez-se necessária porque o trabalho original tinha outros interesses teóricos, de modo que um estudo detalhado do processamento de sentenças foi parcialmente abordado. Apesar de Elman mostrar em trabalhos subseqüentes (1991, 1995) a capacidade da RRS em processar sentenças sintaticamente complexas, observamos a possibilidade de estudar também aspectos semânticos que podem atuar durante o processamento de sentenças.

Exemplificando, nas sentenças *o menino empurrou a menina* e *o menino viu a menina*, temos que *o menino* é *sujeito* e *a menina* é *objeto* de ambas as sentenças. Apesar de essas palavras desempenharem a mesma função sintática, semanticamente elas desempenham diferentes papéis, a saber: *o menino* é *agente* da ação expressa pelo verbo *empurrar*, enquanto *a menina* é *paciente* desta ação. Na outra sentença, *o menino* é *experienciador* e *a menina* é *tema*. Portanto, o processamento adequado dessas sentenças depende do reconhecimento das semelhanças sintáticas e das diferenças semânticas entre as palavras relacionadas pelo verbo.

Como as sentenças que Elman (1990) usou para treinar uma RRS incorporam diferentes relações sintáticas e semânticas entre nomes e verbos, a investigação do aprendizado e representação desse conhecimento mostra-se pertinente e relevante. Um modelo de processamento de sentenças deve dispor de meios para lidar com informações desse tipo. No caso de uma rede neural podemos supor que o espaço de estados (padrões de ativação da camada oculta) aprende a se comportar de maneira a responder adequadamente a esse tipo de informação.

Com essas questões em vista, replicamos o experimento de Elman (1990), com o intuito de realizar testes e análises para investigar a qualidade da representação lingüística, em especial a representação semântica, desenvolvida pela rede. O estudo do processamento de sentenças mostrou que a RRS aprendeu informações semânticas importantes. Por exemplo, a diferença existente entre *boy* de acordo com o tipo de verbo que o precede (*see*

boy vs chase boy). No entanto, as análises também mostraram que a rede não aprendeu informações sintáticas, i.e., noções estruturais como constituinte sintático e sentença.

Esse primeiro insucesso nos remeteu ao experimento de Elman (1991), que foi bem sucedido na aprendizagem da sintaxe. Lá, verificamos o uso de um recurso que torna explícito o final de uma sentença: a última palavra sempre é sucedida pela entrada de um vetor específico, que representa um ‘ponto final’.

Em busca de um modelo capaz de representar tanto informações semânticas quanto sintáticas, no experimento seguinte explicitamos o final das sentenças, de modo a avaliar os efeitos causados unicamente por essa modificação. As análises do processamento das sentenças mostraram a existência de dimensões responsáveis por responder a informações sintáticas e semânticas, ao contrário do resultado parcial obtido no primeiro experimento.

Avaliado o efeito da informação sintática, o próximo passo consistiu em examinar os efeitos da informação semântica, que é dada por vieses distribucionais. Dessa vez, controlamos os padrões de sentenças de forma a ressaltar a grade temática dos verbos. As análises mostraram que as representações desenvolvidas captam principalmente essa informação semântica, havendo a perda da informação sintática mesmo com o uso do ponto final. Apesar disso, esse experimento foi útil para corroborar a idéia que os resultados obtidos na segunda simulação refletem, de fato, informações semânticas presentes sob a forma de perfis distribucionais nas sentenças.

Comparando os resultados obtidos nos três experimentos realizados, podemos considerar que o modelo é capaz de representar conhecimento sintático-semântico que atua no nível de sentenças. Para tanto, conforme observado, o ambiente de treinamento deve conter informações suficientes, seja a pista sintática da fronteira entre as sentenças, seja a pista semântica fornecida pelos diferentes padrões de combinação entre as palavras. Ademais, os resultados são amparados por investigações que mostram os diferentes efeitos causados pelo controle das informações do conjunto de treinamento. Concluimos, portanto, que os experimentos respondem satisfatoriamente às questões propostas.

