

Universidade Estadual de Campinas
Faculdade de Engenharia Elétrica
Departamento de Engenharia da Computação e Automação Industrial

MODELOS PARAMÉTRICOS E NÃO-PARAMÉTRICOS DE REDES NEURAIS ARTIFICIAIS E APLICAÇÕES

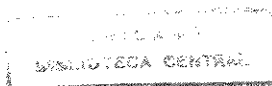
FERNANDO JOSÉ VON ZUBEN

Orientador: PROF. DR. MÁRCIO LUIZ DE ANDRADE NETTO
(DCA-FEE-UNICAMP)

Este exemplar corresponde à versão final da tese
defendida por Fernando José Von Zuben
e aprovada pela Comissão
Julgadora em 27 / 02 / 96.
Orientador

Tese apresentada à Faculdade de Engenharia Elétrica da
Unicamp (FEE-UNICAMP) como parte dos requisitos
exigidos para obtenção do título de Doutor em
Engenharia Elétrica.

CAMPINAS
Estado de São Paulo - Brasil
Fevereiro de 1996



UNIDADE:	01
N.º CHAMADA:	T/UNICAMP
	Z81m
V.	E.
Y	CPD BL/27618
RAZ.	667/96
C	☐
D	☒
PREÇO	R\$ 11,00
DATA	03/03/96
N.º CPD	

CM-00087502-1

FICHA CATALOGRÁFICA ELABORADA PELA
BIBLIOTECA DA ÁREA DE ENGENHARIA - BAE - UNICAMP

Z81m Zuben, Fernando José Von
Modelos paramétricos e não-paramétricos de redes
neurais artificiais e aplicações / Fernando José Von
Zuben.--Campinas, SP: [s.n.], 1996.

Orientador: Márcio Luiz de Andrade Netto.
Tese (doutorado) - Universidade Estadual de Campinas
Faculdade de Engenharia Elétrica.

1. Redes neurais (Computação). 2. Teoria da
aproximação. 3. Otimização matemática. 4. Modelos
matemáticos. I. Netto, Márcio Luiz de Andrade. II.
Universidade Estadual de Campinas. Faculdade de
Engenharia Elétrica. III. Título.

Dedico a meus pais, Loris e Claudio,
e irmãos, Paula Maria e Claudio José

Agradecimentos

Eu agradeço primeiramente ao Professor Márcio Luiz de Andrade Netto, que há seis anos tem sido meu orientador, período em que sempre pude contar com amplo apoio, indispensáveis ensinamentos, profícuas discussões e ambiente altamente propício para o desenvolvimento de pesquisa acadêmica. Foi o Professor Márcio quem me introduziu na área de redes neurais artificiais, tema central da tese. O clima de companheirismo e respeito mútuo existente entre os colegas de departamento foi importante na criação de condições de trabalho adequadas. Resultados importantes apresentados na tese não teriam sido obtidos sem o acesso a material bibliográfico gentilmente fornecido pelo Professor Fernando Antonio Campos Gomide. Sou grato à Faculdade de Engenharia Elétrica da Unicamp, pela oportunidade de realizar o curso de Doutorado, e aos responsáveis pela chefia do Departamento de Engenharia da Computação e Automação Industrial durante o período de desenvolvimento deste trabalho, por permitirem o acesso às instalações. Este trabalho contou com o suporte financeiro do Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPq), processo nº 142615/93-5.

Índice

Resumo	ix
Abstract.....	xi

1 Introdução	1
1.1 Redes neurais artificiais	1
1.2 A razão para a utilização de redes neurais artificiais.....	2
1.3 Motivação para o estudo de redes neurais artificiais	3
1.4 Organização da tese e problemas abordados	4

Parte I: Análise e síntese de modelos de aproximação

2 Interpolação e aproximação polinomial.....	7
2.1 Interpolação polinomial unidimensional	8
2.1.1 A motivação para o uso de funções polinomiais.....	8
2.2 O princípio da unisolvência	9
2.3 A interpolação como um processo de aproximação	13
2.3.1 Condições de convergência	14
2.4 Interpolação polinomial multidimensional	15
2.5 Aproximação polinomial unidimensional.....	16
2.5.1 Comparação com a expansão de Taylor.....	16
2.5.2 Aproximações compostas	17
2.5.2.1 Aproximação simultânea da função e suas derivadas	17
2.5.2.2 Interpolação e aproximação simultâneas	18
2.5.3 Taxas de convergência de processos de aproximação polinomial	18
2.5.4 Aproximação unidimensional utilizando splines polinomiais	20
2.6 Aproximação multidimensional	22
2.6.1 O Teorema de Stone-Weierstrass	22
2.7 Aproximação linear em espaços normados	24

2.7.1 O conceito de melhor aproximação linear	25
2.7.2 A unicidade da melhor aproximação linear.....	26
2.8 Aproximação linear em espaços com produto interno.....	26
2.8.1 A operação de projeção realizada pelo produto interno.....	27
2.8.2 Sistemas ortonormais	28
2.8.3 Funções ortonormais baseadas em polinômios de Hermite	29
2.8.4 Expansão ortonormal: série de Fourier generalizada	29
2.8.5 Expansão ortonormal truncada: a aproximação de quadrados mínimos	30
2.9 Aproximação não-linear: Teorema de Kolmogorov	32
3 Conceitos adicionais da teoria de aproximação de funções	35
3.1 Regressão, aproximação e interpolação	35
3.2 O problema geral de aproximação	36
3.3 A suavidade da função a ser aproximada	38
3.4 A classe de funções de aproximação.....	40
3.5 O pré-processamento dos dados de entrada-saída.....	41
3.6 Avaliação do nível de aproximação	42
3.7 Aproximação suave: regularização	43
3.7.1 O problema de aproximação regularizado	44
3.7.2 Fórmula geral de solução.....	45
3.8 Extensões do problema de aproximação regularizado	48
3.8.1 O problema de aproximação regularizado aditivo	48
3.8.2 A redução no número de parâmetros	49
3.8.3 O pré-processamento das variáveis de entrada	50
3.9 Generalização e sobreajuste	51
4 Aproximação de funções utilizando modelos de projeção.....	55
4.1 Aproximação por expansão ortogonal aditiva	56
4.2 O problema de aproximação resultante.....	59
4.3 Condições para aproximação exata.....	61
4.4 Busca da direção inicial de projeção (<i>projection pursuit</i>).....	62
4.4.1 A motivação para a definição de um índice de projeção	64
4.4.2 Índices de projeção baseados apenas nos dados de entrada	65

4.4.2.1 Índices derivados de análise multivariada	65
4.4.2.2 Índices para busca exploratória.....	66
4.4.3 Índices de projeção completos	67
4.4.3.1 Regressão inversa.....	67
4.4.3.2 Decomposição espectral da média da hessiana	68
4.4.3.3 Derivadas médias.....	69
4.4.3.4 Índice de projeção para associações lineares	70
4.4.4 Casos particulares	71
4.4.4.1 Tomografia Computadorizada	71
4.4.4.2 Partição do vetor de entrada.....	71
4.5 Determinação da função de expansão ortogonal	72
4.5.1 Solução paramétrica.....	73
4.5.2 Solução não-paramétrica	74
4.5.2.1 Splines polinomiais suavizantes.....	74
4.5.2.2 Validação cruzada ordinária.....	75
4.5.2.3 Validação cruzada generalizada	76
4.5.2.4 Obtenção do parâmetro de regularização	77
4.6 O processo de ajuste retroativo	78
4.7 Extensões da aproximação por expansão ortogonal	81
4.7.1 Expansão ortogonal por composição multiplicativa	81
4.7.2 Expansão ortogonal a múltiplas direções	81
4.7.3 O tratamento de problemas mais complexos	82

Parte II: Modelos de aproximação baseados em redes neurais artificiais

5 Redes neurais como modelos de aproximação de funções.....	83
5.1 Capacidade de aproximação e propriedades de convergência.....	84
5.2 Redes neurais artificiais paramétricas.....	85
5.2.1 O problema de aproximação paramétrico.....	88
5.2.2 A capacidade de aproximação universal.....	88
5.2.3 Propriedades de convergência	89
5.2.4 O caso de funções de ativação não-sigmoidais.....	91
5.3 Redes neurais artificiais não-paramétricas via métodos construtivos	91

5.3.1 A motivação para o uso de métodos construtivos	92
5.3.2 O tratamento de múltiplas saídas	94
5.3.3 O algoritmo de aproximação construtivo	95
5.3.4 Procedimentos de aperfeiçoamento do algoritmo.....	97
5.3.4.1 Definição das condições iniciais do algoritmo a partir dos dados	98
5.3.4.2 A iteração em dois grupos de variáveis e solução fechada para o terceiro grupo	98
5.3.4.3 Aproximação e critérios de parada baseados em validação cruzada	98
5.3.5 Otimização da condição de solvabilidade	99
5.3.6 A convergência do processo de aproximação construtivo	103
5.3.7 Comparação com outros métodos construtivos.....	105
5.3.7.1 Algoritmos construtivos para problemas de classificação	107
5.4 Processo de aproximação e treinamento supervisionado	107
5.5 Generalização e o efeito bias-variância	109
5.5.1 A estimativa do nível de generalização.....	110
5.6 Comparação entre a rede neural paramétrica e a não-paramétrica	112
5.7 Extensões	114
5.7.1 Função de ativação de expansão radial	115
5.7.2 O uso de mais de uma camada intermediária.....	115
5.7.3 Modelos multiplicativos com múltiplas saídas	116
6 Redes neurais recorrentes	117
6.1 Arquiteturas básicas	118
6.1.1 Modelagem de sistemas dinâmicos lineares	119
6.1.2 Extensão para o caso não-linear	120
6.1.3 Modelagem de sistemas dinâmicos não-lineares	121
6.2 O efeito da recorrência no processo de adaptação	122
6.3 Treinamento supervisionado para redes neurais recorrentes	123
6.3.1 As propriedades da função objetivo	125
6.3.2 A geração da informação de 1ª e 2ª ordem.....	126
6.3.2.1 O cálculo do vetor gradiente via algoritmo de retropropagação dinâmica	128
6.3.2.2 O cálculo do produto da matriz hessiana por um vetor.....	129
6.4 Análise da capacidade de processamento.....	131
6.5 Tipos de comportamento dinâmico em uma rede recorrente	132

6.5.1 Comportamento de regime	132
6.5.2 Sistemas de funções iterativas não-contrativas	133
6.5.3 Propriedades de estabilidade	134
7 Aplicações de redes neurais	135
7.1 Aproximação de funções	135
7.1.1 Comparação entre modelos paramétricos e não-paramétricos	136
7.1.2 Análise do método de aproximação por busca de projeção	141
7.1.3 Aproximação de funções com múltiplas saídas	142
7.1.4 Robustez a dados não-informativos	143
7.2 Classificação de padrões	144
7.3 Identificação de sistemas dinâmicos de tempo discreto	145
7.3.1 O processo de identificação	146
7.3.2 Aplicações do modelo de identificação	147
7.3.3 Comparação entre modelos de entrada-saída e por espaço de estados	148
7.4 Previsão de séries temporais	152
7.5 Redes neurais e controle adaptativo	158
7.5.1 O paradigma de controle	158
7.5.2 O conceito de realimentação	160
7.5.3 Controle moderno e controle adaptativo	160
7.5.4 Controle adaptativo utilizando redes neurais	162
7.5.5 Dificuldade de análise	163
7.5.6 Resultados	163
8 Conclusão	167
8.1 Contribuições	168
8.2 Extensões	169

Apêndices

A Conceitos básicos de álgebra linear e notação	171
A.1 Escalar	171
A.2 Vetor	171
A.3 Matriz	172

A.4 Conjuntos e operações com conjuntos	172
A.5 Conjuntos especiais de números reais	173
A.6 Espaço vetorial linear	174
A.7 Subespaço vetorial linear	174
A.8 Variedade linear	175
A.9 Conjunto convexo	175
A.10 Combinação linear e combinação convexa	175
A.11 Dependência linear	175
A.12 Base e dimensão de um espaço vetorial linear	176
A.13 Produtos entre vetores	176
A.14 Norma, semi-norma e quase-norma	177
A.15 Vetores ortogonais e ortonormais	178
A.16 Relação entre produto interno e norma euclidiana	178
A.17 Projeção de um vetor em uma determinada direção	178
A.18 Vetores ortonormais gerados a partir de vetores linearmente independentes	178
A.19 Algumas propriedades de matrizes	179
A.20 Teoremas envolvendo propriedades de matrizes	183
B Espaços normados e espaços lineares de funções	185
B.1 Espaço métrico	185
B.2 Espaço linear normado	186
B.3 Espaço de Banach	186
B.4 Espaço com produto interno	186
B.5 Espaço de Hilbert	187
B.6 Espaço linear de funções	187
B.6.1 Espaço de funções mensuráveis	188
B.6.2 Espaço de Lebesgue	188
B.6.3 Espaço de funções limitadas	189
B.6.4 Espaço de funções contínuas e uniformemente contínuas	189
B.6.5 Espaço de funções absolutamente contínuas	189
B.6.6 Espaço de Lipschitz	190
B.6.7 Espaço de funções continuamente diferenciáveis até ordem r	190
B.6.8 Espaço de Sobolev	191

B.6.9 Espaço de funções analíticas	191
B.6.10 Espaço de funções íntegras	192
B.6.11 Espaço de polinômios de ordem até M	192
B.6.12 Espaço de funções constantes	192
C Conceitos básicos de análise funcional e topológica	193
C.1 Conjunto aberto e conjunto fechado	193
C.2 Distância entre um ponto e um espaço vetorial linear normado	194
C.3 Distância entre dois espaços vetoriais lineares normados	194
C.4 Seqüência de números reais	194
C.5 Seqüência de vetores	195
C.6 Seqüência de Cauchy	195
C.7 Transformação	195
C.8 Funcional	196
C.9 Hiperplano	196
C.10 Diferencial	197
C.11 Derivada parcial e derivada total	197
C.12 Convergência ponto-a-ponto	198
C.13 Convergência uniforme	199
C.14 Critério de Cauchy para convergência uniforme	199
C.15 Convergência uniforme e continuidade	199
C.16 Convergência uniforme e diferenciação	200
C.17 Funcionais contínuos e uniformemente contínuos	200
C.18 Módulo de continuidade	201
C.19 Série de Taylor	202
C.20 Teorema diferencial do valor médio	203
C.21 Funcional convexo e quase-convexo	203
D Métodos de otimização paramétrica não-linear irrestrita	205
D.1. Avaliação do nível de aproximação paramétrica	206
D.2. Modelos de aproximação paramétrica	206
D.3. Algoritmos de otimização paramétrica utilizando diferenciação	208
D.3.1. O método do gradiente	209

D.3.2. O método de Newton	211
D.3.3. O método do gradiente conjugado	214
D.3.3.1. Motivação para a origem do método.....	214
D.3.3.2. Método das direções conjugadas.....	215
D.3.3.3. Método do Gradiente Conjugado.....	217
D.3.3.4. Aplicação a problemas não-quadráticos	218
D.3.3.5. A positivação da hessiana	219
D.3.3.6. A garantia de ajustes minimizantes.....	220
D.3.3.7. Cálculo exato da informação de 2ª ordem	220

Bibliografia

Bibliografia	223
---------------------------	------------

Resumo

Esta tese apresenta métodos de análise e síntese de modelos paramétricos e não-paramétricos de redes neurais artificiais, utilizando resultados derivados da teoria de aproximação de funções e análise numérica. A flexibilidade destas estruturas conexionistas não-lineares é explorada com base em técnicas de regularização e métodos de otimização não-linear irrestrita. A rede neural não-paramétrica resultante realiza mapeamentos não-lineares estáticos via métodos construtivos caracterizados por redução de dimensionalidade e propriedades de aproximação bem-definidas. Estruturas genéricas de processamento dinâmico não-linear podem ser obtidas via redes neurais multicamadas recorrentes, que são modelos paramétricos tendo redes neurais multicamadas não-recorrentes como caso particular. É investigado um conjunto de problemas não-lineares, cujas soluções são formuladas de modo a permitir a aplicação direta dos modelos de redes neurais desenvolvidos.

Abstract

This thesis presents methods of analysis and synthesis of parametric and nonparametric artificial neural network models, using results from approximation theory and numerical analysis. The flexibility of these nonlinear connectionist structures is explored based on regularization techniques and nonlinear unconstrained optimization methods. The resulting nonparametric neural network performs arbitrary nonlinear and static mappings via constructive methods characterized by dimensionality reduction and well-defined approximation properties. Generic nonlinear dynamic processing structures can be obtained via recurrent multilayer neural networks, parametric models having nonrecurrent multilayer neural networks as a particular case. A set of nonlinear problems is investigated, and solutions are formulated so that the developed neural network models can be directly employed.

Capítulo 1

Introdução

O expressivo desenvolvimento observado ultimamente na formalização e aplicação de conceitos conexionistas representa uma tentativa de restabelecer associações entre campos de atuação científica que até então vinham se desenvolvendo de forma independente: inteligência artificial, teoria de informação e teoria de controle. Apesar de esta iniciativa de tratar unificadamente controle e informação nos organismos vivos e nas máquinas ter sido formalizada já nos anos 40 por Norbert Wiener (WIENER, 1948), foi com o advento dos computadores com alto poder de processamento e armazenagem de informação que se viabilizou o desenvolvimento da cibernética, sendo que redes neurais artificiais vêm desempenhando um papel significativo neste processo.

O estímulo inicial que conduziu ao desenvolvimento de modelos matemáticos de redes neurais, denominados redes neurais artificiais, foi um esforço para entender mais detalhadamente o funcionamento do cérebro. O objetivo era construir um mecanismo que operasse de modo similar, ou seja, que calculasse, aprendesse, lembrasse e otimizasse da mesma forma que o cérebro humano. Tomando como base os protótipos até aqui desenvolvidos, é de consenso geral que este objetivo ainda está longe de ser atingido.

No entanto, diversas áreas de atuação científica vêm adotando e aperfeiçoando estes protótipos no papel de ferramentas computacionais aplicadas à solução dos mais variados tipos de problema, justificando o forte crescimento de atividades na área de redes neurais artificiais.

1.1 Redes neurais artificiais

Redes neurais artificiais produzem classes particulares de modelos paramétricos e não-paramétricos, representadas pela interconexão de unidades de processamento não-lineares estáticas, com entrada vetorial e saída escalar, denominadas neurônios artificiais. Os coeficientes que caracterizam a intensidade de conexão entre os neurônios são denominados

pesos. Quando algum tipo de comportamento dinâmico deve estar presente, são incluídos elementos de memória em algumas etapas do processamento, como integradores no caso contínuo e atrasadores de tempo no caso discreto.

Os modelos mais simples de uma rede neural assumem neurônios com funções de ativação idênticas, quase sempre do tipo sigmoidal, representando um efeito de saturação. No entanto, para aumentar o poder de representação sem requerer um aumento proporcional da quantidade de neurônios e conexões, é fundamental permitir que diferentes neurônios possam apresentar funções de ativação distintas e específicas para cada aplicação.

Para uma introdução mais detalhada ao tema desta seção, inclusive evidenciando as motivações biológicas da teoria conexionista, recomenda-se uma leitura preliminar dos textos de HECHT-NIELSEN (1990) e HERTZ *et al.* (1991). Além disso, VON ZUBEN (1993) oferece um texto introdutório particularmente adequado aos propósitos deste estudo.

1.2 Razão para a utilização de redes neurais artificiais

Redes neurais artificiais de processamento numérico, arquitetura em camadas e fluxo de informação com e sem realimentação produzem estruturas de processamento de sinais com grande poder de adaptação e capacidade de representação não-linear. A presença de realimentação cria a possibilidade de armazenar informações na forma de representações internas, além de introduzir dinâmica no processamento.

Sempre que métodos tradicionais, derivados da teoria de sistemas lineares, não apresentarem um desempenho satisfatório no tratamento de problemas envolvendo, por exemplo, aproximação de funções multivariáveis ou identificação e controle de sistemas dinâmicos, redes neurais artificiais do tipo descrito acima despontam como alternativas bastante competitivas, devido principalmente à sua capacidade de representação de comportamentos não-lineares arbitrários.

No entanto, para explorar adequadamente esta capacidade de representação, é necessário desenvolver algoritmos de treinamento que permitam tratar eficientemente o poder de adaptação presente nessas estruturas. Para tanto, os mais eficientes métodos de otimização não-linear, bem como resultados derivados de teoria de sistemas não-lineares, são comumente empregados.

É possível antecipar as seguintes características de caráter consensual presentes em redes neurais artificiais:

- capacidade de processamento paralelo;
- capacidade de representação distribuída, aumentando a tolerância a falhas de componentes;
- possibilidade de implementação utilizando "hardware" analógico;
- grande capacidade de adaptação;
- flexibilidade no atendimento de critérios de generalização;
- possibilidade de incorporação de poderosas ferramentas algébricas, facilitando a análise e interpretação dos resultados.

No entanto, estas características não são exclusivas das redes neurais artificiais, já que estão presentes, ao menos em parte, em inúmeros outros modelos aplicados aos mesmos propósitos, incluindo modelos lineares.

É certo que, quando comparado com os tradicionais modelos lineares, qualquer modelo não-linear adequadamente projetado é potencialmente capaz de produzir um melhor desempenho. Esta afirmação, apesar de verdadeira, não contempla nenhum critério de ordem prática, justificando o predomínio histórico de abordagens lineares no tratamento de problemas inerentemente não-lineares.

Conclui-se, então, que qualquer iniciativa de reverter este quadro, ou seja, de explorar a maior capacidade potencial das abordagens não-lineares, incluindo aquelas baseadas em redes neurais artificiais, deve invariavelmente estar associada a motivações de ordem prática. É por esta razão que o atual crescimento de interesse no estudo e implementação de abordagens conexionistas é justificado com base no aumento proporcional de poder de processamento computacional a baixo custo.

1.3 Motivação para o estudo de redes neurais artificiais

A crescente disponibilidade de recursos computacionais tem contribuído decisivamente na viabilização de problemas de síntese de modelos paramétricos e não-paramétricos de redes neurais artificiais submetidas a processos de treinamento supervisionado. No entanto, apesar de ser fundamental para a correta avaliação das potencialidades destes modelos conexionistas e definição das motivações que levam a aplicá-los em substituição a modelos tradicionais, a capacidade de análise não tem acompanhado a evolução verificada no desenvolvimento de modelos de redes neurais artificiais passíveis de implementação computacional.

Este desequilíbrio verificado nas fases de análise e síntese é o responsável pela caracterização de redes neurais artificiais como modelos de processamento do tipo *caixa preta*,

ou seja, capazes de conduzir à solução de variados tipos de problemas não-lineares sem, todavia, permitir interpretar adequadamente os resultados ou verificar a evolução do processo adaptativo associado.

A análise de redes neurais artificiais com base em resultados da teoria de aproximação de funções vem colaborar para reduzir este desequilíbrio, ampliando, assim, a capacidade de interpretação dos resultados. Uma consequência imediata deste estudo é a possibilidade de desenvolver modelos não-paramétricos de redes neurais artificiais, viabilizando a definição automática dos recursos computacionais necessários para a solução de cada problema de aplicação. Contudo, os modelos paramétricos de redes neurais artificiais continuam a desempenhar um papel significativo, principalmente no tratamento de processamento dinâmico não-linear.

Em termos de aplicação, modelos paramétricos e não-paramétricos de redes neurais artificiais podem ser empregados em processos de solução de diversos tipos de problemas não-lineares, desde que a solução possa ser mapeada na forma de estruturas adaptativas adequadamente projetadas para utilizar eficientemente o conjunto de informações disponíveis associado a cada problema.

1.4 Organização da tese e problemas abordados

Os resultados principais da tese, geralmente apresentados na forma de definições e teoremas, utilizam conceitos básicos apresentados, por sua vez, nos apêndices. Portanto, para evitar qualquer dificuldade de entendimento, principalmente com relação à notação, recomenda-se uma leitura preliminar do conteúdo dos apêndices. Referências à literatura são utilizadas indistintamente em substituição a detalhes e resultados não apresentados no texto.

Com o objetivo de posicionar adequadamente os métodos de aproximação a serem desenvolvidos nos capítulos subseqüentes, o capítulo 2 se ocupa exclusivamente da formalização matemática de conceitos de interpolação e aproximação de funções puramente determinísticas a partir de amostras imunes a ruído. Utilizando o fato de que, neste contexto, a interpolação é uma alternativa para o problema de aproximação, as limitações presentes no processo de aproximação por interpolação são exploradas para evidenciar a necessidade de se desenvolver a teoria de aproximação como um ramo distinto da teoria de interpolação. São abordados os casos unidimensional e multidimensional.

No capítulo 3, é formalizado e analisado o problema de aproximação, em uma região compacta, de uma função determinística multivariável desconhecida, a partir de um número finito de dados amostrados. Além de o processo de amostragem estar sujeito a ruído, a utilização de um número finito de dados requer a definição de processos de aproximação regularizados, capazes de encontrar a melhor aproximação com base em conjuntos de dados de aproximação e teste. Resulta um problema de aproximação bem-comportado e capaz de atender a critérios de generalização. A formulação apresentada se baseia no fato de que o tipo de suavidade imposta à função a ser aproximada vai determinar o tipo de modelo de aproximação a ser utilizado, já que o compromisso entre suavidade do processo de aproximação e capacidade de aproximação é completamente análogo ao compromisso existente entre dependência da dimensionalidade e flexibilidade do modelo de aproximação.

O modelo de aproximação regularizado obtido no capítulo 3 é, então, melhor explorado no capítulo 4, na forma de um modelo de projeção por composição aditiva, um caso típico de modelo de redução de dimensionalidade. Os termos da composição aditiva correspondem a funções escalares de expansão ortogonal a uma direção de projeção a ser determinada a partir dos dados disponíveis. Após a apresentação de métodos para busca automática da direção de projeção inicial, procedimentos paramétricos e não-paramétricos de determinação da função de expansão ortogonal são, então, discutidos, tendo como etapa fundamental a aplicação de validação cruzada. A interdependência entre a forma das funções de expansão ortogonal e a correspondente direção de projeção acaba conduzindo a processos iterativos de aproximação, sendo que, a cada iteração, uma transformação deve ser aplicada ao conjunto de dados para que a informação já representada seja removida, permitindo o reinício do processo a partir de uma nova direção de projeção. O processo iterativo é também cíclico, pois sempre que um novo termo é incluído junto à composição aditiva, os termos já existentes são reajustados em sequência até a convergência. Finalmente, extensões para o tratamento de expansão ortogonal por composição multiplicativa são abordadas.

Redes neurais artificiais passam a ser o tema central do estudo apenas a partir do capítulo 5. Primeiramente são apresentadas as estruturas básicas que compõem uma rede neural com uma camada intermediária. A capacidade de aproximação e as propriedades de convergência deste tipo de rede neural, quando empregada como modelo de aproximação paramétrico, são então discutidas. Em seguida, o modelo de aproximação de funções analisado no capítulo 4 é representado na forma de uma rede neural com uma camada intermediária. No entanto, o modelo não-paramétrico resultante não corresponde simplesmente a uma nova

representação de um procedimento de aproximação de funções na forma de uma rede neural. A nova forma de representação do modelo de projeção permite visualizar propriedades algébricas fundamentais para a implementação de um processo de aproximação alternativo, capaz de explorar eficientemente a geometria do problema de aproximação multidimensional e com propriedades de convergência bem definidas. O processo de aproximação resultante apresenta etapas construtivas e de poda, sendo um método eficaz de determinação do número mínimo de neurônios na camada intermediária necessário para resolver a tarefa de aproximação. Um estudo do problema de generalização junto ao processo de aproximação resultante e uma comparação entre os modelos paramétrico e não-paramétrico de redes neurais com uma camada intermediária concluem o capítulo.

Redes neurais recorrentes como modelos paramétricos de aproximação de mapeamentos dinâmicos de tempo discreto são, então, apresentadas no capítulo 6. O processo de aproximação para redes neurais recorrentes é descrito com base em conceitos de análise de sensibilidade, enquanto que a análise da capacidade de processamento de informação é fundamentada em conceitos da teoria de sistemas de funções iterativas. O processo de treinamento supervisionado resultante é válido para redes neurais paramétricas tanto recorrentes como não-recorrentes, envolvendo processos iterativos de otimização não-linear irrestrita utilizando diferenciação até 2ª ordem.

Resultados da aplicação dos modelos de redes neurais artificiais determinísticas para aproximação de mapeamentos estáticos e dinâmicos são descritos no capítulo 7. São abordados problemas de aproximação de funções, classificação de padrões, identificação de sistemas dinâmicos, previsão de séries temporais e controle de processos. Apesar de exemplos de simulação serem incluídos, o objetivo principal deste capítulo é formular cada problema no sentido de permitir a aplicação direta dos modelos e respectivos processos de aproximação desenvolvidos principalmente nos capítulos 5 e 6.

O capítulo 8 conclui o estudo fazendo uma análise geral dos resultados apresentados e das principais contribuições deste trabalho. São também sugeridas possíveis extensões a serem investigadas em trabalhos futuros.

Parte I

Análise e síntese de modelos de aproximação

Capítulo 2

Interpolação e aproximação polinomial

As funções constituem as ferramentas matemáticas básicas para descrição e análise de processos físicos. Enquanto em alguns casos estas funções são conhecidas explicitamente, muitas vezes é necessário extrair um conjunto de informações acerca do processo e construir aproximações baseadas neste conjunto. Portanto, a construção de modelos de aproximação para uma função real desconhecida envolve duas etapas básicas:

- análise: extração de informação a respeito da função;
- síntese: reconstrução da função com base na informação disponível.

O processo de análise não faz parte do escopo deste estudo. Em seu lugar é sempre assumida a disponibilidade de um conjunto de informações amostradas da função a ser aproximada. Portanto, os resultados apresentados nesta tese correspondem ao problema de síntese, ou seja, de como utilizar eficientemente este conjunto de informações. A abordagem do problema de síntese pode ser:

- existencial: estudo da possibilidade de aproximação de uma função com base apenas em suas propriedades qualitativas;
- construtiva: reconstrução da função com base em suas propriedades quantitativas.

A abordagem existencial do problema de síntese independe do tipo de informação disponível, enquanto que a abordagem construtiva se baseia em um conjunto específico de informações. Neste estudo, o conjunto de informações disponíveis é assumido restrito a amostras do valor da função em pontos do espaço de aproximação, de tal forma que o problema de aproximação a ser abordado possui invariavelmente o seguinte enunciado geral:

Problema: Aproximar uma função vetorial f de argumento vetorial utilizando amostras do valor da função em pontos do espaço de aproximação.

Dependendo da função a ser aproximada, da informação disponível e do critério de aproximação especificado, este problema pode ter múltiplas soluções, uma única solução ou nenhuma solução. Sempre que uma solução existir, ela é denominada a melhor aproximação.

Todo método de investigação científica que converge para problemas deste tipo geralmente está procurando atender a objetivos genéricos do tipo:

- predição do comportamento da função em pontos não-amostrados;
- extração de informações a respeito da função original a partir da interpretação da função de aproximação.

Duas hipóteses restritivas motivadas por limitações de ordem computacional, como capacidade de memória e tempo de processamento, e por dificuldades analíticas e algébricas, como inexistência ou não-unicidade de solução de problemas de aproximação definidos a partir de um conjunto infinito de informações (DAVIS, 1975), são consideradas junto ao problema de aproximação proposto:

- o número de amostras do valor da função em pontos do espaço de aproximação (conjunto de informações disponíveis) é finito;
- o espaço de aproximação é fechado e limitado, ou seja, a aproximação vai se dar em uma região compacta.

Com o objetivo de posicionar adequadamente os métodos de aproximação a serem desenvolvidos neste estudo e fornecer subsídios para sua implementação, este capítulo se ocupa exclusivamente da formalização matemática de conceitos de interpolação e aproximação de funções puramente determinísticas a partir de amostras imunes a ruído. Utilizando o fato de que, neste contexto, a interpolação é uma alternativa para o problema de aproximação, as limitações presentes no processo de aproximação por interpolação são exploradas para evidenciar a necessidade de se desenvolver a teoria de aproximação como um ramo distinto da teoria de interpolação. São abordados os casos unidimensional e multidimensional.

2.1 Interpolação polinomial unidimensional

2.1.1 A motivação para o uso de funções polinomiais

Os polinômios têm desempenhado um papel central em teoria de aproximação e análise numérica. Conforme definido no Apêndice B, o espaço de polinômios de ordem até M (M finito) com coeficientes reais é definido no intervalo $[a,b]$ na forma:

$$P_M[a,b] = \left\{ p_M(x) = \sum_{i=0}^M c_i x^i, \quad c_0, \dots, c_M \in \mathbb{R} \text{ e } x \in [a,b] \right\}.$$

RALSTON (1965) procurou justificar a importância dos polinômios salientando as seguintes propriedades existentes em $P_M[a,b]$:

- $P_M[a,b]$ é um espaço linear de dimensão finita;
- os polinômios são funções suaves, sendo infinitamente diferenciáveis, analíticas e íntegras;
- os polinômios são de fácil avaliação e armazenagem em computadores digitais;
- a derivada e a antiderivada de um polinômio são também polinômios cujos coeficientes podem ser determinados algebricamente, mesmo por um computador;
- é possível definir taxas de convergência precisas para aproximação de funções suaves por polinômios.

Estas propriedades parecem indicar que polinômios são funções de interpolação e aproximação ideais. No entanto, $P_M[a,b]$ possui um tipo de inflexibilidade que se manifesta na forma da seguinte propriedade:

- processos de interpolação e aproximação utilizando polinômios produzem funções que oscilam excessivamente, sendo que a oscilação fica necessariamente mais pronunciada com o aumento da ordem do polinômio.

A contenção desta oscilação é fundamental para a garantia de convergência de problemas gerais de aproximação (DAVIS, 1975). Faz-se necessário, então, a substituição de polinômios por outras funções mais flexíveis. Dentre inúmeras alternativas, encontram-se as funções polinomiais por partes, as quais têm a vantagem de preservar boa parte das propriedades apresentadas acima para espaços polinomiais. Dentre as funções polinomiais por partes, destacam-se os splines polinomiais, a serem discutidos mais adiante e que serão utilizados posteriormente na solução de problemas de aproximação unidimensional.

2.2 O princípio da unisolvência

O teorema de interpolação polinomial apresentado na sequência, válido para o caso unidimensional, permite concluir que sempre existe uma reta que passa por dois pontos quaisquer no plano, uma parábola que passa por três pontos quaisquer no plano, uma cúbica que passa por quatro pontos quaisquer no plano, e assim por diante.

Teorema 2.1: Dados $n+1$ escalares reais distintos $x_0, x_1, \dots, x_n$ e $n+1$ escalares reais arbitrários $y_0, y_1, \dots, y_n$, existe um único polinômio $p_n(x)$ de ordem até n tal que

$$p_n(x_i) = y_i, \quad i = 0, 1, \dots, n$$

Prova: Tomando o polinômio $p_n(x) = c_0 + c_1x + \dots + c_nx^n$ com $n+1$ coeficientes indeterminados, resulta um sistema linear com $n+1$ equações em c_i ($i=0, 1, \dots, n$):

$$c_0 + c_1x_j + \dots + c_nx_j^n = y_j, \quad j = 0, 1, \dots, n,$$

que podem ser colocadas na forma matricial:

$$\begin{bmatrix} 1 & x_0 & x_0^2 & \cdots & x_0^n \\ 1 & x_1 & x_1^2 & \cdots & x_1^n \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_n & x_n^2 & \cdots & x_n^n \end{bmatrix} \cdot \begin{bmatrix} c_0 \\ c_1 \\ \vdots \\ c_n \end{bmatrix} = \begin{bmatrix} y_0 \\ y_1 \\ \vdots \\ y_n \end{bmatrix} \Rightarrow \mathbf{V}(x_0, x_1, \dots, x_n) \cdot \mathbf{c} = \mathbf{y}$$

Para calcular $\det(\mathbf{V}(x_0, x_1, \dots, x_n))$, conhecido como determinante de Vandermonde, considere a função:

$$\mathbf{F}(x) = \det(\mathbf{V}(x_0, x_1, \dots, x_{n-1}, x)) = \begin{vmatrix} 1 & x_0 & x_0^2 & \cdots & x_0^n \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{n-1} & x_{n-1}^2 & \cdots & x_{n-1}^n \\ 1 & x & x^2 & \cdots & x^n \end{vmatrix}.$$

Note que $\mathbf{F}(x)$ é um polinômio de ordem n . Além disso, $\mathbf{F}(x)$ se anula para $x = x_j$ ($j=0, 1, \dots, n-1$), pois resultam duas linhas idênticas no determinante. Logo,

$$\mathbf{F}(x) = d(x - x_0)(x - x_1) \cdots (x - x_{n-1}),$$

onde d é o coeficiente de x^n . Expandindo $\mathbf{F}(x)$ com base nos menores associados à última linha do determinante obtém-se $d = \det(\mathbf{V}(x_0, x_1, \dots, x_{n-1}))$. Sendo assim,

$$\mathbf{F}(x) = \det(\mathbf{V}(x_0, x_1, \dots, x_{n-1})) (x - x_0)(x - x_1) \cdots (x - x_{n-1}),$$

conduzindo à fórmula recursiva

$$\mathbf{F}(x_n) = \det(\mathbf{V}(x_0, x_1, \dots, x_{n-1})) (x_n - x_0)(x_n - x_1) \cdots (x_n - x_{n-1}).$$

Partindo de

$$\mathbf{F}(x_1) = \det(\mathbf{V}(x_0)) (x_1 - x_0) = x_1 - x_0,$$

resulta

$$\mathbf{F}(x_2) = \det(\mathbf{V}(x_0, x_1))(x_2 - x_0)(x_2 - x_1) = (x_1 - x_0)(x_2 - x_0)(x_2 - x_1)$$

e, por recursão,

$$\mathbf{F}(x_n) = \prod_{i>j}^n (x_i - x_j)$$

Como, por hipótese, os pontos $x_0, x_1, \dots, x_n$ são distintos, conclui-se que $\mathbf{F}(x_n) = \det(\mathbf{V}(x_0, x_1, \dots, x_n)) \neq 0$. Há, portanto, uma única solução para o sistema linear com $n+1$ equações em c_i ($i=0,1,\dots,n$), dada por

$$\mathbf{c} = [\mathbf{V}(x_0, x_1, \dots, x_n)]^{-1} \mathbf{y},$$

conduzindo ao polinômio $p_n(x)$ desejado. $\square$

Sendo assim, é sempre possível obter um polinômio de ordem até n que interpole $n+1$ pontos distintos arbitrários no plano, resultando em uma mesma dimensão para o conjunto de parâmetros e o conjunto de informações disponíveis. Para a obtenção de c_i ($i=0,1,\dots,n$) com valores elevados de n , em lugar de se recorrer à inversão direta da matriz $\mathbf{V}(x_0, x_1, \dots, x_n)$, é possível empregar processos iterativos baseados na fórmula de Lagrange ou na fórmula de Newton (DAVIS, 1975). Apesar de exigir uma maior quantidade de cálculo que a fórmula de Lagrange, a fórmula de Newton tem a vantagem de não precisar refazer os cálculos já realizados caso um ponto adicional seja incorporado ao problema de interpolação.

A despeito da importância deste resultado, não é possível aplicá-lo a problemas de interpolação envolvendo mais de uma variável, mesmo que o conjunto de parâmetros e o conjunto de informações disponíveis permaneçam com a mesma dimensão.

Exemplo (DAVIS, 1975):

Considere as potências em duas variáveis $q_0(x,y)=1$, $q_1(x,y)=x$, $q_2(x,y)=y$, $q_3(x,y)=x^2$, $q_4(x,y)=xy$, $q_5(x,y)=y^2$, $q_6(x,y)=x^3, \dots$. Dados $n+1$ pontos distintos (x_i, y_i) ($i=0,\dots,n$), nem sempre é possível encontrar $n+1$ parâmetros tal que a combinação linear de $q_0, q_1, \dots, q_n$ assumam valores pré-definidos nestes $n+1$ pontos.

Tendo a interpolação polinomial como um caso particular, um resultado mais geral e também restrito ao problema de interpolação unidimensional é denominado princípio da unisolvência e é derivado com base no seguinte teorema:

Teorema 2.2: Sejam $x_0, x_1, \dots, x_n \in [a, b]$ pontos distintos no intervalo finito $[a, b]$. Dados valores arbitrários $y_0, y_1, \dots, y_n \in \mathfrak{R}$ e funções $f_0, f_1, \dots, f_n: [a, b] \rightarrow \mathfrak{R}$, é possível resolver unicamente o problema de interpolação

$$\sum_{j=0}^n c_j f_j(x_i) = y_i, \quad i = 0, 1, \dots, n$$

se e somente se

$$\begin{vmatrix} f_0(x_0) & f_1(x_0) & \cdots & f_n(x_0) \\ f_0(x_1) & \ddots & & \vdots \\ \vdots & & & \\ f_0(x_n) & \cdots & & f_n(x_n) \end{vmatrix} = |f_j(x_i)| \neq 0.$$

Prova: Expressando o problema de interpolação na forma matricial tem-se

$$\begin{bmatrix} f_0(x_0) & f_1(x_0) & \cdots & f_n(x_0) \\ f_0(x_1) & \ddots & & \vdots \\ \vdots & & & \\ f_0(x_n) & \cdots & & f_n(x_n) \end{bmatrix} \cdot \begin{bmatrix} c_0 \\ c_1 \\ \vdots \\ c_n \end{bmatrix} = \begin{bmatrix} y_0 \\ y_1 \\ \vdots \\ y_n \end{bmatrix} \Rightarrow [f_j(x_i)] \cdot \mathbf{c} = \mathbf{y}.$$

É bem conhecido que este sistema de equações lineares tem sempre uma única solução $\mathbf{c} \in \mathfrak{R}^{n+1}$ para um vetor arbitrário $\mathbf{y} \in \mathfrak{R}^{n+1}$ se e somente se $|f_j(x_i)| \neq 0$, ou seja, se a matriz $[f_j(x_i)]$ for inversível. $\square$

Definição 2.1 (Princípio da unisolvência (DAVIS, 1975)): Um conjunto $\{f_0, f_1, \dots, f_n\}$ de $n+1$ funções definidas em um intervalo $[a, b]$ é denominado unisolvente em $[a, b]$ se $|f_j(x_i)| \neq 0$ para toda escolha de $n+1$ pontos distintos em $[a, b]$. De forma equivalente, $\{f_0, f_1, \dots, f_n\}$ é unisolvente em $[a, b]$ se e somente se a única combinação linear de $\{f_0, f_1, \dots, f_n\}$ que se anula em $n+1$ pontos distintos de $[a, b]$ possui todos os coeficientes nulos.

No caso unidimensional, a obtenção de conjuntos de funções unisolventes é simples. Exemplos:

- O conjunto $\{1, x-d, (x-d)^2, \dots, (x-d)^n\}$ é unisolvente em qualquer intervalo $[a, b]$ e para qualquer $d \in \mathfrak{R}$ (veja demonstração do teorema 2.1 para $d = 0$);
- Se uma função qualquer $w(x)$ não se anula em $[a, b]$, então o conjunto $\{w(x), (x-d)w(x), (x-d)^2 w(x), \dots, (x-d)^n w(x)\}$ é unisolvente em qualquer intervalo $[a, b]$ e para qualquer $d \in \mathfrak{R}$.

2.3 A interpolação como um processo de aproximação

Considere os pontos (x_i, y_i) ($i=0, \dots, n$) amostrados de uma função $f \in C^n[a, b]$, ou seja, uma função contínua e com derivadas contínuas até ordem n em $[a, b]$. Uma vez obtida a solução $p_n(x)$ do processo de interpolação polinomial, o teorema a seguir estabelece uma medida de proximidade entre $f(x)$ e $p_n(x)$ para todo $x \in [a, b]$.

Teorema 2.3 (Teorema de Cauchy para interpolação polinomial): Seja $f(x) \in C^n[a, b]$ e suponha que $f^{(n+1)}(x)$ exista para todo $x \in [a, b]$. Se $a \leq x_0 < x_1 < \dots < x_n \leq b$, então

$$f(x) - p_n(x) = \frac{(x - x_0)(x - x_1) \cdots (x - x_n)}{(n+1)!} f^{(n+1)}(\xi),$$

onde $\min(x, x_0, x_1, \dots, x_n) < \xi < \max(x, x_0, x_1, \dots, x_n)$. O valor de ξ depende de $x, x_0, x_1, \dots, x_n$ e f .

Prova: Veja ATKINSON (1989). $\square$

Como geralmente o valor de ξ não é conhecido, recorre-se a um processo de maximização. Para todo $x \in [a, b]$, considere o resíduo

$$R_n(x) = f(x) - p_n(x).$$

Então, se $f \in C^{n+1}[a, b]$,

$$|R_n(x)| \leq \frac{|x - x_0| \cdot |x - x_1| \cdots |x - x_n|}{(n+1)!} \max_{a \leq t \leq b} |f^{(n+1)}(t)|.$$

Note que se $f \in P_n[a, b]$, então $R_n(x) = 0$ para todo $x \in [a, b]$. O limitante superior para o resíduo é composto por dois termos:

- $\frac{|x - x_0| \cdot |x - x_1| \cdots |x - x_n|}{(n+1)!}$ depende apenas dos pontos de interpolação $x_0, x_1, \dots, x_n$. Para um dado conjunto de pontos de interpolação, este termo pode levar a uma estimativa muito superior ao resíduo efetivamente verificado.
- $\max_{a \leq t \leq b} |f^{(n+1)}(t)|$ depende apenas de f e é um limitante superior para $f^{(n+1)}(\xi)$. Dependendo do valor de ξ , este termo também pode levar a uma estimativa muito superior ao resíduo efetivamente verificado.

Contudo, utilizando-se os zeros do polinômio de Chebychev (DAVIS, 1975), é possível definir o menor limitante superior para o primeiro termo. Com a mudança de variável

$$x = \frac{b+a}{2} + \frac{b-a}{2}t$$

convertendo o intervalo $[a,b]$ em $[-1,1]$, o menor limitante superior para o resíduo é dado por:

$$|R_n(t)| \leq \frac{1}{2^n(n+1)!} \max_{-1 \leq t \leq 1} |f^{(n+1)}(t)|, \quad t \in [-1,1].$$

Esta medida de proximidade entre $f(x)$ e $p_n(x)$ para todo $x \in [a,b]$ tem a desvantagem de requerer o conhecimento de derivadas de f de ordem elevada, principalmente com o aumento da ordem do polinômio de interpolação, a qual está invariavelmente associada ao número de pontos de interpolação.

2.3.1 Condições de convergência

Considere uma função $f \in C[a,b]$ e um polinômio $p_n(x)$ que interpole $f(x)$ exatamente em $n+1$ pontos no intervalo $[a,b]$. Tomando uma sequência de polinômios $p_n(x)$, com $n \rightarrow \infty$, não é possível garantir que $p_n(x)$ convirja para $f(x)$ em $[a,b]$, conforme enunciado a seguir:

Teorema 2.4: Seja $[a,b]$ um intervalo fixo e $x_{n0} < x_{n1} < \dots < x_{nn}$ uma coleção de pontos no intervalo $[a,b]$ para todo inteiro $n \geq 0$. Então existe uma função $f \in C[a,b]$ tal que

$$\|f - p_n\|_{L_\infty[a,b]} \rightarrow \infty \text{ quando } n \rightarrow \infty,$$

onde p_n é o único polinômio de ordem n que interpola f nos pontos $x_{n0}, x_{n1}, \dots, x_{nn}$.

Prova: Veja CHENEY (1966). $\square$

A garantia de convergência de processos de aproximação por interpolação está necessariamente associada à imposição de alguma restrição adicional ao problema de interpolação como, por exemplo:

- Em 1922, Fejér provou que, controlando-se os valores das derivadas da função de interpolação, esta converge para $f(x)$ em $[a,b]$ quando $n \rightarrow \infty$ (FEJÉR, 1922). Como as derivadas de $p_n(x)$ crescem intrinsecamente associadas a n , devem ser empregadas outras funções de aproximação com derivadas que possam ser convenientemente controladas.
- Em 1916, Bernstein provou que, para funções contínuas tais que $\lim_{\delta \rightarrow 0} [\mu_{[a,b]}^f(\delta) \log \delta] = 0$, $p_n(x)$ converge para $f(x)$ em $[a,b]$ quando $n \rightarrow \infty$ (BERNSTEIN, 1916). Neste caso, o custo da manutenção de $p_n(x)$ como função de interpolação é a exigência de uma maior

suavidade para $f(x)$. Veja Apêndice C para definição de $\mu'_{[a,b]}(\cdot)$, denominado módulo de continuidade de f em $[a,b]$.

Os métodos de aproximação a serem apresentados neste estudo, mesmo quando não estruturados na forma de processos de interpolação, adotam técnicas derivadas do resultado de Fejér.

2.4 Interpolação polinomial multidimensional

Pelo fato de o teorema 2.2 não ser extensível a mais de uma dimensão, tem-se:

Teorema 2.5 (Teorema de Haar (BUCK, 1959)): Seja $X \subset \mathbb{R}^m$ ($m \geq 2$) um conjunto que contém um ponto interior $\mathbf{x}$. Seja $\{f_0, \dots, f_n\}$ um conjunto de funções definidas em X e contínuas em uma vizinhança de $\mathbf{x}$. Então, este conjunto de funções não pode ser unisolvente em X .

Prova: Seja U uma hipersfera centrada em $\mathbf{x}$ e contida em X suficientemente pequena de forma a garantir que as funções f_j ($j=0, \dots, n$) sejam contínuas em U . Tome $n+1$ pontos distintos $\mathbf{x}_0, \dots, \mathbf{x}_n \in U$ tal que $|f_j(\mathbf{x}_i)| \neq 0$ (se $|f_j(\mathbf{x}_i)| = 0$ o conjunto de funções é certamente não-unisolvente). Mantenha os pontos $\mathbf{x}_j$ ($j > 1$) fixos e mova os pontos $\mathbf{x}_0$ e $\mathbf{x}_1$ continuamente em U de tal forma que os dois pontos troquem de posição (observe que as posições intermediárias de $\mathbf{x}_0$ e $\mathbf{x}_1$ devem obedecer sempre $\mathbf{x}_0 \neq \mathbf{x}_1$). Com isso, é induzida uma troca de duas colunas do determinante $|f_j(\mathbf{x}_i)|$, com a conseqüente mudança de sinal (veja apêndice A). Já que as funções f_j ($j=0, \dots, n$) são contínuas em U , então existe uma posição intermediária para $\mathbf{x}_0$ e $\mathbf{x}_1$ tal que $|f_j(\mathbf{x}_i)| = 0$, ou seja, $\{f_0, \dots, f_n\}$ não é unisolvente em X . $\square$

Apesar de a obtenção de solução para o problema de interpolação multidimensional não estar necessariamente vinculada à unisolvência do conjunto $\{f_0, f_1, \dots, f_n\}$, este resultado, aliado ao problema de convergência de processos de aproximação por interpolação evidenciados na seção anterior, acaba por justificar o tratamento de problemas de aproximação sem recorrer a técnicas de interpolação.

Portanto, a partir deste ponto, o processo de reconstrução de uma função vai receber um tratamento mais geral, que independe do tipo de informação disponível. Basicamente, investiga-se a possibilidade de aproximação de uma função conhecendo-se apenas suas propriedades qualitativas.

2.5 Aproximação polinomial unidimensional

O teorema de Weierstrass, datado de 1885 e enunciado a seguir, apresenta uma alternativa ao processo de reconstrução completa da função f , sem recorrer a técnicas de aproximação por interpolação. Em linhas gerais, o teorema garante a existência de solução para o problema de aproximação polinomial unidimensional de funções contínuas definidas em um intervalo fechado, conforme ilustrado na figura 2.1.

Teorema 2.6 (Teorema de Weierstrass): Seja $f: [a,b] \subset \mathbb{R} \rightarrow \mathbb{R}$ uma função contínua ($f \in C[a,b]$). Dado um $\varepsilon > 0$, para todo $x \in [a,b]$ existe um polinômio $p_n(x)$ de ordem n suficientemente elevada e dependente de f e ε tal que

$$\|f(x) - p_n(x)\|_{L_\infty[a,b]} \leq \varepsilon.$$

Prova: Veja RUDIN (1976). $\square$

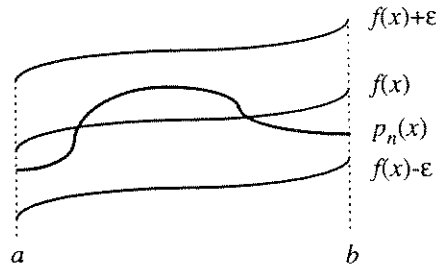


Figura 2.1: Para ε dado, as condições do Teorema de Weierstrass são atendidas por $p_n(x)$

Como continuidade e continuidade uniforme se confundem em intervalos compactos, o teorema de Weierstrass conduz ao seguinte resultado:

Corolário 2.1: Qualquer função $f \in C[a,b]$ pode ser aproximada uniformemente em $[a,b]$ por uma sequência de polinômios $\{p_n(x)\}$.

2.5.1 Comparação com a expansão de Taylor

É conveniente contrastar o resultado do teorema de Weierstrass com a expansão de Taylor (veja apêndice C), explicitando as diferenças entre os dois tipos de aproximação (DAVIS, 1975). Suponha que $f(x)$ seja analítica no intervalo $-c \leq x \leq c$. Então a expansão em série de potência $f(x) = \sum_{k=0}^{\infty} d_k x^k$ converge uniformemente para $|x| \leq c$. Logo, é evidente

que, dado $\varepsilon > 0$, é possível tomar um número n suficiente de termos da expansão de $f(x)$ que produza um polinômio $p_n(x) = \sum_{k=0}^n d_k x^k$ para o qual $|f(x) - p_n(x)| \leq \varepsilon$ para $|x| \leq c$.

