

Universidade Estadual de Campinas
Instituto de Matemática, Estatística e Computação Científica

**Um Estudo Comparativo em Memórias Associativas
com Ênfase em Memórias Associativas Morfológicas**

Autor: Marcos Eduardo Ribeiro do Valle Mesquita

Orientador: Prof. Dr. Peter Sussner

Campinas, SP

Agosto/2005

Este exemplar corresponde à redação final da dissertação devidamente corrigida e defendida por **Marcos Eduardo Ribeiro do Valle Mesquita** e aprovada pela comissão julgadora.

Campinas, 24 de Agosto de 2005

Prof. Dr. Peter Sussner

Banca Examinadora

1. Alvaro Rodolfo De Pierro, Dr. DMA/IMECC/Unicamp
2. Emanuel Pimentel Barbosa, Dr. DE/IMECC/Unicamp
3. Fernando Gomide, Dr. DCA/FEEC/Unicamp
4. Peter Sussner, Dr. DMA/IMECC/Unicamp

Dissertação apresentada ao Instituto de Matemática, Estatística e Computação Científica, UNICAMP, como requisito parcial para obtenção do Título de Mestre em Matemática Aplicada.

**FICHA CATALOGRÁFICA ELABORADA PELA
BIBLIOTECA DO IMECC DA UNICAMP**

Bibliotecário: Maria Júlia Milani Rodrigues - CRB8a / 2116

Mesquita, Marcos Eduardo Ribeiro do Valle

M562e Um estudo comparativo em memórias associativas com ênfase em memórias associativas morfológicas / Marcos Eduardo Ribeiro do Valle -- Campinas, [S.P. :s.n.], 2005.

Orientador : Peter Sussner

Dissertação (mestrado) - Universidade Estadual de Campinas, Instituto de Matemática, Estatística e Computação Científica.

1. Redes neurais (Computação). 2. Morfologia matemática. 3. Sistemas de memória de computadores. I. Sussner, Peter. II. Universidade Estadual de Campinas. Instituto de Matemática, Estatística e Computação Científica. III. Título.

Título em inglês: A comparative study on associative memories with emphasis on morphological associative memories

Palavras-chave em inglês (Keywords): 1. Neural networks (Computer science). 2. Mathematical morphology. 3. Computer memory systems.

Área de concentração: Matemática Aplicada

Titulação: Mestre em Matemática Aplicada

Banca examinadora: Prof. Dr. Peter Sussner (IMECC-UNICAMP)

Prof. Dr. Álvaro Rodolfo de Pierro (IMECC-UNICAMP)

Prof. Dr. Emanuel Pimentel Barbosa (IMECC-UNICAMP)

Prof. Dr. Fernando Antônio Campos Gomide (FEEC-UNICAMP)

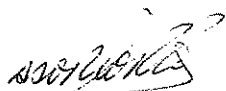
Data da defesa: 24/08/2005

Dissertação de Mestrado defendida em 24 de agosto de 2005 e aprovada

Pela Banca Examinadora composta pelos Profs. Drs.



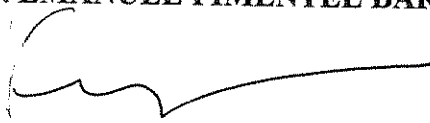
Prof. (a). Dr (a). PETER SUSSNER



Prof. (a). Dr (a). ÁLVARO RODOLFO DE PIERRO



Prof. (a). Dr (a). EMANUEL PIMENTEL BARBOSA



Prof. (a). Dr (a). FERNANDO ANTONIO CAMPOS GOMIDE

Resumo

Memórias associativas neurais são modelos do fenômeno biológico que permite o armazenamento de padrões e a recordação destes após a apresentação de uma versão ruidosa ou incompleta de um padrão armazenado. Existem vários modelos de memórias associativas neurais na literatura, entretanto, existem poucos trabalhos comparando as várias propostas. Nesta dissertação comparamos sistematicamente o desempenho dos modelos mais influentes de memórias associativas neurais encontrados na literatura. Esta comparação está baseada nos seguintes critérios: capacidade de armazenamento, distribuição da informação nos pesos sinápticos, raio da bacia de atração, memórias espúrias e esforço computacional. Especial ênfase é dado para as memórias associativas morfológicas cuja fundamentação matemática encontra-se na morfologia matemática e na álgebra de imagens.

Palavras-chave: Redes Neurais, Morfologia Matemática, Sistemas de Memória de Computadores.

Abstract

Associative neural memories are models of biological phenomena that allow for the storage of pattern associations and the retrieval of the desired output pattern upon presentation of a possibly noisy or incomplete version of an input pattern. There are several models of neural associative memories in the literature, however, there are few works relating them. In this thesis, we present a systematic comparison of the performances of some of the most widely known models of neural associative memories. This comparison is based on the following criteria: storage capacity, distribution of the information over the synaptic weights, basin of attraction, number of spurious memories, and computational effort. The thesis places a special emphasis on morphological associative memories whose mathematical foundations lie in mathematical morphology and image algebra.

Keywords: Neural Networks, Mathematical Morphology, Computer Memory Systems.

À Margarida Ribeiro do Valle e Mercedes Mesquita Silva

Agradecimentos

Aos meus pais, Antônio Marcos de Mesquita Silva e Ana Lúcia Ribeiro do Valle Silva, pela vida e educação que me deram.

À minha irmã, Ana Elisa do Valle Mesquita, e a minha noiva, Luciana Maria Ricci, pelo apoio durante esta jornada.

Ao meu orientador, Peter Sussner, sou grato pela orientação.

A minha família e a todos os meus colegas da Unicamp e São João da Boa Vista. Em particular para Márcio, Rangel, Gustavo, Roberto, Jorge, Renata, André, David, Ederson, Mateus, Homero, Denilson e Ângela. As boas amizades permanecem!

A todos os professores que contribuíram para a minha formação na Unicamp.

À FAPESP, pelo apoio financeiro.

Sumário

1	Introdução	1
1.1	Contexto Histórico	2
1.1.1	Redes Neurais Artificiais	2
1.1.2	Memórias Associativas Neurais	4
1.1.3	Morfologia Matemática e Álgebra de Imagens	5
1.2	Objetivos e Organização da Dissertação	6
2	Conceitos Básicos de Redes Neurais	9
2.1	Introdução	9
2.2	Modelos Neurais	9
2.2.1	Modelo Neural Clássico	10
2.2.2	Modelo Neural Morfológico	11
2.3	Arquiteturas de Redes Neurais	12
2.4	Aprendizagem	15
2.4.1	Aprendizado Supervisionado	15
3	Conceitos Básicos de Memórias Associativas Neurais	17
3.1	Formulação Matemática, Armazenamento e Associação	17
3.2	Classificação das Memórias Associativas Neurais	18
3.3	Características para um Bom Desempenho	20
4	Memórias Associativas Lineares	21
4.1	Armazenamento por Correlação	21
4.1.1	Armazenamento por Correlação Auto-associativo Bipolar	23
4.2	Armazenamento por Projeção	24
5	Memórias Associativas Dinâmicas	29
5.1	Memória Associativa de Hopfield Discreta	29
5.1.1	Arquitetura da Rede	29
5.1.2	Aprendizado	30
5.1.3	Convergência	30
5.2	Memória Associativa Bidirecional	33
5.2.1	Arquitetura	33
5.2.2	Aprendizado	33

5.2.3	Convergência	34
5.3	Memória Associativa de Personnaz	38
5.3.1	Arquitetura da Rede	38
5.3.2	Aprendizado	38
5.3.3	Convergência	39
5.4	Memória Associativa de Kanter-Sompolinsky	39
5.4.1	Arquitetura da Rede	40
5.4.2	Aprendizado	40
5.4.3	Convergência	40
5.5	Memória Associativa Bidirecional Assimétrica	40
5.5.1	Arquitetura	41
5.5.2	Aprendizado	41
5.6	Memória Associativa com Capacidade Exponencial	42
5.6.1	Arquitetura	42
5.6.2	Aprendizado	43
5.6.3	Convergência	44
5.7	Memória Associativa Bidirecional com Capacidade Exponencial	44
5.7.1	Arquitetura	44
5.7.2	Aprendizado	46
5.7.3	Convergência	46
5.8	Modelo do Estado Cerebral numa Caixa (BSB)	46
5.8.1	Arquitetura	46
5.8.2	Aprendizado	47
5.8.3	Convergência	47
6	Memórias Associativas Morfológicas	49
6.1	Memórias Associativas Morfológicas Heteroassociativas	49
6.1.1	Aprendizado	50
6.2	Memórias Auto-Associativas Morfológicas	58
6.3	Memórias Auto-associativas Morfológicas Binárias	61
6.4	Memórias Associativas Morfológicas de Duas Camadas	63
6.4.1	Arquitetura	63
6.4.2	Aprendizado	65
7	Desempenho das Memórias Associativas Binárias	67
7.1	Capacidade de Armazenamento	67
7.1.1	Memória Associativa de Hopfield	69
7.1.2	Memória Associativa Bidirecional	70
7.1.3	Memória Associativa de Personnaz	70
7.1.4	Memória Associativa de Kanter-Sompolinsky	71
7.1.5	Memória Associativa Bidirecional Assimétrica	71
7.1.6	Memória Associativa com Capacidade Exponencial	71
7.1.7	Memória Associativa Bidirecional com Capacidade Exponencial	71
7.1.8	Memórias Associativa Morfológicas	72

7.1.9	Memória Associativa Morfológica de Duas Camadas	72
7.2	Distribuição da Informação	72
7.3	Raio de Atração	73
7.4	Memórias Espúrias	77
7.5	Esforço Computacional	78
7.5.1	Número de Operações na Fase de Armazenamento	80
7.5.2	Número de Operações por Iteração na Fase de Recordação	81
7.5.3	Número de Iterações na Fase de Recordação	82
8	Conclusão	85
	Referencias bibliograficas	88

Capítulo 1

Introdução

A primeira pergunta que surge em nossa mente é, o que é uma memória? A resposta não é simples e não é nosso objetivo discutir este assunto, mas podemos dizer que uma memória é, simplificada-mente, um sistema que possui três funções ou etapas: 1 – *Registro*, processo pelo qual armazenamos informação; 2 – *Preservação*, para garantir que a informação está intacta; 3 – *Recordação*, processo pelo qual uma informação é recuperada [12].

Existem vários tipos de memória. Por exemplo, quando escrevemos o número de um telefone num papel, estamos usando o papel como memória. Depois poderemos ler e usar este número.

Sempre que registramos informações na memória, precisamos de uma chave ou algo que permita recuperar o conteúdo armazenado. Por exemplo, quando deixamos uma bolsa num guarda-volumes, pegamos um ticket indicando o compartimento onde ela ficará. Aqui, o ticket não possui nenhuma relação com o conteúdo da bolsa e podemos dizer que o endereço (ticket) é apenas um símbolo abstrato da memória onde a entidade (bolsa) está, e mais, este endereço não possui nenhuma relação com o conteúdo armazenado. Este tipo de memória é muito usado nos computadores digitais (memória RAM ou ROM). Em muitos casos este tipo de memória apresenta-se eficiente, entretanto possui uma série de limitações. Por exemplo, o que acontecerá se perdermos o ticket?

Suponha que perdemos o pequeno ticket. Teremos que usar um procedimento diferente para recuperar nossa bolsa. Evidentemente procuraremos o responsável e passaremos informações parciais, mas suficientes, sobre a bolsa e seu conteúdo. É comum encontrarmos bolsas semelhantes, mas o conteúdo geralmente é diferente e uma descrição parcial dele é suficiente para que o responsável possa identificá-la, e provavelmente concordará em devolvê-la. Neste exemplo, podemos dizer que o endereço (descrição parcial do que tem na bolsa) é igual ao conteúdo e dizemos que esta é uma memória auto-associativa (também conhecida por memória endereçada por conteúdo), um caso particular das memórias associativas.

Uma memória associativa poderia recuperar um item da memória a partir de informações parciais. Por exemplo, se um item armazenado na memória é “J.J. Hopfield & D.W. Tank, *Biological Cybernetics* 52, 141–152 (1985)”. A entrada “& Tank (1985)” poderia ser suficiente para recuperarmos a informação completa. Além disso, uma memória associativa ideal poderia trabalhar com ruídos (ou erros) e recuperar esta referência mesmo a partir de entradas incorretas como “& Rank, (1985)”. Nos computadores digitais, apenas formas relativamente simples de memória associativa tem sido implementadas em hardware. A maioria dos recursos para tolerância a ruído no acesso da informação são implementados via software [40]. Esta é uma das razões para o estudo das memórias associativas.

As memórias associativas encontram aplicações em vários ramos da ciência. Por exemplo, Zhang *et. al.* utilizaram um modelo de memória associativa para reconhecimento e classificação de padrões [111, 112]. A metodologia para classificação de padrões baseada em memórias associativas também foi aplicada em problemas de detecção de falha em motores [54], segurança de rede [110] e aprendizado de linguagem natural [29]. Hopfield mostrou que seu modelo de memória associativa pode ser usado para resolver problemas de otimização, como por exemplo, o problema do caixeiro viajante [38]. As memórias associativas morfológicas discutidas nesta dissertação foram aplicadas em problemas de localização de faces, auto-localização e análise de imagens hiperespectrais [68, 26]. Recentemente, Valle e Sussner apresentaram uma aplicação das memórias associativas nebulosas em um modelo de previsão [103, 98, 102].

O termo memória associativa veio da psicologia e não da engenharia. Veio da psicologia porque o cérebro humano pode ser visto como uma memória associativa. Ele associa o item a ser lembrado com um fragmento da recordação. Por exemplo, ouvindo um trecho de uma música podemos lembrar da canção inteira, ou sentido um certo perfume podemos associar o cheiro a uma pessoa especial. Não só o cérebro humano, mas moscas de fruta ou lesmas de jardim também possuem memórias associativas. Na verdade, qualquer sistema nervoso relativamente simples apresenta uma memória associativa [12]. Isso sugere que a habilidade de criar associações é natural – praticamente espontânea – em qualquer sistema neural. Portanto, uma fonte de inspirações para os estudos das memórias associativas encontra-se nos estudos do funcionamento de um sistema nervoso, em particular, do cérebro humano.

O cérebro humano é composto por bilhões de neurônios interligados formando uma rede. Dizemos que o cérebro é uma rede neural biológica. Em nossos estudos, apresentaremos modelos que descrevem um neurônio e chamaremos este modelo de neurônio artificial. Chamaremos de rede neural artificial, ou simplesmente rede neural, uma rede formada por neurônios artificiais [33]. A teoria das redes neurais é vasta e possui aplicações em várias áreas, como por exemplo, no reconhecimento de padrões, controle, otimização e previsão de mercados financeiros [8, 45, 61]. Neste trabalho voltaremos nossa atenção para a interseção das redes neurais com as memórias associativas. Especificamente, estudaremos as memórias associativas neurais.

Neste trabalho também discutiremos as memórias associativas morfológicas que são modelos de memória associativa onde usamos a morfologia matemática e a álgebra de imagens como ferramenta [70, 73, 76]. Este tipo particular de memória associativa neural será o enfoque principal do nosso trabalho.

Antes de introduzirmos os conceitos sobre redes neurais e memórias associativas, apresentaremos brevemente a história das *redes neurais*, *memórias associativas*, *morfologia matemática* e da *álgebra de imagens*.

1.1 Contexto Histórico

1.1.1 Redes Neurais Artificiais

Podemos dizer que os estudos das redes neurais artificiais iniciaram em 1943 quando o biólogo Warren McCulloch e o matemático Walter Pitts apresentaram um modelo matemático de um neurônio biológico [55]. Eles assumiram que um neurônio seguia uma lei “tudo ou nada” e acreditavam que

com um número suficiente de neurônios com conexões sinápticas apropriadas operando de forma síncrona (paralelamente), seria possível, a princípio, a computação de qualquer função booleana computável. Este foi um resultado muito significativo e com ele é aceito o surgimento das disciplinas de redes neurais e inteligência artificial [33].

Em 1949, o neurofisiologista Donald Hebb publicou o livro “The Organization of Behavior”[34]. Neste livro foi apresentada pela primeira vez a formulação de uma regra de aprendizagem. Hebb notou que as conexões sinápticas do cérebro são continuamente modificadas conforme um organismo aprende novas tarefas, criando assim agrupamentos neurais. Ele propôs no seu livro o seguinte postulado de aprendizagem: “A eficiência de uma sinapse é aumentada pela interação entre dois neurônios através da sinapse”. Este postulado forma a base do que chamamos hoje de *aprendizagem (ou regra) de Hebb*. Outras regras de aprendizado foram apresentadas posteriormente e muitas são generalizações da regra de Hebb [2, 32].

Cerca de 15 anos após a publicação do clássico artigo de McCulloch e Pitts, uma nova abordagem para o problema de reconhecimento de padrões foi introduzida por Rosenblatt [82] em seu trabalho sobre o *Perceptron*, um método inovador de aprendizagem supervisionada. O coroamento deste trabalho encontra-se no teorema da convergência do perceptron, cuja primeira demonstração foi delineada por Rosenblatt em 1960. Este teorema garante que o perceptron sempre converge para os pesos corretos se os pesos sinápticos que resolvem o problema existirem. Na mesma época, Widrow e Hoff introduziram o *algoritmo do mínimo quadrado médio* (LMS, Least Mean-Square) e usaram-no para formular o *Adaline* (Adaptive Linear Element, Elemento Linear Adaptativo) [107]. A diferença fundamental entre o perceptron e o Adaline está no procedimento de aprendizagem. Infelizmente existem limites fundamentais para aquilo que o perceptron de camada única e o Adaline podem calcular [58]. Rosenblatt e Widrow estavam cientes destas limitações e apresentaram redes neurais de múltiplas camadas que poderiam superar tais limitações, mas não conseguiram estender seus algoritmos de aprendizado para estas redes mais complexas. Uma destas redes é o *Madaline* (Multiple-Adaline), proposta por Widrow e seus estudantes, e uma das primeiras redes neurais em camadas de treinamento com múltiplos elementos adaptativos. Tais dificuldades e a falta de recursos tecnológicos nos anos 60 proporcionou um adormecimento nas redes neurais e poucos pesquisadores como James Anderson, Shunichi Amari, Leon Cooper, Kunihiko Fukushima, Stephen Grossberg e Teuvo Kohonen permaneceram no ramo.

Nos anos 80, a ausência de recursos tecnológicos foi superada e a pesquisa em redes neurais aumentou drasticamente. Computadores pessoais e estações de trabalho, que cresciam em capacidade, tornaram-se vitais para o desenvolvimento da pesquisa em redes neurais artificiais. Além disso, novos conceitos foram introduzidos. Dois deles tiveram grande influência no meio científico. O primeiro foi o uso da mecânica estatística para explicar as operações e convergência de algumas redes neurais recorrentes. Este conceito foi introduzido pelo físico John Hopfield em 1982 [40]. A segunda chave para o desenvolvimento na década de 80 foi o *algoritmo de retropropagação* (back-propagation), usado para treinar o perceptrons de múltiplas camadas. Este algoritmo foi descoberto por pesquisadores diferentes em diferentes pontos do mundo. Bryson talvez tenha sido o primeiro a estudar o algoritmo de retropropagação [13, 83]. Entretanto, a publicação mais influente foi o livro em dois volumes “Parallel Distributed Processing: Explorations in the Microstructures of Cognition”, editado por Rumelhart e McClelland. Este livro exerceu uma grande influência na utilização da aprendizagem por retropropagação, que emergiu como o algoritmo de aprendizagem mais popular para o treinamento de perceptrons de múltiplas camadas devido a sua simplicidade computacional e eficiência.

A história das redes neurais segue com muitos capítulos desafiantes e interessantes, mas ficaremos por aqui, pois acreditamos que você, leitor, já está bem situado no contexto histórico das redes neurais artificiais¹. Agora voltaremos aos anos 50 e falaremos sobre as memórias associativas neurais.

1.1.2 Memórias Associativas Neurais

Os estudos sobre memória associativa iniciaram nos anos 50 por Taylor [100]. Em 1961, Steinbruch introduziu o conceito de *matriz de aprendizagem* [88]. Nos anos seguintes surgiram as *memórias associativas holográficas* [17, 22, 105, 106] que não serão discutidas neste trabalho. Em 1972, Anderson [4], Kohonen [47] e Nakano [60] introduziram, de maneira independente, a idéia de uma memória por matriz de correlação, baseada na regra de aprendizagem por produto externo que pode ser interpretada como uma generalização do postulado de aprendizagem de Hebb. Nestes artigos, os autores apresentaram um modelo linear para os neurônios formando a *memória associativa linear* que estudaremos no capítulo 4.

Em 1977 foi publicado o livro “Associative Memory – A System Theoric Approach” de Kohonen [48]. Este foi um dos primeiros livros sobre memórias associativas e apresenta uma análise detalhada das memórias associativas lineares. No mesmo ano, Anderson et al. introduziram o *modelo do estado cerebral numa caixa* (BSB) [6, 5], uma das primeiras redes neurais que pode ser vista como uma memória auto-associativa dinâmica. Entretanto, o comportamento deste modelo como uma memória associativa só foi amplamente estudado após a publicação dos artigos de Hopfield. Além de ter aplicações como uma memória associativa, a BSB também pode ser vista como um modelo cognitivo [5].

Em 1982, Hopfield publicou o clássico artigo “Neural Networks and Physical Systems with Emergent Collective Computational Abilities” [40] onde introduz a famosa *rede* (ou *memória associativa*) de Hopfield discreta. Neste artigo, Hopfield sugere que um sistema dinâmico pode representar uma memória associativa dinâmica onde cada estado estável do sistema seria um padrão memorizado. Para compreender a computação executada e mostrar a convergência da rede, Hopfield introduziu o conceito de *função energia* (função de Lyapunov) para uma rede neural, um conceito que foi posteriormente usado para a análise de várias outras memórias associativas dinâmicas [23, 24, 41, 42]. Hopfield também apresentou alguns resultados sobre a *capacidade de armazenamento* da memória auto-associativa de Hopfield. Dois anos depois, Hopfield publicou um novo artigo apresentando a *rede* (ou *memória auto-associativa*) de Hopfield contínua, uma extensão do modelo apresentado anteriormente que utiliza uma função sigmóide como função de ativação [37]. Nos anos seguintes, Hopfield e Tank apresentaram aplicações da rede de Hopfield em problemas de otimização, como o clássico *problema do caixeiro viajante* [38, 39, 99].

Na década de 80 e início da década de 90, foram apresentadas variações do modelo de Hopfield. Entre elas, podemos citar a *memória associativa de Hopfield com armazenamento por projeção*. Este modelo foi inicialmente discutido por Personnaz et. al em 1985 [67], e posteriormente por Kanter e Sompolinsky em 1987 [44]. Em 1991, Chiueh e Goodman apresentaram a *memória associativa de capacidade exponencial* (ECAM), um modelo que abrange o modelo de Hopfield com armazenamento por correlação [15]. Em 1987, surgiram modelos de *memórias associativas dinâmicas para heteroassociação* [51, 62]. Entre estes modelos encontramos a *memória associativa bidirecional*

¹Para o leitor interessado em maiores detalhes da história das redes neurais, recomendamos os livros: [7, 28, 33]

(BAM), introduzida por Kosko, que pode ser vista como uma generalização da memória associativa de Hopfield [51, 52].

Na metade dos anos 90 iniciaram-se os estudos sobre as *memórias associativas morfológicas* (MAM), que usam a morfologia matemática e a álgebra de imagens como ferramenta matemática para descrever o neurônio. Antes de entrar nos detalhes do contexto histórico das memórias associativas morfológicas, vamos voltar novamente no tempo e apresentar brevemente a história da morfologia matemática e da álgebra de imagens, que são as ferramentas matemáticas usadas nestas memórias associativas.

1.1.3 Morfologia Matemática e Álgebra de Imagens

A *morfologia matemática* surgiu em 1964 enquanto Matheron e Serra estudavam a geometria de meios porosos e análise de textura. Esta nova ferramenta matemática está baseada nos trabalhos de Minkowski e Hadwiger sobre teoria de medida geométrica e geometria integral [57, 27]. Os primeiros anos, de 1964 a 1968, foram dedicados ao desenvolvimento de um corpo de notações teóricas e de um protótipo para a análise de textura. Deste período podemos citar o artigo “*Eléments pour une théorie des milieux poreux*”, de Matheron, onde aparece a primeira transformação morfológica para investigar a geometria dos objetos de uma imagem. Podemos citar também o trabalho em hardware especializado: “*Texture Analyser*” de J. Serra e J.-C. Kein. Neste período também foi criado o “*Centre de Morphologie Mathématique*” no campus da Escola de Minas de Paris, em Fontainebleu, França. Vários pesquisadores juntaram-se a este grupo e formaram o que chamamos hoje de “*Escola de Fontainebleu*”. Entre eles, podemos citar: Klein, Lantuéjoul, Meyer e Beucher [85, 87].

No ano 1975, Matheron publicou o livro “*Random Sets and Integral Geometry*”, que contém as primeiras fundamentações teóricas da morfologia matemática. Este livro foi bem aceito pelos interessados em geometria estocástica, mas infelizmente levou alguns anos para ser aceito pela comunidade de processamento de imagens. Podemos dizer que a morfologia matemática só passou a fazer parte das ferramentas para processamento de imagens após a publicação do clássico livro “*Image Analysis and Mathematical Morphology*” de Jean Serra, publicado em 1982. Segundo Heijmans, este livro pode ser visto como o primeiro tratamento sistemático da morfologia matemática como uma ferramenta para a análise de imagens [35]. Uma breve revisão sobre a morfologia matemática em tons de cinza, incluindo abordagens usando a teoria dos conjuntos nebulosos, pode ser encontrada em [96].

Na década seguinte houve uma explosão no número de artigos e pesquisadores trabalhando com a morfologia matemática e seria impossível listar todos eles. No Brasil encontramos os núcleos de pesquisa liderados por Barrera, na Universidade de São Paulo, e Banon, no Instituto Nacional de Pesquisas Espaciais em São José do Campos. Uma exposição detalhada contendo os trabalhos mais influentes da morfologia matemática encontra-se em [84]. Esta explosão não atingiu apenas a morfologia matemática, mas todas as atividades envolvendo processamento de imagens e resultou num excesso de técnicas, notações e operações para o processamento de imagens. Surgiu então a necessidade de uma estrutura algébrica padrão, eficiente e com um certo rigor matemático, designado especificamente para o processamento de imagens. Em resposta a esta situação, pesquisadores da Universidade da Flórida desenvolveram uma estrutura matemática para análise e processamento de imagens conhecida como *Álgebra de Imagens* (Image Algebra). Muitas são as vantagens da álgebra de imagens, mas não vamos listá-las aqui. Vamos dizer apenas que ela abrange todas as áreas de processamento de imagens e visão computacional usando uma linguagem comum. Para o leitor inte-

ressado, recomendamos [69, 79, 81].

A álgebra de imagens nasceu na metade da década de 80 e os primeiros artigos apareceram em 1987 [72, 78, 80]. Em 1990, Ritter, Wilson e Davidson publicaram o artigo: “Image Algebra: An Overview” [81]. Este artigo descreve as estruturas algébricas da álgebra de imagens e teve um papel importante na divulgação desta teoria matemática. Nos anos seguintes, a álgebra de imagens começou a ganhar aplicações. As redes neurais e as operações da morfologia matemática foram descritas e estudadas usando elementos e operações da álgebra de imagens [20, 71, 69]. Sendo descritas pela mesma estrutura matemática, foi fácil combiná-las, criando as *redes neurais morfológicas*, antes conhecidas como *redes da álgebra de imagens* (IA networks).

Uma rede neural morfológica é uma rede neural onde os neurônios são descritos pelas operações da morfologia matemática. Estas operações podem ser formuladas usando a álgebra de imagens, mas esta não é a única formulação matemática de uma rede neural morfológica. Uma formulação diferente, que não será discutida aqui, pode ser encontrada em [108]. Entre as redes neurais morfológicas, encontramos as *memórias associativas morfológicas* (Morphological Associative Memory, MAM). O estudo das memórias associativas morfológicas é recente, a primeira publicação apareceu em 1996 [73]. Outras publicações vieram depois [74, 76, 91, 92, 93, 94, 97], mas ainda são poucas e com certeza existem muitos modelos e teorias para serem descobertas e estudadas. Uma nova classe de memórias associativas neurais baseadas na teoria dos conjuntos nebulosos conhecida como “*Memórias Associativas Nebulosas Implicativas*” (Implicative Fuzzy Associative Memories, IFAM) [103, 98, 102]. As memórias associativas nebulosas implicativas generalizam as memórias associativas morfológicas quando a última é aplicada à padrões nebulosos. As memórias associativas nebulosas implicativas, por sua vez, podem ser vistas como um caso particular das memórias associativas morfológicas nebulosas que serão discutidas em trabalhos futuros. Neste trabalho não discutiremos as memórias associativas nebulosas implicativas nem as memórias associativas morfológicas nebulosas.

1.2 Objetivos e Organização da Dissertação

Muitos modelos de memória associativa neural foram apresentados nos últimos anos, entretanto, não encontramos na literatura um trabalho reunindo e comparando os modelos mais influentes. Esta dissertação de mestrado tem como objetivo principal discutir os principais modelos de memória associativa neural, incluindo as memórias associativas morfológicas. Os modelos de memória associativa neural mais influentes são essencialmente modelos binários. Por esta razão, esta dissertação é dedicada principalmente às memórias associativas neurais binárias.

Também não existe na literatura um critério comum de comparação para o desempenho de memórias associativas neurais. Esta dissertação de mestrado também tem como objetivo formalizar conceitos usados empiricamente na literatura para a comparação de modelos de memórias associativas. Formalizados os conceitos, usamos estes para comparar os modelos mais influentes de memória associativa binária.

Resumindo, os objetivos desta dissertação de mestrado são:

1. Reunir os modelos mais influentes de memória associativa neural (binárias) incluindo as memórias associativas morfológicas.
2. Formalizar critérios para comparação do desempenho das memórias associativas neurais binárias.

3. Comparar os modelos mais influentes com base nos critérios discutidos no item anterior.

Esta dissertação de mestrado está dividida em 8 capítulos. Os capítulos 2 e 3 tratam dos conceitos básicos de redes neurais e memórias associativas neurais. Especificamente, no capítulo 2, discutimos o conceito de rede neural artificial e como classificá-la. Apresentamos alguns modelos neurais e falamos brevemente sobre regras de aprendizado. No capítulo 3, apresentamos uma formulação matemática para o problema das memórias associativas, como classificá-las e apresentamos uma lista contendo as características desejáveis para o bom desempenho de uma memória associativa neural.

Os capítulos 4, 5, e 6 tem como objetivo apresentar os modelos mais influentes de memória associativa neural incluindo as memórias associativas neurais morfológicas, isto é, estes três capítulos cobrem o primeiro item dos objetivos desta dissertação. Precisamente, no capítulo 4, discutimos as memórias associativas lineares. Estes são os modelos mais simples de memória associativa neural e possuem um papel fundamental: definem as principais regras para armazenamento de padrões numa memória associativa neural. No capítulo 5 examinamos vários modelos de memórias associativas dinâmicas para auto e heteroassociação. Introduzimos neste capítulo a *Memória Associativa Bidirecional com Capacidade Exponencial* que é uma extensão da *Memória Associativa com Capacidade Exponencial* para o caso hetero-associativo. No capítulo 6 discutimos as memórias associativas morfológicas.

O capítulo 7 cobre os objetivos listados nos itens 2 e 3 acima. Neste capítulo formalizamos alguns conceitos para comparação das memórias associativas neurais binárias. Neste capítulo também apresentamos resultados teóricos e empíricos para a comparação dos modelos mais influentes de memória associativa neural. Terminamos a dissertação no capítulo 8 com a conclusão.

Foram realizados vários experimentos computacionais nesta dissertação de mestrado. Todos eles foram conduzidos no software MATLAB. As implementações dos modelos de memória associativa apresentadas nos capítulos 5 e 6, bem como as rotinas usadas nos experimentos computacionais do capítulo 7, podem ser obtidas com o autor em sua página pessoal [104].

As imagens usadas nesta dissertação de mestrado podem ser encontradas na página eletrônica do grupo CVG (*Computer Vision Group*) [1]. As imagens originais foram convertidas em imagens menores de dimensão 64×64 pixels usando o comando `imresize` do MATLAB com o método de interpolação padrão. Depois transformamos cada imagem num vetor coluna com 4096 componentes usando o comando `reshape` do MATLAB. As imagens binárias apresentadas na figuras 1.1 e 1.2 serão usadas com frequência nos exemplos computacionais com imagens binárias. Estas imagens foram obtidas aplicando um *threshold* nas imagens com tamanho reduzido, antes de serem transformadas num vetor coluna. Os níveis dos thresholds foram calculados usando o método de Otsu que minimiza a variância entre os valores preto e branco dos pixels [63] e pode ser obtido usando o comando `graythresh` no MATLAB. Nesta dissertação, o valor 1 representa o preto (objeto) e o valor 0 representa o branco numa imagem binária. As imagens apresentadas na figura 1.3 serão usadas nos exemplos computacionais com imagens em tons de cinza.

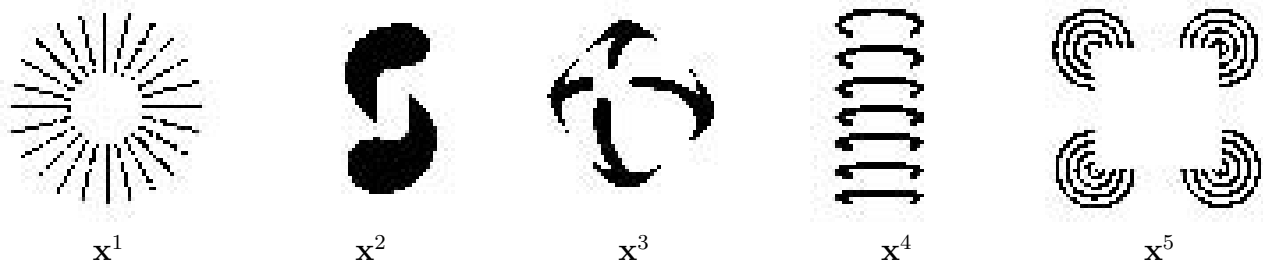


Fig. 1.1: Padrões $x^1, x^2, \dots, x^5$ usados nos exemplos computacionais com imagens binárias.

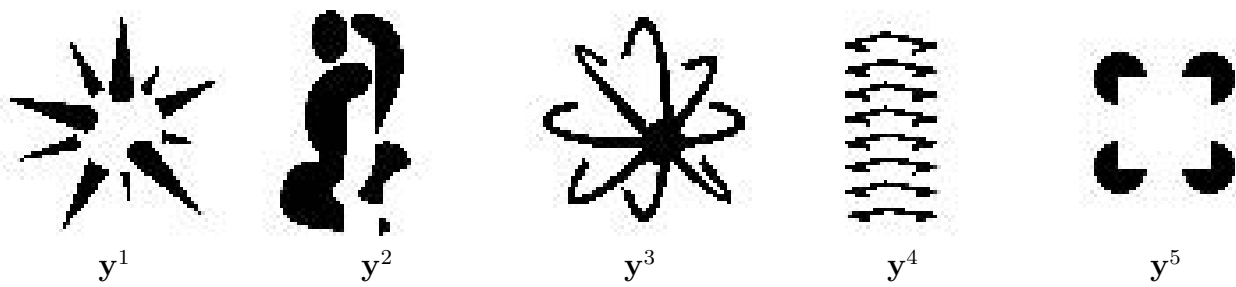


Fig. 1.2: Padrões $y^1, y^2, \dots, y^5$ usados nos exemplos computacionais com imagens binárias.

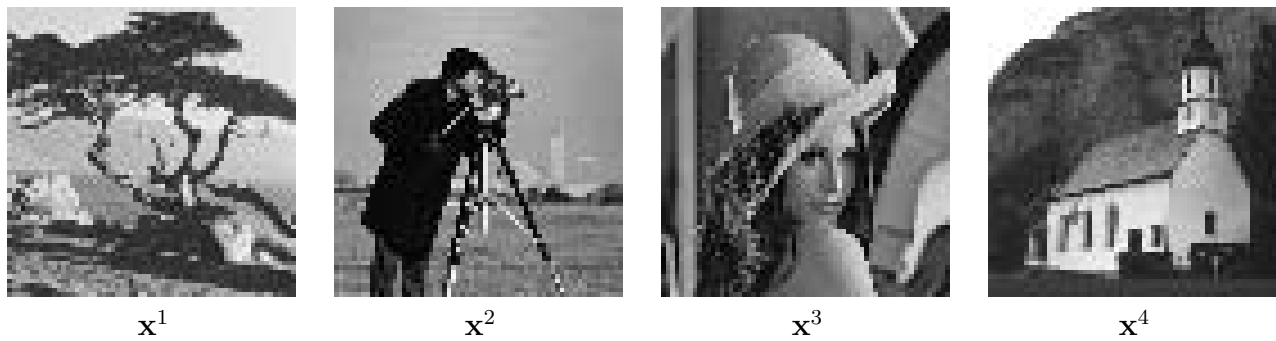


Fig. 1.3: Padrões x^1, x^2, x^3, x^4 usados nos exemplos computacionais com imagens em tons de cinza.

Capítulo 2

Conceitos Básicos de Redes Neurais

Neste capítulo discutimos os conceitos básicos de redes neurais artificiais; a notação e a nomenclatura que será usada durante toda a dissertação.

2.1 Introdução

Podemos dizer que os computadores modernos são retardatários no mundo da computação, pois os computadores biológicos – o cérebro e o sistema nervoso animal e humano – existem por milhões de anos e são extremamente eficientes no processamento de informações sensoriais e no controle da interação entre o animal e o meio em que vive. Tarefas como procurar um sanduíche, reconhecer uma face ou relembrar coisas são operações simples, como somas e multiplicações, para o nosso cérebro.

O fato dos computadores biológicos serem tão efetivos sugere que possamos extrair características similares a partir de um modelo do sistema neural. O modelo matemático que descreve o sistema nervoso biológico é conhecido como rede neural artificial ou simplesmente *rede neural*¹ [33]. Apesar de nossos modelos serem apenas metáforas do sistema nervoso biológico, eles fornecem um modo elegante e diferente de compreendermos o funcionamento de máquinas computacionais, além de oferecer informações úteis sobre o funcionamento do nosso cérebro.

Uma rede neural é caracterizada por três fatores:

1. Modelos (ou características) neurais,
2. Arquitetura (ou topologia) da rede,
3. Regra de aprendizado.

Nas próximas seções discutiremos estes fatores.

2.2 Modelos Neurais

Os neurônios, ou células nervosas, são os elementos computacionais usados pelo sistema nervoso. Podemos identificar três elementos básicos no modelo neural:

¹As redes neurais também são referidas na literatura como neurocomputadores, redes conexionistas ou processadores paralelamente distribuídos.

1. Um conjunto de *sinapses* (ou *elos de conexões*), cada uma caracterizada por um peso. Especificamente, um sinal x_j , na entrada da sinapse j do neurônio i , interage com o peso sináptico w_{ij} . Note que o primeiro índice de w refere-se ao neurônio em questão e o segundo refere-se ao terminal de entrada da sinapse.
2. Uma *regra de propagação*, que define as operações usadas para processar os sinais de entrada, ponderados pelas respectivas sinapses. Normalmente estas operações constituem um *combinador linear*. As operações usadas no modelo clássico é a *multiplicação* seguida da *soma*; entretanto, no modelo morfológico do neurônio, usaremos a *soma* e uma operação de *máximo* ou *mínimo*.
3. Uma *função de ativação*, usada para introduzir não-linearidade no modelo e/ou restringir a amplitude da saída de um neurônio. Neste trabalho usaremos basicamente quatro tipos de funções de ativação: identidade, função linear por partes, função sinal ou limiar, e a função exponencial. Deixaremos a função de ativação implícita quando usarmos a identidade.

É comum encontramos um bias nos modelos neurais artificiais. O bias pode ser visto como uma sinapse (w_{i0}) conectada a uma entrada constante (x_0). Por efeito de simplicidade, não acrescentaremos o bias nos nossos modelos neurais.

Encontramos vários modelos neurais na literatura [55, 73], como por exemplo os modelos neurais nebulosos [65, 103, 98, 102]. Neste trabalho, discutiremos somente o *modelo neural clássico* e o *modelo neural morfológico*. Apresentaremos a seguir cada um destes modelos.