Feito esse balanço dos experimentos realizados no capítulo 2, ainda nos resta relacionar esse trabalho no quadro de investigação delineado no capítulo 1. Lá, vimos que o uso de modelos e simulações, na psicolinguística computacional, não é apenas uma opção

metodológica, como também reflete pressupostos teóricos sobre a natureza dos processos cognitivos. Nesse contexto, introduzimos a abordagem conexionista, segundo a qual a cognição é a emergência de estados globais em uma rede de unidades simples de processamento (Varela, 1994). Discutidos alguns aspectos gerais, estabelecemos as bases teórico-metodológicas necessárias ao entendimento do trabalho que realizamos no capítulo 2, que trata detalhadamente sobre o processamento de sentenças usando redes recorrentes simples. Nesse momento, limitamos-nos a introduzir as questões de interesse, para então explorá-las através de experimentos, que foram apresentados e analisados separadamente, o que facilitou a exposição frente às questões propostas. Dispensamos, estrategicamente, o tratamento de aspectos teóricos de relevância para a psicolinguística, o que deixou o capítulo 2 mais expositivo do que explicativo.

Há dois aspectos teóricos que, para finalizar, merecem ser destacados.

No capítulo 2, fizemos amplo uso do conceito de representação. Por mais controverso que este termo seja no âmbito das ciências cognitivas, seu uso permite que possamos identificar e explicar, conceitualmente, os processos ocorridos no interior do sistema, entre a entrada e saída.

No caso da rede neural que usamos, vimos que o estado do sistema é definido por um conjunto de 150 variáveis, sendo que cada uma descreve a ativação de um neurônio oculto diante das informações recebidas (entrada atual + estado interno anterior). Vimos, também, que através de uma análise estatística é possível identificar, nesse espaço multidimensional, os conjuntos de unidades (que caracterizam dimensões principais) que se comportam de modo análogo a entradas funcionalmente semelhantes. É isto, pois, que caracteriza a chamada representação distribuída.

Além disso, é importante notar que o espaço de estados também reflete estados anteriores, o que caracteriza uma RRS como um sistema dinâmico. Sendo assim, foi possível analisar o espaço de estados durante o processamento de sentenças, o que revelou a existência de trajetórias que mostram regularidades no tratamento das informações processadas.

Em segundo lugar, o processamento de sentenças, tal como realizado por uma RRS, é feito de forma incremental, de modo que o processamento de uma palavra é influenciado

por palavras anteriormente processadas pelo sistema. Em termos práticos, isso significa que a última palavra é processada no contexto da sentença inteira. De fato, diversos estudos na psicolinguística apontam para a importância de uma memória de trabalho que atua durante o processamento de sentenças. Considerando esse aspecto, redes recorrentes simples e, em especial, o modelo de que nos utilizamos, lidam de forma natural e psicologicamente plausível com a incrementalidade.

Somado a isso, temos evidências para acreditar que o processamento incremental é responsável pela qualidade dos resultados obtidos. Tal é o caso, por exemplo, quando uma parte do sistema responde de forma diferenciada à entrada da palavra *boy*, de acordo com o tipo de verbo que a precede. Não fosse a informação do contexto, o sistema não teria condições de captar informações como essa.

Logo, o processamento incremental, inerente ao funcionamento da RRS, a torna um sistema adequado ao estudo do processamento de sentenças. Com efeito, muito já se explorou sobre a capacidade desse tipo de rede em processar sentenças, principalmente no estudo de fenômenos sintáticos. Para citar alguns: Christiansen (1994), Lupyan & Christiansen (2002), além dos trabalhos de Elman (1991, 1993, 1995).

REFERÊNCIAS BIBLIOGRÁFICAS

- ANDERSON, J. A. (1972) A simple neural network generating an interactive memory. *Mathematical Biosciences*, 14, 197-220.
- AUDI, R. (1999) *The Cambridge Dictionary of Philosophy*. New York: Cambridge University Press.
- BECHTEL, W., ABRAHAMSEN, A., & GRAHAM, G. (1998). The life of cognitive science. In Bechtel, W. & Graham, G. (Eds.) *A companion to cognitive science*. Oxford: Blackwell.
- BECHTEL, W. & ABRAHAMSEN, A. (2002) *Connectionism and the Mind – Parallel Processing, Dynamics and Evolution in Networks*. Oxford: Blackwell.
- BERTALANFFY, L. von. (1968) *General system theory: Foundations, development, applications*. New York: George Braziller.
- BLACKBURN, S. (1997) *Dicionário Oxford de Filosofia*. Rio de Janeiro: Jorge Zahar Editor.
- BRESCIANI, E. & D'OTTAVIANO, I. (2000) Conceitos básicos de sistêmica. In D'Ottaviano, I. & Gonzales, M.E. (Org.) *Auto-organização: estudos interdisciplinares*. Campinas: CLE/UNICAMP
- CHOMSKY, N. (1980) *Rules and representations*. Oxford: Blackwell.
- CHRISTIANSEN, M.H. (1994) *Infinite languages, finite minds: connectionism, learning and linguistic structure*. Tese de doutoramento, Universidade de Edinburgo.
- CHRISTIANSEN, M.H. & CHATER, N. (1999) Towards a connectionist model of recursion in human linguistic performance. *Cognitive Science*, 23: 157-205.
- _____ (eds.) (2001) *Connectionist Psycholinguistics*. Westport: Ablex.
- CORRÊA, L. M. S. (2002) Explorando a relação entre língua e cognição na interface: o conceito de interpretabilidade e suas implicações para teorias do processamento e da aquisição da linguagem. In *Veredas: Revista de Estudos Lingüísticos*. Juiz de Fora.
- CROCKER, M. (1996) *Computational Psycholinguistics: An interdisciplinary approach to the study of language*. Dordrecht: Kluwer Academic Publishers.
- CRYSTAL, D. (1988) *Dicionário de Lingüística*. Rio de Janeiro: Jorge Zahar Editor.
- DELCHAMPS, D. (1988) *State Space and Input-Output Linear Systems*. New York: Springer-Verlag.