Mas não há expansão em série de potência para funções não-analíticas. Mesmo assim, o teorema de Weierstrass garante que é possível aproximar uniformemente funções contínuas, inclusive as não-analíticas. Dada uma sequência $\{\varepsilon_k\} \rightarrow 0$ para $k \rightarrow \infty$, existem polinômios

$$\begin{aligned} p_{n_1}(x) &= d_{10} + d_{11}x + \dots + d_{1n_1}x^{n_1} \\ p_{n_2}(x) &= d_{20} + d_{21}x + \dots + d_{2n_2}x^{n_2} \\ &\vdots \qquad \qquad \qquad \vdots \end{aligned}$$

para os quais $|f(x) - p_{n_k}(x)| \leq \varepsilon_k$, $|x| \leq c$, $k=1,2,\dots$.

Conseqüentemente, $\lim_{k \rightarrow \infty} p_{n_k}(x) = f(x)$ uniformemente para $|x| \leq c$. Com isso, é possível expressar $f(x)$ em uma série de polinômios na forma:

$$f(x) = p_{n_1}(x) + (p_{n_2}(x) - p_{n_1}(x)) + (p_{n_3}(x) - p_{n_2}(x)) + \dots$$

que converge uniformemente para $f(x)$, com $|x| \leq c$.

Conclusão: Por expansão de Taylor, uma função analítica pode ser expandida em uma série de potência uniformemente convergente e, pelo teorema de Weierstrass, uma função contínua, analítica ou não-analítica, pode ser expandida em uma série de polinômios uniformemente convergente, cujos termos não podem ser rearranjados, ao menos no caso não-analítico, de forma a produzir uma série de potência uniformemente convergente.

2.5.2 Aproximações compostas

2.5.2.1 Aproximação simultânea da função e suas derivadas

Existem métodos de aproximação polinomial que realizam a aproximação simultânea da função e suas derivadas (DAVIS, 1975). Estes métodos geralmente têm uma taxa de convergência bem menor, mas permitem identificar propriedades qualitativas da função já nas fases iniciais do processo de aproximação. Não se deve confundir aproximação simultânea da função e suas derivadas com aproximação uniforme, já que esta não implica aquela.

2.5.2.2 Interpolação e aproximação simultâneas

Se $f \in C[a,b]$, o teorema de Weierstrass garante a aproximação uniforme de f por um polinômio. Um resultado que pode ser útil em várias aplicações é a possibilidade de interpolação simultânea de um número finito de pontos, conforme ilustrado na figura 2.2.

Teorema 2.7 (Teorema de Walsh): Seja $f: [a,b] \subset \mathbb{R} \rightarrow \mathbb{R}$ uma função definida em $[a,b]$ e uniformemente aproximável por polinômios neste intervalo. Sejam n pontos distintos $x_1, x_2, \dots, x_n \in [a,b]$. Então f é uniformemente aproximável por polinômios p que satisfazem as condições auxiliares

$$p(x_i) = f(x_i), \quad i = 1, \dots, n \quad (n \text{ finito}).$$

Prova: Veja DAVIS (1975). $\square$

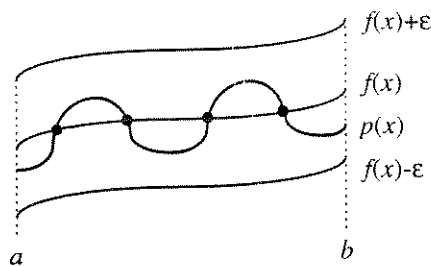


Figura 2.2: O polinômio $p(x)$ realiza interpolação e aproximação simultâneas

2.5.3 Taxas de convergência de processos de aproximação polinomial

Apesar de o teorema de Weierstrass estabelecer que toda função contínua em um intervalo fechado pode ser aproximada uniformemente com precisão arbitrária utilizando-se um polinômio, a ordem deste polinômio não é determinada. De fato, quanto menos suave for a função f , maior deve ser a ordem do polinômio de aproximação, o que conduz a seguinte questão:

- dependendo do grau de suavidade da função f , quão bem esta função pode ser aproximada por um polinômio de ordem fixa?

O estudo deste problema conduz a uma outra questão fundamental:

- com que taxa o erro de aproximação decresce com o aumento da ordem do polinômio?

Este é um problema que vem sendo investigado há muito tempo, sendo que alguns dos resultados já obtidos são enunciados a seguir.

Definição 2.2: Seja E um espaço linear normado com norma $\|\cdot\|_E$. Dado um subespaço linear M -dimensional $E_M \subseteq E$, a distância de $f \in E$ para E_M é definida na forma:

$$\text{dist}(f, E_M)_E = \inf_{g \in E_M} \|f - g\|_E.$$

Em processos de aproximação de funções suaves $f: [a, b] \rightarrow \Re$ por polinômios, é possível determinar, em função de M , o menor limitante superior para a distância em norma de Lebesgue entre a função a ser aproximada e o espaço de polinômios de ordem até M , dada por $\text{dist}(f, P_M[a, b])_{L_p[a, b]}$. Para todo $x \in [a, b]$ tem-se:

- Se $f \in L'_p[a, b]$ então $\text{dist}(f, P_M[a, b])_{L_p[a, b]} \leq c \left(\frac{1}{M} \right)^r \left\| \frac{d^r f}{dx^r} \right\|_{L_p[a, b]}$, com c uma constante real positiva dependente de p e $[a, b]$.

Teorema 2.8: Com $UL'_p[a, b]$ dado por

$$UL'_p[a, b] = \left\{ f: f \in L'_p[a, b] \text{ e } \left\| \frac{d^r f}{dx^r} \right\|_{L_p[a, b]} \leq 1 \right\},$$

para o problema de aproximação de $f \in UL'_p[a, b]$ com norma $L_p[a, b]$, nenhum outro espaço funcional de dimensão finita pode apresentar taxas de convergência de ordem superior à obtida para o espaço dos polinômios.

Prova: Veja SCHUMAKER (1981). $\square$

Este resultado permite concluir que a sequência de espaços polinomiais $P_1, P_2, \dots$ é uma sequência de espaços de aproximação assintoticamente ótima para $f \in UL'_p[a, b]$ com norma $L_p[a, b]$, produzindo

$$\text{dist}(UL'_p[a, b], P_M[a, b])_{L_p[a, b]} \leq c \left(\frac{1}{M} \right)^r.$$

2.5.4 Aproximação unidimensional utilizando splines polinomiais

Como discutido anteriormente, a principal desvantagem do uso de polinômios em processos de aproximação unidimensional em intervalos fechados $[a,b]$ é a produção de funções de aproximação que oscilam excessivamente, principalmente com o aumento da ordem dos polinômios e do intervalo de aproximação considerado.

Portanto, para aumentar a flexibilidade de processos de aproximação utilizando polinômios, é interessante trabalhar com polinômios de ordem reduzida e dividir o intervalo de aproximação em subintervalos menores. Apesar de permitir um ganho natural de flexibilidade, a aproximação polinomial por partes não mais assegura a suavidade, nem mesmo a continuidade da aproximação. Esta foi a razão da participação inexpressiva destes métodos no conjunto de procedimentos da teoria de aproximação e análise numérica anterior aos anos 60 (SCHUMAKER, 1981).

Definição 2.3: Considere $a = x_0 < x_1 < \dots < x_k < x_{k+1} = b$ e seja $\Delta = \{x_i\}_0^{k+1}$. O conjunto Δ divide o intervalo $[a,b]$ em $k+1$ subintervalos na forma:

$$\begin{cases} I_i = [x_i, x_{i+1}) & \text{para } i = 0, 1, \dots, k-1 \\ I_k = [x_k, x_{k+1}] \end{cases}.$$

Dado um inteiro positivo M , o espaço

$$PP_M(\Delta) = \left\{ f: \begin{array}{l} \text{existem polinômios } p_0, p_1, \dots, p_k \in P_M \\ \text{com } f(x) = p_i(x) \text{ para } x \in I_i, i = 0, 1, \dots, k \end{array} \right\}$$

é denominado espaço de funções polinomiais por partes de ordem M .

A terminologia empregada nesta definição é puramente descritiva, de forma que um elemento $f \in PP_M(\Delta)$ consiste de $k+1$ partes, para cada qual está definida um polinômio de ordem M . Como aplicação clássica de funções pertencentes a $PP_M(\Delta)$ têm-se o tradicional método de Euler para solução de equações diferenciais ordinárias.

Com base na definição 2.3, para manter a flexibilidade da aproximação e garantir um certo grau de suavidade global para f , foi definido um espaço de funções polinomiais cuja aplicação teve um avanço rápido a partir dos anos 60, principalmente devido a sua reconhecida praticidade, tornando-se um tópico definitivo em teoria de aproximação e análise numérica.

Funções pertencentes a este espaço são denominadas splines polinomiais, sendo que a atribuição do termo splines às funções de aproximação por partes é devida a SCHOENBERG (1946a, 1946b).

Este espaço de funções polinomiais é definido na forma (SCHUMAKER, 1981):

Definição 2.4: Seja Δ uma partição do intervalo $[a,b]$ e seja M um inteiro positivo. O espaço definido na forma

$$S_M(\Delta) = PP_M(\Delta) \cap C^{M-2}[a,b]$$

é denominado espaço de splines polinomiais de ordem M com nós nos pontos $x_1, \dots, x_k$. Ao conjunto $C^{M-2}[a,b]$ pertencem todas as funções definidas no intervalo $[a,b]$ com derivadas contínuas até ordem $M-2$.

Os splines polinomiais herdam boa parte das propriedades atribuídas aos polinômios, produzindo os seguintes resultados gerais (SCHUMAKER, 1981):

- toda função contínua no intervalo $[a,b]$ pode ser aproximada arbitrariamente bem por splines polinomiais de ordem M (M fixo), desde que um número suficiente de nós sejam fornecidos;
- taxas de convergência podem ser fornecidas precisamente para aproximação simultânea de funções suaves e suas derivadas utilizando splines polinomiais.

Para um estudo abrangente de métodos de implementação e aplicação de splines polinomiais, veja DE BOOR (1978). O procedimento de construção de splines polinomiais de ordem M envolve duas etapas básicas: a partição do espaço de variáveis de entrada em regiões disjuntas e a aproximação por um polinômio de ordem até M em cada uma destas partições. O processo de aproximação é restrito de forma que a função de aproximação global resultante tenha derivadas contínuas até ordem $M-1$ na fronteira entre as partições. Tanto a ordem dos splines polinomiais como a política de partição do espaço de variáveis de entrada devem ser determinados adequadamente para conduzir a bons resultados de aproximação, sendo que a política de partição tende a exercer uma maior influência no resultado final (FRIEDMAN, 1991a).

2.6 Aproximação multidimensional

O teorema de Weierstrass tem sido generalizado de várias formas. Uma das mais importantes generalizações corresponde à sua aplicação a funções multivariáveis. Basicamente, o que se quer provar é a possibilidade de aproximação uniforme de uma função contínua multivariável em uma região compacta. Para tanto, é apresentado o teorema de Stone-Weierstrass e a versão multivariável do teorema de Weierstrass é, então, obtida.

2.6.1 O Teorema de Stone-Weierstrass

Baseado nos resultados de Weierstrass, STONE (1948) procurou encontrar as propriedades gerais de funções de aproximação não restritas a polinômios. Para tanto, foi utilizada uma formalização matemática que passamos a descrever sucintamente na forma de algumas definições e teoremas auxiliares, que precedem à apresentação do teorema de Stone-Weierstrass.

Definição 2.5: Uma região compacta $X \subset \mathbb{R}^m$ é definida na forma:

$$X: [a_1, b_1] \times [a_2, b_2] \times \cdots \times [a_m, b_m], \text{ com } -\infty < a_i \leq x_i \leq b_i < \infty, i = 1, \dots, m,$$

onde x_i representa a i -ésima coordenada do $\mathbb{R}^m$.

Definição 2.6: Uma família $\mathbf{A}$ de funções $f: X \rightarrow \mathbb{R}$ é dita ser uma álgebra de funções se, para todo $f_1, f_2 \in \mathbf{A}$ e c uma constante real, tem-se:

$$(i) f_1 + f_2 \in \mathbf{A}; \quad (ii) f_1 f_2 \in \mathbf{A}; \quad (iii) cf_1 \in \mathbf{A};$$

ou seja, $\mathbf{A}$ é fechada para adição, multiplicação e multiplicação por escalar.

Definição 2.7: Seja $\mathbf{A}$ uma família de funções $f: X \rightarrow \mathbb{R}$. Então $\mathbf{A}$ separa pontos em X se, para cada par de pontos distintos $\mathbf{x}_1, \mathbf{x}_2 \in X$, corresponde uma função $f \in \mathbf{A}$ tal que $f(\mathbf{x}_1) \neq f(\mathbf{x}_2)$.

Definição 2.8: Se para todo $\mathbf{x} \in X$ corresponde uma função $f \in \mathbf{A}$ tal que $f(\mathbf{x}) \neq 0$, é dito que $\mathbf{A}$ não se anula em nenhum ponto de X .

Definição 2.9: Seja $\mathbf{B}$ o conjunto de todas as funções obtidas como limite das seqüências de membros de $\mathbf{A}$ que convergem uniformemente. Então $\mathbf{B}$ é dito ser o fecho uniforme de $\mathbf{A}$, com $\mathbf{A} \subseteq \mathbf{B}$.

Exemplo: O conjunto de todos os polinômios em $[a,b]$ é uma álgebra, e o teorema de Weierstrass (teorema 2.6) permite afirmar que o conjunto de funções contínuas em $[a,b]$ é o fecho uniforme do conjunto dos polinômios em $[a,b]$.

Teorema 2.9 (Teorema de Stone-Weierstrass): Seja $\mathbf{A}$ uma álgebra de funções $f: X \rightarrow \mathbb{R}$. Se $\mathbf{A}$ separa pontos em X e $\mathbf{A}$ não se anula em nenhum ponto de X , então o fecho uniforme de $\mathbf{A}$ consiste de todas as funções reais contínuas em X , ou seja $\mathbf{B} \equiv C[X]$.

Prova: Veja RUDIN (1976) e STONE (1948). $\square$

Este teorema fornece critérios que uma dada função de aproximação deve satisfazer para ser capaz de aproximar uniformemente funções contínuas arbitrárias definidas em regiões compactas.

Lema 2.1: Seja $\mathbf{A}$ uma álgebra de funções $f: X \rightarrow \mathbb{R}$ tal que $\mathbf{A}$ separa pontos em X e $\mathbf{A}$ não se anula em nenhum ponto de X . O fecho uniforme de $\mathbf{A}$, denominado $\mathbf{B}$, apresenta as seguintes propriedades:

- se $f \in \mathbf{A}$, então $|f| \in \mathbf{B}$;
- se $f_1, f_2 \in \mathbf{A}$, então $\max\{f_1, f_2\} \in \mathbf{B}$ e $\min\{f_1, f_2\} \in \mathbf{B}$, onde

$$\max\{f_1, f_2\} = \frac{1}{2}[(f_1 + f_2) + |f_1 - f_2|]$$

$$\min\{f_1, f_2\} = \frac{1}{2}[(f_1 + f_2) - |f_1 - f_2|]$$

Além disso, para qualquer função $g \in C[X]$, quaisquer dois pontos $\mathbf{x}_1, \mathbf{x}_2 \in X$ e qualquer $\varepsilon > 0$, existe $f \in \mathbf{B}$ tal que

$$|f(\mathbf{x}_i) - g(\mathbf{x}_i)| \leq \varepsilon, \quad i = 1, 2.$$

Prova: Veja DAVIS (1975). $\square$

Teorema 2.10: Seja $\mathbf{A}$ uma álgebra de funções $f: X \rightarrow \mathbb{R}$ tal que $\mathbf{A}$ separa pontos em X e $\mathbf{A}$ não se anula em nenhum ponto de X . Para que uma função $g \in C[X]$ seja uniformemente aproximável em X por elementos de $\mathbf{A}$ é necessário e suficiente que, dados quaisquer dois pontos $\mathbf{x}_1, \mathbf{x}_2 \in X$ e qualquer $\varepsilon > 0$, exista $f \in \mathbf{A}$ tal que

$$|f(\mathbf{x}_i) - g(\mathbf{x}_i)| \leq \varepsilon, \quad i = 1, 2.$$

Prova: Se g é uniformemente aproximável em X por elementos de $\mathfrak{A}$, então a condição do teorema é certamente atendida. Por outro lado, assumindo que a condição do teorema é atendida, pelo lema 2.1 tem-se que g pode ser uniformemente aproximável por elementos de $\mathfrak{B}$. Sendo assim, pelo teorema 2.9 tem-se que g pode ser uniformemente aproximável por elementos de $\mathfrak{A}$. $\square$

Teorema 2.11 (Generalização do teorema de Weierstrass para o caso multidimensional): Seja $g(x_1, x_2, \dots, x_m)$ uma função contínua tal que $g: X \rightarrow \mathfrak{R}$. A função g pode ser aproximada uniformemente em X por polinômios em $x_1, x_2, \dots, x_m$.

Prova: Por definição, o conjunto de polinômios em $x_1, x_2, \dots, x_m$ é uma álgebra de funções, denominada $\mathfrak{A}$. Sejam $\mathbf{x}_i = [x_1^{(i)} \ x_2^{(i)} \ \dots \ x_m^{(i)}]^T$, $i = 1, 2$, pontos distintos em X e considere o polinômio

$$p(x_1, x_2, \dots, x_m) = g(\mathbf{x}_1) + \frac{g(\mathbf{x}_2) - g(\mathbf{x}_1)}{\sum_{i=1}^m (x_i^{(2)} - x_i^{(1)})^2} \sum_{i=1}^m (x_i - x_i^{(1)})(x_i^{(2)} - x_i^{(1)}).$$

Como $p(\mathbf{x}_i) = g(\mathbf{x}_i)$, $i = 1, 2$, as condições do teorema 2.10 são atendidas, concluindo a demonstração. $\square$

2.7 Aproximação linear em espaços normados

Considere X um espaço funcional normado, com norma $\|\cdot\|_X$. Sejam $f_0, f_1, \dots, f_n$ elementos de X linearmente independentes.

Definição 2.10: Dados $f, g \in X$, a medida de proximidade entre f e g é dada por:

$$\|g - f\|_X.$$

Considere o problema de aproximação de uma função qualquer $g \in X$ por uma combinação linear de $f_0, f_1, \dots, f_n$. Neste caso, o objetivo é resolver o problema de minimização

$$\min_{d_0, \dots, d_n} \|g - (d_0 f_0 + d_1 f_1 + \dots + d_n f_n)\|_X.$$

2.7.1 O conceito de melhor aproximação linear

Definição 2.11: A melhor aproximação de uma função qualquer $g \in X$ por uma combinação linear de $f_0, f_1, \dots, f_n$ é um elemento $c_0 f_0 + c_1 f_1 + \dots + c_n f_n$ para o qual

$$\|g - (c_0 f_0 + c_1 f_1 + \dots + c_n f_n)\|_X \leq \|g - (d_0 f_0 + d_1 f_1 + \dots + d_n f_n)\|_X$$

para toda escolha de constantes $d_0, d_1, \dots, d_n$.

Logo, a melhor aproximação resolve o problema de minimização da norma $\|g - (d_0 f_0 + d_1 f_1 + \dots + d_n f_n)\|_X$ com $d_0 = c_0, d_1 = c_1, \dots, d_n = c_n$.

Teorema 2.12: Dado $g \in X$ e sendo $f_0, f_1, \dots, f_n$ elementos de X linearmente independentes, o problema de minimização

$$\min_{d_0, \dots, d_n} \|g - (d_0 f_0 + d_1 f_1 + \dots + d_n f_n)\|_X$$

tem uma solução.

Prova: Veja DAVIS (1975). $\square$

Corolário 2.2: Sejam $g \in C[a, b]$ e n um inteiro fixo. Tomando a norma como sendo

$$\|\cdot\|_{C[a, b]} = \max_{x \in [a, b]} |\cdot|,$$

e adotando $f_0 = 1, f_1 = x, \dots, f_n = x^n$, o problema de aproximação polinomial

$$\min_{d_0, \dots, d_n} \max_{x \in [a, b]} |g(x) - (d_0 + d_1 x + \dots + d_n x^n)|$$

tem uma solução.

Corolário 2.3: Sejam $g \in C[a, b]$ e $x_0, x_1, \dots, x_k$ pontos distintos em $[a, b]$. Assumindo que $k \geq n$, o problema de aproximação polinomial

$$\min_{d_0, \dots, d_n} \max_{0 \leq i \leq k} |g(x_i) - (d_0 + d_1 x_i + \dots + d_n x_i^n)|$$

tem uma solução.

Corolário 2.4: Sejam $g \in C[a, b]$ e $x_0, x_1, \dots, x_k$ pontos distintos em $[a, b]$. Assumindo que $k \geq n$ e tomando $\|\cdot\|_{C[a, b]}$ como sendo a norma euclidiana, o problema de aproximação polinomial

$$\min_{d_0, \dots, d_n} \left\| g(x_i) - (d_0 + d_1 x_i + \dots + d_n x_i^n) \right\|_{C[a,b]}^2 = \min_{d_0, \dots, d_n} \sum_{i=0}^k \left[g(x_i) - (d_0 + d_1 x_i + \dots + d_n x_i^n) \right]^2$$

tem uma solução. Este é o bem conhecido problema de ajuste polinomial de dados pelo critério dos quadrados mínimos.

2.7.2 A unicidade da melhor aproximação linear

Teorema 2.13: Assumindo as condições do teorema 2.12, o conjunto dos melhores aproximantes para $g \in X$ é convexo, o que implica que este conjunto é composto por um único elemento ou por um número infinito de elementos. O conjunto é unitário se a norma $\|\cdot\|_X$ for estritamente convexa.

Prova: Veja DAVIS (1975). $\square$

Apesar da importância deste resultado, ele não se aplica ao caso de melhor aproximação *uniforme*. No entanto, com base nos resultados do corolário 2.2 e teorema 2.13 e levando-se em conta que a norma $\|\cdot\|_{C[a,b]} = \max_{x \in [a,b]} |\cdot|$ é estritamente convexa, conclui-se que o problema

$$\min_{p \in P_n} \max_{x \in [a,b]} |g(x) - p(x)|$$

tem uma única solução $p_n(x)$. Esta solução é denominada a aproximação de Chebishev de grau menor ou igual a n .

Teorema 2.14: Seja $E_n(g) = \max_{x \in [a,b]} |g(x) - p_n(x)|$. Se $g \in C[a,b]$, então $E_0(g) \geq E_1(g) \geq \dots$ e

$$\lim_{n \rightarrow \infty} E_n(g) = 0.$$

Prova: Veja DAVIS (1975). $\square$

2.8 Aproximação linear em espaços com produto interno

O processo de aproximação linear mais empregado é o método dos quadrados mínimos, cuja formalização se dá em espaços com produto interno, definidos no Apêndice B. Em outras palavras, a solução do problema de quadrados mínimos representa a melhor aproximação em espaços dotados de produto interno.

Teorema 2.15: Se X é um espaço com produto interno, a equação

$$\|f\|_X = \sqrt{\langle f, f \rangle}, \text{ com } f \in X,$$

define uma norma em X . Neste caso, X torna-se um espaço linear normado.

Prova: Veja DAVIS (1975). $\square$

2.8.1 A operação de projeção realizada pelo produto interno

O produto interno pode ser interpretado geometricamente como sendo uma operação de projeção em espaços multidimensionais, conforme sugerido na figura 2.3, representando o caso bidimensional.

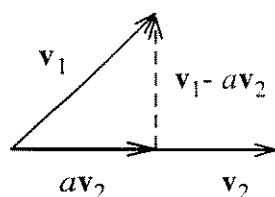


Figura 2.3 - A projeção realizada pelo produto interno no $\mathbb{R}^2$

Sejam $\mathbf{v}_1, \mathbf{v}_2 \in \mathbb{R}^2$ elementos não-nulos. Considere um escalar a tal que $a\mathbf{v}_2$ corresponda a projeção de $\mathbf{v}_1$ na direção de $\mathbf{v}_2$. Então, pode-se afirmar que

$$a\mathbf{v}_2 \perp \mathbf{v}_1 - a\mathbf{v}_2,$$

conduzindo a

$$\langle a\mathbf{v}_2, \mathbf{v}_1 - a\mathbf{v}_2 \rangle = 0.$$

Logo,

$$a\langle \mathbf{v}_2, \mathbf{v}_1 \rangle - a^2\langle \mathbf{v}_2, \mathbf{v}_2 \rangle = 0,$$

permitindo obter a na forma

$$a = \frac{\langle \mathbf{v}_2, \mathbf{v}_1 \rangle}{\langle \mathbf{v}_2, \mathbf{v}_2 \rangle}.$$

Isto significa que a projeção de $\mathbf{v}_1$ na direção de $\mathbf{v}_2$ ($\mathbf{v}_2 \neq \mathbf{0}$) assume a forma:

$$\text{proj}_{\mathbf{v}_2}(\mathbf{v}_1) = \frac{\langle \mathbf{v}_2, \mathbf{v}_1 \rangle}{\langle \mathbf{v}_2, \mathbf{v}_2 \rangle} \mathbf{v}_2. \quad (2.1)$$

2.8.2 Sistemas ortonormais

Definição 2.12: Um conjunto S de elementos de um espaço X com produto interno é denominado ortogonal se

$$\langle f_1, f_2 \rangle = \begin{cases} 0 & \text{se } f_1 \neq f_2 \\ a & \text{se } f_1 = f_2 \end{cases} \quad f_1, f_2 \in S \text{ e } a \in \mathbb{R}.$$

Definição 2.13: Um conjunto S de elementos de um espaço X com produto interno é denominado ortonormal se

$$\langle f_1, f_2 \rangle = \begin{cases} 0 & \text{se } f_1 \neq f_2 \\ 1 & \text{se } f_1 = f_2 \end{cases} \quad f_1, f_2 \in S.$$

Teorema 2.16: Qualquer conjunto finito de elementos ortonormais não-nulos $f_1, f_2, \dots, f_n$ é linearmente independente.

Prova: Veja DAVIS (1975). $\square$

O teorema 2.16 tem uma recíproca parcial muito importante. Qualquer conjunto de elementos linearmente independentes não é necessariamente ortonormal, mas pode ser ortonormalizado, por exemplo, via processo de ortogonalização de Gram-Schmidt, apresentado no apêndice A.

Lema 2.2: Dados $f_1, f_2 \in C[a, b]$, se $w(x)$ é uma função integrável em $[a, b]$, então a integral

$$\langle f_1, f_2 \rangle_w = \int_a^b w(x) f_1(x) f_2(x) dx$$

define um produto interno em $C[a, b]$. A função $w(x)$ é denominada função de ponderação, sendo geralmente assumida positiva em $[a, b]$.

Prova: Veja DAVIS (1975). $\square$

Como as potências $1, x, x^2, \dots$ são linearmente independentes em $C[a, b]$, elas podem ser ortonormalizadas com relação a este produto interno, produzindo um conjunto de polinômios

$$p_n(x) = c_n x^n + \dots, \quad n = 0, 1, 2, \dots; \quad c_n > 0$$

que são ortonormais com base na seguinte relação

$$\int_a^b w(x) p_m(x) p_n(x) dx = \begin{cases} 1 & \text{se } m = n \\ 0 & \text{se } m \neq n \end{cases} \quad m, n = 0, 1, \dots$$

Neste caso, diz-se que as funções $f_n(x) = \sqrt{w(x)} p_n(x)$ ($n = 0, 1, \dots$) são ortonormais em $[a, b]$.

2.8.3 Funções ortonormais baseadas em polinômios de Hermite

De acordo com a escolha do intervalo $[a,b]$ e da função de ponderação $w(x)$, inúmeras funções ortonormais podem ser obtidas em $C[a,b]$ com base no processo de ortonormalização das potências $1, x, x^2, \dots$. Neste estudo são consideradas apenas as funções ortonormais geradas a partir de polinômios de Hermite, que são obtidos tomando-se $a = -\infty, b = +\infty$ e $w(x) = e^{-x^2}$.

Os polinômios de Hermite são definidos recursivamente como segue:

- $p_0(x) = 1$;
- $p_1(x) = 2x$;
- $p_{i+1}(x) = 2xp_i(x) - 2ip_{i-1}(x), \quad i > 0$.

Estes polinômios são ortogonais com base no seguinte produto interno (MAGNUS *et al.*, 1966):

$$\int_{-\infty}^{\infty} e^{-x^2} p_i(x) p_j(x) dx = \begin{cases} 0 & \text{se } i \neq j \\ \sqrt{\pi} 2^i i! & \text{se } i = j \end{cases}$$

de tal forma que as funções

$$h_i(x) = \frac{e^{-\frac{x^2}{2}}}{\sqrt{\sqrt{\pi} 2^i i!}} p_i(x), \quad i = 0, 1, \dots \tag{2.2}$$

são ortonormais em $(-\infty, \infty)$ por produzirem

$$\int_{-\infty}^{\infty} h_i(x) h_j(x) dx = \begin{cases} 0 & \text{se } i \neq j \\ 1 & \text{se } i = j \end{cases} \tag{2.3}$$

2.8.4 Expansão ortonormal: série de Fourier generalizada

Seja $\{h_0, h_1, \dots\}$ a seqüência infinita das funções ortonormais definidas de acordo com as equações (2.2) e (2.3) da seção anterior. A expansão ortonormal (também denominada série de Fourier generalizada) de uma função arbitrária $g: \Re \rightarrow \Re$ é dada na forma:

$$g(x) = \sum_{i=0}^{\infty} c_i h_i(x), \tag{2.4}$$

onde $c_i \ (i=0,1,\dots)$ são os coeficientes de Fourier generalizados. Definindo o produto escalar como sendo

$$\langle f_1, f_2 \rangle = \int_{-\infty}^{\infty} f_1(x) f_2(x) dx ,$$

estes coeficientes podem ser obtidos como segue:

- Para $j=0,1,\dots$, multiplique a equação (2.4) por $h_j(x)$ e integre em todo o intervalo produzindo:

$$\int_{-\infty}^{\infty} g(x) h_j(x) dx = \sum_{i=0}^{\infty} c_i \int_{-\infty}^{\infty} h_i(x) h_j(x) dx ,$$

- Utilizando o resultado da equação (2.3), resulta:

$$c_j = \int_{-\infty}^{\infty} g(x) h_j(x) dx = \langle g, h_j \rangle .$$

Portanto, a equação (2.4) pode então ser reescrita na forma:

$$g(x) = \sum_{i=0}^{\infty} \langle g, h_i \rangle h_i(x) \quad (2.5)$$

Uma outra forma de interpretar esta expansão ortonormal é relacionar a equação (2.5) com a equação (2.1), produzindo:

$$g(x) = \sum_{i=0}^{\infty} \text{proj}_{h_i}(g) ,$$

ou seja, a expansão ortonormal de uma função g , é simplesmente a somatória da projeção desta função em cada elemento do sistema ortonormal.

Exemplo: Tomando $g \in C[-\pi, \pi]$ ou $g \in L_2[-\pi, \pi]$, o sistema ortonormal $\frac{1}{\sqrt{2\pi}}, \frac{\text{sen}(x)}{\sqrt{\pi}}, \frac{\cos(x)}{\sqrt{\pi}}, \frac{\text{sen}(2x)}{\sqrt{\pi}}, \dots$ e o produto escalar $\langle f_1, f_2 \rangle = \int_{-\pi}^{\pi} f_1(x) f_2(x) dx$, resulta a série de Fourier:

- $g(x) = \frac{a_0}{2} + \sum_{k=1}^{\infty} a_k \cos(kx) + b_k \text{sen}(kx) ;$
- $a_k = \frac{1}{\pi} \int_{-\pi}^{\pi} g(x) \cos(kx) dx$ e $b_k = \frac{1}{\pi} \int_{-\pi}^{\pi} g(x) \text{sen}(kx) dx .$

2.8.5 Expansão ortonormal truncada: a aproximação de quadrados mínimos

A expansão ortonormal truncada tem a seguinte propriedade:

Teorema 2.17: Dado um conjunto de funções ortonormais h_i ($i=0,1,\dots$) e uma função arbitrária $g: \mathfrak{R} \rightarrow \mathfrak{R}$, então

$$\left\| g - \sum_{i=0}^P \langle g, h_i \rangle h_i \right\| \leq \left\| g - \sum_{i=0}^P c_i h_i \right\|$$

para qualquer escolha de constantes $c_0, c_1, \dots, c_P$ (P finito).

Prova: Veja DAVIS (1975). $\square$

Particularmente, todo problema de quadrados mínimos pode ser resolvido por uma expansão ortonormal truncada, isto é, a expansão ortonormal truncada é a melhor aproximação em um espaço com produto interno.

Dado um conjunto $\{(x_l, s_l)\}_{l=1}^N$ de dados amostrados da função $g: \mathfrak{R} \rightarrow \mathfrak{R}$ tal que $s_l = g(x_l)$, o problema da melhor aproximação de g por uma série ortonormal truncada pode ser colocado na forma de um problema de quadrados mínimos como segue:

$$\min_{c_0, \dots, c_P} \sum_{l=1}^N \left(s_l - \sum_{i=0}^P c_i h_i(x_l) \right)^2. \quad (2.6)$$

Tomando a norma euclidiana $\|\cdot\|_2$, e construindo a matriz $\mathbf{H}$ e os vetores $\mathbf{c}$ e $\mathbf{s}$ na forma:

$$\mathbf{H} = \begin{bmatrix} h_0(x_1) & h_1(x_1) & \cdots & h_P(x_1) \\ h_0(x_2) & h_1(x_2) & \cdots & h_P(x_2) \\ \vdots & \vdots & \ddots & \vdots \\ h_0(x_N) & h_1(x_N) & \cdots & h_P(x_N) \end{bmatrix}, \quad \mathbf{c} = \begin{bmatrix} c_0 \\ c_1 \\ \vdots \\ c_P \end{bmatrix} \quad \text{e} \quad \mathbf{s} = \begin{bmatrix} s_1 \\ s_2 \\ \vdots \\ s_N \end{bmatrix},$$

a equação (2.6) pode ser reescrita como segue:

$$\min_{\mathbf{c}} \|\mathbf{s} - \mathbf{H}\mathbf{c}\|_2^2 = \min_{\mathbf{c}} (\mathbf{s} - \mathbf{H}\mathbf{c})^T (\mathbf{s} - \mathbf{H}\mathbf{c}) = \mathbf{s}^T \mathbf{s} + \min_{\mathbf{c}} (-\mathbf{s}^T \mathbf{H}\mathbf{c} - \mathbf{c}^T \mathbf{H}^T \mathbf{s} + \mathbf{c}^T \mathbf{H}^T \mathbf{H} \mathbf{c}).$$

Fazendo $J(\mathbf{c}) = -\mathbf{s}^T \mathbf{H}\mathbf{c} - \mathbf{c}^T \mathbf{H}^T \mathbf{s} + \mathbf{c}^T \mathbf{H}^T \mathbf{H} \mathbf{c}$, a seguinte condição necessária deve ser atendida no ponto de mínimo:

$$\frac{\partial J(\mathbf{c})}{\partial \mathbf{c}} = 0 \Rightarrow -2\mathbf{H}^T \mathbf{s} + 2\mathbf{H}^T \mathbf{H} \mathbf{c} = 0 \Rightarrow \mathbf{H}^T \mathbf{H} \mathbf{c} = \mathbf{H}^T \mathbf{s}.$$

Assumindo $P < N$ e tal que $\mathbf{H}$ tenha posto completo, a solução ótima, no sentido dos quadrados mínimos, é denominada $\mathbf{c}^*$ e pode ser expressa na forma:

$$\mathbf{c}^* = (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \mathbf{s}. \quad (2.7)$$

2.9 Aproximação não-linear: Teorema de Kolmogorov

Em 1900, Hilbert (HILBERT, 1935) propôs uma lista de problemas matemáticos até então sem solução e classificou-os como os maiores desafios para os matemáticos do século 20. O 13º problema de Hilbert questionava a existência de funções multivariáveis genuínas. Por exemplo, para $x, y \in \mathfrak{R}$, tanto $f_1(x, y) = x + y$ como $f_2(x, y) = xy = e^{\log x + \log y}$ são superposições de funções monovariáveis e a operação de adição.

Seja X_m o hipercubo unitário $[0, 1]^m$ do espaço $\mathfrak{R}^m$ e considere $C[X_m]$ o conjunto de todas as funções escalares reais e contínuas de m variáveis definidas em X_m . Seja $C[\mathfrak{R}]$ o conjunto das funções escalares reais e contínuas definidas em $\mathfrak{R}$. Como consequência de esforços visando resolver o 13º problema de Hilbert, KOLMOGOROV (1957) publicou um teorema referente à representação exata de qualquer função $g \in C[X_m]$ ($m \geq 2$) por uma superposição finita de funções $f_j \in C[\mathfrak{R}]$, utilizando apenas a operação de adição. O teorema é apenas existencial e pode ser formulado como segue:

Teorema 2.18 (Teorema de Kolmogorov): Para toda função $g \in C[X_m]$, com $m \geq 2$, existe uma representação exata dada na forma:

$$g(x_1, x_2, \dots, x_m) = \sum_{j=1}^{2m+1} f_j \left(\sum_{i=1}^m \varphi_{ij}(x_i) \right), \quad (2.8)$$

onde $f_j(\cdot) \in C[\mathfrak{R}]$, $\varphi_{ij}(\cdot) \in C[\mathfrak{R}]$ e $x_i \in [0, 1]$ ($i=1, \dots, m$; $j=1, \dots, 2m+1$). Além disso, as funções $\varphi_{ij}(\cdot)$ são monotônicas e independentes de g .

Prova: Veja KOLMOGOROV (1957). $\square$

A este teorema seguiram-se os resultados de SPRECHER (1965) e LORENTZ (1966), que conduzem a uma nova representação para $g \in C[X_m]$ ($m \geq 2$) na forma:

Teorema 2.19 (Generalização do Teorema de Kolmogorov): Para toda função $g \in C[X_m]$, com $m \geq 2$, existe uma representação exata dada na forma:

$$g(x_1, x_2, \dots, x_m) = \sum_{j=1}^{2m+1} f \left(\sum_{i=1}^m \lambda_i \varphi_j(x_i) \right), \quad (2.9)$$

onde $f(\cdot) \in C[\mathfrak{R}]$, $\varphi_j(\cdot) \in C[\mathfrak{R}]$, $\lambda_i \in \mathfrak{R}$ e $x_i \in [0, 1]$ ($i=1, \dots, m$). As constantes λ_i são independentes de g e as funções $\varphi_j(\cdot)$ são monotônicas e independentes de g .

Prova: Veja SPRECHER (1965) e LORENTZ (1966). $\square$

É possível estender este resultado para o caso de funções contínuas definidas em outras regiões compactas que não um hipercubo ou em hipercubos $[a_1, b_1] \times [a_2, b_2] \times \dots \times [a_n, b_n]$ através de transformações lineares

$$y_i = \frac{x_i - a_i}{b_i - a_i}, \quad i = 1, \dots, m.$$

KAHANE (1975) mostrou que a representação exata é possível para quase todo conjunto de constantes λ_i racionalmente independentes e quase todo conjunto de funções monotônicas contínuas $\varphi_j(\cdot)$. As funções $\varphi_j(\cdot)$ podem ser tomadas como funções que satisfazem uma condição de Lipschitz de ordem 1 (veja Apêndice B). Mas, apesar de monotônicas e contínuas, devem ser essencialmente não-suaves, como mostra o seguinte lema:

Lema 2.3: Existem funções suaves de m variáveis ($m \geq 2$) que não são representáveis pela superposição linear de funções suaves de um número menor de variáveis. Particularmente, existem funções analíticas que não são representáveis pela superposição linear de funções suaves de uma variável.

Prova: Veja VITUSHKIN (1977). $\square$

Na mesma direção, segue um resultado mais específico:

Lema 2.4: Nem todo polinômio $g: X_m \rightarrow \mathfrak{R}$ pode ser representado exatamente na forma:

$$g(x_1, x_2, \dots, x_m) = \sum_{j=1}^n f_j \left(\sum_{i=1}^m \lambda_{ij} x_i \right),$$

onde $f_j(\cdot) \in C[\mathfrak{R}]$, $\lambda_{ij} \in \mathfrak{R}$ e $x_i \in \mathfrak{R}$ ($i=1, \dots, m; j=1, \dots, n$) e n é um número natural arbitrário e finito.

Prova: Veja SPRECHER (1966). $\square$

Mesmo quando a representação exata é possível, ou seja, com funções não-suaves $\varphi_j(\cdot) \in C[\mathfrak{R}]$ nas equações (2.8) e (2.9), os resultados não são construtivos, pois não auxiliam na determinação das funções e parâmetros que conduzem à aproximação exata, principalmente porque as funções f_j ($j=1, \dots, 2m+1$) da equação (2.8) e f da equação (2.9) dependem de g .

No capítulo 4 são abordados procedimentos alternativos que não se baseiam em aproximação exata, mas em aproximação com convergência assintótica, onde ainda assim os resultados do teorema de Kolmogorov e suas generalizações são fundamentais, por justificarem o emprego de modelos formados por composição aditiva de funções de uma variável na aproximação de mapeamentos suaves.

Capítulo 3

Conceitos adicionais da teoria de aproximação de funções

Como um ramo da análise funcional, a teoria de aproximação de funções tornou-se uma ferramenta teórica fundamental na implementação de métodos de aproximação capazes de aplicarem eficientemente os poderosos recursos computacionais atualmente disponíveis. Mas independente de motivações de ordem computacional, a teoria de aproximação de funções tem sido desenvolvida há vários séculos. A quantidade de informação disponível é extraordinária. Tomando como exemplo apenas tópicos envolvendo funções polinomiais ortogonais, foi preparada ainda no final dos anos 30 uma bibliografia contendo várias centenas de páginas com referências, sendo que este tópico continuou a evoluir substancialmente desde então.

Dentre as várias categorias de problemas de aproximação, destaca-se a aproximação de funções multivariáveis desconhecidas com base em um conjunto finito de dados amostrados, sendo que o processo de amostragem está geralmente sujeito a algum tipo de ruído.

Dependendo do contexto, este problema de matemática aplicada recebe denominações diferentes:

- em estatística: regressão múltipla
- em engenharia: aprendizado supervisionado em redes neurais artificiais.

3.1 Regressão, aproximação e interpolação

Em estatística, se $(\mathbf{x}, s)$ é um par de vetores de variáveis aleatórias ($\mathbf{x} \in \mathfrak{R}^m$ e $s \in \mathfrak{R}$), diz-se que existe uma distribuição conjunta $P^{\mathbf{x},s}$ de $(\mathbf{x}, s)$ à qual corresponde uma distribuição condicional $P^s(\cdot/\mathbf{x})$ de s dado $\mathbf{x}$. Caso a distribuição conjunta $P^{\mathbf{x},s}$ de $(\mathbf{x}, s)$ seja uma distribuição gaussiana multivariável $G(\mu, \Sigma)$, com média conhecida μ e matriz de covariância Σ , existem fórmulas simples para expressar as associações lineares entre as variáveis vinculadas à distribuição condicional $P^s(\cdot/\mathbf{x})$. No entanto, a identificação de associações não-lineares e sua representação em espaços multidimensionais envolvem a aplicação de modelos de regressão

não-linear (paramétricos ou não-paramétricos) ajustados com base em dados amostrados da distribuição condicional $P^s(\cdot/\mathbf{x})$.

No desenvolvimento que se segue, as variáveis são consideradas determinísticas e, portanto, não é utilizado o conceito de distribuição condicional, mas o de mapeamento representado por uma função vetorial e de argumento vetorial. Neste caso, é mais adequado referir-se a métodos de aproximação em lugar de métodos de regressão. No entanto, é importante observar que os métodos de aproximação a serem apresentados neste estudo foram originalmente projetados para problemas de regressão multivariável, sendo que o tratamento de variáveis determinísticas representa um caso particular.

Apesar de as variáveis serem determinísticas, é assumido que elas estão sujeitas a ruído. Como interpolação pode ser considerada como o limite da aproximação quando os dados não estão sujeitos a ruído¹, a possibilidade de ruído no processo de amostragem representa mais uma razão, além daquelas já discutidas no capítulo 2, para descartar estratégias de aproximação baseadas em interpolação.

3.2 O problema geral de aproximação

Seja X uma região compacta do $\mathcal{R}^m$ e seja $g: X \subset \mathcal{R}^m \rightarrow \mathcal{R}^r$ uma função desconhecida que associa a cada vetor de variáveis de entrada $\mathbf{x} \in X$ um vetor de variáveis de saída $\mathbf{s} \in \mathcal{R}^r$. Se g é uma função limitada, então a imagem de X sob g é também uma região compacta $g[X] \subset \mathcal{R}^r$ definida na forma:

$$g[X] = \{g(\mathbf{x}): \mathbf{x} \in X\}.$$

Elementos de um conjunto multidimensional de pares de vetores de entrada-saída $(\mathbf{x}_l, \mathbf{s}_l)$, $l=1, \dots, N$, são gerados considerando-se que os vetores de entrada $\mathbf{x}_l$ estão distribuídos na região compacta $X \subset \mathcal{R}^m$ de acordo com uma função de densidade de probabilidade fixa $d\rho: X \subset \mathcal{R}^m \rightarrow [0,1]$ e que os vetores de saída $\mathbf{s}_l$ assumem valores finitos produzidos pelo mapeamento definido pela função g na forma:

$$\mathbf{s}_l = g(\mathbf{x}_l) + \varepsilon_l, \quad l = 1, \dots, N.$$

¹ Atendendo a níveis de aproximação arbitrariamente especificados, soluções para problemas de aproximação podem ser obtidas sem que se recorra à interpolação de um único ponto pertencente ao conjunto de dados. No entanto, para dados imunes a ruído, a aproximação exata (limite do processo de aproximação) implica a interpolação de todos os pontos pertencentes ao conjunto de dados.

onde $\varepsilon_l \in \mathfrak{R}^r$ é um vetor de erro aditivo, composto por variáveis aleatórias de média zero e variância fixa, geralmente refletindo a dependência de s_l , $l=1, \dots, N$, com relação a quantidades não-controláveis (por exemplo, ruído no processo de amostragem) e não-observáveis (por exemplo, variáveis não-mensuráveis).

Dado o conjunto de dados amostrados $\{(\mathbf{x}_l, \mathbf{s}_l)\}_{l=1}^N$ definidos na região compacta $X \times g[X]$ do $\mathfrak{R}^m \times \mathfrak{R}^r$, o problema de aproximação envolve a determinação de um modelo de aproximação $\hat{g}$ que corresponda à melhor aproximação para o mapeamento realizado pela função g em $X \times g[X]$. Dentre os motivos que justificam a obtenção de um modelo de aproximação $\hat{g}$ para uma função desconhecida g , é possível destacar (SILVERMAN, 1985):

- $\hat{g}$ constitui uma forma de apresentação (para futuras investigações) das associações existentes entre as variáveis de entrada que contribuem para produzir as variáveis de saída;
- com $\hat{g}$ é possível obter estimativas de dados ainda não-observados;
- $\hat{g}$ pode explicitar propriedades importantes presentes na função g .

A abordagem inicial do problema de aproximação é geralmente de natureza exploratória, procurando identificar alguma informação preliminar acerca do fenômeno ou processo que produziu os dados, sem a imposição de nenhum tipo de restrição. A análise multivariável clássica fornece ferramentas poderosas na obtenção de associações lineares (estrutura de correlação) entre as variáveis. Se todas as informações relevantes puderem ser extraídas com base nestas ferramentas clássicas, nenhum passo adicional se faz necessário.

Já nos casos em que associações não-lineares arbitrárias estão presentes, a determinação do tipo de não-linearidade é fundamental para a obtenção do melhor modelo de aproximação a partir dos dados de entrada-saída. Quando a forma da não-linearidade é conhecida e passível de descrição matemática, modelos paramétricos são normalmente empregados, com os parâmetros podendo ser determinados com base em técnicas de regressão não-linear tradicionais (RATKOWSKY, 1989).

Apesar de simplificar o problema de aproximação, a imposição de restrições paramétricas pode dificultar ou até impossibilitar a representação adequada da estrutura presente nos dados de entrada-saída, principalmente quando a função g apresenta não-linearidades que não se restringem a nenhuma família de funções paramétricas. Mesmo no caso de problemas de aproximação passíveis de tratamento paramétrico, é recomendada uma abordagem inicial utilizando modelos não-paramétricos para auxiliar na determinação do tipo

de parametrização que pode ser utilizado. STONE (1977) faz um estudo mais aprofundado desta e outras motivações para o emprego de modelos não-paramétricos.

No entanto, o uso de modelos não-paramétricos provoca uma significativa acentuação de uma característica já presente em alguns problemas de aproximação que utilizam abordagens paramétricas (STONE, 1982): a maldição da dimensionalidade. Esta expressão foi empregada originalmente por BELLMAN (1961) e se refere basicamente à existência de uma relação direta entre a dimensionalidade dos dados e a quantidade de dados necessária para possibilitar o sucesso da tarefa de aproximação. A consequência prática da maldição da dimensionalidade é a redução do poder de aproximação com a redução da densidade de dados amostrados. Portanto, a manutenção do poder de aproximação com um aumento da dimensão do espaço de aproximação está associada a um aumento exponencial no número de dados (TIKHONOV & ARSENIN, 1977).

Apesar de estar invariavelmente associada a uma redução da capacidade de aproximação, a imposição de um conjunto de restrições (inclusive, restrições paramétricas) aos modelos não-paramétricos pode conduzir à independência da dimensionalidade (CHEN, 1991). Daí, conclui-se que os melhores modelos de aproximação são aqueles capazes de conciliar o nível de dependência da dimensionalidade com a flexibilidade do modelo de aproximação.

Antes de um estudo mais detalhado de modelos de aproximação que procuram atender a este compromisso, impõe-se uma discussão sobre o comportamento de problemas de aproximação baseados em dados de entrada-saída.

3.3 A suavidade da função a ser aproximada

Certos tipos de problemas de valores de contorno envolvendo equações diferenciais parciais parecem ter sido os primeiros a motivar uma classificação utilizando os termos *mal-comportado* e *bem-comportado*. Um problema é bem-comportado quando sua solução é única e depende continuamente dos dados disponíveis. Neste sentido, problemas de aproximação baseados em pares de vetores de entrada-saída gerados a partir do mapeamento definido pela função a ser aproximada são problemas mal-comportados (TIKHONOV & ARSENIN, 1977), pois a informação contida nos dados de entrada-saída não são suficientes para reconstruir unicamente o mapeamento em regiões onde os dados não estão disponíveis e, em princípio, não há como assegurar que a solução ótima dependa continuamente dos dados disponíveis.

Portanto, algum tipo de restrição deve ser imposta à função a ser aproximada de forma a produzir um problema de aproximação bem-comportado. Uma das restrições mais gerais e fracas (é atendida pela maior parte dos mapeamentos, especialmente aqueles originados de fenômenos físicos) é assumir que o mapeamento é suave. Como será visto mais adiante, esta restrição já faz com que a aproximação seja possível, pois impõe automaticamente aos dados um nível de redundância, ou seja, pequenas perturbações nas variáveis de entrada determinam pequenas perturbações nas variáveis de saída. Deste modo, a suavidade de uma função reflete o efeito da não-localidade, isto é, o valor da função em um ponto depende do valor da função em uma região que envolve este ponto.

O poder do processo de aproximação é tanto maior quanto maior for o grau de suavidade da função que gera o mapeamento. Como exemplo, as funções a seguir são apresentadas em ordem crescente de suavidade (veja definições no apêndice B): funções mensuráveis em norma de Lebesgue, funções limitadas, funções contínuas, funções que satisfazem uma condição de Lipschitz, funções diferenciáveis, funções analíticas, funções polinomiais de ordem finita, funções constantes.

Em face da impossibilidade de se estimar o grau de suavidade de uma função multidimensional a ser aproximada, STONE (1982) sugere que a *única* opção é impor um grau apropriado de suavidade. Este procedimento contribui para amenizar o efeito da maldição da dimensionalidade, ou seja, permite manter em níveis aceitáveis o número de dados de entrada-saída necessário para realizar a aproximação. No entanto, restrições de suavidade muito acentuadas, como assumir que o mapeamento é linear ou invariante sobre algum tipo de transformação, apesar de simplificarem decisivamente o problema de aproximação, podem acarretar uma redução drástica na capacidade de aproximação quando aplicadas a mapeamentos que não atendem a estas restrições.

Este compromisso entre suavidade do processo de aproximação e capacidade de aproximação é completamente análogo ao compromisso existente entre dependência da dimensionalidade e flexibilidade do modelo de aproximação, discutido na seção anterior. Pode-se então antecipar que o tipo de suavidade imposta à função a ser aproximada vai determinar o tipo de modelo de aproximação a ser utilizado.

Devido à necessidade de processamento em tempo finito, os modelos de aproximação a serem implementados devem apresentar um número limitado de componentes e operadores. No entanto, para efeito de consistência do processo de aproximação, é desejável que estes modelos de aproximação possuam capacidade de aproximação universal. A capacidade de

representar exatamente qualquer tipo de mapeamento é certamente uma propriedade assintótica, pois geralmente só é atendida para modelos de aproximação com um número ilimitado de componentes e operadores.

Dentre todos os modelos que apresentam capacidade de aproximação universal, o melhor modelo de aproximação pode finalmente ser caracterizado como aquele que apresenta restrições de suavidade compatíveis com o mapeamento a ser aproximado, de forma a atingir um nível de aproximação adequado com base no menor número de componentes e operadores.

O nível de suavidade do mapeamento a ser aproximado geralmente não é conhecido *a priori*, mas certamente deve poder ser obtido a partir dos dados amostrados. Caso contrário, a tarefa de aproximação continua mal-condicionada. Neste estudo, o nível de suavidade é determinado a partir dos dados amostrados utilizando-se validação cruzada, como será visto no próximo capítulo.

3.4 A classe de funções de aproximação

Problemas de aproximação são geralmente abordados como segue:

- determina-se uma classe de funções F a serem utilizadas no processo de aproximação;
- define-se um processo de aproximação que selecione elementos de F com base na informação disponível.