2.2.1 Modelo Neural Clássico

O modelo neural clássico foi proposto por McCulloch e Pitts [55]. Neste modelo, um neurônio i é descrito pelo par de equações:

$$v_i = \sum_{j=1}^n w_{ij}x_j, \quad \text{e} \quad y_i = \phi(v_i), \quad (2.1)$$

ou pela única equação

$$y_i = \phi\left(\sum_{j=1}^n w_{ij}x_j\right), \quad (2.2)$$

onde $x_1, x_2, \dots, x_n$ são os sinais de entrada, $w_{i1}, w_{i2}, \dots, w_{in}$ são os pesos sinápticos do neurônio i , $\phi(\cdot)$ é a função de ativação, v_i é a ativação do neurônio e y_i é o sinal de saída. Na figura 2.1 temos uma representação simbólica de um neurônio clássico com n entradas.

Um conjunto com m neurônios em paralelo poder ser escrito na forma matricial através da equação

$$\mathbf{y} = \phi(W\mathbf{x}), \quad (2.3)$$

onde $\mathbf{x}$ é o vetor coluna contendo os sinais de entrada, $W \in \mathbb{R}^{m \times n}$ representa a matriz de pesos sinápticos, ϕ é uma função vetorial com componentes $\phi_i : \mathbb{R} \rightarrow \mathbb{R}$ e $\mathbf{y}$ é o vetor coluna contendo a saída da rede. Note que cada linha da equação (2.3) representa um neurônio descrito por (2.2).

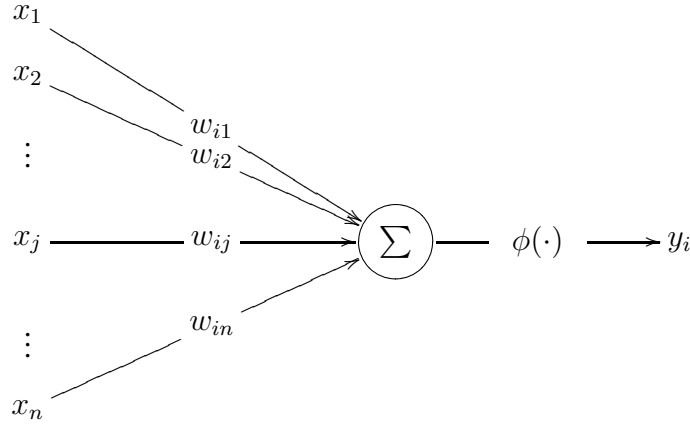


Fig. 2.1: Representação do modelo matemático clássico de um neurônio artificial.

2.2.2 Modelo Neural Morfológico

O modelo neural morfológico foi proposto por Davidson e Ritter no início dos anos 90 [20, 73, 74, 76]. Neste modelo, a soma é substituída pelo máximo ou pelo mínimo e a multiplicação é substituída pela soma. Uma motivação biológica para esta substituição é fornecida em [77]. Matematicamente, um neurônio morfológico i é descrito pela equação

$$v_i = (w_{i1} + x_1) \vee (w_{i2} + x_2) \vee \dots \vee (w_{in} + x_n) = \bigvee_{j=1}^n (w_{ij} + x_j), \quad (2.4)$$

ou

$$v_i = (w_{i1} + x_1) \wedge (w_{i2} + x_2) \wedge \dots \wedge (w_{in} + x_n) = \bigwedge_{j=1}^n (w_{ij} + x_j), \quad (2.5)$$

em conjunto com a equação

$$y_i = \phi(v_i). \quad (2.6)$$

Note que a equação (2.6) é idêntica a segunda equação de (2.1). Logo, a diferença entre o modelo clássico e o modelo morfológico está no cálculo da ativação do neurônio (v_i). Neste trabalho usaremos somente a função identidade e a função limiar como função de ativação no modelo neural morfológico.

O modelo neural descrito pelas equações (2.4) e (2.6), e o modelo descrito pelas equações (2.5) e (2.6) são conhecidos na literatura como *modelo neural morfológico* porque (2.4) e (2.5) representam as operações básicas da morfologia matemática: *dilatação* e *erosão*, respectivamente [79, 85, 87].

Um conjunto com m neurônios morfológicos em paralelo também pode ser escrito na forma matricial de um modo análogo à equação (2.3) do modelo clássico. Para tanto, precisamos definir um produto matricial em termos das operações de máximo ou mínimo, e da soma. Para uma matriz $A \in \mathbb{R}^{m \times p}$, e uma matriz $B \in \mathbb{R}^{p \times n}$, o produto matricial $C = A \boxplus B$, também conhecido como *produto máximo de A por B*, é definido por

$$c_{ij} = \bigvee_{k=1}^p (a_{ik} + b_{kj}). \quad (2.7)$$

O produto mínimo de A por B é definido de modo semelhante. Especificamente, os elementos de $C = A \boxtimes B$ são dados por

$$c_{ij} = \bigwedge_{k=1}^p (a_{ik} + b_{kj}). \quad (2.8)$$

Um conjunto com m neurônios morfológicos, expresso na forma matricial, é dado por:

$$\mathbf{y} = \phi(W \boxtimes \mathbf{x}), \quad (2.9)$$

ou

$$\mathbf{y} = \phi(W \boxdot \mathbf{x}). \quad (2.10)$$

As equações (2.9) e (2.10) generalizam o par de equações (2.4) e (2.6), e o par (2.5) e (2.6), respectivamente.

Os neurônios morfológicos descritos pela equação (2.9) estão baseado em operações da estrutura algébrica $(\mathbb{R}, \vee, +)$, conhecida como *belt* [18, 19]. Analogamente, os neurônios morfológicos descritos pela equação (2.10) estão baseados em operações no belt $(\mathbb{R}, \wedge, +)$. Os belts $(\mathbb{R}, \vee, +)$ e $(\mathbb{R}, \wedge, +)$ podem ser combinados formando a estrutura algébrica de grupo ordenado-reticulado (lattice-ordered group) $(\mathbb{R}, \vee, \wedge, +)$. Existe uma elegante dualidade em $(\mathbb{R}, \vee, \wedge, +)$ obtida a partir da equação $r \wedge u = -((-r) \vee (-u))$, válida para os número reais. Dado $r \in \mathbb{R}$, definimos o *conjugado aditivo* r^* como $r^* = -r$. Desta forma, $(r^*)^* = r$ e $r \wedge u = (r^* \vee u^*)^*$ para todo $r, u \in \mathbb{R}$. Se $A \in \mathbb{R}^{m \times n}$, então a *matriz conjugada* A^* de A é a dada por $A^* = -A^T$. Segue então que

$$A \wedge B = (A^* \vee B^*)^*, \quad (2.11)$$

e

$$A \boxdot B = (B^* \boxtimes A^*)^*, \quad (2.12)$$

para matrizes de tamanho apropriado. Consequentemente, uma rede neural morfológica formulada usando a operação $\boxtimes$ pode ser reformulada em termos da operação $\boxdot$, e vice-versa, usando a relação expressa na equação (2.12). Além disso, toda proposição em $(\mathbb{R}, \vee, \wedge, +)$ induz uma proposição dual obtida substituindo o símbolo $\wedge$ por $\vee$ e vice-versa, e revertendo as desigualdades.

Note que o modelo neural morfológico não envolve operações de multiplicação, mas apenas operações como máximo ou mínimo e soma. Temos então um modelo neural com computações rápidas e de fácil implementação em hardware. Problemas de convergência e longos algoritmos de treinamento praticamente não existem [75]. Além disso, as redes neurais morfológicas são capazes de resolver os problemas computacionais convencionais [73].

2.3 Arquiteturas de Redes Neurais

A *arquitetura* (ou topologia) de uma rede neural consiste na estrutura de interconexão dos seus neurônios que são geralmente organizados em camadas com um ou vários neurônios. Na literatura encontramos a seguinte classificação para as camadas de neurônios:

- *Camada de Entrada*: Camada de neurônios que introduz as entradas externas na rede. Os neurônios desta camada são chamados de *neurônios de entrada*.

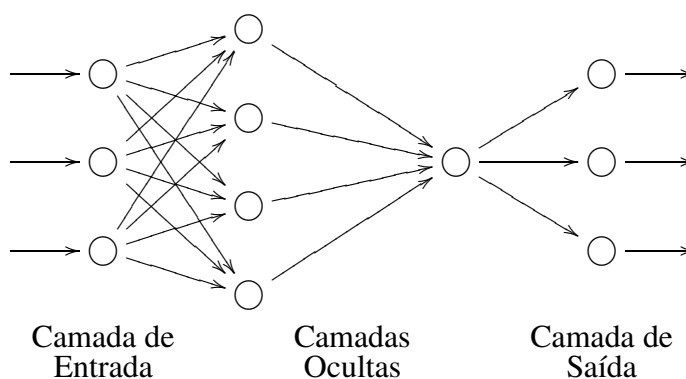


Fig. 2.2: Nomenclatura das camadas de uma rede neural. Rede neural progressiva totalmente conexa de múltiplas camadas.

- *Camada de Saída*: Camada de neurônios que produz a saída da rede. Os neurônios desta camada são chamados *neurônios de saída*.
- *Camada Oculta*: Camada de neurônios que interage com outras camadas da rede. Os neurônios desta camada são aqueles que não pertencem nem a camada de entrada nem a camada de saída da rede.

Na figura 2.2 apresentamos uma rede de múltiplas camadas com duas camadas ocultas, sendo uma delas composta por um único neurônio.

É comum caracterizarmos uma rede pelo número de camadas. A contagem das camadas é feita considerando apenas as camadas com pesos ajustáveis. Não contamos a camada de entrada, pois ela não tem pesos ajustáveis. Uma rede neural que não possui camadas ocultas é chamada *rede de camada única*, pois encontramos pesos ajustáveis somente na camada de saída. Na figura 2.3 apresentamos uma rede de camada única. Numa *rede de múltiplas camadas* encontramos uma ou mais camadas ocultas. A figura 2.4 ilustra como é feita a contagem das camadas.

Uma rede de múltiplas camadas é dita *progressiva*² quando as conexões sinápticas avançam para a saída da rede, ou seja, quando não houver laço de realimentação³ na rede. Apresentamos exemplos de redes progressivas nas figuras 2.2, 2.3 e 2.4. Chamaremos *rede recorrente* quando houver conexões entre neurônios de uma mesma camada e/ou conexões retropropagadas. Apresentamos na figura 2.5 uma rede recorrente com uma camada de neurônios ocultos.

Uma rede de múltiplas camadas é dita *totalmente conexa* quando cada neurônio de uma camada está conectado a todos os neurônios da camada seguinte. As figuras 2.2 e 2.3 apresentam redes totalmente conexas. Se alguns elos de comunicação (conexões sinápticas) estiverem faltando, diremos que a rede é *parcialmente conexa*. Na figura 2.4 encontramos uma rede parcialmente conexa. A mesma nomenclatura vale para uma rede recorrente de camada única. Na figura 2.6 apresentamos uma rede recorrente de camada única totalmente conexa. Neste caso, todos os neurônios possuem laço de alimentação de modo que a saída da iteração t é usada como entrada na iteração $t + 1$.

²Tradução para o termo inglês “feedforward”.

³Existe realimentação quando a saída de um elemento influencia em parte a entrada aplicada àquele elemento particular.

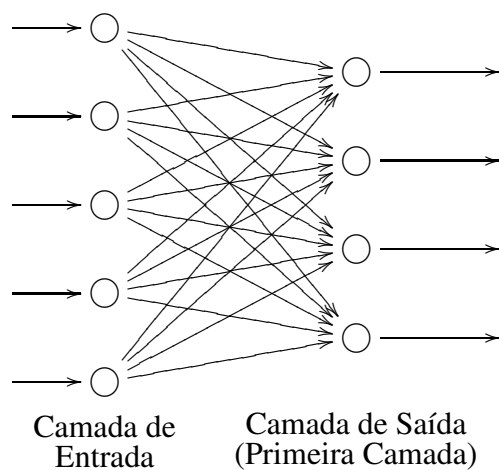


Fig. 2.3: Rede neural progressiva totalmente conexa de camada única.

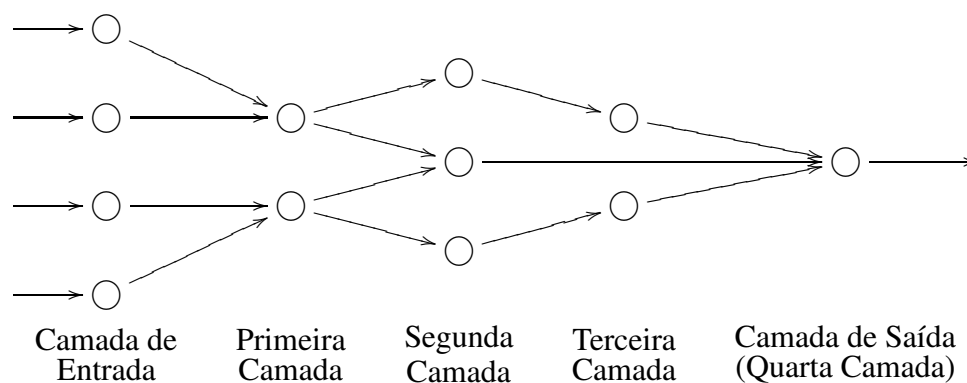


Fig. 2.4: Numeração das camadas de uma rede neural de múltiplas camadas. Rede neural progressiva parcialmente conexa com múltiplas camadas.

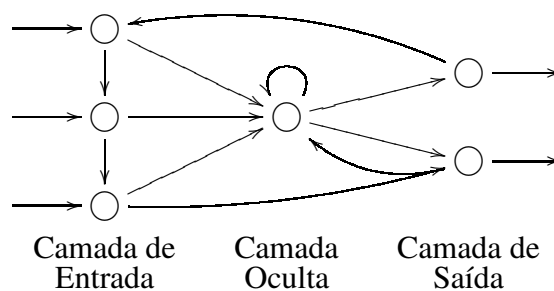


Fig. 2.5: Rede recorrente com uma camada oculta.

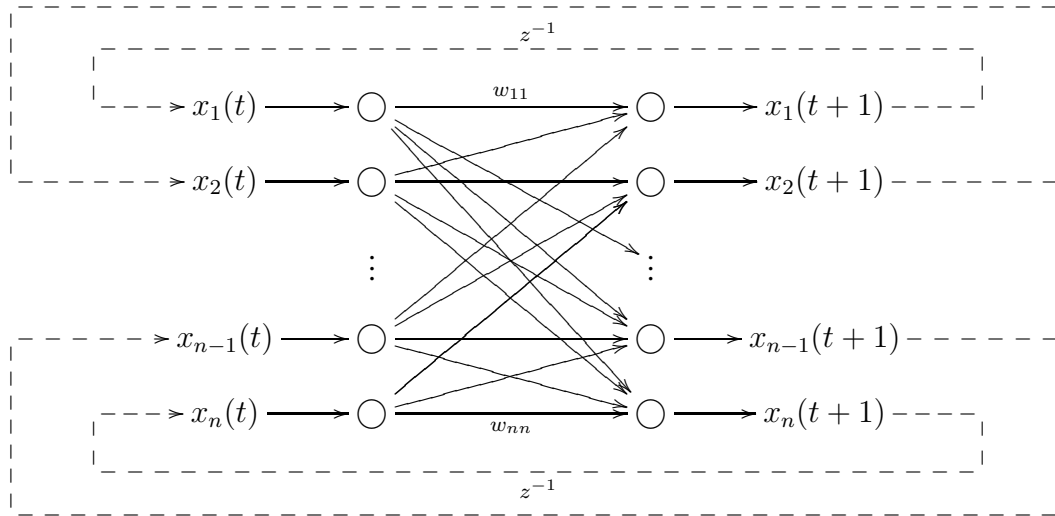


Fig. 2.6: Rede recorrente de camada única totalmente conexa.

2.4 Aprendizagem

A capacidade de aprender é uma das principais características da inteligência. O aprendizado numa rede neural é realizado ajustando-se os pesos das conexões sinápticas. Em outras palavras, aprendizagem é o processo onde os parâmetros livres de uma rede são modificados. O tipo de aprendizagem é determinado pela maneira pela qual ocorre a modificação dos parâmetros.

Existem dois tipos básicos de aprendizagem numa rede neural: *aprendizado supervisionado* (ou aprendizado com professor) e *aprendizado não-supervisionado* (ou aprendizado sem professor). Em ambos os casos precisamos de um conjunto de dados, conhecidos como *dados de treinamento*.

O *aprendizado supervisionado* consiste na apresentação de exemplos de entrada-saída. Durante o processo de aprendizado é feito um ajuste nos pesos de forma a minimizar a diferença (erro) entre a resposta da rede e a resposta desejada⁴. Temos assim a ação de um “professor” que apresenta a resposta correta indicando a ação ótima a ser realizada pela rede neural.

No *aprendizado não supervisionado*, apenas os dados de entrada são fornecidos. Neste caso, o aprendizado está baseado em agrupamentos de padrões. Os pesos são ajustados de modo que padrões semelhantes produzam a mesma saída.

Nos estudos das memórias associativas usaremos somente o aprendizado supervisionado, pois sempre teremos os dados de entrada e as respectivas saídas desejadas.

2.4.1 Aprendizado Supervisionado

Uma rede neural artificial com n neurônios na camada de entrada e m neurônios na camada de saída pode representar uma função $G : \mathbb{R}^n \longrightarrow \mathbb{R}^m$ [21]. Por exemplo, uma rede neural progressiva clássica de camada única pode ser escrita como

$$\mathbf{y} = G(\mathbf{x}) = \phi(W\mathbf{x}).$$

⁴A aprendizagem supervisionada pode ser vista como aprendizagem por correção de erro.

Lembre-se que a recíproca também é verdadeira visto que as redes neurais clássicas são aproximadores universais [33]. Temos que observar, porém, que falta conduzir pesquisas sobre a utilização de redes neurais morfológicas como aproximadores de funções.

No aprendizado supervisionado conhecemos o conjunto de vetores de entrada $\mathbf{x}^\xi$ e suas respectivas saídas $\mathbf{y}^\xi$, para $\xi = 1, 2, \dots, k$. Podemos então interpretar a rede neural como uma função dos pesos sinápticos e resolver um problema de otimização onde minimizamos o erro cometido pela rede neural. Por exemplo, para encontrar a matriz dos pesos sinápticos de uma rede neural clássica progressiva de camada única, resolvemos o problema:

$$\min_W \|\mathbf{y}^\xi - \phi(W\mathbf{x}^\xi)\|, \quad \text{para todo } \xi = 1, \dots, k.$$

Existem vários algoritmos de aprendizado supervisionado que resolvem um problema de otimização. O mais conhecido é o *algoritmo de retropropagação* (backpropagation) [21, 33]. Entretanto, devido ao elevado custo computacional, não utilizaremos nenhum algoritmo de otimização complexo para encontrar a matriz dos pesos sinápticos. Neste trabalho utilizaremos somente regras simples, como o armazenamento por correlação ou o armazenamento por projeção, para treinar nossas redes neurais (veja Capítulo 4).

Capítulo 3

Conceitos Básicos de Memórias Associativas Neurais

Neste capítulo apresentamos os fundamentos matemáticos das memórias associativas neurais e especificamos a notação e a nomenclatura que será usada durante toda a dissertação.

3.1 Formulação Matemática, Armazenamento e Associação

Uma *memória associativa* (Associative Memory, AM) representa um sistema de entrada-saída (input-output) que armazena vários pares de padrões $(\mathbf{x}, \mathbf{y})$, onde $\mathbf{x} \in \mathbb{R}^n$ e $\mathbf{y} \in \mathbb{R}^m$. Numa memória associativa criamos um mapeamento entre a entrada e a saída dado por $\mathbf{y} = G(\mathbf{x})$, onde $G : \mathbb{R}^n \rightarrow \mathbb{R}^m$ é o *mapeamento associativo* da memória. Cada par entrada-saída $(\mathbf{x}, \mathbf{y})$ armazenado na memória é dito uma *associação*. A entrada do sistema (vetor $\mathbf{x}$) é conhecido como *padrão-chave* (ou *memória-chave*) e a saída (vetor $\mathbf{y}$) é chamado *padrão recordado*.

A formulação matemática para um problema de memória associativa pode ser escrita como: Dado um conjunto finito de pares de entrada-saída $\{(\mathbf{x}^\xi, \mathbf{y}^\xi) : \xi = 1, 2, \dots, k\}$ a ser armazenado, nossa tarefa é encontrar um mapeamento que recupere cada um destes pares, isto é, determinar uma função G tal que $G(\mathbf{x}^\xi) = \mathbf{y}^\xi$, para todo $\xi = 1, \dots, k$ [31]. Além disso, desejamos que G tenha tolerância a ruído (capacidade de correção de erros). Assim, se $\tilde{\mathbf{x}}^\xi$ é uma versão ruidosa de $\mathbf{x}^\xi$, isto é, se $\tilde{\mathbf{x}}^\xi \neq \mathbf{x}^\xi$ e $d(\mathbf{x}^\xi, \tilde{\mathbf{x}}^\xi) < \delta$ com $\delta > 0$ pequeno, desejamos que $\mathbf{x}^\xi$ e $\tilde{\mathbf{x}}^\xi$ produzam a mesma saída, ou seja, $G(\tilde{\mathbf{x}}^\xi) = \mathbf{y}^\xi$ (a função G não é injetora). Nas memórias associativas neurais utilizamos uma rede neural para representar o mapeamento associativo G e a fase de armazenamento reduz-se a determinar a(s) matriz(es) dos pesos sinápticos.

Note que, se um par $(\mathbf{x}^\xi, \mathbf{y}^\xi)$ foi corretamente armazenado numa memória associativa e se G é um mapeamento ímpar, então $G(-\mathbf{x}) = -G(\mathbf{x}) = -\mathbf{y}$, ou seja, o par $(-\mathbf{x}^\xi, -\mathbf{y}^\xi)$ também foi armazenado na memória. Podemos mostrar que uma rede neural clássica com funções de ativação ímpares produz um mapeamento G ímpar. Este fato não vale para as redes neurais morfológicas devido as operações de máximo e mínimo usadas no modelo neural morfológico. Várias memórias associativas neurais clássicas utilizam funções de ativação ímpar e, conseqüentemente, representam mapeamentos ímpares. Como exemplo temos a rede de Hopfield, a BAM e a BSB.

O conjunto das associações $\{(\mathbf{x}^\xi, \mathbf{y}^\xi), \xi = 1, \dots, k\}$ é chamado *conjunto das memórias funda-*

mentais. Cada associação $(\mathbf{x}^\xi, \mathbf{y}^\xi)$ neste conjunto é uma *memória fundamental*, cada padrão $\mathbf{x}^\xi$ é uma *chave fundamental* e cada padrão $\mathbf{y}^\xi$ neste conjunto é uma *recordação fundamental*. Quando o conjunto das memórias fundamentais é da forma $\{(\mathbf{x}^\xi, \mathbf{x}^\xi), \xi = 1, 2, \dots, k\}$, dizemos que esta é uma *memória auto-associativa*. Neste caso particular, os termos memória fundamental, chave fundamental e recordação fundamental são sinônimos. No caso geral, quando $\mathbf{y}^\xi$ é diferente de $\mathbf{x}^\xi$, temos uma memória *heteroassociativa*.

O processo usado para determinar (ou sintetizar) uma memória associativa é conhecido como *fase de armazenamento*. Um dos principais objetivos numa memória associativa é criar um mapeamento com uma grande capacidade de armazenamento, isto é, uma vasta quantidade de memórias fundamentais podem ser armazenadas [33]. Um dos maiores problemas em uma memória associativa é a criação de associações que não fazem parte do conjunto das memórias fundamentais. Estas associações, armazenadas indevidamente, são as chamadas *memórias espúrias*.

Quando a fase de armazenamento está completa, inicia-se a *fase de recordação*. Aqui, uma memória pode ser testada para verificar se as memórias fundamentais foram corretamente armazenadas e a capacidade de correção de erro pode ser medida apresentando as chaves fundamentais corrompidas com vários tipos de ruídos e observando a saídas resultantes, i.e., comparamos a saída de uma entrada ruidosa com a saída desejada. O conjunto dos pontos $\mathbf{x} \in \mathbb{R}^n$ tais que $G(\mathbf{x}) = \mathbf{y}$ é chamado *região de recordação* do padrão $\mathbf{y}$.

3.2 Classificação das Memórias Associativas Neurais

As memórias associativas neurais podem ser divididas em duas grandes classes: as *memórias associativas estáticas* e as *memórias associativas dinâmicas* (Dynamic Associative Memory, DAM). As memórias associativas neurais estáticas podem ser descritas por uma rede neural progressiva. Uma rede neural recorrente usada como mapeamento associativo produz uma memória associativa dinâmica. Portanto, a arquitetura da rede neural utilizada (progressiva ou recorrente) define a arquitetura da memória associativa neural (estática ou dinâmica, respectivamente).

Os padrões armazenados numa memória associativa neural podem ser bipolares ($\{-1, 1\}$), binários ($\{0, 1\}$), discretos ($\mathbb{Z}$) ou contínuos ($\mathbb{R}$). Apresentaremos nesta dissertação modelos de memórias associativas neurais para padrões contínuos mas daremos ênfase às memórias associativas bipolares e binárias.

A fase de recordação de uma memória associativa dinâmica pode ser interpretada como um processo temporal que assume valores discretos ou contínuos. Deste modo, classificamos as memórias associativas dinâmicas como sendo *discretas* ou *contínuas no tempo*. Nesta dissertação não discutiremos as memórias associativas dinâmicas contínuas no tempo [33, 37, 36]. Na figura 3.1 apresentamos um diagrama com a classificação das memórias associativas neurais. Os modelos com caixas pontilhadas não serão discutidos neste trabalho.

As memórias associativas dinâmicas discretas no tempo podem ser descritas pelas equações

$$\mathbf{y}(t) = F(\mathbf{x}(t)), \quad (3.1)$$

$$\mathbf{x}(t+1) = H(\mathbf{y}(t)), \quad (3.2)$$

$$\forall t = 0, 1, \dots \quad (3.3)$$

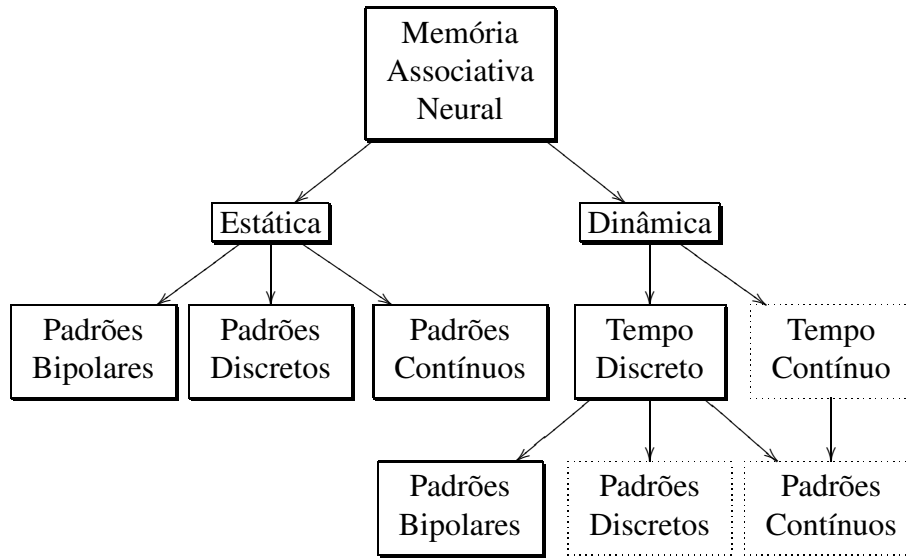


Fig. 3.1: Diagrama para classificação de uma memória associativa neural. Os modelos com caixas pontilhadas não serão discutidos neste trabalho.

onde $F : \mathbb{R}^n \rightarrow \mathbb{R}^m$ e $H : \mathbb{R}^m \rightarrow \mathbb{R}^n$ são funções não lineares. Dizemos que o par $(\mathbf{x}, \mathbf{y})$ é um *ponto estacionário* de uma memória associativa dinâmica se $F(\mathbf{x}) = \mathbf{y}$ e $H(\mathbf{y}) = \mathbf{x}$. Um ponto estacionário é também referido como *ponto fixo* no caso auto-associativo. O conjunto dos pontos $\mathbf{x} = \mathbf{x}(0)$ para o qual a sequência dos pares $\{(\mathbf{x}(0), \mathbf{y}(0)), (\mathbf{x}(1), \mathbf{y}(1)), \dots\}$ gerada através das equações (3.1)-(3.3) não converge é conhecido como *região de indecisão*. Note que uma memória associativa dinâmica sempre converge para um ponto estacionário independente do padrão-chave se e somente se a região de indecisão for vazia. Devemos impor um número máximo de iterações na ausência de informações sobre a região de indecisão ou quando soubermos que esta é não vazia.

A interpretação das equações (3.1) e (3.2) depende do tipo de atualização das componentes dos vetores $\mathbf{y}(t)$ e $\mathbf{x}(t+1)$. As duas formas de atualização mais comum são: *sincronizada* (ou *paralela*) e *assíncrona* (ou *seqüencial*). Numa memória associativa com atualização sincronizada, todas as componentes dos vetores $\mathbf{y}(t)$ e $\mathbf{x}(t+1)$ são atualizadas simultaneamente a cada iteração. Na atualização assíncrona, a cada iteração t , definimos I_t como sendo uma permutação do conjunto $\{1, 2, \dots, m\}$, J_t como sendo uma permutação de $\{1, 2, \dots, n\}$ e computamos

$$y_i(t) = F_i(\mathbf{x}(t)), \quad i \in I_t \quad (3.4)$$

$$x_j(t+1) = H_j(\mathbf{y}(t)), \quad j \in J_t \quad (3.5)$$

seguindo a ordem proposta em I_t e J_t . Em outras palavras, na iteração t , atualizamos uma única componente de cada vez nos vetores $\mathbf{y}(t)$ e $\mathbf{x}(t+1)$ seguindo uma seqüência aleatória de índices até atualizarmos todas as componentes. Na iteração seguinte, repetimos o processo escolhendo seqüências diferentes de índices I_t e J_t . Note que as atualização assíncrona adicionam incerteza na seqüência entre o padrão-chave e o padrão recordado. Por esta razão, podemos dizer que uma memória associativa dinâmica com atualização assíncrona representa um modelo estocástico.

O modo de atualização das componentes pode afetar drasticamente a fase de recordação de uma memória associativa dinâmica. Por exemplo, a memória associativa de Hopfield com atualização assíncrona sempre converge para um ponto fixo se certas condições forem satisfeitas. Por outro lado, a mesma memória associativa com atualização sincronizada pode ter uma região de indecisão não vazia. Nesta dissertação de mestrado, embora usamos com frequência uma notação vetorial semelhante às equações (3.1) e (3.2) para descrever uma memória associativa dinâmica, consideramos apenas atualização assíncrona.

O mapeamento associativo G de uma memória associativa dinâmica é definido como segue. Dado um padrão-chave $\mathbf{x}$, tome $\mathbf{x}(0) = \mathbf{x}$ e compute a seqüência finita $\{(\mathbf{x}(0), \mathbf{y}(0)), \dots, (\mathbf{x}(t_f), \mathbf{y}(t_f))\}$, onde t_f é, ou o menor t tal que $(\mathbf{x}(t), \mathbf{y}(t))$ é um ponto estacionário, ou o número máximo de iterações permitidas. O padrão recordado pela memória associativa após a fase de recordação é dado pela seguinte equação

$$\mathbf{y} = G(\mathbf{x}) := \mathbf{y}(t_f). \quad (3.6)$$

Note que o mapeamento associativo G inclui a dinâmica da memória associativa e considera, indiretamente, o modo de atualização das componentes.

3.3 Características para um Bom Desempenho

É de nosso interesse caracterizar o desempenho de uma memória associativa. Um conjunto de características desejáveis para uma classe de memórias associativas encontra-se em [31, 32, 64].

Uma memória associativa de baixa performance é aquela incapaz de armazenar todas as memórias fundamentais ou com baixa tolerância a ruído. Ela possui um grande número de memórias espúrias, e estas possuem grandes regiões de recordação. No caso das memórias associativas dinâmicas, uma baixa performance também pode ser caracterizada pela presença de oscilações, onde um estado inicial próximo a uma memória armazenada tem grande probabilidade de convergir para uma memória espúria ou para um ciclo limite.

Uma memória associativa dinâmica ideal é aquela que possui grande tolerância a ruído, um número relativamente pequeno de memórias espúrias, e cada memória espúria possui uma pequena região de recordação. No caso das memórias associativas dinâmicas, ela deve ser uma memória associativa estável no sentido de não possuir oscilações e possuir uma convergência rápida para quaisquer padrões-chaves apresentados à rede. Resumindo, diremos que uma memória associativa possui um bom desempenho quando possuir as seguintes características:

1. Grande capacidade de armazenamento,
2. Tolerância a ruído ou entradas incompletas,
3. Existência de poucas memórias espúrias,
4. O armazenamento da informação deve ser distribuído e robusto.
5. Recordação rápida e baixo custo computacional.

Utilizaremos estas características no capítulo 7 para comparar os vários modelos de memórias associativas neurais apresentados nesta dissertação.

Capítulo 4

Memórias Associativas Lineares

Neste capítulo apresentamos as memórias associativas lineares que foram introduzidas em 1972 independentemente por Anderson, Kohonen e Nakano [4, 47, 60]. Numa *memória associativa linear* (Linear Associative Memory, LAM) criamos um mapeamento linear $G : \mathbb{R}^n \rightarrow \mathbb{R}^m$ entre a entrada e a saída. Neste caso, o mapeamento G pode ser representado por uma matriz, digamos $W \in \mathbb{R}^{m \times n}$, e dado um padrão-chave $\mathbf{x} \in \mathbb{R}^n$, encontramos o padrão recordado $\mathbf{y} \in \mathbb{R}^m$ através da equação

$$\mathbf{y} = W\mathbf{x}. \quad (4.1)$$

Uma memória associativa linear é descrita por uma rede neural clássica progressiva de camada única (veja Figura 2.3 no Capítulo 2) com a função identidade como função de ativação. Discutimos dois procedimentos diferentes utilizados para obter a matriz W da equação (4.1). Precisamente, o *armazenamento por correlação* (Correlation Recipe) e o *armazenamento por projeção* (Projection Recipe) [31, 32]. Estes dois procedimentos definem a base para o aprendizado das memórias associativas discutidas nos próximos capítulos.

4.1 Armazenamento por Correlação

O *armazenamento por correlação* (Correlation Recording), também conhecido como *aprendizado de Hebb*, é um dos procedimentos mais usados para obter a matriz dos pesos sinápticos W e está baseado no postulado de aprendizagem de Hebb [34]. O postulado de Hebb afirma que se um neurônio A é ativado por um neurônio B repetidas vezes, então o neurônio A se tornará mais sensível aos estímulos do neurônio B e a conexão sináptica entre A e B será aumentada [21, 33]. Em resumo, o peso sináptico w_{ij} sofrerá uma variação dada pela correlação entre a entrada x_j e a saída y_i . Se temos um conjunto finito de pares $\{(\mathbf{x}^\xi, \mathbf{y}^\xi) : \xi = 1, \dots, k\}$ a ser armazenado numa memória associativa linear, o armazenamento por correlação fornecerá uma matriz $W \in \mathbb{R}^{m \times n}$ onde

$$w_{ij} = \sum_{\xi=1}^k y_i^\xi x_j^\xi. \quad (4.2)$$

Usando uma notação matricial temos

$$W = YX^T = \sum_{\xi=1}^k \mathbf{y}^\xi (\mathbf{x}^\xi)^T, \quad (4.3)$$



Fig. 4.1: Padrões recordados pela memória auto-associativa linear com armazenamento por correlação quando usamos como entrada as chaves fundamentais $\mathbf{x}^1, \dots, \mathbf{x}^5$.

onde $X = [\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^k] \in \mathbb{R}^{n \times k}$ é a matriz obtida tomando as chaves fundamentais como coluna e $Y = [\mathbf{y}^1, \mathbf{y}^2, \dots, \mathbf{y}^k] \in \mathbb{R}^{m \times k}$ é obtida concatenando as recordações fundamentais.

A memória associativa linear com armazenamento por correlação pode ser facilmente determinada, mas possui sérias limitações. Por exemplo, substituindo (4.3) em (4.1) e assumindo que $\mathbf{x}^h$ é uma chave fundamental, encontramos a seguinte expressão para o padrão recordado $\tilde{\mathbf{y}}^h$:

$$\tilde{\mathbf{y}}^h = \left[\sum_{\xi=1}^k \mathbf{y}^\xi (\mathbf{x}^\xi)^T \right] \mathbf{x}^h = \|\mathbf{x}^h\|^2 \mathbf{y}^h + \sum_{\xi \neq h}^k \mathbf{y}^\xi (\mathbf{x}^\xi)^T \mathbf{x}^h. \quad (4.4)$$

O segundo termo do lado direito de (4.4) é um vetor ruído e surge devido à *interferência cruzada* (cross-talk) entre o padrão $\mathbf{x}^h$ e as demais chaves fundamentais. Note que este termo será zero se os vetores $\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^k$ forem ortogonais. O primeiro termo do lado direito de (4.4) é proporcional à recordação fundamental $\mathbf{y}^h$, com constante de proporcionalidade $\|\mathbf{x}^h\|^2$. Para evitar esta constante, podemos impor $\|\mathbf{x}^\xi\|^2 = 1$ para $\xi = 1, 2, \dots, k$. Assim, uma condição suficiente para recordar uma memória perfeitamente é ter um conjunto de vetores $\{\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^k\}$ ortonormal. Dificilmente teremos uma recordação perfeita das memórias fundamentais se os padrões $\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^k$ não forem ortonormais. Note que não impomos condições sobre as recordações fundamentais, portanto, os vetores coluna $\mathbf{y}^1, \mathbf{y}^2, \dots, \mathbf{y}^k$ podem ser quaisquer.

Exemplo 4.1.1. Considere as imagens com 256 tons de cinza apresentadas na figura 1.3. Estas imagens foram convertidas em padrões (vetores coluna) $\mathbf{x}^1, \dots, \mathbf{x}^4 \in [0, 1]^{4096}$ e armazenadas na memória auto-associativa linear com armazenamento por correlação. Usando as chaves fundamentais como entrada, encontramos como resposta os padrões apresentados na figura 4.1, respectivamente. Neste exemplo percebemos claramente o efeito da interferência cruzada discutido anteriormente. O erro quadrático médio normalizado (EQMN) calculado através da equação

$$EQMN = \frac{1}{k} \sum_{\xi=1}^k \frac{\|W\mathbf{x}^\xi - \mathbf{y}^\xi\|}{\|\mathbf{x}^\xi\|} \quad (4.5)$$

foi aproximadamente 3100.

Exemplo 4.1.2. A memória associativa linear com armazenamento por correlação também pode ser usada para armazenar padrões bipolares. Considere as chaves fundamentais $\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^5 \in$



Fig. 4.2: Padrões recordados pela memória associativa linear com armazenamento por correlação quando usamos como entrada as chaves fundamentais $\mathbf{x}^1, \dots, \mathbf{x}^5$.

$\{-1, 1\}^{4096}$ apresentados na figura 1.1 e as recordações fundamentais $\mathbf{y}^1, \mathbf{y}^2, \dots, \mathbf{y}^5 \in \{-1, 1\}^{4096}$ apresentados na figura 1.2. Armazenamos estes cinco pares de padrões na memória associativa linear usando o armazenamento por correlação. Apresentando as chaves fundamentais $\mathbf{x}^1, \dots, \mathbf{x}^5$ como entrada, encontramos os padrões recordados apresentados na figura 4.2, respectivamente. Percebemos novamente o efeito da interferência cruzada neste exemplo. Note que, embora as recordações fundamentais sejam padrões bipolares, os padrões recordados não são padrões bipolares. De fato, $W\mathbf{x}^1, W\mathbf{x}^2, \dots, W\mathbf{x}^5 \in [-14354, 14354]^{4096}$. O erro quadrático médio normalizado (EQMN) foi aproximadamente 10471, um valor grande pois além da interferência cruzada temos o valor $\|\mathbf{x}^\xi\|^2 = 4096$ multiplicando o vetor coluna $\mathbf{y}^\xi$ na equação (4.4).

O armazenamento por correlação parece não ser muito eficiente, visto que dificilmente armazenará o conjunto das memórias fundamentais se existir uma chave fundamental que não é ortogonal as demais. Todavia, ele possui grandes vantagens. A primeira delas está na motivação biológica do postulado de Hebb. A segunda vantagem está no custo computacional realizado para encontrar a matriz W , pois neste armazenamento realizamos $(2k - 1)mn$ operações¹.