- DIJKSTRA, T. & SMEDT, K. (1996a) *Computational Psycholinguistics*. London: Taylor & Francis.
- DIJKSTRA, T. & SMEDT, K. (1996b) Computer models in psycholinguistics: an introduction. In Dijkstra, T. & Smedt, K. (1996) *Computational Psycholinguistics*. London: Taylor & Francis.
- DORON, R. & PAROT, F (1998) *Dicionário de Psicologia*. São Paulo: Editora Ática.
- ELMAN, J. (1990) Finding structure in time. *Cognitive Science*, 14: 179-211.
- _____ (1991) Distributed representation, simple recurrent networks, and grammatical structure. *Machine Learning*, 7: 195-225.
- _____ (1993) Learning and development in neural networks: the importance of starting small. *Cognition*, 48: 71-99.
- _____ (1995) Language as a dynamical system. In Port, R.F. & van Gelder, T. (Eds.), *Mind as Motion: Explorations in the Dynamics of Cognition*. Cambridge, MA: MIT Press.
- ELMAN, J., BATES, E., JOHNSON, M., KARMILOFF-SMITH, A., PARISI, D. & PLUNKETT, K. (1996) *Rethinking Innateness. A Connectionist Perspective on Development*. Cambridge, MA: MIT Press.
- FRANÇOZO, E. & ALBANO, E. C. (2003) Virtudes e Vicissitudes do Cognitivismo, revisitadas. *In Cognito*, 1(1):45 - 54.
- GARDNER, H. (1996) *A nova ciência da mente: uma história da revolução cognitiva*. São Paulo: EDUSP.
- GERNSBACHER, M. (1994) *Handbook of Psycholinguistics*. Academic Press.
- HAEGEMAN, L. (1991) *Introduction to Government and Binding Theory*. Oxford: Blackwell.
- HAIR, J., ANDERSON, R., TATHAM, R., BLACK, W. (2005) *Análise multivariada de dados*. Porto Alegre: Bookman.
- HARE, M., MC RAE, K. AND ELMAN, J.L. (2003) Sense and structure: Meaning as a determinant of verb subcategorization preferences. *Journal of Memory and Language*, 48, 281-303
- HAYKIN, S. (2001) *Redes Neurais: Princípios e prática*. Porto Alegre: Bookman.
- KHEIR, N. D. (1988) *Systems Modelling and Computer Simulation*. New York: Marcel Dekker, Inc.
- KLIR, G. (1969) *Approach to General Systems Theory*. New York: van Nostrand.