O sucesso do processo de aproximação depende da existência de funções em F compatíveis com o problema de aproximação. Assumindo que a função g a ser aproximada é suave (hipótese razoável tendo em vista o vínculo com processos físicos), a garantia de factibilidade e consistência do processo de aproximação está diretamente (mas não unicamente) associada à presença das seguintes propriedades básicas junto à classe de funções F :

- as funções em F devem ser suaves;
- como a implementação é geralmente realizada utilizando-se computadores digitais, as funções em F devem ser de fácil armazenagem e manipulação, incluindo neste contexto a avaliação da função e de suas derivadas e antiderivadas;
- a classe de funções F deve ser suficientemente abrangente de forma que funções suaves arbitrárias possam ser aproximadas utilizando-se elementos de F .

O estudo de classes de funções de aproximação F é o tema central da teoria de aproximação, enquanto que o projeto e análise de algoritmos de implementação do processo de aproximação utilizando elementos de F são procedimentos derivados da teoria de análise numérica.

3.5 O pré-processamento dos dados de entrada-saída

Uma característica importante dos métodos de aproximação a serem considerados é a invariância com relação à escolha da origem (translação) e dos eixos de coordenadas ortonormais (rotação) do espaço de entrada. Portanto, é possível padronizar os vetores de entrada $\mathbf{x}_l$ na forma:

$$\tilde{\mathbf{x}}_l = \mathbf{C}_x^{-1/2}(\mathbf{x}_l - \bar{\mathbf{x}}), \quad l=1, \dots, N, \quad (3.1)$$

onde

$$\bar{\mathbf{x}} = \frac{1}{N} \sum_{l=1}^N \mathbf{x}_l$$

é a média de $\mathbf{x}$ e

$$\mathbf{C}_x = \frac{1}{N} \sum_{l=1}^N (\mathbf{x}_l - \bar{\mathbf{x}})(\mathbf{x}_l - \bar{\mathbf{x}})^T$$

é a matriz de covariância de $\mathbf{x}$. A média e a matriz de covariância de $\tilde{\mathbf{x}}$ são dadas por:

$$\bar{\tilde{\mathbf{x}}} = \mathbf{0} \quad \text{e} \quad \mathbf{C}_{\tilde{\mathbf{x}}} = \mathbf{I}.$$

A transformação afim dada pela equação (3.1) permite fixar a análise dos dados em associações entre variáveis que não sejam apenas de natureza elipsoidal (TUKEY & TUKEY, 1981), além de fazer com que toda operação de projeção junto aos dados padronizados também produza dados padronizados.

A aplicação deste processo de padronização a dados já centrados, ou seja, com $\bar{\mathbf{x}} = \mathbf{0}$, é equivalente a utilização dos resultados da análise de componentes principais para escalonar os eixos principais de forma que assumam norma unitária. Neste processo é geralmente preferível operar com a matriz de correlação em lugar da matriz de covariância (JOLLIFFE, 1986).

No caso dos métodos de aproximação a serem adotados, a invariância com relação à transformação afim de padronização descrita acima se aplica apenas aos conceitos derivados da teoria de aproximação, já que, em termos de análise numérica, é sabido que os algoritmos de otimização a serem empregados não são necessariamente invariantes afins (HUBER, 1985).

3.6 Avaliação do nível de aproximação

Para uma avaliação do nível de aproximação é necessário introduzir o conceito de distância entre a função a ser aproximada e sua aproximação:

$$J(\hat{g}) = \text{dist}[g(\mathbf{x}), \hat{g}(\mathbf{x})],$$

que é geralmente induzida por uma norma (veja apêndice C).

Como, por hipótese, na geração dos pares de vetores de entrada-saída $\{(\mathbf{x}_l, \mathbf{s}_l) \in \mathfrak{R}^m \times \mathfrak{R}^r\}_{l=1}^N$ os vetores de entrada $\mathbf{x}_l$ estão distribuídos na região compacta $X \subset \mathfrak{R}^m$ de acordo com uma função de densidade de probabilidade $d_P: X \subset \mathfrak{R}^m \rightarrow [0,1]$, a distância entre g e $\hat{g}$ pode ser calculada na forma:

$$J(\hat{g}) = \int_X (g(\mathbf{x}) - \hat{g}(\mathbf{x}))^2 d_P(\mathbf{x}) d\mathbf{x}. \quad (3.2)$$

Esta medida de aproximação é denominada erro quadrático médio (EQM) e é bem-definida sempre que a expressão $(g(\mathbf{x}) - \hat{g}(\mathbf{x}))^2 d_P(\mathbf{x})$ for integrável em X , fazendo com que $J(\hat{g})$ assuma valores finitos (note que $J(\hat{g}) \geq 0, \forall \hat{g}$). Apesar de não ser a única medida de distância, o EQM é sem dúvida a medida mais utilizada, principalmente por duas razões (HECHT-NIELSEN, 1990):

- por ponderar uniformemente o quadrado da magnitude do erro em toda a região compacta X , esta medida assegura que erros maiores tenham uma maior participação;
- por levar em consideração a frequência de ocorrência de um determinado vetor de entrada, esta medida é mais sensível a erros provenientes de entradas mais frequentes.

No entanto, sempre que para uma determinada aplicação estas características não forem desejadas, outras medidas de distância (normas) devem ser utilizadas, como erro absoluto máximo e erro absoluto médio. No entanto, métodos de solução para o problema geral de aproximação baseados em outras medidas que não o EQM apresentam uma das seguintes particularidades (WIDROW & STEARNS, 1985):

- não foram ainda implementados ou são pouco eficientes;
- foram desenvolvidos para a solução de problemas específicos.

Uma forma alternativa de cálculo do EQM dado pela equação (3.2) é a partir de um número infinito de dados de entrada saída:

$$J(\hat{g}) = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{l=1}^N (s_l - \hat{g}(\mathbf{x}_l))^2. \quad (3.3)$$

Observe que, se os vetores de entrada $\mathbf{x}_l$ forem gerados de acordo com a distribuição de probabilidade $d_P(\mathbf{x})$, o limite da expressão (3.3) converge sempre para o mesmo valor, independente dos pares de vetores $(\mathbf{x}_l, s_l)$ utilizados (HECHT-NIELSEN, 1990). Certamente existem conjuntos de pares $(\mathbf{x}_l, s_l)$ para os quais a série não converge, mas a probabilidade de ocorrência de tais conjuntos é zero. O comportamento de convergência pode ser observado já para valores de N finitos, mas suficientemente elevados.

Como, na prática, deve-se operar com valores de N finitos, um procedimento para se decidir um valor adequado para N é avaliar $J(\hat{g})$ para pares de vetores $(\mathbf{x}_l, s_l)$, $l=1, \dots, N$, aumentando N progressivamente, até que se verifique a tendência de convergência para um valor fixo, dado pela equação (3.3).

No entanto, nem sempre é possível tomar tantos pares de vetores de entrada-saída quanto desejado, já que geralmente existe um número fixo N de pares $(\mathbf{x}_l, s_l)$ a serem utilizados em um determinado problema de aproximação. Além disso, é recomendado que se utilize conjuntos de pares diferentes nas fases de aproximação (conjunto de aproximação) e avaliação do nível de aproximação (conjunto de teste ou validação), pois o objetivo da aproximação não é determinar uma função $\hat{g}$ capaz de aproximar g apenas para os vetores $\mathbf{x}_l$, $l=1, \dots, N$, pertencentes ao conjunto de aproximação, mas em toda a região de aproximação, ou seja, para todo $\mathbf{x} \in X$.

3.7 Aproximação suave: regularização

Nesta seção é analisado como o tipo de suavidade imposta à função a ser aproximada vai determinar o tipo de modelo de aproximação a ser utilizado. Técnicas que exploram restrições de suavidade em problemas de aproximação com base em dados de entrada-saída são conhecidas como regularização, tendo sido introduzidas por TIKHONOV (1963) para o tratamento de problemas de reconstrução de superfícies. Regularização é o procedimento de obtenção de um problema de aproximação bem-comportado a partir de um problema de aproximação mal-comportado pela imposição de restrições de suavidade ao modelo de aproximação.

A regularização conduz a problemas de aproximação bem-comportados por permitir a definição de um processo contínuo, mas não necessariamente único, que produza a melhor aproximação com base nos dados de entrada-saída. Este processo converge para um problema de otimização que deve minimizar uma função objetivo apresentando dois termos básicos: um expressando o erro de aproximação e o outro o nível de distanciamento da condição de suavidade.

Se as amostras da função a ser aproximada (superfície a ser reconstruída) forem totalmente aleatórias, não é possível estimar o valor da função (altura da superfície) em pontos não-amostrados, ou seja, não é possível realizar o processo de generalização. Ao impedir que o comportamento da função seja estritamente local, a existência de regularidade na função a ser aproximada acaba por viabilizar o processo de generalização (POGGIO & GIROSI, 1990b).

3.7.1 O problema de aproximação regularizado

Como já apresentado, o problema geral de aproximação é a definição de um modelo de aproximação $\hat{g}$ para uma função $g: X \subset \mathbb{R}^m \rightarrow \mathbb{R}^r$ a partir de um conjunto de dados $\{(\mathbf{x}_l, \mathbf{s}_l) \in \mathbb{R}^m \times \mathbb{R}^r\}_{l=1}^N$ amostrados na presença de ruído com base no mapeamento definido pela função em uma região compacta do espaço de aproximação².

Como este problema é mal-condicionado, a escolha de uma solução particular dentre as infinitas soluções possíveis está vinculada à atribuição de um grau apropriado de suavidade à função g . O problema geral de aproximação passa então a ser um problema de aproximação regularizado, que deve encontrar uma solução utilizando técnicas variacionais definidas com base no conjunto de dados e no tipo de suavidade imposta.

O grau de suavidade pode ser definido por um funcional $\phi(\hat{g})$ que assume valores tanto menores quanto mais suave for a função $\hat{g}$. Um procedimento comumente utilizado é medir o grau de suavidade com base no comportamento oscilatório da função $\hat{g}$ (GIROSI *et al.*, 1995). Deste modo, uma análise da função no domínio da frequência, permite afirmar que uma função é mais suave que outras quando apresenta menor quantidade de energia em altas frequências. O conteúdo de alta frequência de uma função pode ser medido adotando-se dois procedimentos básicos em sequência:

² Nesta seção é assumido $r = 1$, ou seja, $\mathbf{s} = s \in \mathbb{R}$, sendo que a extensão dos resultados desta seção para $r > 1$ é apresentada no capítulo 5, já no contexto de redes neurais artificiais.

- realizar uma filtragem da função por meio de um filtro passa-alta;
- medir a quantidade de energia ainda presente na função filtrada.

Sendo assim, a função de regularização ϕ pode ser dada na forma (GIROSI *et al.*, 1995):

$$\phi(\hat{g}) = \int_{\mathbf{z}} \frac{|\hat{G}(\mathbf{z})|^2}{F(\mathbf{z})} d\mathbf{z}, \quad (3.4)$$

onde $\hat{G}$ e F são as transformadas de Fourier de $\hat{g}$ e f , respectivamente. F é uma função definida positiva (MICCHELLI, 1986) previamente especificada que se aproxima de 0 quando $\|\mathbf{z}\| \rightarrow \infty$, de tal forma que $1/F$ corresponda a um filtro passa-alta. A função F é assumida ser simétrica (função par) e definida de forma a garantir a existência de $\hat{g}$ tal que a integral na equação (3.4) exista.

Neste caso, para N finito, o problema de aproximação regularizado é encontrar o modelo de aproximação $\hat{g}$ que minimize o seguinte funcional:

$$J(\hat{g}) = \frac{1}{N} \sum_{l=1}^N (s_l - \hat{g}(\mathbf{x}_l))^2 + \lambda \phi(\hat{g}), \quad (3.5)$$

onde λ é um número positivo denominado parâmetro de regularização. O primeiro termo, como discutido anteriormente, representa uma medida de distância entre g e $\hat{g}$, enquanto o segundo termo representa o grau de suavidade da função $\hat{g}$. O compromisso entre suavidade do processo de aproximação e fidelidade aos dados pertencentes ao conjunto de aproximação é, portanto, controlado pelo parâmetro de regularização λ . Certamente vai existir um valor ótimo λ^* para o parâmetro de regularização que conduza a um melhor desempenho em termos de generalização.

3.7.2 Fórmula geral de solução

Uma formulação detalhada e matematicamente rigorosa da solução do problema de minimização do funcional dado pela equação (3.5) pode ser encontrado em WAHBA (1990). Uma derivação simplificada é apresentada em (GIROSI *et al.*, 1995) e parcialmente reproduzida a seguir.

Considere a transformada de Fourier de $\hat{g}$ dada por

$$\hat{g}(\mathbf{x}) = c \int_{\mathbf{z}} \hat{G}(\mathbf{z}) e^{j\mathbf{x}^T \mathbf{z}} d\mathbf{z}. \quad (3.6)$$

Substituindo na equação (3.5) as expressões para ϕ e $\hat{g}$ dadas respectivamente pelas equações (3.4) e (3.6) tem-se:

$$J(\hat{G}) = \frac{1}{N} \sum_{l=1}^N \left(s_l - c \int_X \hat{G}(\mathbf{z}) e^{j\mathbf{x}_l^T \mathbf{z}} d\mathbf{z} \right)^2 + \lambda \int_X \frac{\hat{G}(-\mathbf{z}) \hat{G}(\mathbf{z})}{F(\mathbf{z})} d\mathbf{z}, \quad (3.7)$$

onde também se aplicou a igualdade $\hat{G}^*(\mathbf{z}) = \hat{G}(-\mathbf{z})$, derivada do fato de $\hat{g}$ ser real.

O mínimo de $J(\hat{G})$ é dado por $\hat{G}$ tal que

$$\frac{\partial J(\hat{G})}{\partial \hat{G}(\mathbf{t})} = 0, \quad \forall \mathbf{t} \in X. \quad (3.8)$$

Para cada termo no lado direito da expressão (3.7) tem-se:

$$\begin{aligned} \bullet \quad & \frac{\partial}{\partial \hat{G}(\mathbf{t})} \left[\frac{1}{N} \sum_{l=1}^N \left(s_l - c \int_X \hat{G}(\mathbf{z}) e^{j\mathbf{x}_l^T \mathbf{z}} d\mathbf{z} \right)^2 \right] = \frac{-2c}{N} \sum_{l=1}^N \left[(s_l - \hat{g}(\mathbf{x}_l)) \int_X \frac{\partial \hat{G}(\mathbf{z})}{\partial \hat{G}(\mathbf{t})} e^{j\mathbf{x}_l^T \mathbf{z}} d\mathbf{z} \right] \\ & = \frac{-2c}{N} \sum_{l=1}^N \left[(s_l - \hat{g}(\mathbf{x}_l)) \int_X \delta(\mathbf{z} - \mathbf{t}) e^{j\mathbf{x}_l^T \mathbf{z}} d\mathbf{z} \right] = \frac{-2c}{N} \sum_{l=1}^N \left[(s_l - \hat{g}(\mathbf{x}_l)) e^{j\mathbf{x}_l^T \mathbf{t}} \right] \\ \bullet \quad & \frac{\partial}{\partial \hat{G}(\mathbf{t})} \left[\lambda \int_X \frac{\hat{G}(-\mathbf{z}) \hat{G}(\mathbf{z})}{F(\mathbf{z})} d\mathbf{z} \right] = 2\lambda \int_X \frac{\hat{G}(-\mathbf{z})}{F(\mathbf{z})} \cdot \frac{\partial \hat{G}(\mathbf{z})}{\partial \hat{G}(\mathbf{t})} d\mathbf{z} = 2\lambda \int_X \frac{\hat{G}(-\mathbf{z})}{F(\mathbf{z})} \delta(\mathbf{z} - \mathbf{t}) d\mathbf{z} = 2\lambda \frac{\hat{G}(-\mathbf{t})}{F(\mathbf{t})} \end{aligned}$$

Com isso, a equação (3.8) pode ser reescrita na forma:

$$\frac{-c}{N} \sum_{l=1}^N \left[(s_l - \hat{g}(\mathbf{x}_l)) e^{j\mathbf{x}_l^T \mathbf{t}} \right] + \lambda \frac{\hat{G}(-\mathbf{t})}{F(\mathbf{t})} = 0, \quad \forall \mathbf{t} \in X. \quad (3.9)$$

Mudando a variável $\mathbf{t}$ para $-\mathbf{t}$, é possível expressar $\hat{G}(\mathbf{t})$ como segue:

$$\hat{G}(\mathbf{t}) = \frac{c}{N} F(\mathbf{t}) \sum_{l=1}^N \left[\frac{s_l - \hat{g}(\mathbf{x}_l)}{\lambda} e^{-j\mathbf{x}_l^T \mathbf{t}} \right], \quad \forall \mathbf{t} \in X. \quad (3.10)$$

A anti-transformada de Fourier de $\hat{G}$ é, então, dada por:

$$\hat{g}(\mathbf{x}) = \frac{c}{N} \sum_{l=1}^N \frac{s_l - \hat{g}(\mathbf{x}_l)}{\lambda} \delta(\mathbf{x}_l - \mathbf{x}) * f(\mathbf{x}) = \frac{c}{N} \sum_{l=1}^N \frac{s_l - \hat{g}(\mathbf{x}_l)}{\lambda} f(\mathbf{x} - \mathbf{x}_l). \quad (3.11)$$

Se f é uma função definida positiva, a função de regularização ϕ é uma norma. No entanto, se f for apenas condicionalmente definida positiva, ϕ é uma semi-norma (veja apêndice A) que apresenta um espaço nulo ϕ de dimensão finita k (MADYCH & NELSON, 1988;

MADYCH & NELSON, 1990). Portanto, para f condicionalmente definida positiva, existe uma classe de funções que não são detectáveis por ϕ . Neste caso, todas as funções $\hat{g}$ que diferem apenas em um termo que pertença ao espaço nulo de ϕ são consideradas equivalentes. Desta forma, a solução mais geral da minimização do funcional dado pela equação (3.5) é:

$$\hat{g}(\mathbf{x}) = \sum_{l=1}^N a_l f(\mathbf{x} - \mathbf{x}_l) + p(\mathbf{x}), \quad (3.12)$$

onde

$$a_l = \frac{c}{N} \cdot \frac{s_l - \hat{g}(\mathbf{x}_l)}{\lambda}$$

e

$$p(\mathbf{x}) = \begin{cases} 0 & \text{se } f \text{ for definida positiva} \\ \sum_{j=1}^k b_j p_j(\mathbf{x}) & \text{se } f \text{ for condicionalmente definida positiva} \end{cases}$$

O conjunto $\{p_j\}_{j=1}^k$ é uma base de ϕ .

Os coeficientes a_l ($l=1, \dots, N$) e b_j ($j=1, \dots, k$) satisfazem o seguinte sistema de equações lineares:

$$\begin{aligned} (\mathbf{F} + \lambda \mathbf{I})\mathbf{a} + \mathbf{P}^T \mathbf{b} &= \mathbf{y}, \\ \mathbf{P}\mathbf{a} &= \mathbf{0} \end{aligned} \quad (3.13)$$

onde

- $\mathbf{I}$ é a matriz identidade de ordem N ;
- $\mathbf{y} = [s_1 \ s_2 \ \dots \ s_N]^T$;
- $\mathbf{a} = [a_1 \ a_2 \ \dots \ a_N]^T$;
- $\mathbf{b} = [b_1 \ b_2 \ \dots \ b_k]^T$;
- $\mathbf{F} = \begin{bmatrix} f(\mathbf{0}) & f(\mathbf{x}_1 - \mathbf{x}_2) & \dots & f(\mathbf{x}_1 - \mathbf{x}_N) \\ f(\mathbf{x}_2 - \mathbf{x}_1) & f(\mathbf{0}) & & \vdots \\ \vdots & & \ddots & \\ f(\mathbf{x}_N - \mathbf{x}_1) & \dots & & f(\mathbf{0}) \end{bmatrix}$;
- $\mathbf{P} = \begin{bmatrix} p_1(\mathbf{x}_1) & p_1(\mathbf{x}_2) & \dots & p_1(\mathbf{x}_N) \\ p_2(\mathbf{x}_1) & p_2(\mathbf{x}_2) & & \vdots \\ \vdots & & \ddots & \\ p_k(\mathbf{x}_1) & \dots & & p_k(\mathbf{x}_N) \end{bmatrix}$.

Observe que, para $\mathbf{P} = \mathbf{0}$, se a função f é definida positiva, ou seja, se $\det(\mathbf{F}) > 0$ para todo conjunto $\{\mathbf{x}_l\}_{l=1}^N$, sempre existe solução para o sistema linear (3.13). O caso $\lambda = 0$ corresponde a interpolação simples.

3.8 Extensões do problema de aproximação regularizado

Ainda assumindo um conjunto de dados amostrados do tipo $\{(\mathbf{x}_l, s_l) \in \mathcal{R}^m \times \mathcal{R}^r\}_{l=1}^N$, com $r = 1$, esta seção apresenta três extensões do problema de aproximação regularizado, seguindo ainda os resultados de GIROSI *et al.* (1995). Estas extensões, apesar de poderem ocorrer isoladamente, são introduzidas em sequência junto ao problema de aproximação original, convergindo para uma solução a ser melhor explorada no próximo capítulo.

3.8.1 O problema de aproximação regularizado aditivo

Justificada pelos resultados apresentados no capítulo anterior, uma extensão importante do problema de aproximação regularizado é assumir que o modelo de aproximação é formado por uma composição aditiva de funções monovariáveis na forma:

$$\hat{g}(\mathbf{x}) = \sum_{j=1}^m \hat{g}_j(x_j), \quad (3.14)$$

onde x_j é o j -ésimo componente do vetor de entrada $\mathbf{x}$. Sendo assim, é possível impor condições de suavidade para cada componente aditivo $\hat{g}_j$ ($j=1, \dots, m$) ao invés de aplicar o critério de regularização diretamente sobre a função de aproximação $\hat{g}$. Neste caso, o problema de minimização da expressão apresentada na equação (3.5) pode ser reformulado como segue:

$$\min_{\hat{g}} J(\hat{g}) = \min_{\hat{g}} \frac{1}{N} \sum_{l=1}^N (s_l - \hat{g}(\mathbf{x}_l))^2 + \lambda \phi(\hat{g}) = \min_{\hat{g}_1, \dots, \hat{g}_m} \frac{1}{N} \sum_{l=1}^N \left(s_l - \sum_{j=1}^m \hat{g}_j(x_{jl}) \right)^2 + \lambda \sum_{j=1}^m \frac{1}{d_j} \int_{x_j} \frac{|\hat{G}_j(z)|^2}{F(z)} dz, \quad (3.15)$$

onde, para $j=1, \dots, m$:

- $\hat{G}_j$ é a transformada de Fourier de $\hat{g}_j$;
- d_j são constantes positivas dadas, que permitem a imposição de graus de suavidade distintos para cada componente aditivo de $\phi(\hat{g})$;

- X_j são intervalos fechados onde a variável x_j assume valores, de tal forma que $X_1 \times \dots \times X_m \equiv X$.

Seguindo os mesmos passos adotados na seção anterior, a solução deste problema de aproximação regularizado, desconsiderando os termos pertencentes ao espaço nulo de ϕ , assume a forma:

$$\hat{g}(\mathbf{x}) = \sum_{l=1}^N a_l f_{tot}(\mathbf{x} - \mathbf{x}_l), \quad (3.16)$$

onde

$$f_{tot}(\mathbf{x} - \mathbf{x}_l) = \sum_{j=1}^m d_j f(x_j - x_{jl}) \quad (3.17)$$

Substituindo a equação (3.17) na equação (3.16) e comparando com a equação (3.14), os componentes aditivos $\hat{g}_j$ ($j=1, \dots, m$) assumem a forma:

$$\hat{g}_j(x_j) = \sum_{l=1}^N c_{jl} f(x_j - x_{jl}), \text{ com } c_{jl} = d_j a_l.$$

Observe que, com as constantes d_j fixas, os componentes aditivos não são independentes. Por outro lado, assumindo d_j ($j=1, \dots, m$) como parâmetros livres, a função de aproximação $\hat{g}$ pode ser expressa na forma:

$$\hat{g}(\mathbf{x}) = \sum_{l=1}^N \sum_{j=1}^m c_{jl} f(x_j - x_{jl}), \quad (3.18)$$

onde agora os coeficientes c_{jl} são independentes.

3.8.2 A redução no número de parâmetros

Ainda desconsiderando os termos referentes ao espaço nulo de ϕ , tem-se que, na solução do problema de regularização original, dada pela equação (3.12), o número N de parâmetros a determinar é igual ao número de dados. Por outro lado, ao assumir uma função de aproximação aditiva na forma da equação (3.14), resulta uma solução com $N.m$ parâmetros, dada pela equação (3.18). Este número de parâmetros, além de exceder o número de dados para $d > 1$, ainda depende da dimensão do vetor de entrada.

Embora ainda seja possível encontrar uma solução minimizante, modelos de aproximação com tal proporção de parâmetros frente aos dados disponíveis perdem a

capacidade de generalização e são diretamente atingidos pela maldição da dimensionalidade (STONE, 1982). É recomendada, portanto, uma redução do número de funções básicas, mantendo a forma do modelo de aproximação. Com isso, a nova função de aproximação, obtida a partir da equação (3.18), assume a forma:

$$\hat{g}(\mathbf{x}) = \sum_{l=1}^{N'} \sum_{j=1}^m c_{jl} f(x_j - t_{jl}) \quad (3.19)$$

onde $N' < N$ e t_{jl} ($j=1, \dots, m; l=1, \dots, n$) são novos parâmetros a serem determinados. Logicamente, N' deve ser escolhido tal que $2N'm < Nm$.

Observe que os componentes aditivos de $\hat{g}$ continuam a ser independentes, correspondendo a funções unidimensionais descritas como segue:

$$\hat{g}_j(x_j) = \sum_{l=1}^{N'} c_{jl} f(x_j - t_{jl})$$

3.8.3 O pré-processamento das variáveis de entrada

Em problemas de aproximação de funções multivariáveis, uma etapa fundamental é a escolha adequada do conjunto de variáveis de entrada. É importante considerar a possibilidade de ocorrência dos seguintes fatores (GIROSI *et al.*, 1995):

- algumas variáveis podem ser mais relevantes que outras;
- algumas variáveis podem ser totalmente irrelevantes;
- dentre as variáveis relevantes, existem aquelas que podem ser expressas pela combinação linear de um conjunto de variáveis originais.

Neste caso, a solução do problema de aproximação pode ser melhor representada não a partir de $\mathbf{x}$, o vetor original de variáveis de entrada, mas a partir de $\mathbf{Vx}$, uma transformação linear deste vetor, onde $\mathbf{V}$ é uma matriz retangular de dimensão $n \times m$. A transformação $\mathbf{x} \rightarrow \mathbf{Vx}$ pode ser inserida diretamente no problema de aproximação regularizado, produzindo uma solução equivalente à equação (3.19). Novamente desconsiderando os termos pertencentes ao espaço nulo de Φ , esta solução assume a forma:

$$\hat{g}(\mathbf{x}) = \sum_{l=1}^{N'} \sum_{j=1}^n c_{jl} f(\mathbf{v}_j^T \mathbf{x} - t_{jl}), \quad (3.20)$$

onde $\mathbf{v}_j$ ($j=1, \dots, n$) é o vetor coluna formado pelos componentes da j -ésima linha da matriz $\mathbf{V}$.

Definindo as seguintes funções unidimensionais

$$f_j(y) = \sum_{l=1}^{N'} c_{jl} f(y - t_{jl}), \text{ com } j=1, \dots, n, \quad (3.21)$$

a equação (3.20) pode então ser reescrita como segue:

$$\hat{g}_n(\mathbf{x}) = \sum_{j=1}^n f_j(\mathbf{v}_j^T \mathbf{x}). \quad (3.22)$$

Portanto, a função de aproximação resulta em uma composição aditiva de funções monovariáveis a determinar, com argumento representado por uma projeção das variáveis de entrada em uma direção também a determinar. A escolha da função de base f na equação (3.21) está associada ao tipo de suavidade imposta ao problema de regularização.

3.9 Generalização e sobreajuste

Como já discutido anteriormente, $\hat{g}$ deve aproximar g em todo o espaço compacto X , ou seja, o problema de aproximação original assume a forma:

$$\min_{\hat{g}} J(\hat{g}) = \min_{\hat{g}} \int_X (g(\mathbf{x}) - \hat{g}(\mathbf{x}))^2 d_p(\mathbf{x}) d\mathbf{x},$$

onde $d_p(\mathbf{x})$ é a densidade de probabilidade da distribuição de $\mathbf{x}$ em X . Tomando amostras de $\mathbf{x} \in X$ de acordo com a distribuição $d_p(\mathbf{x})$, este problema pode ser descrito de forma totalmente análoga como segue:

$$\min_{\hat{g}} J(\hat{g}) = \min_{\hat{g}} \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{l=1}^N (s_l - \hat{g}(\mathbf{x}_l))^2.$$

Por limitações de ordem prática, toma-se N finito, levando a um problema de aproximação mal-comportado. Neste caso, introduz-se um termo adicional que penaliza a ausência de suavidade na função de aproximação $\hat{g}$, convergindo para um problema variacional bem-comportado na forma:

$$\min_{\hat{g}} J(\hat{g}) = \min_{\hat{g}} \frac{1}{N} \sum_{l=1}^N (s_l - \hat{g}(\mathbf{x}_l))^2 + \lambda \phi(\hat{g}).$$

Como será visto no próximo capítulo para um caso específico, aplicando-se um procedimento denominado validação cruzada generalizada, é possível obter um valor ótimo para λ que garanta um bom desempenho em termos de generalização.

No entanto, foi também incorporado ao problema de aproximação, dentre outras modificações, a característica aditiva da função de aproximação, conduzindo a um problema de aproximação regularizado aditivo da forma:

$$\min_{\hat{g}_n} J(\hat{g}_n) = \min_{f_1, \dots, f_n} \frac{1}{N} \sum_{l=1}^N \left(s_l - \sum_{j=1}^n f_j(\mathbf{v}_j^T \mathbf{x}_l) \right)^2 + \sum_{j=1}^n \lambda_j \phi_j(f_j).$$

Neste caso, existe um parâmetro de regularização e uma função de penalização para ausência de suavidade para cada componente aditivo da função de aproximação $\hat{g}_n$, que assume a forma:

$$\hat{g}_n(\mathbf{x}) = \sum_{j=1}^n f_j(\mathbf{v}_j^T \mathbf{x}).$$

Com relação ao desempenho na generalização, o mesmo procedimento de validação cruzada generalizada mencionado acima pode ser aplicado para determinar os valores ótimos de λ_j ($j=1, \dots, n$). No entanto, isto é insuficiente para garantir uma boa capacidade de generalização, pois, com a adoção de um modelo aditivo, criou-se mais um parâmetro livre: o número n de funções básicas f_j ($j=1, \dots, n$).

Valores de n muito elevados conduzem ao fenômeno de sobreajuste do processo de aproximação, ou seja, a função de aproximação passa a incorporar o ruído como informação relevante para o modelo de aproximação. Por outro lado, valores de n muito pequenos impedem que o modelo de aproximação tenha flexibilidade suficiente para representar adequadamente todas as associações não-lineares presentes nos dados de aproximação. Portanto, deve existir um valor intermediário ótimo para n , que conduza ao melhor desempenho na generalização.

Novamente, um método de validação cruzada vem fornecer uma alternativa para a obtenção deste valor ótimo para n . Para tanto, é necessário que exista não apenas o conjunto de aproximação $\{(\mathbf{x}_l, s_l)\}_{l=1}^N$, mas também um conjunto de teste $\{(\mathbf{x}'_l, s'_l)\}_{l=1}^{N'}$ obtido a partir do mesmo processo de amostragem. Por exemplo, é possível dividir aleatoriamente o conjunto de dados disponíveis para formar estes dois conjuntos, assumidos independentes.

Sendo assim, cria-se um problema de aproximação implícita e explicitamente regularizado na forma:

$$\begin{aligned} \min_n J'(\hat{g}_n) &= \min_{\hat{g}_n} \frac{1}{N'} \sum_{l=1}^{N'} \left(s'_l - \hat{g}_n(\mathbf{x}'_l) \right)^2 \\ \text{sujeito a } \hat{g}_n &= \arg \min_{\hat{g}_n} J(\hat{g}_n) = \arg \min_{f_1, \dots, f_n} \frac{1}{N} \sum_{l=1}^N \left(s_l - \sum_{j=1}^n f_j(\mathbf{v}_j^T \mathbf{x}_l) \right)^2 + \sum_{j=1}^n \lambda_j \phi_j(f_j) \end{aligned}$$

Este problema de aproximação resultante é coerente no sentido de requerer que a aproximação seja medida não apenas em relação ao conjunto de aproximação, mas também em relação a outros dados igualmente representativos do comportamento da função a ser aproximada, mas que não participam na definição do modelo de aproximação.

Capítulo 4

Aproximação de funções utilizando modelos de projeção

Considerando problemas de aproximação de funções multidimensionais definidas em uma região compacta de um espaço vetorial de dimensão finita, modelos não-paramétricos de aproximação são aqueles que não impõem qualquer tipo de restrição com relação ao modelo de aproximação, garantindo a flexibilidade necessária para aproximar qualquer tipo de função.

Para evitar uma degradação de desempenho com o aumento da dimensão do espaço compacto de aproximação, esta total flexibilidade do processo de aproximação requer um aumento exponencial no número de amostras da função, o que acarreta um aumento proporcional na carga computacional. Diz-se então que o problema de aproximação não-paramétrica está sujeito à maldição da dimensionalidade.

De fato, existe um compromisso entre flexibilidade do modelo de aproximação e imunidade à maldição da dimensionalidade (STONE, 1982), e esse compromisso pode ser representado na forma de um critério de aproximação que contemple a fidelidade aos dados e a imposição de um certo nível de suavidade junto à função de aproximação. Tendo que atender a restrições de suavidade, a função de aproximação passa a ser menos flexível, podendo ser considerada um caso intermediário entre modelos totalmente paramétricos e modelos totalmente não-paramétricos. Estes modelos intermediários geralmente continuam a ser denominados não-paramétricos, pois a estrutura final do processo de aproximação não pode ser antecipada, passando a depender fundamentalmente do conjunto de amostras da função na região de aproximação e do tipo de restrição de suavidade imposta ao modelo.

Conforme discutido no capítulo 3, as restrições impostas ao modelo de aproximação geralmente confiam na existência de um certo nível de redundância no conjunto de amostras. Neste sentido, modelos não-paramétricos de aproximação baseados em técnicas de redução de dimensionalidade têm sido empregados com sucesso na aproximação de várias classes de funções multidimensionais definidas em uma região compacta de espaços de dimensão finita

(DONOHO & JOHNSTONE, 1989). Não está descartada, no entanto, a possibilidade de se realizar também uma redução de dimensionalidade prévia, empregando, por exemplo, análise de componentes principais para descartar componentes pouco significativos (JOLLIFFE, 1986).

Modelos de redução de dimensionalidade que utilizam uma composição aditiva de funções de expansão ortogonal a uma direção adequadamente fixada podem ser prontamente obtidos a partir da imposição de restrições de suavidade específicas. Tendo sido deduzidos no capítulo anterior e a serem melhor investigados neste capítulo, estes modelos de aproximação também recebem a denominação de modelos de projeção e assumem a forma:

$$\hat{g}_n(\mathbf{x}) = \sum_{j=1}^n f_j(\mathbf{v}_j^T \mathbf{x}), \quad (4.1)$$

onde $\mathbf{x} \in \mathfrak{R}^m$ é o vetor de variáveis de entrada, $\mathbf{v}_j \in \mathfrak{R}^m$ é a direção de projeção ($j=1, \dots, n$) e o produto escalar $\mathbf{v}_j^T \mathbf{x}$ pode ser tomado, a menos de uma constante, como uma projeção de $\mathbf{x}$ na direção $\mathbf{v}_j$. Observe que o j -ésimo termo $f_j(\cdot)$ do somatório é constante para valores de $\mathbf{x}$ em hiperplanos do $\mathfrak{R}^m$ definidos na forma $\mathbf{v}_j^T \mathbf{x} = c$, para $c \in \mathfrak{R}$ constante.

A utilização de modelos na forma da equação (4.1) conduz a processos de aproximação por expansão ortogonal que apresentam inúmeras vantagens estatísticas e computacionais, como insensibilidade a fatores de escala e à estrutura de correlação dos dados (FRIEDMAN, 1987). Além disso, é sabido que modelos de projeção são menos afetados pela presença de dados não-informativos, isto é, corrompidos por ruídos espúrios (HUBER, 1985; CHENTOUF & JUTTEN, 1995).

4.1 Aproximação por expansão ortogonal aditiva

Com base no conjunto de dados de aproximação, ou seja, pontos no espaço de aproximação, um bem-sucedido modelo de aproximação utilizando funções aditivas foi proposto por FRIEDMAN & STUETZLE (1981), assumindo exatamente a forma da equação (4.1). Neste modelo, os termos da composição aditiva correspondem a funções escalares de expansão ortogonal a projeções unidirecionais, sendo denominado modelo de aproximação por busca de projeção.

Técnicas de busca de projeção para análise exploratória de dados foram originalmente propostas por KRUSKAL (1969) e SWITZER (1970). No entanto, foram FRIEDMAN & TUKEY (1974) que primeiramente apresentaram uma implementação eficiente, que só então passou a receber a denominação de busca de projeção (*projection pursuit*).

A projeção consiste de operações lineares em que um mapeamento de uma determinada dimensão tem suprimidas algumas de suas estruturas de modo a tornar possível sua representação em espaços de menor dimensão. Neste caso, a estrutura projetada pode ser considerada como uma *sombra* da estrutura original, fazendo com que as projeções mais *interessantes* sejam aquelas que preservam parcelas representativas da estrutura original.

A busca destas direções de projeção envolve uma série de manipulações do conjunto de dados de aproximação disponível. É importante que as projeções, ao contrário daquelas obtidas via análise de componentes principais, sejam invariantes afins (HUBER, 1985), isto é, sejam insensíveis à estrutura de correlação dos dados, a qual é capaz de representar todas as associações lineares, mas geralmente tem pouco a ver com as associações não-lineares existentes. Portanto, baseadas nos dados de entrada-saída, as direções de projeção devem enfatizar as relações não-lineares existentes entre as variáveis do problema de aproximação.

A questão que surge é como obter de forma automática estas direções de projeção. Uma alternativa foi apresentada por FRIEDMAN & TUKEY (1974), em que a direção de projeção corresponde à solução que maximiza um determinado índice numérico de projeção. A partir de então, uma série de índices foram apresentados na literatura, cada qual evidenciando um conjunto particular de características a serem atendidas pelos dados projetados.

Em virtude da inexistência de um índice de projeção que se aplique a todos os casos, e como geralmente não existe um conhecimento prévio a respeito das características presentes no conjunto de dados de aproximação, a utilização de um índice de desempenho na determinação da direção de projeção, ao invés de determiná-la arbitrariamente, permite assegurar apenas um aumento da probabilidade de se encontrar direções de projeção *interessantes*. Na grande maioria dos problemas de aproximação, este aumento da probabilidade é bastante significativo (HUBER, 1985).

Uma vez definida a direção de projeção, funções monovariáveis são então determinadas de tal forma que sua expansão em direções ortogonais à direção de projeção forneça a melhor aproximação possível com base nos dados disponíveis. No caso de projeções unidirecionais, a expansão ortogonal ocorre em hiperplanos do espaço de aproximação.

A interdependência entre as funções de expansão ortogonal e a correspondente direção de projeção acaba conduzindo a processos iterativos de aproximação. Para que estes processos sejam computacionalmente eficientes, é geralmente necessário que:

- cada passo do processo iterativo demande a menor quantidade de cálculo possível, o que geralmente conduz à necessidade de se recorrer a informações variacionais de 1ª e 2ª ordem (LUENBERGER, 1989);
- propriedades teóricas que garantam redução de dimensionalidade estejam presentes, fazendo com que a aproximação em espaços multidimensionais apresentem taxas de convergência típicas de problemas de aproximação de menor dimensão (HÄRDLE & STOKER, 1989; STONE, 1982).

Dentre as vantagens de se utilizarem projeções unidirecionais têm-se a manutenção de uma maior simplicidade do processo de aproximação e a possibilidade de visualizar graficamente o comportamento da função de aproximação na direção de projeção. Além disso, é possível explicitar o tipo de associação não-linear existente entre as variáveis (associações lineares entre as variáveis são extraídas diretamente da matriz de covariância) e uma série de outras informações que não estão diretamente disponíveis considerando-se a dimensão completa do espaço de aproximação.

No entanto, apesar de não ser o objetivo deste estudo, é importante salientar que a interpretação dos dados projetados geralmente não se apresenta como uma tarefa simples. Estruturas interessantes presentes na projeção dos dados não necessariamente correspondem à projeção de estruturas interessantes, da mesma forma que estruturas interessantes presentes nos dados podem conduzir a nenhuma projeção interessante (GOVER, 1987).

Além disso, é improvável que exista apenas uma única direção de projeção capaz de explicitar todo tipo de informação a respeito do conjunto multidimensional de dados de entrada-saída, o que seria equivalente a assumir que o problema de aproximação é monovariável. Mesmo que este seja o caso, nem sempre é possível garantir a determinação exata desta direção de projeção. Portanto, justifica-se o estabelecimento de uma sequência de direções de projeção, cada qual explicitando a maior parcela possível de informação necessária para o sucesso da tarefa de aproximação.

Com isso, após a definição de uma direção de projeção e da correspondente função de expansão ortogonal, uma transformação deve ser aplicada ao conjunto de dados para que a informação já representada seja removida, permitindo o reinício do processo a partir de uma nova direção de projeção. Logo, a busca sequencial de direções de projeção pode ser implementada na forma (FRIEDMAN *et al.*, 1984): encontra-se um direção de projeção *ótima*, remove-se do conjunto de dados a estrutura resultante da projeção dos dados nesta direção e

reinicia-se o processo até que nenhuma outra projeção revele qualquer estrutura, ou seja, até que o modelo de aproximação concorde com os dados amostrados em todas as projeções. Ao invés de operar com subtração de vetores (remoção junto ao conjunto de dados das estruturas já representadas), FUJITA (1992) resolve um problema semelhante (no contexto de redes neurais artificiais) utilizando um método alternativo, baseado em complementação ortogonal de vetores, sendo, no entanto, mais custoso em termos computacionais e numericamente mal-condicionado.

Este procedimento iterativo e construtivo, em que cada novo sub-problema de aproximação deve representar apenas informações não representadas pelos sub-problemas de aproximação anteriores, produz funções de aproximação multivariável utilizando composição aditiva de funções monovariáveis, na forma da equação (4.1).

Como as funções $f_j(\cdot)$ ($j=1,\dots,n$) são constantes para valores de $\mathbf{x}$ em hiperplanos do $\mathbb{R}^m$, elas são denominadas funções de expansão ortogonal a uma determinada direção – *ridge function* (DAHMEN & MICCHELLI, 1987). Tomando $m = 2$ e $f_j(\cdot)$ arbitrário, as figuras 1(a) e 1(b) permitem verificar esta propriedade. Observe que a expansão é ortogonal à direção de projeção $\mathbf{v} = [1 \ 0]^T$.

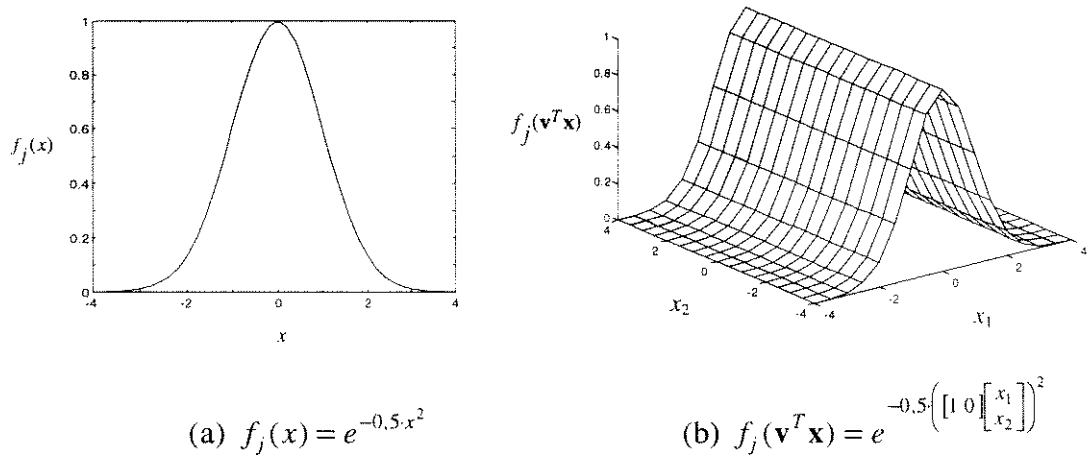


Figura 4.1 - Função de expansão ortogonal em que $\mathbf{v} = [1 \ 0]^T$ e $\mathbf{x} = [x_1 \ x_2]^T$

4.2 O problema de aproximação resultante

O problema de aproximação por expansão ortogonal aditiva pode ser completamente descrito na forma (FRIEDMAN & STUETZLE, 1981; HUBER, 1985):

- Seja X uma região compacta do $\mathbb{R}^m$ e seja $g: X \subset \mathbb{R}^m \rightarrow \mathbb{R}$ a função a ser aproximada;

- O conjunto de dados de aproximação $\{(\mathbf{x}_l, s_l) \in \mathfrak{R}^m \times \mathfrak{R}\}_{l=1}^N$ é gerado considerando-se que os vetores de entrada $\mathbf{x}_l$ estão distribuídos na região compacta $X \subset \mathfrak{R}^m$ de acordo com uma função de densidade de probabilidade fixa $d_P: X \subset \mathfrak{R}^m \rightarrow [0,1]$ e que os vetores de saída s_l são produzidos pelo mapeamento definido pela função g na forma:

$$s_l = g(\mathbf{x}_l) + \varepsilon_l, \quad l = 1, \dots, N,$$

onde $\varepsilon_l \in \mathfrak{R}$ é uma variável aleatória de média zero e variância fixa.

- A função g que associa a cada vetor de entrada $\mathbf{x} \in X$ uma saída escalar $s \in \mathfrak{R}$ pode ser aproximada com base no conjunto de dados de aproximação $\{(\mathbf{x}_l, s_l) \in \mathfrak{R}^m \times \mathfrak{R}\}_{l=1}^N$ por uma composição aditiva de funções de expansão ortogonal na forma:

$$g(\mathbf{x}) \approx \hat{g}_n(\mathbf{x}) = \sum_{j=1}^n f_j(\mathbf{v}_j^T \mathbf{x}),$$

sendo que as funções de expansão ortogonal $f_j(\cdot)$, por serem constantes em hiperplanos do espaço de aproximação, são consideradas como uma generalização de funções lineares. Por motivações de ordem numérica e por analogia com operadores de projeção, é interessante, sempre que possível, tomar direções de projeção de norma unitária tal que $\mathbf{v}_j^T \mathbf{v}_j = 1$ ($j=1, \dots, n$).

- Assuma que os primeiros $n-1$ termos já foram determinados, ou seja, os vetores $\mathbf{v}_j$ e as funções $f_j(\cdot)$ ($j=1, \dots, n-1$). Sejam

$$d_l = s_l - \sum_{j=1}^{n-1} f_j(\mathbf{v}_j^T \mathbf{x}_l), \quad l=1, \dots, N,$$

os resíduos do processo de aproximação. Obtenha a direção de projeção $\mathbf{v}_n$ e a função $f_n(\cdot)$, soluções do seguinte problema variacional com restrição de suavidade (veja capítulo 3):

$$\min_{\mathbf{v}_n, f_n} \frac{1}{N} \sum_{l=1}^N \left(d_l - f_n(\mathbf{v}_n^T \mathbf{x}_l) \right)^2 + \lambda_n \phi(f_n) \quad (4.2)$$

- Faça $n = n+1$ e repita o processo a partir do cálculo dos novos resíduos enquanto o nível de aproximação desejado ainda não foi atingido.

Este processo de aproximação tem algumas propriedades importantes:

- a aproximação por expansão ortogonal aditiva apresenta robustez a dados não-informativos (HUBER, 1985; CHENTOUF & JUTTEN, 1995);

- assumindo que a função g a ser aproximada é quadraticamente integrável, uma condição quase sempre satisfeita em regiões compactas de espaços multidimensionais, HUBER (1985) conjecturou a convergência absoluta da aproximação dada pela equação (4.1), o que mais tarde foi demonstrado por JONES (1987). Além disso, HALL (1989b) demonstrou que a taxa de convergência do processo é $\sqrt{n}$ -consistente e independente da dimensão m do espaço de entrada;
- ainda tomando g quadraticamente integrável, DONOHO & JOHNSTONE (1989) demonstraram, para o caso bidimensional ($m = 2$), que a taxa de convergência deste processo é igual ou superior àquela fornecida por métodos tradicionais de aproximação. A convergência é tanto mais rápida quanto mais estruturalmente aditiva for a não-linearidade presente nas associações entre as variáveis do problema de aproximação. Os resultados de DONOHO & JOHNSTONE (1989) foram estendidos para $m > 2$ por ZHAO & ATKESON (1992).

4.3 Condições para aproximação exata

A partir de resultados apresentados no capítulo 2, especialmente aqueles vinculados ao teorema de Stone-Weierstrass, é sempre possível obter uma aproximação para g na forma dada pela equação (4.1). Por exemplo, seja $X = [0,1]^m$:

- Para escalares reais a_j , a função $\sum_j a_j e^{\mathbf{v}_j^T \mathbf{x}}$ é densa em $C[X]$, ou seja, a função $\sum_j a_j e^{\mathbf{v}_j^T \mathbf{x}}$ tem $C[X]$ como seu fecho uniforme;
- Para escalares complexos a_j, b_j , a função $\sum_j [a_j \cos(\mathbf{v}_j^T \mathbf{x}) + b_j \sin(\mathbf{v}_j^T \mathbf{x})]$ é densa em $L_2[X]$.

Apesar da possibilidade de aproximar arbitrariamente bem determinadas famílias de funções em intervalos compactos, para efeito de avaliação de algoritmos de aproximação é interessante determinar sob que condições a aproximação pode vir a ser exata. Sendo possível a aproximação exata, ou seja, $g(\mathbf{x}) \equiv \hat{g}(\mathbf{x}) \quad \forall \mathbf{x} \in X$, outro aspecto importante é verificar se a representação de $\hat{g}(\mathbf{x})$ é única.

DIACONIS & SHAHSHAHANI (1984) apresentaram resultados que estabelecem condições necessárias e suficientes para que uma função $g: X \subset \mathbb{R}^m \rightarrow \mathbb{R}$ possa ser representada exatamente como uma composição aditiva de n funções monovariáveis na forma da equação (4.1). Como consequência desses resultados têm-se:

- Todo funcional polinomial pode ser representado exatamente por uma aproximação dada pela equação (4.1), com n finito.
- Os vetores $\mathbf{v}_j$ devem ser determinados em função do conjunto de dados de aproximação, pois a utilização de vetores $\mathbf{v}_j$ arbitrários pode inviabilizar a obtenção da representação exata, se esta for possível.
- De todos os funcionais que podem ser representados exatamente por uma aproximação dada pela equação (4.1), com n finito, os funcionais polinomiais são os únicos que não possuem representação única. Por exemplo, para $\mathbf{x} = [x_1 \ x_2]^T$, com $g(\mathbf{x}) = x_1 x_2$ e $\hat{g}_2(\mathbf{x}) = \frac{1}{4ab}(ax_1 + bx_2)^2 - \frac{1}{4ab}(ax_1 - bx_2)^2$, tem-se que $g(\mathbf{x}) \equiv \hat{g}_2(\mathbf{x})$ para todo $a, b \in \mathfrak{R}$ tal que $ab \neq 0$ e $a^2 + b^2 = 2$.
- Para $\mathbf{x} = [x_1 \ x_2]^T$, funcionais como $g(\mathbf{x}) = e^{x_1 x_2}$ e $g(\mathbf{x}) = \sin(x_1 x_2)$ não podem ser representados exatamente por uma aproximação dada pela equação (4.1), com n finito.

Na verdade, o número de funcionais que não podem ser representados exatamente por uma aproximação dada pela equação (4.1), com n finito, é bem maior que o número de funcionais que podem assumir exatamente tal representação. Como discutido no capítulo 3, na prática não existe a preocupação de resolver o problema de aproximação exatamente, mas sim a de obter a melhor aproximação que atenda a uma restrição de suavidade previamente definida. Sendo assim, o capítulo 2 apresenta as motivações para a aplicação de modelos aditivos no tratamento do problema de aproximação em discussão. Resultados específicos para modelos aditivos de expansão ortogonal na forma da equação (4.1) podem ser encontrados em LOGAN (1975) e LOGAN & SHEPP (1975), restritos ao caso $m = 2$.

4.4 Busca da direção inicial de projeção (*projection pursuit*)

Conforme descrito na seção 4.2, o problema de aproximação multidimensional por composição aditiva de funções monovariáveis de expansão ortogonal pode ser formulado como uma composição de problemas mais simples na forma da equação (4.2), em que devem ser definidas funções monovariáveis de forma a aproximar projeções de g em planos do espaço de aproximação.

Uma projeção linear do $\mathfrak{R}^m$ para o $\mathfrak{R}^k$ ($k \leq m$) pode ser definida como uma transformação linear realizada por qualquer matriz $\mathbf{V}$ de dimensão $m \times k$ e posto k na forma:

$$\mathbf{z} = \mathbf{V}^T \mathbf{x}, \quad \mathbf{x} \in \mathfrak{R}^m \text{ e } \mathbf{z} \in \mathfrak{R}^k.$$

A projeção é ortogonal se as colunas da matriz $\mathbf{V}$ forem ortogonais e de norma unitária.