4.1.1 Armazenamento por Correlação Auto-associativo Bipolar

No caso auto-associativo bipolar, os elementos da diagonal de W são $w_{ii} = \sum_{\xi=1}^k (\mathbf{x}_i^\xi)^2$. Estes elementos são sempre positivos e seus valores aumentam indefinidamente quando adicionamos novos padrões. Deste modo, quando armazenamos muitos padrões, encontramos uma matriz cujos elementos da diagonal são muito maiores que os demais elementos e se fizéssemos uma normalização dos elementos da matriz, encontraríamos uma matriz parecida com uma matriz diagonal. Entretanto, uma matriz diagonal não possui capacidade de correção de erro e não teremos uma memória associativa eficiente. Este é um problema comum no aprendizado de Hebb e pode ser resolvido impondo $w_{ii} = 0$. Quando impomos $w_{ii} = 0$, encontramos:

$$w_{ij} = \begin{cases} \sum_{\xi=1}^k \mathbf{x}_i^\xi \mathbf{x}_j^\xi, & \text{se } i \neq j, \\ 0 & \text{se } i = j, \end{cases} \quad \text{para } 1 \leq i, j \leq n. \quad (4.6)$$

Usando uma notação matricial temos

$$W = XX^T - kI, \quad (4.7)$$

¹Neste trabalho $+$, $-$, $\times$ e $\div$ representam uma operação (1 flop).

onde I é a matriz identidade $n \times n$ e k é o número de padrões armazenados na memória associativa linear. A regra de aprendizado descrita pelas equações 4.6 e 4.7 é conhecido como *armazenamento por correlação com diagonal nula* ou *aprendizado de Hebb sem auto-conexão* (ou auto-realimentação). Este aprendizado será usado nas memórias auto-associativas dinâmicas apresentadas nos capítulo 5.

4.2 Armazenamento por Projeção

O *armazenamento por projeção* (Projection Recording) foi proposto por Kohonen e Ruohonen [50] e tem como objetivo resolver o problema

$$\min_W \|Y - WX\|_F^2 = \min_W \sum_{\xi=1}^k \|\mathbf{y}^\xi - W\mathbf{x}^\xi\|_2^2, \quad (4.8)$$

onde $\|\cdot\|_F$ representa a norma de Frobenius² e $\|\cdot\|_2$ representa a norma Euclidiana (para vetores). Uma memória associativa linear treinada usando o armazenamento por projeção é chamada *memória associativa linear ótima* (Optimal Linear Associative Memory, OLAM) pois a matriz dos pesos sinápticos $W \in \mathbb{R}^{m \times n}$ é a matriz que minimiza o erro entre as recordações fundamentais $\mathbf{y}^\xi$ e os padrões recordados $W\mathbf{x}^\xi$. A solução de (4.8) é

$$W = YX^\dagger, \quad (4.9)$$

onde $X^\dagger$ é a pseudo-inversa (ou inversa generalizada de Moore-Penrose) de X [25, 101]. A matriz W é a matriz de projeção no espaço gerado pelas recordações fundamentais $\mathbf{y}^\xi$, por isso chamamos este aprendizado de *armazenamento por projeção*.

Se o conjunto $\{\mathbf{x}^\xi, \xi = 1, 2, \dots, k\}$ é linearmente independente, i.e., se X é uma matriz de posto completo, então a pseudo-inversa de X é dada por

$$X^\dagger = (X^T X)^{-1} X^T, \quad (4.10)$$

e

$$W = Y(X^T X)^{-1} X^T. \quad (4.11)$$

Se os padrões $\{\mathbf{x}^\xi, \xi = 1, \dots, k\}$ forem ortonormais, então $X^T X = I$ e $W = YX^T$ é a matriz encontrada usando o armazenamento por correlação. Logo, o armazenamento por correlação é um caso particular do armazenamento por projeção se os padrões armazenados forem ortonormais.

Existe uma versão iterativa do procedimento de armazenamento por projeção baseada no teorema de Greville que não será discutida aqui mas pode ser encontrada em [49]. Esta versão iterativa pode ser usada para adicionar uma nova associação na memória associativa linear ótima.

No caso auto-associativo, $Y = X$, temos $W^2 = W$ (matriz de projeção) e $W = W^T$ (matriz simétrica). Estes fatos podem ser obtidos diretamente das propriedades da matriz pseudo-inversa [53, 48, 101]. Note que $W = I$ é solução do problema $\min_W \|X - WX\|_F$ independente da matriz X . Logo podemos armazenar um número ilimitado de padrões na OLAM no caso auto-associativo. No caso hetero-associativo, $\min_W \|Y - WX\|_F$ pode ser maior que zero e não garantimos sucesso

²Também referida como norma Euclidiana para matrizes.

na fase de armazenamento. Se X possui posto completo, pela equação (4.11) concluímos que a matriz W do armazenamento por projeção é tal que $WX = Y$. Por outro lado, se $k > n$, então certamente teremos um conjunto linearmente dependente, X provavelmente não terá posto completo e dificilmente conseguiremos armazenar o conjunto das memórias fundamentais na OLAM.

Exemplo 4.2.1. Considere o caso auto-associativo onde armazenamos os padrões em tons de cinza apresentados na figura 1.3. Vimos no exemplo 4.1 que a memória associativa linear com armazenamento por correlação não é capaz de armazenar o conjunto das memórias fundamentais. Armazenamos estas memórias fundamentais na OLAM e encontramos como saída as recordações fundamentais após apresentar as respectivas chaves fundamentais como entrada. De fato, todas as memórias fundamentais foram armazenadas com sucesso posto que o conjunto $\{\mathbf{x}^1, \dots, \mathbf{x}^4\}$ é linearmente independente. O EQMN deste experimento foi $2,7 \times 10^{-15}$ por causa de erros numéricos.

Exemplo 4.2.2. Vamos considerar o caso hetero-associative bipolar. Considere os padrões de entrada apresentados na figura 1.1 e os padrões de saída apresentados na figura 1.2. Primeiramente, notamos que o posto da matriz X é 5. Portanto, a OLAM deve ser capaz de armazenar o conjunto de memórias fundamentais $\{(\mathbf{x}^\xi, \mathbf{y}^\xi) : \xi = 1, \dots, 5\}$. De fato, armazenando este conjunto e apresentando os padrões $\mathbf{x}^1, \dots, \mathbf{x}^5$ como entrada, encontramos um EQMN de $5,6 \times 10^{-14}$. Novamente, o valor encontrado para EQMN é diferente de zero devido a erros de arredondamento.

Teorema 4.2.1. *Sejam $X \in \mathbb{R}^{n \times k}$ e $Y \in \mathbb{R}^{m \times k}$ as matrizes concatenando as chaves e recordações fundamentais, respectivamente. Considere a decomposição em valores singulares (SVD)*

$$X = U\Sigma V^T = \sum_{j=1}^r \sigma_j \mathbf{u}_j \mathbf{v}_j^T. \quad (4.12)$$

Seja $N = [\mathbf{n}^1, \mathbf{n}^2, \dots, \mathbf{n}^k] \in \mathbb{R}^{n \times k}$ uma matriz gerada aleatoriamente com distribuição gaussiana com média zero e variancia³ σ^2 , isto é, $E(\mathbf{n}_i^\xi) = 0$ e

$$E(\mathbf{n}_i^\xi \mathbf{n}_j^\xi) = \begin{cases} \sigma^2 & \text{se } i = j, \\ 0 & \text{caso contrário,} \end{cases} \quad (4.13)$$

onde $E(X)$ representa a esperança da variável aleatória X . Se $W = YX^\dagger$ é a matriz das conexões sinápticas da OLAM, então o erro quadrático médio (Mean Square Error, MSE) total das recordações da memória associativa será

$$MSE = E(\|Y - W(X + N)\|_F^2) \quad (4.14)$$

$$= \sum_{j=r+1}^k \|Y \mathbf{v}_j\|_2^2 + k\sigma^2 \sum_{j=1}^r \frac{\|Y \mathbf{v}_j\|_2^2}{\sigma_j^2}, \quad (4.15)$$

onde $\mathbf{v}_1, \dots, \mathbf{v}_k$ são os vetores singulares da direita e $\sigma_1, \dots, \sigma_r$ são os valores singulares de X .

Com base neste teorema concluímos que:

³Note que usamos σ_j , com índice, para representar os valores singulares e σ^2 , sem índice, para representar a variância de $\mathbf{n}_i^\xi$.

1. O primeiro termo de (4.15) surge devido a dependência linear das chaves fundamentais. Se os vetores $\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^k$ forem linearmente independentes, então o posto r da matriz X será k e o primeiro termo de (4.15) será zero. Se os padrões-chave forem linearmente dependentes, o erro do primeiro termo será inevitável e podemos chamar este termo de *erro da dependência linear*.
2. O segundo termo de (4.15) é obtido devido ao ruído nos padrões de entrada e podemos chamar este termo de *erro do ruído*. Quando apresentamos como entrada um padrão-chave sem ruído, então $\sigma^2 = 0$ e este termo também será zero. Por outro lado, se fornecermos uma entrada ruidosa, necessariamente teremos a interferência do erro do ruído na recordação. Note que, no somatório do erro do ruído, se um dos valores singulares for muito pequeno, então $k\sigma\|Yv_j\|^2/\sigma_j \gg 0$ e o erro de recordação será grande. Logo, quando uma entrada ruidosa é apresentada a memória associativa, o erro de recordação pode ser descrito pelos valores singulares σ_j de X . Observe também que, se X tem posto completo, então $r = k$ e quanto mais elementos são armazenados, maior será o erro devido ao termo do ruído.

Corolário. *O erro quadrático médio total das recordações da OLAM no caso auto-associativo é*

$$MSE = k\sigma^2 r. \quad (4.16)$$

Note que o MSE da memória associativa linear ótima no caso auto-associativo não possui o termo devido ao erro da dependência linear.

Sabemos que a média da soma total do ruído é

$$E\left(\sum_{\xi=1}^k \|\mathbf{n}^\xi\|_2^2\right) = k\sigma^2 n. \quad (4.17)$$

Com base nas equações (4.16) e (4.17), concluímos que a medida de correção de erro desta memória associativa é:

$$\frac{k\sigma^2 r}{k\sigma^2 n} = \frac{r}{n}. \quad (4.18)$$

Logo, se $r < n$, então a memória auto-associativa linear com armazenamento por projeção reduz o ruído da entrada. O pior caso ocorre quando $r = n$ onde o ruído não diminui. Note também que, quanto menor for r (ou k , pois $r \leq k$), melhor será a correção de erro esta memória. Entretanto, não temos uma restrição quanto ao número de padrões armazenados. Podemos armazenar um número de padrões maior que a dimensão dos padrões de entrada, isto é, $k > n$, mas perdemos com isso a capacidade de correção de erro, pois a matriz W tende para a matriz identidade.

O teorema 4.2.1 e o corolário deste teorema podem ser encontrados no artigo de Murakami e Aibara [59]. Outros resultados sobre a capacidade de correção de erro da memória associativa linear ótima (usando armazenamento por projeção) foram apresentados por Kohonen [49], Stiles e Denq [89] e Casasent e Telfer [14], entre outros.

Exemplo 4.2.3. Considere os padrões em tons de cinza apresentados na figura 1.3. Pelo teorema 4.2.1 e pelo exemplo 4.2.1, sabemos que a memória auto-associativa linear ótima é capaz de armazenar os padrões $\mathbf{x}^1, \dots, \mathbf{x}^4$. Vamos verificar agora a tolerância a ruído deste modelo. Na figura 4.3 apresentamos versões ruidosas $\tilde{\mathbf{x}}^1, \dots, \tilde{\mathbf{x}}^4$ das memórias fundamentais geradas com distribuição gaussiana

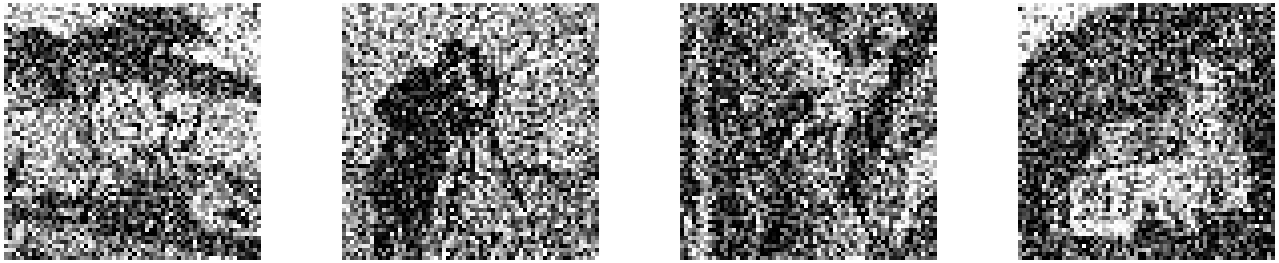


Fig. 4.3: Memórias-chave apresentada à OLAM no exemplo 4.2.3.

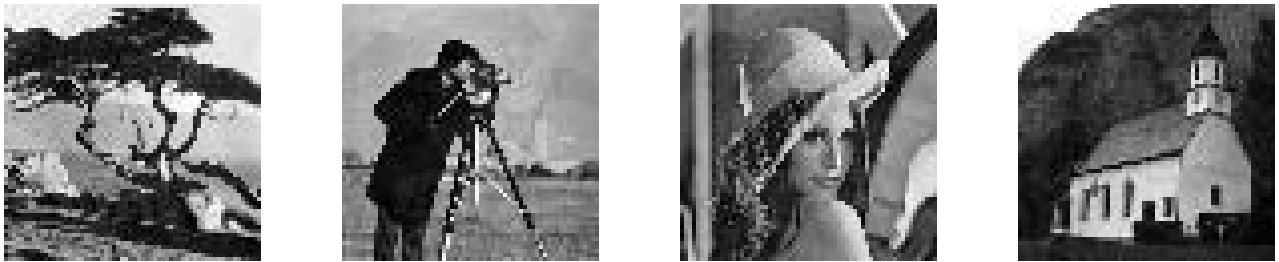


Fig. 4.4: Padrões recordados pela OLAM no exemplo 4.2.3 quando apresentamos as memórias-chave apresentadas na figura 4.3.

com média zero e variancia 0, 1. Na figura 4.4 apresentamos os padrões recordados após apresentar os padrões da figura 4.3 como entrada. O erro quadrático médio normalizado calculado após 1000 simulações foi aproximadamente 0,08.

Exemplo 4.2.4. Considere os padrões bipolares $\mathbf{x}^1, \dots, \mathbf{x}^5 \in \{-1, +1\}^{4096}$ e $\mathbf{y}^1, \dots, \mathbf{y}^5 \in \{-1, +1\}^{4096}$ apresentados nas figuras 1.1 e 1.2, respectivamente. Vimos no exemplo 4.2.2 que a OLAM é capaz de armazenar este conjunto de memórias fundamentais. Armazenamos estes padrões na OLAM e usamos como entrada os padrões $\tilde{\mathbf{x}}^1, \dots, \tilde{\mathbf{x}}^5$ apresentados na figura 4.5, respectivamente. Estes padrões ruidosos foram gerados a partir das chaves fundamentais revertendo o valor de um pixel seguindo uma distribuição uniforme com probabilidade 0,3. Na figura 4.6 apresentamos os respectivos padrões recordados. O erro quadrático médio normalizado (EQMN) foi aproximadamente 0,6. Note que os padrões recordados pela OLAM foram imagens em tons de cinza $[-1, +1]^{4096}$ com ruído vindo de outras recordações fundamentais mas dominado pelo saída desejada (erro do ruído). Encontramos as memórias fundamentais $\mathbf{y}^1, \dots, \mathbf{y}^5$ e um EQMN igual a zero aplicando um threshold com corte no nível de cinza fornecido pelo método de Otsu [63].



Fig. 4.5: Versões corrompidas dos padrões $x^1, \dots, x^5$ da figura 1.1. Estes padrões foram gerados introduzindo ruído uniforme com probabilidade 0,3 de reverter o valor de um pixel.

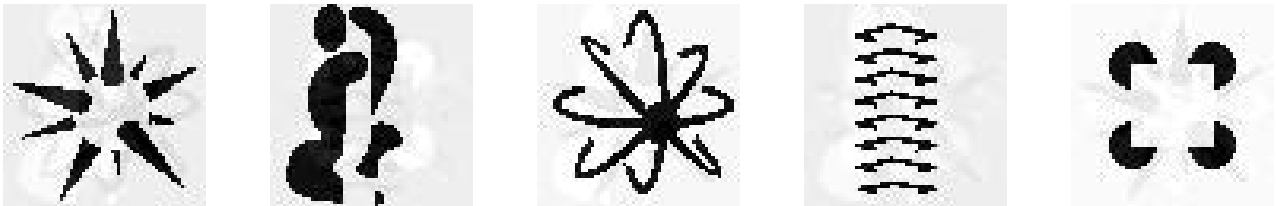


Fig. 4.6: Padrões recordados pela OLAM no exemplo 4.2.4 quando apresentamos as memórias-chave apresentadas na figura 4.5.

Capítulo 5

Memórias Associativas Dinâmicas

Neste capítulo apresentamos os principais modelos de memórias associativas dinâmicas para padrões bipolares e introduzimos a memória associativa bidirecional com capacidade exponencial. Para cada modelo fornecemos uma breve introdução, especificamos a arquitetura da rede neural e a regra de aprendizado, apresentamos uma breve análise sobre a convergência e exemplos computacionais.

5.1 Memória Associativa de Hopfield Discreta

A memória associativa de Hopfield foi introduzida em 1982 pelo físico J.J. Hopfield e é o modelo de memória associativa neural mais conhecido e estudado na literatura [40, 31, 33]. Este modelo forma a base para os demais modelos de memórias associativas discutidas neste capítulo.

5.1.1 Arquitetura da Rede

A memória associativa de Hopfield é descrita por uma rede neural clássica recorrente de camada única totalmente conexa com função sinal como função de ativação [40]. A arquitetura da rede de Hopfield está apresentada na figura 2.6. O modelo de Hopfield é descrito pela equação

$$x_i(t+1) = \text{sinal} \left(\sum_{j=1}^n w_{ij} x_j(t) \right) = \text{sinal} \left(\mathbf{w}_i^T \mathbf{x}(t) \right), \quad \text{para } i \in I_n \text{ e } t = 0, 1, \dots \quad (5.1)$$

onde $\mathbf{w}_i^T$ é a i -ésima linha da matriz W , $\mathbf{x}(t)$ é a entrada da rede no tempo t e I_n é uma permutação do conjunto de índices $\{1, 2, \dots, n\}$. A função sinal usada nas memórias associativas dinâmicas é definida como segue:

$$\text{sinal} \left(\mathbf{w}_i^T \mathbf{x}(t) \right) = \begin{cases} +1 & \text{se } \mathbf{w}_i^T \mathbf{x}(t) > 0, \\ x_i(t) & \text{se } \mathbf{w}_i^T \mathbf{x}(t) = 0, \\ -1 & \text{se } \mathbf{w}_i^T \mathbf{x}(t) < 0. \end{cases} \quad (5.2)$$

Podemos descrever a rede de Hopfield simplificadamente através da equação

$$\mathbf{x}(t+1) = \text{sinal}(W\mathbf{x}(t)), \quad t = 0, 1, \dots \quad (5.3)$$

Neste caso, devemos lembrar que a atualização das componentes é feita no modo assíncrono.

Pode-se usar atualização sincronizada na rede de Hopfield. Entretanto, a rede com atualização sincronizada pode apresentar ciclo limite, isto é, uma região de indecisão não nula [3]. Lembre-se que o tipo de atualização altera a trajetória dos pontos $\mathbf{x}(t)$, mas não altera os pontos fixos da memória associativa dinâmica.

5.1.2 Aprendizado

Na memória associativa de Hopfield usamos o armazenamento por correlação com diagonal nula. Neste caso, a matriz dos pesos sinápticos é computada através da equação 4.6, ou pela equação 4.7. Note que a matriz dos pesos sinápticos W é simétrica ($W = W^T$) com diagonal nula ($w_{ii} = 0, \forall i = 1, \dots, n$).

Repare na semelhança entre a memória associativa de Hopfield e a memória associativa linear com armazenamento por correlação discutida no capítulo 4. A diferença entre estas duas memórias associativas está na existência da função sinal, que força a saída a ser $+1$ ou -1 , e na recursividade presentes na rede de Hopfield. A seguir apresentaremos as propriedades da memória associativa de Hopfield e veremos como esta simples mudança (função sinal + recursividade) produz melhoras significativas no desempenho da memória associativa.

5.1.3 Convergência

O seguinte teorema, introduzido por Hopfield em [40], garante a convergência da memória associativa dinâmica descrita pelas equação 5.1 (ou 5.3). Em outras palavras, o seguinte teorema garante que a região de indecisão da memória associativa de Hopfield com atualização assíncrona é vazia. Outros resultados sobre a convergência desta memória associativa dinâmica podem ser encontrados em [9].

Teorema 5.1.1 (Teorema da Convergência de Hopfield). *Seja W uma matriz simétrica ($W = W^T$) com diagonal não negativa. ($w_{ii} \geq 0, i = 1, \dots, n$). A memória associativa dinâmica descrita pela equação 5.3 com atualização assíncrona converge para um ponto fixo e minimiza a função energia*

$$Energia(\mathbf{x}) = -\frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n w_{ij} x_i x_j = -\frac{1}{2} \mathbf{x}^T W \mathbf{x}. \quad (5.4)$$

O seguinte teorema fornece uma estimativa para o número máximo de memórias fundamentais que podem ser armazenadas na memória associativa de Hopfield. Este teorema foi introduzido por McEliece *et. al.* em [56]. Uma demonstração simplificada pode ser encontrada em [3].

Teorema 5.1.2 (Teorema de McEliece *et. al.*). *Sejam $\mathbf{x}^1, \dots, \mathbf{x}^k \in \{-1, 1\}^n$, com n suficientemente grande, padrões não-correlacionados gerados aleatoriamente com distribuição uniforme.*

1. *Uma memória fundamental $\mathbf{x}^\xi$ terá grande probabilidade de ser um ponto fixo da memória associativa de Hopfield se*

$$k \leq n/(2 \log n). \quad (5.5)$$

2. *Todos os padrões $\mathbf{x}^1, \dots, \mathbf{x}^k$ serão pontos fixos da memória associativa de Hopfield com grande probabilidade se*

$$k \leq n/(4 \log n). \quad (5.6)$$

O teorema 5.1.2 está baseado na convergência forte em probabilidade. Neste caso, dada a sequência de variáveis aleatórias $X_1, X_2, \dots$ e a variável aleatória X , dizemos que $X_n \rightarrow X$ quando $n \rightarrow \infty$ com grande probabilidade se $Pr(X_n \rightarrow X \text{ quando } n \rightarrow \infty) = 1$ [43].

Exemplo 5.1.1. Considere uma rede de Hopfield com 100 neurônios. Geramos k padrões bipolares aleatoriamente com distribuição uniforme e armazenamos todos eles na memória associativa de Hopfield. Depois verificamos se o primeiro padrão é um ponto fixo e se todos os padrões são pontos fixos. Repetimos o experimento 1000 vezes para diferentes valores de k (número de memórias fundamentais) e calculamos a probabilidade empírica do primeiro e de todas as recordações fundamentais serem pontos fixos da memória associativa de Hopfield. Na figura 5.1 apresentamos com linha marcada com $\circ$ a probabilidade empírica de uma certa memória fundamental (no nosso caso $\mathbf{x}^1$) ser um ponto fixo e com linha marcada com $\square$ a probabilidade empírica de todas as memórias fundamentais serem pontos fixos. A linha tracejada representa uma estimativa teórica para a capacidade de armazenamento obtida usando a equação 7.12 que será introduzida no capítulo 7. No eixo horizontal colocamos k , o número de memórias fundamentais. As linhas pontilhadas verticais indicam os valores $n/(2 \log n)$ e $n/(4 \log n)$. Note que a probabilidade empírica de todos os padrões serem pontos fixos deixa de ser 1 quando $k > n/(n \log n)$. A probabilidade empírica de todas as memórias fundamentais serem pontos fixos para $k = n/(2 \log(n))$ foi menor que 0,8. Entretanto, a probabilidade empírica de uma certa memória fundamental ser ponto fixo foi muito próxima de 1. Estes resultados numéricos conferem com o teorema 5.1.

Note que a matriz dos pesos sinápticos $W = XX^T$ fornecida pelo armazenamento por correlação com auto-alimentação também satisfaz as condições do teorema 5.1.1. Logo, este procedimento também pode ser usado para treinar a memória associativa descrita pela equação 5.3. O comportamento desta nova memória associativa dinâmica será muito parecido com o modelo com diagonal nula porque não podemos armazenar muitos padrões e, com poucos padrões armazenados, não temos uma matriz W com elementos na diagonal muito maior, em módulo, que os demais (veja seção 4.1.1 sobre armazenamento por correlação auto-associativo).

Exemplo 5.1.2. Considere os padrões bipolares apresentados na figura 1.1. Armazenamos estes padrões na memória associativa de Hopfield usando o armazenamento por correlação (com diagonal nula) e depois apresentamos as chaves fundamentais como entrada. A memória associativa de Hopfield encontrou os pontos fixos com no máximo 2 iterações. Na figura 5.2 apresentamos passo a passo os padrões recordados pela memória associativa de Hopfield até o final da segunda iteração. Note que $k = 5$ é muito menor que o valor fornecido pelo teorema 5.1.2, $(4096/(4 \log(4096))) = 283,5$. Entretanto, nenhuma das memórias fundamentais é um ponto fixo. Isso acontece porque os padrões $\mathbf{x}^1, \dots, \mathbf{x}^5$ possuem 2173 componentes em comum (no background) e uma correlação média $(1/20) \sum_{\xi \neq \eta} \langle \mathbf{x}^\xi, \mathbf{x}^\eta \rangle = 2344$. Um resultado similar será obtido pela memória associativa dinâmica descrita pela equação 5.3 treinada com o armazenamento por correlação com auto-alimentação.

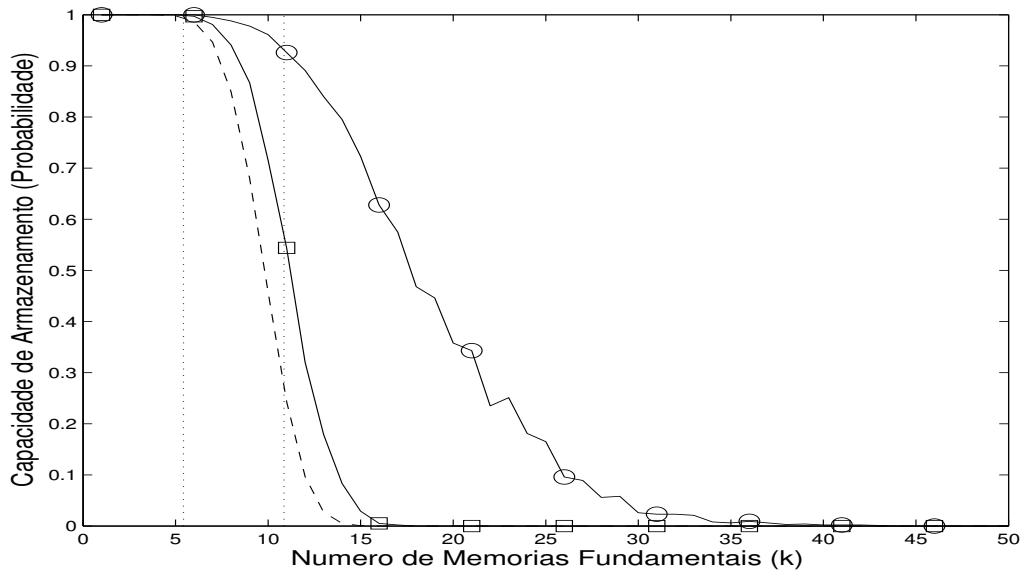
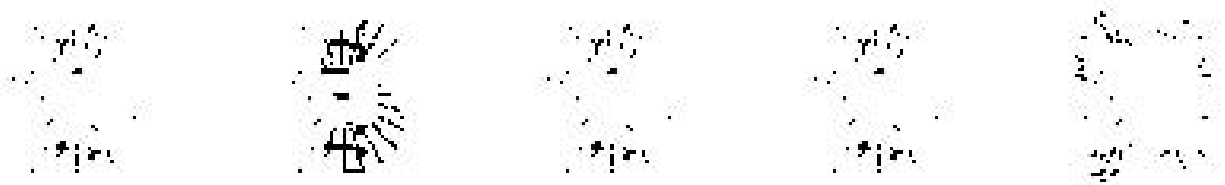
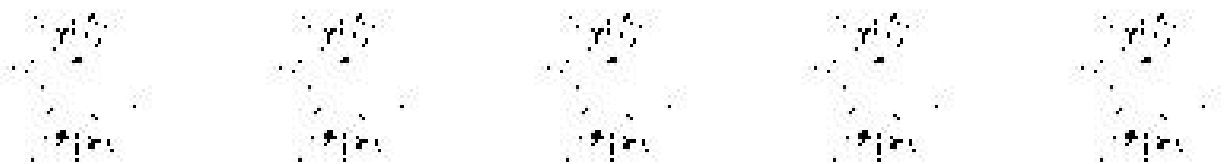


Fig. 5.1: Probabilidade das memórias fundamentais serem pontos fixos na memória associativa de Hopfield por k . A linha marcada com $\circ$ representa a probabilidade de uma dada memória fundamental ser ponto fixo e a linha marcada com $\square$ representa a probabilidade de todas as memórias fundamentais serem pontos fixos. A linha tracejada representa a capacidade de armazenamento que será discutida no capítulo 7. As linhas pontilhadas verticais representam os valores $n/(2 \log n)$ e $n/(4 \log n)$.



Padrões $\mathbf{x}^\xi(1) = \text{sinal}(W\mathbf{x}^\xi(0))$, $\xi = 1, \dots, 5$ obtidos no final da primeira iteração.



Pontos fixos $\mathbf{x}^\xi(2) = \text{sinal}(W\mathbf{x}^\xi(1))$, $\xi = 1, \dots, 5$ obtidos no final da segunda iteração.

Fig. 5.2: Padrões recordados pela memória associativa de Hopfield no exemplo 5.1.2 quando apresentamos as chaves fundamentais como entrada.

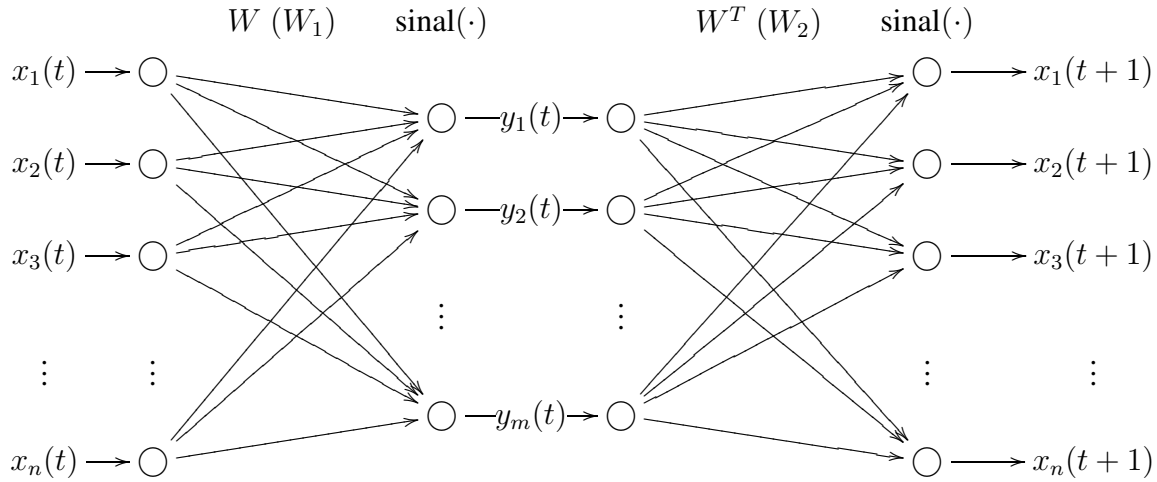


Fig. 5.3: Arquitetura da Memória Associativa Bidirecional (Assimétrica).

5.2 Memória Associativa Bidirecional

A *Memória Associativa Bidirecional* (Bidirectional Associative Memory, BAM), proposta por Kosko em 1987, é uma generalização da memória auto-associativa de Hopfield discreta com armazenamento por correlação com auto-alimentação [51, 52].

5.2.1 Arquitetura

A BAM é descrita por uma rede neural clássica recorrente totalmente conexa com duas camadas e função sinal como função de ativação, como apresentado na figura 5.3. A matriz dos pesos sinápticos da segunda camada da BAM é a transposta da matriz dos pesos sinápticos da primeira cada. A BAM é descrita pelo par de equações

$$\mathbf{y}(t) = \text{sinal}(W\mathbf{x}(t)), \quad (5.7)$$

$$\mathbf{x}(t+1) = \text{sinal}(W^T\mathbf{y}(t)), \quad \text{para } t = 0, 1, 2, \dots \quad (5.8)$$

com atualização assíncrona.

5.2.2 Aprendizado

As matrizes dos pesos sinápticos da BAM são obtidas usando o armazenamento por correlação (aprendizado de Hebb). A matriz dos pesos sinápticos da primeira camada é $W = YX^T$ e a matriz da segunda camada é $W^T = XY^T$. Note que W^T é a matriz das conexões sinápticas da memória associativa linear com armazenamento por correlação com as memórias fundamentais $\{(\mathbf{y}^\xi, \mathbf{x}^\xi), \xi = 1, 2, \dots, k\}$, onde $\mathbf{y}^\xi \in \{-1, 1\}^m$ é a entrada e $\mathbf{x}^\xi \in \{-1, 1\}^n$ é a saída, para $\xi = 1, \dots, k$. No caso auto-associativo, isto é, se $\mathbf{y}^\xi = \mathbf{x}^\xi$, para $\xi = 1, \dots, k$, teremos a memória associativa de Hopfield.

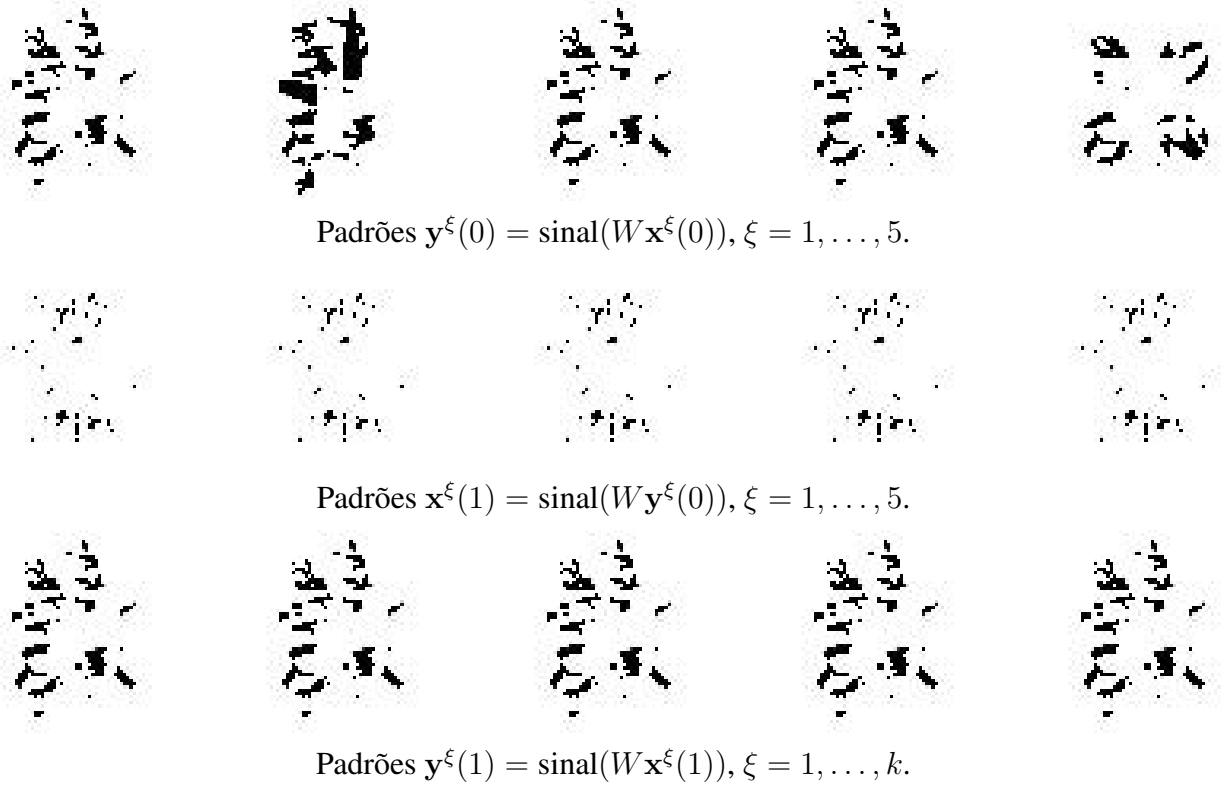


Fig. 5.4: Padrões recordados pela BAM quando apresentamos as chaves fundamentais $x^1, \dots, x^5$ como entrada.

5.2.3 Convergência

A BAM pode ser convertida numa memória auto-associativa discreta de Hopfield com vetores de estados $\mathbf{x}^T = [\mathbf{x}^T, \mathbf{y}^T]^T$ e matriz das conexões sinápticas

$$W^T = \begin{bmatrix} 0 & W^T \\ W & 0 \end{bmatrix}. \quad (5.9)$$

A matriz W^T é simétrica com diagonal nula. Logo, pelo teorema da convergência de Hopfield, esta memória associativa sempre converge para um ponto estacionário com atualização assíncrona. A convergência da BAM, para ambos os casos, também pode ser verificada mostrando que a função energia

$$\text{Energia}(\mathbf{x}, \mathbf{y}) = -\frac{1}{2} (\mathbf{x}^T W \mathbf{y} + \mathbf{y}^T W^T \mathbf{x}) = -\mathbf{x}^T W \mathbf{y}, \quad (5.10)$$

é minimizada.

Exemplo 5.2.1. Considere os padrões $x^1, \dots, x^5 \in \{-1, 1\}^{4096}$ e $y^1, \dots, y^5 \in \{-1, 1\}^{4096}$ apresentados nas figuras 1.1 e 1.2. Armazenamos o conjunto das memórias fundamentais $\{(x^\xi, y^\xi), \xi = 1, \dots, 5\}$ e verificamos se as memórias fundamentais são pontos estacionários da BAM. Começamos apresentando as chaves fundamentais $x^1, \dots, x^5$ como entrada e encontramos como resposta os padrões apresentados na figura 5.4. Os padrões $x^1(2), \dots, x^5(2)$ obtidos no final da segunda iteração

são iguais aos padrões $\mathbf{x}^1(1), \dots, \mathbf{x}^5(1)$ entrados no final da primeira iteração. Neste exemplo a BAM convergiu com 2 iterações, entretanto nenhuma das memórias fundamentais é um ponto estacionário.

A seguinte conjectura, introduzida por nós nesta dissertação de mestrado, é uma extensão do teorema de McEliece *et. al.* para a BAM. Lembre-se que $X_n \rightarrow X$ quando $n \rightarrow \infty$ com grande probabilidade se $P(X_n \rightarrow X \text{ quando } n \rightarrow 1) = 1$ [43].

Conjectura 5.2.1. Sejam $\mathbf{x}^1, \dots, \mathbf{x}^k \in \{-1, 1\}^n$ e $\mathbf{y}^1, \dots, \mathbf{y}^k \in \{-1, 1\}^m$, n e m suficientemente grandes, padrões não-correlacionados gerados aleatoriamente com distribuição uniforme.

1. Uma memória fundamental $(\mathbf{x}^\xi, \mathbf{y}^\xi)$ terá uma grande probabilidade de ser um ponto estacionário da BAM se

$$k \leq \frac{1}{2} \min \left(\frac{n}{\log m}, \frac{m}{\log n} \right). \quad (5.11)$$

2. Todas as memórias fundamentais $(\mathbf{x}^1, \mathbf{y}^1), \dots, (\mathbf{x}^k, \mathbf{y}^k)$ serão pontos estacionários da BAM com grande probabilidade se

$$k \leq \frac{1}{4} \min \left(\frac{n}{\log m}, \frac{m}{\log n} \right). \quad (5.12)$$

Acreditamos que uma demonstração formal do resultado acima pode ser obtida fazendo as devidas modificações na demonstração do teorema principal (The Big Theorem) apresentada em [56]. Nosso objetivo nesta dissertação é fazer um estudo comparativo em memórias associativas. Por esta razão apresentamos apenas um esboço da demonstração desta conjectura. Uma demonstração formal requer detalhes que serão omitidos.