- KLIR, G. (1991) *Facets of Systems Science*. New York: Plenum Press.
- KOHONEN, T. (1972) Correlation matrix memories. *IEEE Transactions on Computers*, C-21, 353-359.
- KOVÁCS, Z. (1996) *Redes Neurais Artificiais: Fundamentos e Aplicações*, 2ª. ed. São Paulo: Edição Acadêmica.
- LACHMAN, R. (1960) The model in theory construction. *Psychological Review*, 67, 113-129.
- LANDY, D. (2004) Representations in Simple Recurrent Networks Which are Always Compositional. *Proceedings of the 26th Annual Conference of the Cognitive Science Society*.
- LUPIAN, G. & CHRISTIANSEN, M.H. (2002) Case, word order and language learnability: insights from connectionist modeling. In *Proceedings of the 24th Annual Conference of the Cognitive Science Society* p. 596-601). Mahwah, NJ: Lawrence Erlbaum.
- MCCULLOCH, W.S. & PITTS, W. (1943) A logical calculus of the ideas immanent in nervous activity. *Bulletin of Mathematical Biophysics*, 5: 115-133.
- NAKANO, K. (1972) Association – a model of associative memory. *IEEE Transactions on Systems, Man and Cybernetics*, SMC-2, 380-388.
- NEELANKAVIL, F. (1987) *Computer simulation and modelling*. Chichester: John Wiley and sons.
- NEWELL, A. & SIMON, H. A. (1976) Computer science as empirical inquiry: Symbols and search. *Communications of the Association for Computing Machinery*, 19(3), 113-126.
- NEWPORT, E.L. (1988) Constraints on learning and their role in language acquisition: Studies of the acquisition of American Sign Language. *Language Sciences*, 10, 147-172.
- NEWPORT, E.L. (1990) Maturational constraints on language learning. *Cognitive Science*, 14, 11-28.
- NORMAN, D. A. (1986) Reflections on Cognition and Parallel Distributed Processing. In McClelland, J. L., Rumelhart, D. E. & the PDP Research Group (Eds.), *Parallel distributed processing: Explorations in the microstructure of cognition. Volume 2: Psychological and biological models* (p. 531-546). Cambridge, MA: MIT Press.
- PITTS, W. & MCCULLOCH, W. S. (1947) How we know universals: The perception of auditory and visual forms. *Bulletin of Mathematical Biophysics*, 9: 127-147.

- PLUNKETT, K. & ELMAN, J. (1997) *Exercises in Rethinking Innateness. A handbook for connectionist simulations*. Cambridge, MA: MIT Press.
- PORT, R. & VAN GELDER, T. (eds) (1995) *Mind as Motion: Explorations in the Dynamics of Cognition*. Cambridge, MA: MIT Press.
- ROSENBLATT, F. (1958) The perceptron: A probabilistic model for information storage and organization in the brain. *Psychological Review*, 65, 368-408.
- ROSSA, A. ROSSA, C. (2004) *Rumo à psicolinguística conexionista*. Porto Alegre: EDIPUCRS.
- RUMELHART, D., HINTON, G., WILLIAMS, R. (1986) Learning internal representations by error propagation. In Rumelhart, D. E., McClelland, J. L. & the PDP Research Group (Eds.) *Parallel Distributed Processing: Explorations in the Microstructure of Cognition. Volume 1*. Cambridge, MA: MIT Press.
- SELFRIDGE, O. G. (1959) Pandemonium: A paradigm for learning. In *Symposium on the Mechanization of Thought Processes*. London: HMSO.
- SHANNON, R. E. (1975). *Systems simulation: the art and science*. Englewood Cliffs: Prentice Hall, Inc.
- STILLINGS, N., FEINSTEIN, M., GARFIELD, J., RISSLAND, E., ROSENBAUM, D., WEISLER, S. & BAKER-WARD, L. (1987) *Cognitive Science: an Introduction*. Cambridge, MA: MIT Press.
- TAYLOR, W. K. (1956) Electrical simulation of some nervous system functional activities. *Information Theory*, vol. 3, E. C. Cherry, ed., pp.314-328.
- TRAXLER, M. & GERNSBACHER, M. (2006) *Handbook of Psycholinguistics*. Academic Press; 2nd edition.
- TURING, A. (1936) On computable numbers, with an application to the Entscheidungsproblem. *Proceedings of the London Mathematical Society*, Series 2, 42, 230-265.
- VAN GELDER, T. & PORT, R. (1995) It's About Time: an Overview of the Dynamical Approach to Cognition. In Port, R. & van Gelder, T. (Eds.). *Mind as Motion: Explorations in the Dynamics of Cognition*. Cambridge, MA: MIT Press.
- VARELA, F. (1994) *Conhecer: As ciências cognitivas, tendências e perspectivas*. Lisboa: Instituto Piaget.
- WEINBERG, G. (1975) *An Introduction to General Systems Thinking*. New York, John Wiley and Sons.

ZEIGLER, B. P., PRAEHOFER, H. & KIM, T. G. (2000) *Theory of Modeling and Simulation*. London: Academic Press.