Para uma única direção de projeção, a matriz $\mathbf{V}$ se reduz a um vetor $\mathbf{v} \in \mathbb{R}^{m \times 1}$. Neste caso, o método de aproximação por busca de projeção tem por objetivo encontrar a direção de projeção $\mathbf{v}$ que resolva o seguinte problema:

$$\max_{\mathbf{v}} I_P(\mathbf{v}, \mathbf{X}, \mathbf{S}), \quad (4.3)$$

onde $I_P(\mathbf{v}, \mathbf{X}, \mathbf{S})$ é um índice de projeção, $\mathbf{X} \in \mathbb{R}^{m \times N}$ é a matriz de dados de entrada e $\mathbf{S} \in \mathbb{R}^{N \times r}$ é a matriz de dados de saída. Com $r = 1$, $\mathbf{X}$ e $\mathbf{S}$ são dados respectivamente por:

$$\mathbf{X} = [\mathbf{x}_1 \ \mathbf{x}_2 \ \cdots \ \mathbf{x}_N] \quad \text{e} \quad \mathbf{S} = [s_1 \ s_2 \ \cdots \ s_N]^T.$$

Três peculiaridades do processo de otimização de $I_P(\mathbf{v}, \mathbf{X}, \mathbf{S})$ contribuem para aumentar sua eficiência em termos de implementação:

- pelo fato de a operação de projeção realizar um tipo de suavização dos dados (HUBER, 1985), o resultado do processo de projeção não vai sofrer modificações acentuadas em face de pequenos ajustes na direção de projeção. Conclui-se então que não é fundamental a determinação precisa da direção de projeção ótima, embora continue a ser interessante a rapidez no processo de busca, ou seja, o índice $I_P(\mathbf{v}, \mathbf{X}, \mathbf{S})$ deve ser uma função de fácil avaliação;
- apesar de geralmente o problema de otimização apresentado na equação (4.3) não ser convexo, é possível definir índices de projeção e suas derivadas que variem continuamente com a direção de projeção $\mathbf{v}$, permitindo a aplicação dos mais eficientes algoritmos de otimização disponíveis com a garantia de bom comportamento numérico e convergência para ótimos locais. A característica local dos algoritmos de otimização a serem empregados não se manifesta forçosamente como uma limitação, pois existe a necessidade de se explorar boa parte dos ótimos locais, principalmente aqueles que não se encontram tão distantes da otimalidade global (JONES, 1987);
- por operar com *direções* de projeção, o índice deve ser tal que $I_P(\mathbf{v}, \mathbf{X}, \mathbf{S}) = I_P(-\mathbf{v}, \mathbf{X}, \mathbf{S})$.

É importante salientar que, neste estudo, as direções fornecidas por índices de projeção são tomadas como condições iniciais para algoritmos iterativos de construção de modelos de aproximação por busca de projeção, a serem apresentados mais adiante. Portanto, os resultados apresentados neste seção se referem à busca da direção *inicial* de projeção. $I_P(\mathbf{v}, \mathbf{X}, \mathbf{S})$ pode também ser empregado como critério de parada, pois ele é capaz de indicar quando todas as estruturas presentes nos dados já foram identificadas.

4.4.1 A motivação para a definição de um índice de projeção

O método original de aproximação por busca de projeção proposto por FRIEDMAN & TUKEY (1974) tinha um objetivo modesto. Seu desenvolvimento foi motivado pela necessidade de automatização de um dispositivo até então de ajuste manual que apresentava o gráfico da projeção de um conjunto de dados multivariáveis em uma determinada direção. A direção de projeção era ajustada manualmente e de forma contínua. A apresentação gráfica dos pontos projetados permitia simular o deslocamento destes pontos em conformidade com a variação imposta à direção de projeção. Desta forma, por inspeção visual, o usuário poderia determinar qual a direção de projeção que fornecia o gráfico mais *interessante*, um critério de natureza predominantemente subjetiva.

DIACONIS (1985) apresentou estimativas de número de projeções em função da dimensão do espaço de aproximação, demonstrando a impraticabilidade de se explorar exaustivamente, na forma descrita acima, as projeções em espaços de dimensão maior ou igual a três. Justifica-se, portanto, a implementação de um procedimento automatizado de busca da direção de projeção capaz de emular a estratégia utilizada pelo usuário e substituí-lo principalmente no caso de problemas de aproximação de dimensão elevada.

No entanto, não existe um consenso com relação às consequências da aplicação de índices numéricos e das transformações necessárias para permitir sua aplicação a procedimentos automáticos de busca de direções de projeção (HUBER, 1985). Além disso, é difícil evitar que diferentes conceitos do que venha a ser uma direção de projeção interessante sejam confundidos na definição de um determinado índice (JONES & SIBSON, 1987). Em vista disto, é muitas vezes preferível definir o índice de projeção dual, ou seja, expressando o que vem a ser uma direção *não-interessante* e otimizar no sentido da divergência deste conceito. Logicamente, este procedimento não garante que se obtenha necessariamente uma direção de projeção interessante, mas estas certamente não estão descartadas.

Em vista da dificuldade de modelagem de um único índice de projeção que satisfaça todas as expectativas, recorre-se a diversos índices de projeção, cada qual fornecendo uma solução específica. Em seguida, as soluções obtidas são comparadas e é selecionada aquela que conduz ao melhor resultado.

4.4.2 Índices de projeção baseados apenas nos dados de entrada

Os primeiros índices de projeção foram desenvolvidos com base em conceitos puramente heurísticos e considerando apenas os dados de entrada representados pela matriz $\mathbf{X} \in \Re^{m \times N}$, ou seja,

$$I_p(\mathbf{v}, \mathbf{X}, \mathbf{S}) = I_p(\mathbf{v}, \mathbf{X}).$$

Desde o primeiro tratamento formal, apresentado por HUBER (1985), os índices $I_p(\mathbf{v}, \mathbf{X})$ mais difundidos são aqueles que apresentam a propriedade de invariância afim, ou seja,

$$I_p(a\mathbf{v}+b, \mathbf{X}) = I_p(\mathbf{v}, \mathbf{X}), \text{ para } a, b \in \Re.$$

4.4.2.1 Índices derivados de análise multivariada

Inicialmente, é importante mencionar que, com base em métodos clássicos de análise multivariada, direções de projeção interessantes podem ser determinadas diretamente a partir dos dados de entrada, sem recorrer a processos de otimização de índices de projeção (ANDREWS *et al.*, 1971; ANDREWS, 1972). A desvantagem de métodos deste tipo é a inclusão de direções espúrias no conjunto de direções interessantes, principalmente com o aumento da dimensão m do espaço de entrada. Além disso, o pré-processamento adequado dos dados de entrada pode ser crucial para o desempenho deste método.

Outros índices de projeção baseados apenas nos dados de entrada também podem ser considerados como casos particulares de métodos clássicos de análise multivariada. Por exemplo, é possível tomar o componente principal da matriz de entrada como a direção de projeção. Para tanto, basta adotar (JOLLIFFE, 1986):

$$I_p(\mathbf{v}, \mathbf{X}) = \pm \mathbf{v}^T \left(\frac{1}{N} \mathbf{X} \mathbf{X}^T - \bar{\mathbf{x}} \bar{\mathbf{x}}^T \right) \mathbf{v},$$

onde $\bar{\mathbf{x}} = \frac{1}{N} \sum_{l=1}^N \mathbf{x}_l$ e resolver o seguinte problema:

$$\begin{aligned} & \max_{\mathbf{v}} I_p(\mathbf{v}, \mathbf{X}) \\ & \text{s. a. } \|\mathbf{v}\| = 1 \end{aligned}$$

Este índice de projeção expressa a proporção da variância total presente nos dados projetados na direção $\mathbf{v}$.

Neste contexto, pode-se detectar duas possíveis desvantagens do uso da análise de componentes principais na determinação da direção de projeção:

- o índice de projeção $I_P(\mathbf{v}, \mathbf{X})$ resultante não é invariante afim, ou seja, a solução depende do escalonamento das variáveis de entrada;
- é estabelecido um vínculo direto entre a direção que fornece a máxima variância e a direção que fornece o máximo possível de informação a respeito das associações existentes entre as variáveis de entrada. Este vínculo não é necessário e frequentemente não é verificado na prática (JONES & SIBSON, 1987).

4.4.2.2 Índices para busca exploratória

Apesar de não ser o objetivo deste estudo, a busca de projeção é também utilizada para detectar estruturas específicas presentes na associação entre as variáveis de entrada. O objetivo deste procedimento de busca exploratória não é propriamente o de aproximação, mas o de exploração do espaço de amostragem para detecção de comportamentos específicos na distribuição dos dados amostrados (FRIEDMAN, 1987).

Neste sentido, são requeridos índices de projeção que, além de serem funções contínuas e diferenciáveis da direção de projeção, expressem o nível de não-normalidade dos dados projetados. Com isso, maximizar este índice de projeção significa encontrar a direção de projeção que apresenta o maior grau de estruturação (não-normalidade) possível para os dados projetados. Uma solução para este problema foi apresentada por HALL (1989a), utilizando polinômios de Hermite.

A justificativa para a escolha deste tipo de índice de projeção para propósitos de busca exploratória é baseada no seguinte argumento heurístico: um conjunto de dados multivariáveis apresenta distribuição normal se e somente se todas as projeções unidimensionais destes dados apresentarem distribuição normal. Portanto, se a projeção unidimensional que produz a distribuição menos normal resulta em uma distribuição normal, então todas as outras projeções também resultarão em distribuições normais (HUBER, 1985).

No caso de problemas de busca de projeção em que os dados só podem assumir valores discretos, é geralmente mais adequado considerar índices de projeção que meçam não-uniformidade (DIACONIS, 1985). Neste sentido, certamente vão existir casos em que índices de projeção que meçam outros tipos de não-unimodalidade sejam mais adequados.

4.4.3 Índices de projeção completos

Os índices de projeção a serem apresentados nesta seção são denominados completos no sentido de utilizarem também a matriz de saída como informação necessária para obter a direção de projeção mais interessante, ou seja, são funções do tipo $I_P(\mathbf{v}, \mathbf{X}, \mathbf{S})$.

Além disso, é assumido que os dados projetados estão distribuídos normalmente ao longo da direção de projeção. Esta hipótese de distribuição normal é justificada com base nos resultados de DIACONIS & FRIEDMAN (1984), estabelecendo que as projeções unidirecionais de dados multidimensionais tendem a assumir uma distribuição normal. Esta tendência se acentua com o aumento da dimensionalidade dos dados, conforme demonstrado por HALL & LI (1993).

Com isso, índices que assumem distribuição normal dos dados projetados têm uma maior eficiência com o aumento da dimensionalidade dos dados, justamente quando é mais necessário automatizar a busca de projeção.

4.4.3.1 Regressão inversa (HSING & CARROLL, 1992; LI, 1991)

Índices de projeção baseados em regressão inversa realizam a inversão de papéis entre $\mathbf{X}$ e $\mathbf{S}$ ($\mathbf{X} \in \mathbb{R}^{m \times N}$ e $\mathbf{S} \in \mathbb{R}^{N \times 1}$). Um benefício imediato desta inversão de papéis é a possibilidade de resolver, em lugar de um problema de aproximação multivariável, m problemas de aproximação monovariáveis, com dados (s_l, x_{il}) ($i=1, \dots, m; l=1, \dots, N$), onde m é a dimensão do vetor de entrada. Com isso, resulta um método de fácil implementação, em que a estimativa da curva de regressão inversa não requer nenhum procedimento de aproximação mais sofisticado, embora splines suavizantes ou outros métodos de regularização possam vir a ser aplicados.

O método de regressão inversa opera junto aos dados (s_l, x_{il}) ($i=1, \dots, m; l=1, \dots, N$) da seguinte forma (LI, 1991):

1. ordene os dados em $\mathbf{S}$, de tal forma que $s_1 \leq s_2 \leq \dots \leq s_N$;
2. em seguida, agrupe os dados em conjuntos de dois;
3. faça uma estimativa $\hat{\Sigma}_{xx}$ da matriz de covariância Σ_{xx} dentro de cada conjunto;
4. extraia a média $\bar{\Sigma}_{xx}$, padronize em relação a Σ_{xx} e tome o autovetor $\mathbf{d}$ correspondente ao autovalor máximo (em valor absoluto) de $\bar{\Sigma}_{xx}$;
5. a direção de projeção é dada por $\mathbf{v} = (\Sigma_{xx})^{-1/2} \mathbf{d}$.

Portanto, o índice de projeção assume a forma:

$$I_P(\mathbf{v}, \mathbf{X}, \mathbf{S}) = \left(\Sigma_{\mathbf{xx}}^{1/2} \mathbf{v} \right)^T \left(\pm \bar{\Sigma}_{\mathbf{xx}} \right) \left(\Sigma_{\mathbf{xx}}^{1/2} \mathbf{v} \right)$$

Os passos 2 e 3 fornecem uma estimativa simplificada da curva de regressão inversa, enquanto que a direção $\mathbf{d}$ é denominada a direção mais efetiva para redução de dimensão, ou seja, é a direção mais interessante para a projeção dos dados (HSING & CARROLL, 1992). O passo 5 promove a reversão do processo de padronização.

No entanto, índices definidos a partir de regressão inversa não são capazes de fornecer nenhuma direção interessante para funções g simétricas em relação à média de $\mathbf{x}$ (LI, 1991). Neste caso, LI (1992) propõe como alternativa o método de decomposição espectral da média da hessiana, apresentado a seguir.

4.4.3.2 Decomposição espectral da média da hessiana (LI, 1992)

Dada a função $g(\mathbf{x})$ a ser aproximada, a matriz hessiana $\mathbf{H}_g(\mathbf{x})$ correspondente varia com $\mathbf{x}$ sempre que a função g não for quadrática. Seja $\bar{\mathbf{H}}_g$ a média de $\mathbf{H}_g(\mathbf{x})$ e

$$\Sigma_{\mathbf{x}} = \frac{1}{N} \sum_{l=1}^N (\mathbf{x}_l - \bar{\mathbf{x}})(\mathbf{x}_l - \bar{\mathbf{x}})^T, \text{ com } \bar{\mathbf{x}} = \frac{1}{N} \sum_{l=1}^N \mathbf{x}_l,$$

a matriz de covariância de $\mathbf{x}$. Os autovetores $\mathbf{d}_1, \dots, \mathbf{d}_m$ da matriz $\bar{\mathbf{H}}_g \Sigma_{\mathbf{x}}$ tais que

$$\bar{\mathbf{H}}_g \Sigma_{\mathbf{x}} \mathbf{d}_i = \lambda_i \mathbf{d}_i, \quad \mathbf{d}_i^T \Sigma_{\mathbf{x}} \mathbf{d}_i = 1, \quad i=1, \dots, m,$$

são denominados as direções principais da hessiana com relação à distribuição de $\mathbf{x}$. Estas são as direções ao longo das quais a curvatura média da função $g(\mathbf{x})$ é maior em termos de derivada de segunda ordem. Neste caso, a direção correspondente ao maior autovalor em valor absoluto é aquela de maior curvatura média.

Neste caso,

$$I_P(\mathbf{v}, \mathbf{X}, \mathbf{S}) = \mathbf{v}^T (\pm \bar{\mathbf{H}}_g \Sigma_{\mathbf{x}}) \mathbf{v}$$

é um índice de projeção que expressa a curvatura média da função $g(\mathbf{x})$ em termos de derivada de segunda ordem. Este índice é invariante afim e a direção de projeção mais interessante é então dada por:

$$\mathbf{v}^* = \arg \left\{ \max_{\mathbf{v}} \mathbf{v}^T (\pm \bar{\mathbf{H}}_g \Sigma_{\mathbf{x}}) \mathbf{v} \right\} \text{ s. a. } \|\mathbf{v}\| = 1$$

Note que outras matrizes poderiam ser utilizadas em lugar de $\Sigma_{\mathbf{x}}$ para incorporar outros critérios de importância relativa entre as variáveis de entrada.

Utilizando um resultado derivado do lema de Stein (STEIN, 1981), LI (1992) provou que a matriz $\bar{\mathbf{H}}_g$ está relacionada com a matriz de covariância ponderada

$$\Sigma_{\mathbf{sx}} = \frac{1}{N} \sum_{l=1}^N (s_l - \bar{s})(\mathbf{x}_l - \bar{\mathbf{x}})(\mathbf{x}_l - \bar{\mathbf{x}})^T, \text{ com } \bar{s} = \frac{1}{N} \sum_{l=1}^N s_l.$$

na forma:

$$\bar{\mathbf{H}}_g = \Sigma_{\mathbf{x}}^{-1} \Sigma_{\mathbf{sx}} \Sigma_{\mathbf{x}}^{-1}.$$

Então, os autovetores $\mathbf{d}_1, \dots, \mathbf{d}_m$ da matriz $\bar{\mathbf{H}}_g \Sigma_{\mathbf{x}}$ são obtidos a partir de:

$$\Sigma_{\mathbf{sx}} \mathbf{d}_i = \lambda_i \Sigma_{\mathbf{x}} \mathbf{d}_i, i=1, \dots, m.$$

Logo, a direção de projeção $\mathbf{v}^*$ que maximiza $I_P(\mathbf{v}, \mathbf{X}, \mathbf{S})$ é dada por:

$$\mathbf{v}^* = \mathbf{d}_{i^*}, \text{ onde } i^* = \arg \max_i |\lambda_i|.$$

4.4.3.3 Derivadas médias (HÄRDLE & STOKER, 1989; SAMAROV, 1993)

Em sua forma mais geral, o método das derivadas médias realiza estimativas da direção de projeção sem assumir nenhum tipo de distribuição para $\mathbf{x}$. No entanto, isto acarreta um custo computacional elevado e a sujeição à maldição da dimensionalidade (LI, 1992).

Por outro lado, ao assumir distribuição normal, o processo é bastante simplificado. Considere as seguintes definições:

- $\mathbf{I}_N$ é a matriz identidade de ordem N ;
- $\mathbf{1}_{N \times N}$ é a matriz $N \times N$ com todos os elementos iguais a 1;
- $\mathbf{X} = [\mathbf{x}_1 \ \mathbf{x}_2 \ \dots \ \mathbf{x}_N]$ é a matriz de dados de entrada;
- $\mathbf{S} = [s_1 \ s_2 \ \dots \ s_N]^T$ é o vetor de dados de saída;
- $\bar{s} = \frac{1}{N} \sum_{l=1}^N s_l$ é a média dos dados de saída;

- $\mathbf{H} = \mathbf{I}_N - \frac{1}{N} \mathbf{1}_{N \times N}$ é a matriz de centralização (MARDIA *et al.*, 1979).

Então, a matriz de covariância de $\mathbf{x}$ pode ser expressa na forma matricial como:

$$\Sigma_{\mathbf{x}} = \frac{1}{N} \mathbf{X} \mathbf{H} \mathbf{X}^T.$$

Tomando os elementos da diagonal de $\Sigma_{\mathbf{x}}$, é possível construir uma matriz diagonal na forma:

$$\mathbf{B} = \text{diag}(\Sigma_{\mathbf{x}}),$$

permitindo produzir a seguinte matriz:

$$\mathbf{D} = \mathbf{H} \mathbf{X}^T \mathbf{B}^{-1/2}.$$

Se $\mathbf{d}_l$ é o vetor-coluna formado pelos elementos da l -ésima linha de $\mathbf{D}$, então:

$$\mathbf{E} = \frac{1}{N} \sum_{l=1}^N s_l \mathbf{d}_l \mathbf{d}_l^T.$$

Finalmente, tomando

$$\mathbf{C} = \mathbf{B}^{-1/2} \Sigma_{\mathbf{x}} \mathbf{B}^{-1/2},$$

a direção de projeção dada pelo critério das derivadas médias é o autovetor correspondente ao autovalor máximo em valor absoluto da matriz

$$\mathbf{T} = \mathbf{C}^{-1} \mathbf{E} \mathbf{C}^{-1} - \mathbf{C}^{-1} \bar{s},$$

sendo que o correspondente índice de projeção pode ser obtido prontamente.

Em termos de simulação computacional e assumindo distribuição normal, o índice obtido a partir de derivadas médias produz resultados muito próximos daqueles fornecidos pelo índice obtido a partir da decomposição espectral da média da hessiana.

4.4.3.4 Índice de projeção para associações lineares

Assumindo associações lineares entre as variáveis de entrada, a direção de projeção ótima $\mathbf{v}^*$ é aquela que resolve o seguinte problema de otimização:

$$\min_{\mathbf{v}} \frac{1}{2} (\mathbf{X}^T \mathbf{v} - \mathbf{S})^T (\mathbf{X}^T \mathbf{v} - \mathbf{S}).$$

Portanto, tomando como índice de projeção $I_P(\mathbf{v}, \mathbf{X}, \mathbf{S}) = -\frac{1}{2} (\mathbf{X}^T \mathbf{v} - \mathbf{S})^T (\mathbf{X}^T \mathbf{v} - \mathbf{S})$, a solução ótima $\mathbf{v}^*$ assume a forma:

$$\mathbf{v}^* = (\mathbf{X}\mathbf{X}^T)^{-1} \mathbf{X}\mathbf{S}.$$

A inversão de $\mathbf{X}\mathbf{X}^T$ sempre vai existir se a matriz de dados de entrada $\mathbf{X}$ tiver posto completo, uma condição que deve ser naturalmente atendida para que os dados sejam realmente representativos (WAHBA, 1990).

4.4.4 Casos particulares

Contrastando com a busca de uma única direção de projeção a cada iteração, em alguns casos não é necessário realizar esta busca, enquanto que em outros casos é preciso realizar buscas de múltiplas direções em partições do espaço de aproximação.

4.4.4.1 Tomografia Computadorizada

Uma das ramificações do método de busca de projeção de FRIEDMAN & TUKEY (1974) é a tomografia computadorizada (SHEPP & KRUSKAL, 1978). O objetivo continua sendo a reconstrução de estruturas de dimensões elevadas a partir de projeções em espaços de menor dimensão. As principais diferenças são:

- o conjunto de informações para aproximação se restringem às projeções unidirecionais, ou seja, não se tem acesso à informação na forma de amostras no espaço de dimensão completa;
- ausência de busca das direções de projeção: as direções são obtidas por convolução entre direções igualmente espaçadas e filtros lineares;
- aplicação restrita a estruturas tridimensionais.

4.4.4.2 Partição do vetor de entrada

Além da possibilidade de utilizar-se projeções multidimensionais, uma alternativa ao uso de uma única direção de projeção para cada iteração é realizar a partição do vetor de entrada em q vetores na forma:

$$\mathbf{x} = [\mathbf{x}_1 \vdots \mathbf{x}_2 \vdots \dots \vdots \mathbf{x}_q].$$

Com uma partição análoga para a direção de projeção, o processo de aproximação assume a forma:

$$\hat{g}_n(\mathbf{x}) = \sum_{j=1}^n \hat{f}_j(\mathbf{v}_{1j}^T \mathbf{x}_1, \mathbf{v}_{2j}^T \mathbf{x}_2, \dots, \mathbf{v}_{qj}^T \mathbf{x}_q).$$

Este procedimento foi proposto por BOOKSTEIN (1985) e é apropriado para o caso de múltiplas associações disjuntas entre as variáveis de entrada, as quais podem ser encontradas por técnicas de regressão alternada (LYTTKENS, 1972).

4.5 Determinação da função de expansão ortogonal

Uma vez definida a direção de projeção $\mathbf{v}_n \in \mathfrak{R}^m$, o problema de aproximação regularizado apresentado na equação (4.2) tem por objetivo aproximar uma versão suave da projeção dos resíduos da função desconhecida $g: \mathfrak{R}^m \rightarrow \mathfrak{R}$ na direção $\mathbf{v}_n$. O fator fundamental que continua caracterizando todo o processo de implementação é a tentativa de aproximar a função em regiões onde não se dispõe de informação suficiente para implementar um processo totalmente não-paramétrico.

Para $\mathbf{v}_n$ ($n \geq 1$) fixo, é possível renomear as projeções unidirecionais dos dados de entrada na forma:

$$z_l = \mathbf{v}_n^T \mathbf{x}_l, l=1, \dots, N. \quad (4.4)$$

Com isso, a função monovariável $f_n(\cdot)$ deve resolver o seguinte problema de aproximação regularizado:

$$\min_{f_n} \frac{1}{N} \sum_{l=1}^N [d_l - f_n(z_l)]^2 + \lambda_n \phi(f_n), \quad (4.5)$$

onde d_l são os resíduos do processo de aproximação dados por:

$$d_l = s_l - \sum_{j=1}^{n-1} f_j(\mathbf{v}_j^T \mathbf{x}_l), l=1, \dots, N.$$

Mesmo que o tipo de suavidade da função g seja compatível com aquele imposto pela função de regularização $\phi(\cdot)$ às funções $f_j(\cdot)$ ($j=1, \dots, n-1$), o comportamento dos pontos $\{(z_l, d_l)\}_{l=1}^N$ pode ser bastante errático devido a variação de $g(\mathbf{x}) - \sum_{j=1}^{n-1} f_j(\mathbf{v}_j^T \mathbf{x})$ em outras direções que não $\mathbf{v}_n$. Com isso, o valor ótimo do parâmetro de regularização λ_n não pode ser determinado a priori, sendo função do conjunto de dados projetados $\{(z_l, d_l)\}_{l=1}^N$.

Embora outras técnicas de determinação da função de expansão ortogonal tenham sido apresentadas na literatura, como, por exemplo, funções lineares por partes (FRIEDMAN & STUETZLE, 1981) e método da vizinhança mais próxima (HALL, 1989b), este estudo vai se restringir ao emprego de bases de funções ortonormais (solução paramétrica) e splines polinomiais suavizantes (solução não-paramétrica).

4.5.1 Solução paramétrica

Como uma solução paramétrica para o problema de aproximação dos dados $\{(z_l, d_l)\}_{l=1}^N$, é possível utilizar uma base de funções ortonormais de dimensão $P+1$ (P finito), dada por $\{h_0, ..., h_P\}$. Tomando $\{h_0, ..., h_P\}$ como as funções de Hermite apresentadas no capítulo 2, o problema de aproximação paramétrica se restringe a encontrar um vetor de coeficientes $\mathbf{c} = [c_0 \dots c_P]^T$ tal que f_j seja dado pela seguinte combinação linear:

$$f_j(z_l) = \sum_{i=0}^P c_i h_i(z_l).$$

Definindo

$$\mathbf{H} = \begin{bmatrix} h_0(z_1) & h_1(z_1) & \dots & h_P(z_1) \\ h_0(z_2) & h_1(z_2) & \dots & h_P(z_2) \\ \vdots & \vdots & \ddots & \vdots \\ h_0(z_N) & h_1(z_N) & \dots & h_P(z_N) \end{bmatrix}, \quad \mathbf{c} = \begin{bmatrix} c_0 \\ c_1 \\ \vdots \\ c_P \end{bmatrix} \quad \text{e} \quad \mathbf{y} = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_N \end{bmatrix},$$

a solução ótima $\mathbf{c}^*$ no sentido dos quadrados mínimos para o problema de otimização

$$\min_{\mathbf{c}} \frac{1}{2} (\mathbf{y} - \mathbf{H}\mathbf{c})^T (\mathbf{y} - \mathbf{H}\mathbf{c})$$

produz, assumindo $P < N$,

$$\mathbf{c}^* = (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \mathbf{y}.$$

Observe que sempre é possível encontrar $P \geq 0$, ou seja, um número de colunas de $\mathbf{H}$ que forneça um número de condição (BRONSON, 1989) para $\mathbf{H}^T \mathbf{H}$ que garanta sua inversão sem a ocorrência de dificuldades em termos numéricos. Problemas numéricos poderiam surgir para valores elevados de P , pois sendo N finito e z_l ($l=1, ..., N$) restrito a um intervalo limitado, a ortonormalidade entre as colunas de $\mathbf{H}$ deixa de ser satisfeita. Neste caso, pode-se considerar $\{h_0, ..., h_P\}$ como um subconjunto de funções ortonormais com componentes extraídos de um conjunto de maior dimensão, onde as $P+1$ funções escolhidas são aquelas que conduzem à melhor solução para o problema de aproximação paramétrico.

Dado o procedimento de geração da base de funções ortonormais, procedimentos para a definição de um valor ótimo para P podem estar baseados em validação cruzada (CRAVEN & WAHBA, 1979) ou em procedimentos análogos, como apresentado por LAY *et al.* (1994).

4.5.2 Solução não-paramétrica

Conforme observado no capítulo 3, pelo fato de N ser finito, os melhores modelos de aproximação são aqueles capazes de conciliar o nível de dependência da dimensionalidade com a flexibilidade de aproximação, resultando modelos que devem atender a restrições de suavidade, não sendo, portanto, totalmente não-paramétricos.

Definindo

$$R(\lambda_n) = \frac{1}{N} \sum_{l=1}^N \left(g(z_l) - f_{n,\lambda_n}(z_l) \right)^2 \quad (4.6)$$

onde $f_{n,\lambda_n}(\cdot)$ é a solução do problema (4.5) para um dado λ_n , o valor ótimo λ_n^* é tal que:

$$\lambda_n^* = \arg \min_{\lambda_n} R(\lambda_n). \quad (4.7)$$

O problema (4.7) não pode ser diretamente resolvido, pois é apresentado em função de g , a função a ser aproximada e, portanto, desconhecida. No entanto, boas estimativas para λ_n^* podem ser obtidas utilizando métodos de validação cruzada (CRAVEN & WAHBA, 1979), conforme descrito mais adiante.

4.5.2.1 Splines polinomiais suavizantes

Sejam z_l ($l=1, \dots, N$) as projeções da variável da entrada definidas na equação (4.4).

Assumindo

$$z_l \in [a, b], \quad l=1, \dots, N,$$

e tomando $f_n \in L_2^2[a, b]$, o espaço de Sobolev de ordem 2 associado ao espaço de Lebesgue $L_2[a, b]$ (veja apêndice B), então, ao definir a função de regularização na forma:

$$\Phi(f_n) = \int_a^b \left(\frac{\partial^2 f_n(u)}{\partial u^2} \right)^2 du,$$

a função $f_{n,\lambda_n}(\cdot)$, definida para um dado λ_n e tal que:

$$f_{n,\lambda_n} = \arg \min_{f_n} \frac{1}{N} \sum_{l=1}^N [d_l - f_n(z_l)]^2 + \lambda_n \int_a^b \left(\frac{\partial^2 f_n(u)}{\partial u^2} \right)^2 du, \quad (4.8)$$

é composta por splines polinomiais suavizantes de ordem 3 (veja capítulo 2 e REINSCH (1967), REINSCH (1971), SCHOENBERG (1964) e WAHBA (1975)).

Os splines polinomiais suavizantes podem ser considerados generalizações de filtros passa-baixa, tendo os dados de aproximação $\{(z_l, d_l)\}_{l=1}^N$ como entrada, já que a função de regularização $\phi(f_n)$ penaliza componentes de alta frequência em f_n (WAHBA, 1975).

4.5.2.2 Validação cruzada ordinária

Um método prático e efetivo de estimar o parâmetro de regularização λ_n^* , solução do problema (4.7), é a validação cruzada generalizada. Nesta seção, é apresentado o método de validação cruzada ordinária, ponto de partida para se chegar ao método de validação cruzada generalizada, a ser discutido na próxima seção.

Como já evidenciado na capítulo 3, o parâmetro de regularização λ_n representa um compromisso entre a fidelidade aos dados, medida pela expressão

$$\frac{1}{N} \sum_{l=1}^N [d_l - f_n(z_l)]^2,$$

e o nível de suavidade de f_n em $[a, b]$, avaliado na forma:

$$\int_a^b \left(\frac{\partial^2 f_n(u)}{\partial u^2} \right)^2 du.$$

A estimativa do valor ótimo λ_n^* tem soluções mais simples nos seguintes casos:

- quando a variância dos dados de aproximação $\{(z_l, d_l)\}_{l=1}^N$ é conhecida (REINSCH, 1967);
- quando os pontos z_l ($l=1, \dots, N$) são igualmente espaçados (WAHBA, 1975).

Como o problema (4.7) em estudo é assumido corresponder a casos mais gerais, é necessário buscar soluções mais elaboradas.

Para um dado λ_n , seja $f_{n, \lambda_n}^{[k]}(\cdot)$ a função composta por splines polinomiais suavizantes que resolvem o seguinte problema:

$$\min_{f_n} \frac{1}{N} \sum_{\substack{l=1 \\ l \neq k}}^N [d_l - f_n(z_l)]^2 + \lambda_n \int_a^b \left(\frac{\partial^2 f_n(u)}{\partial u^2} \right)^2 du,$$

onde são considerados todos os pontos do conjunto $\{(z_l, d_l)\}_{l=1}^N$, exceto (z_k, d_k) , para $1 \leq k \leq N$.

Observe que $f_{n, \lambda_n}^{[k]}(\cdot) \equiv f_{n, \lambda_n}(\cdot)$ quando d_k é substituído por $f_{n, \lambda_n}^{[k]}(z_k)$.

Considere, então, a função de validação cruzada ordinária

$$V_0(\lambda_n) = \frac{1}{N} \sum_{k=1}^N \left[d_k - f_{n,\lambda_n}^{[k]}(z_k) \right]^2. \quad (4.9)$$

O método de validação cruzada ordinária, desenvolvido por WAHBA & WOLD (1975a, 1975b), permite encontrar $\hat{\lambda}_n^*$ que resolve o problema

$$\min_{\lambda_n} V_0(\lambda_n). \quad (4.10)$$

Em WAHBA & WOLD (1975a), $\hat{\lambda}_n^*$ foi mostrado ser uma boa estimativa para λ_n^* , solução do problema (4.7). No entanto, com V_0 dado pela equação (4.9), a obtenção de $\hat{\lambda}_n^*$ sem assumir periodicidade de g e equidistância dos dados fica bastante dificultada (WAHBA, 1990).

4.5.2.3 Validação cruzada generalizada

Considere a matriz $\mathbf{A}(\lambda_n)$ tal que

$$\begin{bmatrix} f_{n,\lambda_n}(z_1) \\ \vdots \\ f_{n,\lambda_n}(z_N) \end{bmatrix} = \mathbf{A}(\lambda_n) \begin{bmatrix} d_1 \\ \vdots \\ d_N \end{bmatrix}. \quad (4.11)$$

Como $\mathbf{A}(\lambda_n)$ é uma matriz simétrica (CRAVEN & WAHBA, 1979), ela sempre pode ser transformada em uma matriz circulante $\tilde{\mathbf{A}}(\lambda_n)$ (BRONSON, 1989) por rotação do sistema de coordenadas (CRAVEN & WAHBA, 1979), conduzindo à função de validação cruzada generalizada

$$V(\lambda_n) = \frac{1}{N} \frac{\|[\mathbf{I} - \tilde{\mathbf{A}}(\lambda_n)]\mathbf{d}\|^2}{\left(\frac{1}{N} \text{Tr}\{\mathbf{I} - \tilde{\mathbf{A}}(\lambda_n)\} \right)^2}, \quad (4.12)$$

onde $\mathbf{d} = [d_1 \dots d_N]^T$.

O método de validação cruzada generalizada, desenvolvido por CRAVEN & WAHBA (1979), permite encontrar $\tilde{\lambda}_n^*$ que resolve o problema

$$\min_{\lambda_n} V(\lambda_n). \quad (4.13)$$

Em CRAVEN & WAHBA (1979), $\tilde{\lambda}_n^*$ foi mostrado ser uma boa estimativa para λ_n^* , solução do problema (4.7). Além disso, $\tilde{\lambda}_n^*$, ao contrário de $\hat{\lambda}_n^*$ (solução do problema (4.10)), pode ser obtido de forma bastante eficiente, como demonstrado a seguir.

4.5.2.4 Obtenção do parâmetro de regularização

Dos resultados de REINSCH (1967) é sabido que a matriz $\tilde{\mathbf{A}}(\lambda_n)$ pode ser expressa na forma:

$$\tilde{\mathbf{A}}(\lambda_n) = \mathbf{I} - \mathbf{Q}(\mathbf{Q}^T \mathbf{Q} + p\mathbf{T})^{-1} \mathbf{Q}^T,$$

onde

- $p = \frac{1}{N\lambda_n}$;
- $\mathbf{Q}$ é a matriz tridiagonal de dimensão $N \times (N-2)$ com elementos q_{ij} , $i=1, \dots, N$, $j=1, \dots, (N-2)$, dados por:

$$q_{i,i+1} = \frac{1}{\Delta z_{i+1}}, \quad q_{ii} = -\frac{1}{\Delta z_i} - \frac{1}{\Delta z_{i+1}}, \quad q_{i+1,i} = \frac{1}{\Delta z_{i+1}}$$

onde $\Delta z_i = z_{i+1} - z_i$;

- $\mathbf{T}$ é a matriz tridiagonal de dimensão $(N-2) \times (N-2)$ com elementos t_{ij} , $i,j=1, \dots, (N-2)$, dados por:

$$t_{ii} = \frac{2}{3}(\Delta z_i + \Delta z_{i+1}), \quad t_{i+1,i} = t_{i,i+1} = \frac{\Delta z_{i+1}}{3},$$

onde $\Delta z_i = z_{i+1} - z_i$.

Para $\Delta z_i > 0$, a matriz $\mathbf{T}$ é estritamente definida positiva, de tal forma que sempre é possível obter

$$\mathbf{F} = \mathbf{Q}\mathbf{T}^{-1/2}.$$

É possível definir matrizes ortogonais $\mathbf{B}$, de dimensão $N \times (N-2)$, e $\mathbf{D}$, de dimensão $(N-2) \times (N-2)$, de modo a produzir a decomposição de $\mathbf{F}$ em valores singulares na forma (GOLUB & REINSCH, 1970):

$$\mathbf{F} = \mathbf{B}\mathbf{C}\mathbf{D}^T,$$

onde $\mathbf{C}$ é uma matriz diagonal contendo os valores singulares de $\mathbf{F}$, denominados $c_1, c_2, \dots, c_{N-2}$. Observe que todos os valores singulares de $\mathbf{F}$ são não-nulos (CRAVEN & WAHBA, 1979).

Com isso, é possível expressar $\tilde{\mathbf{A}}(\lambda_n)$ na forma:

$$\tilde{\mathbf{A}}(\lambda_n) = \mathbf{I} - \mathbf{B} \begin{bmatrix} \frac{c_1^2}{c_1^2 + p} & & 0 \\ & \ddots & \\ 0 & & \frac{c_{N-2}^2}{c_{N-2}^2 + p} \end{bmatrix} \mathbf{B}^T.$$

Fazendo

$$\begin{bmatrix} e_1 \\ \vdots \\ e_{N-2} \end{bmatrix} = \mathbf{B}^T \mathbf{d}$$

onde $\mathbf{d} = [d_1 \dots d_N]^T$, $V(\lambda_n)$ dado pela equação (4.12) pode ser reescrito como segue:

$$V(p) = \frac{1}{N} \frac{\sum_{j=1}^{N-2} \left(\frac{c_j^2}{c_j^2 + p} \right)^2 e_j^2}{\left[\frac{1}{N} \sum_{j=1}^{N-2} \left(\frac{c_j^2}{c_j^2 + p} \right) \right]^2}, \quad (4.14)$$

permitindo aplicar os mais eficientes métodos de otimização para resolver

$$\min_p V(p), \quad (4.15)$$

cujas solução p^* permite obter $\tilde{\lambda}_n^*$ na forma:

$$\tilde{\lambda}_n^* = \frac{1}{Np^*}. \quad (4.16)$$

4.6 O processo de ajuste retroativo

O processo de aproximação por expansão ortogonal aditiva segue, então, os seguintes passos básicos:

- dado o conjunto de dados de aproximação $\{(\mathbf{x}_l, s_l) \in X \times \mathfrak{R}\}_{l=1}^N$, com $X \subset \mathfrak{R}^m$, e seja $g: X \subset \mathfrak{R}^m \rightarrow \mathfrak{R}$ a função a ser aproximada;
- partindo de $n = 1$, e construindo o vetor $\mathbf{d} = [d_1 \dots d_N]^T$ na forma:

$$d_l = s_l - \sum_{j=1}^{n-1} f_j(\mathbf{v}_j^T \mathbf{x}_l), \quad l=1, \dots, N,$$

resolva o seguinte problema de otimização:

$$\min_{\mathbf{v}_n, f_n} \frac{1}{N} \sum_{l=1}^N \left(d_l - f_n(\mathbf{v}_n^T \mathbf{x}_l) \right)^2 + \lambda_n \phi(f_n),$$

pela obtenção sucessiva de valores ótimos para f_n e $\mathbf{v}_n$ na forma:

1. defina um valor inicial para $\mathbf{v}_n$ (empregando os resultados da seção 4.4);
2. para $\mathbf{v}_n$ fixo, resolva o seguinte problema de otimização (empregando os resultados da seção 4.5):

$$\min_{f_n} \frac{1}{N} \sum_{l=1}^N \left(d_l - f_n(\mathbf{v}_n^T \mathbf{x}_l) \right)^2 + \lambda_n \phi(f_n);$$

3. para f_n fixo, partindo do valor atual de $\mathbf{v}_n$, resolva iterativamente o seguinte problema de otimização (este é um problema de otimização paramétrica que pode ser resolvido empregando-se o algoritmo de Newton modificado, apresentado no Apêndice D):

$$\min_{\mathbf{v}_n} \frac{1}{N} \sum_{l=1}^N \left(d_l - f_n(\mathbf{v}_n^T \mathbf{x}_l) \right)^2;$$

4. enquanto não houver convergência (medida por algum critério de parada), retorne ao item 2.
- Faça $n = n+1$ e repita o processo a partir do cálculo dos novos valores para o vetor de resíduos $\mathbf{d}$ enquanto o nível de aproximação desejado ainda não foi atingido (medido por algum critério de parada).

Deste processo de aproximação resulta, então, um modelo de aproximação por composição aditiva de funções de expansão ortogonal, na forma:

$$g(\mathbf{x}) \approx \hat{g}_n(\mathbf{x}) = \sum_{j=1}^n f_j(\mathbf{v}_j^T \mathbf{x}).$$

No entanto, o processo de construção deste modelo de aproximação – descrito acima – apresenta uma limitação advinda da estratégia de aproximação empregada, a qual é descrita a seguir:

- 1) com base nos dados de aproximação referentes ao problema de aproximação original;
- 2) encontre uma única direção de projeção e uma única função de expansão ortogonal a esta direção (sujeita a restrições de suavidade) que melhor aproxime os dados;
- 3) remova do conjunto de dados a informação representada no item 2);
- 4) enquanto o nível de aproximação desejado ainda não foi atingido, retorne ao item 2).

Observe que cada um dos n termos da composição aditiva resultou de um processo de aproximação que tinha por objetivo representar *toda* a informação presente nos dados de aproximação e que *ainda não tinham sido representadas* pelos termos anteriores. Isto implica que cada novo termo da composição aditiva não leva em conta a possibilidade de que, posteriormente, *novos* termos possam vir a compor o processo de aproximação.

Tomando qualquer termo da composição aditiva, com exceção do n -ésimo termo, e denominando-o k ($1 \leq k < n$), surge a seguinte questão: o que ocorreria com f_k e $\mathbf{v}_k$ se, na solução do problema:

$$\min_{\mathbf{v}_k, f_k} \frac{1}{N} \sum_{l=1}^N \left(d_l - f_k(\mathbf{v}_k^T \mathbf{x}_l) \right)^2 + \lambda_k \phi(f_k),$$

em lugar de $\mathbf{d} = [d_1 \dots d_N]^T$ tal que

$$d_l = s_l - \sum_{j=1}^{k-1} f_j(\mathbf{v}_j^T \mathbf{x}_l), \quad l=1, \dots, N,$$

se tomasse

$$d_l = s_l - \sum_{\substack{j=1 \\ j \neq k}}^n f_j(\mathbf{v}_j^T \mathbf{x}_l), \quad l=1, \dots, N?$$

Se as duas escolhas para o vetor de resíduos $\mathbf{d}$ produzirem dados distintos, então f_k e $\mathbf{v}_k$ podem ser (e geralmente são) diferentes em cada caso. Conclui-se, portanto, que a solução produzida pelo processo de aproximação descrito acima pode deixar de ser ótima para os termos já calculados sempre que um termo adicional for incorporado. Sendo assim, é recomendável a aplicação de um processo de ajuste retroativo (*backfitting*):

- para cada j ($1 \leq j \leq n$), omite-se $f_j(\mathbf{v}_j^T \mathbf{x})$ do somatório e determina-se novos valores ótimos para f_j e $\mathbf{v}_j$. Repita este processo de ajuste retroativo até convergência, medida por algum critério de parada.

A demonstração de convergência do processo de ajuste retroativo foi apresentada por BREIMAN & FRIEDMAN (1985). Vale salientar também que o processo de ajuste retroativo foi originalmente proposto para reajustar apenas a função f_j , mantendo-se fixa a direção $\mathbf{v}_j$ (FRIEDMAN & STUETZLE, 1981), sendo que a estratégia iterativa do método é totalmente análoga àquela empregada por métodos numéricos de solução de sistemas de equações lineares, como o método de Gauss-Seidel (BUJA & STUETZLE, 1985). Além disso, o processo de ajuste retroativo é indispensável na implementação de métodos de poda.

4.7 Extensões da aproximação por expansão ortogonal

4.7.1 Expansão ortogonal por composição multiplicativa

Em termos de redução de dimensionalidade, DONOHO & JOHNSTONE (1989) demonstraram, para espaços de aproximação bidimensionais, que as funções de expansão ortogonal aditivas nem sempre são as mais eficazes, sendo em alguns casos superadas pelas funções de expansão ortogonal multiplicativas. Além disso, métodos de aproximação utilizando funções de expansão ortogonal aditivas e multiplicativas são complementares no sentido de que quando o uso de um tipo de função permite uma redução de dimensão, o uso do outro tipo não a permite. ZHAO & ATKESON (1992) estenderam estes resultados para espaços de aproximação multidimensionais.

A adaptação dos resultados obtidos neste capítulo para o caso de expansão ortogonal multiplicativa vai requerer as seguintes modificações básicas (FRIEDMAN, 1987):

- o modelo de aproximação vai assumir a forma: $g(\mathbf{x}) \approx \hat{g}_n(\mathbf{x}) = \prod_{j=1}^n f_j(\mathbf{v}_j^T \mathbf{x})$
- por ocasião do ajuste do k -ésimo termo da composição, o cálculo do vetor de resíduos $\mathbf{d} = [d_1 \dots d_N]^T$ é dado por ($l=1, \dots, N$):

$$- d_l = \frac{s_l}{\prod_{j=1}^{k-1} f_j(\mathbf{v}_j^T \mathbf{x}_l)}, \text{ para a adição de novos termos à composição;}$$

$$- d_l = \frac{s_l}{\prod_{\substack{j=1 \\ j \neq k}}^n f_j(\mathbf{v}_j^T \mathbf{x}_l)}, \text{ para o processo de ajuste retroativo.}$$

Motivado pelos resultados de DONOHO & JOHNSTONE (1989), processos de aproximação que procurem conciliar composição aditiva e multiplicativa podem conduzir a resultados bastante promissores. No entanto, este processo de conciliação ainda carece de embasamento teórico e certamente deve considerar uma série de fatores vinculados aos problemas de implementação, não discutidos neste estudo.

4.7.2 Expansão ortogonal a múltiplas direções

As funções de expansão ortogonal a uma determinada direção, tanto para o caso aditivo como multiplicativo, podem ser generalizadas para operarem com expansões

ortogonais a múltiplas direções. ou seja, ao invés de a expansão se dar em um hiperplano do espaço de aproximação, ela ocorre em espaços de menor dimensão. O tratamento de múltiplas direções de expansão ortogonal não será detalhado neste estudo pelos seguintes motivos:

- inerente aumento de complexidade do processo de aproximação e de carga computacional com a utilização simultânea de múltiplas direções. Por exemplo, se forem k as direções de projeção, o número de parâmetros aumenta k vezes em relação ao caso unidirecional;
- impossibilidade de garantir um aumento correspondente do poder de aproximação (FRIEDMAN, 1987);
- sendo k as direções de projeção, o método fornece apenas um subespaço de dimensão k , quando é geralmente interessante que as k direções estivessem ordenadas de acordo com sua contribuição no processo de aproximação neste espaço (HUBER, 1985). Note que este item perde significado no caso de direções ortogonais, mas as direções ótimas geralmente não atendem a este critério.

Apesar de funções descontínuas não serem consideradas neste estudo, deve-se salientar que, no caso de aproximação de funções descontínuas, projeções multidirecionais podem revelar estruturas que não são observáveis em projeções unidirecionais. Por exemplo, é muito difícil perceber a presença de "buracos" (regiões fechadas onde os dados estão ausentes) utilizando-se projeções unidirecionais.

4.7.3 O tratamento de problemas mais complexos

Uma consequência do tipo de implementação do processo de aproximação por busca de projeção é a possibilidade de inserção do algoritmo desenvolvido neste estudo como parte de algoritmos iterativos mais complexos, como o caso de aproximação envolvendo múltiplas variáveis de saída, a ser tratado no capítulo 5, e o caso de aproximação híbrida, onde é assumido que parte do modelo é adequadamente representado por estruturas inerentemente paramétricas (WAHBA, 1985).

Métodos de aproximação por busca de projeção apresentam-se, portanto, como ferramentas apropriadas para estimular o surgimento de novas idéias e para propiciar novas perspectivas de análise de outros métodos já existentes.

Parte II

Modelos de aproximação baseados em redes neurais artificiais

Capítulo 5

Redes Neurais como modelos de aproximação de funções

Redes neurais artificiais multicamadas são consideradas poderosas estruturas de aproximação de funções, produzindo modelos de aproximação paramétricos, equivalentes aos tradicionais modelos de regressão em estatística, ou não-paramétricos, produtos de processos de aproximação construtivos ou de poda. Na verdade, quase todos os modelos paramétricos e não-paramétricos presentes na literatura são passíveis de implementação na forma de uma rede neural artificial (CHENG & TITTERINGTON, 1994), conduzindo, muitas vezes, a novas perspectivas de análise e síntese.

A rede neural representa um modelo paramétrico sempre que sua arquitetura (estrutura de conexão e número de neurônios) for definida previamente, independente do problema de aproximação. Por outro lado, se a arquitetura da rede neural puder ser definida em função do problema de aproximação, ela representa um modelo não-paramétrico. Sendo assim, é possível, apesar de ser pouco provável, que o resultado da aplicação de dois processos de aproximação, um paramétrico e o outro não-paramétrico, produza uma mesma rede neural.

Mesmo no caso de definição prévia da arquitetura da rede neural, existem casos em que não se pode prever a forma final da solução produzida pelo modelo de aproximação paramétrico. Esta flexibilidade do modelo está invariavelmente associada à presença de um grande número de parâmetros e nunca a uma suposta natureza não-paramétrica do modelo. Apesar de muitos autores insistirem em classificar erroneamente redes neurais deste tipo como modelos não-paramétricos, não se pode confundir estruturas ricamente parametrizadas com estruturas não-paramétricas.

Este capítulo se ocupa especificamente do estudo de dois tipos de redes neurais com uma camada intermediária: uma paramétrica e outra não-paramétrica. É mostrado que ambas apresentam capacidade de aproximação universal, ou seja, é sempre possível obter uma boa aproximação utilizando qualquer um dos dois tipos de redes neurais. No entanto, apenas a rede

neural não-paramétrica pode se adaptar no sentido de produzir modelos que representam a melhor aproximação, também denominada de arquitetura ótima para o problema de aproximação considerado.

Além de não corresponderem ao caso de melhor aproximação, existe ainda o agravante de que a obtenção de boas aproximações no caso de redes neurais paramétricas geralmente exige a implementação de processos exaustivos de busca das arquiteturas mais adequadas, ou então a aplicação de métodos mais eficientes baseados em procedimentos de computação evolutiva (BALAKRISHNAN & HONAVAR, 1995; MILLER *et al.*, 1989). Como será visto mais adiante, computação evolutiva também pode ser utilizada na implementação de modelos não-paramétricos.

Já no caso de redes neurais não-paramétricas, dois procedimentos são utilizados: métodos construtivos e métodos de poda. Comparando-se estes dois métodos aplicados a estruturas não-paramétricas, verifica-se que os métodos puramente construtivos apresentam inúmeras vantagens sobre os métodos puramente de poda, mas ambos são superados por métodos híbridos predominantemente construtivos, envolvendo etapas de construção e poda intercaladas.

É discutida também a relação entre as redes neurais paramétricas e os resultados do teorema de Stone-Weierstrass, e também entre as redes neurais não-paramétricas e uma combinação dos resultados do teorema de Stone-Weierstrass e do teorema de Kolmogorov e suas generalizações (veja capítulo 2).

5.1 Capacidade de aproximação e propriedades de convergência

A importância e flexibilidade de classes de modelos de aproximação podem ser medidas pela sua capacidade de aproximação e pelas propriedades de convergência que estes modelos apresentam quando empregados na aproximação de funções pertencentes a subespaços funcionais específicos. Dada a função $g: X \subset \mathbb{R}^m \rightarrow \mathbb{R}'$ a ser aproximada, surgem duas questões fundamentais:

1. a família de funções de aproximação é consistente, ou seja, é suficientemente ampla para conter g ou, pelo menos, uma boa aproximação de g ?
2. o processo de aproximação é capaz de produzir uma sequência de funções de aproximação que, no limite, aproxime g tão bem quanto desejado?

Famílias de modelos de aproximação que permitem uma resposta afirmativa para a primeira questão, com g sendo qualquer função em um subespaço funcional específico, possuem capacidade de aproximação universal neste subespaço. A capacidade de aproximação universal é, na verdade, um pré-requisito para as classes de funções que permitem uma resposta afirmativa para a segunda questão.

Com relação à segunda questão, modelos de aproximação podem aproximar g tão bem quanto desejado sempre que for possível definir uma taxa de convergência (limitante inferior para a velocidade de decaimento do erro de aproximação) para o processo de aproximação.

Neste capítulo, no decorrer da apresentação de redes neurais paramétricas e não-paramétricas a serem utilizadas como modelos de aproximação, resultados da teoria de análise funcional são utilizados para mostrar que estas redes neurais produzem famílias de funções que respondem afirmativamente às duas questões enunciadas acima, pelo menos para $g \in C[X]$, onde $C[X]$ é o subespaço das funções contínuas em X .