Seja $X = [\mathbf{x}^1, \dots, \mathbf{x}^k] \in \{-1, 1\}^{n \times k}$ e $Y = [\mathbf{y}^1, \dots, \mathbf{y}^k] \in \{-1, 1\}^{m \times k}$ onde x_j^ξ e y_i^η são variáveis aleatórias independentes com probabilidade 1/2 de ser +1 ou -1 para todo $j = 1, \dots, n$, $i = 1, \dots, m$ e $\xi, \eta = 1, \dots, k$. Sabemos que $y_i^\eta = \text{sin}(\mathbf{w}_i^T \mathbf{x}^\eta)$ para $\eta \in \{1, \dots, k\}$ se e somente se $y_i^\eta \mathbf{w}_i^T \mathbf{x}^\eta \geq 0$ para todo $i = 1, \dots, m$ e $\eta \in \{1, \dots, k\}$. A matriz dos pesos sinápticos da BAM é $W = YX^T$ e

$$\mathbf{w}_i^T \mathbf{x}^\eta = \sum_{j=1}^n \sum_{\xi=1}^k y_i^\xi x_j^\xi x_j^\eta. \quad (5.13)$$

Logo,

$$y_i^\eta (\mathbf{w}_i^T \mathbf{x}^\eta) = \sum_{\xi=1}^k \sum_{j=1}^n y_i^\eta y_i^\xi x_j^\xi x_j^\eta = n + \sum_{\xi \neq \eta}^k \sum_{j=1}^n y_i^\eta y_i^\xi x_j^\xi x_j^\eta = n + r, \quad (5.14)$$

onde r representa o *ruído*.

O termo do ruído pode ser positivo ou negativo e será nosso objetivo estimar o seu valor. Para tanto, vamos tomar

$$z_j = \sum_{\xi \neq \eta}^k y_i^\eta y_i^\xi x_j^\xi x_j^\eta. \quad (5.15)$$

Deste modo, o termo do ruído será

$$r = \sum_{j=1}^n z_j. \quad (5.16)$$

Sabemos que $y_i^\eta, y_i^\xi, x_j^\xi$ e x_j^η são variáveis aleatórias independentes e identicamente distribuídas (i.i.d.), portanto $y_i^\eta y_i^\xi x_j^\xi x_j^\eta = \pm 1$ com a mesma probabilidade. Sendo assim, podemos seguir a demonstração rigorosa de McEliece *et. al.* [56].

Após aplicar o teorema central do limite, encontramos

$$Pr [y_i^\eta = \text{sinal}(\mathbf{w}_i^T \mathbf{x}^\eta)] = \frac{1}{2} \left[1 + \text{erf} \left(\sqrt{\frac{n}{2k}} \right) \right], \quad (5.17)$$

para n, k suficientemente grandes com k/n pequeno. A função $\text{erf}(x)$ é a *função erro* dada por

$$\text{erf}(x) = \frac{2}{\sqrt{\pi}} \int_0^x e^{-t^2} dt. \quad (5.18)$$

Para valores grandes no argumento, a função erro pode ser aproximada por

$$\text{erf}(x) \approx 1 - \frac{1}{x\sqrt{\pi}} e^{-x^2}. \quad (5.19)$$

Finalmente,

$$Pr [\mathbf{y}^\eta = \text{sinal}(W\mathbf{x}^\eta)] = Pr [y_i^\eta = \text{sinal}(\mathbf{w}_i^T \mathbf{x}^\eta), \forall i = 1, \dots, m] \quad (5.20)$$

$$\approx 1 - m \sqrt{\frac{k}{2\pi n}} \exp\left(-\frac{n}{2k}\right), \quad (5.21)$$

pois as componentes de X e Y são variáveis aleatórias i.i.d. Dizemos que $f \approx g$ se $f(x)/g(x) \rightarrow 1$ quando $x \rightarrow \infty$. Seguindo o resultado apresentado em [3], teremos $\mathbf{y}^\eta = \text{sinal}(W\mathbf{x}^\eta)$ com probabilidade próxima de 1 se $k \leq n/(2 \log m)$. Analogamente, deveremos ter $k \leq m/(2 \log n)$ para termos $\mathbf{x}^\eta = \text{sinal}(W^T \mathbf{y}^\eta)$ com probabilidade próxima de 1. Logo, a memória fundamental $(\mathbf{x}^\eta, \mathbf{y}^\eta)$ será um ponto estacionário se

$$k \leq \frac{1}{2} \min \left(\frac{n}{\log m}, \frac{m}{\log n} \right). \quad (5.22)$$

Se queremos $\mathbf{y}^\xi = \text{sinal}(W\mathbf{x}^\xi)$, para todo $\xi = 1, \dots, k$, então deveremos ter

$$Pr(Y = \text{sinal}(WX)) = [Pr(y_i^\eta = \text{sinal}(\mathbf{w}_i^T \mathbf{x}^\eta))]^{mk}. \quad (5.23)$$

Esta probabilidade será próxima de 1 se $k \leq n/(4 \log m)$. Analogamente, $Pr(X = \text{sinal}(W^T Y))$ será próxima de 1 se $k \leq m/(4 \log n)$. Finalmente, todos as memórias fundamentais serão pontos fixos da BAM com probabilidade próxima de 1 se

$$k \leq \frac{1}{4} \min \left(\frac{n}{\log m}, \frac{m}{\log n} \right), \quad (5.24)$$

para n, m, k suficientemente grandes.

Note que a conjectura 5.2.1 coincide com o teorema 5.1 se $m = n$.

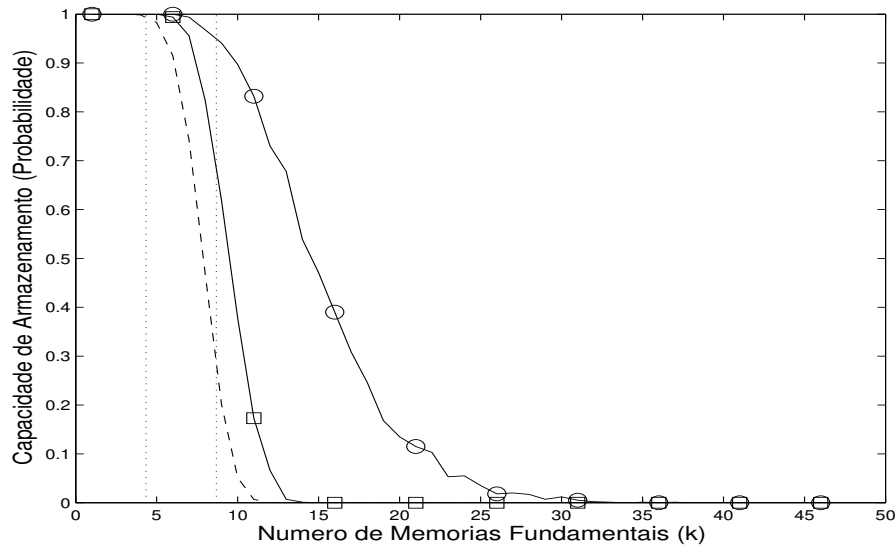


Fig. 5.5: Probabilidade das memórias fundamentais serem pontos fixos na BAM por k . A linha marcada com $\circ$ representa a probabilidade empírica de uma dada memória fundamental ser ponto fixo e a linha marcada com $\square$ representa a probabilidade empírica de todas as memórias fundamentais serem pontos fixos. A linha tracejada representa a capacidade de armazenamento obtida usando a equação 7.15 (probabilidade teórica). As linhas pontilhadas verticais representam os valores $\min(m/\log(n), n/\log(m))/2$ e $\min(m/\log(n), n/\log(m))/4$.

Exemplo 5.2.2. Considere o conjunto $\{(\mathbf{x}^\xi, \mathbf{y}^\xi), \mathbf{x}^\xi \in \{-1, 1\}^{100}, \mathbf{y}^\xi \in \{-1, 1\}^{80}, \xi = 1, \dots, k\}$ gerado aleatoriamente com distribuição uniforme. Armazenamos este conjunto de memórias fundamentais na BAM. Depois verificamos se a associação $(\mathbf{x}^1, \mathbf{y}^1)$ é um ponto estacionário e se todas as memórias fundamentais $(\mathbf{x}^\xi, \mathbf{y}^\xi), \xi = 1, \dots, k$, são pontos estacionários. Repetimos o experimento 1000 vezes para diferentes valores de k (número de memórias fundamentais) e calculamos as probabilidades empíricas de termos pontos estacionários. Na figura 5.5 apresentamos com linha marcada com $\circ$ a probabilidade empírica de uma certa memória fundamental ser um ponto estacionário e com linha marcada com $\square$ a probabilidade empírica de todas as memórias fundamentais serem pontos estacionários. A linha tracejada representa uma estimativa teórica para esta probabilidade, que será chamada capacidade de armazenamento no capítulo 7. No eixo horizontal colocamos k , o número de memórias fundamentais. As linhas pontilhadas verticais indicam os valores $\min(m/\log(n), n/\log(m))/2$ e $\min(m/\log(n), n/\log(m))/4$, respectivamente. Note que a probabilidade empírica de todos os padrões serem pontos fixos deixa de ser 1 quando $k \geq \min(m/\log(n), n/\log(m))/4$. A probabilidade empírica de todas as memórias fundamentais serem pontos fixos para $k = 8$, sendo $\min(m/\log(n), n/\log(m))/2 = 8,7$ foi 0,83, entretanto, a probabilidade empírica de uma certa memória fundamental ser ponto fixo foi 0,97. Estes resultados numéricos conferem com a conjectura 5.2.1. Lembre-se que $n = 100$ e $m = 80$ neste exemplo.

5.3 Memória Associativa de Personnaz

Personnaz *et. al.* apresentaram uma variação da memória associativa de Hopfield que utiliza o armazenamento por projeção para obter a matriz dos pesos sinápticos [67]. Este modelo é referido na literatura simplesmente como “rede de Hopfield com armazenamento por projeção” [3, 31], mas será referido neste trabalho como *Memória Associativa de Personnaz* (Personnaz Associative Memory, Personnaz AM) porque apresenta características diferentes do modelo clássico de Hopfield introduzido na seção anterior e também porque caracterizamos uma rede neural pelo modelo dos neurônios, arquitetura e aprendizado (veja Capítulo 2). Logo, regras de aprendizado distintas produzirão redes neurais (ou memórias associativas) distintas.

5.3.1 Arquitetura da Rede

A memória associativa de Personnaz é descrita por uma rede neural clássica recorrente de camada única totalmente conexa com a função sinal como função de ativação. A arquitetura da memória associativa de Personnaz está apresentada na figura 2.6. Este modelo é descrito pela equação 5.1 (ou equação 5.3). A diferença entre a memória associativa de Personnaz e a memória associativa de Hopfield está na regra de armazenamento.

5.3.2 Aprendizado

A matriz dos pesos sinápticos da memória associativa de Personnaz é

$$W = XX^\dagger, \quad (5.25)$$

onde $X = [\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^k] \in \{-1, 1\}^{n \times k}$ é a matriz com as memórias fundamentais.

Proposição 1. *Todas as memórias fundamentais $\mathbf{x}^1, \dots, \mathbf{x}^k$ serão pontos fixos da memória associativa de Personnaz.*

Demonstração. O problema $\min_W \|X - WX\| = 0$ sempre tem solução (em particular $W = I$) e $\text{sinal}(WX) = \text{sinal}(X) = X$. $\square$

Proposição 2. *Seja $X \in \{-1, 1\}^{n \times k}$. Se $W = XX^\dagger$ então $W = W^T$ e $0 \leq w_{ii} \leq 1$, para todo $i = 1, \dots, n$.*

Demonstração. Considere a decomposição em valores singulares (decomposição SVD)

$$X = U\Sigma V^T = \sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T, \quad (5.26)$$

onde $\mathbf{u}_1, \dots, \mathbf{u}_n$ e $\mathbf{v}_1, \dots, \mathbf{v}_n$ são os vetores singulares da esquerda e da direita, respectivamente, σ_i são os valores singulares e r é o posto da matriz X [101, 25]. Sabemos que

$$X^\dagger = V\Sigma^\dagger U^T = \sum_{j=1}^r \sigma_j^{-1} \mathbf{v}_j \mathbf{u}_j^T, \quad (5.27)$$

e portanto,

$$W = XX^\dagger = \left(\sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T \right) \left(\sum_{j=1}^r \sigma_j^{-1} \mathbf{v}_j \mathbf{u}_j^T \right) = \sum_{i=1}^r \sum_{j=1}^r \sigma_i \sigma_j^{-1} \mathbf{u}_i \mathbf{v}_i^T \mathbf{v}_j \mathbf{u}_j^T. \quad (5.28)$$

Os vetores $\mathbf{v}_1, \dots, \mathbf{v}_n$ são ortogonais, logo $\mathbf{v}_i^T \mathbf{v}_j = \delta_{ij}$ e

$$W = XX^\dagger = \sum_{j=1}^r \sigma_j \sigma_j^{-1} \mathbf{u}_j \mathbf{u}_j^T = \sum_{j=1}^r \mathbf{u}_j \mathbf{u}_j^T. \quad (5.29)$$

Logo, $w_{ii} = \sum_{j=1}^r u_{ij}^2 \geq 0$, onde u_{ij} é a i -ésima componente do vetor $\mathbf{u}_j$. Por outro lado, sabemos que a matriz U é ortogonal, logo $I = UU^T$ e

$$1 = \sum_{j=1}^n u_{ij}^2 = \sum_{j=1}^r u_{ij}^2 + \sum_{j=r+1}^n u_{ij}^2 = w_{ii} + \sum_{j=r+1}^n u_{ij}^2. \quad (5.30)$$

Logo,

$$0 \leq w_{ii} = 1 - \sum_{j=r+1}^n u_{ij}^2 \leq 1. \quad (5.31)$$

Pela Equação (5.29) concluímos também que

$$W^T = \left(\sum_{j=1}^r \mathbf{u}_j \mathbf{u}_j^T \right)^T = \sum_{j=1}^r \mathbf{u}_j \mathbf{u}_j^T = W, \quad (5.32)$$

ou seja, $W = W^T$. □

5.3.3 Convergência

A Proposição 2 e o Teorema 5.1.1 mostram que a memória associativa de Personnaz com atualização assíncrona sempre converge para um ponto fixo.

Exemplo 5.3.1. Considere os padrões bipolares apresentados na figura 1.1. Armazenamos estes padrões na memória associativa de Personnaz e verificamos que todas as memórias fundamentais são pontos fixos conforme a proposição 1.

Considere agora os padrões ruidosos apresentados na figura 4.5. Estes padrões foram gerados a partir dos padrões $\mathbf{x}^1, \dots, \mathbf{x}^5 \in \{-1, 1\}^{4096}$ da figura 1.1 introduzindo ruído uniforme com probabilidade 0,3 de reverter o valor do pixel. A memória associativa de Personnaz encontrou as memórias fundamentais no final da primeira iteração.

5.4 Memória Associativa de Kanter-Sompolinsky

Kanter e Sompolinsky discutiram outra variação da memória associativa de Hopfield que utiliza o armazenamento por projecção, neste caso com diagonal nula, para obter a matriz dos pesos sinápticos [44]. Este modelo, também referido na literatura como “memória associativa de hopfield com armazenamento por projecção” [3, 31], será referido neste trabalho como *Memória Associativa de Kanter-Sompolinsky* (Kanter-Sompolinsky Associative Memory, Kanter-Sompolinsky AM).

5.4.1 Arquitetura da Rede

A memória associativa de Kanter-Sompolinsky é descrita por uma rede neural clássica recorrente de camada única totalmente conexa com a função sinal como função de ativação. A arquitetura da memória associativa de Kanter-Sompolinsky está apresentada na figura 2.6. Este modelo é descrito pela equação 5.1 (ou pela equação 5.3).

5.4.2 Aprendizado

A matriz dos pesos sinápticos da memória associativa de Kanter-Sompolinsky é

$$M = XX^\dagger, \quad \text{com} \quad m_{ii} = 0 \quad \forall i = 1, \dots, n, \quad (5.33)$$

onde $X = [\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^k] \in \{-1, 1\}^{n \times k}$ é a matriz com as memórias fundamentais.

Proposição 3. *Todas as memórias fundamentais $\mathbf{x}^1, \dots, \mathbf{x}^k$ serão pontos fixos da memória associativa de Kanter-Sompolinsky.*

Demonstração. A matriz dos pesos sinápticos da memória associativa de Kanter-Sompolinsky é $M = W - D$, onde D é uma matriz diagonal $n \times n$ com $0 \leq d_{ii} \leq 1$ e W é obtida resolvendo o problema da equação 4.8. Note que $M\mathbf{x}^\xi = W\mathbf{x}^\xi - D\mathbf{x}^\xi = (I - D)\mathbf{x}^\xi$ para toda memória fundamental $\mathbf{x}^\xi$. Logo, $1 - d_{ii} \geq 0$, para todo $i = 1, \dots, n$ e pela definição da função sinal temos que

$$\text{sinal}(\mathbf{m}_i^T \mathbf{x}^\xi) = \text{sinal}\left((1 - d_{ii})x_i^\xi\right) = x_i^\xi, \quad (5.34)$$

para todo $i = 1, \dots, n$ e para todo $\xi = 1, \dots, k$. $\square$

Note que a diferença entre a memória associativa de Kanter-Sompolinsky e a memória associativa de Personnaz está na diagonal nula da matriz dos pesos sinápticos do primeiro modelo. Por causa da diagonal nula, a memória associativa de Kanter-Sompolinsky possui uma tolerância a ruído melhor que a memória associativa de Personnaz [44]. Verificamos este fato através de experimentos computacionais no capítulo 7.

5.4.3 Convergência

A Proposição 2 e o Teorema 5.1.1 mostram que a memória associativa de Kanter-Sompolinsky com atualização assíncrona sempre converge para um ponto fixo.

Exemplo 5.4.1. Repetimos o mesmo experimento realizado no exemplo 5.3.1. Primeiro, verificamos que todas as memórias fundamentais são pontos fixos da memória associativa de Kanter-Sompolinsky. Depois apresentamos os padrões da figura 4.5 como entrada e verificamos que as recordações fundamentais $\mathbf{x}^1, \dots, \mathbf{x}^k$ foram encontradas no final da primeira iteração.

5.5 Memória Associativa Bidirecional Assimétrica

A *Memória Associativa Bidirecional Assimétrica* (Asymmetric Bidirectional Associative Memories, ABAM), também conhecida como *Projection HDAM* é uma extensão da memória associativa de Personnaz para o caso heteroassociativo [109, 32].

5.5.1 Arquitetura

A memória associativa bidirecional assimétrica é uma rede neural clássica recorrente totalmente conexa descrita pelas equações

$$\mathbf{y}(t) = \text{sinal}(W_1 \mathbf{x}(t)), \quad (5.35)$$

$$\mathbf{x}(t+1) = \text{sinal}(W_2 \mathbf{y}(t)), \quad \text{para } t = 0, 1, 2, \dots \quad (5.36)$$

com atualização assíncrona. A arquitetura desta rede está apresentada na figura 5.3.

5.5.2 Aprendizado

Considere os pares de padrões $(\mathbf{x}^\xi, \mathbf{y}^\xi)$, onde $\mathbf{x}^\xi \in \{-1, 1\}^n$ e $\mathbf{y}^\xi \in \{-1, 1\}^m$, para todo $\xi = 1, \dots, k$. Tome $X = [\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^k]$ e $Y = [\mathbf{y}^1, \mathbf{y}^2, \dots, \mathbf{y}^k]$. As matrizes dos pesos sinápticos, W_1 e W_2 , são encontradas usando o armazenamento por projeção, isto é,

$$W_1 = YX^\dagger \quad \text{e} \quad W_2 = XY^\dagger. \quad (5.37)$$

Teorema 5.5.1. *Sejam $X = [\mathbf{x}^1, \dots, \mathbf{x}^k] \in \{-1, 1\}^{n \times k}$ e $Y = [\mathbf{y}^1, \dots, \mathbf{y}^k] \in \{-1, 1\}^{m \times k}$ as matrizes das memórias fundamentais $\{(\mathbf{x}^\xi, \mathbf{y}^\xi), \xi = 1, \dots, k\}$. Se X e Y possuem ambas posto completo, então todas as memórias fundamentais serão pontos estacionários.*

Demonstração. Se X tem posto completo, então $X^\dagger = (X^T X)^{-1} X^T$ e

$$\text{sinal}(W_1 X) = \text{sinal}(Y(X^T X)^{-1} X^T X) = \text{sinal}(Y) = Y. \quad (5.38)$$

Analogamente, substituindo X por Y e vice-versa, encontramos $\text{sinal}(W_2 Y) = X$ se Y tem posto completo. $\square$

Não garantimos a convergência da ABAM porque

$$W' = \begin{bmatrix} 0 & W_2 \\ W_1 & 0 \end{bmatrix}, \quad (5.39)$$

não é uma matriz simétrica ($W_2 \neq W_1^T$). Resultados empíricos mostraram que este modelo apresenta um comportamento oscilatório principalmente quando o número de padrões armazenados está próximo ou é maior que $\min(m, n)$ [30].

Exemplo 5.5.1. Considere os padrões $\mathbf{x}^1, \dots, \mathbf{x}^5 \in \{-1, 1\}^{4096}$ e $\mathbf{y}^1, \dots, \mathbf{y}^5 \in \{-1, 1\}^{4096}$ apresentados nas figuras 1.1 e 1.2. Neste caso, as matrizes X e Y possuem ambas posto completo. Armazenamos o conjunto das memórias fundamentais $\{(\mathbf{x}^\xi, \mathbf{y}^\xi), \xi = 1, \dots, 5\}$ na ABAM e verificamos que todas as memórias fundamentais são pontos estacionários segundo o teorema 5.5.1. Depois apresentamos como entrada os padrões corrompidos $\tilde{\mathbf{x}}^1, \dots, \tilde{\mathbf{x}}^5$ apresentados na figura 4.5 e verificamos que os padrões recordados pela ABAM no final da primeira iteração são exatamente as recordações fundamentais.

5.6 Memória Associativa com Capacidade Exponencial

A *Memória Associativa com Capacidade Exponencial* (Exponential Correlation Associative Memory, ECAM) pode ser vista como um caso particular das *Memórias Associativas Recorrentes por Correlação* (Recurrent Correlation Associative Memories, RCAM) que serão discutidas nesta seção [15, 16, 32]. As memórias associativas recorrentes por correlação podem ser vistas como generalizações da memória associativa de Hopfield discreta com auto-alimentação.

5.6.1 Arquitetura

Na memória associativa de Hopfield discreta com auto-alimentação, a i -ésima componente do vetor de estados, $\mathbf{x}(t+1)$, é calculado através da equação

$$\mathbf{x}_i(t+1) = \text{sinal}(\mathbf{w}_i^T \mathbf{x}(t)), \quad (5.40)$$

onde $\mathbf{w}_i^T$ é a i -ésima linha da matriz W . No armazenamento por correlação temos

$$\mathbf{w}_i^T = \sum_{\xi=1}^k x_i^\xi (\mathbf{x}^\xi)^T. \quad (5.41)$$

Portanto, a i -ésima componente de $\mathbf{x}(t+1)$ é

$$\mathbf{x}_i(t+1) = \text{sinal} \left(\sum_{\xi=1}^k x_i^\xi (\mathbf{x}^\xi)^T \mathbf{x}(t) \right) = \text{sinal} \left(\sum_{\xi=1}^k \langle \mathbf{x}^\xi, \mathbf{x}(t) \rangle x_i^\xi \right), \quad (5.42)$$

onde $\langle \cdot, \cdot \rangle$ representa o produto interno Euclidiano. O termo $\langle \mathbf{x}^\xi, \mathbf{x}(t) \rangle$ representa a correlação entre o estado atual da memória $\mathbf{x}(t)$ e o ξ -ésimo padrão armazenado $\mathbf{x}^\xi$. Podemos dizer, de um modo intuitivo, que quanto mais “parecido” o vetor $\mathbf{x}(t)$ for de $\mathbf{x}^\xi$, maior será o valor de $\langle \mathbf{x}^\xi, \mathbf{x}(t) \rangle$. Para enfatizar esta correlação, podemos aplicar uma função f no resultado de $\langle \mathbf{x}^\xi, \mathbf{x}(t) \rangle$. A função f deve ser monótona não-decrescente, pois desejamos que vetores pouco correlacionados apresentem um valor menor do que aqueles mais correlacionados.

Ao introduzir a função f , encontramos a seguinte expressão para a i -ésima componente do vetor $\mathbf{x}(t+1)$:

$$\mathbf{x}_i(t+1) = \text{sinal} \left(\sum_{\xi=1}^k f(\langle \mathbf{x}^\xi, \mathbf{x}(t) \rangle) x_i^\xi \right). \quad (5.43)$$

Numa notação matricial temos

$$\mathbf{x}(t+1) = \text{sinal} \left(X F \left(X^T \mathbf{x}(t) \right) \right), \quad (5.44)$$

onde $F(\cdot) = [f(\cdot), \dots, f(\cdot)]^T$. Observando a equação (5.44), podemos dizer que as RCAM's são redes neurais recorrentes de duas camadas totalmente conexa com neurônios clássicos com função de ativação f na primeira camada e função sinal na camada de saída. Na figura 5.6 apresentamos a arquitetura das RCAM's.

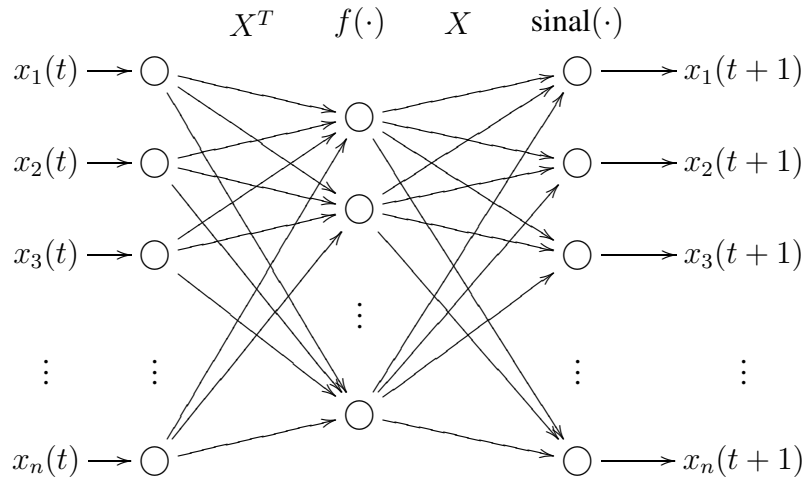


Fig. 5.6: Arquitetura das Memórias Associativas Recorrentes por Correlação.

A função f introduzida representa um papel importante na capacidade e dinâmica da RCAM. Existem infinitas funções que podem ser usadas, mas precisamos que esta função seja fácil de ser implementada e forneça uma RCAM assintoticamente estável com uma capacidade de armazenamento grande. Apresentamos a seguir algumas funções usadas na RCAM. Depois apresentaremos condições suficientes para obtermos uma memória assintoticamente estável.

1. Para $f(v) = v$, encontramos a *memória associativa de Hopfield discreta com armazenamento por correlação*.
2. Para $f(v) = (n+v)^q$, onde q é um inteiro maior que 1 e n é a dimensão dos vetores $\mathbf{x}^\xi$, encontramos a *memória associativa de ordem alta por correlação* (High-Order Correlation Associative Memory). A capacidade de armazenamento desta memória é assintoticamente proporcional a n^q [16].
3. Para $f(v) = (n - v)^{-p}$, onde $p \geq 1$, encontramos a *memória associativa com função potencial por correlação* (Potential-Function Correlation Associative Memory). A capacidade de armazenamento deste modelo aumenta exponencialmente com a dimensão n dos padrões armazenados. Este modelo foi proposto independentemente por Dembo e Zeitouni, e por Sayeh e Han [16]. Apresentamos aqui uma versão com tempo discreto e estados bipolares, embora, originalmente tenha sido proposto com tempo contínuo e padrões com valores reais.
4. Quando $f(v) = a^v$, com $a > 1$, encontramos a *memória associativa com capacidade exponencial*. Esta é a única RCAM discutida neste trabalho.

5.6.2 Aprendizado

O aprendizado da RCAM é direto pois a matriz dos pesos sinápticos da camada de entrada é X^T e a matriz dos pesos sinápticos da camada de saída é X . O armazenamento de um novo padrão $\mathbf{x}^h$ é feito

concatenando as matrizes X^T e X com $(\mathbf{x}^h)^T$ e $\mathbf{x}^h$, respectivamente.

5.6.3 Convergência

As quatro RCAMs apresentadas anteriormente (itens (1) – (4)) com atualização assíncrona possuem regiões de indecisão vazias. O teorema sobre a convergência das RCAM's foi enunciado e demonstrado por Chiueh e Goodman em [15].

Teorema 5.6.1 (Teorema da Convergência de Chiueh e Goodman). *Seja $f(v)$ uma função contínua monótona não-decrescente definida sobre um intervalo fechado. A RCAM descrita pela equação (5.43) com atualização assíncrona sempre converge para um ponto fixo para qualquer padrão-chave apresentado a RCAM.*

O seguinte teorema, introduzido por Chiueh e Goodman em [15], nos diz que o número de memórias fundamentais que podem ser armazenadas como pontos fixos da ECAM aumenta exponencialmente com n , o número de componentes dos padrões.

Teorema 5.6.2. *Sejam $\mathbf{x}^1, \dots, \mathbf{x}^k \in \{-1, 1\}^n$, com n suficientemente grande, padrões não correlacionados gerados aleatoriamente com distribuição uniforme. Todos as memórias fundamentais serão pontos fixos com grande probabilidade se $k \leq \alpha c^n$, onde $c \in [1, 2]$ é uma constante fixa que depende de a .*

É importante observar que o valor da constante c do teorema acima decresce quando o valor de a diminui [16].

Exemplo 5.6.1. Repetimos o experimento realizado no exemplo 5.3.1. Verificamos que todas as memórias fundamentais são pontos fixos da ECAM. Depois apresentamos os padrões ruidosos da figura 4.5 como entrada e verificamos que os padrões recordados são as respectivas recordações fundamentais. Neste exemplo utilizamos $a = 1,007$ (ou $f(x) = 2^{(x/100)}$) para evitar overflow. A ECAM encontrou as recordações fundamentais com uma única iteração.

5.7 Memória Associativa Bidirecional com Capacidade Exponencial

A *Memória Associativa Bidirecional com Capacidade Exponencial* (Bidirectional Exponential Capacity Associative Memory, BECAM) é uma generalização da memória ECAM para o caso heteroassociativo. Este modelo de memória associativa foi introduzido por nós nesta dissertação de mestrado.

5.7.1 Arquitetura

A BECAM é descrita por uma rede neural clássica recorrente com quatro camadas totalmente conexa com função sinal e função exponencial como funções de ativação.

Seja $\{(\mathbf{x}^\xi, \mathbf{y}^\xi), \xi = 1, 2, \dots, k\}$ com $\mathbf{x}^\xi \in \{-1, 1\}^n$ e $\mathbf{y}^\xi \in \{-1, 1\}^m$, para todo $\xi = 1, \dots, k$. Na BAM, a i -ésima componente do vetor de estados $\mathbf{y}(t)$ é dada por

$$\mathbf{y}(t)_i = \text{sinal}(\mathbf{w}_i^T \mathbf{x}(t)), \quad (5.45)$$

onde $\mathbf{w}_i^T$ é a i -ésima linha da matriz W . Quando usamos o armazenamento por correlação, temos $\mathbf{w}_i^T = \sum_{\xi=1}^k y_i^\xi (\mathbf{x}^\xi)^T$. Assim, a i -ésima componente de $\mathbf{y}(t)$ é

$$y(t)_i = \text{senal} \left(\sum_{\xi=1}^k y_i^\xi (\mathbf{x}^\xi)^T \mathbf{x}(t) \right) = \text{senal} \left(\sum_{\xi=1}^k \langle \mathbf{x}^\xi, \mathbf{x}(t) \rangle y_i^\xi \right), \quad (5.46)$$

onde $\langle \cdot, \cdot \rangle$ representa o produto interno Euclidiano. Podemos aplicar uma função f no valor $\langle \mathbf{x}^\xi, \mathbf{x}(t) \rangle$ a fim de enfatizar esta correlação. As quatro funções apresentadas na seção 5.6 podem ser usadas neste modelo e cada uma produz uma memória heteroassociativa diferente. Voltamos nossa atenção para a função $f(v) = a^v$, para $a > 1$, que chamaremos de BECAM (Bidirectional Exponential Capacity Associative Memory).

Encontramos a seguinte expressão para a i -ésima componente do vetor $\mathbf{y}(t)$ quando introduzimos a função f :

$$y_i(t) = \text{senal} \left(\sum_{\xi=1}^k f(\langle \mathbf{x}^\xi, \mathbf{x}(t) \rangle) y_i^\xi \right). \quad (5.47)$$

Numa notação matricial temos

$$\mathbf{y}(t) = \text{senal} \left(Y F \left(X^T \mathbf{x}(t) \right) \right), \quad (5.48)$$

onde $F(\cdot) = [f(\cdot), \dots, f(\cdot)]^T$. Analogamente, podemos encontrar a seguinte expressão para $\mathbf{x}(t+1)$:

$$\mathbf{x}(t+1) = \text{senal} \left(X F \left(Y^T \mathbf{y}(t) \right) \right). \quad (5.49)$$

Observando as equações (5.48) e (5.49) percebemos que a BECAM é uma rede recorrente com quatro camadas. Na figura 5.7.1 apresentamos a arquitetura da BECAM.

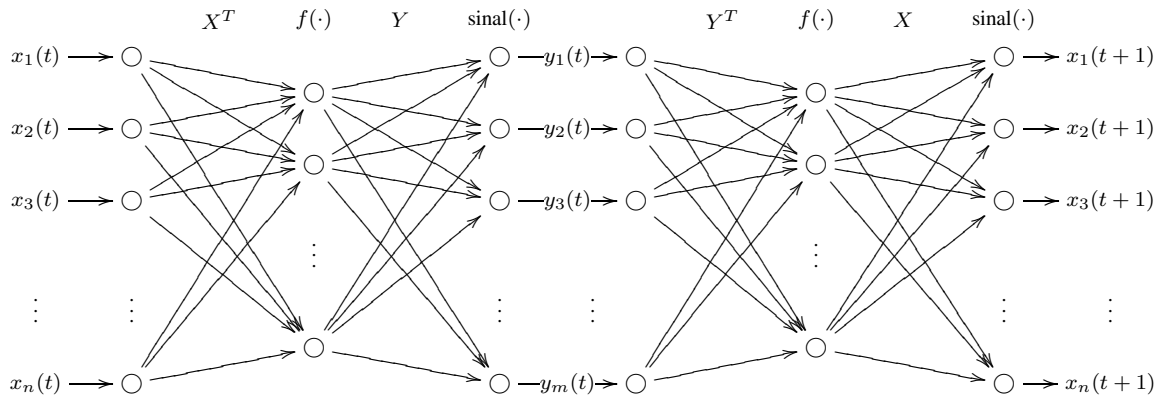


Fig. 5.7: Arquitetura da Memória Associativa Bidirecional com Capacidade Exponencial.

5.7.2 Aprendizado

As matrizes dos pesos sinápticos da BECAM são: X^T , Y , Y^T e X . Note que o armazenamento de um novo par $(\mathbf{x}^h, \mathbf{y}^h)$ é feito concatenando as matrizes X^T , X , Y^T e Y com $(\mathbf{x}^h)^T$, $\mathbf{x}^h$, $(\mathbf{y}^h)^T$ e $\mathbf{y}^h$, respectivamente.

5.7.3 Convergência

A convergência da BECAM pode ser verificada convertendo ela numa ECAM. Tome

$$\mathbf{x}' = \begin{bmatrix} \mathbf{x} \\ \mathbf{y} \end{bmatrix} \quad \text{e} \quad X' = \begin{bmatrix} 0 & X \\ Y & 0 \end{bmatrix}, \quad (5.50)$$

como sendo os padrões e a matriz dos pesos sinápticos, respectivamente. Pela equação (5.44), teremos

$$\begin{aligned} \begin{bmatrix} \mathbf{x}(t+1) \\ \mathbf{y}(t+1) \end{bmatrix} &= \text{senal} \left(\begin{bmatrix} 0 & X \\ Y & 0 \end{bmatrix} f \left(\begin{bmatrix} 0 & Y^T \\ X^T & 0 \end{bmatrix} \begin{bmatrix} \mathbf{x}(t) \\ \mathbf{y}(t) \end{bmatrix} \right) \right) \\ &= \text{senal} \left(\begin{bmatrix} 0 & X \\ Y & 0 \end{bmatrix} f \left(\begin{bmatrix} Y^T \mathbf{y}(t) \\ X^T \mathbf{x}(t) \end{bmatrix} \right) \right) = \text{senal} \left(\begin{bmatrix} X f(Y^T \mathbf{y}(t)) \\ Y f(X^T \mathbf{x}(t)) \end{bmatrix} \right). \end{aligned}$$

Exemplo 5.7.1. Repetimos o experimento realizado no exemplo 5.5.1. Verificamos que todas as memórias fundamentais são pontos fixos da BECAM. Depois apresentamos os padrões ruidosos $\mathbf{x}^1, \dots, \mathbf{x}^5$ da figura 4.5 como entrada e verificamos que os padrões $\mathbf{y}^1(1), \dots, \mathbf{y}^5(1)$ recordados no final da primeira iteração são as respectivas recordações fundamentais. Neste exemplo utilizamos $a = 1,007$ (ou $f(x) = 2^{(x/100)}$) para evitar overflow. A BECAM encontrou as recordações fundamentais com uma única iteração.

5.8 Modelo do Estado Cerebral numa Caixa (BSB)

O *Modelo do Estado Cerebral numa Caixa* (Brain-State-in-a-Box, BSB) foi proposto por Anderson et al. em 1977 [6] como uma rede recorrente de tempo discreto, não-linear e totalmente conexa. Como nos modelos anteriores, o modelo do estado cerebral numa caixa pode ser visto como uma memória auto-associativa que minimiza a energia de um sistema dinâmico não-linear.

5.8.1 Arquitetura

O modelo BSB é uma rede neural clássica recorrente com tempo discreto de camada única. Seja $\mathbf{x}(t+1)$ o vetor de estado no tempo $t+1$. A recorrência é dada por $\alpha W \mathbf{x}(t)$, onde W é a matriz dos pesos sinápticos e α é uma constante positiva que controla o peso deste termo. Na BSB, adicionamos um termo de decaimento ao termo de recorrência dado pelo vetor $\mathbf{x}(t)$, o estado anterior. Assim, o próximo estado, $\mathbf{x}(t+1)$, será dado por

$$\mathbf{x}(t+1) = L(\mathbf{x}(t) + \alpha(W\mathbf{x}(t) + \boldsymbol{\theta})). \quad (5.51)$$

Neste trabalho consideraremos $\theta = 0$. A função de ativação usada no modelo BSB é a função linear por partes definida como

$$L(\mathbf{x}) = 1 \wedge [(-1) \vee \mathbf{x}]. \quad (5.52)$$

Note que esta função limita os vetores de estados numa região restrita do espaço, especificamente, no hipercubo $H_n = [-1, 1]^n$. Daí o nome: Modelo do Estado Cerebral numa Caixa [46].

5.8.2 Aprendizado

O modelo original da BSB proposto por Anderson et al. utiliza o aprendizado de Hebb (armazenamento por correlação) onde tomamos $W = XX^T$ [6]. Neste caso, a matriz dos pesos sinápticos é simétrica, semi-definida positiva e com diagonal positiva.

Outras regras de aprendizado podem ser utilizadas na BSB, por exemplo, o armazenamento por projeção ou o armazenamento iterativo proposto em [66]. Neste trabalho discutiremos somente o armazenamento por correlação.

5.8.3 Convergência

A BSB sempre converge para um ponto fixo independente da memória-chave apresentada e minimiza a função energia

$$\text{Energia}(\mathbf{x}) = -\frac{1}{2}\mathbf{x}^T W \mathbf{x}. \quad (5.53)$$

se impormos pequenas condições ao modelo [41, 23, 24, 66].

Teorema 5.8.1 (Teorema de Golden da Minimização da Função Energia do Modelo BSB). *Considere o modelo neural descrito pela equação (5.51). Se $W = W^T$ é semi-definida positiva ou $\alpha < (2/|\lambda_{\min}|)$, onde $\lambda_{\min}$ é o menor autovalor da matriz simétrica W , então:*

1. $\text{Energia}(\mathbf{x}(t+1)) < \text{Energia}(\mathbf{x}(t))$ se $\mathbf{x}(t+1) \neq \mathbf{x}(t)$,
2. $\text{Energia}(\mathbf{x}(t+1)) = \text{Energia}(\mathbf{x}(t))$ se e somente se $\mathbf{x}(t+1) = \mathbf{x}(t)$,

onde $\text{Energia}(\mathbf{x})$ é a função definida na equação (5.53).

A demonstração do teorema 5.8.1 pode ser encontrada em [24].