ANEXO A

Código em C do programa gerador de sentenças (Experimento 1)

```
/* programa desenvolvido por Rafael Dias Santos */
#include <stdio.h>
#include <stdlib.h>
#include <time.h>

/* número de padroes de sentenças: 16 (0-15)*/
#define NFRAMES 16

/* categorias */
char *HUM[]={"MAN","BOY","WOMAN","GIRL"}; int nHUM=4;
char *ANI[]={"CAT","MOUSE","DOG","MONSTER","DRAGON","LION"}; int nANI=6;
char *AGR[]={"MONSTER","DRAGON","LION"}; int nAGR=3;
char *OBJ[]={"CAR","BOOK","ROCK","PLATE","GLASS","SANDWICH","BREAD","COOKIE"}; int
nOBJ=8;
char *ALIM[]={"SANDWICH","BREAD","COOKIE"}; int nALIM=3;
char *QUEB[]={"PLATE","GLASS"}; int nQUEB=2;

char *TRAN[]={"CHASE","LIKE","EAT","SEE","SMELL","MOVE","BREAK","SMASH"}; int
nTRAN=8;
char *INTR[]={"THINK","SLEEP","EXIST"}; int nINTR=3;
char *AGPAT[]={"MOVE","BREAK","SMASH"}; int nAGPAT=3;
char *DEST[]={"BREAK","SMASH"}; int nDEST=2;
char *PERC[]={"SEE","SMELL"}; int nPERC=2;
char *COME[]={"EAT"}; int nCOME=1;

void main(){
    int i;
    int frame;
    FILE *fp;

    /* início do gerador aleatório */
    randomize();

    /* arquivo de saída */
    fp = fopen("sentences.txt", "wt");

    /* número de sentenças = 10000*/
    for (i=0;i<10000;i++){

        /* escolhe um padrão de sentença aleatório */
        frame = random(NFRAMES);

        /* preencha o padrão */
        switch(frame){
            case(0): /* N-HUM    V-COME    N-ALIM */
                fprintf(fp,"%s ",HUM[random(nHUM)]);
                fprintf(fp,"%s ",COME[random(nCOME)]);
                fprintf(fp,"%s\n",ALIM[random(nALIM)]);
                break;
            case(1): /* N-HUM    V-PERC    N-OBJ */
```

```

fprintf(fp, "%s ", HUM[random(nHUM)]);
fprintf(fp, "%s ", PERC[random(nPERC)]);
fprintf(fp, "%s\n", OBJ[random(nOBJ)]);
break;
case(2): /* N-HUM    V-DEST      N-QUEB */
fprintf(fp, "%s ", HUM[random(nHUM)]);
fprintf(fp, "%s ", DEST[random(nDEST)]);
fprintf(fp, "%s\n", QUEB[random(nQUEB)]);
break;
case(3): /* N-HUM    V-INTR */
fprintf(fp, "%s ", HUM[random(nHUM)]);
fprintf(fp, "%s\n", INTR[random(nINTR)]);
break;
case(4): /* N-HUM    V-TRAN      N-HUM */
fprintf(fp, "%s ", HUM[random(nHUM)]);
fprintf(fp, "%s ", TRAN[random(nTRAN)]);
fprintf(fp, "%s\n", HUM[random(nHUM)]);
break;
case(5): /* N-HUM    V-AGPAT      N-OBJ */
fprintf(fp, "%s ", HUM[random(nHUM)]);
fprintf(fp, "%s ", AGPAT[random(nAGPAT)]);
fprintf(fp, "%s\n", OBJ[random(nOBJ)]);
break;
case(6): /* N-HUM    V-AGPAT      */
fprintf(fp, "%s ", HUM[random(nHUM)]);
fprintf(fp, "%s\n", AGPAT[random(nAGPAT)]);
break;
case(7): /* N-ANI    V-EAT      N-ALIM */
fprintf(fp, "%s ", ANI[random(nANI)]);
fprintf(fp, "%s ", COME[random(nCOME)]);
fprintf(fp, "%s\n", ALIM[random(nALIM)]);
break;
case(8): /* N-ANI    V-TRAN      N-ANI */
fprintf(fp, "%s ", ANI[random(nANI)]);
fprintf(fp, "%s ", TRAN[random(nTRAN)]);
fprintf(fp, "%s\n", ANI[random(nANI)]);
break;
case(9): /* N-ANI    V-AGPAT      N-OBJ */
fprintf(fp, "%s ", ANI[random(nANI)]);
fprintf(fp, "%s ", AGPAT[random(nAGPAT)]);
fprintf(fp, "%s\n", OBJ[random(nOBJ)]);
break;
case(10): /* N-ANI    V-AGPAT */
fprintf(fp, "%s ", ANI[random(nANI)]);
fprintf(fp, "%s\n", AGPAT[random(nAGPAT)]);
break;
case(11): /* N-OBJ    V-AGPAT */
fprintf(fp, "%s ", OBJ[random(nOBJ)]);
fprintf(fp, "%s\n", AGPAT[random(nAGPAT)]);
break;
case(12): /* N-AGR    V-DEST      N-QUEB */
fprintf(fp, "%s ", AGR[random(nAGR)]);
fprintf(fp, "%s ", DEST[random(nDEST)]);
fprintf(fp, "%s\n", QUEB[random(nQUEB)]);
break;
case(13): /* N-AGR    V-COME      N-HUM */