5.2 Redes neurais artificiais paramétricas

Considere um neurônio artificial (veja seção 1.1) com função de ativação $f: \mathfrak{R} \rightarrow \mathfrak{R}$ e um vetor de parâmetros $\mathbf{p} = [p_0 \ p_1 \ \cdots \ p_m]^T$. Este neurônio recebe como sinal de entrada um vetor $\mathbf{x} \in \mathfrak{R}^m$, com $\mathbf{x} = [x_1 \ \cdots \ x_m]^T$, e produz um sinal de ativação $y \in \mathfrak{R}$, conforme apresentado na figura 5.1 e de acordo com a equação:

$$y = f(\mathbf{p}, \mathbf{x}) = f\left(\sum_{i=1}^m p_i x_i + p_0\right). \quad (5.1)$$

A entrada adicional $x_0 = 1$ é denominada entrada de polarização do neurônio. Esta entrada é particularmente importante em redes neurais paramétricas, podendo ser desconsiderada no caso de redes neurais não-paramétricas. Uma simples inspeção da equação (5.1) permite verificar que a atribuição prática da entrada de polarização é permitir uma translação da função f , operação dispensável no caso não-paramétrico por já ser parte inerente do processo de aproximação associado.

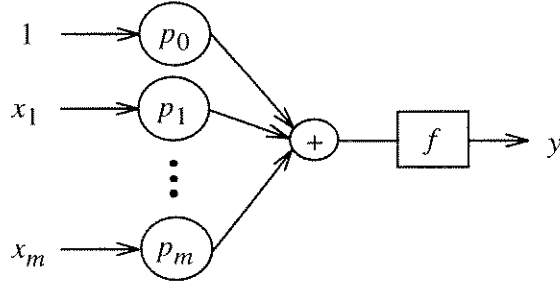


Figura 5.1 - Neurônio artificial com entrada de polarização

Considere, agora, a rede neural paramétrica apresentada na figura 5.2, com três camadas de neurônios descritas na forma:

- camada de entrada: composta por m neurônios sensoriais, responsáveis pela recepção dos sinais de entrada e sua propagação, sem qualquer modificação, para a camada seguinte da rede neural;
- camada intermediária: composta por n neurônios análogos aos da figura 5.1;
- camada de saída: composta por r neurônios análogos aos da figura 5.1, mas com f fixo e igual à função identidade;

Os parâmetros da rede neural são matrizes definidas na forma:

$$\mathbf{V} = \begin{bmatrix} v_{10} & v_{11} & \cdots & v_{1m} \\ v_{20} & \ddots & & \vdots \\ \vdots & & & \\ v_{n0} & \cdots & & v_{nm} \end{bmatrix} \quad \text{e} \quad \mathbf{W} = \begin{bmatrix} w_{10} & w_{11} & \cdots & w_{1n} \\ w_{20} & \ddots & & \vdots \\ \vdots & & & \\ w_{r0} & \cdots & & w_{rn} \end{bmatrix}.$$

Esta rede neural recebe como sinal de entrada um vetor $\mathbf{x} \in \Re^m$, com $\mathbf{x} = [x_1 \cdots x_m]^T$, e produz na saída um vetor $\hat{\mathbf{s}} \in \Re^r$, com $\hat{\mathbf{s}} = [\hat{s}_1 \cdots \hat{s}_r]^T$.

Neste caso, o sinal de ativação do k -ésimo neurônio de saída ($k=1, \dots, r$) é dado por:

$$\begin{aligned} \hat{s}_k &= RN_k^{[n]}(\mathbf{V}, \mathbf{w}_k, \mathbf{x}) = \sum_{j=1}^n w_{kj} y_j + w_{k0} = \sum_{j=1}^n w_{kj} f_j \left(\sum_{i=1}^m v_{ji} x_i + v_{j0} \right) + w_{k0} \\ &= \sum_{j=1}^n w_{kj} f_j (\mathbf{v}_j^T \mathbf{x} + v_{j0}) + w_{k0} = \begin{bmatrix} 1 & f_1(\mathbf{v}_1^T \mathbf{x} + v_{10}) & \cdots & f_n(\mathbf{v}_n^T \mathbf{x} + v_{n0}) \end{bmatrix} \begin{bmatrix} w_{k0} \\ w_{k1} \\ \vdots \\ w_{kn} \end{bmatrix} \end{aligned} \quad (5.2)$$

onde $\mathbf{w}_k \in \Re^n$ é o vetor-coluna formado pelos elementos da k -ésima linha da matriz $\mathbf{W}$ e $\mathbf{v}_j \in \Re^m$ é o vetor-coluna formado pelos elementos da j -ésima linha da matriz $\mathbf{V}$, com exceção dos elementos da primeira coluna de $\mathbf{V}$.

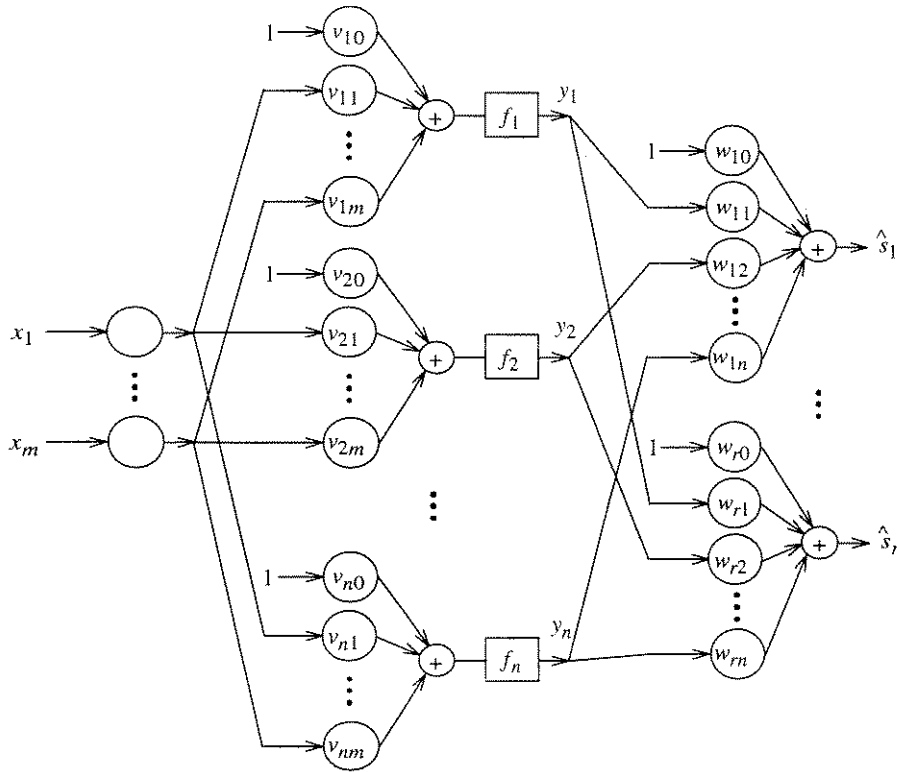


Figura 5.2 - Rede neural artificial com uma camada intermediária

Neste estudo de redes neurais paramétricas, consideram-se funções de ativação idênticas, tais que

$$f_j(\cdot) \equiv f(\cdot), j=1, \dots, n, \quad (5.3)$$

onde $f(\cdot)$ assume uma forma específica. A razão para não se utilizar funções de ativação distintas em redes neurais paramétricas está associada à maior dificuldade de análise do poder de aproximação e à influência exercida por métodos clássicos de expansão em funções básicas, como no caso da série de Fourier. Simplificações junto a modelos de natureza biológica também conduzem a redes neurais com neurônios dotados de funções de ativação idênticas (HECHT-NIELSEN, 1990).

Por outro lado, em problemas de aproximação em que é possível antecipar algum tipo de associação existente entre as variáveis de entrada, este conhecimento inicial pode ser introduzido na rede neural na forma de neurônios com diferentes funções de ativação previamente especificadas.

5.2.1 O problema de aproximação paramétrico

Conforme apresentado no capítulo 3, o problema de aproximação de uma função desconhecida $g: X \subset \mathfrak{R}^m \rightarrow \mathfrak{R}^r$ está baseado em um conjunto finito de dados amostrados $\{(\mathbf{x}_l, \mathbf{s}_l)\}_{l=1}^N$ gerados na forma:

$$\mathbf{s}_l = g(\mathbf{x}_l) + \varepsilon_l, \quad l=1, \dots, N,$$

onde $\mathbf{x}_l \in X$, $\mathbf{s}_l \in \mathfrak{R}^r$ e $\varepsilon_l \in \mathfrak{R}^r$ é um vetor de variáveis aleatórias expressando o ruído, de média zero e variância fixa, do processo de amostragem. Tomando uma rede neural paramétrica com uma camada intermediária como o modelo de aproximação $\hat{g}: X \subset \mathfrak{R}^m \rightarrow \mathfrak{R}^r$, resulta:

$$\hat{g}(\mathbf{x}) = RN^{[n]}(\mathbf{V}, \mathbf{W}, \mathbf{x}) = \begin{bmatrix} RN_1^{[n]}(\mathbf{V}, \mathbf{w}_1, \mathbf{x}) \\ \vdots \\ RN_r^{[n]}(\mathbf{V}, \mathbf{w}_r, \mathbf{x}) \end{bmatrix}. \quad (5.4)$$

conduzindo ao seguinte problema de minimização:

$$\min_{\mathbf{V}, \mathbf{W}} \frac{1}{N} \sum_{k=1}^r \sum_{l=1}^N (\hat{s}_{kl} - s_{kl})^2, \quad (5.5)$$

onde $\hat{s}_{kl} = RN_k^{[n]}(\mathbf{V}, \mathbf{w}_k, \mathbf{x}_l)$ ($k=1, \dots, r; l=1, \dots, N$) são fornecidos pela equação (5.2).

Para o caso de as funções f_j ($j=1, \dots, n$) serem diferenciáveis ao menos até 2ª ordem, os mais eficientes métodos de otimização podem ser aplicados na determinação de $\mathbf{V}^*$ e $\mathbf{W}^*$, soluções do problema paramétrico de aproximação (5.5), conforme apresentado no apêndice D.

5.2.2 A capacidade de aproximação universal

Com funções de ativação idênticas dadas pela equação (5.3), em que $f(\cdot): \mathfrak{R} \rightarrow \mathfrak{R}$ é uma função sigmoidal mensurável tal que

$$f(z) = \begin{cases} a & \text{se } z \rightarrow +\infty, \\ b & \text{se } z \rightarrow -\infty, \end{cases} \text{ com } a > b, \quad (5.6)$$

a teoria de análise funcional é utilizada a seguir para demonstrar que uma função $g: X \subset \mathfrak{R}^m \rightarrow \mathfrak{R}^r$, com $g \in C[X]$, pode ser aproximada uniformemente por uma rede neural do tipo apresentado na figura 5.2, com n finito.

CYBENKO (1989) provou que, se $f(\cdot)$ for contínua, g pode ser aproximada uniformemente por uma rede neural do tipo apresentado na figura 5.2. Os resultados de CYBENKO (1989) foram posteriormente estendidos para $f(\cdot)$ limitada e possivelmente descontínua (JONES, 1990). Utilizando outros procedimentos matemáticos, resultados semelhantes foram obtidos assumindo que $f(\cdot)$ é uma função não-decrescente (FUNAHASHI, 1989; HORNIK *et al.*, 1989). O teorema de Stone-Weierstrass (veja capítulo 2), base das demonstrações de HORNIK *et al.* (1989), foi aplicado de forma mais ampla por COTTER (1990), que apresentou uma série de configurações para redes neurais com uma camada intermediária atendendo às condições do teorema.

No entanto, estes resultados não são construtivos no sentido de não informarem o número n de neurônios na camada intermediária necessário para garantir um determinado nível de aproximação, ou seja, não permitem relacionar o erro de aproximação com o número de neurônios na camada intermediária. Basicamente, se um operador $dist(\cdot, \cdot)$ (específico para cada tipo de norma definida junto ao espaço de aproximação) mede a distância entre a função g a ser aproximada e a função $RN^{[n]}(\mathbf{V}, \mathbf{W}, \mathbf{x})$ produzida pela rede neural, então os resultados enunciados acima podem ser reunidos no teorema existencial a seguir, onde a função sigmoidal $f(\cdot)$ é assumida respeitar as restrições impostas a cada caso.

Teorema 5.1: Se $g: X \subset \mathbb{R}^m \rightarrow \mathbb{R}^r$ é tal que $g \in C[X]$, com X uma região compacta, então, para qualquer $\varepsilon > 0$, existe n , $\mathbf{V}$ e $\mathbf{W}$ tal que

$$dist(g(\mathbf{x}), RN^{[n]}(\mathbf{V}, \mathbf{W}, \mathbf{x})) < \varepsilon.$$

Conclui-se, portanto, que a função $RN^{[n]}(\mathbf{V}, \mathbf{W}, \mathbf{x})$ é densa em $C[X]$.

Como um resultado mais específico, HORNIK *et al.* (1990) e HORNIK (1991) demonstraram que $RN^{[n]}(\mathbf{V}, \mathbf{W}, \mathbf{x})$ é densa em $C^q[X]$ ($q \geq 0$), o espaço das funções continuamente diferenciáveis até ordem q , sendo capazes de aproximar simultaneamente uma classe de funções mensuráveis em espaços de Sobolev (ADAMS, 1975) e suas derivadas.

5.2.3 Propriedades de convergência

No sentido de estimar n , BARRON (1993) utilizou um resultado atribuído a B. Maurey (PISIER, 1981) para provar que, utilizando a representação $RN^{[n]}(\mathbf{V}, \mathbf{W}, \mathbf{x})$, a melhor aproximação de uma função g , que atende a uma condição de suavidade obtida a partir de sua

própria representação de Fourier, apresenta um erro de aproximação que decresce com o inverso da raiz quadrada de n . Esta taxa de convergência é independente da dimensão m do espaço de entrada, desde que a matriz de parâmetros $\mathbf{V}$ seja ajustável.

Os resultados apresentados a seguir são obtidos tomando $r = 1$, sendo diretamente extensíveis ao caso $r > 1$, conforme demonstrado por CHEN *et al.* (1995).

Considere as seguintes definições:

- A rede neural $RN^{[n]}(\mathbf{V}, \mathbf{w}, \mathbf{x})$ tem uma única saída e funções de ativação sigmoidais idênticas a $f(\cdot)$, definida na equação (5.6);
- As funções $g: X \subset \Re^m \rightarrow \Re$ a serem aproximadas são tais que existe sua transformada de Fourier na forma:

$$g(\mathbf{x}) = \int_X e^{j\mathbf{z}^T \mathbf{x}} \tilde{f}(\mathbf{z}) d\mathbf{z},$$

onde $\mathbf{z} \tilde{f}(\mathbf{z})$ é integrável;

- Seja c_g o primeiro momento da distribuição de magnitude de Fourier da função g a ser aproximada, ou seja,

$$c_g = \int_X |\mathbf{z}| |\tilde{f}(\mathbf{z})| d\mathbf{z}, \text{ com } |\mathbf{z}| = \sqrt{\mathbf{z}^T \mathbf{z}}.$$

Se c_g é finito, então a função g é continuamente diferenciável em X .

Com base nestas definições e utilizando aproximação por combinação convexa de famílias de funções (PISIER, 1981; BARRON, 1993), chega-se ao seguinte teorema:

Teorema 5.2: Para toda função $g: X \subset \Re^m \rightarrow \Re$ tal que c_g é finito e para todo $n \geq 1$, existe uma representação $RN^{[n]}(\mathbf{V}, \mathbf{w}, \mathbf{x})$ tal que

$$\text{dist}(g(\mathbf{x}), RN^{[n]}(\mathbf{V}, \mathbf{w}, \mathbf{x})) \leq \frac{h(c_g)}{\sqrt{n}},$$

onde $h(\cdot)$ é uma função contínua e monotonicamente crescente e $\text{dist}(\cdot, \cdot)$ é a norma L_2 do erro de aproximação.

Prova: Veja BARRON (1993).

HORNIK *et al.* (1994) estenderam o resultado deste teorema para aproximação simultânea da função e suas derivadas.

5.2.4 O caso de funções de ativação não-sigmoidais

Ainda considerando funções de ativação idênticas, é possível adotar $f(\cdot)$ como uma função não-sigmoidal. Neste caso, vale ressaltar que os resultados de HORNIK *et al.* (1994), mencionados na seção anterior, também se aplicam a classes de funções de ativação não-sigmoidais, como funções senoidais e gaussianas.

Por exemplo, redes neurais com função de ativação senoidal podem ser diretamente comparadas com a aproximação por expansão em série de Fourier. Em termos de taxa de convergência, se $SF^{[n]}(\mathbf{a}, \mathbf{x})$ é a expansão em série de Fourier contendo n termos, onde $\mathbf{a}$ é o vetor de coeficientes da série, então (BARRON, 1993):

$$\text{dist}(g(\mathbf{x}), SF^{[n]}(\mathbf{a}, \mathbf{x})) \propto \frac{1}{n^{1/m}},$$

onde m é a dimensão do vetor de entrada.

Verifica-se, portanto, que, para $m \geq 2$, a aproximação utilizando a rede neural paramétrica da figura 5.2, com n neurônios com função de ativação senoidal, é pelo menos tão eficiente quanto a aproximação utilizando n termos de uma série de Fourier. De fato, enquanto os termos da expansão por série de Fourier têm frequência fixa, a rede neural com função de ativação senoidal pode ser considerada uma expansão em série de Fourier em que a frequência é um parâmetro ajustável (LAPÉDES & FARBER, 1987).

5.3 Redes neurais artificiais não-paramétricas via métodos construtivos

As redes neurais artificiais não-paramétricas a serem abordadas neste estudo, podem ser derivadas a partir dos modelos de projeção amplamente analisados no capítulo 4, sendo que BARRON & BARRON (1988) foram os primeiros a introduzir modelos de projeção no contexto de estruturas com aprendizado.

No entanto, os resultados a seguir não correspondem simplesmente a uma nova representação de um procedimento de aproximação de funções na forma de uma rede neural. A nova forma de representação do modelo de projeção permite visualizar propriedades algébricas fundamentais para a implementação de um processo de aproximação alternativo, capaz de explorar eficientemente a geometria do problema de aproximação. Este processo

alternativo vai promover a otimização da condição de solvabilidade do problema de aproximação, a ser definida mais adiante, transformando o problema de obtenção de modelos de projeção em um problema de aprendizado construtivo para redes neurais com uma camada intermediária.

O fato de tanto modelos connexionistas como modelos de projeção terem evoluído simultaneamente e de forma exponencial a partir do início dos anos 80 não decorre de uma mera coincidência. A aplicação destes métodos ganhou interesse e factibilidade pela mesma e única razão: a crescente disponibilidade de recursos computacionais a baixo custo. A conjugação desses dois tipos de modelo de aproximação, evidenciando os seus aspectos complementares, produz um modelo de aproximação não-paramétrico multivariável bastante flexível.

Enquanto que o processo de aproximação para uma rede neural paramétrica envolve apenas ajuste de parâmetros, mantendo a estrutura de aproximação fixa, o processo de aproximação para uma rede neural não-paramétrica é predominantemente construtivo, pois parte de uma rede neural com apenas um neurônio na camada intermediária e promove a adição de novos neurônios e subtração de neurônios já existentes, até que uma estrutura de aproximação suficientemente flexível seja obtida. Neste caso, o processo de aproximação pode ser interpretado como um método de otimização do número de neurônios na camada intermediária da rede neural.

5.3.1 A motivação para o uso de métodos construtivos

A motivação para o uso de métodos construtivos pode ser apresentada levando-se em conta o fato de eles operarem no sentido contrário dos métodos de poda.

Conforme descrito por KARNIN (1990), LE CUN *et al.* (1990), HASSIBI & STORK (1993) e REED (1993), os métodos de poda assumem que a arquitetura inicial da rede neural contém pelo menos tanta estrutura quanto a necessária para realizar a tarefa de aproximação. Por exemplo, é comum estabelecer que a arquitetura inicial apresenta uma dimensão elevada e conexões entre todos os neurônios ou, pelo menos, entre todos os neurônios de camadas adjacentes. Neste caso, os recursos considerados em excesso por não estarem sendo utilizados ativamente no processo de aproximação podem ser gradativamente desativados ou simplesmente eliminados. Os recursos em excesso devem ser adequadamente identificados,

podendo corresponder a conexões, neurônios ou até camadas de neurônios. A todo procedimento de poda deve seguir-se um processo de reajuste da estrutura ainda ativa.

No entanto, os métodos de poda apresentam invariavelmente os seguintes problemas (GHOSH & TUMER, 1994; KWOK & YEUNG, 1995):

- não existe um método prático de se determinar diretamente uma arquitetura inicial para a rede neural que contenha garantidamente tanta estrutura quanto a necessária para realizar a tarefa de aproximação. Com isso, para aumentar a probabilidade de se escolher uma arquitetura com tal característica, geralmente adotam-se arquiteturas iniciais fortemente sobredimensionadas;
- já que a maior parte do processo de aproximação é realizado considerando-se redes neurais sobredimensionadas, a demanda por recursos computacionais é grande e parte dos recursos computacionais utilizados acaba sendo desperdiçada toda vez que a poda elimina estruturas que já passaram por alguma fase de processamento;
- como geralmente inúmeras redes neurais de diferentes dimensões são capazes de representar soluções aceitáveis para o problema de aproximação, a aplicação de métodos de poda não favorece a escolha da solução de menor dimensão, ou seja, aquela com menor número de componentes e operadores;
- para que métodos de poda sejam computacionalmente factíveis, eles devem estimar o efeito da eliminação de cada recurso individualmente, mas devem eliminar múltiplos recursos simultaneamente. Com isso, não é possível obter uma estimativa confiável do efeito que cada operação de poda possa vir a causar junto ao erro de aproximação.

O tratamento de parte destes problemas tem conduzido a soluções específicas, como em WEIGEND *et al.* (1991), embora os métodos utilizados sejam mal-condicionados e ainda pouco eficientes computacionalmente, conforme observado por KWOK & YEUNG (1995).

Por operarem no sentido contrário dos métodos de poda, os métodos construtivos podem evitar a ocorrência de problemas como os mencionados acima. No entanto, pelo fato de não ser possível garantir que toda inclusão de estrutura por parte do algoritmo construtivo venha contribuir para a solução do problema de aproximação, métodos de poda representam um procedimento complementar importante no sentido de promover a eliminação de estruturas desnecessariamente incluídas.

Conclui-se, portanto, que um método híbrido de aproximação é o mais adequado, contendo etapas construtivas seguidas de etapas de poda. Em razão de procedimentos de poda

só entrarem em operação caso o método construtivo falhe na definição de estruturas adicionais junto ao modelo de aproximação, o método híbrido de aproximação é predominantemente construtivo. Sendo assim, no decorrer da apresentação é preservada a denominação de método construtivo de aproximação.

5.3.2 O tratamento de múltiplas saídas

O modelo de projeção apresentado no capítulo anterior foi desenvolvido considerando-se que a função a ser aproximada é do tipo $g: X \subset \mathfrak{R}^m \rightarrow \mathfrak{R}^r$, com $r = 1$. Duas possíveis extensões para o tratamento do caso $r > 1$ são discutidas a seguir: múltiplas aproximações monovariáveis e aproximação multivariável.

A forma mais simples de se obter a aproximação considerando múltiplas variáveis de saída é utilizar modelos de aproximação independentes, cada um desenvolvido para tratar uma única variável de saída. Sendo r o número de variáveis de saída, determinam-se r modelos de aproximação como segue:

$$\hat{g}_{n_k,k}(\mathbf{x}) = \sum_{j=1}^{n_k} f_{kj}(\mathbf{v}_{kj}^T \mathbf{x}), \quad k=1, \dots, r, \quad (5.7)$$

um para cada variável de saída, produzindo um modelo de aproximação na forma:

$$\hat{g}(\mathbf{x}) = \begin{bmatrix} \hat{g}_{n_1,1}(\mathbf{x}) \\ \vdots \\ \hat{g}_{n_r,r}(\mathbf{x}) \end{bmatrix}. \quad (5.8)$$

Por outro lado, a possível existência de associações entre as variáveis de saída pode ser explorada na obtenção de modelos de aproximação que demandem menos recursos computacionais ao aproximarem múltiplas variáveis de saída simultaneamente. Além disso, esta possibilidade de se utilizar um menor número de parâmetros e operadores pode auxiliar na obtenção de melhores resultados em termos de generalização e também conduzir a modelos mais facilmente interpretáveis (FRIEDMAN, 1985). Sendo assim, cada variável de saída é aproximada como segue:

$$\hat{g}_{n,k}(\mathbf{x}) = \sum_{j=1}^n w_{kj} f_j(\mathbf{v}_j^T \mathbf{x}), \quad k=1, \dots, r, \quad (5.9)$$

produzindo um modelo de aproximação na forma:

$$\hat{g}(\mathbf{x}) = \begin{bmatrix} \hat{g}_{n,1}(\mathbf{x}) \\ \vdots \\ \hat{g}_{n,r}(\mathbf{x}) \end{bmatrix}. \quad (5.10)$$

Observe que o modelo representado pela equação (5.8) corresponde a um caso particular do modelo representado pela equação (5.10), bastando assumir:

- $n = \sum_{k=1}^r n_k$;
- $w_{kj} = 0$, para $j \leq \sum_{i=1}^{k-1} n_i$ e $j > \sum_{i=1}^k n_i$;
- $f_{kj} = w_{kl} f_l$, com $l = \sum_{i=1}^{k-1} n_i + j$.

BREIMAN & FRIEDMAN (1996) apresentam um estudo mais aprofundado das vantagens e desvantagens destes dois modelos de aproximação multivariável, mostrando que modelos na forma da equação (5.10) geralmente produzem melhores resultados. Utilizando exemplos de simulação computacional, MALTHOUSE (1995) compara o desempenho dos dois modelos, chegando a resultados que concordam com aqueles obtidos por BREIMAN & FRIEDMAN (1996). Outros modelos de aproximação multivariável são apresentados por FRIEDMAN (1994).

5.3.3 O algoritmo de aproximação construtivo

Reproduzindo o modelo de aproximação de uma rede neural com uma camada intermediária (veja equação (5.2)) na forma:

$$RN_k^{[n]}(\mathbf{V}, \mathbf{w}_k, \mathbf{x}) = \sum_{j=1}^n w_{kj} f_j(\mathbf{v}_j^T \mathbf{x} + v_{j0}) + w_{k0}, \quad k=1, \dots, r, \quad (5.11)$$

uma comparação com a equação (5.9) permite verificar que a única diferença entre os dois modelos é a presença de termos adicionais na equação (5.11), representando a polarização dos neurônios.

Com isso, conclui-se que uma rede neural não-polarizada com uma camada intermediária pode ser considerada o resultado da implementação de um método estatístico especialmente desenvolvido para a geração de modelos de aproximação na forma da equação (5.10), denominado SMART ('Smooth Multiple Additive Regression Technique') (FRIEDMAN,

1984). Quando necessário, a polarização pode ser prontamente adicionada, produzindo um modelo de aproximação $\hat{g}$ formado por uma composição aditiva de funções f_j ($j=1, \dots, n$) adequadamente transladadas (por v_{j0}), rotacionadas (por $\mathbf{v}_j$) e escalonadas (por $\|\mathbf{v}_j\|$). Para cada saída k , as avaliações de f_j em projeções de $\mathbf{x}$ na direção $\mathbf{v}_j$ são, por sua vez, adequadamente transladadas (por w_{k0}) e escalonadas (por w_{kj}).

Tomando como ponto de partida o algoritmo de retroajuste apresentado no capítulo anterior, seção 4.6, o processo de aprendizado construtivo para uma rede neural com uma camada intermediária, m entradas e r saídas, pode ser apresentado na forma:

Algoritmo de aproximação construtivo para múltiplas saídas (HWANG *et al.*, 1994)

1. Dados $\mathbf{X} \in \mathbb{R}^{m \times N}$ e $\mathbf{S} \in \mathbb{R}^{r \times N}$, tome $j = 0$, $\mathbf{D} = \mathbf{S}$;
2. Faça $j = j+1$ e atribua um valor inicial para $\mathbf{v}_j \in \mathbb{R}^m$, uma forma inicial para f_j e um valor inicial para $\mathbf{w}_j \in \mathbb{R}^r$;
3. Utilizando $\mathbf{X}$ e $\mathbf{D}$, resolva os seguintes problemas em sequência até convergência (medida por algum critério de parada):
 - 3.1. fixe f_j e $\mathbf{w}_j$ e obtenha um valor ótimo para $\mathbf{v}_j$;
 - 3.2. fixe $\mathbf{v}_j$ e $\mathbf{w}_j$ e obtenha um f_j ótimo;
 - 3.3. fixe $\mathbf{v}_j$ e f_j , obtenha um valor ótimo para $\mathbf{w}_j$ e retorne ao passo 3.1;
4. Para cada b tal que $1 \leq b < j$, calcule

$$\mathbf{D} = \mathbf{S} - \sum_{\substack{k=1 \\ k \neq b}}^j \mathbf{w}_k f_k(\mathbf{v}_k^T \mathbf{X}),$$

e repita o passo 3, com $j = b$;

5. Por avaliação da participação de cada neurônio $\mathbf{w}_k f_k(\mathbf{v}_k^T \mathbf{X})$, $k = 1, \dots, j$, na representação da matriz $\mathbf{S}$, aplique um procedimento de poda de neurônios que não apresentem um nível de participação mínima;
6. Enquanto um determinado nível de aproximação não for atingido (medido por algum critério de parada), retorne ao passo 2.

Observe que, no passo 2, HWANG *et al.* (1994) não utilizam os dados disponíveis $\mathbf{X}$ e $\mathbf{S}$ para definir um valor inicial para os parâmetros e uma forma inicial para a função de ativação, recorrendo a uma definição arbitrária. O passo 3 é a etapa fundamental do processo de ajuste,

sendo que HWANG *et al.* (1994) aplicam o método de minimização de Gauss-Newton para resolver o item 3.1, uma função paramétrica de suavização dos dados unidimensionais como solução do item 3.2 (veja capítulo 4) e pseudo-inversão para resolver o item 3.3. O passo 4 é o responsável pelo processo de retroajuste, enquanto que o passo 5 executa, quando necessário, a poda de neurônios já introduzidos.

A implicação prática do método de retroajuste é promover algum tipo de adaptação por partes de toda a rede neural sempre que um novo neurônio for acrescentado. Com isso, enquanto no caso da rede neural paramétrica o ajuste era mais custoso (se aplicava a todos os neurônios ao mesmo tempo) mas feito uma única vez, aqui o ajuste é menos custoso (se aplica a um neurônio de cada vez) mas deve ser aplicado várias vezes e de forma cíclica. Outra característica importante é o fato de que o ajuste de cada neurônio, mantendo os outros fixos, é feito por camada e também de forma cíclica:

- ajusta-se a direção de projeção ($\mathbf{v}_j$);
- ajusta-se a função de ativação ou função de expansão ortogonal (f_j);
- ajustam-se os parâmetros correspondentes da camada de saída ($\mathbf{w}_j$);
- repete-se o processo até convergência.

Apesar de não apresentarem resultados comparáveis a este processo de retroajuste, principalmente por demandarem excessivos recursos computacionais, ASH (1989) e HIROSE *et al.* (1991) apresentam métodos construtivos que realizam um ajuste simultâneo de todos os neurônios sempre que um novo neurônio é acrescentado à estrutura da rede neural. MOODY (1994) propõe, por sua vez, um processo híbrido, ou seja, a aplicação do processo de retroajuste seguido de um ajuste simultâneo de todos os neurônios.

5.3.4 Procedimentos de aperfeiçoamento do algoritmo

O algoritmo de aproximação construtivo descrito acima pode ainda ser substancialmente aperfeiçoado, produzindo um processo mais eficiente e menos custoso computacionalmente (VON ZUBEN & NETTO, 1995a; VON ZUBEN & NETTO, 1995c). A versão final do algoritmo de aproximação construtivo é apresentada ao final da seção 5.3.5., contendo três importantes modificações discutidas a seguir.

5.3.4.1 Definição das condições iniciais do algoritmo a partir dos dados

Primeiramente, utilizando os resultados do capítulo 4, é possível empregar as matrizes de dados disponíveis $\mathbf{X}$ e $\mathbf{S}$ para definir um valor inicial para o vetor de parâmetros $\mathbf{v}_j$, representando a direção de projeção mais interessante. Com isso, utilizando um resultado a ser apresentado na seção 5.3.5, apenas com o conhecimento de $\mathbf{v}_j$ inicial pode-se determinar uma função f_j inicial também a partir de $\mathbf{X}$ e $\mathbf{S}$. Dadas condições iniciais para $\mathbf{v}_j$ e f_j , viabiliza-se a definição, em forma fechada e utilizando $\mathbf{X}$ e $\mathbf{S}$, de um valor inicial ótimo para $\mathbf{w}_j$.

Pelo fato de o problema de aproximação como um todo não ser convexo, a definição de condições iniciais mais próximas do ótimo global ou de ótimos locais próximos da otimalidade global (levando-se em conta os dados disponíveis, $\mathbf{X}$ e $\mathbf{S}$) é importante para reduzir a influência de ótimos locais distantes da otimalidade global e proporcionar uma convergência mais rápida (LUENBERGER, 1989).

5.3.4.2 A iteração em dois grupos de variáveis e solução fechada para o terceiro grupo

Com base no mesmo resultado a ser apresentado na seção 5.3.5, a solução iterativa para o passo 3 do algoritmo de HWANG *et al.* (1994) pode ser construída sem a necessidade de iteração em três grupos de variáveis. Basta iteragir em $\mathbf{v}_j$ e f_j (sem utilizar qualquer valor para $\mathbf{w}_j$) e, após convergência, obter $\mathbf{w}_j$ otimamente e de forma fechada.

Além de requerer um menor custo computacional, este procedimento de solução iterativa em dois grupos de variáveis tem propriedades de convergência análogas às aquelas apresentadas no capítulo 4 para o caso de uma única saída. Por outro lado, não existem resultados que garantam convergência para o caso de solução iterativa em três grupos de variáveis.

5.3.4.3 Aproximação e critérios de parada baseados em validação cruzada

Tanto na determinação da função de ativação f_j para cada neurônio j , como também na definição do número n de neurônios na camada intermediária, procedimentos de validação cruzada ordinária e validação cruzada generalizada (derivada a partir de conceitos de regularização, conforme apresentados no capítulo 4) podem ser aplicados para aumentar a capacidade de generalização da rede neural resultante.

5.3.5 Otimização da condição de solvabilidade

O processo de retroajuste descrito na seção 5.3.3 pode ainda ser substancialmente melhorado. Tirando proveito de sua representação na forma de uma rede neural com uma camada intermediária, é possível visualizar propriedades algébricas fundamentais para a implementação de um processo capaz de explorar eficientemente a geometria do problema de aproximação (VON ZUBEN & NETTO, 1995a; VON ZUBEN & NETTO, 1995c).

Com base na equação (5.11), representando o modelo de aproximação de uma rede neural com uma camada intermediária, e no conjunto de dados de aproximação $\{(\mathbf{x}_l, \mathbf{s}_l)\}_{l=1}^N$ gerado na forma:

$$\mathbf{s}_l = g(\mathbf{x}_l) + \varepsilon_l, \quad l=1, \dots, N,$$

onde $\mathbf{x}_l \in X \subset \mathfrak{R}^m$, $\mathbf{s}_l, \varepsilon_l \in \mathfrak{R}^r$, o sistema de equações

$$\hat{s}_{kl} = RN_k^{[n]}(\mathbf{V}, \mathbf{w}_k, \mathbf{x}_l) = \sum_{j=1}^n w_{kj} f_j(\mathbf{v}_j^T \mathbf{x}_l + v_{j0}) + w_{k0}, \quad k=1, \dots, r; \quad l=1, \dots, N,$$

pode ser colocado na forma matricial (BÄRMANN & BIEGLER-KÖNIG, 1992)

$$\Sigma \mathbf{W}^T = \hat{\mathbf{S}}^T, \quad (5.12)$$

onde

- $\mathbf{V} = \begin{bmatrix} v_{10} & v_{11} & \cdots & v_{1n} \\ v_{20} & \ddots & & \vdots \\ \vdots & & & \\ v_{n0} & \cdots & & v_{nm} \end{bmatrix}$, com $\mathbf{v}_j \in \mathfrak{R}^m$ o vetor-coluna formado pelos elementos da j -ésima

linha da matriz $\mathbf{V}$, com exceção dos elementos da primeira coluna de $\mathbf{V}$;

- $\Sigma = \begin{bmatrix} 1 & \sigma_{11} & \cdots & \sigma_{1n} \\ 1 & \sigma_{21} & \ddots & \vdots \\ \vdots & \vdots & & \\ 1 & \sigma_{N1} & \cdots & \sigma_{Nn} \end{bmatrix}$, com $\sigma_{lj} = f_j(\mathbf{v}_j^T \mathbf{x}_l + v_{j0})$;

- $\mathbf{W} = \begin{bmatrix} w_{10} & w_{11} & \cdots & w_{1n} \\ w_{20} & \ddots & & \vdots \\ \vdots & & & \\ w_{r0} & \cdots & & w_{rn} \end{bmatrix}$;

- $\hat{\mathbf{S}} = \begin{bmatrix} \hat{s}_{11} & \hat{s}_{12} & \cdots & \hat{s}_{1N} \\ \hat{s}_{21} & \ddots & & \vdots \\ \vdots & & & \\ \hat{s}_{r1} & \cdots & & \hat{s}_{rN} \end{bmatrix}$.

Definindo a matriz $\mathbf{S}$ na forma:

$$\mathbf{S} = [\mathbf{s}_1 \ \mathbf{s}_2 \ \cdots \ \mathbf{s}_N] = \begin{bmatrix} s_{11} & s_{12} & \cdots & s_{1N} \\ s_{21} & \ddots & & \vdots \\ \vdots & & & \\ s_{r1} & \cdots & & s_{rN} \end{bmatrix},$$

e o vetor de funções de ativação $\mathbf{f} = [f_1 \ f_2 \ \cdots \ f_n]^T$, o problema de aproximação pode ser expresso como segue:

$$\min_{n, \mathbf{V}, \mathbf{W}, \mathbf{f}} \|\hat{\mathbf{S}}^T - \mathbf{S}^T\| + \lambda \phi(\mathbf{V}, \mathbf{W}, \mathbf{f}), \quad (5.13)$$

onde $\phi(\cdot)$ é a função de regularização adotada e λ é o parâmetro de regularização (veja capítulos 3 e 4). Definindo a norma matricial na forma:

$$\|\hat{\mathbf{S}}^T - \mathbf{S}^T\| = \frac{1}{N} \sum_{k=1}^r \sum_{l=1}^N (\hat{s}_{kl} - s_{kl})^2 \quad (5.14)$$

é possível estabelecer uma comparação entre a equação (5.5), representando o problema de aproximação paramétrico, e a equação (5.13), representando o problema de aproximação não-paramétrico regularizado. Apesar de os dois tipos de problemas utilizarem uma rede neural com uma camada intermediária, fica claro que o processo não-paramétrico promove ajustes estruturais na rede neural, além do ajuste paramétrico, única etapa do processo paramétrico.

Sem considerar restrições de suavidade, o problema de aproximação (5.13) pode ser colocado na forma:

$$\min_{n, \mathbf{V}, \mathbf{W}, \mathbf{f}} \|\hat{\mathbf{S}}^T - \mathbf{S}^T\| = \min_{n, \mathbf{V}, \mathbf{W}, \mathbf{f}} \|\Sigma \mathbf{W}^T - \mathbf{S}^T\|, \quad (5.15)$$

com a norma matricial dada pela equação (5.14) e já substituindo $\hat{\mathbf{S}}^T$ dado pela equação (5.12). Portanto, a condição para solução exata do problema de aproximação irrestrito (5.15), denominada condição de solvabilidade, requer que cada coluna de $\mathbf{S}^T$ pertença ao subespaço linear gerado pelas colunas da matriz Σ de dimensão $N \times (n+1)$ (SKELTON, 1988). Observe que a matriz de parâmetros $\mathbf{W}$ não contribui para a definição da condição de solvabilidade.

Para qualquer problema de aproximação, ou seja, para qualquer escolha de $\mathbf{S}$, um requisito necessário para se atender a condição de solvabilidade em (5.15) é tomar $(n+1) \geq N$ (veja apêndice A). Isto implica em adotar redes neurais com um número de neurônios na camada intermediária igual ou maior que o número de dados de aproximação. Além da

inviabilidade prática, o atendimento desta condição necessária acarretaria problemas em termos de generalização, já abordados no capítulo 3 e a serem ainda discutidos neste capítulo.

Portanto, a motivação para a utilização de $n+1 < N$ está vinculada à imposição de restrições de suavidade junto ao modelo de aproximação, diretamente associada ao termo de regularização presente no problema (5.13). Com isso, o atendimento da condição de solvabilidade não é o objetivo do processo de aproximação que busca resolver o problema (5.13). O objetivo, na verdade, é *otimizar* a condição de solvabilidade sem violar restrições de suavidade impostas pelo termo de regularização.

Utilizando resultados de geometria analítica, o método de otimização da condição de solvabilidade para o processo de aproximação construtivo apresentado em seções anteriores é descrito a seguir.

Denominando a j -ésima coluna de Σ por σ_j e a j -ésima coluna de $\mathbf{W}$ por $\mathbf{w}_j$, tem-se:

$$\hat{\mathbf{S}}^T = \Sigma \cdot \mathbf{W}^T = \sum_{j=0}^n \sigma_j \cdot \mathbf{w}_j^T, \quad (5.16)$$

onde σ_j ($j=1, \dots, n$) é função de $\mathbf{v}_j$ e f_j .

Com $n = 0$, resolve-se o seguinte problema de minimização

$$\min_{\mathbf{w}_0} \left\| \sigma_0 \cdot \mathbf{w}_0^T - \mathbf{S}^T \right\| = \min_{\mathbf{w}_0} \left\| \begin{bmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{bmatrix} \cdot [w_{10} \ w_{20} \ \cdots \ w_{r0}] - \mathbf{S}^T \right\|. \quad (5.17)$$

Para um dado $n > 0$, seja $\mathbf{D}$ a matriz contendo os dados ainda não representados pela rede neural com $n-1$ neurônios na camada intermediária, ou seja,

$$\mathbf{D} = \mathbf{S} - \sum_{j=0}^{n-1} \sigma_j \cdot \mathbf{w}_j^T. \quad (5.18)$$

Sendo assim, resolve-se o seguinte problema de minimização

$$\min_{\sigma_n, \mathbf{w}_n} \left\| \sigma_n \cdot \mathbf{w}_n^T - \mathbf{D}^T \right\| + \lambda_n \phi(f_n). \quad (5.19)$$

O primeiro termo da função objetivo do problema (5.19) força σ_n a se alinhar otimamente com as colunas de $\mathbf{D}^T$, sendo que a direção que atende a este objetivo no sentido dos quadrados mínimos é a fornecida pelo componente principal da matriz $\mathbf{D}^T \mathbf{D}$ (MARDIA *et al.*, 1979).

Se $a_1 \geq a_2 \geq \dots \geq a_N \geq 0$ são os autovalores de $\mathbf{D}^T \mathbf{D}$, com autovetores respectivamente dados por $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_N$ ($\mathbf{u}_l \in \Re^{N \times 1}$; $l=1, \dots, N$), então o componente principal de $\mathbf{D}^T \mathbf{D}$ é $\mathbf{u}_1$ (ou equivalentemente $-\mathbf{u}_1$). Assim, σ_n é ótimo para algum $\mathbf{v}_n \in \Re^m$ e $f_n: \Re \rightarrow \Re$ que maximize o produto escalar normalizado (BÄRMANN & BIEGLER-KÖNIG, 1992):

$$\sigma_n^* = \arg \max_{\sigma_n} \frac{(\sigma_n^T \mathbf{u}_1)^2}{\sigma_n^T \sigma_n}. \quad (5.20)$$

A solução deste problema envolve uma iteração em $\mathbf{v}_n$ e f_n até convergência, dada na forma (VON ZUBEN & NETTO, 1995a; VON ZUBEN & NETTO, 1995c):

- 1) com f_n fixo, encontre $\mathbf{v}_n^*$ tal que $\mathbf{v}_n^* = \arg \max_{\mathbf{v}_n} \frac{(\sigma_n^T \mathbf{u}_1)^2}{\sigma_n^T \sigma_n}$;
- 2) Com $\mathbf{v}_n$ fixo, encontre f_n^* tal que $f_n^* = \arg \min_{f_n} \sum_{l=1}^N (u_{ln} - f_n(\mathbf{v}_n^T \mathbf{x}_l))^2 + \lambda_n \phi(f_n)$.

O problema de maximização do item 1) pode ser resolvido aplicando-se o método de Newton modificado, apresentado no apêndice D, enquanto que o problema de minimização do item 2) pode ser resolvido aplicando-se diretamente os resultados do capítulo 4, seja com f_n uma função paramétrica ou não-paramétrica. A condição inicial para o problema do item 1) é obtida a partir do conjunto de índices de projeção apresentados na capítulo 4. Mesmo não sendo empregados neste estudo, vale mencionar que já existem índices de projeção especificamente projetados para operarem no contexto de redes neurais artificiais, utilizando técnicas de treinamento não-supervisionado (INTRATOR, 1993a; 1993b; 1993c).

Uma vez obtido σ_n^* , ou seja, $\mathbf{v}_n^*$ e f_n^* após convergência, o valor ótimo de $\mathbf{w}_n$ é dado por:

$$\mathbf{w}_n^* = (\sigma_n^{*T} \sigma_n^*)^{-1} \sigma_n^{*T} \mathbf{D}^T. \quad (5.21)$$

Observe que a norma $\|\sigma_n^* \cdot \mathbf{w}_n^{*T} - \mathbf{D}^T\|$ só se anula se as colunas de $\mathbf{D}^T$ forem linearmente dependentes.

Conclui-se, então, que pode-se obter $\mathbf{v}_n^*$, f_n^* e $\mathbf{w}_n^*$ iteragindo apenas em $\mathbf{v}_n$ e f_n , e calculando $\mathbf{w}_n^*$ posteriormente e de forma fechada. A introdução deste resultado no algoritmo de aproximação retroativo para múltiplas saídas, apresentado na seção 5.3.3, permite melhorar consideravelmente sua eficiência, ao evitar a iteração em $\mathbf{v}_n, f_n$ e $\mathbf{w}_n$, produzindo:

Algoritmo de aproximação construtivo para múltiplas saídas (Versão Final)

1. Dados $\mathbf{X} \in \mathbb{R}^{m \times N}$ e $\mathbf{S} \in \mathbb{R}^{r \times N}$, tome $j = 0$, $\mathbf{D} = \mathbf{S}$;
2. Utilizando $\mathbf{X}$ e $\mathbf{D}$, faça $j = j+1$ e atribua:
 - um valor inicial para $\mathbf{v}_j \in \mathbb{R}^m$, empregando os índices de projeção apresentados no capítulo 4 (tomar a melhor solução);
 - uma forma inicial para f_j , empregando os resultados apresentados no capítulo 4 (f_j pode ser uma função paramétrica ou não-paramétrica);
 - um valor inicial para $\mathbf{w}_j \in \mathbb{R}^r$, já empregando a equação (5.21);
3. Utilizando $\mathbf{X}$ e $\mathbf{D}$, resolva os seguintes problemas em seqüência até convergência (medida por algum critério de parada):
 - 3.1. fixe f_j e obtenha um valor ótimo para $\mathbf{v}_j$;
 - 3.2. fixe $\mathbf{v}_j$, obtenha um f_j ótimo via técnicas de regularização e retorne ao passo 3.1;
4. Obtenha um valor ótimo para $\mathbf{w}_j$ empregando a equação (5.21);
5. Para cada b tal que $1 \leq b < j$, calcule

$$\mathbf{D} = \mathbf{S} - \sum_{\substack{k=1 \\ k \neq b}}^j \mathbf{w}_k f_k(\mathbf{v}_k^T \mathbf{X}),$$
 e repita os passos 3 e 4, com $j = b$;
6. Por avaliação da participação de cada neurônio $\mathbf{w}_k f_k(\mathbf{v}_k^T \mathbf{X})$, $k = 1, \dots, j$, na representação da matriz $\mathbf{S}$, aplique um procedimento de poda de neurônios que não apresentem um nível de participação mínima;
7. Enquanto um determinado nível de aproximação não for atingido (medido por algum critério de parada, por exemplo, utilizando validação cruzada), retorne ao passo 2.

5.3.6 A convergência do processo de aproximação construtivo

O processo de aproximação de $g: X \subset \mathbb{R}^m \rightarrow \mathbb{R}^r$ descrito nas seções anteriores parte de um modelo inicial $\hat{g}_{0,k}(\mathbf{x}) = 0$ e implementa a aproximação na forma:

$$\hat{g}_{n,k}(\mathbf{x}) = \hat{g}_{n-1,k}(\mathbf{x}) + w_{k,n} f_n(\mathbf{v}_n^T \mathbf{x}), \quad k=1, \dots, r; \quad n > 0, \quad (5.22)$$

onde $w_{k,n}f_n(\mathbf{v}_n^T \mathbf{x})$ é definido de tal forma a produzir uma solução minimizante para o problema:

$$\min_{w,f,v} dist(g_k(\mathbf{x}) - \hat{g}_{n-1,k}(\mathbf{x}), wf(\mathbf{v}^T \mathbf{x})), k=1, \dots, r,$$

onde o operador $dist(\cdot, \cdot)$ mede a distância entre as funções em todo o espaço X . Como $\hat{g}_{n-1,k}(\mathbf{x})$ é mantido constante durante este processo de minimização, diz-se que o ajuste é realizado neurônio a neurônio, procedimento também adotado em outros contextos por FAHLMAN & LEBIERE (1990), SETIONO & HUI (1995) e ZHANG (1994), cujos resultados experimentais atestam que o ajuste neurônio a neurônio produz processos de aproximação mais eficazes computacionalmente. De fato, torna-se mais fácil representar informações menos significativas após a representação das informações mais significativas. Além disso, como cada neurônio da camada intermediária é otimizado de forma a representar o máximo possível de estrutura ainda não representada pelos outros neurônios, a correlação entre as informações representadas por diferentes neurônios tende a ser a menor possível.

No entanto, a convergência do ajuste neurônio a neurônio não pode ser demonstrada diretamente a partir da capacidade de aproximação universal da rede neural resultante. A demonstração de convergência foi obtida por JONES (1987). Como mostrado por CHEN *et al.* (1995), toda a dedução pode ser realizada tomando-se $r = 1$.

Neste caso, um resultado bem conhecido em teoria de aproximação é a capacidade apresentada por séries de Fourier com termos na forma $\cos(\theta_j^T \mathbf{x} + \theta_{0j})$, $\mathbf{x}, \theta_j \in \mathfrak{R}^m$, $j=1, \dots$, em aproximar uniformemente funções $g: X \subset \mathfrak{R}^m \rightarrow \mathfrak{R}$ contínuas (DAVIS, 1975). Basicamente, aproxima-se g por uma função $\hat{g}_\infty$ tal que

$$\frac{\partial^{2m} \hat{g}_\infty}{\partial x_1^2 \partial x_2^2 \cdots \partial x_m^2}$$

é contínua. Sendo assim, $\hat{g}_\infty$ tem uma série cosenoidal

$$\hat{g}_\infty(\mathbf{x}) = \sum_{j=1}^{\infty} \cos(\theta_j^T \mathbf{x} + \theta_{0j})$$

uniformemente convergente em X , a qual pode ser truncada para aproximar g . Observe que, se n_1 for o número de termos da série truncada, sempre existe um modelo de projeção tal que (JONES, 1990):

$$\hat{g}_n(\mathbf{x}) = \sum_{j=1}^{n_1} \cos(\theta_j^T \mathbf{x} + \theta_{0,j}) = \sum_{j=1}^n f_j(\mathbf{v}_j^T \mathbf{x}). \quad (5.23)$$

com f_j ($j=1, \dots, n$) funções contínuas, $\mathbf{v}_j \in \mathcal{R}^m$ ($j=1, \dots, n$) e n finito.

A aproximação usando uma série cosenoidal finita é adequada apenas para funções g com uma banda de frequência limitada. Se o espectro de frequência de g apresentar componentes de alta frequência, ou seja, se g tem energia em altas frequências, o número n de termos da série truncada pode ser muito elevado. Esta explosão de n pode ser evitada utilizando-se modelos de projeção, como os apresentados na equação (5.23), em que $\mathbf{v}_j$ e f_j ($j=1, \dots, n$) são escolhidos sequencialmente de maneira a incluir o máximo de energia possível para um dado n , tomando um número finito de amostras de g . Neste caso, JONES (1992) mostra que o erro de aproximação de g é da ordem $O(1/\sqrt{n})$.

Associados a estes resultados formais, resultados heurísticos atestam que procedimentos construtivos que acrescentam um neurônio de cada vez conduzem a processos de aproximação com taxas de convergência maiores, principalmente por evitarem o problema de não-estacionariedade da função de minimização, comum em processos de aproximação paramétricos, ou seja, com estrutura de aproximação estática previamente definida (KWOK & YEUNG, 1995). A não-estacionariedade da função de minimização está associada a processos de ajuste simultâneo em que todos os parâmetros estão sujeitos a modificação. Neste caso, a função de minimização associada a um determinado parâmetro pode ser influenciada por outros parâmetros e, se todos os parâmetros estão sujeitos a modificação, esta função de minimização acaba sofrendo modificações mesmo que o parâmetro com que está diretamente associada permaneça fixo. Com isso, levando-se em conta apenas o erro na fase de treinamento, conclui-se que, enquanto os procedimentos construtivos vão produzir um erro de aproximação certamente menor com um número maior de neurônios na camada intermediária (a menos que o erro já tenha se anulado), o mesmo não se pode afirmar com relação a processos de ajuste simultâneo de parâmetros em procedimentos não-construtivos (BIEGLER-KÖNIG & BÄRMANN, 1993).

5.3.7 Comparação com outros métodos construtivos

Inúmeros outros métodos construtivos podem ser utilizados para definir a arquitetura ótima de uma rede neural para um determinado problema de aproximação. Neste estudo, o método construtivo se baseia no procedimento sistemático de obtenção do modelo de projeção

apresentado no capítulo anterior, que constrói o modelo de aproximação com base na equação (5.22), definida acima.