Teorema 5.8.2. *Seja $W \in \mathbb{R}^{n \times n}$ com $w_{ii} \geq 0$ para $i = 1, \dots, n$. Se $\mathbf{x} \in [-1, 1]^n$ é um ponto fixo da BSB então $\mathbf{x} \in \{-1, 1\}^n$, isto é, $\mathbf{x}$ é um vértice do hipercubo $[-1, 1]^n$.*

A demonstração do teorema 5.8.2 encontra-se em [66]. Note que o teorema 5.8.2 não impõe nenhuma hipótese sobre a simetria da matriz dos pesos sinápticos. O seguinte teorema foi introduzido por nós e relaciona o modelo BSB com a memória associativa de Hopfield.

Teorema 5.8.3. *Um padrão $\mathbf{x} \in \{-1, 1\}^n$ será um ponto fixo da BSB (com $\theta = 0$) se e somente se $\mathbf{x}$ for um ponto fixo da memória associativa de Hopfield com a mesma matriz de pesos sinápticos.*

Demonstração. Seja $\mathbf{x} \in \{-1, 1\}^n$. Pela definição da função linear por partes, da função sinal e lembrando que $\alpha > 0$, temos:

$$\mathbf{x} = L(\alpha W\mathbf{x} + \mathbf{x}) \Leftrightarrow (\alpha(W\mathbf{x})_i + x_i) x_i \geq 1 \quad \forall i = 1, \dots, n \quad (5.54)$$

$$\Leftrightarrow \alpha x_i (W\mathbf{x})_i + 1 \geq 1 \quad \forall i = 1, \dots, n \quad (5.55)$$

$$\Leftrightarrow x_i (W\mathbf{x})_i \geq 0 \quad \forall i = 1, \dots, n \quad (5.56)$$

$$\Leftrightarrow \mathbf{x} = \text{sinal}(W\mathbf{x}). \quad (5.57)$$

□

Note que o teorema 5.8.3 não impõe nenhuma hipótese sobre a matriz dos pesos sinápticos W . Portanto, este resultado permanece válido para outras regras de aprendizado, por exemplo, para o armazenamento por projeção. Note também que não impomos nenhuma condição sobre a constante α , exceto $\alpha > 0$. Finalmente, observe que este resultado só é válido para padrões nos vértices de $[-1, 1]^n$. Poderemos ter pontos fixos na BSB que estão no interior de $[-1, 1]^n$ se as hipóteses do teorema 5.8.2 não forem satisfeitas. Estes pontos fixos da BSB obviamente não serão pontos fixos da memória associativa de Hopfield e suas variações.

Capítulo 6

Memórias Associativas Morfológicas

Neste capítulo discutimos as *memórias associativas morfológicas* (Morphological Associative Memories, MAM). Discutimos primeiro o caso heteroassociativo em tons de cinza apresentando exemplos e resultados sobre a capacidade de armazenamento e tolerância a ruído. Depois voltamos nossa atenção para o caso auto-associativo. Terminamos o capítulo discutindo o caso auto-associativo binário e as memórias associativas de duas camadas.

6.1 Memórias Associativas Morfológicas Heteroassociativas

As memórias associativas morfológicas são redes neurais descritas pelo modelo neural morfológico e foram introduzidas em [73, 74, 76]. O caso mais simples ocorre quando temos uma rede alimentada adiante, totalmente conexa e de camada única, cuja função de ativação é a função identidade. Neste caso, o mapeamento associativo da memória $G : \mathbb{R}^n \rightarrow \mathbb{R}^m$ é descrito por uma matriz $W \in \mathbb{R}^{m \times n}$ ou $M \in \mathbb{R}^{m \times n}$ e o produto-mínimo ou o produto-máximo. Dado um padrão-chave $\mathbf{x} \in \mathbb{R}^n$, encontramos o padrão recordado $\mathbf{y} \in \mathbb{R}^m$ através da equação

$$\mathbf{y} = W \boxtimes \mathbf{x}, \quad (6.1)$$

ou

$$\mathbf{y} = M \boxtimes \mathbf{x}. \quad (6.2)$$

Repare na semelhança entre o mapeamento associativo das memórias associativas morfológicas heteroassociativas e o mapeamento associativo das memórias associativas lineares. A diferença está no produto-máximo ou produto-mínimo usado nos modelos morfológicos. Lembre-se que os modelos morfológicos apresentam um comportamento não-linear devido a estas duas operações.

Podemos fazer uma interpretação das equações (6.1) e (6.2) usando a morfologia matemática [85, 86, 87]. Podemos verificar que os operadores $\delta(\mathbf{x}) = W \boxtimes \mathbf{x}$ e $\varepsilon(\mathbf{x}) = M \boxtimes \mathbf{x}$ são operações de dilatação e erosão na morfologia matemática, respectivamente. Desta forma, a tarefa na fase de armazenamento das memórias associativas morfológicas seria encontrar uma dilatação (erosão) tal que $\delta(\mathbf{x}^\xi) = \mathbf{y}^\xi$ ($\varepsilon(\mathbf{x}^\xi) = \mathbf{y}^\xi$) para todo $\xi = 1, \dots, k$. Não usaremos esta interpretação na fase de armazenamento, mas podemos extrair resultados interessantes (e intuitivos) para a fase de recordação. Por exemplo, a dilatação (erosão) é um operador extensivo (anti-extensivo), isto é, a dilatação (erosão) expande (reduz) um objeto preto de uma imagem com fundo branco. A memória

associativa morfológica descrita pela equação (6.1) (Equação (6.2)) também possui esta característica. Além disso, uma dilatação (erosão) é usada para remover ruído *erosivo* (*dilatativo*) de uma imagem. Dizemos que uma imagem $\tilde{\mathbf{x}}$ é uma versão de $\mathbf{x}$ corrompida com ruído *erosivo* (*dilatativo*) se $\tilde{\mathbf{x}} \leq \mathbf{x}$ ($\tilde{\mathbf{x}} \geq \mathbf{x}$).

6.1.1 Aprendizado

Suponha que armazenamos um único par $(\mathbf{x}, \mathbf{y})$ na memória. Pelas equações (6.1) e (6.2), encontramos, respectivamente,

$$y_i = (W \boxtimes \mathbf{x})_i = \bigvee_{j=1}^n (w_{ij} + x_j), \quad (6.3)$$

e

$$y_i = (M \boxtimes \mathbf{x})_i = \bigwedge_{j=1}^n (w_{ij} + x_j). \quad (6.4)$$

Pela equação (6.3) temos que $y_i \leq w_{ij} + x_j$, ou seja, $y_i - x_j \leq w_{ij}$. Podemos definir W através da equação $w_{ij} = y_i - x_j$. Usando a notação matricial da álgebra de imagens, temos

$$W_{XY} = \begin{bmatrix} y_1 - x_1 & \dots & y_1 - x_n \\ \vdots & \ddots & \vdots \\ y_m - x_1 & \dots & y_m - x_n \end{bmatrix} = \mathbf{y} \boxtimes \mathbf{x}^*, \quad (6.5)$$

onde $\mathbf{x}^* = -\mathbf{x}^T$ é o conjugado aditivo do vetor $\mathbf{x}$.

Podemos verificar que W_{XY} recupera o padrão $\mathbf{y}$ quando apresentamos $\mathbf{x}$ como entrada. De fato,

$$W_{XY} \boxtimes \mathbf{x} = \begin{bmatrix} \bigvee_{i=1}^n (y_1 - x_i + x_i) \\ \vdots \\ \bigvee_{i=1}^n (y_m - x_i + x_i) \end{bmatrix} = \mathbf{y}. \quad (6.6)$$

Repare na semelhança entre as equações (6.5) e (4.3) do armazenamento por correlação quando armazenamos um único padrão. Por analogia a (4.3), quando armazenamos o conjunto de memórias fundamentais $\{(\mathbf{x}^\xi, \mathbf{y}^\xi), \xi = 1, \dots, k\}$, temos

$$W_{XY} = Y \boxtimes X^*, \quad (6.7)$$

onde $X = [\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^k]$ e $Y = [\mathbf{y}^1, \mathbf{y}^2, \dots, \mathbf{y}^k]$. Note que usamos $\boxtimes$ para gerar a matriz W_{XY} (equação (6.7)) e usamos $\boxtimes$ na fase de recordação (equação (6.6)).

Partindo de (6.4) e seguindo um raciocínio análogo, encontramos

$$M_{XY} = Y \boxtimes X^*, \quad (6.8)$$

como sendo a matriz dos pesos sinápticos da memória associativa morfológica descrita pela equação (6.2).

Exemplo 6.1.1. Seja

$$X = \begin{bmatrix} 0 & 0 & 0 \\ 0 & -2 & -3 \\ 0 & -4 & 0 \end{bmatrix} \quad \text{e} \quad Y = \begin{bmatrix} 0 & -1 & 0 \\ 1 & -1 & -2 \\ 0 & 0 & 0 \end{bmatrix}. \quad (6.9)$$

De acordo com a equação (6.7) temos

$$W_{XY} = Y \boxtimes X^* = \begin{bmatrix} 0 & 0 & 0 \\ 1 & 1 & 1 \\ 0 & 0 & 0 \end{bmatrix} \wedge \begin{bmatrix} -1 & 1 & 3 \\ -1 & 1 & 3 \\ 0 & 2 & 4 \end{bmatrix} \wedge \begin{bmatrix} 0 & 3 & 0 \\ -2 & 1 & -2 \\ 0 & 3 & 0 \end{bmatrix} = \begin{bmatrix} -1 & 0 & 0 \\ -2 & 1 & -2 \\ 0 & 0 & 0 \end{bmatrix}, \quad (6.10)$$

e pela equação (6.8) encontramos

$$M_{XY} = Y \boxtimes X^* = \begin{bmatrix} 0 & 0 & 0 \\ 1 & 1 & 1 \\ 0 & 0 & 0 \end{bmatrix} \vee \begin{bmatrix} -1 & 1 & 3 \\ -1 & 1 & 3 \\ 0 & 2 & 4 \end{bmatrix} \vee \begin{bmatrix} 0 & 3 & 0 \\ -2 & 1 & -2 \\ 0 & 3 & 0 \end{bmatrix} = \begin{bmatrix} 0 & 3 & 3 \\ 1 & 1 & 3 \\ 0 & 3 & 4 \end{bmatrix}. \quad (6.11)$$

Podemos verificar que $W_{XY} \boxtimes \mathbf{x}^\xi = \mathbf{y}^\xi = M_{XY} \boxtimes \mathbf{x}^\xi$, para todo $\xi = 1, 2, 3$. Por exemplo,

$$W_{XY} \boxtimes \mathbf{x}^1 = \begin{bmatrix} -1 & 0 & 0 \\ -2 & 1 & -2 \\ 0 & 0 & 0 \end{bmatrix} \boxtimes \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} = \mathbf{y}^1, \quad (6.12)$$

e

$$M_{XY} \boxtimes \mathbf{x}^2 = \begin{bmatrix} 0 & 3 & 3 \\ 1 & 1 & 3 \\ 0 & 3 & 4 \end{bmatrix} \boxtimes \begin{bmatrix} 0 \\ -2 \\ -4 \end{bmatrix} = \begin{bmatrix} -1 \\ -1 \\ 0 \end{bmatrix} = \mathbf{y}^2. \quad (6.13)$$

Chamamos a atenção do leitor para o fato de que $\mathbf{y}^\xi \boxtimes (\mathbf{x}^\xi)^* = \mathbf{y}^\xi \boxtimes (\mathbf{x}^\xi)^*$. Podemos reduzir a notação denotando estes produtos externos morfológicos por $\mathbf{y}^\xi + (\mathbf{x}^\xi)^*$. Desta forma, as equações (6.7) e (6.8) possuem um análogo à equação (4.3), isto é

$$W_{XY} = \bigwedge_{\xi=1}^k \left(\mathbf{y}^\xi + (\mathbf{x}^\xi)^* \right), \quad \text{e} \quad M_{XY} = \bigvee_{\xi=1}^k \left(\mathbf{y}^\xi + (\mathbf{x}^\xi)^* \right). \quad (6.14)$$

Podemos adicionar facilmente um novo par de padrões nas memórias W_{XY} e M_{XY} usando as equações de (6.14). Se armazenamos os pares $(\mathbf{x}^\xi, \mathbf{y}^\xi)$, para $\xi = 1, \dots, k$, e desejamos adicionar um novo par $(\mathbf{x}^{k+1}, \mathbf{y}^{k+1})$, então definimos

$$W_{XY}^{nova} = W_{XY}^{atual} \wedge \left[\mathbf{y}^{k+1} + (-\mathbf{x}^{k+1})^T \right], \quad (6.15)$$

ou

$$M_{XY}^{nova} = M_{XY}^{atual} \vee \left[\mathbf{y}^{k+1} + (-\mathbf{x}^{k+1})^T \right]. \quad (6.16)$$

Entretanto, não podemos excluir um padrão armazenado.

A seguinte definição, introduzida em [76], diz respeito ao armazenamento das memórias fundamentais em uma memória associativa morfológica.

Definição 6.1.1. Uma matriz A é uma memória $\boxminus$ -perfeita para $(\mathbf{x}^\xi, \mathbf{y}^\xi)$, com $\xi = 1, \dots, k$, se e somente se, $A \boxminus \mathbf{x}^\xi = \mathbf{y}^\xi$, para todo $\xi = 1, \dots, k$. Analogamente, uma matriz B é uma memória $\boxplus$ -perfeita para $(\mathbf{x}^\xi, \mathbf{y}^\xi)$, com $\xi = 1, \dots, k$, se e somente se, $B \boxplus \mathbf{x}^\xi = \mathbf{y}^\xi$, para todo $\xi = 1, \dots, k$.

Pela definição, se $X = [\mathbf{x}^1, \dots, \mathbf{x}^k]$, $Y = [\mathbf{y}^1, \dots, \mathbf{y}^k]$ e A é uma memória $\boxminus$ -perfeita, então $A \boxminus X = Y$. Se B é $\boxplus$ -perfeita, então $B \boxplus X = Y$. Temos também que, se A é $\boxminus$ -perfeita, então $(A \boxminus \mathbf{x}^\xi)_i = y_i^\xi$ para todo $\xi = 1, \dots, k$ e todo $i = 1, \dots, m$. Equivalentemente,

$$\bigvee_{j=1}^n (a_{ij} + x_j^\xi) = y_i^\xi, \quad \forall \xi = 1, \dots, k \text{ e } i = 1, \dots, m. \quad (6.17)$$

Da equação (6.17) obtemos a seguinte desigualdade para um índice $j \in \{1, \dots, n\}$ arbitrário:

$$a_{ij} + x_j^\xi \leq y_i^\xi, \quad \forall \xi = 1, \dots, k \quad (6.18)$$

$$\Leftrightarrow a_{ij} \leq y_i^\xi - x_j^\xi, \quad \forall \xi = 1, \dots, k \quad (6.19)$$

$$\Leftrightarrow a_{ij} \leq \bigwedge_{\xi=1}^k (y_i^\xi - x_j^\xi) = w_{ij}. \quad (6.20)$$

Temos com isso que $A \leq W_{XY}$ e consequentemente

$$Y = A \boxminus X \leq W_{XY} \boxminus X. \quad (6.21)$$

Pela equação (6.14), temos $W_{XY} \leq \mathbf{y}^\xi \times (\mathbf{x}^\xi)^* \leq M_{XY}$ para todo $\xi = 1, \dots, k$ e usando a equação (6.6) concluímos que $W_{XY} \boxminus \mathbf{x}^\xi \leq (\mathbf{y}^\xi \times (\mathbf{x}^\xi)^*) \boxminus \mathbf{x}^\xi = \mathbf{y}^\xi = (\mathbf{y}^\xi \times (\mathbf{x}^\xi)^*) \boxminus \mathbf{x}^\xi \leq M_{XY} \boxminus \mathbf{x}^\xi$, para $\xi = 1, \dots, k$, ou equivalentemente,

$$W_{XY} \boxminus X \leq Y \leq M_{XY} \boxminus X. \quad (6.22)$$

Finalmente, pelas equações (6.21) e (6.22) concluímos que

$$Y = A \boxminus X \leq W_{XY} \boxminus X \leq Y, \quad (6.23)$$

e portanto $W_{XY} \boxminus X = Y$. Um argumento similar mostra que se B é $\boxplus$ -perfeita para $(\mathbf{x}^\xi, \mathbf{y}^\xi)$, para todo $\xi = 1, \dots, k$, então $M_{XY} \leq B$ e $M_{XY} \boxplus X = Y$.

Teorema 6.1.1. Se A é $\boxminus$ -perfeita e B é $\boxplus$ -perfeita para $(\mathbf{x}^\xi, \mathbf{y}^\xi)$, com $\xi = 1, \dots, k$, então

$$A \leq W_{XY} \leq M_{XY} \leq B \quad \text{e} \quad W_{XY} \boxminus X = Y = M_{XY} \boxplus X. \quad (6.24)$$

Este teorema mostra que W_{XY} é o maior elemento (máximo) do conjunto das memórias $\boxminus$ -perfeitas e M_{XY} é o menor elemento (mínimo) do conjunto das memórias $\boxplus$ -perfeitas. Além disso, se existe uma memória $\boxminus$ -perfeita, então W_{XY} é também $\boxminus$ -perfeita. O mesmo vale para M_{XY} , substituindo $\boxminus$ por $\boxplus$.

O seguinte teorema fornece uma condição necessárias e suficientes para o perfeito armazenamento das memórias fundamentais em uma memória associativa morfológica no caso hetero-associativo [73, 76].

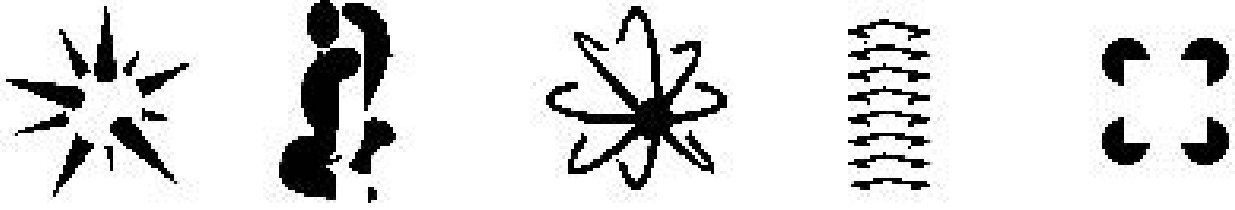


Fig. 6.1: Padrões recordados pela memória associativa morfológica M_{XY} após apresentarmos as chaves fundamentais $\mathbf{x}^1, \dots, \mathbf{x}^5$ como entrada.

Teorema 6.1.2. A matriz W_{XY} é $\boxtimes$ -perfeita para $(\mathbf{x}^\xi, \mathbf{y}^\xi)$, com $\xi = 1, \dots, k$, se e somente se, para cada $\xi = 1, \dots, k$, cada linha da matriz $[\mathbf{y}^\xi \times (\mathbf{x}^\xi)^*] - W_{XY}$ contém um elemento nulo. Por dualidade, a matriz M_{XY} é $\boxtimes$ -perfeita para $(\mathbf{x}^\xi, \mathbf{y}^\xi)$, com $\xi = 1, \dots, k$, se e somente se, para cada $\xi = 1, \dots, k$, cada linha da matriz $M_{XY} - [\mathbf{y}^\xi \times (\mathbf{x}^\xi)^*]$ contém um elemento nulo.

Note que o teorema 6.1.2 não contém nenhuma hipótese sobre o número máximo de memórias fundamentais. Logo, podemos armazenar quantos padrões desejarmos, desde que, cada linha de $\mathbf{y}^\xi \times (\mathbf{x}^\xi)^* - W_{XY}$ e $M_{XY} - \mathbf{y}^\xi \times (\mathbf{x}^\xi)^*$ contenha um elemento nulo para cada $\xi = 1, \dots, k$.

Exemplo 6.1.2. No exemplo 6.1.1, verificamos que $W_{XY} \boxtimes \mathbf{x}^\xi = \mathbf{y}^\xi = M_{XY} \boxtimes \mathbf{x}^\xi$, para $\xi = 1, 2, 3$, ou seja, W_{XY} e M_{XY} são memórias $\boxtimes$ e $\boxtimes$ -perfeitas, respectivamente. Logo, elas satisfazem as condições do teorema 6.1.2. Por exemplo, para $\xi = 2$, temos que

$$\begin{aligned} w_{11} &= y_1^2 - x_1^2 = [\mathbf{y}^2 \times (\mathbf{x}^2)^*]_{11}, & m_{13} &= y_1^2 - x_3^2 = [\mathbf{y}^2 \times (\mathbf{x}^2)^*]_{13}, \\ w_{22} &= y_2^2 - x_2^2 = [\mathbf{y}^2 \times (\mathbf{x}^2)^*]_{22}, & m_{22} &= y_2^2 - x_2^2 = [\mathbf{y}^2 \times (\mathbf{x}^2)^*]_{22}, \\ w_{31} &= y_3^2 - x_1^2 = [\mathbf{y}^2 \times (\mathbf{x}^2)^*]_{31}, & m_{33} &= y_3^2 - x_3^2 = [\mathbf{y}^2 \times (\mathbf{x}^2)^*]_{33}. \end{aligned} \quad (6.25)$$

Exemplo 6.1.3. Considere os padrões binários $\mathbf{x}^1, \dots, \mathbf{x}^5$ e $\mathbf{y}^1, \dots, \mathbf{y}^5$ apresentados nas figuras 1.1 e 1.2. Armazenamos estes padrões nas memórias associativas morfológicas M_{XY} e W_{XY} . Verificamos que $W_{XY} \boxtimes \mathbf{x}^\xi = \mathbf{y}^\xi$, para todo $\xi = 1, \dots, k$. Entretanto somente as equações $W_{XY} \boxtimes \mathbf{x}^1 = \mathbf{y}^1$ e $W_{XY} \boxtimes \mathbf{x}^5 = \mathbf{y}^5$ valem para a memória associativa morfológica M_{XY} . Na figura 6.1 apresentamos os padrões recordados pela memória associativa morfológica M_{XY} após apresentarmos as chaves fundamentais $\mathbf{x}^1, \dots, \mathbf{x}^5$ como entrada. O erro quadrático normalizado

$$\frac{\|\mathbf{y}^\xi - M_{XY} \boxtimes \mathbf{x}^\xi\|}{\|\mathbf{y}^\xi\|} \quad (6.26)$$

calculado sobre os padrões $\mathbf{x}^2, \mathbf{x}^3$ e $\mathbf{x}^4$ foi 0,068, 0,034 e 0,016, respectivamente.

Teorema 6.1.3. Sejam $X = [\mathbf{x}^1, \dots, \mathbf{x}^k] \in \{0, 1\}^{n \times k}$ e $Y = [\mathbf{y}^1, \dots, \mathbf{y}^k] \in \{0, 1\}^{m \times k}$ matrizes geradas aleatoriamente com probabilidade 1/2 de um elemento ser 0 ou 1. Se $W_{XY} = Y \boxtimes X^*$ e $M_{XY} = Y \boxtimes X^*$, então

$$\begin{aligned} Pr(W_{XY} \boxtimes X = Y) &= Pr(M_{XY} \boxtimes X = Y) \\ &= \left[1 - \frac{1}{4} \left(2 - \frac{3^{k-1} + 1}{4^{k-1}} \right) \left(1 - \left(\frac{1}{2} \right)^{k-1} \right)^{n-1} - \frac{1}{4} \left(1 - \left(\frac{3}{4} \right)^{k-1} \right) \left(1 - \left(\frac{7}{8} \right)^{k-1} \right)^{n-1} \right]^{mk}. \end{aligned}$$

Demonstração. Vamos demonstrar o teorema 6.1.3 somente para a memória associativa morfológica W_{XY} . O resultado para M_{XY} pode ser obtido de modo análogo. Durante a demonstração usaremos o fato de x_j^ξ e y_i^ξ serem variáveis aleatórias independentes e identicamente distribuídas para todo $i = 1, \dots, m, j = 1, \dots, n$ e $\xi = 1, \dots, k$.

Sabemos que $W_{XY} \boxtimes X \leq Y$ para todo X e Y . Assim, $W_{XY} \boxtimes X = Y$ se e somente se $W_{XY} \boxtimes X \geq Y$. A probabilidade de $W_{XY} \boxtimes X \geq Y$ será obtida através da probabilidade de $\bigvee_{j=1}^n (w_{1j} + x_j^1) \geq y_1^1$. O cálculo desta última será dividida em quatro casos distintos, disjuntos e equiprováveis. No final teremos $Pr(\bigvee_{j=1}^n (w_{1j} + x_j^1) = y_1^1) = (Pr(\text{Caso 1}) + Pr(\text{Caso 2}) + Pr(\text{Caso 3}) + Pr(\text{Caso 4}))/4$.

Caso (1). Se $x_1^1 = 1$ e $y_1^1 = 1$, teremos

$$\begin{aligned} Pr(w_{11} \geq 0) &= Pr\left(\bigwedge_{\xi=1}^k (y_1^\xi - x_1^\xi) \geq 0\right) = Pr\left(0 \wedge \bigwedge_{\xi=2}^k (y_1^\xi - x_1^\xi) \geq 0\right) = [Pr(y_1^\xi - x_1^\xi \geq 0)]^{k-1} = \left(\frac{3}{4}\right)^{k-1}, \\ Pr(w_{1j} + x_j^\xi \leq 0) &= Pr\left(\bigwedge_{\xi=1}^k (y_1^\xi - x_j^\xi + x_j^1) \leq 0\right) = Pr\left(1 \wedge \bigwedge_{\xi=2}^k (y_1^\xi - x_j^\xi + x_j^1) \leq 0\right) \\ &= 1 - Pr\left(\bigwedge_{\xi=2}^k (y_1^\xi - x_j^\xi + x_j^1) > 0\right) = 1 - [Pr((y_1^\xi - x_j^\xi + x_j^1) > 0)]^{k-1} = 1 - \left(\frac{1}{2}\right)^{k-1}, \\ Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) \geq 1\right) &= 1 - Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) < 1\right) = 1 - [Pr(w_{1j} + x_j^1 \leq 0)]^{n-1} = 1 - \left[1 - \left(\frac{1}{2}\right)^{k-1}\right]^{n-1}. \end{aligned}$$

Assim,

$$\begin{aligned} Pr\left(\bigvee_{j=1}^n (w_{1j} + x_j^1) \geq y_1^1\right) &= Pr\left((w_{11} + 1) \vee \left(\bigvee_{j=2}^n (w_{1j} + x_j^1)\right) \geq 1\right) \\ &= Pr((w_{11} + 1) \geq 1) + Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) \geq 1\right) - Pr((w_{11} + 1) \geq 1) Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) \geq 1\right) \\ &= 1 - \left(1 - \left(\frac{3}{4}\right)^{k-1}\right) \left(1 - \left(\frac{1}{2}\right)^{k-1}\right)^{n-1}. \end{aligned}$$

Caso (2). Se $x_1^1 = 1$ e $y_1^1 = 0$, teremos

$$\begin{aligned} Pr\left(\bigvee_{j=1}^n (w_{1j} + x_j^1) \geq y_1^1\right) &= Pr\left((w_{11} + 1) \vee \left(\bigvee_{j=2}^n (w_{1j} + x_j^1)\right) \geq 0\right) \\ &= Pr((w_{11} + 1) \geq 0) + Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) \geq 0\right) - Pr((w_{11} + 1) \geq 0) Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) \geq 0\right) \\ &= 1 + Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) \geq 0\right) - Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) \geq 0\right) = 1. \end{aligned}$$

Caso (3). Se $x_1^1 = 0$ e $y_1^1 = 1$, teremos

$$Pr(w_{11} \geq 1) = Pr\left(\bigwedge_{\xi=1}^k (y_1^\xi - x_1^\xi) \geq 1\right) = Pr\left(1 \wedge \bigwedge_{\xi=2}^k (y_1^\xi - x_1^\xi) \geq 1\right) = [Pr(y_1^\xi - x_1^\xi \geq 1)]^{k-1} = \left(\frac{1}{4}\right)^{k-1},$$

Assim,

$$\begin{aligned} Pr\left(\bigvee_{j=1}^n (w_{1j} + x_j^1) \geq y_1^1\right) &= Pr\left(w_{11} \vee \left(\bigvee_{j=2}^n (w_{1j} + x_j^1)\right) \geq 1\right) \\ &= Pr(w_{11} \geq 1) + Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) \geq 1\right) - Pr(w_{11} \geq 1) Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) \geq 1\right) \\ &= 1 - \left(1 - \left(\frac{1}{4}\right)^{k-1}\right) \left(1 - \left(\frac{1}{2}\right)^{k-1}\right)^{n-1}. \end{aligned}$$

Caso (4). Se $x_1^1 = 0$ e $y_1^1 = 0$, teremos

$$Pr(w_{11} \geq 0) = Pr\left(\bigwedge_{\xi=1}^k (y_1^\xi - x_1^\xi) \geq 0\right) = Pr\left(0 \wedge \bigwedge_{\xi=2}^k (y_1^\xi - x_1^\xi) \geq 0\right) = [Pr(y_1^\xi - x_1^\xi \geq 0)]^{k-1} = \left(\frac{3}{4}\right)^{k-1},$$

$$\begin{aligned} Pr(w_{1j} + x_j^\xi \leq -1) &= Pr\left(\bigwedge_{\xi=1}^k (y_1^\xi - x_j^\xi + x_j^1) \leq -1\right) = Pr\left(0 \wedge \bigwedge_{\xi=2}^k (y_1^\xi - x_j^\xi + x_j^1) \leq -1\right) \\ &= 1 - Pr\left(\bigwedge_{\xi=2}^k (y_1^\xi - x_j^\xi + x_j^1) > -1\right) = 1 - [Pr((y_1^\xi - x_j^\xi + x_j^1) \geq 0)]^{k-1} = 1 - \left(\frac{7}{8}\right)^{k-1}, \end{aligned}$$

$$Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) \geq 0\right) = 1 - Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) < 0\right) = 1 - [Pr(w_{1j} + x_j^1 \leq -1)]^{n-1} = 1 - \left[1 - \left(\frac{7}{8}\right)^{k-1}\right]^{n-1}.$$

Assim,

$$\begin{aligned} Pr\left(\bigvee_{j=1}^n (w_{1j} + x_j^1) \geq y_1^1\right) &= Pr\left(w_{11} \vee \left(\bigvee_{j=2}^n (w_{1j} + x_j^1)\right) \geq 0\right) \\ &= Pr(w_{11} \geq 0) + Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) \geq 0\right) - Pr(w_{11} \geq 0) Pr\left(\bigvee_{j=2}^n (w_{1j} + x_j^1) \geq 0\right) \\ &= 1 - \left(1 - \left(\frac{3}{4}\right)^{k-1}\right) \left(1 - \left(\frac{7}{8}\right)^{k-1}\right)^{n-1}. \end{aligned}$$

Finalmente teremos

$$\begin{aligned} Pr \left(\bigvee_{j=1}^n (w_{1j} + x_j^1) = y_1^1 \right) &= \frac{1}{4} \left[1 - \left(1 - \left(\frac{3}{4} \right)^{k-1} \right) \left(1 - \left(\frac{1}{2} \right)^{k-1} \right)^{n-1} + 1 \right. \\ &\quad \left. + 1 - \left(1 - \left(\frac{1}{4} \right)^{k-1} \right) \left(1 - \left(\frac{1}{2} \right)^{k-1} \right)^{n-1} + 1 - \left(1 - \left(\frac{3}{4} \right)^{k-1} \right) \left(1 - \left(\frac{7}{8} \right)^{k-1} \right)^{n-1} \right] \\ &= 1 - \frac{1}{4} \left(2 - \frac{3^{k-1} + 1}{4^{k-1}} \right) \left(1 - \left(\frac{1}{2} \right)^{k-1} \right)^{n-1} - \frac{1}{4} \left(1 - \left(\frac{3}{4} \right)^{k-1} \right) \left(1 - \left(\frac{7}{8} \right)^{k-1} \right)^{n-1}. \end{aligned}$$

Portanto,

$$\begin{aligned} Pr(W_{XY} \boxtimes X = Y) &= [Pr(W_{XY} \boxtimes \mathbf{x}^\xi = \mathbf{y}^\xi)]^k = \left[Pr \left(\bigvee_{j=1}^n (w_{ij} + x_j^\xi) = y_i^\xi \right) \right]^{mk} \\ &= \left[1 - \frac{1}{4} \left(2 - \frac{3^{k-1} + 1}{4^{k-1}} \right) \left(1 - \left(\frac{1}{2} \right)^{k-1} \right)^{n-1} - \frac{1}{4} \left(1 - \left(\frac{3}{4} \right)^{k-1} \right) \left(1 - \left(\frac{7}{8} \right)^{k-1} \right)^{n-1} \right]^{mk}. \end{aligned}$$

□

Exemplo 6.1.4. Considere o conjunto de padrões binários $\{(\mathbf{x}^\xi, \mathbf{y}^\xi), \mathbf{x}^\xi \in \{0, 1\}^{100}, \mathbf{y}^\xi \in \{0, 1\}^{80}, \xi = 1, \dots, k\}$ gerado aleatoriamente com distribuição uniforme. Armazenamos este conjunto de memórias fundamentais nas memórias associativas morfológicas W_{XY} e M_{XY} . Depois verificamos se $\mathbf{y}^\xi = W \boxtimes \mathbf{x}^\xi$ para todo $\xi = 1, \dots, k$. Repetimos o experimento 1000 vezes para diferentes valores de k (número de memórias fundamentais armazenados). Na figura 6.2 apresentamos com ∇ a probabilidade empírica de $\mathbf{y}^\xi = W \boxtimes \mathbf{x}^\xi$ e com $\triangle$ a probabilidade empírica de $\mathbf{y}^\xi = M \boxtimes \mathbf{x}^\xi$, para todo $\xi = 1, \dots, k$. A linha tracejada representa a probabilidade teórica dada pela equação (6.27). No eixo horizontal colocamos k , o número de memórias fundamentais. Note que as probabilidades empíricas $Pr(W \boxtimes X = Y)$ e $Pr(M \boxtimes X = Y)$ coincidiram e ambas são menores que 1 para $k \geq 4$. Este resultado mostra que as memórias associativas morfológicas binárias heteroassociativas não conseguem armazenar muitas memórias fundamentais.

O seguinte teorema, introduzido em [73, 76], caracteriza a recordação de um padrão $\mathbf{y}^\gamma$ após apresentarmos como entrada uma versão corrompida da chave fundamental $\mathbf{x}^\gamma$. Um resultado análogo pode ser obtido para a memória associativa M_{XY} usando o conceito de dualidade.

Teorema 6.1.4. *Seja $\tilde{\mathbf{x}}^\gamma$ uma versão ruidosa do padrão $\mathbf{x}^\gamma$. $W_{XY} \boxtimes \tilde{\mathbf{x}}^\gamma = \mathbf{y}^\gamma$ se e somente se*

$$\tilde{x}_j^\gamma \leq x_j^\gamma \vee \left[\bigwedge_{i=1}^m \bigvee_{\xi \neq \gamma}^k (y_i^\gamma - y_i^\xi + x_{j_i}^\xi) \right], \quad \forall j = 1, \dots, m, \quad (6.27)$$

e para cada índice de linha $i \in \{1, \dots, m\}$, existe um índice de coluna $j_i \in \{1, \dots, n\}$ tal que

$$\tilde{x}_{j_i}^\gamma = x_{j_i}^\gamma \vee \left[\bigvee_{\xi \neq \gamma}^k (y_i^\gamma - y_i^\xi + x_{j_i}^\xi) \right]. \quad (6.28)$$

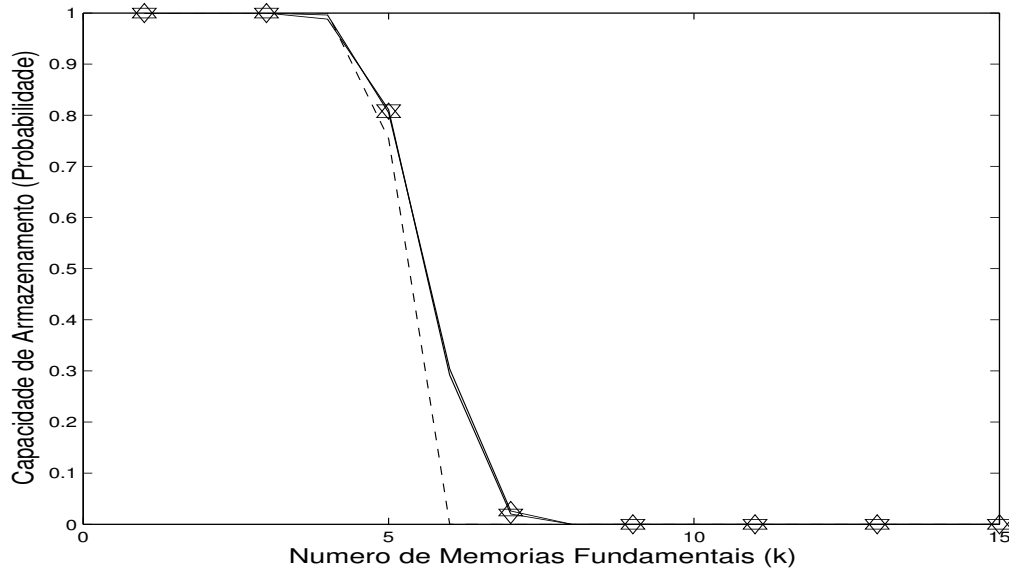


Fig. 6.2: A linha com ∇ representa a probabilidade empírica $\mathbf{y}^\xi = W \square \mathbf{x}^\xi$ para todo $\xi = 1, \dots, k$. A linha com $\triangle$ representa a probabilidade empírica de $\mathbf{y}^\xi = M \square \mathbf{x}^\xi$ para todo $\xi = 1, \dots, k$. A linha tracejada representa a probabilidade dada pela equação 6.27.

O seguinte teorema, introduzido por Sussner em [93], determina precisamente a ação da memória associativa morfológica W_{XY} no caso binário. Um resultado análogo pode ser obtido para W_{XY} usando o conceito de dualidade.

Teorema 6.1.5. *Seja $X \in \{0, 1\}^{n \times k}$, $Y \in \{0, 1\}^{m \times k}$ e $\mathbf{x} \in \{0, 1\}^n$. Se $\mathbf{0} \neq \mathbf{x} \leq \bigvee_{\xi=1}^k \mathbf{x}^\xi$ então*

$$W_{XY} \square \mathbf{x} = \bigvee_{t=1}^q \bigvee_{t \in \Theta_t} \mathbf{y}^\xi, \quad (6.29)$$

onde $\{\Theta_1, \Theta_2, \dots, \Theta_q\}$ é o conjunto de todos $\Theta_t \subseteq \{1, \dots, k\}$ tais que

$$\bigvee_{\xi \in \Theta_t} \mathbf{x}^\xi \geq \mathbf{x}. \quad (6.30)$$

Exemplo 6.1.5. Considere os padrões binários $\mathbf{x}^1, \dots, \mathbf{x}^5$ e $\mathbf{y}^1, \dots, \mathbf{y}^5$ apresentados nas figuras 1.1 e 1.2. Armazenamos estes padrões na memória associativa morfológica W_{XY} . No exemplo 6.1.3 verificamos que $W_{XY} \square \mathbf{x}^\xi = \mathbf{y}^\xi$ para todo $\xi = 1, \dots, k$. Considere agora os padrões $\tilde{\mathbf{x}}^1, \dots, \tilde{\mathbf{x}}^5$ apresentados na figura 6.3. Estes padrões foram obtidos introduzindo ruído salt (ruído erosivo) com probabilidade 0,3. Verificamos que as recordações fundamentais $\mathbf{y}^1, \dots, \mathbf{y}^5$ foram encontradas como saída após apresentarmos os padrões corrompidos $\tilde{\mathbf{x}}^1, \dots, \tilde{\mathbf{x}}^5$ como entrada. Repetimos o experimento 1000 vezes e verificamos que o erro quadrático médio normalizado

$$EQMN = E \left(\frac{1}{5} \sum_{\xi=1}^5 \frac{\|\mathbf{y}^\xi - W_{XY} \square \tilde{\mathbf{x}}^\xi\|}{\|\mathbf{y}^\xi\|} \right) = 0. \quad (6.31)$$



Fig. 6.3: Padrões corrompidos com ruído salt com probabilidade 0,3 usado no exemplo 6.1.5.