```

```

    fprintf(fp, "%s ", AGR[random(nAGR)]);
    fprintf(fp, "%s ", COME[random(nCOME)]);
    fprintf(fp, "%s\n", HUM[random(nHUM)]);
    break;
case(14): /* N-AGR    V-COME        N-ANI */
    fprintf(fp, "%s ", AGR[random(nAGR)]);
    fprintf(fp, "%s ", COME[random(nCOME)]);
    fprintf(fp, "%s\n", ANI[random(nANI)]);
    break;
case(15): /* N-AGR    V-COME        N-FOOD */
    fprintf(fp, "%s ", AGR[random(nAGR)]);
    fprintf(fp, "%s ", COME[random(nCOME)]);
    fprintf(fp, "%s\n", ALIM[random(nALIM)]);
    break;
}
}

fclose(fp);
}

```

ANEXO B

Código em C do programa gerador de sentenças (Experimento 3)

```
#include <stdio.h>
#include <stdlib.h>
#include <time.h>

/* número de padrões 5 (0-4) */
#define NFRAMES 5

/* categorias */

char *ANI[]={ "WOMAN", "MAN", "GIRL", "BOY", "WOLF", "FOX", "CAT", "DOG" }; int nANI=8;
char *INA[]={ "PLATE", "GLASS", "VASE", "BOX", "CASE", "WATCH", "TOY", "BALL" }; int nINA=8;
char
*NOMES[]={ "WOMAN", "MAN", "GIRL", "BOY", "WOLF", "FOX", "CAT", "DOG", "PLATE", "GLASS", "V
ASE", "BOX", "CASE", "WATCH", "TOY", "BALL" }; int nNOMES=16;

char *VI[]={ "RUN", "WALK", "FALL", "SINK" }; int nVI=4;          /* Verbos I */
char *VII[]={ "BEAT", "HURT" }; int nVII=2;                     /* Verbos II */
char *VIII[]={ "BREAK", "SMASH" }; int nVIII=2;                /* Verbos III */
char *VIV[]={ "MOVE", "CARRY" }; int nVIV=2;                   /* Verbos IV */
char *VV[]={ "FALL", "SINK", "BREAK", "SMASH" }; int nVV=4;     /* Verbos V */

void main(){
    int i;
    int frame;
    FILE *fp;

    /* início do gerador aleatório */
    randomize();

    /* arquivo de saída */
    fp = fopen("sentences.txt", "wt");

    for (i=0; i<10000; i++){
        /* escolhe um padrão de sentença aleatoriamente */
        frame = random(NFRAMES);

        /* preenche o padrão */
        switch(frame){
            case(0): /* N-ANI      VI */
                fprintf(fp, " %s ", ANI[random(nANI)]);
                fprintf(fp, "%s\n", VI[random(nVI)]);
                break;
            case(1): /* N-ANI      VII      N-ANI */
                fprintf(fp, " %s ", ANI[random(nANI)]);
                fprintf(fp, "%s ", VII[random(nVII)]);
                fprintf(fp, "%s\n", ANI[random(nANI)]);
                break;
            case(2): /* N-ANI      VIII      N-INA */
                fprintf(fp, " %s ", ANI[random(nANI)]);
```

```

    fprintf(fp, "%s ", VIII[random(nVIII)]);
    fprintf(fp, "%s\n", INA[random(nINA)]);
    break;
case(3): /* N-ANI      VIV N-NOMES*/
    fprintf(fp, " %s ", ANI[random(nANI)]);
    fprintf(fp, "%s ", VIV[random(nVIV)]);
    fprintf(fp, "%s\n", NOMES[random(nNOMES)]);
    break;
case(4): /* N-INA      VV      */
    fprintf(fp, " %s ", INA[random(nINA)]);
    fprintf(fp, "%s\n", VV[random(nVV)]);
    break;
    }
}
fclose(fp);
}

```