Quando aplicado a redes neurais artificiais, este método parte de uma rede neural com uma camada intermediária, do tipo apresentado na figura 5.2, com $n = 0$. A partir desta estrutura inicial, o algoritmo construtivo passa a introduzir novos neurônios na camada intermediária, um a um, até que a rede neural seja capaz de realizar a tarefa de aproximação. Caso neurônios já introduzidos não estejam contribuindo significativamente para aumentar o poder de representação da rede neural, um procedimento de poda é acionado para suprimi-los da estrutura da rede neural. Logo, este método construtivo (com fases de poda) é um procedimento de geração e teste na forma:

1. incrementa-se a estrutura de representação;
2. gera-se uma solução para o problema de aproximação com base nesta estrutura;
3. verifica-se se a solução é aceitável;
4. Se for aceitável, fim. Caso contrário, verifica-se a viabilidade da supressão de estruturas pouco representativas, executa-se o procedimento de poda e retorna-se ao passo 1.

Deduz-se, portanto, que este método construtivo só pode produzir um tipo de arquitetura, como a maior parte dos métodos construtivos (ASH, 1989; HANSON, 1990; HIROSE *et al.*, 1991; LEE & KIL, 1991; SHIN & GHOSH, 1995). No caso em estudo, resulta sempre uma rede neural com uma camada intermediária, do tipo apresentado na figura 5.2.

Certamente, procedimentos construtivos mais complexos podem realizar uma busca mais ampla considerando outras arquiteturas de redes neurais, por exemplo, redes neurais com múltiplas camadas intermediárias e até redes neurais em que todos os neurônios possam se conectar entre si. No entanto, estes procedimentos construtivos mais complexos requerem custos computacionais elevados e têm propriedades de convergência desconhecidas ou muito inferiores às apresentadas na seção anterior (KWOK & YEUNG, 1995). Isto não impede, no entanto, estudos envolvendo métodos construtivos alternativos, como os realizados por FAHLMAN & LEBIERE (1990), NABHAN & ZOMAYA (1994), SANGER (1991), TENG & WAH (1994) e TENORIO & LEE (1990).

Os fatores que contribuem para a complexidade do processo de busca da melhor arquitetura da rede neural podem ser enunciados na forma (MILLER *et al.*, 1989):

- o espaço de busca é infinito e o número de neurônios e conexões entre neurônios é ilimitado;

- o espaço de busca não é diferenciável, pois as mudanças no número de neurônios e conexões são discretas, além de promoverem um efeito descontínuo no comportamento da rede neural;
- redes neurais com estruturas similares podem apresentar comportamento muito diferentes, enquanto que comportamentos similares podem ser produzidos por redes neurais estruturalmente muito diferentes;
- o resultado do processo de busca depende da condição inicial.

Métodos de computação evolutiva são bastante apropriados para tratar problemas com estas características (BALAKRISHNAN & HONAVAR, 1995; WEISS, 1994; YAO, 1993), permitindo a geração de redes neurais com arquiteturas bastante variadas. Observe, no entanto, que a aplicação eficaz de algoritmos genéticos a métodos construtivos requer um processo complexo de codificação, conduzindo à necessidade de utilização de operadores genéticos específicos (KWOK & YEUNG, 1995).

5.3.7.1 Algoritmos construtivos para problemas de classificação

Apesar de um problema de classificação de padrões poder ser visto como um problema de aproximação de funções, onde a função a ser aproximada é a hipersuperfície que separa as diferentes classes de padrões, para efeito de implementação de métodos de classificação construtivos é mais adequado considerar a seguinte distinção: enquanto problemas de aproximação de funções mapeiam um espaço discreto ou contínuo em um espaço contínuo, problemas de classificação de padrões mapeiam um espaço discreto ou contínuo em um espaço discreto.

Neste sentido, algoritmos construtivos especificamente projetados para resolver problemas de classificação podem ser encontrados em FREAN (1990), MARCHAND *et al.* (1990) e PAREKH *et al.* (1995).

5.4 Processo de aproximação e treinamento supervisionado

A aproximação de funções utilizando modelos ajustáveis a partir de um conjunto de dados amostrados pode ser considerada um problema de reconstrução de hipersuperfícies (POGGIO & GIROSI, 1990b). Neste sentido, a generalização é a estimativa da altura desta hipersuperfície em pontos não-pertencentes ao conjunto de dados de aproximação. Utilizando redes neurais com uma camada intermediária como modelos de aproximação, processos de

aproximação deste tipo correspondem exatamente ao processo de treinamento supervisionado (POGGIO & GIROSI, 1990a).

Empregando então a nomenclatura associada a treinamento supervisionado, a mesma noção de generalização pode ser vista como a habilidade de prever a ocorrência de um determinado evento baseado na experiência adquirida na fase de treinamento (HECHT-NIELSEN, 1990). Esta experiência decorre da apresentação repetida do conjunto de dados para treinamento, ou seja, o conjunto de dados amostrados da função a ser aproximada, denominados exemplos. GIROSI *et al.* (1991) fazem inclusive a distinção entre exemplos positivos e negativos.

A capacidade de generalização é medida tomando-se como base o desempenho da rede neural treinada na predição de dados pertencentes a um conjunto de teste, dissociado do conjunto de treinamento. Por trás deste processo está logicamente a hipótese de que o conjunto de treinamento é representativo do conjunto de teste, ou seja, que ambos são gerados a partir de uma mesma distribuição, produzida pela função a ser aproximada.

Conclui-se, portanto, que o processo de aproximação promove ajustes junto ao modelo de aproximação com o objetivo de aumentar sua capacidade de generalização. Conforme discutido em capítulos anteriores, estes ajustes do modelo de aproximação estão sujeitos a critérios de regularização. Como exemplo, observe na figura 5.3 o efeito da regularização no desempenho do processo de aproximação em termos de generalização, onde 'x' representa dados pertencentes ao conjunto de treinamento e 'o' representa dados pertencentes ao conjunto de teste.

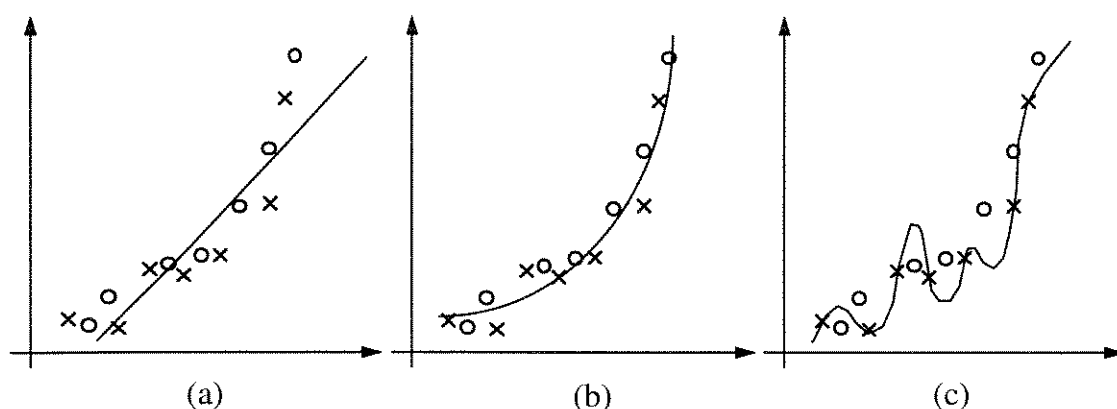


Figura 5.3 - Aproximação para dados sujeitos a ruído, com modelo de aproximação
(a) sobre-regularizado; (b) otimamente regularizado; (c) sub-regularizado

5.5 Generalização e o efeito bias-variância

A capacidade de generalização pode ser tratada sob um ponto de vista predominantemente estatístico, principalmente levando-se em conta que o processo de amostragem está sujeito a ruído e que os dados são gerados a partir de uma função de densidade de probabilidade definida em X (veja capítulo 3). Tomando $g: X \subset \Re^m \rightarrow \Re$ como a função a ser aproximada, considere o conjunto de dados amostrados $\{\mathbf{x}_l, s_l\}_{l=1}^N$ na forma:

$$s_l = g(\mathbf{x}_l) + \varepsilon_l, \quad l=1, \dots, N,$$

onde ε_l é uma variável aleatória de média zero e variância fixa. Observe que novamente é assumida uma única saída ($r = 1$).

Se $E[\cdot]$ é o operador de esperança matemática considerando as N amostras, então $E[s / \mathbf{x}]$ é a melhor aproximação de g em X , o que implica que, para qualquer função de aproximação $\hat{g}$, vale a seguinte relação (GEMAN *et al.*, 1992):

$$E[(s - \hat{g})^2 / \mathbf{x}] \geq E[(s - E[s / \mathbf{x}])^2 / \mathbf{x}].$$

Com isso, um procedimento válido é minimizar a função

$$erro(\hat{g}) = E[(\hat{g} - E[s / \mathbf{x}])^2].$$

A função $erro(\hat{g})$ pode ser dividida em duas partes na forma (GEMAN *et al.*, 1992):

$$erro(\hat{g}) = (E[\hat{g}] - E[s / \mathbf{x}])^2 + E[(\hat{g} - E[\hat{g}])^2], \quad (5.24)$$

onde $E[\hat{g}] - E[s / \mathbf{x}]$ é denominado bias e $E[(\hat{g} - E[\hat{g}])^2]$ é denominado variância da função de aproximação $\hat{g}$.

Portanto, a minimização de $erro(\hat{g})$ representa um compromisso entre os níveis de bias e variância associados à função de aproximação $\hat{g}$, de tal forma que uma função $\hat{g}$ muito flexível geralmente vai apresentar um bias muito baixo, mas uma variância elevada. Por outro lado, uma função $\hat{g}$ pouco flexível geralmente conduz a bias elevado e variância baixa. Desse modo, surgem as seguintes questões: como reduzir o bias sem aumentar proporcionalmente a variância e como reduzir a variância sem aumentar proporcionalmente o bias?

Considerando apenas a informação representada pelo conjunto de dados de aproximação $\{\mathbf{x}_l, s_l\}_{l=1}^N$, é possível definir o erro quadrático médio do problema de aproximação na forma (veja capítulo 3):

$$EQM(\hat{g}) = \frac{1}{N} \sum_{l=1}^N (s_l - \hat{g}(\mathbf{x}_l))^2. \quad (5.25)$$

Um comportamento típico de $erro(\hat{g})$, dado pela equação (5.24), e $EQM(\hat{g})$, dado pela equação (5.25), é apresentado no gráfico da figura 5.4.

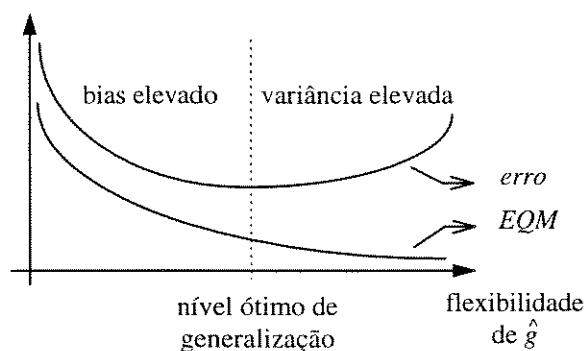


Figura 5.4 - Curvas típicas para $erro(\hat{g})$ e $EQM(\hat{g})$

Tomando $\hat{g}$ como sendo uma rede neural artificial com uma camada intermediária, sua flexibilidade no caso paramétrico é proporcional ao número de ciclos de treinamento¹, enquanto que no caso não-paramétrico é proporcional ao número n de neurônios, tomando funções de ativação fixas, ou ao número n de neurônios e aos correspondentes parâmetros de regularização de cada função de ativação, tomando funções de ativação ajustáveis (veja capítulo 4).

5.5.1 A estimativa do nível de generalização

O dilema entre bias e variância não provém da natureza do problema, mas sim do fato de que o número N de amostras é finito e o modelo de aproximação não é totalmente flexível, mesmo no caso não-paramétrico em estudo.

Certamente, a capacidade de generalização tende a aumentar com o aumento de N (BARRON, 1993). Na verdade, o nível de confiança no modelo de aproximação pode ser

¹ Interromper o processo de treinamento mais cedo pode realmente melhorar a capacidade de generalização em redes neurais sobreparametrizadas, mas não elimina o problema associado ao desperdício de recursos computacionais (MORGAN & BOURLARD, 1990).

baseado na estimativa da diferença de comportamento da função de aproximação com relação aos conjuntos de treinamento e teste, independente do conjunto de teste escolhido.

Para modelos de aproximação lineares, é sabido que os erros de aproximação relativos aos conjuntos de treinamento (ϵ_{tr}) e teste (ϵ_{te}) estão relacionados por (EUBANK, 1988):

$$E[\epsilon_{te}] = E[\epsilon_{tr}] + 2\sigma^2 \frac{p}{N}$$

onde σ^2 é a variância do erro de amostragem, p é o número de parâmetros e N é o número de dados amostrados que compõem o conjunto de treinamento.

Para modelos de aproximação não-lineares, MOODY (1992; 1994) mostrou que uma relação semelhante pode ser obtida utilizando aproximação de segunda ordem:

$$E[\epsilon_{te}(\lambda)] \approx E[\epsilon_{tr}(\lambda)] + 2\sigma_{ef}^2 \frac{p_{ef}(\lambda)}{N} \quad (5.26)$$

onde σ_{ef}^2 é a variância efetiva do erro de amostragem, $p_{ef}(\lambda)$ é o número efetivo de parâmetros e λ é o parâmetro de regularização do modelo de aproximação.

A equação (5.26) permite uma interpretação transparente do comportamento das curvas na figura 5.4 (tome *erro* $\leftrightarrow E[\epsilon_{te}]$ e *EQM* $\leftrightarrow E[\epsilon_{tr}]$): no início do processo de aproximação, um aumento em $p_{ef}(\lambda)$ é capaz de reduzir $E[\epsilon_{te}]$, pois $E[\epsilon_{tr}]$ tem um decréscimo mais acentuado que o crescimento do segundo termo da equação (5.26). Quando o decréscimo de $E[\epsilon_{tr}]$ com um aumento de $p_{ef}(\lambda)$ começa a diminuir conforme o processo vai evoluindo, o segundo termo da equação (5.26) passa a prevalecer, provocando um aumento cada vez mais acentuado de $E[\epsilon_{te}]$.

A equação (5.26) sugere que uma medida de complexidade do modelo de aproximação pode ser utilizada para estimar o nível de generalização. A dimensão de Vapnik-Chervonenkis (VAPNIK & CHERVONENKIS, 1971) é um índice que mede a capacidade computacional do modelo de aproximação, podendo também ser utilizada para avaliar a capacidade de generalização do modelo de aproximação (VAPNIK, 1992), embora a obtenção da dimensão de Vapnik-Chervonenkis seja ainda uma tarefa não-trivial (GHOSH & TUMER, 1994). Para um estudo mais aprofundado a respeito de capacidade de generalização, veja BAUM & HAUSSLER (1989) e COHN & TESAURO (1992).

5.6 Comparação entre a rede neural paramétrica e a não-paramétrica

Os modelos paramétricos e não-paramétricos de redes neurais com uma camada intermediária apresentados neste capítulo podem produzir estruturas do tipo apresentado na figura 5.5. Uma inspeção nos processos de aproximação paramétrico e não-paramétrico permite constatar o seguinte:

- processo de aproximação paramétrico: para n e $f_j(\cdot)$ ($j=1, \dots, n$) predeterminados, determinar os parâmetros v_{ij} ($i=1, \dots, m; j=1, \dots, n$) e w_{jk} ($j=1, \dots, n; k=1, \dots, r$) que resolvam o problema de aproximação;
- processo de aproximação não-paramétrico: determinar n , $f_j(\cdot)$ ($j=1, \dots, n$) e parâmetros v_{ij} ($i=1, \dots, m; j=1, \dots, n$) e w_{jk} ($j=1, \dots, n; k=1, \dots, r$) que resolvam o problema de aproximação.

Portanto, é possível concluir que, no caso de modelos paramétricos, existem duas etapas bem definidas e independentes:

- representação paramétrica: imposição de uma forma específica para a função paramétrica a ser ajustada aos dados;
- otimização paramétrica: ajuste dos parâmetros (definidos na fase de representação) a partir dos dados disponíveis e no sentido de minimizar uma medida do erro de aproximação.

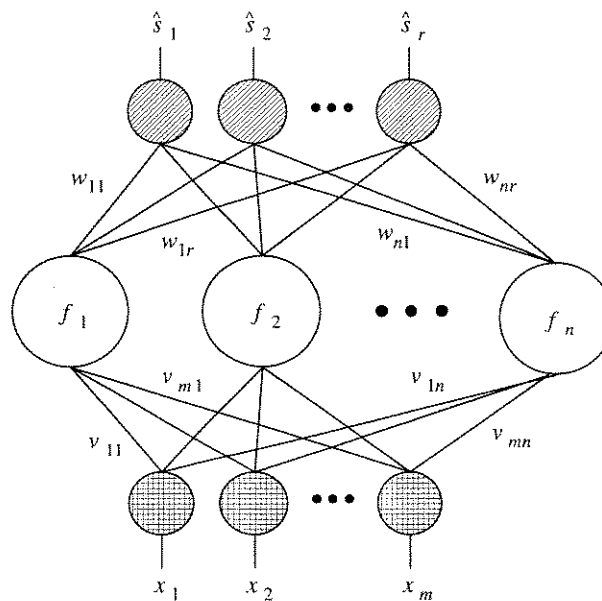


Figura 5.5 - Rede neural genérica com uma camada intermediária

Enquanto isso, no caso de modelos não-paramétricos, estas duas fases também podem ser detectadas, apesar de estarem totalmente interligadas, ou seja, o processo de aproximação promove uma interação de duas fases:

- construção do modelo de aproximação: definição de uma estrutura de representação a partir dos dados;
- otimização do modelo: ajuste paramétrico ou não da estrutura de representação (definida na fase de construção do modelo de representação) também a partir dos dados disponíveis e no sentido de minimizar uma medida do erro de aproximação.

Uma comparação destes métodos com resultados apresentados no capítulo 2, permite concluir que:

- o processo de aproximação paramétrico é a aplicação direta dos resultados do teoremas de Stone-Weierstrass;
- o processo de aproximação não-paramétrico é um procedimento intermediário entre os resultados do teorema de Stone-Weierstrass e do teorema de Kolmogorov e suas generalizações.

Para uma discussão mais aprofundada das conseqüências da aplicação do teorema de Kolmogorov no contexto de redes neurais artificiais, encontram-se na literatura uma série de artigos, dentre os quais é possível destacar: HECHT-NIELSEN (1987), GIROSI & POGGIO (1989), KURKOVÁ (1991), KURKOVÁ (1992), LIN & UNBEHAUEN (1993) e SPRECHER (1993). Com base nos resultados destes artigos, é possível concluir que, restringindo-se a funções de ativação sigmoidais, o teorema de Kolmogorov conduz a redes neurais com duas camadas intermediárias, que, no entanto, podem ser representadas na forma de uma rede neural com uma camada intermediária, desde que os neurônios da camada de entrada tenham também função de ativação sigmoideal e existam múltiplos neurônios associados a cada elemento do vetor de entrada. ITO (1994) demonstrou a capacidade de aproximação universal deste tipo de rede neural, empregando métodos construtivos elementares. Por outro lado, permitindo a adoção de funções de ativação arbitrárias, os resultados do teorema de Kolmogorov conduzem a um problema sem solução conhecida, expresso na forma:

- dado n , determinar $f_j(\cdot)$ ($j=1,\dots,n$) e parâmetros v_{ij} ($i=1,\dots,m; j=1,\dots,n$) e w_{jk} ($j=1,\dots,n; k=1,\dots,r$) que resolvam o problema de aproximação.

Portanto, é possível concluir que o processo de aproximação não-paramétrico apresentado neste estudo, utilizando funções de ativação arbitrárias, representa a *versão factível* do teorema de Kolmogorov para redes neurais não-paramétricas.

É importante mencionar também que tanto a rede neural paramétrica como a não-paramétrica produzem modelos de aproximação que não apenas apresentam um comportamento não-linear de entrada-saída, mas também conduzem a processos de aproximação não-lineares, em contraste com os bem-desenvolvidos métodos de regressão linear, como análise de Fourier e aproximação polinomial. Portanto, os métodos de aproximação de funções utilizando redes neurais artificiais vão além das técnicas de regressão clássicas, representando uma tendência já evidenciada em outro contexto por WERBOS (1974).

Uma consequência imediata é a menor sensibilidade a dados não-informativos ou espúrios apresentada por redes neurais paramétricas ou não-paramétricas, quando comparadas com técnicas de regressão clássicas. Entre os dois tipos de redes neurais, o modelo não-paramétrico é sabido ser menos sensível a dados não-informativos ou espúrios que o modelo paramétrico (HWANG *et al.*, 1992; CHENTOUF & JUTTEN, 1995).

5.7 Extensões

O algoritmo construtivo de aproximação apresentado neste capítulo produz uma rede neural não-paramétrica com uma camada intermediária, um modelo multivariável de redução de dimensionalidade.

O problema de aproximação é, portanto, um problema de otimização em um espaço funcional, onde o estado inicial e a estratégia de busca dependem do procedimento construtivo adotado. Como a busca exaustiva em todo o espaço funcional é computacionalmente proibitiva, o espaço de busca é tomado como um subespaço do espaço funcional, isto é, o subespaço formado pelas redes neurais artificiais com uma camada intermediária. Sendo assim, alguns tipos de problemas de aproximação são favorecidos em detrimento de outros, pois se o procedimento construtivo de aproximação é especialmente projetado para buscar soluções em um subespaço do espaço funcional, ele não pode ser considerado o mais eficiente em todo o espaço funcional (FRIEDMAN, 1994).

5.7.1 Função de ativação de expansão radial

Redes neurais artificiais com uma camada intermediária composta por neurônios dotados de funções de ativação de expansão radial são denominadas de redes neurais de base radial e também são utilizadas como modelos de aproximação. Este modelo de rede neural também pode ser prontamente deduzido a partir de resultados da teoria de regularização (POGGIO & GIROSI, 1990a).

Quando o argumento da função de expansão radial é deduzido com base em uma norma ponderada, uma função de expansão ortogonal pode ser considerada como o limite de uma função radial quando um dos eixos da hiperelipse de expansão radial tende a infinito e todos os outros tendem a zero. Mesmo que a expansão radial se dê com base em uma norma ponderada, como esta ponderação não é ajustável, não é possível selecionar de forma automática o melhor conjunto de variáveis de entrada que contribuem efetivamente para a aproximação em cada passo, ou seja, a rede neural com funções de ativação de expansão radial não é um modelo de redução de dimensionalidade, sendo geralmente aplicada a problemas de aproximação de baixa dimensionalidade (GIROSI, 1992).

Mesmo estando fortemente susceptíveis à maldição da dimensionalidade, é possível demonstrar a capacidade de aproximação universal de redes neurais de base radial (HARTMAN *et al.*, 1990; PARK & SANDBERG, 1991), desenvolver métodos construtivos de aproximação (DINGANKAR, 1995; LEE & KIL, 1991) e obter taxas de convergência (DINGANKAR, 1995).

5.7.2 O uso de mais de uma camada intermediária

O fato de terem sido estudados neste capítulo apenas modelos de redes neurais com uma camada intermediária merece dois comentários adicionais:

- a restrição a uma camada intermediária se justifica pela existência de propriedades de consistência (capacidade de aproximação universal) e pela possibilidade de estimar a taxa de convergência do processo de aproximação em função do número de neurônios na camada intermediária. Estes resultados não estão disponíveis considerando-se múltiplas camadas intermediárias;
- um modelo não-paramétrico de rede neural com uma camada intermediária, que utiliza funções de ativação suavizantes determinadas a partir do conjunto de dados de aproximação e, portanto, não restritas à forma sigmoideal, pode ser teoricamente

representado por uma rede neural completamente equivalente com duas camadas intermediárias que utiliza funções de ativação sigmoidal. Esta afirmação é imediatamente verificada observando-se que cada uma das funções de ativação suavizantes pode ser representada por uma rede neural com uma camada intermediária que utilizam funções de ativação sigmoidal (VERKOOIJEN & DANIELS, 1994). O mesmo resultado pode ser deduzido utilizando-se funções de ativação de expansão radial em lugar de funções sigmoidais (SAHA *et al.*, 1993).

5.7.3 Modelos multiplicativos com múltiplas saídas

Conforme já discutido no capítulo anterior por ocasião da apresentação de modelos de projeção aditivos, o processo de aproximação para modelos aditivos pode ser facilmente adaptado para produzir modelos de projeção multiplicativos. Como as redes neurais estudadas neste capítulo correspondem a modelos aditivos, as mesmas adaptações continuam a ser válidas. Novamente, analogias com modelos desenvolvidos com base em conceitos estatísticos, tais como PPDS ("Projection Pursuit Density eStimation") (FRIEDMAN *et al.*, 1984), MARS ("Multivariate Adaptive Regression Splines") (FRIEDMAN, 1991b) e CART ("Classification and Regression Tree") (BREIMAN *et al.*, 1984) podem ser exploradas. Conforme descrito por FRIEDMAN (1991a), modelos baseados em composição multiplicativa são mais eficientes na aproximação de mapeamentos caracterizados por complexidades locais.

Capítulo 6

Redes Neurais Recorrentes

Redes neurais recorrentes são estruturas de processamento capazes de representar uma grande variedade de comportamentos dinâmicos. A presença de realimentação de informação permite a criação de representações internas e dispositivos de memória capazes de processar e armazenar informações temporais e sinais sequenciais. Por exemplo, vários tipos de redes recorrentes são capazes de inferir gramáticas regulares a partir de exemplos (CLEEREMANS *et al.*, 1989; GILES *et al.*, 1992).

Como um resultado bem conhecido da teoria de processamento de sinais, a presença de conexões recorrentes ou realimentação de informação em sistemas de processamento pode conduzir a comportamentos complexos, mesmo com um número reduzido de parâmetros (LJUNG, 1987; OPPENHEIM & SCHAFER, 1989). Como estruturas de processamento de sinais, redes neurais recorrentes se assemelham a filtros não-lineares com resposta ao impulso infinita (NERRAND *et al.*, 1993).

Em contrapartida à possibilidade de representação de comportamento dinâmico de tempo discreto ou contínuo, o processo de adaptação e a análise da capacidade de processamento de informação presente em redes neurais recorrentes são geralmente mais complexos que no caso não-recorrente (WILLIAMS & ZIPSER, 1989b).

Este estudo se restringe ao tratamento da dinâmica discreta, utilizando o operador de atraso de tempo. Neste caso, redes neurais paramétricas parcialmente recorrentes são apresentadas como modelos capazes de representar dinâmicas discretas não-lineares genéricas (VON ZUBEN & NETTO, 1994a), o processo de adaptação de parâmetros é desenvolvido com base em conceitos de análise de sensibilidade (CRUZ JR. & PERKINS, 1964), e a análise da capacidade de processamento de informação é fundamentada em conceitos de teoria de funções iterativas (BARNESLEY, 1988).

6.1 Arquiteturas básicas

As redes neurais recorrentes podem ser classificadas como totalmente ou parcialmente recorrentes. Redes totalmente recorrentes possuem todo tipo de conexão recorrente, todas ajustáveis. Quando apenas uma parte das possíveis conexões recorrentes é admitida, ou então quando as conexões recorrentes não são ajustáveis, resultam redes parcialmente recorrentes.

Dentre as arquiteturas de redes neurais recorrentes, a rede de Elman é a mais simples (ELMAN, 1990). Um diagrama de blocos da rede de Elman é apresentado na figura 6.1.

Observando-se a figura 6.1, pode-se constatar que a rede de Elman, além das camadas de entrada, intermediária e de saída, possui uma camada de contexto. Os neurônios pertencentes às camadas de entrada e saída interagem com o ambiente externo, enquanto que os neurônios das camadas intermediária e de contexto não o fazem. Além disso, é possível definir:

- camada de entrada: composta por neurônios sensoriais que recebem um sinal externo e o propagam sem alterá-lo;
- camada de saída: composta por neurônios lineares cujas saídas são somas ponderadas de seus respectivos sinais de entrada;
- camada intermediária: composta por neurônios que apresentam funções de ativação lineares ou não-lineares, dependendo do tipo de aplicação;
- camada de contexto: composta por neurônios que apenas memorizam as ativações dos neurônios da camada intermediária, operando como atrasadores de um instante de tempo.

A intensidade de cada conexão é indicada por um valor numérico denominado peso. As conexões não-recorrentes são ajustáveis, enquanto que as conexões recorrentes são fixas, fazendo com que a rede de Elman seja uma rede neural parcialmente recorrente.

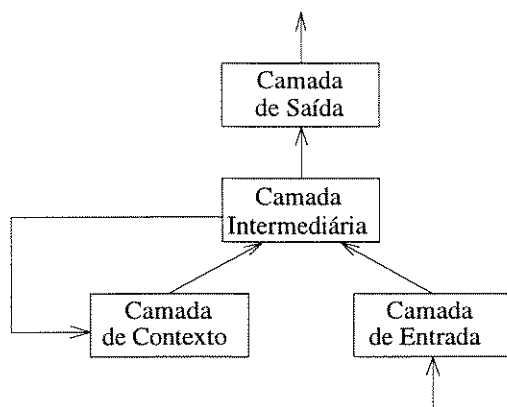


Figura 6.1 - Diagrama de blocos da rede de Elman

Em um instante de tempo específico t , os sinais de ativação dos neurônios da camada de contexto, juntamente com os sinais de ativação dos neurônios sensoriais da camada de entrada, formam o vetor de entrada para a camada intermediária. Neste estágio de processamento, a rede é puramente não-recorrente, sendo que estes sinais são propagados de forma a produzir o vetor de saída no instante t . Após a propagação dos sinais até a saída, a ativação atual dos neurônios da camada intermediária é armazenada pelos neurônios da camada de contexto através das conexões recorrentes, permitindo a repetição do ciclo de processamento, agora para o instante $t+1$.

No início da operação, a ativação dos neurônios da camada intermediária não é conhecida. Com isso, um procedimento usual é tomar o intervalo em que a função de ativação de cada um destes neurônios assume valores e adotar como ativação inicial o valor médio do respectivo intervalo.

A arquitetura detalhada da rede de Elman é apresentada na figura 6.2.

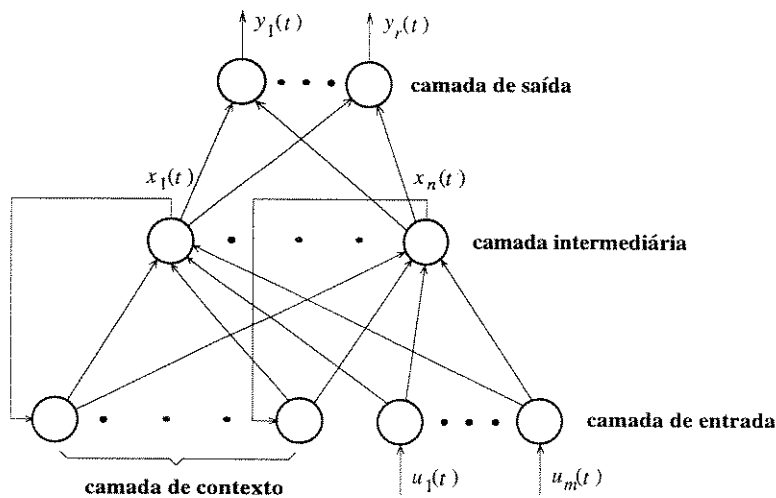


Figura 6.2 - Estrutura detalhada da rede de Elman

6.1.1 Modelagem de sistemas dinâmicos lineares

A rede de Elman linear mostra-se adequada para a representação de sistemas dinâmicos lineares invariantes no tempo (PHAM & LIU, 1992), conforme apresentado a seguir.

Analisando-se o processamento de sinais na rede da figura 6.2, pode-se concluir que:

$$\begin{cases} \mathbf{x}(t) = \mathbf{f}(\mathbf{W}_{xx}\mathbf{x}(t-1), \mathbf{W}_{xu}\mathbf{u}(t)) \\ \mathbf{y}(t) = \mathbf{W}_{yx}\mathbf{x}(t) \end{cases} \quad (6.1)$$

onde

- $\mathbf{W}_{xx} \in \mathbb{R}^{n \times n}$, $\mathbf{W}_{xu} \in \mathbb{R}^{n \times m}$ e $\mathbf{W}_{yx} \in \mathbb{R}^{r \times n}$ são matrizes de pesos;
- n é o número de neurônios nas camadas intermediária e de contexto, m é o número de entradas e r é o número de saídas;
- $\mathbf{f}: \mathbb{R}^n \times \mathbb{R}^m \rightarrow \mathbb{R}^n$ é uma função vetorial e de argumento vetorial.

Em particular, se as funções de ativação dos neurônios da camada intermediária são lineares, a equação (6.1) assume a forma:

$$\begin{cases} \mathbf{x}(t) = \mathbf{W}_{xx}\mathbf{x}(t-1) + \mathbf{W}_{xu}\mathbf{u}(t) \\ \mathbf{y}(t) = \mathbf{W}_{yx}\mathbf{x}(t) \end{cases} \quad (6.2)$$

Quando o vetor de entrada $\mathbf{u}(t)$ é defasado um instante de tempo na camada de entrada, resulta:

$$\begin{cases} \mathbf{x}(t) = \mathbf{W}_{xx}\mathbf{x}(t-1) + \mathbf{W}_{xu}\mathbf{u}(t-1) \\ \mathbf{y}(t) = \mathbf{W}_{yx}\mathbf{x}(t) \end{cases} \quad (6.3)$$

A equação (6.3) corresponde à representação geral de tempo discreto de sistemas dinâmicos lineares por espaço de estados, assumindo que o vetor de entrada $\mathbf{u}(t)$ não participa diretamente na definição do vetor de saída $\mathbf{y}(t)$.

A ordem do modelo vai depender, portanto, do número de neurônios na camada intermediária da rede de Elman. Assim, uma rede de Elman linear pode ser capaz de modelar um sistema dinâmico linear de qualquer ordem, desde que as matrizes de pesos possam ser adequadamente ajustadas. Exemplos de simulação ilustrando esta propriedade são apresentados por VON ZUBEN (1993).

6.1.2 Extensão para o caso não-linear

Permitindo a definição de funções de ativação não-lineares para os neurônios da camada intermediária, a rede de Elman apresentada na figura 6.2 também pode aproximar mapeamentos não-lineares. Se estas funções de ativação forem lineares, os resultados da seção anterior mostram que existe uma analogia perfeita com representações de sistemas dinâmicos lineares por espaço de estados. No entanto, esta analogia não se estende ao caso não-linear.

Apesar da capacidade de aproximar alguns tipos de mapeamentos não-lineares, redes de Elman cujos neurônios da camada intermediária apresentam funções de ativação não-lineares arbitrárias não podem ser descritas na forma geral de representação por espaço de estados de sistemas dinâmicos não-lineares (veja equação (6.4) a seguir). Isto ocorre devido à

impossibilidade de aumentar a flexibilidade da rede de Elman pela adição de mais neurônios na camada intermediária sem aumentar simultaneamente o número de estados. Com isso, é sempre possível encontrar um sistema dinâmico não-linear cujo comportamento dinâmico não possa ser identificado por uma rede de Elman não-linear de mesma ordem (KOLEN, 1994). Esta limitação é compreensível, já que a rede de Elman representa uma das arquiteturas mais simples de rede neural recorrente. A alternativa apresentada na próxima seção está baseada na desvinculação da flexibilidade do modelo de representação em relação ao número de estados.

6.1.3 Modelagem de sistemas dinâmicos não-lineares

Nesta seção é desenvolvida uma nova arquitetura de rede recorrente capaz de realizar a modelagem de sistemas dinâmicos não-lineares de tempo discreto. Para tanto, é suficiente a obtenção de um modelo de rede neural cujo comportamento dinâmico possa ser descrito na forma geral de representação por espaço de estados de sistemas dinâmicos não-lineares, conforme apresentado por VON ZUBEN & NETTO (1994a, 1994b).

A representação por espaço de estados mais geral de sistemas dinâmicos não-lineares de tempo discreto é dada por:

$$\begin{cases} \mathbf{x}(t+1) = f(\mathbf{x}(t), \mathbf{u}(t)) \\ \mathbf{y}(t) = g(\mathbf{x}(t), \mathbf{u}(t)) \end{cases} \tag{6.4}$$

onde $\mathbf{u}(t) \in \mathbb{R}^m$, $\mathbf{x}(t) \in \mathbb{R}^n$, $\mathbf{y}(t) \in \mathbb{R}^r$, $f: \mathbb{R}^{n \times m} \rightarrow \mathbb{R}^n$ e $g: \mathbb{R}^{n \times m} \rightarrow \mathbb{R}^r$.

Expressando a equação (6.4) na forma de um diagrama de fluxo de sinais, resulta o esquema da figura 6.3. Em princípio, os modelos de redes neurais desenvolvidos no capítulo anterior podem então ser utilizados para aproximar as funções f e g . No entanto, conforme apresentado por VON ZUBEN (1993), a presença de conexões recorrentes requer a adoção de modificações com relação ao processo de aproximação, ao menos com relação à função f .

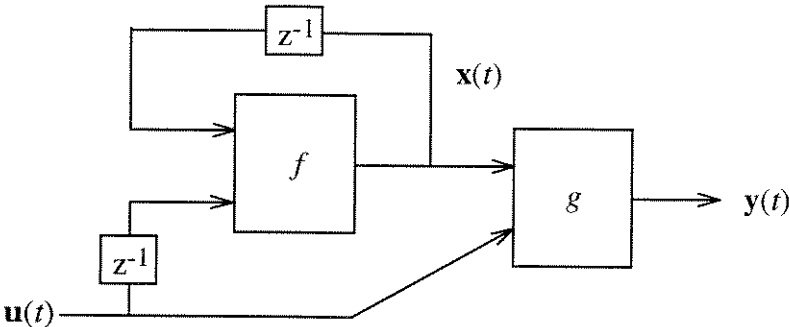


Figura 6.3 - Representação por espaço de estados de um sistema dinâmico não-linear

6.2 O efeito da recorrência no processo de adaptação

Redes neurais não-recorrentes podem ser interpretadas como poderosos operadores de transformação de representação, mas não são capazes de reutilizar a informação transformada, produzindo apenas mapeamentos estáticos. Esta é a razão de este tipo de rede neural encontrar dificuldade em representar comportamento dinâmico, já que o vetor de saída da rede neural, denominado $\hat{s}(t)$, depende apenas do vetor de entrada definido no mesmo instante, denominado $\mathbf{x}(t)$. Isto conduz a mapeamentos do tipo:

$$\hat{s}(t) = RN(\mathbf{x}(t), \theta(t)), \quad (6.5)$$

onde $\theta(t)$ denota o vetor de parâmetros no instante t , sendo de dimensão fixa no caso de redes neurais paramétricas e de dimensão variável no caso de redes neurais não-paramétricas. Como será visto no capítulo 7, a representação de comportamento dinâmico utilizando a estrutura da equação (6.5) se dá por representação espacial do tempo, ou seja, pela adição de componentes, junto ao vetor de entrada $\mathbf{x}(t)$, que expressem comportamento temporal.

Por outro lado, em redes neurais recorrentes o tempo é representado pelo seu efeito real no processamento. Assumindo apenas a existência de recorrência externa, ou seja, realimentação apenas da informação de saída da rede neural, resultam modelos do tipo:

$$\hat{s}(t) = RN_{rec}(\mathbf{x}(t), \hat{s}(t-1), \theta(t)). \quad (6.6)$$

Com isso, por substituição recursiva da equação (6.6) nela mesma, fica claro que o vetor de saída da rede neural depende da história do vetor de entrada e do vetor de parâmetros.

A motivação para restringir a análise ao caso de redes neurais com recorrência externa é a analogia com as redes não-recorrentes paramétricas com uma camada intermediária apresentadas no capítulo 5 (veja figura 5.2) e a constatação, por simples inspeção da figura 6.3, que recorrências externas são suficientes para possibilitar a representação de qualquer comportamento dinâmico de tempo discreto (veja também SIEGELMANN & SONTAG (1992) e NERRAND *et al.* (1993)). Resulta, então, uma rede neural paramétrica com recorrência externa, apresentada na figura 6.4.

Uma consequência prática deste processo de realimentação externa é a maior dificuldade de aplicação de métodos construtivos para redes neurais recorrentes. Não existe, por exemplo, uma formulação prática aceitável para a definição de que tipo de estrutura ainda não foi representada pelo modelo, passo fundamental para a aplicação do método construtivo desenvolvido no capítulo anterior.

Logo, métodos construtivos que conduzam a uma arquitetura ótima para redes neurais recorrentes só poderão ser viabilizados a partir do momento em que for possível avaliar a real influência que a introdução de determinado tipo de estrutura vai exercer junto ao comportamento dinâmico global do modelo, um resultado análogo ao princípio da superposição de efeitos.

O processo construtivo para redes neurais recorrentes apresentado por FAHLMAN (1991) e suas generalizações (LEE GILES *et al.*, 1995) não atendem a estas condições. Por exemplo, um problema crucial do algoritmo de LEE GILES *et al.* (1995) é a ausência da etapa de retroajuste (veja capítulos 4 e 5).

Esta é a razão de se utilizar neste estudo apenas redes neurais recorrentes paramétricas como modelos de representação de comportamento dinâmico.

6.3 Treinamento supervisionado para redes neurais recorrentes

A disponibilidade de redes neurais recorrentes paramétricas de importância prática está associada à existência de algoritmos de otimização eficientes para o ajuste de parâmetros do processo de aproximação resultante, denominado treinamento supervisionado.

Como o modelo de rede neural não-recorrente paramétrica apresentado no capítulo 5 (veja figura 5.2) é um caso particular do modelo de rede neural recorrente paramétrica com recorrência externa (veja figura 6.4), justifica-se a iniciativa de deixar para este capítulo a descrição do processo de treinamento supervisionado para redes neurais paramétricas, tanto recorrentes como não-recorrentes.

Para o caso paramétrico, é possível substituir $\theta(t)$ por θ nas equações (6.5) e (6.6). Com isso, fica claro que, ao contrário do modelo de rede neural não-recorrente, o modelo de rede neural recorrente é uma função composta em θ . Logo, a análise variacional dos dois modelos produz os seguintes resultados:

- rede neural não-recorrente:

$$\hat{s}(t) = RN(\mathbf{x}(t), \theta)$$

$$\frac{\partial \hat{s}(t)}{\partial \theta} = \frac{\partial RN}{\partial \theta}$$
- rede neural recorrente:

$$\hat{s}(t) = RN_{rec}(\mathbf{x}(t), \hat{s}(t-1), \theta)$$

$$\frac{\partial \hat{s}(t)}{\partial \theta} = \frac{\partial RN_{rec}}{\partial \theta} = \frac{\partial RN_{rec}}{\partial \theta} + \underbrace{\frac{\partial RN_{rec}}{\partial \hat{s}(t-1)} \frac{\partial \hat{s}(t-1)}{\partial \theta}}_{\text{termo adicional}}$$

Baseado nas hipóteses levantadas no capítulo 3 com relação a problemas de aproximação a partir de dados amostrados da função a ser aproximada, o problema de modelagem de sistemas dinâmicos de tempo discreto utilizando redes neurais recorrentes paramétricas pode ser colocado na forma:

- Seja $\{(\mathbf{x}(t), \mathbf{s}(t))\}_{t=1}^N \in X \times \mathfrak{R}^r$, onde $X \subset \mathfrak{R}^m$, um conjunto de dados amostrados de um sistema dinâmico de tempo discreto a ser identificado (o capítulo 7 apresenta a formalização necessária para se chegar a este conjunto de dados);
- Seja $RN_{rec}(\mathbf{x}(t), \hat{\mathbf{s}}(t-1), \theta) : X \times \mathfrak{R}^r \times \mathfrak{R}^P \rightarrow \mathfrak{R}^r$ o mapeamento realizado pela rede neural recorrente, onde $\theta \in \mathfrak{R}^P$ é um vetor contendo todos os P parâmetros da rede neural, ordenados de forma arbitrária, mas fixa;
- Resolva o seguinte problema de otimização:

$$\min_{\theta} J(\theta) = \min_{\theta} \frac{1}{2} \sum_{t=1}^N \|RN_{rec}(\mathbf{x}(t), \hat{\mathbf{s}}(t-1), \theta) - \mathbf{s}(t)\|^2, \quad (6.7)$$

onde $\|\cdot\|$ é a norma euclidiana. Vale ressaltar que critérios de regularização (específicos para o caso paramétrico) podem ser prontamente empregados.

A solução do problema de otimização não-linear (6.7) só pode ser obtida via processos numéricos iterativos, já que soluções analíticas não estão disponíveis, principalmente devido à natureza não-linear do modelo de aproximação $RN_{rec}(\mathbf{x}(t), \hat{\mathbf{s}}(t-1), \theta)$. A idéia básica deste processo numérico é realizar ajustes iterativos no vetor de parâmetros $\theta \in \mathfrak{R}^P$ sempre em direções em que a função objetivo $J(\theta)$ decresça. Observe que $J(\theta)$ define uma superfície no espaço de parâmetros $\mathfrak{R}^P$, ou seja, a dimensão do espaço de busca é igual ao número de parâmetros a serem ajustados.

O processo iterativo obedece a uma lei de adaptação definida na forma:

$$\theta_{i+1} = \theta_i + \alpha_i \mathbf{d}_i, \quad (6.8)$$

onde θ_i é o vetor de parâmetros, α_i é um escalar que define o passo de ajuste e $\mathbf{d}_i$ é a direção de ajuste, todos definidos na iteração i . Processos iterativos deste tipo vêm ganhando cada vez mais importância prática devido à crescente disponibilidade de recursos computacionais.

É importante distinguir dois tipos de situações:

- ajuste em batelada: para a i -ésima iteração, a direção de ajuste $\mathbf{d}_i$ é definida com base em todo o conjunto de dados $\{(\mathbf{x}(t), \mathbf{s}(t))\}_{t=1}^N$;

- ajuste recursivo ou sequencial: para a i -ésima iteração, a direção de ajuste $\mathbf{d}_i$ é definida com base no dado $(\mathbf{x}(t), \mathbf{s}(t))$, com $t = i \bmod N$.

Uma análise das vantagens e desvantagens destes dois modos de implementação do ajuste iterativo, incluindo exemplos de simulação, é feita por VON ZUBEN (1993). Vale salientar ainda que o método de ajuste recursivo ou sequencial é uma forma de aproximação análoga ao método de Robbins-Monro (ROBBINS & MONRO, 1951; WHITE, 1989).

Neste estudo, é sempre utilizado o ajuste em batelada, sendo que algoritmos iterativos para a determinação de α_i e $\mathbf{d}_i$ são desenvolvidos no apêndice D, assumindo que a função $RN_{rec}(\mathbf{x}(t), \hat{\mathbf{s}}(t-1), \theta)$ é diferenciável pelo menos até 2ª ordem.

Os resultados apresentados no apêndice D sugerem que a solução iterativa do problema de otimização (6.7) seja obtida pelo método do gradiente conjugado, devidamente adaptado para tratar a otimização de funções não-quadráticas. Além de estar entre os mais eficientes métodos de otimização utilizando diferenciação, o custo computacional associado à aplicação do método do gradiente conjugado é mostrado ser apenas duas vezes maior que o custo associado aos métodos de 1ª ordem, comprovadamente menos eficientes. Exemplos e comparações entre métodos de otimização para solução de problemas envolvendo treinamento de redes neurais paramétricas, restritos ao caso não-recorrente, podem ser encontrados em BATTTI (1992) e VAN DER SMAGT (1994).

Um aspecto fundamental da formulação apresentada acima é a flexibilidade de notação, permitindo aplicar o algoritmo, sem modificações, tanto para redes neurais recorrentes como não-recorrentes. A análise no contexto de funções compostas dispensa a aplicação do algoritmo de retropropagação no tempo (WERBOS, 1990), que exige o uso do conceito de derivada ordenada ou então, em uma forma aproximada, requer a criação de réplicas da rede neural a cada iteração (MCCLELLAND & RUMELHART, 1986).

6.3.1 As propriedades da função objetivo

Como as aplicações de redes neurais paramétricas geralmente envolvem um número elevado de parâmetros, a função objetivo (ou superfície de erro) $J: \mathbb{R}^P \rightarrow \mathbb{R}$, por estar definida em um espaço de dimensão igual ao número de parâmetros, é de difícil análise. CHEN *et al.* (1993) e COETZEE & STONICK (1995) apresentam um estudo abrangente desta superfície de erro, donde se destacam as seguintes propriedades básicas:

- se as funções de ativação forem não-lineares, a função J pode apresentar mínimos locais.

- se as funções de ativação dos neurônios são contínuas e diferenciáveis até ordem n , então J também o será;

Como os métodos de otimização de importância prática se restringem ao uso de informações locais a respeito da função J , então, sempre que J não for uma função convexa, a solução do problema de otimização (6.7) pode depender da condição inicial θ_0 (LUENBERGER, 1969). Iniciativas voltadas para aumentar a probabilidade de se atingir o mínimo global que atenda inclusive a critérios de regularização são importantes para a obtenção do melhor modelo de rede neural. Neste sentido, é recomendada a aplicação do algoritmo de otimização complementado pelos seguintes procedimentos:

- determinação de melhores candidatos para condição inicial (BABA *et al.*, 1994);
- utilização de várias condições iniciais distintas (BALAKRISHNAN & HONAVAR, 1995);
- utilização da técnica de recozimento simulado (*simulated annealing*) (KIRKPATRICK *et al.*, 1983; STYBLINSKI & TANG, 1990).

6.3.2 A geração da informação de 1ª e 2ª ordem

Os resultados desta seção, assim como as aplicações apresentadas no capítulo 7, embora se concentrem em redes neurais com uma camada intermediária, são diretamente extensíveis a múltiplas camadas intermediárias. Além disso, os algoritmos de geração da informação de 1ª e 2ª ordem, a serem apresentados a seguir, são válidos tanto para redes neurais multicamadas recorrentes como não-recorrentes, já que a rede não-recorrente é um caso particular da rede recorrente em estudo.

Considere a rede neural com uma camada intermediária e com recorrência externa apresentada na figura 6.4, onde o índice t denota instante discreto de tempo. O vetor de entrada desta rede neural é dado por

$$\mathbf{x}(t) = [u_1^t \cdots u_m^t]^T, \quad (6.9)$$

o vetor de saída por

$$\hat{\mathbf{s}}(t) = [x_1^{t+1} \cdots x_r^{t+1}]^T, \quad (6.10)$$

e o processamento da rede neural assume a forma:

$$x_k^{t+1} = \sum_{j=0}^n w_{kj}^o y_j^t = \sum_{j=1}^n [w_{kj}^o f(s_j^t)] + w_{k0}^o = \sum_{j=1}^n \left[w_{kj}^o f \left(\sum_{i=1}^m w_{ji}^e u_i^t + w_{j0}^e + \underbrace{\sum_{i=m+1}^{m+r} w_{ji}^e x_{i-m}^t}_{\text{termo adicional}} \right) \right] + w_{k0}^o. \quad (6.11)$$

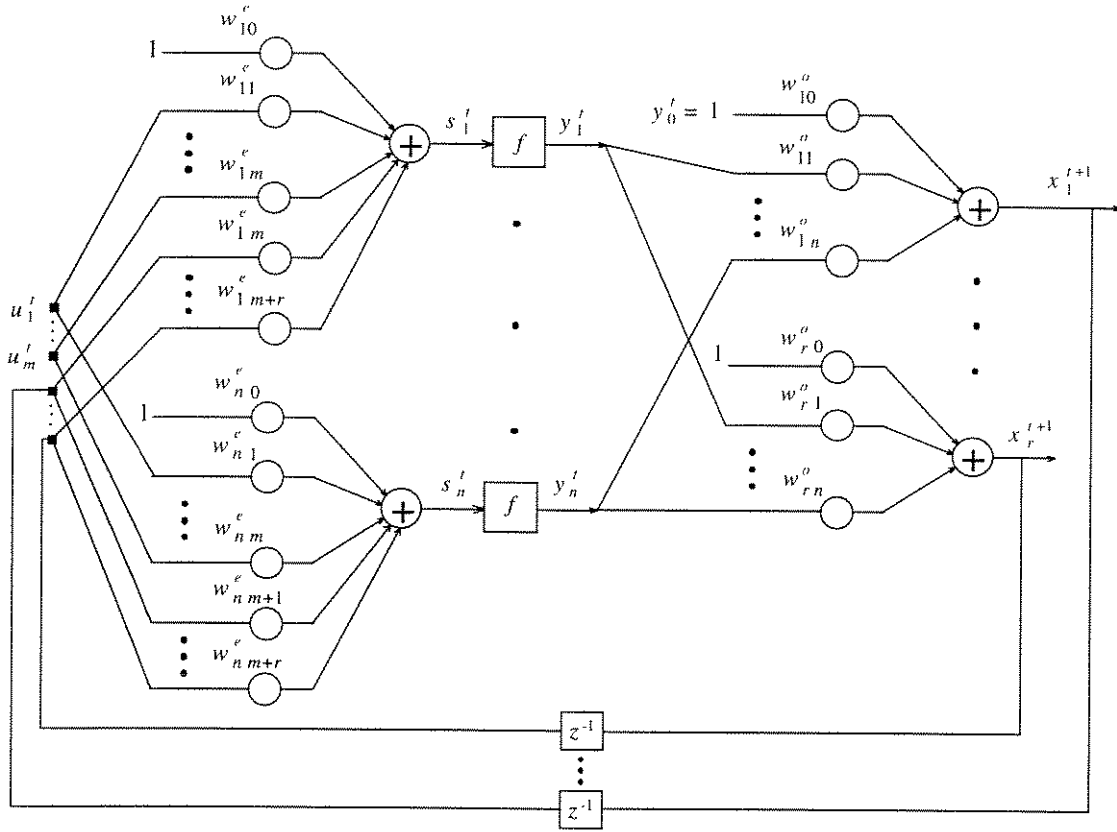


Figura 6.4 - Rede neural recorrente com uma camada intermediária

A presença de fluxo de informação recorrente provocou apenas a introdução do termo adicional indicado, responsável direto pela representação de comportamento dinâmico. Observe também que a rede neural da figura 6.4 tem todos os neurônios da camada intermediária com a mesma função de ativação f .