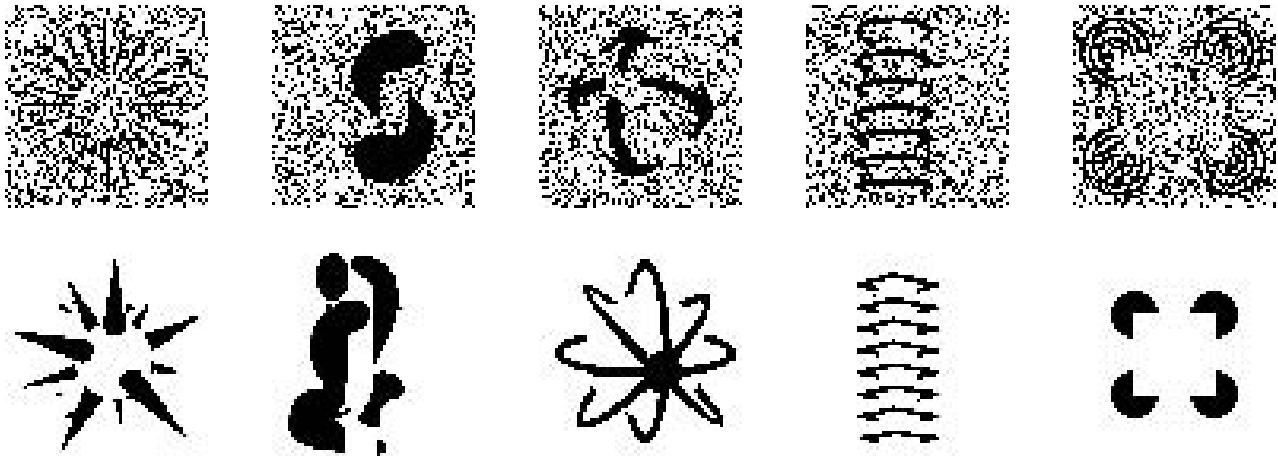


Fig. 6.4: Padrões corrompidos com ruído pepper com probabilidade 0,3 usados como entrada, abaixo os respectivos padrões recordados pela memória associativa morfológica M_{XY} .

Também armazenamos o conjunto das memórias fundamentais $\{(\mathbf{x}^\xi, \mathbf{y}^\xi), \xi = 1, \dots, 5\}$ na memória associativa M_{XY} . Utilizamos como entrada os padrões corrompidos com ruído pepper (ruído dilativo) com probabilidade 0,3 apresentados no topo da figura 6.4 e encontramos como resposta da memória os padrões apresentados na segunda linha da figura 6.4. O erro quadrático normalizado $\|\mathbf{y}^\xi - M_{XY} \boxtimes \tilde{\mathbf{x}}^\xi\| / \|\mathbf{y}^\xi\|$ para $\xi = 1, \dots, 5$ foi, respectivamente, 0,081, 0,068, 0,035, 0,016 e 0. Repetimos o experimento 1000 vezes e encontramos um $EQMN = 0,04$.

6.2 Memórias Auto-Associativas Morfológicas

Nesta seção consideramos o caso auto-associativo, isto é, $Y = X$. Uma discussão detalhada das memórias auto-associativas morfológica para padrões em tons de cinza encontra-se em [97]. Neste caso, as matrizes W_{XX} e M_{XX} são dadas por

$$w_{ij} = \bigwedge_{\xi=1}^k (x_i^\xi - x_j^\xi) \quad \text{e} \quad m_{ij} = \bigvee_{\xi=1}^k (x_i^\xi - x_j^\xi), \quad (6.32)$$

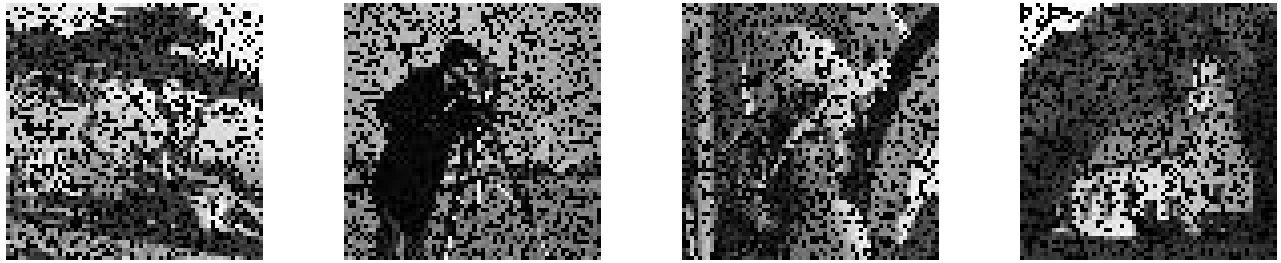


Fig. 6.5: Padrões corrompidos com ruído pepper (ruído erosivo) utilizados no exemplo 6.2.1.

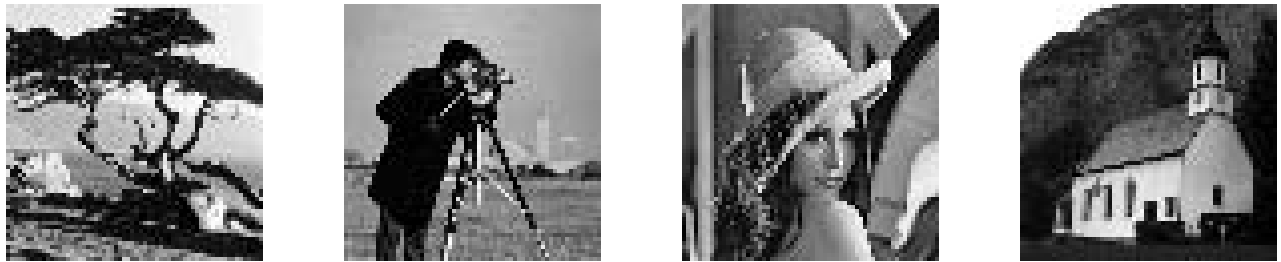


Fig. 6.6: Padrões recordados pela memória auto-associativa morfológica W_{XX} quando apresentamos os padrões da figura 6.5 como entrada.

e na fase de recordação usamos as equações (6.1) e (6.2), respectivamente.

A fase de recordação das memórias auto-associativas morfológicas são descritas em termos dos seus pontos fixos e regiões de recordação [90, 94]. Um resultado análogo pode ser obtido para a memória associativa morfológica M_{XX} usando o conceito de dualidade.

Teorema 6.2.1. *Para todo $X = [\mathbf{x}^1, \dots, \mathbf{x}^k] \in \mathbb{R}^{n \times k}$, o conjunto dos pontos fixos de W_{XX} inclui os padrões $\mathbf{x}^1, \dots, \mathbf{x}^k$. Além disso, para todo $\mathbf{x} \in \mathbb{R}^n$, temos*

$$W_{XX} \boxtimes \mathbf{x} = \hat{\mathbf{x}}, \quad (6.33)$$

onde $\hat{\mathbf{x}}$ denota o supremo de $\mathbf{x}$ no conjunto dos pontos fixos de W_{XX} .

O seguinte corolário afirma que a fase de recordação das memórias auto-associativas morfológicas é efetuada em um único passo.

Corolário. *Seja $X \in \mathbb{R}^n$. O conjunto dos pontos fixos de W_{XX} consiste de todos $W_{XX} \boxtimes \mathbf{x}$ tais que $\mathbf{x} \in \mathbb{R}^n$. Além disso, se $\mathbf{x}^\xi$, para $\xi = 1, \dots, k$ é o padrão recordado por W_{XX} após apresentarmos o padrão-chave $\mathbf{x}$, então $\mathbf{x} \leq \mathbf{x}^\xi$.*

Lembre-se que uma versão ruidosa $\tilde{\mathbf{x}}$ de um padrão $\mathbf{x}$ é uma versão *erodida* de $\mathbf{x}$ quando $\tilde{\mathbf{x}} \leq \mathbf{x}$ e é uma versão *dilatada* de $\mathbf{x}$ quando $\tilde{\mathbf{x}} \geq \mathbf{x}$. Usando esta terminologia temos que se $W_{XX} \boxtimes \tilde{\mathbf{x}}^\gamma = \mathbf{x}^\gamma$, então pelo teorema 6.2.1, $\tilde{\mathbf{x}}^\gamma$ deve ser uma versão erodida de $\mathbf{x}^\gamma$. Analogamente, se $M_{XX} \boxtimes \tilde{\mathbf{x}}^\gamma = \mathbf{x}^\gamma$, então $\tilde{\mathbf{x}}^\gamma$ deve ser uma versão dilatada de $\mathbf{x}^\gamma$.

Exemplo 6.2.1. Considere os padrões $\mathbf{x}^1, \dots, \mathbf{x}^4$ apresentados na figura 1.3. Armazenamos estes padrões na memória associativa morfológicas W_{XX} usando a equação (6.32). Depois geramos os

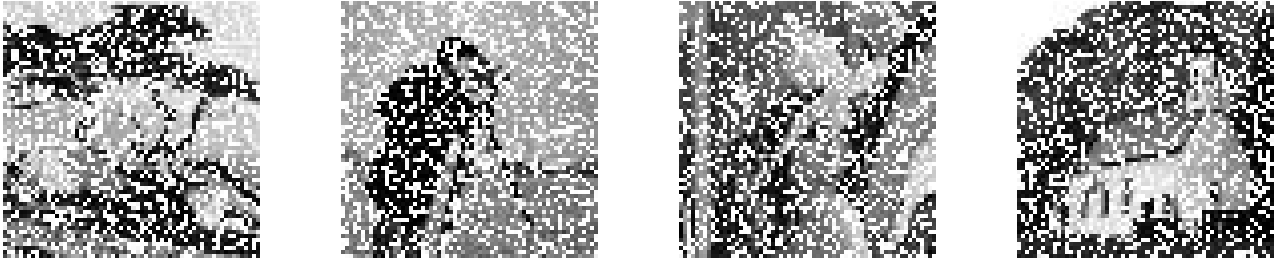


Fig. 6.7: Padrões corrompidos com ruído salt (ruído dilativo) utilizados no exemplo 6.2.1.

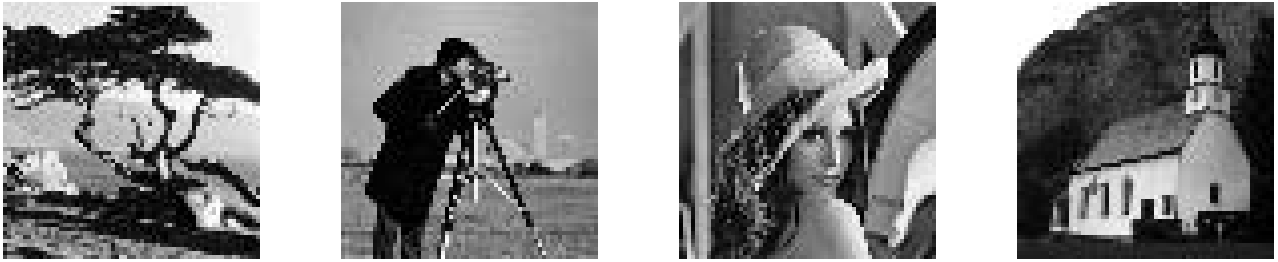


Fig. 6.8: Padrões recordados pela memória auto-associativa morfológica M_{XX} quando apresentamos os padrões da figura 6.7 como entrada.

padrões $\tilde{x}^1, \dots, \tilde{x}^4$ introduzindo *ruído pepper* (ruído erosivo) com distribuição uniforme e probabilidade 0,3. Na figura 6.5 apresentamos os padrões corrompidos $x^1, \dots, x^4$ e na figura 6.6 apresentamos os padrões recordados pela memória associativa morfológica W_{XX} após apresentarmos os padrões corrompidos $\tilde{x}^1, \dots, \tilde{x}^4$. Repetimos o experimento acima 1000 vezes e encontramos um erro quadrático médio normalizado (EQMN) de 0,0028.

Realizamos um experimento análogo usando a memória associativa morfológica M_{XX} e os padrões corrompidos com *ruído salt* (ruído dilativo) gerados com distribuição uniforme com probabilidade 0,3 apresentados na figura 6.7. Na figura 6.8 apresentamos os padrões recordados pela memória associativa morfológica M_{XX} após apresentarmos os padrões corrompidos com ruído dilativo da figura 6.7. Repetimos o experimento 1000 vezes e encontramos um $EQMN = 0,0039$.

Um padrão $\tilde{x}$ contendo ruído dilativo não pode ser recordado usando W_{XX} , pois pelo corolário acima, se $\tilde{x}_i^\gamma > x_i^\gamma$ para algum i , então $W_{XX} \boxtimes \tilde{x}^\gamma > x^\gamma$. Analogamente, um padrão $\hat{x}^\gamma$ contendo ruído erosivo não pode ser recordado usando M_{XX} pois se $\hat{x}_i^\gamma < x_i^\gamma$ para algum i , então $M_{XX} \boxtimes \hat{x}^\gamma < x^\gamma$.

Exemplo 6.2.2. Considere os padrões binários $x^1, \dots, x^{10} \in \{0, 1\}^{35}$ apresentados na figura 6.9. Armazenamos estes padrões na memória associativa morfológica W_{XX} e verificamos que todas as recordações fundamentais são pontos fixos. Depois, introduzimos ruído dilativo nas memórias fundamentais revertendo um único pixel do fundo da imagem obtendo os padrões apresentados na figura 6.10. Na figura 6.11 apresentamos os padrões recordados pela memória associativa morfológica W_{XX} quando apresentamos os padrões da figura 6.10 como entrada. Note que nenhuma das memórias fundamentais foi recordada. Por outro lado, a memória associativa morfológica M_{XX} é eficiente para recordar padrões corrompidos com ruído dilativo e todas as memórias fundamentais foram recordadas

por esta memória após apresentarmos os padrões da figura 6.10 como entrada.

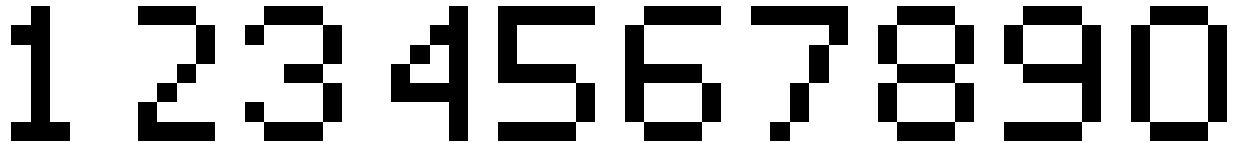


Fig. 6.9: Padrões binários $\mathbf{x}^1, \dots, \mathbf{x}^{10}$ usados nos exemplos 6.2.2 e 6.3.1.



Fig. 6.10: Padrões $\tilde{\mathbf{x}}^1, \dots, \tilde{\mathbf{x}}^{10}$ corrompidos com ruído dilativo.



Fig. 6.11: Padrões recordados pela memória associativa M_{XX} após apresentar os padrões da figura 6.10 como entrada.

6.3 Memórias Auto-associativas Morfológicas Binárias

Nesta seção discutimos as memórias associativas morfológicas binárias. Lembramos que um padrão é binário se $\mathbf{x}^\xi \in \{0, 1\}^n$ para todo $\xi = 1, \dots, k$. Sabemos que as memórias auto-associativas morfológicas podem armazenar infinitos padrões. No caso binário poderemos armazenar 2^n padrões, onde n é a dimensão do padrões.

Dado um conjunto de padrões binários $\{\mathbf{x}^1, \dots, \mathbf{x}^k\}$ e um índice l , denotaremos por L_l o maior subconjunto de $\{1, \dots, k\}$ tal que $x_l^\xi = 1$ para todo $\xi \in L_l$. Em outras palavras, o conjunto L_l diz

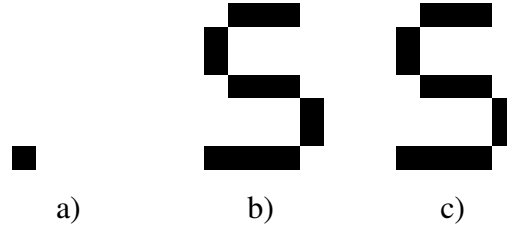


Fig. 6.12: Nos itens a), b) e c) temos e_{31} , $W_{XX} \boxtimes e_{31}$ e $x^5 \wedge x^9$, respectivamente.

quais são as memórias fundamentais x^ξ que possuem valor 1 na l -ésima componente. Usando esta notação podemos enunciar o seguinte teorema [94].

Teorema 6.3.1. *Seja y um ponto fixo de W_{XX} tal que $y_l = 1$ e $y \leq \bigwedge_{\xi \in L_l} x^\xi$, então $y = \bigwedge_{\xi \in L_l} x^\xi$.*

Corolário. *Se e_l é a l -ésima coluna da matriz identidade, então*

$$W_{XX} \boxtimes e_l = \bigwedge_{\xi \in L_l} x^\xi. \quad (6.34)$$

Exemplo 6.3.1. Considere os padrões $x^1, \dots, x^{10}$ apresentados na figura 6.9. Na figura 6.12 a) apresentamos o padrão $e_{31} \in \{0, 1\}^{35}$ que foi usado como entrada da memória W_{XX} . Na figura 6.12 b) apresentamos o padrão recordado pela memória auto-associativa morfológica, isto é $W_{XX} \boxtimes e_{31}$. Note que os padrões x^5 e x^9 possuem valor 1 na componente 31. Logo, $L_{31} = \{5, 9\}$. Na figura 6.12 c) apresentamos $\bigwedge_{\xi \in L_{31}} x^\xi = x^5 \wedge x^9$. Note que as imagens dos itens b) e c) são iguais, verificando a validade do corolário 6.3.

Na demonstração do corolário, usamos o fato de $\bigwedge_{\xi \in L_l} x^\xi$ ser um ponto fixo de W_{XX} . Na verdade, podemos mostrar um resultado muito mais forte. Para isso, vamos apresentar a seguinte definição [11]:

Definição 6.3.1. Uma expressão envolvendo $x^1, \dots, x^k$ e os símbolos $\vee$ e $\wedge$ é chamado *polinômio reticulado* em $x^1, \dots, x^k$.

Pela definição acima, temos que $\bigvee_{l=1}^p \bigwedge_{\xi \in L_l} x^\xi$ é a forma geral de um polinômio reticulado. Em particular, $\bigwedge_{\xi \in L_l} x^\xi$ é um polinômio reticulado em $x^1, \dots, x^k$ e é um ponto fixo de W_{XX} . O seguinte teorema relaciona o conjunto dos pontos fixos das memórias associativas morfológicas com os polinômios reticulados em $x^1, \dots, x^k$ [95].

Teorema 6.3.2 (Teorema dos Pontos Fixos das Memórias Associativas Morfológicas Binárias). *Seja $X = [x^1, \dots, x^k] \in \{0, 1\}^{n \times k}$. Um padrão binário y , diferente dos padrões 0 e 1, é um ponto fixo de W_{XX} se e somente se y é um polinômio reticulado em $x^1, \dots, x^k$. Analogamente, um padrão binário $z \neq 0, 1$ é um ponto fixo de M_{XX} se e somente se z é um polinômio reticulado em $x^1, \dots, x^k$.*

O teorema 6.3.2 mostra que uma memória auto-associativa morfológica binária tem grande probabilidade de ter um grande número de estados espúrios. Com base nos teoremas apresentados anteriormente podemos dizer:

1. A capacidade absoluta das memórias auto-associativas binárias é 2^n se n for a dimensão dos padrões armazenados.
2. Todo padrão recordado é um ponto fixo da memória (convergência com uma única iteração).
3. Os pontos fixos de W_{XX} e M_{XX} incluem os padrões originais bem como um grande número de estados espúrios.
4. A memória W_{XX} exibe uma excelente tolerância com respeito a padrões erodidos e M_{XX} , com respeito a padrões dilatados.
5. W_{XX} e M_{XX} não são eficientes na recordação de padrões corrompidos por ambos ruído dilati-
vos e erosivos.

Na próxima seção apresentaremos um modelo com características melhores com respeito aos itens 3, 4 e 5, que mantém as características dos itens 1 e 2.

6.4 Memórias Associativas Morfológicas de Duas Camadas

As memórias W_{XY} e M_{XY} possuem excelente tolerância com respeito a padrões corrompidos com ruído somente erosivo ou somente dilativo, respectivamente, mas não são eficientes quando o padrão apresentado possui ambos ruídos. Para evitar este problema, podemos fazer combinações das memórias W_{XY} e M_{XY} . Estas combinações produzem as *memórias associativas morfológicas de duas camadas*. Como veremos, as memórias associativas de duas camadas possuem um número menor de pontos fixos, aumentando assim a tolerância a ruídos e diminuindo o número de memórias espúrias.

6.4.1 Arquitetura

Primeiramente vamos lembrar a definição do núcleo de uma memória associativa morfológica W .

Definição 6.4.1. Sejam $X \in \{0, 1\}^{n \times k}$ e $Y \in \{0, 1\}^{m \times k}$. Uma matriz $Z = [z^1, z^2, \dots, z^k]$ de dimensões $n \times k$ é um *núcleo* para (X, Y) se e somente se existe uma memória W tal que

$$W \boxtimes (M_{ZZ} \boxtimes X) = Y. \quad (6.35)$$

Se $X = Y$, dizemos que Z é um núcleo para X .

Suponha que existe Z (núcleo) tal que $W_{ZY} \boxtimes (M_{ZZ} \boxtimes X) = Y$. Como $M_{XX} \boxtimes X = X$ para todo $X \in \{0, 1\}^{n \times k}$, então podemos dizer que:

$$W_{ZY} \boxtimes (M_{ZZ} \boxtimes (M_{XX} \boxtimes X)) = Y. \quad (6.36)$$

Assim, dado um padrão $\mathbf{x}$, encontramos

$$\mathbf{y} = W_{ZY} \boxtimes (M_{ZZ} \boxtimes (M_{XX} \boxtimes \mathbf{x})) = W_{ZY} \boxtimes ((M_{ZZ} \boxtimes M_{XX}) \boxtimes \mathbf{x}) = W_{ZY} \boxtimes (M \boxtimes \mathbf{x}), \quad (6.37)$$

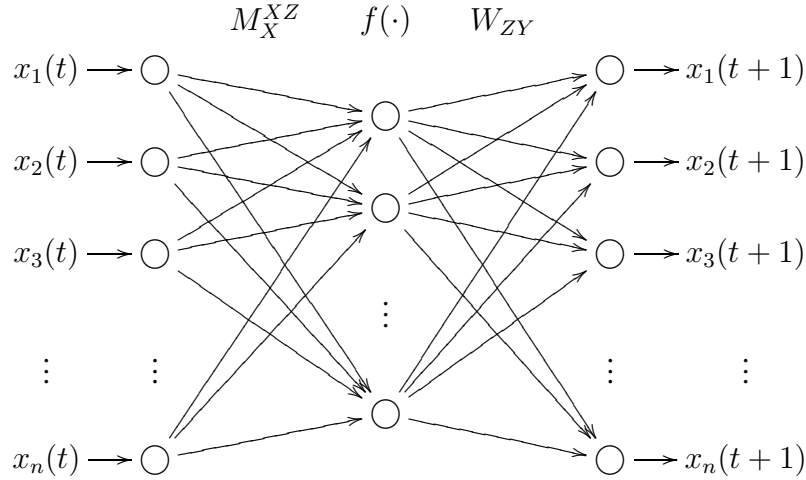


Fig. 6.13: Arquitetura das Memória Associativa de Duas Camadas descrita pela equação (6.40).

onde M é uma matriz que depende de Z e X . Na dedução acima, $M = M_{ZZ} \boxtimes M_{XX}$, mas faremos uma pequena mudança na definição da matriz M .

Nossa dificuldade está em encontrar um núcleo para (X, Y) . Sabemos que o núcleo depende somente de X . Precisamente, devemos ter $Z \leq X$, com $\mathbf{z}^\gamma \wedge \mathbf{z}^\xi = \mathbf{0}$ e $\mathbf{z}^\gamma \not\leq \mathbf{x}^\xi$ para todo $\xi, \gamma, \xi \neq \gamma$. Entretanto, vamos supor apenas que $\mathbf{z}^\gamma \wedge \mathbf{z}^\xi = \mathbf{0}$ e $\mathbf{z}^\gamma \not\leq \mathbf{z}^\xi$ para todo $\gamma, \xi, \xi \neq \gamma$. Note que Z não depende mais de X . Se $X \in \{0, 1\}^{n \times k}$, $Y \in \{0, 1\}^{m \times k}$ e $Z \in \{0, 1\}^{p \times k}$, então W_{ZY} e M_{XX} são matrizes $m \times p$ e $n \times n$ respectivamente. Para manter uma coerência nas operações, M deve ser uma matriz $p \times n$. Logo, tomaremos

$$M = M_X^{XZ} = M_{XZ} \boxtimes M_{XX}. \quad (6.38)$$

Assim, dado um padrão-chave $\mathbf{x}$, tomamos

$$\mathbf{y} = W_{ZY} \boxtimes (M_X^{XZ} \boxtimes \mathbf{x}), \quad (6.39)$$

como sendo o padrão recordado.

A operação $M_X^{XZ} \boxtimes \mathbf{x}$ pode produzir como resposta um vetor que não pertence ao conjunto $\{0, 1\}^n$ e isso pode ser ruim quando trabalhamos com memórias binárias. Introduzimos uma função limiar aplicada no resultado de $M_X^{XZ} \boxtimes \mathbf{x}$ na equação (6.39) para evitar este problema. Precisamente, definimos a memória associativa:

$$\mathbf{y} = W_{ZY} \boxtimes f(M_X^{XZ} \boxtimes \mathbf{x}), \quad (6.40)$$

onde $M_X^{XZ} = M_{XZ} \boxtimes M_{XX}$ e $f(\mathbf{x}) = [f_1(\mathbf{x}), f_2(\mathbf{x}), \dots, f_p(\mathbf{x})]'$ com

$$f_i(\mathbf{x}) = \begin{cases} 1 & \text{se } x_i > 0 \\ 0 & \text{caso contrário.} \end{cases} \quad (6.41)$$

Na figura 6.4.1 apresentamos a arquitetura da rede deste modelo neural. O modelo dual pode ser obtido de um modo análogo.

6.4.2 Aprendizado

O aprendizado da memória associativa morfológica de duas camadas consiste em escolher $Z = [z^1, z^2, \dots, z^k] \in \{0, 1\}^{p \times k}$ com $z^\xi \wedge z^\gamma = \mathbf{0}$ e $z^\xi \not\leq z^\gamma, \forall \xi, \gamma, \xi \neq \gamma$ e construir as matrizes $W_{ZY} = Y \boxtimes Z^*$ e $M_X^{XZ} = M_{XZ} \boxtimes M_{XX} = (Z \boxtimes X^*) \boxtimes (X \boxtimes X^*)$. A escolha de Z é arbitrária desde que satisfaça as condições apresentadas acima. Note que $I = [e_1, e_2, \dots, e_k] \in \{0, 1\}^{p \times k}$, onde e_i é a i -ésima base do espaço $\mathbb{R}^p$, pode ser usada como a matriz Z neste modelo. Neste trabalho usaremos somente $Z = I_{p \times k}$.

Teorema 6.4.1 (Teorema dos Pontos Fixos da Memória Associativa Morfológica de Duas Camadas). *Seja $X \in \{0, 1\}^{n \times k}$ tal que $M_{XX} \in \{0, 1\}^{n \times k}$. Seja $Y \in \{0, 1\}^{m \times k}$, $Z = [z^1, z^2, \dots, z^k] \in \{0, 1\}^{p \times k}$ e $\mathbf{x} \in \{0, 1\}^n$ arbitrário. Usaremos a notação $\Theta = \{\xi \mid \mathbf{x} \geq \mathbf{x}^\xi\}$ e $M_X^{XZ} = M_{XZ} \boxtimes M_{XX}$. Se $\Theta \neq \emptyset$, $\mathbf{z}^\gamma \geq \mathbf{z}^\xi$ e $\mathbf{z}^\gamma \wedge \mathbf{z}^\xi = \mathbf{0}$ para todo ξ, γ com $\gamma \neq \xi$. Então, para todo $\mathbf{x} \neq \mathbf{1} = [1, 1, \dots, 1]' \in \{0, 1\}^n$, temos*

$$W_{ZY} \boxtimes f(M_X^{XZ} \boxtimes \mathbf{x}) = \bigvee_{\xi \in \Theta} \mathbf{y}^\xi. \quad (6.42)$$

Além disso, para todo $\xi = 1, 2, \dots, k$ e todo $\mathbf{x} \in \{0, 1\}^n$,

$$M_{XX} \boxtimes \mathbf{x} = \mathbf{x}^\xi \Rightarrow W_{ZY} \boxtimes f(M_X^{XZ} \boxtimes \mathbf{x}) = \mathbf{y}^\xi. \quad (6.43)$$

A demonstração deste teorema encontra-se em [94]. A memória associativa morfológica de duas camadas apresentada acima possui menos estados espúrios que o modelo W_{XX} e portanto, possui uma melhor tolerância a ruído. Um resultado análogo pode ser obtido para a versão dual da memória associativa morfológica de duas camadas usando o conceito de dualidade.

Exemplo 6.4.1. Considere os padrões binários $\mathbf{x}^1, \dots, \mathbf{x}^5$ e $\mathbf{y}^1, \dots, \mathbf{y}^5$ apresentados nas figuras 1.1 e 1.2. Armazenamos estes padrões na versão dual da memória associativa morfológica de duas camadas usando $Z = I_{4096 \times 5}$ e verificamos que $\mathbf{y}^\xi = W_{ZY} \boxtimes f(M_X^{XZ} \boxtimes \mathbf{x}^\xi)$ para todo $\xi = 1, \dots, k$. Depois verificamos que as recordações fundamentais $\mathbf{y}^1, \dots, \mathbf{y}^5$ são recordadas após apresentarmos como entrada os padrões $\tilde{\mathbf{x}}^1, \dots, \tilde{\mathbf{x}}^5$ apresentados na figura 6.3. Repetimos o experimento 1000 vezes e verificamos que o erro quadrático médio normalizado

$$EQMN = E \left(\frac{1}{5} \sum_{\xi=1}^5 \frac{\|\mathbf{y}^\xi - M_{XY} \boxtimes \tilde{\mathbf{x}}^\xi\|}{\|\mathbf{y}^\xi\|} \right) = 0. \quad (6.44)$$

Finalmente, utilizamos com entrada os padrões corrompidos com ruído pepper (ruído dilativo) com probabilidade 0,3 apresentados no topo da figura 6.4. Encontramos como resposta da memória memória associativa de duas camadas os padrões apresentados na figura 6.14. O erro quadrático normalizado para $\xi = 1, \dots, 5$ foi, respectivamente, 0,081, 0,068, 0,035, 0,016 e 0. Repetimos o experimento 1000 vezes e encontramos um $EQMN = 0,04$.



Fig. 6.14: padrões recordados pela memória associativa morfológica de duas camadas após apresentarmos os padrões corrompidos apresentados no topo da figura 6.4

Capítulo 7

Desempenho das Memórias Associativas Binárias

Neste capítulo apresentamos comparações do desempenho das memórias associativas binárias introduzidas nos capítulos 5 e 6. Inspirados nas 5 características desejáveis em uma memória associativa apresentadas no capítulo 3, formalizamos os conceitos: *capacidade de armazenamento*, *distribuição da informação*, *raio de atração* e *probabilidade de memória espúria*. É importante observar que o conceito de raio de atração foi utilizado de um modo empírico em [44]. Nesta dissertação apresentamos uma definição rigorosa deste conceito. O conceito de capacidade de armazenamento apresentado nesta dissertação de mestrado difere do conceito de *capacidade absoluta* introduzido em [56]. Discutimos a diferença entre estes dois conceitos na seção 7.1. Não encontramos na literatura um trabalho discutindo o esforço computacional das memórias associativas neurais. Também não encontramos um trabalho mencionando os conceitos de distribuição da informação e a probabilidade de memória espúria de uma memória associativa.

Não discutimos neste capítulo o modelo BSB devido a semelhança deste com o modelo de Hopfield.

7.1 Capacidade de Armazenamento

A capacidade de armazenamento, denotada por $\mathcal{C}_a$, é uma função que diz a probabilidade de armazenarmos um conjunto de memórias fundamentais numa memória associativa. Esta função depende do mapeamento associativo (memória associativa), da dimensão dos padrões de entrada (n), da dimensão dos padrões de saída (m) e do número de memórias fundamentais (k). Usaremos a notação $\mathcal{C}_a(\text{Memória}, n, m; k)$ no caso heteroassociativo e $\mathcal{C}_a(\text{Memória}, n; k)$ no caso auto-associativo.

Seja

$$\Omega = \{(X, Y) : X = [\mathbf{x}^1, \mathbf{x}^2, \dots] \in \mathbb{R}^{n \times \infty}, Y = [\mathbf{y}^1, \mathbf{y}^2, \dots] \in \mathbb{R}^{m \times \infty}\}, \quad (7.1)$$

o conjunto dos pares (X, Y) de matrizes X (gerada aleatoriamente) com n linhas e infinitas colunas e das matrizes Y (geradas aleatoriamente) com m linhas e infinitas colunas. Considere a variável aleatória $C : \Omega \rightarrow \mathbb{N}$ dada por

$$C = \max\{p \in \mathbb{N} : G_p(\mathbf{x}^\xi) = \mathbf{y}^\xi, \quad \forall \xi = 1, \dots, p\}, \quad (7.2)$$

onde $G_p : \mathbb{R}^n \longrightarrow \mathbb{R}^m$ é o mapeamento associativo da memória associativa neural treinada usando os pares $(\mathbf{x}^\xi, \mathbf{y}^\xi)$ para $\xi = 1, \dots, p$ (primeiras p colunas de X e Y).

Definição 7.1.1 (Capacidade de Armazenamento). Considere o espaço amostral Ω definido na equação (7.1) e a variável aleatória C definida pela equação (7.2). A *capacidade de armazenamento* de uma memória associativa é a função dada pela equação

$$C_a(\text{Memória}, m, n; k) := 1 - F_c(k) = Pr(C > k), \quad (7.3)$$

onde F_c é a *função de distribuição* da variável aleatória C .

Note que $Pr(C > k)$ é a probabilidade de $G_k(\mathbf{x}^\xi) = \mathbf{y}^\xi$ para todo $\xi \leq k$. Logo,

$$C_a(\text{Memória}, m, n; k) = Pr(G_k(\mathbf{x}^\xi) = \mathbf{y}^\xi, \xi = 1, \dots, k), \quad (7.4)$$

onde G_k é o mapeamento associativo obtido após armazenar os pares $(\mathbf{x}^\xi, \mathbf{y}^\xi)$ para $\xi = 1, \dots, k$.

Podemos estimar a capacidade de armazenamento de uma memória associativa usando o conceito de *função de distribuição empírica* em conjunto com o *teorema de Glivenko-Cantelli* [10]. A função de distribuição empírica para variáveis aleatórias $C_1, C_2, \dots, C_s$ é a função $F_E(x)$ dada por

$$F_E(x) = \frac{1}{s} \sum_{l=1}^s Ind[C_l \leq x], \quad (7.5)$$

onde Ind é a função indicadora dada por

$$Ind[C_l \leq x] = \begin{cases} 1 & \text{se } C_l \leq x, \\ 0 & \text{caso contrário,} \end{cases} \quad (7.6)$$

O teorema de Glivenko-Cantelli afirma que se $C_1, C_2, \dots$ são variáveis aleatórias independentes com uma função de distribuição comum F_C , então $\sup_x |F_E(x) - F_C(x)| \longrightarrow 0$ com probabilidade 1. O teorema de Glivenko-Cantelli implica que, com probabilidade 1, $\lim_s F_E(x) = F_C(x)$, para cada x onde F_C é contínua. Assim, com probabilidade 1, teremos

$$C_a(\text{Memória}, n; k) = 1 - \lim_s F_E(k) = \lim_s (1 - F_E(k)) = \lim_s \frac{1}{s} \sum_{l=1}^s Ind[C_l > k]. \quad (7.7)$$

O conceito da capacidade de armazenamento foi inspirado no conceito de capacidade absoluta introduzido por McEliece *et. al.* em [56]. A *capacidade absoluta* é definida como $\max\{x \in \mathbb{N} : C_a(\text{Memória}, m, n; x) \geq 1 - \epsilon\}$, com $\epsilon > 0$ pequeno. Por exemplo, o teorema 5.1.2 apresenta a capacidade absoluta da rede de Hopfield e a conjectura 5.2.1 apresenta a capacidade absoluta da BAM. Acreditamos que existe uma perda de informação na medida capacidade absoluta como apresentado no exemplo abaixo.

Exemplo 7.1.1. Considere duas memórias associativas A e B . Na figura 7.1 apresentamos com linha contínua a capacidade de armazenamento da memória associativa A e com linha tracejada a capacidade de armazenamento da memória associativa B . A linha vertical pontilhada representa a capacidade absoluta da memória associativa A e a linha vertical com ponto-traço representa a capacidade

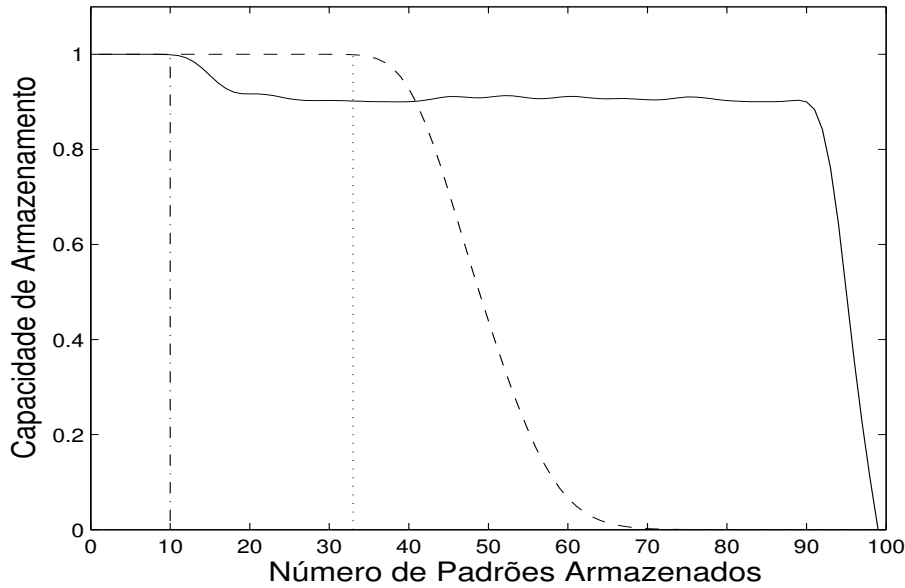


Fig. 7.1: Capacidade de armazenamento das memórias associativas A e B do exemplo 7.1.1. A linha contínua representa capacidade de armazenamento da memória associativa A e a linha tracejada representa a capacidade de armazenamento da memória associativa B . A linha vertical com pontilhada representa a capacidade absoluta da memória associativa A e a linha vertical com ponto-traço representa a capacidade de armazenamento da memória associativa B .

de armazenamento da memória associativa B . Note que a memória associativa A tem probabilidade 1 de armazenar um conjunto com menos de 33 memórias fundamentais, mas tem probabilidade próxima de zero para armazenar um conjunto com mais de 70 memórias fundamentais. Por outro lado, a memória associativa B tem probabilidades 0,911 e 0,9 de armazenar um conjunto com 33 e 90 memórias fundamentais, respectivamente. Neste exemplo, a capacidade absoluta da memória A é maior que a capacidade absoluta da memória B . Entretanto, este resultado não é suficiente para afirmar que a memória A pode armazenar mais padrões que B . De fato, temos probabilidade 0,9 de armazenar um conjunto com 90 memórias fundamentais em B . Por outro lado, certamente não poderemos armazenar o mesmo conjunto em A .

7.1.1 Memória Associativa de Hopfield

Seja $X = [\mathbf{x}^1, \dots, \mathbf{x}^k] \in \{-1, 1\}^{n \times k}$ a matriz das memórias fundamentais gerada aleatoriamente com distribuição uniforme onde $Pr(x_i^\xi = 1) = 1/2$ para todo $i = 1, \dots, n$ e $\xi = 1, \dots, k$. Daniel Amit mostrou em [3] que

$$Pr(x_i^1 = \text{sign}(W\mathbf{x}^1)_i) = \frac{1}{2} \left[1 + \text{erf} \left(\sqrt{\frac{n}{2k}} \right) \right], \quad (7.8)$$

onde erf é a função erro definida pela equação

$$\text{erf}(x) = \frac{2}{\sqrt{\pi}} \int_0^x e^{-t^2} dt. \quad (7.9)$$

Logo, a capacidade de armazenamento da memória associativa de Hopfield é

$$\mathcal{C}_a(\text{MA Hofield}, n; k) = \Pr(\text{sinal}(W\mathbf{x}^\xi) = \mathbf{x}^\xi, \xi = 1, \dots, k) \quad (7.10)$$

$$= \left[\Pr(\text{sinal}(W\mathbf{x}^\xi)_i = x_i^\xi) \right]^{kn} \quad (7.11)$$

$$= \left\{ \frac{1}{2} \left[1 + \text{erf} \left(\sqrt{\frac{n}{2p}} \right) \right] \right\}^{kn}. \quad (7.12)$$

Na figura 5.1 apresentamos o gráfico da capacidade de armazenamento pelo número de memórias fundamentais. A linha tracejada representa a capacidade de armazenamento obtida pela equação (7.12) e a linha contínua com $\square$ representa a capacidade de armazenamento estimada usando a equação (7.7). Neste exemplo usamos $n = 100$ e $s = 1000$ para estimar a função de distribuição empírica.

7.1.2 Memória Associativa Bidirecional

Substituindo o resultado da equação (5.17) na equação (5.23) obtemos

$$\Pr(\mathbf{y}^\xi = \text{sinal}(W\mathbf{x}^\xi), \forall \xi \in \{1, \dots, k\}) = \left\{ \frac{1}{2} \left[1 + \text{erf} \left(\sqrt{\frac{n}{2k}} \right) \right] \right\}^{km}. \quad (7.13)$$

Analogamente,

$$\Pr(\mathbf{x}^\xi = \text{sinal}(W^T\mathbf{y}^\xi), \forall \xi \in \{1, \dots, k\}) = \left\{ \frac{1}{2} \left[1 + \text{erf} \left(\sqrt{\frac{m}{2k}} \right) \right] \right\}^{kn}. \quad (7.14)$$

Logo,

$$\mathcal{C}_a(\text{BAM}, n, m; k) = \left\{ \frac{1}{2} \left[1 + \text{erf} \left(\sqrt{\frac{n}{2k}} \right) \right] \right\}^{km} \left\{ \frac{1}{2} \left[1 + \text{erf} \left(\sqrt{\frac{m}{2k}} \right) \right] \right\}^{kn}. \quad (7.15)$$

Na figura 5.5 apresentamos com linha tracejada a capacidade de armazenamento da BAM dada pela equação (7.15) e com linha contínua com $\square$ a capacidade de armazenamento estimada usando a equação (7.7). Neste exemplo tomamos $n = 100$, $m = 80$ e usamos $s = 1000$ para calcular a função de distribuição empírica.

7.1.3 Memória Associativa de Personnaz

A matriz dos pesos sinápticos da memória associativa de Personnaz é obtida através da equação (4.8) que sempre terá solução no caso auto-associativo ($W = I$ é uma solução). Logo, a capacidade de armazenamento desta memória associativa é

$$\mathcal{C}_a(\text{MA Personnaz}, n; k) = 1. \quad (7.16)$$

7.1.4 Memória Associativa de Kanter-Sompolinsky

A proposição 3 garante que todas as recordações fundamentais são pontos fixos da memória associativa de Kanter-Sompolinsky. Portanto

$$\mathcal{C}_a(\text{MA Kanter}, n; k) = 1. \quad (7.17)$$

7.1.5 Memória Associativa Bidirecional Assimétrica

As matrizes dos pesos sinápticos da ABAM são obtidos usando o armazenamento por projeção. Pelo teorema 4.2.1, o erro da dependencia linear será nulo somente se o vetores $\mathbf{x}^1, \dots, \mathbf{x}^k$ forem linearmente independentes. Entretanto, a probabilidade de um conjunto $\{\mathbf{x}^1, \dots, \mathbf{x}^k\}$ ser linearmente independente é 1 se $k \leq n$ e 0 se $k > n$. Analogamente para os vetores $\mathbf{y}^1, \dots, \mathbf{y}^k$. Podemos afirmar que

$$\mathcal{C}_a(\text{ABAM}, n, m; k) = 1 - f(k - (n \wedge m)), \quad (7.18)$$

onde f é a função limiar definida na equação (6.41).

7.1.6 Memória Associativa com Capacidade Exponencial

Chiueh e Goodman mostraram em [15] que

$$Pr \left(x_i^\eta \neq \text{sinal} \left(\sum_{\xi}^k \exp [< \mathbf{x}^\xi, \mathbf{x}^\eta >] x_i^\xi \right) \right) \leq \frac{e^{-T}}{\sqrt{4\pi T}}, \quad (7.19)$$

para $\eta \in \{1, \dots, k\}$, $i \in \{1, \dots, k\}$ e T grande. Logo,

$$\mathcal{C}_a(\text{ECAM}, n, m; k) = Pr \left(\mathbf{x}^\xi = \text{sinal} \left(\sum_{\xi}^k \exp [< \mathbf{x}^\xi, \mathbf{x}^\eta >] \mathbf{x}^\xi \right), \forall \xi \in \{1, \dots, k\} \right) \quad (7.20)$$

$$\geq \left(1 - \frac{e^{-T}}{\sqrt{4\pi T}} \right)^{nk} = c^{kn}, \quad (7.21)$$

onde c é uma constante próxima de 1.

7.1.7 Memória Associativa Bidirecional com Capacidade Exponencial

Podemos interpretar a BECAM como uma ECAM usando a equação (5.50). Concluimos que a capacidade de armazenamento da BECAM satisfaz

$$\mathcal{C}_a(\text{BECAM}, n, m; k) \geq c^{k(n+m)}, \quad (7.22)$$

onde c é uma constante próxima de 1.

7.1.8 Memórias Associativa Morfológicas

No teorema 6.2.1 mostramos que $W_{XX} \boxtimes X = M_{XX} \boxtimes X = X$. Logo, a capacidade de armazenamento das memórias auto-associativas morfológicas W_{XX} e M_{XX} é

$$\mathcal{C}_a(W_{XX}, n; k) = \mathcal{C}_a(M_{XX}, n; k) = 1. \quad (7.23)$$

Pelo teorema 6.1.3 concluímos que

$$\begin{aligned} \mathcal{C}_a(W_{XY}, n, m; k) &= \mathcal{C}_a(M_{XY}, n, m; k) \\ &= \left[1 - \frac{1}{4} \left(2 - \frac{3^{k-1} + 1}{4^{k-1}} \right) \left(1 - \left(\frac{1}{2} \right)^{k-1} \right)^{n-1} - \frac{1}{4} \left(1 - \left(\frac{3}{4} \right)^{k-1} \right) \left(1 - \left(\frac{7}{8} \right)^{k-1} \right)^{n-1} \right]^{mk}. \end{aligned} \quad (7.24)$$

Na figura 6.27 apresentamos o gráfico da capacidade de armazenamento pelo número de memórias fundamentais armazenadas. Este gráfico foi construído tomando $n = 100$, $m = 80$ e repetindo 1000 vezes cada experimento para obter as probabilidades. As linhas contínuas com ∇ e $\triangle$ representam as estimativas da capacidade de armazenamento através da equação (7.7) para as memórias associativas W_{XY} e M_{XY} , respectivamente. A linha tracejada representa a estimativa dada pela equação (7.24).

7.1.9 Memória Associativa Morfológica de Duas Camadas

O teorema 6.4.1 afirma que $W_{ZY} \boxtimes f(M_X^{XZ} \boxtimes \mathbf{x}) = \mathbf{y}^\xi$ sempre que $M_{XX} \boxtimes \mathbf{x} = \mathbf{x}^\xi$. Como $W_{XX} \boxtimes \mathbf{x}^\xi = \mathbf{x}^\xi$ para todo $\xi = 1, \dots, k$, então a capacidade de armazenamento das memórias associativas morfológicas de duas camadas é

$$\mathcal{C}_a(\text{MAM Duas Camadas}, n, m; k) = 1. \quad (7.25)$$

7.2 Distribuição da Informação

A *Função Distribuição da Informação* (FDI) é uma medida da distribuição da informação armazenada numa memória associativa neural. Esta medida é uma função que depende da memória, da dimensão dos padrões de entrada e saída, do número de memórias fundamentais armazenadas e da porcentagem do número das conexões sinápticas excluídas da memória associativa neural.

Considere o espaço amostral $\Omega = \{(X, Y) | X = [\mathbf{x}^1, \dots, \mathbf{x}^k] \in \mathbb{R}^{n \times k}, Y = [\mathbf{y}^1, \dots, \mathbf{y}^k] \in \mathbb{R}^{m \times k}\}$ e defina sobre Ω a variável aleatória $D : \Omega \rightarrow [0, 100]$ dada por

$$D = \max\{x \in [0, 100] : G_k^x(\mathbf{x}^\xi) = \mathbf{y}^\xi, \forall \xi = 1, \dots, k\}, \quad (7.26)$$

onde $G_k^x : \mathbb{R}^n \rightarrow \mathbb{R}^m$ é o mapeamento associativo da memória associativa neural treinada com as memórias fundamentais $(\mathbf{x}^\xi, \mathbf{y}^\xi)$, $\xi = 1, \dots, k$ e com $x\%$ do número total de conexões sinápticas excluídas (primeiro treinamos a memória associativa e depois excluimos aleatoriamente as conexões sinápticas). Note que $G_k = G_k^0$ é o mapeamento associativo com todas as conexões sinápticas.

Definição 7.2.1 (Função da Distribuição da Informação). A função da distribuição da informação (FDI) de uma memória associativa com k padrões armazenados é

$$\mathcal{D}(\text{Memória}, n, m, k; x) := 1 - F_D(x) = \Pr(D > x) \quad (7.27)$$

$$= \Pr(G_k^x(\mathbf{x}^\xi) = \mathbf{y}^\xi, \xi = 1, \dots, k), \quad (7.28)$$

onde F_D é a função de distribuição da variável aleatória D definida na equação (7.26).

Usamos o conceito de função de distribuição empírica e o teorema de Glivenko-Cantelli para estimar a distribuição da informação de uma memória associativa neural. Nas figuras 7.2 e 7.3 apresentamos a FDI empírica das memórias associativas binárias apresentadas nesta dissertação. As memórias associativas binárias foram treinadas com 6 padrões gerados aleatoriamente. Nos experimentos computacionais usamos $n = 100$, $m = 80$ e estimamos a FDI realizando 100 simulações.

Na figura 7.2 apresentamos a FDI empírica para o caso auto-associativo. A linha com $\circ$ representa a ECAM e a linha com $\times$ representa a memória associativa morfológica de duas camadas. Note que estes dois modelos forneceram a mesma FDI empírica. A linha com $\square$ representa a memória auto-associativa morfológica W_{XX} . A memória auto-associativa M_{XX} produziu um resultado semelhante e não foi apresentada no gráfico. A linha com $*$ representa a memória associativa de Hopfield, a linha com ∇ representa a memória associativa de Personnaz e a linha com $\triangle$ representa a memória associativa de Kanter-Sompolinsky. Note que a memória associativa de Personnaz apresentou a maior FDI empírica ponto a ponto. A informação armazenada na ECAM e na memória associativa morfológica de duas camadas não é bem distribuída nas conexões sinápticas pois estas duas memórias apresentaram a menor FDI empírica ponto a ponto.

Na figura 7.3 apresentamos a FDI empírica para o caso heteroassociativo binário. A linha com $\circ$ representa a BECAM e a linha com $\times$ representa a memória associativa morfológica de duas camadas. A FDI empírica destes dois modelos coincidem. A linha com $\square$ representa a memória associativa morfológica W_{XY} . Note que a memória associativa morfológica W_{XY} teve uma probabilidade 0,33 de armazenar todas as memórias fundamentais como pontos fixos. Este resultado confere com o gráfico da capacidade de armazenamento discutido na seção anterior. A memória associativa M_{XY} produziu um resultado semelhante. A linha com $*$ representa a BAM e a linha com $\triangle$ representa a ABAM. Note que a ABAM apresentou a maior FDI empírica ponto a ponto.

7.3 Raio de Atração

O raio da bacia de atração, ou simplesmente, *raio de atração* é uma medida para a tolerância a ruído de uma memória associativa. Este conceito foi usado empiricamente por Kanter-Sompolinsky em [44]. Nesta seção fornecemos uma definição rigorosa deste conceito.

Definição 7.3.1 (Raio de Atração). Sejam $\Omega = \{(X, Y) : X = [\mathbf{x}^1, \dots, \mathbf{x}^k] \in \mathbb{R}^{n \times k}, Y = [\mathbf{y}^1, \dots, \mathbf{y}^k] \in \mathbb{R}^{m \times k}\}$ e $R^k : \Omega \rightarrow \mathbb{R}$ a variável aleatória

$$R^k = \sup\{r \in \mathbb{R} : d(\mathbf{x}, \mathbf{x}^1) \leq r, G_k(\mathbf{x}) = \mathbf{y}^1\}, \quad (7.29)$$

onde $G_k : \mathbb{R}^n \rightarrow \mathbb{R}^m$ é o mapeamento associativo da memória treinada com as memórias fundamentais $(\mathbf{x}^\xi, \mathbf{y}^\xi)$, para $\xi = 1, \dots, k$, e $d : \mathbb{R}^n \times \mathbb{R}^n \rightarrow [0, +\infty)$ é uma métrica. O raio de atração de

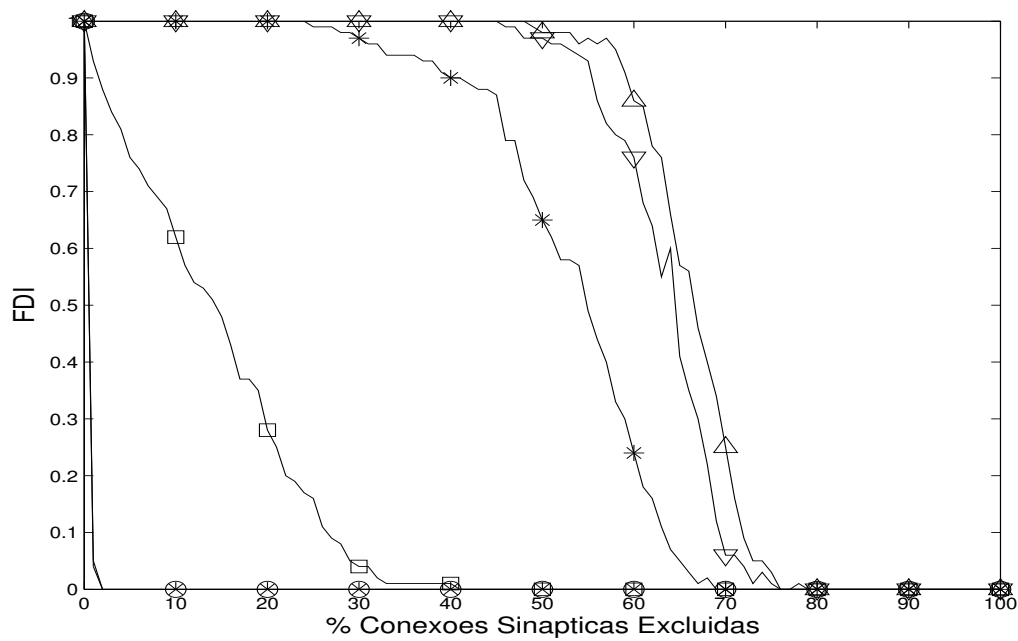


Fig. 7.2: FDI empírica das memórias auto-associativa neurais pela porcentagem de conexões sinápticas excluídas. As linhas representam: ECAM (○), MAM Duas Camadas (×), MAM W_{XX} (□), MA Hopfield (*), MA Personnaz (Δ) e MA Kanter (▽).

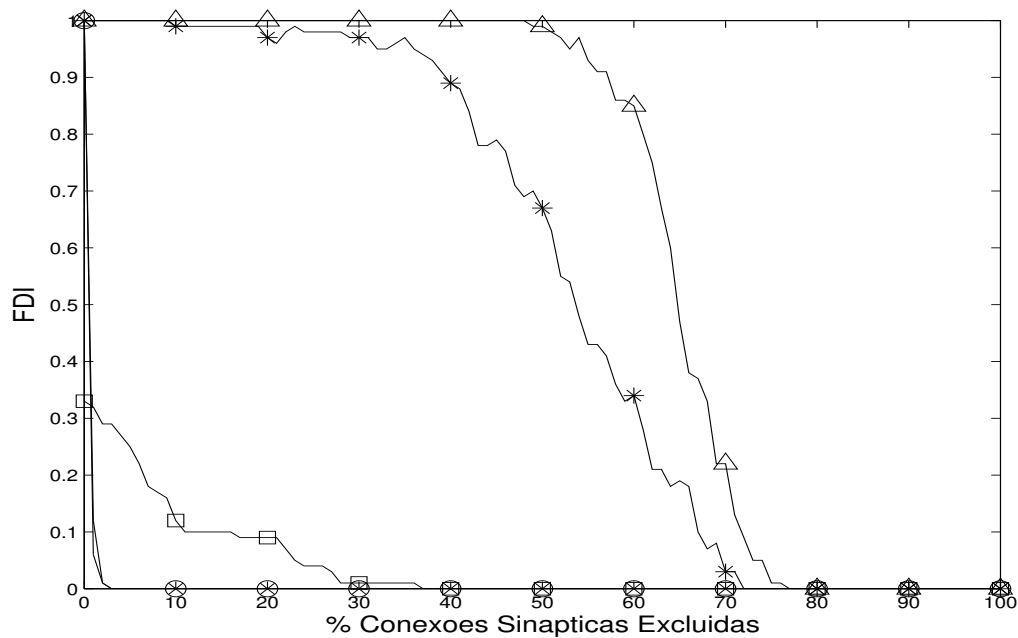


Fig. 7.3: FDI empírica das memórias heteroassociativa neurais pela porcentagem de conexões sinápticas excluídas. As linhas representam: BECAM (○), MAM Duas Camadas (×), MAM W_{XY} (□), BAM (*) e ABAM (Δ).

uma memória associativa será

$$\rho(\text{Memória}, m, n; k) := E(R^k), \quad (7.30)$$

onde $E(R^k)$ representa a esperança da variável aleatória R^k .

Na equação (7.29), se $\mathbf{x}$ é uma versão ruidosa de $\mathbf{x}^1$ com $d(\mathbf{x}, \mathbf{x}^1) < R^k$, então o padrão recordado pela memória é $\mathbf{y}^1$, que é o mesmo padrão recordado quando apresentamos a memória fundamental $\mathbf{x}^1$. No caso binário, podemos substituir $\sup$ por $\max$ pois $d(\mathbf{x}, \mathbf{x}^1)$ assume somente valores discretos. Nesta dissertação usamos como métrica a distância de Hamming definida como sendo o número de componentes distintas de dois padrões binários ou bipolares.

Na prática estimamos R^k da seguinte forma:

1. Tome $\mathbf{x} = \mathbf{x}^1$,
2. Enquanto $G_k(\mathbf{x}) = \mathbf{y}^1$ faça:
 - (a) Escolha aleatoriamente um índice $i \in \{1, \dots, n\}$ que ainda não foi escolhido,
 - (b) Defina $x_i = -x_i$ no caso bipolar (ou $x_i = 1 - x_i$ no caso binário).
3. Defina $R^k = d_H(\mathbf{x}, \mathbf{x}^1) - 1$.

O procedimento acima fornece um valor R^k maior que o valor teórico definido na equação (7.29). Entretanto, esta diferença é irrelevante pois usaremos R^k somente para comparar os modelos de memórias associativas binárias ou bipolares e usaremos sempre o procedimento descrito acima. Tendo R^k , podemos estimar o raio de atração usando a lei dos grandes números.

Lembre-se que o padrão recordado por uma memória associativa dinâmica é obtido somente após a convergência para um ponto fixo e a recursividade da fase de recordação está implícita no mapeamento associativo da memória.

Sabemos que a memória associativa morfológica M_{XY} e a memória associativa morfológica de duas camadas apresentam tolerância a ruído somente se o padrão-chave $\mathbf{x} \geq \mathbf{x}^1$. Por esta razão impomos $\mathbf{x}^1 = \mathbf{0}$ na hora de gerar as memórias fundamentais $(\mathbf{x}^\xi, \mathbf{y}^\xi)$, para $\xi = 1, \dots, k$. Analogamente, impomos $\mathbf{x}^1 = \mathbf{1}$ para a memória associativa morfológica W_{XY} e a versão dual da memória associativa de duas camadas. Portanto, o raio de atração obtido para as memórias associativas morfológicas só fará sentido para padrões corrompidos somente com ruído dilativo ou erosivo.

Na figura 7.4 apresentamos o gráfico do raio de atração pelo número de memórias fundamentais armazenadas nas memórias auto-associativas bipolares e binárias discutidas nesta dissertação. Os gráficos foram obtidos usando $\mathbf{x}^\xi \in \{-1, 1\}^{100}$ (ou $\mathbf{x}^\xi \in \{0, 1\}^{100}$) e calculando a média após 100 simulações. A linha com $\times$ representa a memória associativa morfológica de duas camadas e a linha com $\circ$ representa a ECAM. Estes dois modelos possuem os maiores raios de atração. A linha com $\square$ representa a memória associativa morfológica W_{XX} . A memória associativa M_{XX} produziu um resultado semelhante e não foi apresentada na figura 7.4. Note que o raio de atração da memória associativa morfológica W_{XX} (e M_{XX}) apresentou um grande decaimento. Este resultado é uma consequência do teorema 6.3.2. A linha com $*$ representa a memória associativa de Hopfield, a linha com ∇ representa a memória associativa de Personnaz e a linha com $\triangle$ representa a memória associativa de Kanter-Sompolinsky. Note que a memória associativa de Kanter-Sompolinsky apresentou um raio

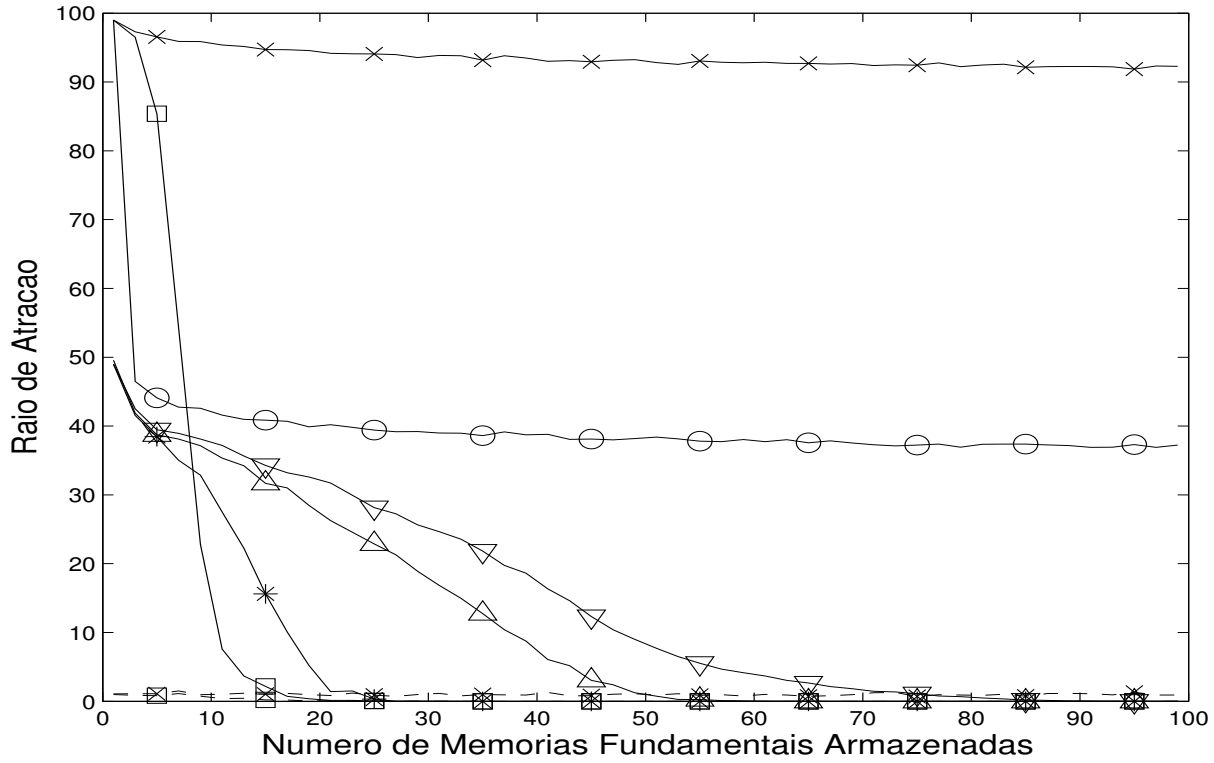


Fig. 7.4: Raio de atração das memórias auto-associativas binárias pelo número de memórias fundamentais. As linhas representam: ECAM ($\circ$), MAM Duas Camadas ($\times$), MAM W_{XX} ($\square$), MA Hopfield ($*$), MA Personnaz (Δ) e MA Kanter (∇). As linhas tracejadas com $\times$ e $\square$ representam as memórias associativas morfológicas de duas camadas e de camada única, respectivamente, sem impor $\mathbf{x}^1 = \mathbf{0}$ ou $\mathbf{x}^1 = \mathbf{1}$.

de atração maior que o raio de atração da memória associativa de Personnaz. Esta é a vantagem de impor $w_{ii} = 0$, para $i = 1, \dots, n$ [44]. Note que a memória associativa morfológica de duas camadas apresentou o maior raio de atração, entretanto, sua correção de erro é limitada para memórias-chave $\mathbf{x} \leq \mathbf{x}^1$. O segundo maior raio de atração foi da ECAM que é válido para ambos os tipos de ruído, dilativo e erosivo. As linhas tracejadas com $\times$ e $\square$ representam as memórias associativas morfológicas de duas camadas e de camada única, respectivamente, sem impor $\mathbf{x}^1 = \mathbf{0}$ ou $\mathbf{x}^1 = \mathbf{1}$. Note que ambas memórias associativas morfológicas tiveram raio de atração próximo de zero.

Na figura 7.5 apresentamos o gráfico do raio de atração pelo número de memórias fundamentais armazenadas. Os gráficos foram obtidos usando $\mathbf{x}^\xi \in \{-1, 1\}^{100}$ (ou $\mathbf{x}^\xi \in \{0, 1\}^{100}$), $\mathbf{y}^\xi \in \{-1, 1\}^{80}$ (ou $\mathbf{y}^\xi \in \{0, 1\}^{80}$) e calculando a média após 100 simulações. A linha com $\times$ representa a memória associativa morfológica de duas camadas e a linha com $\circ$ representa a BECAM. A linha com $\square$ representa a memória associativa morfológica W_{XY} . A memória associativa M_{XY} produziu um resultado semelhante. A linha com $*$ representa a BAM e a linha com Δ representa a ABAM. Em analogia ao caso auto-associativo, a memória associativa morfológica de duas camadas apresentou o maior raio de atração, entretanto, sua correção de erro é limitada para memórias-chave $\mathbf{x} \leq \mathbf{x}^1$. O segundo maior raio de atração foi da BECAM que é a generalização da ECAM para o caso heteroassociativo.

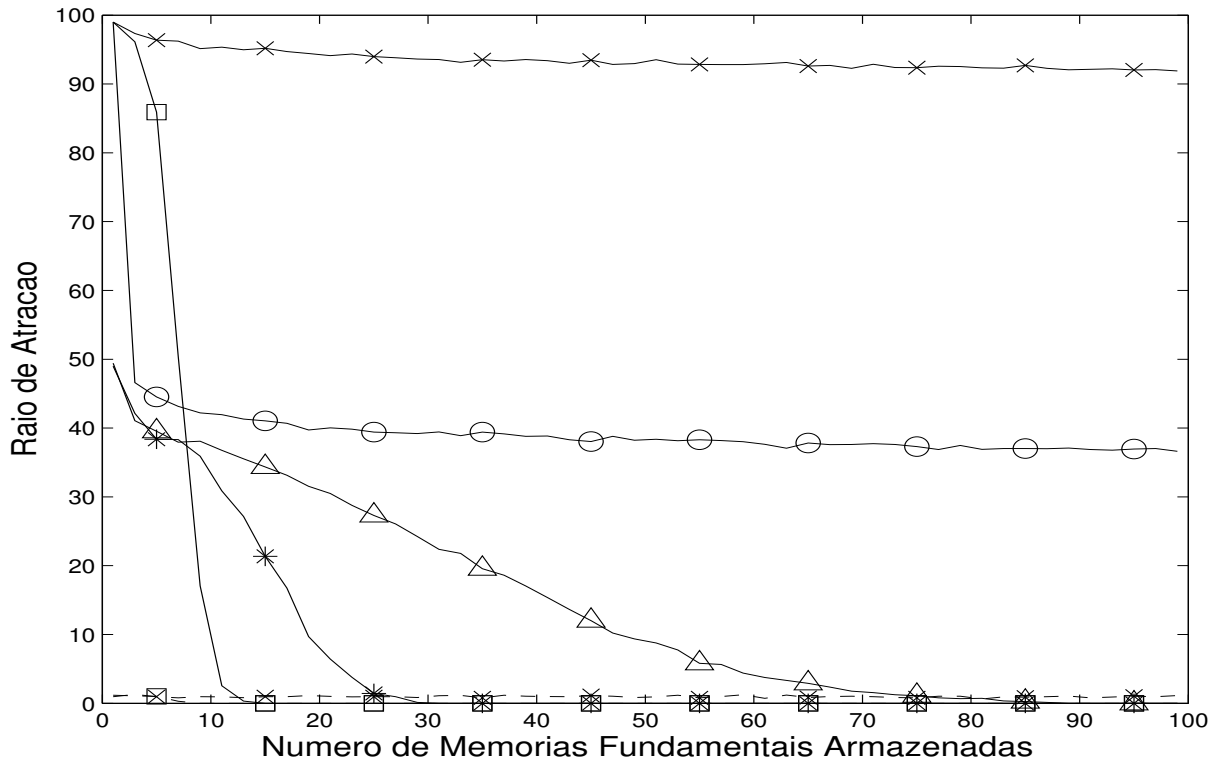


Fig. 7.5: Raio de atração das memórias heteroassociativa pelo número de memórias fundamentais. As linhas representam: BECAM (○), MAM Duas Camadas (×), MAM W_{XY} (□), BAM (*) e ABAM (Δ). As linhas tracejadas com × e □ representam as memórias associativas morfológicas de duas camadas e de camada única, respectivamente, sem impor $\mathbf{x}^1 = \mathbf{0}$ ou $\mathbf{x}^1 = \mathbf{1}$.

As linhas tracejadas com × e □ representam as memórias associativas morfológicas de duas camadas e de camada única, respectivamente, sem impor $\mathbf{x}^1 = \mathbf{0}$ ou $\mathbf{x}^1 = \mathbf{1}$. Note que ambas memórias associativas morfológicas tiveram raio de atração próximo de zero.

7.4 Memórias Espúrias

Nesta seção medimos a probabilidade de um padrão-chave convergir para um padrão que não faz parte do conjunto das memórias fundamentais de uma memória associativa treinada com k padrões.

Definição 7.4.1 (Probabilidade de Memória Espúria, PME). Seja $\Omega = \{(X, Y, \mathbf{x}) | X = [\mathbf{x}^1, \dots, \mathbf{x}^k] \in \mathbb{R}^{n \times k}, Y = [\mathbf{y}^1, \dots, \mathbf{y}^k] \in \mathbb{R}^{m \times k}, \mathbf{x} \in \mathbb{R}^n\}$. A probabilidade de uma memória associativa treinada com k memórias fundamentais convergir para uma memória espúria é

$$\mathcal{E}(\text{Memória}, n, m; k) := Pr(G_k(\mathbf{x}) \notin \{\mathbf{y}^1, \dots, \mathbf{y}^k\}) = 1 - Pr(G_k(\mathbf{x}) \in \{\mathbf{y}^1, \dots, \mathbf{y}^k\}), \quad (7.31)$$

onde $G_k: \mathbb{R}^n \rightarrow \mathbb{R}^m$ é o mapeamento associativo da memória associativa treinada com as memórias fundamentais $(\mathbf{x}^\xi, \mathbf{y}^\xi), \xi = 1, \dots, k$.

Usaremos a lei dos grandes números para estimar a probabilidade de memória espúria (PME) de uma memória associativa neural. Devemos ter um cuidado especial com as memórias associativas morfológicas pois sabemos que a MAM M_{XY} e a MAM de duas camadas sempre fornecerão uma memória espúria como resposta se o padrão-chave $\mathbf{x} \leq \bigwedge_{\xi=1}^k \mathbf{x}^\xi$, onde $\mathbf{x}^\xi, \xi = 1, \dots, k$ são as memórias fundamentais armazenadas. Para evitar este problema, definimos $\mathbf{x}^1 = \mathbf{0}$. Assim, para todo padrão-chave $\mathbf{x} \in \{0, 1\}^n$, teremos $\mathbf{x} \geq \bigwedge_{\xi=1}^k \mathbf{x}^\xi$ e podemos interpretar $\mathbf{x}$ como sendo uma versão corrompida com ruído dilativo. Analogamente, a MAM W_{XY} e a versão dual da MAM de duas camadas sempre fornecerão uma memória espúria se o padrão-chave $\mathbf{x} \geq \bigvee_{\xi=1}^k \mathbf{x}^\xi$ e definiremos $\mathbf{x}^1 = \mathbf{1}$ para evitar este problema.

Na figura 7.6 apresentamos a PME pelo número de memórias fundamentais armazenadas nas memórias auto-associativas binárias apresentadas nos capítulos 5 e 6. Neste experimento usamos $n = 100$ e realizamos 1000 simulações para calcular as probabilidades empíricas. A linha com $\times$ representa a memória associativa morfológica de duas camadas e a linha com $\circ$ representa a ECAM. Note que a PME destes dois modelos foi sempre nula. A linha com $\square$ representa a memória associativa morfológica W_{XX} . A memória associativa M_{XX} produziu um resultado semelhante e não foi apresentada na figura 7.6. Note que a PME da memória associativa morfológica W_{XX} (e M_{XX}) tende rapidamente para 1. Este resultado é uma consequência do teorema 6.3.2 sobre os pontos fixos das memórias auto-associativas morfológicas binárias. A linha com $*$ representa a memória associativa de Hopfield, a linha com ∇ representa a memória associativa de Personnaz e a linha com Δ representa a memória associativa de Kanter-Sompolinsky. Note que a PME é sempre maior que $1/2$ pois o mapeamento associativo destes três últimos modelos é um mapeamento ímpar e portanto, se $\mathbf{x}^\xi$ é um ponto fixo, então $-\mathbf{x}^\xi$ também é um ponto fixo da memória associativa. As memórias associativas morfológicas apresentaram ambas uma probabilidade empírica próxima de 1 se não impormos a condição discutida no parágrafo anterior sobre o padrão $\mathbf{x}^1$.

Na figura 7.7 apresentamos a probabilidade de memória espúria pelo número de memórias fundamentais armazenadas nas memórias heteroassociativas binárias. Neste experimento tomamos $n = 100, m = 80$ e efetuamos 1000 simulações para calcular as probabilidades empíricas. A linha com $\times$ representa a memória associativa morfológica de duas camadas e a linha com $\circ$ representa a BECAM. Ambos modelos apresentaram uma probabilidade de memória espúria nula. A linha com $\square$ representa a memória associativa morfológica W_{XY} . A linha com $*$ representa a BAM e a linha com Δ representa a ABAM. A probabilidade de memória espúria da BAM e da ABAM é sempre maior que $1/2$ porque estes modelos possuem um mapeamento associativo ímpar. Novamente, as memórias associativas morfológicas apresentaram ambas uma probabilidade empírica próxima de 1 se não impormos a condição discutida anteriormente sobre o padrão $\mathbf{x}^1$.

7.5 Esforço Computacional

O esforço computacional pode ser medido pelo número de operações e execuções de funções não lineares realizadas pela memória associativa neural na fase de armazenamento e na fase de recordação. Nas memórias associativas dinâmicas também devemos considerar o número de iterações necessárias para encontrar a saída da rede na fase de recordação.

Nesta seção vamos considerar $X = [\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^k] \in \mathbb{R}^{n \times k}$ e $Y = [\mathbf{y}^1, \mathbf{y}^2, \dots, \mathbf{y}^k] \in \mathbb{R}^{m \times k}$.

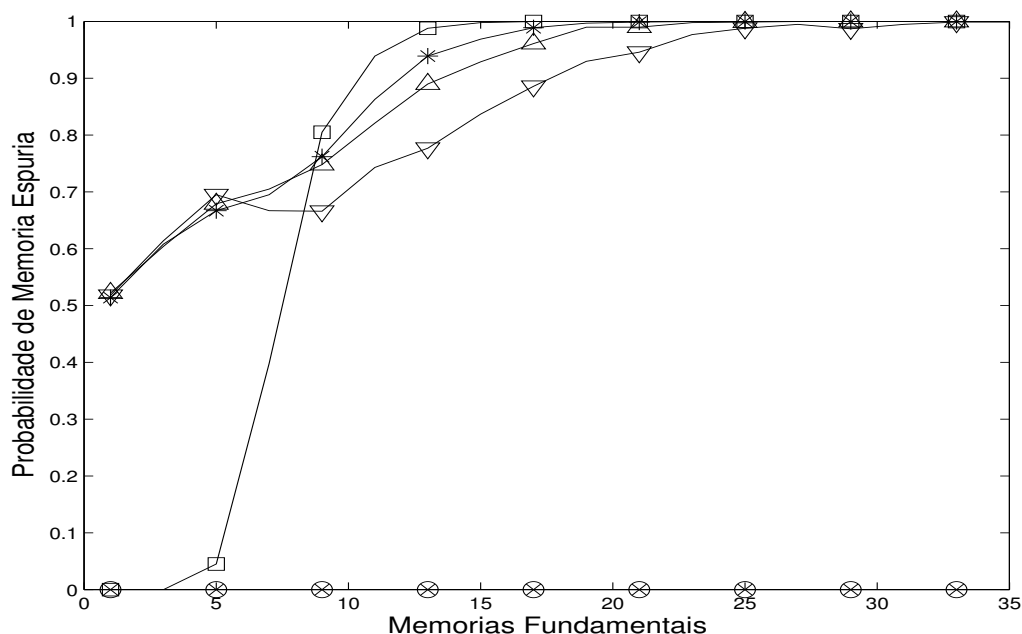


Fig. 7.6: Probabilidade de memória espúria pelo número de memórias fundamentais. As linhas representam: ECAM (○), MAM Duas Camadas (×), MAM W_{XX} (□), MA Hopfield (*), MA Personnaz (△) e MA Kanter (▽).

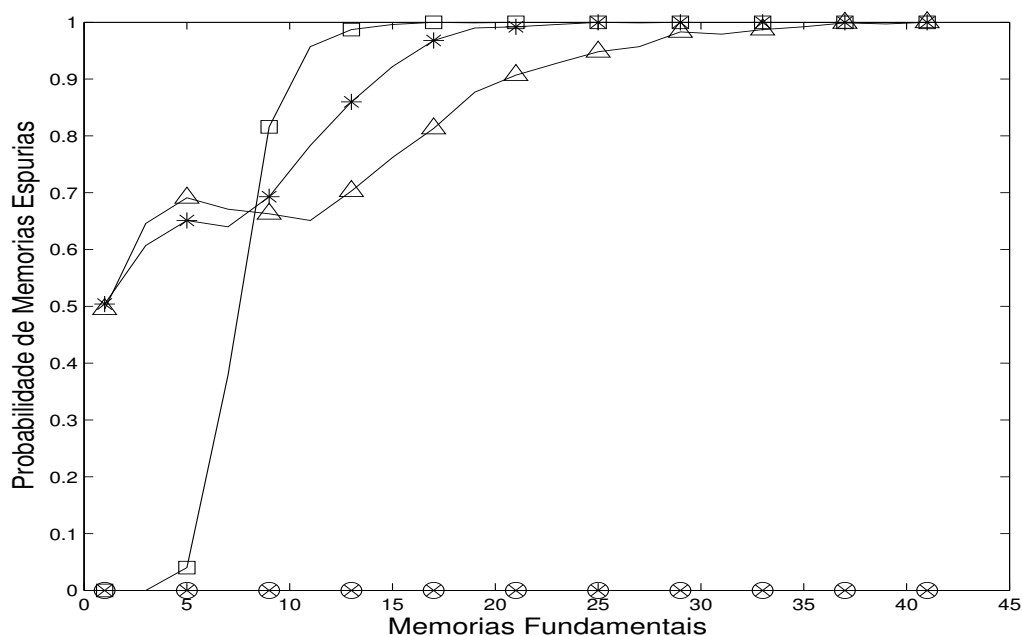


Fig. 7.7: Probabilidade de memória espúria pelo número de memórias fundamentais armazenadas. As linhas representam: BECAM (○), MAM Duas Camadas (×), MAM W_{XY} (□), BAM (*) e ABAM (△).

7.5.1 Número de Operações na Fase de Armazenamento

A matriz de pesos sinápticos da BAM é $W = YX^T$. Logo, serão necessárias $(2k-1)mn$ operações de soma e de multiplicação para obtermos W . Em particular, a matriz dos pesos sinápticos da memória associativa de Hopfield requer $(2k-1)n^2$ operações de soma e multiplicação.