Os parâmetros w_{ji}^e ($j=1, \dots, n$; $i=0, \dots, m+r$) e w_{kj}^o ($k=1, \dots, r$; $j=0, \dots, n$) são arranjados em um vetor de parâmetros $\theta \in \mathfrak{R}^P$, com $P = n(m+r+1) + r(n+1)$, na forma:

$$\theta = [w_{10}^e \cdots w_{1(m+r)}^e \cdots w_{n0}^e \cdots w_{n(m+r)}^e \ w_{10}^o \cdots w_{1n}^o \cdots w_{r0}^o \cdots w_{rn}^o]^T, \quad (6.12)$$

sendo que a política de ordenamento adotada pode ser arbitrária, desde que seja preservada em toda a implementação. O mapeamento multidimensional realizado pela rede neural da figura 6.4 pode, então, ser expresso por:

$$\hat{\mathbf{s}}(t) = RN_{rec}(\mathbf{x}(t), \hat{\mathbf{s}}(t-1), \theta), \quad (6.13)$$

com $RN_{rec}(\cdot, \cdot, \cdot): X \times \mathfrak{R}^r \times \mathfrak{R}^P \rightarrow \mathfrak{R}^r$, onde $X \subset \mathfrak{R}^m$. Observe que, como em todo sistema dinâmico, é necessário definir uma condição inicial $CI = [x_1^0 \cdots x_r^0]^T$ para o processamento da rede neural.

6.3.2.1 O cálculo do vetor gradiente via algoritmo de retropropagação dinâmica

Com base no conjunto de dados $\{(\mathbf{x}(t), \mathbf{s}(t))\}_{t=1}^N$ e na equação (6.13), a função objetivo $J(\theta)$, definida na equação (6.7), pode ser reescrita na forma:

$$J(\theta) = \frac{1}{2} \sum_{t=1}^N \|\hat{\mathbf{s}}(t) - \mathbf{s}(t)\|^2 = \frac{1}{2} \sum_{t=1}^N \sum_{k=1}^r (x_k^{t+1} - s_k(t))^2. \quad (6.14)$$

Com isso, o vetor gradiente de $J(\theta)$ é dado por:

$$\frac{\partial J(\theta)}{\partial \theta} = \sum_{t=1}^N \sum_{k=1}^r (x_k^{t+1} - s_k(t)) \frac{\partial x_k^{t+1}}{\partial \theta}, \quad (6.15)$$

onde

$$\frac{\partial x_k^{t+1}}{\partial \theta} = \left[\frac{\partial x_k^{t+1}}{\partial w_{10}^e} \dots \frac{\partial x_k^{t+1}}{\partial w_{l(m+r)}^e} \dots \frac{\partial x_k^{t+1}}{\partial w_{n0}^e} \dots \frac{\partial x_k^{t+1}}{\partial w_{n(m+r)}^e} \frac{\partial x_k^{t+1}}{\partial w_{10}^o} \dots \frac{\partial x_k^{t+1}}{\partial w_{1n}^o} \dots \frac{\partial x_k^{t+1}}{\partial w_{r0}^o} \dots \frac{\partial x_k^{t+1}}{\partial w_{rn}^o} \right]^T. \quad (6.16)$$

Para w_{kj}^o , $k=1, \dots, r$, $j=0, \dots, n$:

- $\frac{\partial x_k^{t+1}}{\partial w_{kj}^o} = y_j^t + \sum_{a=1}^n w_{ka}^o \frac{\partial y_a^t}{\partial w_{kj}^o}; \quad (6.17)$

- Para $k_1, k_2=1, \dots, r$ ($k_1 \neq k_2$): $\frac{\partial x_{k_1}^{t+1}}{\partial w_{k_2j}^o} = \sum_{a=1}^n w_{k_1a}^o \frac{\partial y_a^t}{\partial w_{k_2j}^o}; \quad (6.18)$

- $\frac{\partial y_a^t}{\partial w_{kj}^o} = \frac{\partial y_a^t}{\partial s_a^t} \frac{\partial s_a^t}{\partial w_{kj}^o} = f'(s_a^t) \sum_{i=m+1}^{m+r} w_{ai}^e \frac{\partial x_{i-m}^t}{\partial w_{kj}^o}. \quad (6.19)$

Para w_{ji}^e , $j=1, \dots, n$, $i=0, \dots, m+r$:

- $\frac{\partial x_k^t}{\partial w_{ji}^e} = \sum_{a=1}^n w_{ka}^o \frac{\partial y_a^t}{\partial w_{ji}^e}; \quad (6.20)$

- $\frac{\partial y_a^t}{\partial w_{ji}^e} = \frac{\partial y_a^t}{\partial s_a^t} \frac{\partial s_a^t}{\partial w_{ji}^e} = f'(s_a^t) \cdot \begin{cases} \sum_{b=m+1}^{m+r} w_{ab}^e \frac{\partial x_{b-m}^t}{\partial w_{ji}^e} & \text{se } a \neq j \\ u_i^t + \sum_{b=m+1}^{m+r} w_{ab}^e \frac{\partial x_{b-m}^t}{\partial w_{ji}^e} & \text{se } a = j \text{ e } 0 \leq i \leq m \\ x_{i-m}^t + \sum_{b=m+1}^{m+r} w_{ab}^e \frac{\partial x_{b-m}^t}{\partial w_{ji}^e} & \text{se } a = j \text{ e } (m+1) \leq i \leq (m+r) \end{cases} \quad (6.21)$

Fazendo as substituições, em seqüência, das equações (6.16) a (6.21) na equação (6.15), resulta o vetor gradiente $\frac{\partial J(\theta)}{\partial \theta}$, avaliado exatamente. Este é conhecido algoritmo de retropropagação (RUMELHART *et al.*, 1986), baseado na aplicação sucessiva da regra da cadeia. Por estar adaptado, no caso em estudo, para tratar funções de tempo discreto, originadas a partir de redes neurais com recorrência externa, é mais conveniente denominá-lo algoritmo de retropropagação dinâmica, seguindo a notação de NARENDRA & PARTHASARATHY (1990; 1991).

Note que o vetor gradiente também poderia ser obtido por algoritmos de propagação direta (WILLIAMS & ZIPSER, 1989a), de forma análoga aos algoritmos empregados em filtragem adaptativa. No entanto, como demonstrado por PINEDA (1989), a retropropagação é menos custosa computacionalmente que a propagação direta, sendo que os dois procedimentos conduzem exatamente aos mesmos resultados numéricos.

6.3.2.2 O cálculo do produto da matriz hessiana por um vetor

De acordo com os resultados do apêndice D, uma propriedade fundamental do método do gradiente conjugado é a ausência de necessidade de se determinar e armazenar a matriz hessiana elemento a elemento (procedimentos para o cálculo exato da matriz hessiana, elemento a elemento, são apresentados por BISHOP (1992)). É suficiente determinar e armazenar o produto da matriz hessiana por um vetor conhecido. Portanto, dado um vetor arbitrário $\mathbf{v} \in \mathbb{R}^P$, o cálculo de $\nabla^2 J(\theta)\mathbf{v}$ requer a aplicação do operador diferencial $\Psi_{\mathbf{v}}\{\cdot\}$ a toda operação realizada para obter $\nabla J(\theta)$, pois $\Psi_{\mathbf{v}}\{\nabla J(\theta)\} = \nabla^2 J(\theta)\mathbf{v}$ (veja apêndice D). Por conveniência notacional, como $\Psi_{\mathbf{v}}\{\theta\} = \mathbf{v}$, é possível expressar $\mathbf{v}$ de forma equivalente à expressão para θ apresentada na equação (6.12), ou seja,

$$\mathbf{v} = \begin{bmatrix} v_{10}^e \cdots v_{l(m+r)}^e \cdots v_{n0}^e \cdots v_{n(m+r)}^e & v_{10}^o \cdots v_{ln}^o \cdots v_{r0}^o \cdots v_{rn}^o \end{bmatrix}^T. \quad (6.22)$$

Aplicando o operador $\Psi_{\mathbf{v}}\{\cdot\}$ à equação (6.11), resulta:

$$\bullet \quad \Psi_{\mathbf{v}}\{x_k^{t+1}\} = \Psi_{\mathbf{v}}\left\{\sum_{j=0}^n w_{kj}^o y_j^t\right\} = \sum_{j=0}^n \left(\Psi_{\mathbf{v}}\{w_{kj}^o\} y_j^t + w_{kj}^o \Psi_{\mathbf{v}}\{y_j^t\}\right) = \sum_{j=0}^n \left(v_{kj}^o y_j^t + w_{kj}^o \Psi_{\mathbf{v}}\{y_j^t\}\right); \quad (6.23)$$

$$\bullet \quad \Psi_{\mathbf{v}}\{y_j^t\} = \Psi_{\mathbf{v}}\{f(s_j^t)\} = f'(s_j^t) \Psi_{\mathbf{v}}\{s_j^t\}; \quad (6.24)$$

- $$\Psi_v \{s_j^t\} = \Psi_v \left\{ \sum_{i=1}^m w_{ji}^e u_i^t + w_{j0}^e + \sum_{i=m+1}^{m+r} w_{ji}^e x_{i-m}^t \right\}$$

$$= \sum_{i=1}^m (v_{ji}^e u_i^t + w_{ji}^e \Psi_v \{u_i^t\}) + v_{j0}^e + \sum_{i=m+1}^{m+r} (v_{ji}^e x_{i-m}^t + w_{ji}^e \Psi_v \{x_{i-m}^t\}). \quad (6.25)$$

Agora, aplicando o operador $\Psi_v \{ \cdot \}$ à equação (6.15), resulta:

- $$\Psi_v \left\{ \frac{\partial J(\theta)}{\partial \theta} \right\} = \Psi_v \left\{ \sum_{t=1}^N \sum_{k=1}^r (x_k^{t+1} - s_k(t)) \frac{\partial x_k^{t+1}}{\partial \theta} \right\}$$

$$= \sum_{t=1}^N \sum_{k=1}^r \left[\Psi_v \{x_k^{t+1} - s_k(t)\} \frac{\partial x_k^{t+1}}{\partial \theta} + (x_k^{t+1} - s_k(t)) \Psi_v \left\{ \frac{\partial x_k^{t+1}}{\partial \theta} \right\} \right]; \quad (6.26)$$

- $$\Psi_v \{x_k^{t+1} - s_k(t)\} = \Psi_v \{x_k^{t+1}\}; \quad (6.27)$$

Para $w_{kj}^o, k=1, \dots, r, j=0, \dots, n$:

- $$\Psi_v \left\{ \frac{\partial x_k^{t+1}}{\partial w_{kj}^o} \right\} = \Psi_v \{y_j^t\} + \sum_{a=1}^n \left(v_{ka}^o \frac{\partial y_a^t}{\partial w_{kj}^o} + w_{ka}^o \Psi_v \left\{ \frac{\partial y_a^t}{\partial w_{kj}^o} \right\} \right); \quad (6.28)$$

- Para $k_1, k_2=1, \dots, r (k_1 \neq k_2)$:
$$\Psi_v \left\{ \frac{\partial x_{k_1}^{t+1}}{\partial w_{k_2 j}^o} \right\} = \sum_{a=1}^n \left(v_{k_1 a}^o \frac{\partial y_a^t}{\partial w_{k_2 j}^o} + w_{k_1 a}^o \Psi_v \left\{ \frac{\partial y_a^t}{\partial w_{k_2 j}^o} \right\} \right); \quad (6.29)$$

- $$\Psi_v \left\{ \frac{\partial y_a^t}{\partial w_{kj}^o} \right\} = \Psi_v \left\{ f'(s_a^t) \sum_{i=m+1}^{m+r} w_{ai}^e \frac{\partial x_{i-m}^t}{\partial w_{kj}^o} \right\} = f''(s_a^t) \Psi_v \{s_a^t\} \sum_{i=m+1}^{m+r} \left(w_{ai}^e \frac{\partial x_{i-m}^t}{\partial w_{kj}^o} \right) +$$

$$+ f'(s_a^t) \sum_{i=m+1}^{m+r} \left(v_{ai}^e \frac{\partial x_{i-m}^t}{\partial w_{kj}^o} + w_{ai}^e \Psi_v \left\{ \frac{\partial x_{i-m}^t}{\partial w_{kj}^o} \right\} \right). \quad (6.30)$$

Para $w_{ji}^e, j=1, \dots, n, i=0, \dots, m+r$:

- $$\Psi_v \left\{ \frac{\partial x_k^t}{\partial w_{ji}^e} \right\} = \Psi_v \left\{ \sum_{a=1}^n w_{ka}^o \frac{\partial y_a^t}{\partial w_{ji}^e} \right\} = \sum_{a=1}^n \left(v_{ka}^o \frac{\partial y_a^t}{\partial w_{ji}^e} + w_{ka}^o \Psi_v \left\{ \frac{\partial y_a^t}{\partial w_{ji}^e} \right\} \right); \quad (6.31)$$

- $$\Psi_v \left\{ \frac{\partial y_a^t}{\partial w_{ji}^e} \right\} = \Psi_v \left\{ f'(s_a^t) \frac{\partial s_a^t}{\partial w_{ji}^e} \right\} = f''(s_a^t) \Psi_v \{s_a^t\} \frac{\partial s_a^t}{\partial w_{ji}^e} + f'(s_a^t) \Psi_v \left\{ \frac{\partial s_a^t}{\partial w_{ji}^e} \right\}; \quad (6.32)$$

- Se $a \neq j$:
$$\Psi_v \left\{ \frac{\partial s_a^t}{\partial w_{ji}^e} \right\} = \sum_{b=m+1}^{m+r} \left(v_{ab}^e \frac{\partial x_{b-m}^t}{\partial w_{ji}^e} + w_{ab}^e \Psi_v \left\{ \frac{\partial x_{b-m}^t}{\partial w_{ji}^e} \right\} \right); \quad (6.33)$$

- Se $a = j$ e $0 \leq i \leq m$: $\Psi_v \left\{ \frac{\partial s'_a}{\partial w_{ji}^e} \right\} = \sum_{b=m+1}^{m+r} \left(v_{ab}^e \frac{\partial x'_{b-m}}{\partial w_{ji}^e} + w_{ab}^e \Psi_v \left\{ \frac{\partial x'_{b-m}}{\partial w_{ji}^e} \right\} \right);$ (6.34)

- Se $a = j$ e $(m+1) \leq i \leq (m+r)$: $\Psi_v \left\{ \frac{\partial s'_a}{\partial w_{ji}^e} \right\} = \Psi_v \{x'_{i-m}\} + \sum_{b=m+1}^{m+r} \left(v_{ab}^e \frac{\partial x'_{b-m}}{\partial w_{ji}^e} + w_{ab}^e \Psi_v \left\{ \frac{\partial x'_{b-m}}{\partial w_{ji}^e} \right\} \right).$ (6.35)

Fazendo as substituições, em sequência, das equações (6.23) a (6.25) e (6.27) a (6.35) na equação (6.26), resulta o vetor $\Psi_v \left\{ \frac{\partial J(\theta)}{\partial \theta} \right\} = \nabla^2 J(\theta) \mathbf{v}$, avaliado exatamente.

6.4 Análise da capacidade de processamento

Ao contrário das redes de Hopfield (HOPFIELD, 1982) e máquinas de Boltzmann (ACKLEY *et al.*, 1985; HERTZ *et al.*, 1991), o importante nas redes neurais recorrentes apresentadas acima não é apenas o seu estado estacionário, mas também, e principalmente, o seu comportamento transitório, ou seja, a fase transitória da resposta deve ser considerada parte da especificação funcional. Sendo assim, o objetivo de se utilizar uma rede neural recorrente não está associado necessariamente a processos de armazenagem e restauração de dados, podendo envolver processos de representação de comportamentos possivelmente não-estacionários. Logo, a dinâmica das redes neurais recorrentes empregadas neste estudo não resulta simplesmente de sua dinâmica interna a partir de um estado inicial, mas podem estar sujeitas à influência contínua de dados de entrada externa.

Neste caso, conforme sugerido por KOLEN (1994), a teoria de sistemas dinâmicos indica que as redes neurais recorrentes de tempo discreto estão mais intimamente relacionadas com sistemas de funções iterativas (BARNSELY, 1988) do que com máquinas de estado finito determinísticas (CLEEREMANS *et al.*, 1989). Dito de outra forma, o comportamento da rede recorrente está governado mais por mapeamentos iterativos do que por um conjunto finito de regras discretas.

A análise via teoria de sistemas de funções iterativas permite interpretar as estratégias e esquemas de manipulação de informação inerentes às redes neurais recorrentes de tempo discreto, fornecendo a fundamentação teórica necessária para o entendimento de suas propriedades dinâmicas. Para o caso de redes neurais recorrentes de tempo contínuo, veja FUNAHASHI & NAKAMURA (1993).

6.5 Tipos de comportamento dinâmico em uma rede recorrente

Um sistema dinâmico consiste de duas partes: um estado e uma dinâmica. O estado descreve a condição atual do sistema na forma de um vetor de variáveis parametrizadas em relação ao tempo, sendo que o conjunto de estados possíveis é denominado espaço de estados do sistema. A dinâmica descreve como o estado do sistema evolui no tempo, sendo que a sequência de estados exibida por um sistema dinâmico durante sua evolução no tempo é denominada trajetória no espaço de estados.

Neste estudo é assumido que a dinâmica é determinística, ou seja, para cada estado do sistema, a dinâmica especifica unicamente o próximo estado (dinâmica discreta) ou então a direção de variação do estado (dinâmica contínua). A tabela 6.1 apresenta os tipos possíveis de sistemas dinâmicos, sendo que a rede neural recorrente descrita neste capítulo corresponde a um sistema de funções iterativas.

Tabela 6.1 - Taxionomia dos sistemas dinâmicos (KOLEN, 1994)

		ESPAÇO DE ESTADOS	
		contínuo	discreto
DINÂMICA	contínua	sistema de equações diferenciais	vidros de spin ¹
	discreta	sistema de funções iterativas	autômata

6.5.1 Comportamento de regime

A necessidade de se considerar a fase transitória da dinâmica da rede neural recorrente como parte importante de seu comportamento para efeito de aplicação não reduz a significância da fase de regime. Para um dado sistema dinâmico, o comportamento característico da trajetória em regime é uma região do espaço de estados denominada atrator. Se toda trajetória de um sistema dinâmico não-perturbado, pertencente a um subconjunto do espaço de estados, converge para um atrator, diz-se que este subconjunto é a base de atração. Tanto para dinâmica discreta como contínua, existem apenas quatro tipos qualitativamente distintos de atratores (OTT, 1993):

- ponto fixo: é um estado invariante sob a dinâmica do sistema, de forma que o sistema não-perturbado permanece indefinidamente neste estado uma vez que o tenha atingido;

¹ Vidros de spin são modelos desenvolvidos para descrever sistemas magnéticos, sendo extensivamente investigados em mecânica estatística (BINDER & YOUNG, 1986).

- ciclo limite: é uma trajetória fechada invariante sob a dinâmica do sistema, de forma que o sistema não-perturbado permanece indefinidamente nesta trajetória uma vez que a tenha atingido em algum de seus pontos;
- atrator quase-periódico: é um ciclo-limite com peridiocidade múltipla, de tal forma que a trajetória nunca retorna ao mesmo ponto;
- caos: é um ciclo limite aperiódico, só podendo ocorrer se a dinâmica do sistema for não-linear, sendo o resultado da presença de infinitos ciclos-limite periódicos instáveis.

Pontos fixos estáveis resultam da aplicação repetida de um operador de contração, de tal forma que, no limite, estas operações reduzem subespaços do espaço de estados em um único ponto. Ciclos-limite resultam da aplicação repetida de compressão e rotação do espaço de estados. Caos, por sua vez, resulta da aplicação repetida de operações sucessivas de expansão e dobra do espaço de estados. A dobra é um método de reorganização topológica do espaço de estados capaz de criar novas vizinhanças (STEWART, 1989).

Em termos de espectro de frequência, enquanto trajetórias quase-periódicas apresentam múltiplos picos nas frequências constituintes, trajetórias caóticas apresentam bandas de frequência com distribuição de energia decrescente com o inverso da frequência, o que permite diferenciá-las do ruído branco, o qual tem distribuição uniforme de frequência através do espectro (BARTLETT, 1990; CASDAGLI, 1989).

De forma equivalente aos atratores, é possível definir os repelentes, ou seja, regiões do espaço de estados que repelem as trajetórias que por lá atravessam.

6.5.2 Sistemas de funções iterativas não-contrativas

BARNSELY (1988) propôs os fundamentos da teoria de sistemas de funções iterativas na forma de um método de descrição do comportamento limite de sistemas de transformações. Estes sistemas são responsáveis por uma série de estruturas fractais e são definidos como um conjunto finito de mapeamentos contrativos sobre um espaço métrico. Embora o comportamento limite de uma única transformação contrativa seja apenas um ponto, o conjunto de pontos resultantes do comportamento limite da união de transformações pode ser muito complexo.

Restritos a arquiteturas particulares de redes neurais recorrentes (incluindo a rede de Elman) e exigindo que os parâmetros da rede neural assumam sempre valores que garantam o mapeamento contrativo dos sinais realimentados, resultados de convergência para treinamento de redes recorrentes são encontrados em KUAN *et al.* (1994).

STARK (1991) e KOLEN (1994) estudaram a relação entre redes neurais recorrentes e sistemas de funções iterativas contrativas, concluindo que redes neurais recorrentes genéricas são sistemas de funções iterativas híbridos, apresentando funções contrativas e expansivas. A principal diferença está na natureza dos atratores, pois, enquanto sistemas de funções iterativas contrativas apresentam um único atrator em todo o espaço de estados, sistemas de funções iterativas não-contrativas podem apresentar número e variedade ilimitados de atratores. Por exemplo, a possibilidade de existirem múltiplas bases de atração cria condições para a implementação de memórias endereçáveis por conteúdo, como no caso da rede de Hopfield (HOPFIELD, 1982).

Talvez o resultado mais expressivo associado a redes neurais recorrentes tenha sido obtido por SIEGELMANN & SONTAG (1992) (veja também KOLEN, 1994), demonstrando que redes neurais recorrentes apresentam capacidade de computação universal, sendo totalmente equivalentes a uma máquina de Turing (TURING, 1936).

Motivado pelos resultados de KOLEN & POLLACK (1991), é possível afirmar que o próprio processo iterativo de ajuste de parâmetros, seja para redes neurais recorrentes como não-recorrentes (veja equação (6.8)), é, por si só, um sistema de funções iterativas híbrido que pode apresentar comportamentos dinâmicos muito complexos.

6.5.3 Propriedades de estabilidade

A determinação de propriedades de estabilidade para sistemas dinâmicos lineares é uma tarefa simples, existindo vasta bibliografia associada ao tema (ÄSTROM & WITTENMARK, 1990; CHEN, 1984; KAILATH, 1980). Por outro lado, o mesmo não pode ser dito com relação ao caso não-linear, que engloba também os sistemas dinâmicos representados pelas redes neurais recorrentes descritas neste capítulo. VIDYASAGAR (1993) apresenta um tratamento geral de estabilidade para o caso de sistemas dinâmicos não-lineares, evidenciando as principais características responsáveis pela dificuldade de análise:

- o sistema é estável para alguns tipos de entradas e instável para outros tipos;
- a estabilidade é geralmente uma propriedade local, ou seja, pequenas perturbações podem não ser capazes de retirar o sistema de uma região de estacionariedade, enquanto que grandes perturbações geralmente conduzem o sistema para fora da região de estacionariedade;
- o sistema pode ter vários pontos de estacionariedade, que podem ser estáveis ou instáveis.

Capítulo 7

Aplicações de Redes Neurais

Este capítulo apresenta resultados referentes à aplicação dos modelos de redes neurais desenvolvidos nos capítulos anteriores a problemas de aproximação de funções, classificação de padrões, identificação de sistemas dinâmicos, previsão de séries temporais e controle de processos. Em todos os casos, os modelos são obtidos pela aplicação de métodos de treinamento supervisionado, a partir de um conjunto de dados de treinamento e teste.

A utilização de redes neurais não-recorrentes envolve os casos paramétrico e não-paramétrico, enquanto que apenas redes neurais recorrentes paramétricas são empregadas. Os problemas propostos a seguir exigem soluções dotadas invariavelmente de estruturas não-lineares. Isto não significa, no entanto, que redes neurais puramente lineares não tenham importância prática, já que estas vêm sendo empregadas com sucesso em inúmeras aplicações (BALDI & HORNIK, 1995).

Apesar de alguns exemplos de simulação serem incluídos, o objetivo principal deste capítulo é sugerir uma formulação adequada para cada tipo de problema, de modo a possibilitar a aplicação direta dos resultados obtidos em capítulos anteriores.

7.1 Aproximação de funções

Problemas de aproximação de funções utilizando modelos de aproximação baseados em redes neurais paramétricas e não-paramétricas sujeitas a treinamento supervisionado não requerem aqui nenhum tipo de formulação especial, pois foram alvo de intenso estudo em praticamente todos os capítulos anteriores. Justifica-se, portanto, a passagem direta para a apresentação de resultados de simulação.

Os resultados a seguir se referem apenas a métodos de aproximação de funções, sendo que métodos de aproximação simultânea de uma função e suas derivadas podem ser encontrados em CARDALIAGUET & EUVRARD (1992) e GALLANT & WHITE (1992).

7.1.1 Comparação entre modelos paramétricos e não-paramétricos

O processo de aproximação utilizando as redes neurais paramétricas e não-paramétricas descritas no capítulo 5 foi testado em cinco problemas de aproximação com duas variáveis independentes ($m = 2$) e uma variável de saída ($r = 1$), dimensões apropriadas para apresentação de resultados gráficos. Estes exemplos foram também utilizados por HWANG *et al.* (1994) para comparar métodos utilizando redes neurais paramétricas e o método de busca de projeção em sua forma original (veja seção 5.3.3), sem explorar as propriedades geométricas advindas do estudo da condição de solvabilidade e também a escolha criteriosa da direção inicial de busca (veja seção 5.3.5). Cada exemplo considera para o problema de aproximação uma das cinco funções não-lineares $g^{(K)}: [0,1]^2 \rightarrow \mathfrak{R}$, $K=1, \dots, 5$, apresentadas a seguir e ilustradas graficamente na figura 7.1.

- 1) função de interação simples: $g^{(1)}(x_1, x_2) = 10,391 \cdot [(x_1 - 0,4) \cdot (x_2 - 0,6) + 0,36]$;
- 2) função radial: $g^{(2)}(x_1, x_2) = 24,234 \cdot [r^2(0,75 - r^2)]$, com $r^2 = (x_1 - 0,5)^2 + (x_2 - 0,5)^2$;
- 3) função harmônica: $g^{(3)}(x_1, x_2) = 42,659 \cdot [0,1 + \tilde{x}_1(0,05 + \tilde{x}_1^4 - 10\tilde{x}_1^2\tilde{x}_2^2 + 5\tilde{x}_2^4)]$, com $\tilde{x}_{1/2} = x_{1/2} - 0,5$;
- 4) função aditiva:

$$g^{(4)}(x_1, x_2) = 1,3356 \cdot [1,5(1 - x_1) + e^{2x_1-1} \sin(3\pi(x_1 - 0,6)^2) + e^{3x_2-1,5} \sin(4\pi(x_2 - 0,9)^2)]$$
;
- 5) função de interação complicada:

$$g^{(5)}(x_1, x_2) = 1,9 \cdot [1,35 + e^{x_1-x_2} \sin(13(x_1 - 0,6)^2) \cdot \sin(7x_2)]$$
.

As funções $g^K(\mathbf{x}) = g^K(x_1, x_2)$ ($K=1, \dots, 5$) são assumidas desconhecidas, sendo possível apenas obter amostras do valor de cada função para vetores de entrada $\mathbf{x} = [x_1 \ x_2]^T$ definidos de forma a explorar adequadamente o espaço de aproximação. Seguindo o mesmo procedimento adotado por HWANG *et al.* (1994), para produzir um conjunto de dados de treinamento de dimensão $N_{tr} = 225$, foram gerados vetores de entrada $\mathbf{x}_l = [x_{1l} \ x_{2l}]^T$ ($l=1, \dots, 225$) com base em uma distribuição uniforme de valores aleatórios no intervalo $[0,1]$.

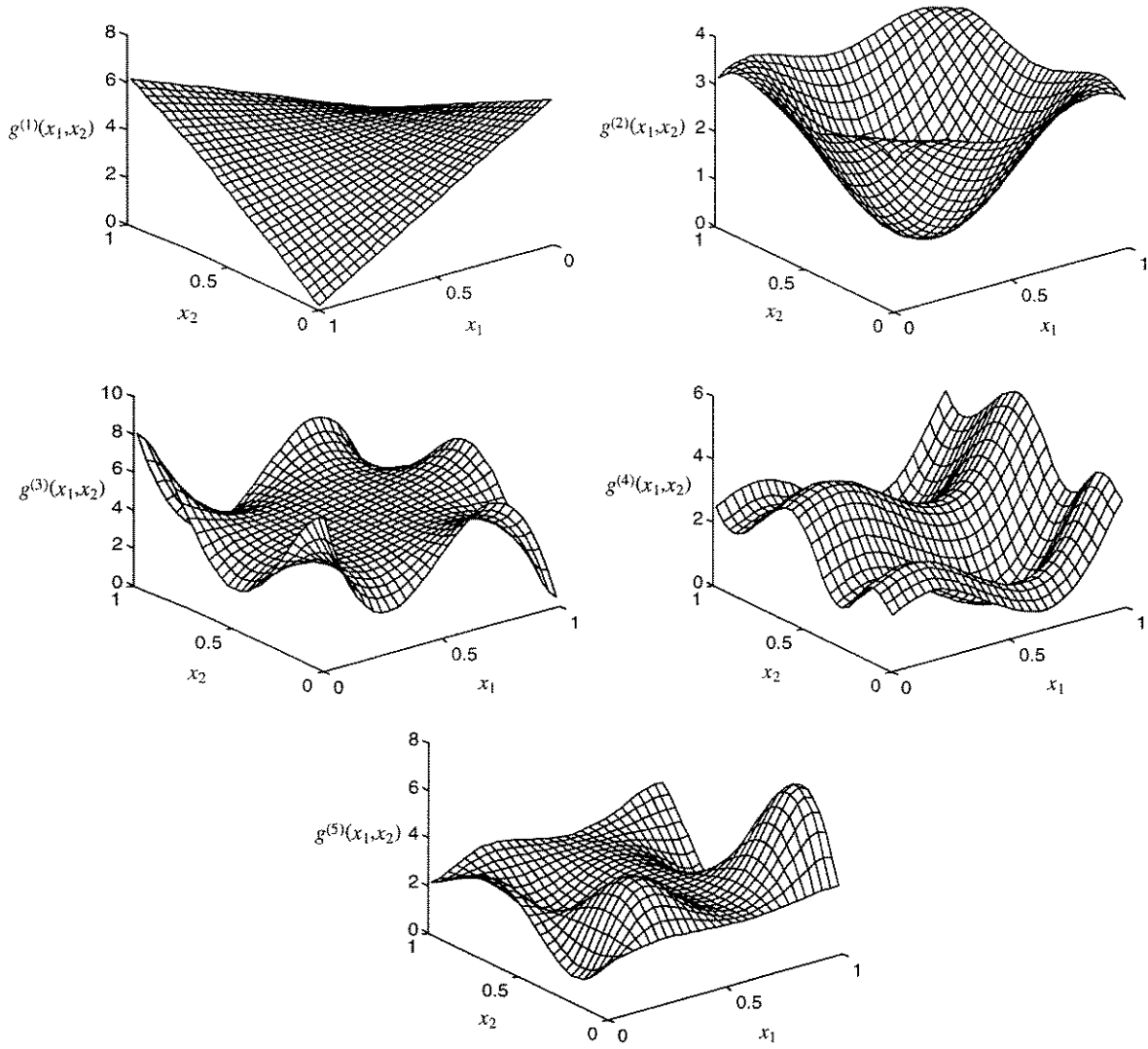


Figura 7.1 - Funções a serem aproximadas

A matriz de variáveis de entrada $\mathbf{X} \in \mathbb{R}^{2 \times 225}$ é a mesma para os cinco exemplos, mas as amostras de cada função avaliada em $\mathbf{x}_l = [x_{1l} \ x_{2l}]^T$ ($l=1, \dots, 225$) podem ou não estar sujeitas a ruído de média zero e variância fixa. O nível de aproximação é avaliado através da comparação dos valores preditos pelo modelo $\hat{g}$ com os valores de um conjunto de dados de entrada-saída para teste (independente dos dados para treinamento), com dimensão $N_{te} = 10000$. A comparação é baseada na fração de variância não-explicada (FVNE) dada por:

$$FVNE_K = \frac{\sum_{l=1}^{N_{te}} (\hat{g}^{(K)}(x_{1l}, x_{2l}) - g^{(K)}(x_{1l}, x_{2l}))^2}{\sum_{l=1}^{N_{te}} (g^{(K)}(x_{1l}, x_{2l}) - \bar{g}^{(K)})^2}, \quad K=1, \dots, 5, \quad (7.1)$$

onde

$$\bar{g}^{(K)} = \frac{1}{N_{te}} \sum_{l=1}^{N_{te}} g^{(K)}(x_{1l}, x_{2l}). \quad (7.2)$$

A fração de variância não-explicada é um critério de avaliação do nível de aproximação que também foi utilizado por FRIEDMAN & STUETZLE (1981), ROOSEN & HASTIE (1994) e HWANG *et al.* (1994). Este critério fornece uma avaliação do erro quadrático médio normalizado por uma estimativa da variância da função a ser aproximada.

Para cada um dos cinco exemplos, foram tomados seis tipos de modelos de aproximação $\hat{g}^{(K)}: \mathcal{R}^2 \rightarrow \mathcal{R}$, todos passíveis de representação na forma de uma rede neural com uma camada intermediária (veja figura 5.2, capítulo 5) e descritos na forma:

- 1) RP5: rede neural paramétrica com uma camada intermediária de 5 neurônios com funções de ativação idênticas $f_j(\cdot) = f(\cdot) = \tanh(\cdot)$ ($j=1, \dots, 5$) e parâmetros ajustados via método do gradiente conjugado, descrito no apêndice D;
- 2) RP15: rede neural paramétrica com uma camada intermediária de 15 neurônios com funções de ativação idênticas $f_j(\cdot) = f(\cdot) = \tanh(\cdot)$ ($j=1, \dots, 15$) e parâmetros ajustados via método do gradiente conjugado, descrito no apêndice D;
- 3) RNPP1: rede neural não-paramétrica submetida ao algoritmo de aproximação construtivo apresentado na seção 5.3.5 (capítulo 5), mas com funções de ativação paramétricas idênticas e definidas previamente, tais que $f_j(\cdot) = f(\cdot) = \tanh(\cdot)$ ($j=1, \dots$);
- 4) RNPP2: rede neural não-paramétrica submetida ao algoritmo de aproximação construtivo apresentado na seção 5.3.3 (capítulo 5), com funções de ativação paramétricas na forma de expansão ortonormal truncada, utilizando funções de Hermite como funções básicas, conforme descrito na seção 2.8 (capítulo 2) e seção 4.5.1 (capítulo 4). A ordem máxima das funções de Hermite utilizadas é $P = 9$, seguindo o procedimento adotado por HWANG *et al.* (1994);
- 5) RNPP3: rede neural não-paramétrica submetida ao algoritmo de aproximação construtivo apresentado na seção 5.3.5 (capítulo 5), com funções de ativação paramétricas na forma de expansão ortonormal truncada, utilizando funções de Hermite como funções básicas, conforme descrito na seção 2.8 (capítulo 2) e seção 4.5.1 (capítulo 4). A ordem máxima das funções de Hermite utilizadas é $P = 9$;
- 6) RNPNP: rede neural não-paramétrica submetida ao algoritmo de aproximação construtivo apresentado na seção 5.3.5 (capítulo 5), com funções de ativação não-paramétricas na forma de splines polinomiais suavizantes, conforme descrito na seção 2.5.4 (capítulo 2) e seção 4.5.2 (capítulo 4).

Os resultados de simulação obtidos são apresentados nas tabelas 7.1 e 7.2, onde n representa o número de neurônios na camada intermediária e o ruído no processo de amostragem associado aos dados presentes na tabela 7.2 é $0,25\varepsilon$, com ε uma variável aleatória de distribuição normal, média zero e variância unitária. Detalhes a respeito da implementação computacional podem ser encontrados em VON ZUBEN & NETTO (1996), tendo sido utilizado o ambiente de simulação computacional MATLAB® (MATLAB, 1992). Os dados referentes ao modelo RNPP2 foram extraídos de HWANG *et al.* (1994).

Tabela 7.1 - Desempenho dos modelos de aproximação para processo de amostragem sem ruído

	$g^{(1)}$		$g^{(2)}$		$g^{(3)}$		$g^{(4)}$		$g^{(5)}$	
	n	FVNE	n	FVNE	n	FVNE	n	FVNE	n	FVNE
RP5	5	0,00008	5	0,01921	5	0,48603	5	0,03593	5	0,18221
RP15	15	0,00001	15	0,00938	15	0,05771	15	0,00169	15	0,03223
RNPP1	15	0,02175	15	0,07448	15	0,22898	15	0,11239	15	0,24627
RNPP2	3	0,00000	3	0,00860	5	0,00001	3	0,00076	5	0,01531
RNPP3	2	0,00000	3	0,00000	5	0,00002	3	0,00060	5	0,01102
RNPNP	2	0,00000	3	0,00262	5	0,00259	2	0,00027	5	0,00745

Tabela 7.2 - Desempenho dos modelos de aproximação para processo de amostragem com ruído

	$g^{(1)}$		$g^{(2)}$		$g^{(3)}$		$g^{(4)}$		$g^{(5)}$	
	n	FVNE	n	FVNE	n	FVNE	n	FVNE	n	FVNE
RP5	5	0,00798	5	0,01970	5	0,48891	5	0,04252	5	0,23827
RP15	15	0,00245	15	0,01160	15	0,07954	15	0,01025	15	0,08124
RNPP1	15	0,03138	15	0,14789	15	0,46982	15	0,12972	15	0,21670
RNPP2	3	0,00767	3	0,03072	5	0,01045	3	0,01841	5	0,04273
RNPP3	2	0,00384	3	0,00356	5	0,00723	2	0,01544	5	0,02439
RNPNP	2	0,00341	3	0,01566	5	0,02581	2	0,00538	5	0,03194

Alguns resultados teóricos discutidos em capítulos anteriores podem ser constatados por simples inspeção dos dados de simulação apresentados nas tabelas 7.1 e 7.2., dentre os quais é possível destacar:

- os modelos RP5, RP15 e RNPP1 são aqueles que têm a tangente hiperbólica como função de ativação. Os resultados para o modelo RNPP1, quando comparado com os modelos

RP5 e RP15, mostram que a tangente hiperbólica não apresenta flexibilidade suficiente para permitir a aproximação neurônio-a-neurônio (modelo RNPP1), já que o ajuste de todos os neurônios simultaneamente produz resultados melhores em todos os exemplos, pelo menos para o modelo RP15;

- comparando-se os modelos RP5, RP15 e RNPP1 (funções de ativação previamente definidas) com os modelos RNPP2, RNPP3 e RNPNP (funções de ativação definidas a partir dos dados de aproximação disponíveis), verifica-se que a fixação prévia da função de ativação acaba por exigir um número bem superior de neurônios para obter níveis de aproximação equivalentes, conduzindo a um custo computacional bem maior;
- a diferença básica entre os modelos RNPP2 e RNPP3 está no fato de o algoritmo da seção 5.3.5 (VON ZUBEN & NETTO, 1995a; 1995c), aplicado ao modelo RNPP3, ser mais eficiente que o algoritmo da seção 5.3.3 (HWANG *et al.*, 1994), aplicado ao modelo RNPP2. Além de utilizar índices de projeção para definir as direções de projeção iniciais, o algoritmo de VON ZUBEN & NETTO (1995a; 1995c) realiza uma iteração em dois grupos de variáveis e calcula exatamente o terceiro grupo, contra iteração em três grupos de variáveis no algoritmo de HWANG *et al.* (1994). Com isso, o algoritmo de VON ZUBEN & NETTO (1995a; 1995c) vai fornecer uma solução mais precisa (por avaliar exatamente o terceiro grupo de variáveis) e exigindo um número inferior de iterações e de cálculos computacionais por iteração;
- comparando-se o modelo com função de ativação paramétrica (RNPP3) e o modelo com função de ativação não-paramétrica (RNPNP), ambos submetidos ao mesmo algoritmo de aproximação construtivo e apresentando os melhores desempenhos dentre todos os modelos, é possível concluir que RNPNP pode fornecer soluções melhores (veja o caso da função $g^{(4)}$) ou piores (veja o caso das funções $g^{(2)}$ e $g^{(3)}$) que o modelo RNPP3. Como, por definição, o modelo RNPNP é mais flexível que o modelo RNPP3, ele sempre tende a produzir uma melhor solução, desde que o número e a distribuição de amostras seja compatível com o nível de suavidade da função a ser aproximada. Com isso, nos casos em que o desempenho do modelo RNPNP é inferior àquele produzido pelo modelo RNPP3, conclui-se que o procedimento de amostragem utilizado não foi adequado. De fato, aumentando-se o número de amostras com a manutenção da mesma densidade de distribuição dos pontos, RNPNP é capaz de superar RNPP3 em todos os exemplos (VON ZUBEN & NETTO, 1996).

Outro aspecto importante é a limitação de poder de aproximação apresentada por expansão ortonormal truncada (modelo RNPP3) quando comparada com aproximação polinomial por partes via splines suavizantes (modelo RNPNP). Considerando os exemplos de simulação apresentados, estas limitações não devem ter sido evidenciadas, principalmente pela baixa dimensionalidade dos dados e pela simplicidade da tarefa de aproximação, nunca exigindo mais que cinco neurônios na camada intermediária para fornecer soluções adequadas utilizando-se estes dois modelos. Além disso, procedimentos de validação cruzada podem ser diretamente incorporados ao modelo RNPNP, conforme discutido no capítulo 4 e também em BATES *et al.* (1987), POPE & GADH (1988) e ROOSEN & HASTIE (1994). Ou seja, a suavidade do modelo RNPNP é definida automaticamente e continuamente a partir dos dados via validação cruzada, enquanto que o modelo RNPP3 (e também o RNPP2) só admite ajuste, também via validação cruzada, no número P de funções básicas a serem utilizadas, um parâmetro que assume valores inteiros, portanto discretos.

7.1.2 Análise do método de aproximação por busca de projeção

Para ilustrar a forma como o método de aproximação por busca de projeção realiza a aproximação de funções, considere o problema de aproximação da função $g^{(5)}$ utilizando o modelo RNPP3, cujo desempenho é apresentado na tabela 7.1. As direções de projeção escolhidas para cada um dos neurônios da camada intermediária, as respectivas funções de ativação e a expansão ortogonal de quatro das funções de ativação são apresentadas respectivamente nas figuras 7.2, 7.3 e 7.4.

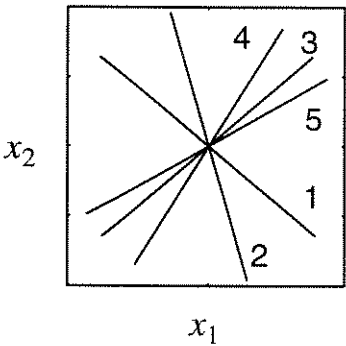


Figura 7.2 - As direções de projeção escolhidas

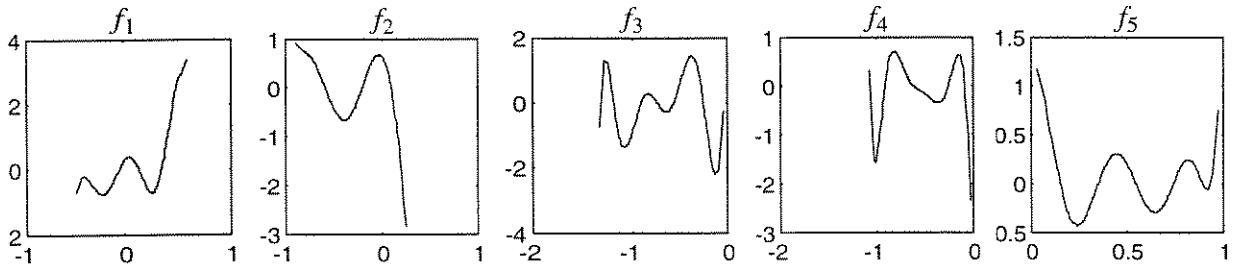


Figura 7.3 - As funções de ativação dos neurônios da camada intermediária

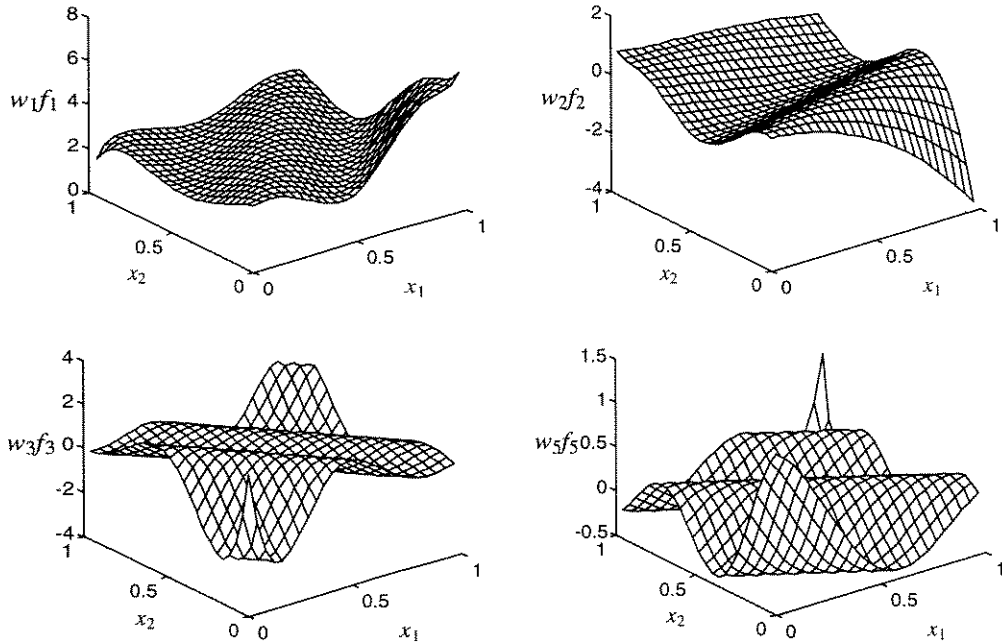


Figura 7.4 - A expansão ortogonal de quatro das funções de ativação obtidas

7.1.3 Aproximação de funções com múltiplas saídas

A flexibilidade associada à possibilidade de escolha da melhor função de ativação para cada neurônio a partir dos dados de aproximação é fundamental para o bom desempenho do método de aproximação construtivo, principalmente quando aplicado a problemas envolvendo múltiplas saídas. A seguir, os resultados de BÄRMANN & BIEGLER-KÖNIG (1992), utilizando um modelo RNPP4, equivalente ao modelo RNPP1, mas com iteração em três grupos de variáveis e com definição aleatória das direções de projeção iniciais, são comparados com o RNPP3 em um problema de aproximação com $g: \mathbb{R}^8 \rightarrow \mathbb{R}^3$, na forma:

$$s = g(\mathbf{x}), \quad \text{com } g = \begin{bmatrix} 0,25(x_1x_2 + x_3x_4 + x_5x_6 + x_7x_8) \\ 0,125(x_1 + x_2 + x_3 + x_4 + x_5 + x_6 + x_7 + x_8) \\ [1 - 0,25(x_1x_2 + x_3x_4 + x_5x_6 + x_7x_8)]^{1/2} \end{bmatrix}. \quad (7.3)$$

Para o nível de aproximação fixado por BÄRMANN & BIEGLER-KÖNIG (1992), foram necessários $n = 12$ neurônios na camada intermediária para o modelo RNPP4, $n = 10$ para o RNPP1 e $n = 3$ para RNPP3.

7.1.4 Robustez a dados não-informativos

Como discutido no capítulo 3 por ocasião da apresentação de problemas de aproximação regularizados (veja também GIROSI *et al.* (1995)), nem todas as variáveis de entrada participam na definição do mapeamento de entrada-saída a ser aproximado. Neste caso, o método de aproximação por busca de projeção deve ser capaz de detectar variáveis não-informativas, posicionando as direções de projeção em espaços ortogonais aos eixos de coordenadas associados a estas variáveis.

Como um exemplo de simulação, a função de aproximação

$$g(x_1, x_2, x_3) = \text{sen}(x_1 + 3x_2) \tag{7.4}$$

é assumida desconhecida e um conjunto de dados de aproximação é gerado na forma:

$$s_l = g(x_{1l}, x_{2l}, x_{3l}) + \delta \epsilon_l, \quad l=1, \dots, 100, \tag{7.5}$$

com x_{il} ($i=1, \dots, 3$; $l=1, \dots, 100$) variáveis aleatórias com distribuição uniforme no intervalo $[-\pi, \pi]$ e ϵ_l uma variável aleatória de distribuição normal, média zero e variância unitária.

Utilizando o modelo de aproximação RNPNP, foram produzidos os resultados de simulação apresentados na tabela 7.3.

Tabela 7.3 - Desempenho do modelo RNPNP para processo de amostragem com e sem ruído

	n	FVNE	$\mathbf{v}$ (direção de projeção)
$\delta = 0$ (sem ruído)	1	0,00003	$[0,316 \ 0,949 \ 0,000]^T$
$\delta = 0,05$ (com ruído)	1	0,00048	$[-0,317 \ -0,948 \ -0,000]^T$

Sem ruído ou mesmo com ruído no processo de amostragem, o método de aproximação por busca de projeção foi capaz de definir automaticamente um direção de projeção ortogonal ao eixo de coordenadas de x_3 , restando ajustar a única função de ativação utilizada (splines suavizantes) no sentido de aproximar a função seno.

7.2 Classificação de padrões

A relação existente entre um problema de aproximação de funções e um problema de classificação de padrões já foi discutida anteriormente, sendo que o primeiro é considerado um mapeamento de um espaço contínuo para um outro espaço contínuo, enquanto que o segundo pode ser visto como um mapeamento de um espaço contínuo para um espaço discreto. Com isso, todo problema de classificação de padrões pode ser representado na forma de uma composição de um problema de aproximação de funções com um mapeamento elementar de um espaço contínuo para um espaço discreto.

Assuma que todo padrão definido no espaço $\mathcal{R}^m$ pertença a uma dentre r classes. Neste caso, o problema de classificação pode ser definido na forma:

- dado o valor de y associado a um determinado padrão de entrada, defina um mapeamento elementar h que associe y à classe a qual pertence o padrão de entrada;
- uma vez definido o mapeamento elementar h , encontre uma função g tal que a composição $g \circ h$ associe o padrão à sua respectiva classe, conforme ilustrado na figura 7.5.

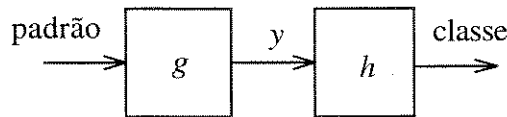


Figura 7.5 - O problema de classificação de padrões

Por exemplo, considere o problema de classificação de padrões pertencentes ao $\mathcal{R}^2$ em uma de duas classes, denominadas **A** e **B**, conforme apresentado na figura 7.6. A função f desconhecida representa a fronteira que delimita as classes, de forma que todo ponto (x_1, x_2) "abaixo" de f pertence à classe **A** e "acima" de f pertence à classe **B**. Dependendo do caso, pontos exatamente na fronteira podem ser atribuídos a nenhuma das classes, a ambas ou a apenas uma delas.

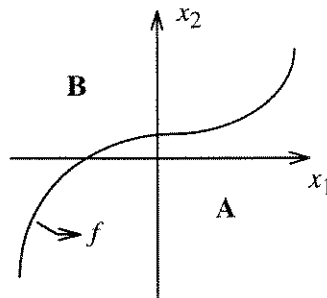


Figura 7.6 - Problema de classificação de padrões (x_1, x_2) nas classes **A** ou **B**

Este problema de classificação de padrões pode ser resolvido conforme indicado na figura 7.7. Uma vez definido o mapeamento elementar h , resulta um problema de aproximação de uma função g tal que:

$$g(x_1, x_2) \begin{cases} < 0 & \text{se } x_2 < f(x_1) \\ = 0 & \text{se } x_2 = f(x_1) \\ > 0 & \text{se } x_2 > f(x_1) \end{cases} \tag{7.6}$$

Redes neurais não-paramétricas podem, então, ser utilizadas na geração do modelo de aproximação $\hat{g}$ a partir de um conjunto de dados de aproximação. Uma comparação entre redes neurais artificiais e outros métodos aplicados a classificação de padrões é apresentada por RIPLEY (1994).

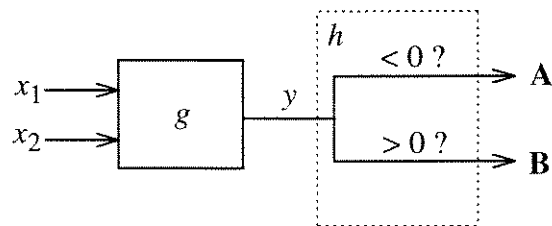


Figura 7.7 - Problema de classificação de padrões (x_1, x_2) nas classes **A** ou **B**

7.3 Identificação de sistemas dinâmicos de tempo discreto

A teoria de identificação de sistemas engloba o estudo de modelos de aproximação para sistemas dinâmicos a partir de dados amostrados. A escolha da estrutura do modelo de aproximação é guiada por conhecimentos preliminares acerca do sistema que gerou os dados, conduzindo, por exemplo, a uma classe de modelos paramétricos. Na ausência de conhecimento preliminar, modelos não-paramétricos são geralmente empregados.

É sabido que, em aplicações práticas, é impossível obter um modelo capaz de descrever o sistema exatamente (SJÖBERG, 1995). O que se busca, na verdade, é obter a melhor aproximação do sistema real, baseada em algum critério. Apesar de quase nunca corresponder à melhor aproximação, é comum assumir que o sistema a ser identificado é linear, viabilizando a aplicação de resultados bem sedimentados e abrangentes da teoria de sistemas lineares.