Na memória associativa de Personnaz, a matriz dos pesos sinápticos é $W = XX^\dagger = \hat{U}\hat{U}^T$, onde $X = \hat{U}\hat{\Sigma}V^T$ é a decomposição SVD reduzida de X (Thin SVD Decomposition). O cálculo da matriz $\hat{U}$ usando o método R-SVD requer $6nk^2 + 11k^3$ operações e o produto $\hat{U}\hat{U}^T$ requer $(2k-1)n^2$ operações de soma e multiplicação [25]. O número total de operações de soma e multiplicação necessárias para obter a matriz dos pesos sinápticos da memória associativa de Personnaz é $6nk^2 + 11k^3 + (2k-1)n^2$. Na memória associativa de Kanter-Sompolinsky tomamos $W = XX^\dagger$ e impomos $w_{ii} = 0$ para todo $i = 1, \dots, n$. Logo, o número de operações necessárias para encontrar a matriz dos pesos sinápticos da memória associativa de Kanter-Sompolinsky é também $6nk^2 + 11k^3 + (2k-1)n^2$ operações de soma e multiplicação.

Na ABAM definimos $W_1 = YX^\dagger$ e $W_2 = XY^\dagger$. O número de operações necessária para calcular $X^\dagger$ usando a decomposição SVD reduzida é $6nk^2 + 20k^3$ operações. Calculado a decomposição SVD reduzida $X = \hat{U}\hat{\Sigma}V^T$, computamos W_1 através do produto $W_1 = \left(Y(V\hat{\Sigma}^\dagger)\right)\hat{U}^T$ que requer $(2m+1)k^2 + 2mnk$ operações de soma e multiplicação. Analogamente, o calculo da matriz W_2 requer $6mk^2 + 20k^3$ operações para calcular a decomposição SVD reduzida de Y e $(2n+1)k^2 + 2mnk$ operações para calcular o produto matricial em W_2 . O número total de operações de soma e multiplicação necessárias para computar as matrizes de pesos sinápticos da ABAM é $40k^3 + 2(4n+4m+1)k^2 + 4mnk$.

Na BECAM usamos as matrizes das memórias fundamentais como matriz dos pesos sinápticos. Portanto, nenhuma operação é efetuada na fase de armazenamento desta memória associativa. Em particular, a ECAM também não efetua nenhuma operação na fase de armazenamento.

Nas memórias associativas morfológicas tomamos $W_{XY} = Y \boxtimes X^*$ e $M_{XY} = Y \boxtimes X^*$. Considerando as operações de máximo ou mínimo como operações binárias, serão necessárias kmn operações de soma e $(k-1)mn$ operações de mínimo para calcular W_{XY} ou kmn operações de soma e $(k-1)mn$ operações de máximo para calcular M_{XY} . O número total de operações necessárias para calcular a matriz dos pesos sinápticos de uma memória associativa morfológica é $(2k-1)mn$ operações de máximo ou mínimo e soma.

Nas memórias associativas morfólicas de duas camadas precisamos da matriz W_{ZY} que requer $(2k-1)pm$ operações de mínimo e soma. A matriz $M_X^{XY} = M_{XZ} \boxtimes M_{XX}$ que requer $(2k-1)pn$ operações para o cálculo de M_{XZ} , $(2k-1)n^2$ operações para o cálculo de M_{XX} e $(2n-1)pn$ operações para o calculo do produto $M_{XZ} \boxtimes M_{XX}$. Como as operações de máximo e mínimo requerem o mesmo esforço computacional, o número total de operações necessárias para obter as matrizes dos pesos sinápticos da memória associativa morfológica de duas camadas será $2pn(k+n-1) + (2k-1)(pm+n^2)$. Nos nossos experimentos computacionais tomamos $Z = I_{k \times k}$, ou seja, $p = k$.

Na tabela 7.1 apresentamos um resumo do que foi dito anteriormente.

Note que o número de operações necessárias para calcular a matriz dos pesos sinápticos da BAM e de uma MAM são os mesmos. Entretanto, teremos uma esforço computacional menor na fase de armazenamento das MAM pois as operações de máximo ou mínimo e soma requerem um esforço computacional menor que as operações de soma e multiplicação.

Memória Associativa	Número Aproximado de Operações	Tipo das Operações
MA Hopfield	$(2k - 1)n^2$	soma e multiplicação
BAM	$(2k - 1)nm$	soma e multiplicação
MA Personnaz	$6nk^2 + 11k^3 + (2k - 1)n^2$	soma e multiplicação
MA Kanter-Sompolinsky	$6nk^2 + 11k^3 + (2k - 1)n^2$	soma e multiplicação
ABAM	$40k^3 + 2(4n + 4m + 1)k^2 + 4mnk$	soma e multiplicação
ECAM	—	—
BECAM	—	—
MAM	$(2k - 1)nm$	máximo ou mínimo e soma
MAM Duas Camadas	$2pn(k + n - 1) + (2k - 1)(pm + n^2)$	máximo ou mínimo e soma

Tab. 7.1: Esforço computacional na fase de armazenamento das memórias associativas neurais.

7.5.2 Número de Operações por Iteração na Fase de Recordação

Na fase de recordação devemos considerar o esforço computacional realizado pela memória por iteração e o número de iterações necessário para recordar um padrão. Discutiremos primeiro o número de operações por iteração, depois apresentaremos uma estimativa do número de iterações das memórias associativas dinâmicas. O esforço computacional será descrito pelo número de operações efetuadas e o número de chamadas de funções não-lineares por iteração na memória associativa.

Nas memórias associativas de Hopfield, Personnaz e Kanter-Sompolinsky realizamos um produto matriz-vetor e depois aplicamos a função sinal em cada componente do resultado do produto matriz-vetor. Estas operações requerem $(2n - 1)n$ operações de soma e multiplicação e n efetuções da função sinal por iteração.

Na ABAM computamos $\text{sinal}(W_1 \mathbf{x})$ e $\text{sinal}(W_2 \mathbf{y})$ por iteração. Na BAM realizamos os mesmos cálculos com $W_1 = W$ e $W_2 = W^T$ na BAM. O total de operações efetuadas por iteração na BAM ou na ABAM é $4mn - (m + n)$ operações de soma e multiplicação e $m + n$ efetuções da função sinal.

Na ECAM computamos $\text{sinal}(X \exp(X^T \mathbf{x}))$. Por iteração realizamos $4kn - (k + n)$ operações, efetuamos a função exponencial k vezes e a efetuamos a função sinal n vezes. Na BECAM, computamos $\text{sinal}(Y \exp(X^T \mathbf{x}))$ e $\text{sinal}(X \exp(Y^T \mathbf{y}))$ realizando $4k(n + m) - (2k + m + n)$ operações de soma e multiplicação, $2k$ efetuções da função exponencial e $m + n$ efetuções da função sinal por iteração.

Nas memórias associativas morfológicas realizamos o produto $W_{XY} \boxtimes \mathbf{x}$ ou $M_{XY} \boxtimes \mathbf{y}$. Ambos requerem $(2n - 1)m$ operações de máximo ou mínimo e soma. O padrão recordado pela memória associativa morfológica de duas camadas é dado por $W_{ZY} \boxtimes f(M_X^{XZ} \boxtimes \mathbf{x})$, onde M_X^{XZ} é uma matriz $p \times n$ e W_{ZY} é uma matriz $m \times p$. Logo, na fase de armazenamento das memórias associativas morfológicas de duas camadas efetuamos $2p(m + n) - (p + m)$ operações de máximo ou mínimo e soma, e computamos a função f definida na equação (6.41) p vezes. Nos nossos experimentos tomamos $Z = I_{k \times k}$ ($p = k$).

Na tabela 7.2 apresentamos o que foi dito anteriormente. A segunda coluna da tabela 7.2 representa o número total de operações binárias incluindo soma, produto, máximo e mínimo. A terceira coluna contém o número de efetuções da função sinal ou da função f definida na equação (6.41). A

Memória Associativa	N. Operações	$f(x)$ ou $\text{sin}(x)$	$\exp(x)$
MA Hopfield	$2n^2 - n$	n	—
BAM	$4nm - (n + m)$	$m + n$	—
MA Personnaz	$2n^2 - n$	n	—
MA Kanter-Sompolinsky	$2n^2 - n$	n	—
ABAM	$4nm - (m + n)$	$m + n$	—
ECAM	$4kn - (k + n)$	n	k
BECAM	$4k(m + n) - (2k + m + n)$	$m + n$	$2k$
MAM	$2mn - m$	—	—
MAM Duas Camadas	$2p(m + n) - (p + m)$	k	—

Tab. 7.2: Esforço computacional por iteração na fase de recordação das memórias associativas neurais.

quarta coluna representa o número de efetuações da função exponencial.

7.5.3 Número de Iterações na Fase de Recordação

O número de iterações na fase de recordação de uma memória associativa dinâmica treinada com k memórias fundamentais é uma função da dimensão dos padrões de entrada e saída e do número de padrões armazenados. Precisamente, o número de iterações na fase de recordação é definido como sendo a média do número de iterações necessárias para a convergência do sistema dinâmico para um ponto estacionário quando iniciamos a memória com um padrão-chave qualquer. As memórias associativas estáticas serão consideradas como sistemas dinâmicos que convergem para um ponto estacionário com 1 iteração.

Na figura 7.8 apresentamos o gráfico do número de iterações pelo número de padrões armazenados de várias memórias auto-associativas dinâmicas. Esta figura foi gerada usando padrões com 100 componentes e o gráfico foi construído calculando a média após 1000 simulações. A linha marcada com * representa a MA de Hopfield, a linha com Δ representa a MA de Personnaz, a linha com ∇ representa a MA de Kanter-Sompolinsky e a linha com $\circ$ representa a ECAM. Note que o número de iterações da MA Personnaz e da MA de Kanter-Sompolinsky tendem para 1 quando o número de memórias fundamentais armazenadas tende para a dimensão dos padrões armazenados. De fato, $W \rightarrow I$ na MA de Personnaz e $W \rightarrow 0$ na MA de Kanter-Sompolinsky quando $k \rightarrow n$. Em ambos os casos $\text{sin}(Wx) = x$ para todo padrão-chave x quando $k \rightarrow n$. Logo, todos os pontos do espaço são pontos fixos da MA de Personnaz e da MA de Kanter-Sompolinsky quando $k \rightarrow n$.

Na figura 7.9 apresentamos o gráfico do número de iterações pelo número de padrões armazenados em memórias heteroassociativas dinâmicas. Neste exemplo tomamos padrões de entrada com 100 componentes, padrões de saída com 80 componentes e calculamos a média após 1000 simulações. O número máximo de iterações permitido foi 1000. A linha com * representa a BAM, a linha com Δ representa a ABAM e a linha com $\circ$ representa a BECAM. Note que a ABAM atingiu o número máximo de iterações permitido. Isso mostra que a ABAM pode convergir para um ciclo limite. Lembre-se que não temos um resultado que garante a convergência da ABAM. Pelo gráfico, a ABAM atingiu o número máximo de iterações para $k \geq 40$. Neste caso, podem haver ciclos-limite na ABAM

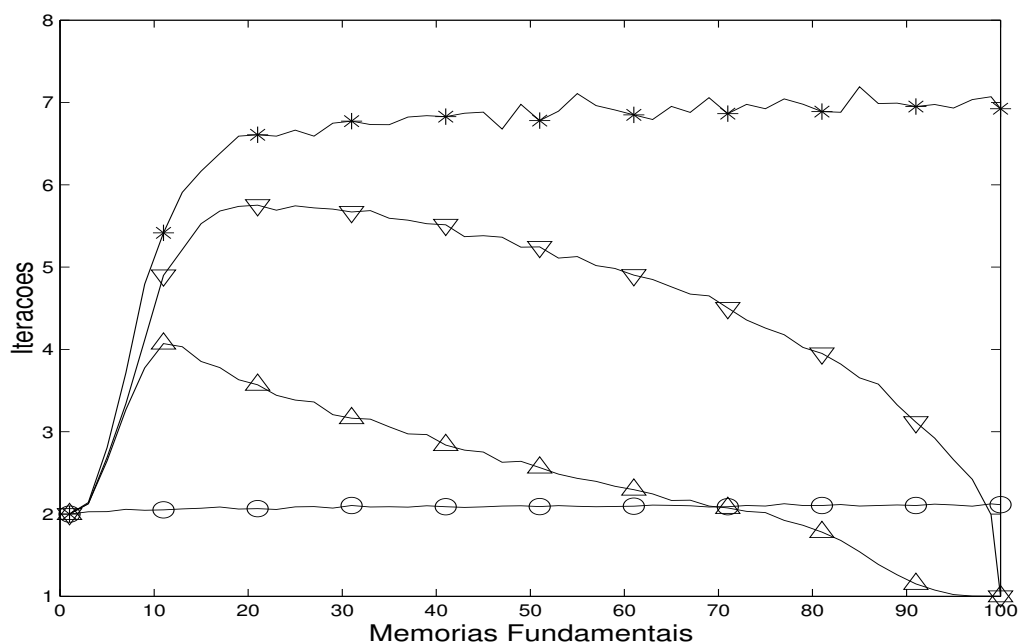


Fig. 7.8: Número de iterações na fase de recordação pelo número de padrões armazenados. As linhas marcadas representam: ECAM (○), MA Hopfield (*), MA Personnaz (△) e MA Kanter (▽).

que impedem a convergência do padrão-chave para um ponto de equilíbrio.

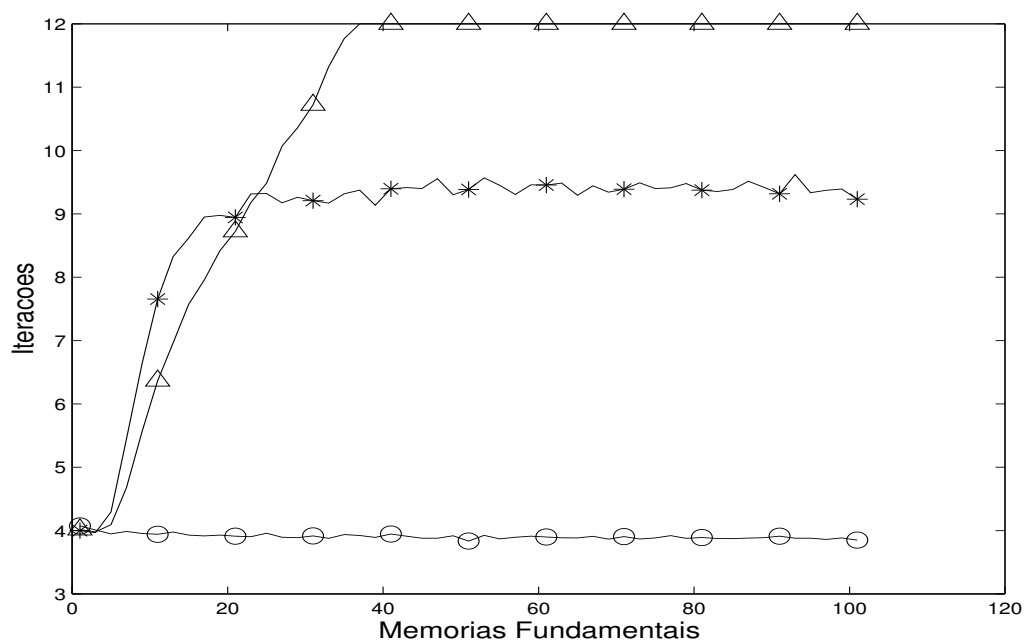


Fig. 7.9: Número de iterações na fase de recordação pelo número de padrões armazenados. As linhas marcadas representam: BECAM (○), BAM (*) e ABAM (△).

Capítulo 8

Conclusão

Nesta dissertação apresentamos um estudo comparativo em memórias associativas neurais com ênfase nas memórias associativas morfológicas. Nos concentramos nos modelos de memórias associativas binárias que são citados frequentemente na literatura de redes neurais.

No capítulo 1 apresentamos uma revisão história e bibliográfica sobre redes neurais, morfologia matemática, álgebra de imagens e principalmente sobre memórias associativas neurais. Nos capítulos 2 e 3 discutimos conceitos básicos de redes neurais e memórias associativas. Existe uma variedade grande de notações para modelos de memórias associativas neurais e estes capítulos básicos esclarecem a notação usada durante a dissertação. Adotamos uma notação matricial comum em ambos livros de álgebra linear e redes neurais artificiais, como por exemplo, nas referências [101] e [33]. As cinco características para um bom desempenho apresentadas por nós no capítulo 3 foram inspiradas nos trabalhos de Hassoun e Pao [31, 64].

No capítulo 4 discutimos as memórias associativas lineares. Este capítulo tem um objetivo didático visto que as memórias associativas lineares definem as principais regras de aprendizado utilizadas nas memórias associativas neurais discutidas nesta dissertação. Verificamos que o armazenamento por correlação possui limitações devido à interferência cruzada. No armazenamento por projeção temos erro devido à dependência linear das memórias fundamentais e devido ao ruído introduzido no padrão-chave. Um número ilimitado de padrões podem ser armazenados na OLAM no caso auto-associativo.

No capítulo 5 discutimos as memórias associativas dinâmicas. Este é o maior capítulo da dissertação devido ao grande número de modelos apresentados na literatura. Para cada modelo apresentamos a arquitetura, regra de aprendizado, exemplos computacionais e uma breve análise sobre a convergência. A ABAM foi o único modelo que não possui nenhum resultado garantindo sua convergência para um ponto fixo. A conjectura 5.2.1 é uma proposta nossa baseada no resultado do artigo de McEliece *et. al.* [56] e pode ser vista como uma generalização do teorema 5.1.2 apresentado para a rede de Hopfield. As memórias associativas de Personnaz e Kanter-Somplinsky são ambas referidas como “rede de Hopfield com armazenamento por projeção”. Verificamos que existem diferenças entre estes dois modelos, principalmente com respeito ao raio de atração (tolerância a ruído). As proposições que garantem a convergência da memória associativa de Personnaz foram introduzidas por nós. A BECAM foi introduzida por nós como uma generalização da ECAM inspirada na BAM. Todos os resultados relativos a BECAM são novos. O teorema 5.8.3 é inédito e relaciona o modelo BSB com o modelo de Hopfield.

No capítulo 6 apresentamos um estudo detalhado das memórias associativas morfológicas. Começamos discutindo o caso hetero-associativo onde apresentamos a arquitetura, regra de aprendizado, exemplos e teoremas que garantem o armazenamento e a recordação de padrões. O teorema 6.1.3 foi introduzido e demonstrado nesta dissertação. Concluimos que as memórias associativas morfológicas hetero-associativas de camada única não são capazes de armazenar muitos padrões binários e possuem uma tolerância a ruído restrita a padrões corrompidos com ruído dilatativo ou ruído erosivo. O caso auto-associativo também foi discutido detalhadamente. Mostramos que as memórias auto-associativas morfológicas de camada única podem armazenar um número ilimitado de padrões, convergem para um ponto fixo com uma única iteração, e também apresentam restrições no tipo de ruído introduzido nos padrões-chave. Apresentamos um teorema que caracteriza todos os pontos fixos das memórias auto-associativas binárias e verificamos que estas apresentam um grande número de padrões espúrios. Discutimos brevemente o método do núcleo e depois introduzimos as memórias associativas morfológicas de duas camadas. Terminamos o capítulo com um teorema que caracteriza os pontos fixos deste último modelo.

No capítulo 7 formalizamos os conceitos para a medida do desempenho de uma memória associativa e usamos estes conceitos para comparar os vários modelos de memória associativa binária apresentados nos capítulos 5 e 6. Com base nos resultados obtidos concluimos que:

- As memórias auto-associativas morfológicas, de Personnaz e Kanter-Sompolinsky apresentaram uma capacidade de armazenamento constante igual a 1, isto é, podemos armazenar infinitos padrões. A ECAM e a BECAM apresentam uma capacidade de armazenamento igual à c^{kn} e $c^{k(n+m)}$ com c próximo de 1, respectivamente.
- A função de distribuição informa quantos pesos sinápticos podemos excluir sem perder a informação armazenada numa memória associativa neural. As memórias associativas de Kanter-Sompolinsky, Personnaz e Hopfield apresentaram os melhores resultados no caso auto-associativo. A ABAM e a BAM apresentaram os melhores resultados no caso heteroassociativo. Os piores resultados para a distribuição da informação foram obtidas pela ECAM, BECAM e as memórias associativas de duas camadas.
- As memórias associativas morfológicas de duas camadas apresentaram a maior tolerância a ruído (maior raio de atração), entretanto, este modelo é restrito a padrões-chave corrompidos com ruído dilatativo ou erosivo. A ECAM e a BECAM são os modelos mais recomendados para padrões-chave corrompidos com ambos ruído dilatativo e erosivo.
- As memórias associativas morfológicas de duas camadas, a ECAM e a BECAM apresentaram as menores probabilidades de convergir para uma memória espúria. Lembramos que este resultado é válido para as memórias associativas morfológicas supondo que o padrão-chave representa uma versão erodida ou dilatada de uma memória fundamental.
- O esforço computacional depende da dimensão dos padrões de entrada e saída (n e m) e do número de memórias fundamentais armazenadas (k). Se $k \ll n, m$, a ECAM e a BECAM serão os modelos que requerem o menor esforço computacional, supondo que as memórias convergem rapidamente para um ponto-fixo. Se k é próximo de n ou m , então as memórias associativas morfológicas serão os modelos que efetuam o menor número de operações.

Finalmente, com base nos resultados apresentados no capítulo 7, não podemos afirmar qual é o melhor modelo de memória associativa neural, pois cada modelo apresenta pontos positivos e negativos. Esta dissertação de mestrado serve como um guia para a escolha do modelo de memória associativa que melhor se enquadra à um dado problema. Por exemplo, as memórias associativas morfológicas de duas camadas serão os modelos recomendados para um problema onde queremos armazenar um grande número de memórias fundamentais buscando um baixo custo computacional e sabendo-se que os padrões-chave serão versões corrompidas somente com ruído dilativo ou somente com ruído erosivo. Lembre-se que no capítulo 1 citamos várias referências contendo aplicações de memórias associativas neurais.

Referências Bibliográficas

- [1] Computer Vision Group Image Database, Available at <http://decsai.ugr.es/cvg/index2.php>.
- [2] AMARI, S.-I. Neural theory of association and concept-formulation. *Biological Cybernetics* 26 (1977), 175–185.
- [3] AMIT, D. J. *Modeling Brain Function – The World of Attractor Neural Network*. Cambridge University Press, 1989.
- [4] ANDERSON, J. A simple neural network generating interactive memory. *Mathematical Biosciences* 14 (1972), 197–220.
- [5] ANDERSON, J. *An Introduction to Neural Networks*. MIT Press, MA, 1995.
- [6] ANDERSON, J., SILVERSTEIN, J., RITZ, S., AND JONES, R. Distinctive features, categorical perception, and probability learning: Some applications of a neural model. *Psychology Review* 84 (July 1977), 413–415.
- [7] ANDERSON, J. A., AND ROSENFELD, E. *Neurocomputing: Foundations of Research*, vol. 1. MIT Press, Cambridge, MA, 1989.
- [8] BEALE, R., AND FIESLER, E., Eds. *Handbook of Neural Computation*. Institute of Physics Publishing and Oxford University Press, 1997.
- [9] BHAYA, A., KASZKUREWICZ, E., AND KOZYAKIN, V. Existence and stability of a unique equilibrium in continuous-valued discrete-time asynchronous hopfield neural networks. *IEEE Trans. on Neural Networks* 7, 3 (May 1996), 620 – 628.
- [10] BILLINGSLEY, P. *Probability and Measure*, 2 ed. John Wiley and Sons, 1986.
- [11] BIRKHOFF, G. *Lattice Theory*, 3 ed. American Mathematical Society, Providence, 1993.
- [12] BRUNAK, S., AND LAUTRUP, B. *Neural Networks: Computers with Intuition*. World Scientific, 1990.
- [13] BRYSON, A., AND HO, Y. *Applied Optimal Control*, rev ed edition ed. John Wiley and Sons, October 1979.

- [14] CASASANT, D., AND TELFER, B. Associative memory synthesis, performance, storage capacity, and updating: new heteroassociative memory results. *SPIE, Int. Robots Comput. Vision* 848 (1987), 313–333.
- [15] CHIUEH, T., AND GOODMAN, R. Recurrent correlation associative memories. *IEEE Trans. on Neural Networks* 2 (Feb. 1991), 275–284.
- [16] CHIUEH, T., AND GOODMAN, R. *Recurrent Correlation Associative Memories and their VLSI Implementation*. In Hassoun [31], 1993, ch. 16, pp. 276–287.
- [17] COLLIER, R. J. Some current views on holography. *IEEE Spectrum* 3 (July 1966), 67–74.
- [18] CUNINGHAME-GREEN, R. *Minimax Algebra: Lecture Notes in Economics and Mathematical Systems* 166. Springer-Verlag, New York, 1979.
- [19] CUNINGHAME-GREEN, R. Minimax algebra and applications. In *Advances in Imaging and Electron Physics*, P. Hawkes, Ed., vol. 90. Academic Press, New York, NY, 1995, pp. 1–121.
- [20] DAVIDSON, J., AND RITTER, G. A theory of morphological neural networks. In *Digital Optical Computing II* (July 1990), vol. 1215 of *Proceedings of SPIE*, pp. 378–388.
- [21] FU, L. *Neural Networks in Computer Intelligence*. McGraw-Hill, New York, NY, 1994.
- [22] GABOR, D. Associative holographic memories. *IBM J. Res. Develop* 13 (1969), 156–159.
- [23] GOLDEN, R. M. The brain-state-in-a-box neural model is a gradient descent algorithm. *Journal of Mathematical Psychology* 30 (1986), 73–80.
- [24] GOLDEN, R. M. Stability and optimization analysis of the generalized brain-state-in-a-box neural network model. *Journal of Mathematical Psychology* 30 (1993), 73–80.
- [25] GOLUB, G. H., AND LOAN, C. F. V. *Matrix Computations*, 3rd ed. Johns Hopkins University Press, 1996.
- [26] GRAÑA, M., GALLEGÓ, J., TORREALDEA, F. J., AND D’ANJOU, A. On the application of associative morphological memories to hyperspectral image analysis. *Lecture Notes in Computer Science* 2687 (2003), 567–574.
- [27] HADWIGER, H. *Vorlesungen Über Inhalt, Oberfläche und Isoperimetrie*. Springer-Verlag, Berlin, 1957.
- [28] HAGAN, M., DEMUTH, H., AND BEALE, M. *Neural Network Desing*. PWS Publishing Company, Boston, 1996.
- [29] HANSON, S., AND KEGL, J. Parnip: a connectionist network that learns natural language grammar from exposure to natural language sentences. In *Proc. 9th Annu. Conf. Cognitive Science* (1987), pp. 106–119.

- [30] HASSOUN, M. H. Dynamic heteroassociative neural memories. *Neural Networks* 2, 4 (1989), 275–287.
- [31] HASSOUN, M. H. Dynamic associative neural memories. In *Associative Neural Memories: Theory and Implementation*, M. H. Hassoun, Ed. Oxford University Press, Oxford, U.K., 1993.
- [32] HASSOUN, M. H. *Fundamentals of Artificial Neural Networks*. MIT Press, Cambridge, MA, 1995.
- [33] HAYKIN, S. *Neural Networks: A Comprehensive Foundation*. Prentice Hall, Upper Saddle River, NJ, 1999.
- [34] HEBB, D. *The Organization of Behavior*. John Wiley & Sons, New York, 1949.
- [35] HEIJMANS, H. *Morphological Image Operators*. Academic Press, New York, NY, 1994.
- [36] HIRSCH, M. W. Convergent activation dynamics in continuous time networks. *Neural Networks* 2, 5 (1989), 331–349.
- [37] HOPFIELD, J. Neurons with graded response have collective computational properties like those of two-state neurons. *Proceedings of the National Academy of Sciences* 81 (May 1984), 3088–3092.
- [38] HOPFIELD, J., AND TANK, D. Neural computation of decisions in optimization problems. *Biological Cybernetics* (1985).
- [39] HOPFIELD, J., AND TANK, D. Computing with neural circuits: A model. *Proceedings of the National Academy of Sciences* 233 (August 1986), 625–633.
- [40] HOPFIELD, J. J. Neural networks and physical systems with emergent collective computational abilities. *Proceedings of the National Academy of Sciences* 79 (Apr. 1982), 2554–2558.
- [41] HUI, S., LILLO, W. E., AND ŽAK, S. H. *Dynamics and Stability Analysis of the Brain-State-in-a-Box (BSB) Neural Models*. In Hassoun [31], 1993, ch. 11, pp. 212–224.
- [42] HUI, S., AND ŽAK, S. H. Dynamical analysis of the brain-state-in-a-box (bsb) neural models. *IEEE Transactions on Neural Networks* 3, 1 (January 1992), 86–94.
- [43] JAMES, B. *Probabilidade: Um Curso em Nível Intermediário*. Publicação IMPA, 1996.
- [44] KANTER, I., AND SOMPOLINSKY, H. Associative recall of memory without errors. *Physical Review* 35 (1987), 380–392.
- [45] KAPPEN, B., AND GIELEN, S. *Neural Networks: Best Practice in Europe*. Progress in Neural Processing, 8. World Scientific, Amsterdam, 1997.
- [46] KAWAMOTO, A. H., AND ANDERSON, J. A. A neural network model of multistable perception. *Acta Psychologica* 59 (1985), 35–65.

- [47] KOHONEN, T. Correlation matrix memory. *IEEE Transactions on Computers C-21* (1972), 353–359.
- [48] KOHONEN, T. *Associative Memory – A System Theoric Approach*. Berlin: Springer-Verlag, 1977.
- [49] KOHONEN, T. *Self-Organization and Associative Memory*. Springer Verlag, 1984.
- [50] KOHONEN, T., AND RUOHONEN, M. Representation of associated data by computers. *IEEE Transactions on Computers C-22* (1973), 701–702.
- [51] KOSKO, B. Adaptive bidirectional associative memories. *Applied Optics* 26, 23 (Dec. 1987), 4947–4960.
- [52] KOSKO, B. Bidirectional associative memories. *IEEE Transactions on Systems, Man, and Cybernetics* 18 (1988), 49–60.
- [53] LIMA, E. L. *Álgebra Linear*, 3 ed. Instituto de Matemática Pura e Aplicada, Rio de Janeiro, RJ, 1998.
- [54] MARCANTONIO, A., DARKEN, C., KUHN, G. M., SANTOSO, I., HANSON, S. J., AND PETSCH, T. A neural network autoassociator for induction motor failure prediction. In *Adv. Neural Inf. Process. Syst.* (1996), vol. 8, pp. 924–930.
- [55] MCCULLOCH, W., AND PITTS, W. A logical calculus of the ideas immanent in nervous activity. *Bulletin of Mathematical Biophysics* 5 (1943), 115–133.
- [56] McELIECE, R. J., POSNER, E. C., RODEMICH, E. R., AND VENKATESH, S. The capacity of the Hopfield associative memory. *IEEE Transactions on Information Theory* 1 (1987), 33–45.
- [57] MINKOWSKI, H. *Gesammelte Abhandlungen*. Teubner Verlag, Leipzig-Berlin, 1911.
- [58] MINSKY, M., AND PAPERT, S. *Perceptrons*. MIT Press, Cambridge, MA, 1969.
- [59] MURAKAMI, K., AND AIBARA, T. An improvement on the moore-penrose generalized inverse associative memory. *IEEE Transactions on Systems, Man, and Cybernetics SMC-17*, 4 (July/August 1987), 699–707.
- [60] NAKANO, K. Associatron: A model of associative memory. *IEEE Trans. on Systems, Man, Cybernetics SMC-2* (1972), 380–388.
- [61] NELSON, M. M., AND ILLINGWORTH, W. *A Practical Guide to Neural Nets*. Addison-Wesley Publishing Company, 1994.
- [62] OKAJIMA, K., TANAKA, S., AND FUJIWARA, S. A heteroassociative memory with feedback connection. In *Proceedings of the IEEE First International Conference on Neural Networks* (San Diego, 1987), vol. II, pp. 711–718.

- [63] OTSU, N. A threshold selection method from gray-level histograms. *IEEE Transactions on Systems, Man, and Cybernetics* 9, 1 (1979), 62–66.
- [64] PAO, Y. H. *Adaptive Pattern Recognition and Neural Networks*. Addison-Wesley, Reading, MA, 1989.
- [65] PEDRYCZ, W. Heterogeneous fuzzy logic networks: Fundamentals and development studies. *IEEE TRANSACTIONS ON NEURAL NETWORKS* 15, 6 (NOV 2004), 1466–1481.
- [66] PERFITTI, R. A synthesis procedure for brain-state-in-a-box neural networks. *IEEE Transactions on Neural Networks* 6, 5 (Sept 1995), 1071–1080.
- [67] PERSONNAZ, L., GUYON, I., AND DREYFUS, G. Information storage and retrieval in spin glass like neural networks. *Journal of Physics Letter* 46 (1985), L359–L365.
- [68] RADUCANU, B., GRAÑA, M., AND ALBIZURI, X. F. Morphological scale spaces and associative morphological memories: Results on robustness and practical applications. *Journal of Mathematical Imaging and Vision* 19, 2 (2003), 113–131.
- [69] RITTER, G. X. Image algebra with applications. Unpublished manuscript, available via anonymous ftp from <ftp://ftp.cis.ufl.edu/pub/src/ia/documents>, 1997.
- [70] RITTER, G. X., DE LEON, J. L. D., AND SUSSNER, P. Morphological bidirectional associative memories. *Neural Networks* 6, 12 (1999), 851–867.
- [71] RITTER, G. X., LI, D., AND WILSON, J. N. Image algebra and its relationship to neural networks. In *Technical Symposium Southeast on Optics, Electro-Optics, and Sensors* (Orlando, FL, Mar. 1989), Proceedings of SPIE.
- [72] RITTER, G. X., SHRADER-FRECHETTE, M. A., AND WILSON, J. N. Image algebra: A rigorous and translucent way of expressing all image processing operations. In *Technical Symposium Southeast on Optics, Electro-Optics, and Sensors* (Orlando, FL, May 1987), Proceedings of SPIE.
- [73] RITTER, G. X., AND SUSSNER, P. An introduction to morphological neural networks. In *Proceedings of the 13th International Conference on Pattern Recognition* (Vienna, Austria, 1996), pp. 709–717.
- [74] RITTER, G. X., AND SUSSNER, P. Morphological neural networks. In *Intelligent Systems: A Semiotic Perspective; Proceedings of the 1996 International Multidisciplinary Conference* (Gaithersburg, Maryland, 1996), pp. 221–226.
- [75] RITTER, G. X., AND SUSSNER, P. Morphological perceptrons. In *ISAS'97, Intelligent Systems and Semiotics* (Gaithersburg, Maryland, 1997).
- [76] RITTER, G. X., SUSSNER, P., AND DE LEON, J. L. D. Morphological associative memories. *IEEE Transactions on Neural Networks* 9, 2 (1998), 281–293.

- [77] RITTER, G. X., AND URCID, G. Lattice algebra approach to single-neuron computation. *IEEE Transactions on Neural Networks* 14, 2 (March 2003), 282–295.
- [78] RITTER, G. X., AND WILSON, J. N. Image algebra: A unified approach to image processing. In *Medical Imaging* (Newport Beach, CA, Feb. 1987), vol. 767 of *Proceedings of SPIE*.
- [79] RITTER, G. X., AND WILSON, J. N. *Handbook of Computer Vision Algorithms in Image Algebra*, 2 ed. CRC Press, Boca Raton, 2001.
- [80] RITTER, G. X., WILSON, J. N., AND DAVIDSON, J. L. AFATL standard image algebra. Tech. Rep. TR 87–04, University of Florida CIS Department, Oct. 1987.
- [81] RITTER, G. X., WILSON, J. N., AND DAVIDSON, J. L. Image algebra: An overview. *Computer Vision, Graphics, and Image Processing* 49, 3 (Mar. 1990), 297–331.
- [82] ROSENBLATT, F. The perceptron: A probabilistic model for information storage and organization in the brain. *Psychological Review* 65 (1958), 386–408.
- [83] RUSSELL, S. J., AND NORVIG, P. *Artificial Intelligence: A Modern Approach*, 2nd edition ed. Prentice Hall, December 2002.
- [84] SERRA, J. Mathematical morphology and cmm : a historical overview. Available at: <http://cmm.enscm.fr/Recherche/pages/nav0b.htm>.
- [85] SERRA, J. *Image Analysis and Mathematical Morphology*. Academic Press, London, 1982.
- [86] SERRA, J. *Image Analysis and Mathematical Morphology, Volume 2: Theoretical Advances*. Academic Press, New York, 1988.
- [87] SOILLE, P. *Morphological Image Analysis*. Springer Verlag, Berlin, 1999.
- [88] STEINBRUCH, K. Die lernmatrix. *Kybernetik 1* (1961), 36–45.
- [89] STILES, G., AND DENQ, D. On the effect of noise on the moore-penrose generalized inverse associative memory. *IEEE Transaction on Pattern Analysis and Machine Intelligence PAMI*–7, 3 (May 1985), 358–360.
- [90] SUSSNER, P. Fixed points of autoassociative morphological memories. In *Proceedings of the International Joint Conference on Neural Networks* (Como, Italy, July 2000), pp. 611–616.
- [91] SUSSNER, P. Observations on morphological associative memories and the kernel method. *Neurocomputing* 31 (Mar. 2000), 167–183.
- [92] SUSSNER, P. A relationship between binary morphological autoassociative memories and fuzzy set theory. In *Proceedings of the IEEE/INNS International Joint Conference on Neural Networks 2001* (Washington, July 2001).
- [93] SUSSNER, P. Associative morphological memories based on variations of the kernel and dual kernel methods. *Neural Networks* 16, 5 (July 2003), 625–632.

- [94] SUSSNER, P. Generalizing operations of binary morphological autoassociative memories using fuzzy set theory. *Journal of Mathematical Imaging and Vision* 9, 2 (Sept. 2003), 81–93. Special Issue on Morphological Neural Networks.
- [95] SUSSNER, P. New results on binary auto- and heteroassociative morphological memories. In *Proceedings of the International Joint Conference on Neural Networks 2005* (Montreal, Canada, 2005), pp. 1199–1204.
- [96] SUSSNER, P., AND VALLE, M. A brief account of the relations between gray-scale mathematical morphologies. In *Proceedings of the Brazilian Symposium on Computer Graphics and Image Processing (SIBGRAPI)* (Natal, Brazil, October 2005), pp. 79 – 86.
- [97] SUSSNER, P., AND VALLE, M. Gray-scale morphological associative memories. Accepted for publication in *IEEE Transactions on Neural Networks*, June 2005.
- [98] SUSSNER, P., AND VALLE, M. Implicative fuzzy associative memories. Accepted for publication in *IEEE Transactions on Fuzzy Systems*, May 2005.
- [99] TANK, D., AND HOPFIELD, J. J. Collective computation in neuronlike circuits. *Scientific American* 257, 6 (December 1987), 104–114.
- [100] TAYLOR, W. Electrical simulation of some nervous system functional activities. *Information Theory* 3 (1956), 314–328.
- [101] TREFETHEN, L. N., AND BAU III, D. *Numerical Linear Algebra*. SIAM Publications, Philadelphia, PA, 1997.
- [102] VALLE, M., AND SUSSNER, P. IFAMs - memórias associativas baseadas no aprendizado nebuloso implicativo. In *Anais do VII Congresso Brasileiro de Redes Neurais* (Natal, October 2005).
- [103] VALLE, M., SUSSNER, P., AND GOMIDE, F. Introduction to implicative fuzzy associative memories. In *Proceedings of the IEEE International Joint Conference on Neural Networks* (Hungary, July 2004), pp. 925 – 931.
- [104] VALLE, M. E. MATLAB Source Code for Associative Memory Models, Available at <http://www.ime.unicamp.br/~mevalle/>.
- [105] VAN HEERDEN, P. J. A new optical method of storing and retrieving information. *Appl. Opt.* 2 (1963), 387–392.
- [106] VAN HEERDEN, P. J. Theory of optical information storage in solids. *Appl. Opt.* 2 (1963), 393–400.
- [107] WIDROW, W., AND HOFF, M. Adaptive switching circuits. *WESCON Convention Record* (1960), 96–104.
- [108] WILSON, S. Morphological networks. In *Visual Communication and Image Processing IV* (Philadelphia, PA, Nov. 1989), Proceedings of SPIE.

- [109] XU ZB, LEUNG Y, H. X. Assymmetric bidirectional associative memories. *IEEE Trans. on Systems, Man, and Cybernetics* 24, 10 (OCT 1994), 1558–1564.
- [110] YEUNG, D., AND CHOW, C. Parzen window network intrusion detectors. In *Int. Conf. Pattern Recognit.* (2002), pp. 385–388.
- [111] ZHANG, B.-L., ZHANG, H., AND GE, S. S. Face recognition by applying wavelet subband representation and kernel associative memory. *IEEE Transactions on Neural Networks* 15, 1 (Jan. 2004), 166–177.
- [112] ZHANG, H., HUANG, W., HUANG, Z., AND ZHANG, B. A kernel autoassociator approach to pattern classification. *IEEE Transactions on Systems, Man and Cybernetics, Part B* 35, 3 (June 2005), 593– 606.