Em muitos casos, no entanto, existe a necessidade de se utilizar modelos não-lineares (paramétricos ou não), tornando a tarefa de identificação muito mais complexa, já que a maior parte dos resultados válidos para o caso linear não mais se aplicam (VIDYASAGAR, 1993). Uma razão para a maior complexidade associada ao uso de modelos de aproximação não-lineares é

a existência de um número infinitamente maior de alternativas para a estrutura do modelo, já que modelos não-lineares correspondem a todos os modelos que não são lineares, ou seja, a não-linearidade pode se manifestar de inúmeras formas.

Com isso, modelos de aproximação não-lineares precisam ser mais flexíveis, o que requer um número maior de parâmetros que o caso linear, de forma que classes abrangentes de não-linearidades possam ser adequadamente representadas. Redes neurais artificiais como modelos paramétricos e não-paramétricos submetidos a técnicas de regularização, conforme descrito em capítulos anteriores, representam alternativas mais flexíveis e confiáveis quando comparadas a modelos baseados na definição arbitrária de estruturas paramétricas com não-linearidades previamente definidas. BILLINGS (1980) apresenta uma coletânea de modelos de identificação deste último tipo, enquanto que NARENDRA & PARTHASARATHY (1990) e CHEN & BILLINGS (1992) descrevem alguns modelos de identificação utilizando redes neurais artificiais.

A discussão acima é válida tanto para sistemas dinâmicos de tempo contínuo como de tempo discreto, mas em função do tipo de modelos de redes neurais a serem utilizados, os resultados a seguir envolvem apenas sistemas dinâmicos de tempo discreto.

7.3.1 O processo de identificação

A partir de dados de entrada externa, $\mathbf{u}(t) \in \mathbb{R}^q$, e saída, $\mathbf{y}(t) \in \mathbb{R}^r$, amostrados de um sistema dinâmico de tempo discreto h na forma:

$$\mathbf{y}(t) = h(\mathbf{u}^{t-1}, \mathbf{y}^{t-1}) + \mathbf{v}(t), \quad (7.7)$$

com $\mathbf{y}(0)$ dado e t um inteiro positivo, é possível compor os seguintes conjuntos de dados

$$\begin{aligned} \mathbf{u}^t &= \{\mathbf{u}(1), \mathbf{u}(2), \dots, \mathbf{u}(t)\} \\ \mathbf{y}^t &= \{\mathbf{y}(1), \mathbf{y}(2), \dots, \mathbf{y}(t)\} \end{aligned}$$

sendo que o ruído aditivo $\mathbf{v}(t)$ indica que a saída não vai ser exatamente uma função dos dados passados.

O processo de identificação associado ao sistema dinâmico da equação (7.7) pode seguir os mesmos procedimentos desenvolvidos em capítulos anteriores para o caso de aproximação de funções a partir de dados amostrados, desde que h seja tomada como uma função multivariável. Para tanto, as seguintes hipóteses devem ser consideradas (SJÖBERG, 1995):

- a identificação de sistemas não-lineares pode ser vista como uma concatenação entre um pré-processamento dos dados amostrados, gerando um novo espaço de entrada, e um mapeamento não-linear do novo espaço de entrada para o espaço de saída. A escolha do pré-processamento e a aproximação do mapeamento não-linear de entrada-saída podem ser tratadas de forma independente (modelos de entrada-saída) ou conjunta (modelos por espaço de estados).
- O novo vetor de entrada $\mathbf{x}(t) \in \mathcal{R}^m$, resultante do pré-processamento do vetor de entrada original, é uma função de entradas externas defasadas no tempo, $\mathbf{u}^{t-1}$, de saídas defasadas no tempo, $\mathbf{y}^{t-1}$, e de ruídos aleatórios expressando os resíduos do processo de identificação, $\boldsymbol{\varepsilon}^{t-1}$, na forma:

$$\mathbf{x}(t) = \boldsymbol{\varphi}(\mathbf{u}^{t-1}, \mathbf{y}^{t-1}, \boldsymbol{\varepsilon}^{t-1}). \quad (7.8)$$

Com isso, o sistema dinâmico da equação (7.7) pode ser reescrito como segue:

$$\mathbf{y}(t) = g(\mathbf{x}(t)) + \mathbf{v}(t). \quad (7.9)$$

SJÖBERG (1995) apresenta estratégias de determinação do vetor de entrada $\mathbf{x}(t)$, dado pela equação (7.8), utilizando procedimentos originalmente desenvolvidas para o caso linear, suplementados por técnicas de validação (veja também LJUNG (1987) e SJÖBERG *et al.* (1995)). Uma vez determinado $\mathbf{x}(t)$, resulta um problema de aproximação de uma função multivariável $g: \mathcal{R}^m \rightarrow \mathcal{R}^r$ a partir de dados amostrados $\{\mathbf{x}(t), \mathbf{y}(t)\}_{t=1}^N$. Este é um problema de identificação de sistemas com formulação bastante apropriada para aplicação dos modelos de redes neurais apresentados nos capítulos 5 e 6 (veja também LEVIN & NARENDRA (1995)).

Sempre que o vetor de entrada $\mathbf{x}(t)$ considerar informações realimentadas do próprio modelo de identificação (por exemplo, estimativas do vetor de saída $\mathbf{y}(t)$ fornecidas pelo modelo de identificação são geralmente empregadas na definição dos resíduos do processo de identificação), uma rede neural recorrente deve ser utilizada como modelo de identificação. Caso contrário, o modelo de identificação pode ser adequadamente representado por uma rede neural não-recorrente. Como de hábito, o modelo de identificação é uma função $\hat{g}: \mathcal{R}^m \rightarrow \mathcal{R}^r$.

7.3.2 Aplicações do modelo de identificação

Existem quatro aplicações básicas para os modelos de identificação $\hat{g}: \mathcal{R}^m \rightarrow \mathcal{R}^r$:

- utilização em tarefas de predição ou simulação;

- auxílio no desenvolvimento de controladores para sistemas dinâmicos (método de controle indireto);
- participação na composição de um controlador adaptativo para o sistema dinâmico (controle adaptativo direto);
- utilização em avaliação de desempenho e diagnóstico de comportamento anormal do sistema dinâmico.

Muitas aplicações de engenharia requerem a solução de problemas de controle e otimização que dependem fundamentalmente da obtenção de soluções inversas para sistemas não-lineares, ou seja, identificação da dinâmica inversa. Ao contrário do caso linear, para o qual existem métodos padrões para solução de mapeamentos inversos, tais como inversão e pseudo-inversão de matrizes (RAO & MITRA, 1971), não existe um método geral e padronizado de solução de problemas de identificação da dinâmica não-linear inversa. Neste caso, por motivos já abordados anteriormente, o modelo de identificação $\hat{g}$ pode ser prontamente aplicado a processos de identificação de dinâmica inversa (BARTO, 1990; HEERMANN, 1992).

7.3.3 Comparação entre modelos de entrada-saída e por espaço de estados

No caso de modelos de entrada-saída, a escolha do conjunto de informações disponíveis até o instante $t-1$, que vai estar representado no vetor de entrada $\mathbf{x}(t)$ dado pela equação (7.8), não é uma tarefa simples, tendo influência direta no desempenho do processo de identificação de g . Esta influência pode ser determinada com base na comparação entre a identificação utilizando modelos lineares e não-lineares:

- utilizando-se modelos lineares de identificação, cada componente adicional do vetor de entrada $\mathbf{x}(t)$ implica mais um parâmetro, ou conjunto de parâmetros, a ser estimado. Esta interpretação também é válida no caso de modelos não-lineares, com o agravante de que um componente adicional em $\mathbf{x}(t)$ representa uma dimensão a mais no espaço de entrada $\mathcal{R}^m$ a ser mapeado no espaço de saída $\mathcal{R}$ pelo modelo de identificação $\hat{g}$;

Os elementos do vetor de entrada $\mathbf{x}(t)$ são funções dos elementos dos seguintes conjuntos:

1. $\mathbf{u}^t$: valores do vetor de entrada externa até o instante t (também denominadas variáveis exógenas);
2. $\mathbf{y}^t$: valores do vetor de saída até o instante t , medidos do sistema a ser identificado;
3. $\hat{\mathbf{y}}^t$: valores preditos do vetor de saída até o instante t , fornecidos pelo modelo de identificação $\hat{\mathbf{g}}$;

O tipo de modelo de identificação está associado à escolha dos conjuntos que vão auxiliar na definição dos elementos de $\mathbf{x}(t)$, produzindo:

Modelo não-linear de resposta ao impulso finita (NFIR)

- utiliza apenas elementos do conjunto 1;
- tem a vantagem de produzir um modelo sempre estável;
- pelo fato de assumir que a dinâmica do sistema depende apenas de entradas externas passadas, tem a desvantagem de requerer uma dimensão elevada para o vetor de entrada $\mathbf{x}(t)$ no caso de dinâmicas com resposta lenta. Ou seja, se o tempo máximo de resposta a qualquer mudança na entrada externa é T e o período de amostragem é δ , então a dimensão de $\mathbf{x}$ deve ser da ordem de T/δ , que pode ser um valor elevado (SJÖBERG *et al.*, 1995).

Modelo não-linear autoregressivo com entrada externa (NARX)

- utiliza elementos dos conjuntos 1 e 2;
- qualquer sistema não-linear genérico pode ser identificado arbitrariamente bem utilizando um modelo NARX (LEVIN & NARENDRA, 1995; SJÖBERG, 1995);
- quando comparado com o modelo NFIR, tem a vantagem de admitir representações mais adequadas para dinâmicas com resposta lenta (SJÖBERG *et al.*, 1995);
- é mais flexível no tratamento do processo de modelagem do ruído (LJUNG, 1987; SJÖBERG, 1995), apesar de não existir um modelo específico para o ruído, o qual é modelado juntamente com a dinâmica do sistema. Com isso, modelos NARX geralmente requerem dimensões elevadas para o vetor de entrada $\mathbf{x}(t)$ como condição necessária para uma boa identificação das propriedades do ruído (SJÖBERG, 1995);
- na impossibilidade de se gerar os elementos do conjunto 2, estes podem ser aproximados pelas correspondentes estimativas presentes no conjunto 3. No entanto, a substituição por valores preditos em lugar dos efetivamente produzidos pelo sistema pode conduzir a modelos instáveis, incapazes de filtrar o erro de estimação (SJÖBERG *et al.*, 1995);

Modelo não-linear de erro de saída (NOE):

- utiliza elementos dos conjuntos 1 e 3;
- tem a vantagem de dispensar a identificação das propriedades do ruído, pois esta etapa passa a estar inerentemente associada à identificação do sistema dinâmico;
- tem a desvantagem de requerer processos de ajuste de parâmetros mais complexos, pois estes devem considerar também o efeito da recorrência (realimentação da informação de saída).

Modelo não-linear autoregressivo, de média móvel e com entrada externa (NARMAX):

- utiliza elementos dos conjuntos 1, 2 e 3;
- tem a vantagem de permitir uma maior flexibilidade no processo de modelagem do ruído, o qual pode ser estimado tomando-se o resíduo $\varepsilon(t) = \mathbf{y}(t) - \hat{\mathbf{y}}(t)$. A presença de entradas diretamente vinculadas ao ruído evita a necessidade de se utilizar um grande número de elementos dos conjuntos 1 e 2 para incorporar a descrição do ruído, como ocorre no caso dos modelos NARX (SJÖBERG, 1995);
- tem a desvantagem de requerer processos de ajuste de parâmetros mais complexos, pois estes devem considerar também o efeito da recorrência (realimentação da informação de saída).

Observe que a diferença básica entre os modelos de identificação está na forma de modelagem do ruído. Além disso, enquanto no caso de identificação linear a escolha do vetor de entrada $\mathbf{x}(t)$ (pré-processamento dos dados) define completamente a estrutura do modelo de identificação (LJUNG, 1987), no caso não-linear existe ainda a necessidade de se aproximar o mapeamento não-linear g (veja equação 7.9). Isto significa que os modelos NFIR, NARX, NOE e NARMAX correspondem a famílias de modelos de entrada-saída para identificação de sistemas.

Uma generalização destas famílias de modelos de entrada-saída pode ser obtida por meio de modelos por espaço de estados, dados na forma:

$$\begin{cases} \mathbf{z}(t+1) = h(\mathbf{z}(t), \mathbf{u}(t), \omega(t)) \\ \mathbf{y}(t) = g(\mathbf{z}(t), \mathbf{u}(t)) + v(t) \end{cases}, \quad (7.10)$$

onde $\mathbf{z}$ é o vetor de estados, $\mathbf{u}$ é o vetor de entrada e $\mathbf{y}$ é o vetor de saída. Observe que agora existem duas fontes de ruído: o ruído associado à dinâmica, $\omega(t)$, e o ruído associado ao processo de medição, $v(t)$.

A equação de estado $\mathbf{z}(t+1) = h(\mathbf{z}(t), \mathbf{u}(t), \omega(t))$ pode ser vista como uma forma de se selecionar o vetor de entrada, e equivale a um pré-processamento dos dados para a equação de saída $\mathbf{y}(t) = g(\mathbf{z}(t), \mathbf{u}(t)) + \varepsilon(t)$. Com isso, enquanto os modelos de entrada-saída têm um vetor de entrada $\mathbf{x}(t)$ pré-especificado, os modelos por espaço de estados apresentam um vetor de estados $\mathbf{z}(t)$, equivalente ao vetor de entrada $\mathbf{x}(t)$ e estimado a partir dos dados. Isto significa que os modelos por espaço de estados têm sua flexibilidade compartilhada entre a definição do vetor de entrada e o mapeamento não-linear g , enquanto que a flexibilidade dos modelos de entrada-saída está concentrada no mapeamento não-linear g (veja equação 7.9). Portanto, a maior flexibilidade de modelos por espaço de estados pode levar a descrições mais sucintas do sistema dinâmico, utilizando um menor número de parâmetros (SJÖBERG, 1995).

A figura 7.8 ilustra os vários tipos de modelos de identificação apresentados, assumindo a utilização de redes neurais artificiais no processo de síntese. O mapeamento não-linear $\hat{g}$ corresponde a redes neurais paramétricas ou não-paramétricas, conforme discutido no capítulo 5, e $\hat{g}_{rec}$ a redes neurais recorrentes, do tipo apresentado no capítulo 6.

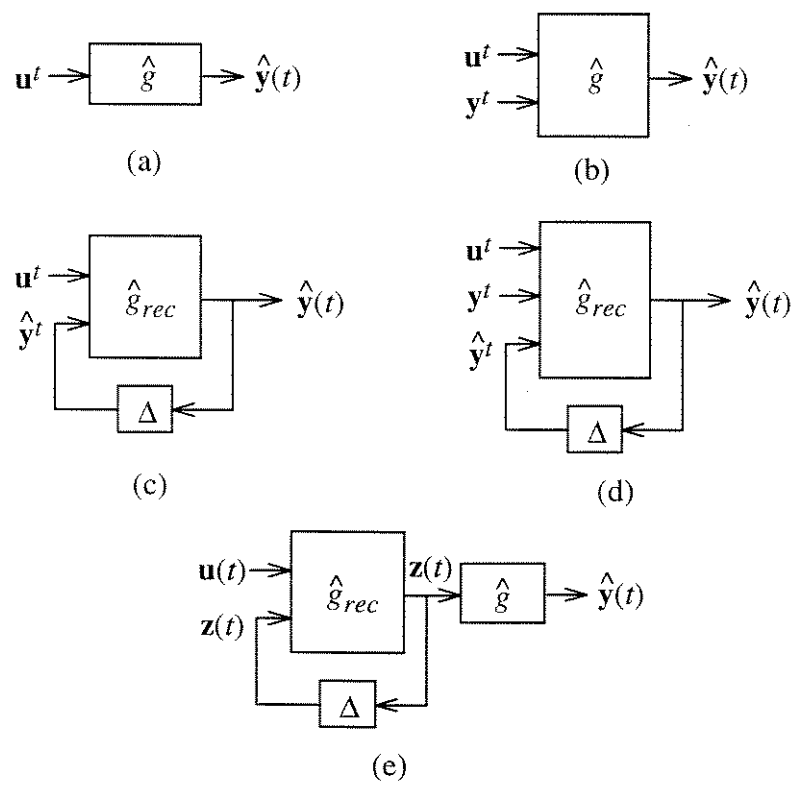


Figura 7.8 - Modelos não-lineares de identificação: (a) NFIR; (b) NARX; (c) NOE; (d) NARMAX; (e) espaço de estados

Dependendo do problema de identificação a ser tratado, os modelos de identificação utilizando redes neurais artificiais apresentados na figura 7.8 podem ser adequadamente simplificados para representar modelos mais específicos. Além disso, as técnicas não-lineares de identificação associadas a cada modelo podem ser utilizadas apenas como procedimentos subseqüentes e complementares à aplicação de modelos lineares. Neste caso, a estrutura não representada pelo modelo linear pode então ser adequadamente representada por um modelo não-linear baseado em redes neurais artificiais. Iniciativas deste tipo garantem de antemão a obtenção de um modelo final pelo menos tão eficaz quanto o modelo linear isolado (SJÖBERG, 1995).

Como um exemplo de aplicação, VON ZUBEN (1993) utilizou uma rede neural recorrente paramétrica na identificação da dinâmica de uma máquina de indução, permitindo desenvolver uma versão não-linear de um observador de estados como parte de uma malha de controle adaptativo direto. Veja seção 7.5.6 e também VON ZUBEN *et al.* (1994a), VON ZUBEN *et al.* (1994b) e VON ZUBEN & NETTO (1994c).

7.4 Previsão de séries temporais

Uma série temporal é uma seqüência de valores medidos em unidades discretas de tempo. Muitas técnicas disponíveis para análise de séries temporais assumem associações lineares entre as variáveis (BOX & JENKINS, 1970), mas, na prática, as variações temporais nos dados não exibem regularidades simples e são de difícil análise e predição. TONG (1983) e TIAO & TSAY (1989) apresentam algumas desvantagens de modelos lineares para análise de séries temporais, dentre as quais destaca-se a dificuldade de explicar mudanças bruscas a intervalos de tempo irregulares. Por outro lado, CHAKRABORTY *et al.* (1992) mostram que redes neurais artificiais constituem modelos não-lineares eficientes na predição de séries temporais, produzindo processos totalmente baseados nos dados amostrados que compõem a série.

Considere o seguinte sistema não-linear auto-regressivo e de média móvel (NARMA):

$$x(t+1) = G(x(t-(n-1)), x(t-(n-2)), \dots, x(t), \varepsilon(t-(m-1)), \varepsilon(t-(m-2)), \dots, \varepsilon(t)), \quad (7.11)$$

onde $x(t)$ representa a saída do sistema no instante t e $\varepsilon(t)$ é um ruído aleatório presente no instante t . Associados respectivamente à saída e ao ruído, os inteiros n e m indicam os números de defasagens no tempo a serem considerados na composição do argumento da função $G(\cdot)$, assumida ser não-linear.

A equação (7.11) pode ser colocada na forma:

$$\mathbf{x}(t+1) = g(\mathbf{x}(t), \boldsymbol{\varepsilon}(t)) \quad (7.12)$$

desde que se torne:

$$\begin{aligned} \bullet \quad \mathbf{x}(t) &= \begin{bmatrix} x(t-(n-1)) \\ x(t-(n-2)) \\ \vdots \\ x(t) \end{bmatrix} & \bullet \quad \boldsymbol{\varepsilon}(t) &= \begin{bmatrix} \varepsilon(t-(m-1)) \\ \varepsilon(t-(m-2)) \\ \vdots \\ \varepsilon(t) \end{bmatrix}; \\ \bullet \quad g(\mathbf{x}(t), \boldsymbol{\varepsilon}(t)) &= \begin{bmatrix} x(t-(n-2)) \\ \vdots \\ x(t) \\ G(x(t-(n-1)), x(t-(n-2)), \dots, x(t), \varepsilon(t-(m-1)), \varepsilon(t-(m-2)), \dots, \varepsilon(t)) \end{bmatrix}. \end{aligned}$$

Com base nos resultados obtidos em capítulos anteriores para o sistema dado pela equação (7.12), redes neurais recorrentes paramétricas podem ser utilizadas como aproximadores universais para o sistema dado pela equação (7.11). Vale salientar que LAPEDES & FARBER (1987) foram os primeiros a propor redes neurais recorrentes como modelos para predição de séries temporais. Para uma discussão abrangente envolvendo outros modelos para predição de séries temporais, veja PRIESTLEY (1988) e HARVEY (1990).

Logo, para uma série temporal gerada pelo sistema dinâmico da equação (7.11), a predição de múltiplos passos à frente utilizando uma rede neural recorrente pode ser obtida a partir do modelo de aproximação apresentado na figura 7.9(a), por simples recursão a partir de uma condição inicial em um instante t .

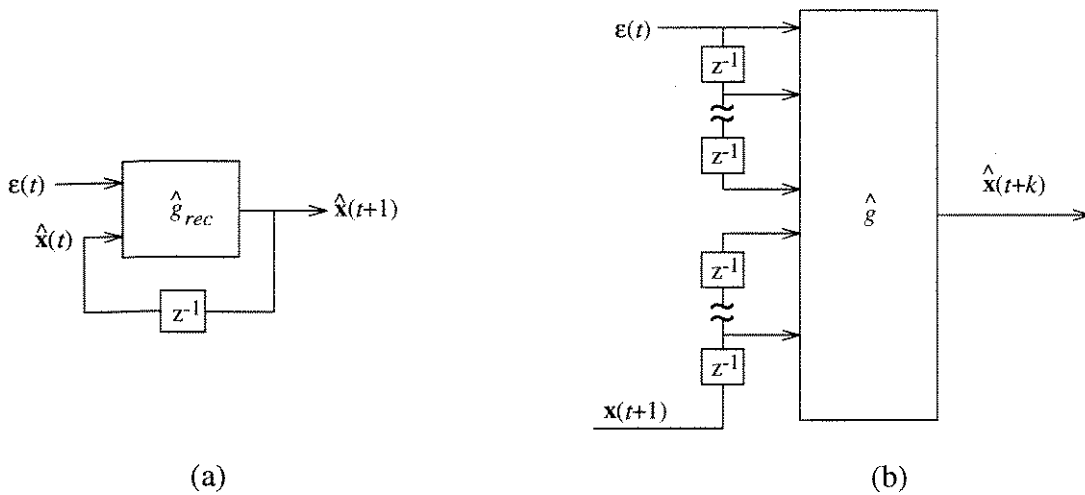


Figura 7.9 - Modelos para predição de séries temporais

(a) previsão de múltiplos passos a frente

(b) previsão de k passos à frente (k fixo, mas arbitrário)

Este modelo de predição é estruturalmente distinto do modelo por linha de derivação de atraso apresentado na figura 7.9(b). O modelo por linha de derivação de atraso utiliza uma rede neural não-recorrente e é geralmente empregado na literatura para predição de k passos à frente, com k arbitrário mas fixo (BAUER & GEISEL, 1990).

As desvantagens do modelo por linha de derivação de atraso quando comparado com a rede neural recorrente são enunciadas a seguir, com o auxílio da figura 7.10, em que os modelos das figuras 7.9(a) e 7.9(b) são apresentados numa mesma estrutura.

- o processamento não é temporal, pois o tempo é mapeado espacialmente pela linha de derivação de atraso, aumentando consideravelmente a dimensão do espaço de aproximação;
- o uso de redes não-recorrentes requer treinamento supervisionado com "teacher forcing" (JORDAN, 1985), feito com a chave em A (veja figura 7.10). Após o treinamento, o processo de predição é feito com a chave em B (veja figura 7.10), provocando uma sensível degradação de desempenho do modelo, pois este não é ajustado para filtrar ruído associado com a possível ocorrência de erro de predição (VON ZUBEN & NETTO, 1995b). Por outro lado, a rede neural recorrente deve se adaptar, já na fase de treinamento, para filtrar eficientemente o ruído proveniente da realimentação de previsões incorretas, de forma que este ruído realimentado não tenha efeito cumulativo (VON ZUBEN & NETTO, 1995b). Isto significa que, no caso da rede neural recorrente, a chave vai estar sempre posicionada em B (veja figura 7.10), tanto na fase de treinamento como de operação. É importante mencionar que, na prática, o processo de aproximação não vai fornecer um solução exata, ou seja, o erro de predição vai estar sempre presente, mesmo que este tenha média zero;
- enquanto a rede neural recorrente da figura 7.9(a) aproxima *a dinâmica que gera a série*, a rede neural da figura 7.9(b) aproxima *a série* (LAPEDES & FARBER, 1987), ou seja, só a rede neural recorrente é capaz de operar o tempo com base em seu efeito real no processamento. Aproximar a dinâmica que gera a série é descobrir as leis fundamentais que regem o comportamento da série, enquanto aproximar a série é descobrir as regularidades ou periodicidades dominantes na série (WEIGEND *et al.*, 1990). O principal problema associado ao processo de aproximar a série é que as regularidades ou periodicidades geralmente não são evidentes e, além disso, podem estar mascaradas por ruído (CHAKRABORTY *et al.*, 1992).

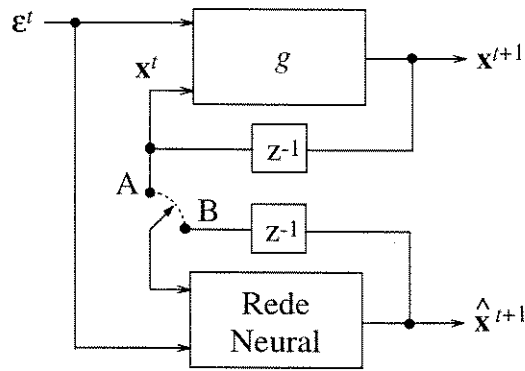


Figura 7.10 - Predição de séries temporais utilizando uma rede neural:

- não-recorrente – chave em A para treinamento e chave em B para operação (não está traçada a linha de derivação de atraso);
- recorrente – chave em B tanto para treinamento como para operação.

Portanto, no contexto de previsão de séries temporais, as redes neurais recorrentes têm vantagens sobre as redes não-recorrentes comparáveis às vantagens de modelos lineares auto-regressivos e de média móvel sobre modelos lineares auto-regressivos (CONNOR *et al.*, 1992; CONNOR *et al.*, 1994).

O processo de treinamento para redes recorrentes adotado neste estudo (também válido para identificação de sistemas dinâmicos) é equivalente ao algoritmo semi-direto empregado para ajuste de parâmetros em filtragem recursiva (WIDROW & STEARNS, 1985), assumindo conhecimento do estado inicial do processo que gera a série temporal (veja figura 7.11). Caso o estado inicial não seja conhecido completamente ou seja totalmente desconhecido (situações praticamente inexistentes em previsão de séries temporais, mas muito comuns em identificação de sistemas dinâmicos), NERRAND *et al.* (1993) apresenta alternativas via algoritmo de propagação direta (WILLIAMS & ZIPSER, 1989a), que não serão abordadas neste estudo.

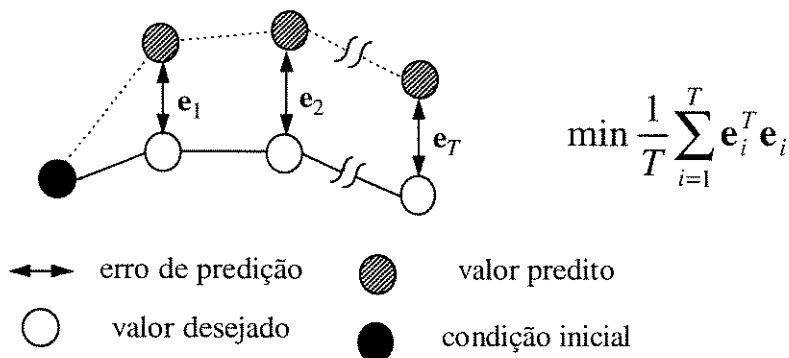


Figura 7.11 - Princípio de operação e problema de otimização a ser resolvido pelo algoritmo de treinamento supervisionado para redes neurais recorrentes

Como um exemplo de simulação, considere uma série temporal gerada a partir do seguinte sistema dinâmico não-linear homogêneo:

$$x(t + 1) = \alpha x(t)[1 - x(t)]. \tag{7.13}$$

Conforme apresentado na figura 7.12, para $\alpha = 4$, este sistema dinâmico apresenta dinâmica caótica, com a série temporal assumindo valores em todo o intervalo $[0,1]$. A mesma rede neural paramétrica com 15 neurônios na camada intermediária foi treinada para aproximar esta série temporal em duas situações: considerando recorrência (figura 7.9(a)) e sem considerar recorrência, ou seja, empregando "teacher-forcing" (figura 7.9(b)). Não foi utilizada linha de derivação de atraso no caso não-recorrente para possibilitar a comparação entre os dois modelos, que apresentam assim o mesmo número de entradas. Também não foi considerada a entrada correspondente ao ruído aleatório, ou seja, os modelos de predição não são de média móvel.

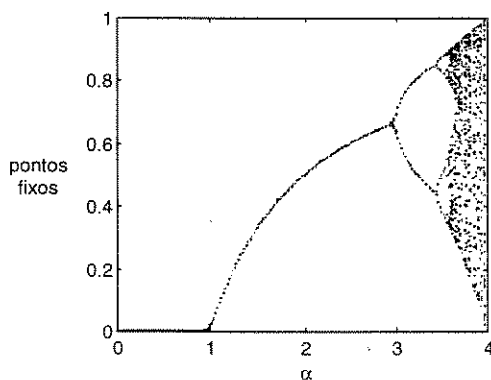


Figura 7.12 - Diagrama de bifurcação dos pontos fixos da dinâmica da equação (7.13)

O mapeamento a ser aproximado pela rede neural nas duas situações é apresentado na figura 7.13. Por ser uma série temporal caótica, qualquer modelo de predição que não realize a identificação exata do mapeamento da figura 7.13 vai apresentar divergência exponencial a partir de uma condição inicial coincidente.

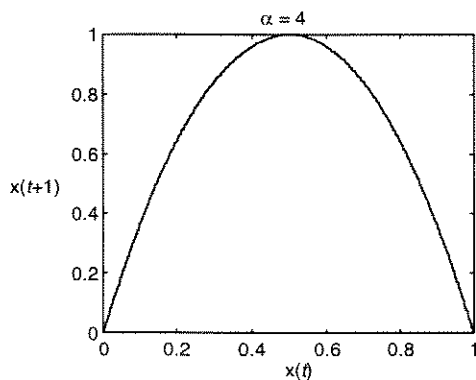


Figura 7.13 - Mapeamento entre estados subsequentes produzido pela equação (7.13)

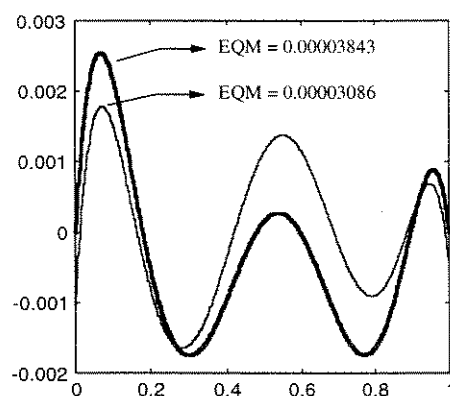


Figura 7.14 - Erro de aproximação referente ao mapeamento da figura 7.13 para a rede neural recorrente (linha em negrito) e não-recorrente com "teacher forcing" (linha normal)

No entanto, embora a rede neural não-recorrente com "teacher forcing" apresentasse um processo de treinamento menos custoso computacionalmente e um erro quadrático médio menor com relação ao mapeamento da figura 7.13 quando comparada à rede neural recorrente (veja figura 7.14), a figura 7.15 mostra o resultado da média do erro de predição de múltiplos passos para diversas condições iniciais, utilizando os dois modelos de redes neurais, mostrando que a rede recorrente é capaz de retardar o processo de divergência na predição da série.

Isto significa que o processo de treinamento levando-se em conta a recorrência não se ocupa apenas da redução do erro de aproximação, mas também da distribuição deste erro por toda a região de aproximação de forma a reduzir seu efeito cumulativo no processo de recursão. Com isso, mesmo com um erro de aproximação maior (erro quadrático médio), seu desempenho é superior (veja figura 7.15). Vale salientar que o ajuste de parâmetros da rede recorrente considerou exemplos de treinamento compostos apenas por dados até a oitava iteração ($T = 8$ na figura 7.11).

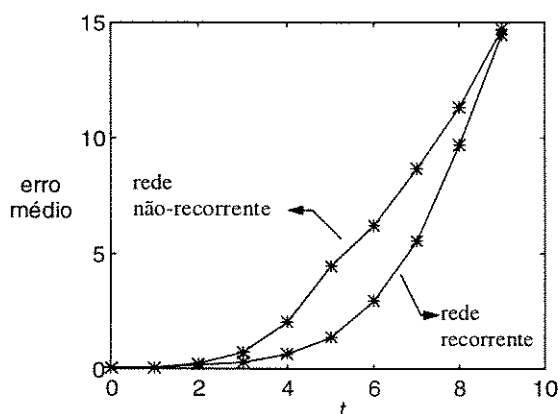


Figura 7.15 - Erro médio de previsão de múltiplos passos à frente

7.5 Redes neurais e controle adaptativo

Foi discutido em capítulos anteriores a capacidade de aprendizado e de generalização presente em redes neurais artificiais. Na aplicação de redes neurais a processos de controle, o aprendizado está diretamente associado ao conceito de adaptação em teoria de controle.

O aprendizado pode ser realizado anteriormente à fase de aplicação ("off-line") ou durante a aplicação ("on-line") da rede neural. Estes dois tipos de aprendizado são essencialmente idênticos do ponto de vista matemático, diferindo apenas do ponto de vista de modelagem, já que a informação disponível e o seu processamento em cada caso são completamente distintos.

Não é a intenção deste estudo fornecer uma descrição detalhada de métodos de aplicação de redes neurais a controle de processos, inclusive por haver uma vasta bibliografia tratando deste tema (BARTO, 1990; HUNT *et al.*, 1992; SONTAG, 1993; VON ZUBEN, 1993; ZBIKOWSKI, 1994; NARENDRA & MUKHOPADHYAY, 1994). No entanto, é importante mencionar uma distinção comumente encontrada na literatura envolvendo controle adaptativo e controle por aprendizado:

- controle adaptativo: tendo sido feita a distinção entre variáveis e parâmetros do sistema a ser controlado, o objetivo do processo adaptativo é ajustar os parâmetros do controlador de forma a compensar possíveis perturbações dinâmicas no sistema causadas por variações de parâmetros, assumidas suaves. Sendo assim, na presença de variações de parâmetros bruscas, é comum a ocorrência de uma degradação de desempenho até que o processo adaptativo produza o efeito de compensação desejado;
- controle por aprendizado: o objetivo do aprendizado é definir estratégias de controle apropriadas para cada região do espaço de parâmetros, de forma a produzir, no limite, uma ação de controle específica para cada ponto de operação do sistema a ser controlado. Dessa forma, substitui-se a tarefa de adaptação pela de identificação do ponto de operação, um processo típico de classificação de padrões.

7.5.1 O paradigma de controle

A aplicação de redes neurais a controle de processos requer a definição preliminar de um paradigma básico de controle, sendo adotada neste estudo a formalização utilizada por SONTAG (1993) e apresentada na figura 7.16. No caso linear, este paradigma de controle é

muito difundido na descrição de sistemas realimentados (DOYLE *et al.*, 1989; IGLESIAS & GLOVER, 1991).

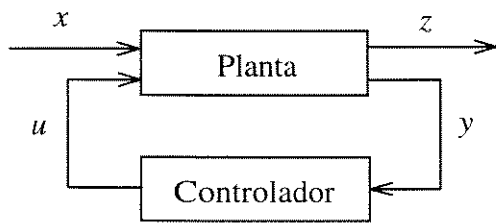


Figura 7.16 - O paradigma de controle

A planta representa o sistema a ser controlado. O sinal x contém todas as entradas que não podem ser diretamente afetadas pelo controlador, incluindo perturbações e sinais de referência. O sinal y representa as variáveis de saída medidas, podendo conter também sinais de entrada via conexão direta. O sinal u corresponde à parte da entrada que pode ser manipulada, enquanto que o sinal z engloba os objetivos de controle. Por exemplo, o sinal z pode representar a diferença entre uma certa função dos estados do sistema e uma função de referência, sendo que o objetivo de controle é a anulação deste sinal.

A introdução de redes neurais como componentes da estrutura de controle está associada a um maior detalhamento do bloco denominado controlador, conforme apresentado na figura 7.17.

O controlador é dividido em duas partes: o sintonizador e o sistema ajustável. O sistema ajustável vai operar sobre a entrada y e representa um sistema dinâmico ou um mapeamento estático com parâmetros ajustáveis. Os parâmetros ajustáveis estão agrupados em um vetor w , atualizado pelo sintonizador. Há também um parâmetro discreto k que determina a flexibilidade do controlador, seja em termos da dimensão do vetor w ou em termos do número de variáveis de estado do controlador. O sintonizador ajusta os parâmetros de acordo com uma lei de adaptação.

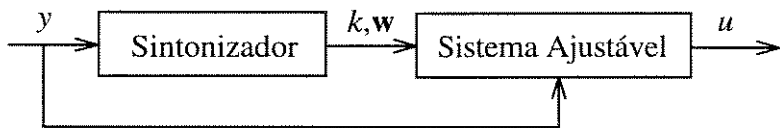


Figura 7.17 - O controlador adaptativo

Tipicamente, o sistema ajustável inclui uma classe de sistemas suficientemente flexíveis para permitir a representação de uma grande variedade de comportamentos estáticos ou dinâmicos. É neste caso que as redes neurais artificiais recorrentes e não-recorrentes descritas em capítulos anteriores surgem como alternativas bastante competitivas, principalmente quando associações não-lineares entre os parâmetros se mostram necessárias. Neste contexto, k teria a função de determinar o número de neurônios ou, de forma mais abrangente, a arquitetura da rede neural, enquanto w incluiria todos os parâmetros ajustáveis da rede neural resultante.

Vale ressaltar, no entanto, que esta estratégia de controle não-linear adaptativo descrita acima, embora permita a utilização de redes neurais em sua implementação, não possui nenhuma propriedade especificamente vinculada às particularidades presentes em redes neurais artificiais.

7.5.2 O conceito de realimentação

A adoção de um paradigma de controle em malha fechada ou por realimentação permite a definição de estratégias de controle mais flexíveis, baseadas no uso apropriado da informação realimentada. A teoria de realimentação passou a ser melhor fundamentada a partir dos estudos realizados por James C. Maxwell. Importantes resultados foram obtidos a partir dos trabalhos de Black, Nyquist e Bode no decorrer dos anos 20 (BLACK, 1977). Atualmente, métodos de controle por realimentação tornam viáveis projetos de alto desempenho, inclusive para sistemas originalmente instáveis ou sujeitos a perturbações não-modeladas.

7.5.3 Controle moderno e controle adaptativo

A teoria de controle moderno engloba as técnicas de controle desenvolvidas a partir do início dos anos 60, tendo como aspecto fundamental a análise no domínio do tempo. Dentre os conceitos que influenciaram substancialmente a pesquisa na área de controle neste período, pode-se citar o princípio do mínimo de Pontryagin, o princípio da programação dinâmica de Bellman e os filtros de Kalman-Bucy e Wiener. Combinados com os conceitos de realimentação, sensibilidade e estabilidade dinâmica, resultados abrangentes e de grande importância prática foram obtidos, conduzindo à teoria de controle adaptativo (ÄSTROM, 1983; ÄSTROM & WITTENMARK, 1989).

Adaptação pode ser definida como uma mudança de comportamento com o objetivo de ajustar-se a novas circunstâncias. Intuitivamente, um controlador adaptativo é um sistema que pode modificar sua própria dinâmica em resposta a mudanças na dinâmica do processo a ser controlado, denominado planta, e em resposta a perturbações externas. Este conceito é de muito interesse para projetistas de sistemas de controle, já que possibilita a um sistema adaptativo ajustar-se a erros no projeto de engenharia e a possíveis falhas e incertezas moderadas de componentes secundários do sistema, aumentando portanto a confiabilidade do projeto.

Em sistemas de controle adaptativo ocorre uma medição contínua e automática das características dinâmicas do processo, uma comparação com as características dinâmicas desejadas ou previamente definidas e a utilização dos resultados desta comparação para ajuste de parâmetros variáveis do controlador ou geração de um sinal atuante de forma que o desempenho ótimo do processo de controle possa ser mantido. Tomando como base uma estrutura de controle convencional em malha fechada, composta por um processo e um controlador com parâmetros ajustáveis, uma representação em diagrama de blocos de um sistema geral de controle adaptativo é apresentada na figura 7.18. Neste sistema, o processo é identificado e o índice de desempenho medido contínua ou periodicamente. Tendo isto sido feito, o índice de desempenho é comparado com o ótimo e uma decisão é tomada. A adaptação do controlador é, portanto, uma operação em malha fechada.

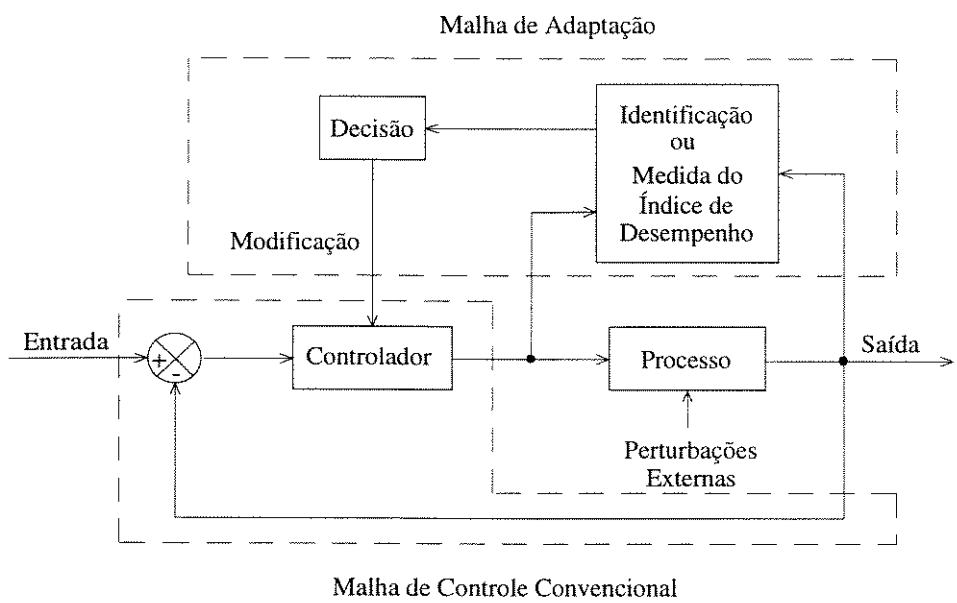


Figura 7.18 - Diagrama de blocos de um sistema de controle adaptativo

A idéia central presente no desenvolvimento de controladores adaptativos é aperfeiçoar o desempenho do controle pela incorporação de mais conhecimento acerca do sistema físico a ser controlado. Dependendo da quantidade de conhecimento incorporada pelo controlador, ele pode prever ou se antecipar ao sistema físico, habilidade importante no cumprimento de determinadas tarefas de controle. Neste contexto, um controlador adaptativo é simplesmente um controlador capaz de adquirir conhecimento acerca do sistema físico em operação.

Os projetistas de sistemas de controle defrontam-se atualmente com a necessidade de expandir seus conceitos e métodos de modo a torná-los aplicáveis a sistemas cujos modelos são incompletos ou inicialmente mal definidos, mas que podem ser aperfeiçoados em fase de operação. Em outras palavras, estes projetistas buscam incorporar a capacidade de modelagem como parte do projeto do sistema. Dessa forma, os modelos passam a ser considerados como entidades em evolução, a partir de uma estrutura inicial.

7.5.4 Controle adaptativo utilizando redes neurais

A teoria linear de controle adaptativo faz uso de resultados da teoria de controle de sistemas lineares invariantes no tempo para definir estruturas de controle para sistemas lineares ou fracamente não-lineares que variam lentamente no tempo (LEVIS *et al.*, 1987).

A principal dificuldade em estender estes resultados para o caso de sistemas não-lineares é a ausência de uma metodologia de controle não-linear genérica (HABER & UNBEHAUEN, 1990).

Conforme discutido no capítulo 6, baseado no modelo linear genérico:

$$\begin{cases} \mathbf{x}(k+1) = \mathbf{Ax}(k) + \mathbf{Bu}(k) \\ \mathbf{y}(k) = \mathbf{Cx}(k) + \mathbf{Du}(k) \end{cases}, \tag{7.14}$$

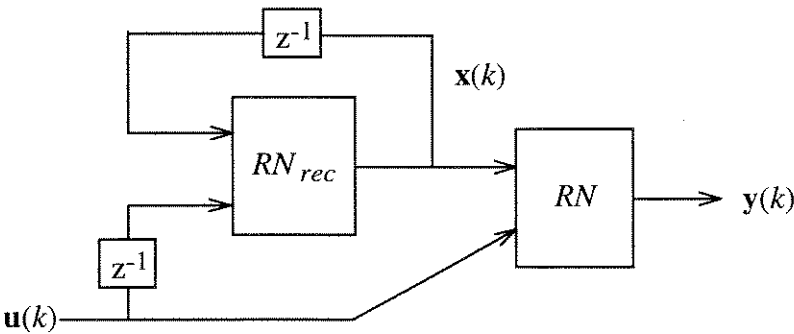


Figura 7.19 - Modelo de um sistema não-linear utilizando redes neurais

com $\mathbf{x} \in \mathbb{R}^n$, $\mathbf{u} \in \mathbb{R}^m$, $\mathbf{y} \in \mathbb{R}^r$ e matrizes $\mathbf{A}$, $\mathbf{B}$, $\mathbf{C}$ e $\mathbf{D}$ de dimensões apropriadas, a rede neural da figura 7.19 se apresenta como uma forte candidata a exercer a função de um modelo não-linear na forma:

$$\begin{cases} \mathbf{x}(k+1) = RN_{rec}(\mathbf{x}(k), \mathbf{u}(k)) \\ \mathbf{y}(k) = RN(\mathbf{x}(k), \mathbf{u}(k)) \end{cases} \quad (7.15)$$

Como sistemas de controle adaptativo são caracterizados pela monitoração da planta e ajuste do controlador em resposta a mudanças de comportamento da planta, a característica adaptativa do modelo da figura 7.19 o coloca como um candidato natural para a implementação de estruturas de controle não-linear adaptativo (VON ZUBEN, 1993; NARENDRA & MUKHOPADHYAY, 1994).

7.5.5 Dificuldade de análise

Enquanto a teoria de controle linear está bastante desenvolvida, apresentando transparência e profundidade em suas formulações, o mesmo não se pode afirmar com relação à teoria de controle não-linear, mesmo resolvendo o problema de representação genérica. Em teoria de controle não-linear, deixam de valer propriedades analíticas importantes, como o princípio da superposição de efeitos e a natureza global de propriedades do sistema detectadas localmente.

Apesar de ser uma representação genérica de sistemas não-lineares, o modelo apresentado na equação (7.15) é de difícil tratamento matemático. Casos particulares deste modelo, na forma:

$$\begin{cases} \mathbf{x}(k+1) = f_1(\mathbf{x}(k)) + f_2(\mathbf{x}(k))^T \mathbf{u}(k) \\ \mathbf{y}(k) = f_3(\mathbf{x}(k)) \end{cases}, \quad (7.16)$$

com f_1 , f_2 e f_3 funções analíticas, já possibilitam a obtenção de resultados de análise mais consistentes, envolvendo propriedades de linearização exata via realimentação, observabilidade e controlabilidade (ISIDORI, 1989; NIJMEIJER & VAN DER SCHAFT, 1990).

7.5.6 Resultados

O processo de adaptação, passo fundamental em teoria de controle adaptativo, pode ser visto como um processo de controle ótimo, com o objetivo de otimizar um índice de desempenho específico para cada problema (ZBIKOWSKI, 1994). Portanto, o processo de

adaptação deve estar baseado na aplicação de técnicas de otimização, como programação dinâmica, programação linear e cálculo variacional.

VON ZUBEN (1993) utilizou parte dos resultados apresentados acima para desenvolver, em termos de simulação computacional, o controle adaptativo de um processo não-linear e não-estacionário (planta) representado por uma máquina de indução com rotor tipo gaiola de esquilo, submetida ao método de controle vetorial de velocidade (veja também VON ZUBEN *et al.*, 1994a; VON ZUBEN *et al.*, 1994b; VON ZUBEN & NETTO, 1994c). A não-estacionariedade da dinâmica da máquina de indução está principalmente associada à variação da resistência de rotor em função da temperatura do material e frequência de operação.

Seja α um parâmetro variante no tempo associado a variáveis de rotor. Sabendo-se que o comando de referência (comando de torque eletromagnético T_e^*) só pode ser seguido com precisão caso se conheça o valor de α a todo instante, a malha de controle deve detectar a variação em α pela observação do efeito desta variação sobre o comportamento da planta (problema de identificação de sistemas, tratado na seção 7.3). A figura 7.20(a) mostra que, se a variação em α não for acompanhada (a estimativa $\hat{\alpha}$ assume que α permanece em seu valor nominal), o desempenho do processo de controle é inadequado para realizar controle de velocidade (o erro na velocidade vai ser proporcional à integral do erro no torque). No entanto, conforme apresentado na figura 7.20(b), a atuação da malha de controle adaptativo na detecção da variação em α garante o acompanhamento do comando de referência, viabilizando o processo de controle de velocidade. A tabela 7.4 apresenta parâmetros da máquina de indução.

Tabela 7.4 - Parâmetros da máquina de indução

Parâmetro	Valor
Número de pólos	8
Frequência	60 Hz
Momento de inércia	0,0016 kg.m ²
Resistência do rotor por fase	3,0167 Ω
Resistência do estator por fase	4,9833 Ω
Indutância de dispersão de rotor	15,7017 mH
Indutância de dispersão do estator	10,4678 mH
Indutância de magnetização	138,4648 mH
Corrente nominal	4,1 A
Tensão nominal	75 V
Potência nominal	1/4 cv
Velocidade nominal de escorregamento	2,8275 rad/s
Torque nominal	2,0113 N.m

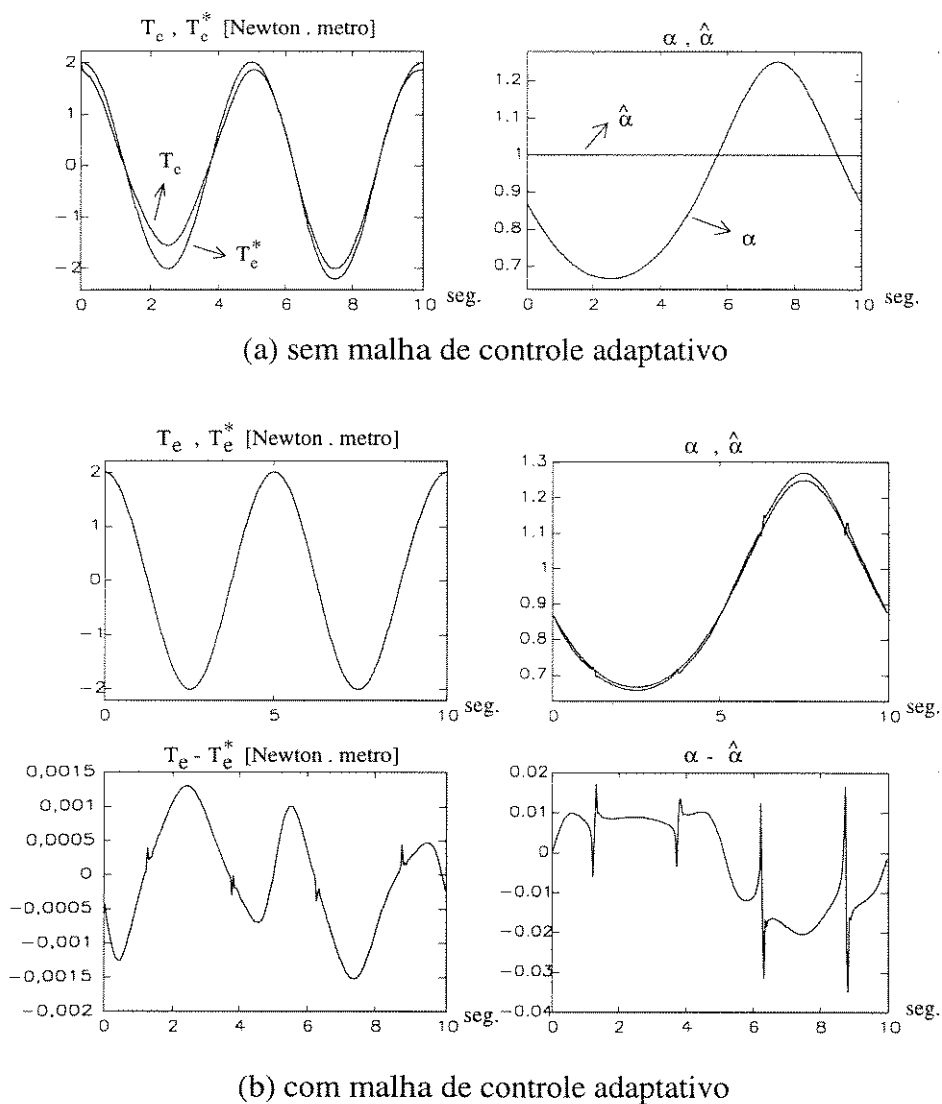


Figura 7.20 - Desempenho do controle de um sistema dinâmico não-estacionário

Para uma análise da aplicação de redes neurais recorrentes de tempo contínuo a controle adaptativo de processos, veja FUNAHASHI & NAKAMURA (1993), SONTAG (1993) e ZBIKOWSKI (1994).

Capítulo 8

Conclusão

O estudo de redes neurais artificiais com base em conceitos de teoria de aproximação de funções e análise numérica conduz a resultados de análise e síntese fundamentais para a aplicação eficaz destas estruturas conexionistas a variados tipos de problemas não-lineares, envolvendo basicamente aproximação de funções e identificação de sistemas dinâmicos. A incorporação de resultados teóricos existenciais e construtivos com relação à capacidade de aproximação e propriedades de convergência de redes neurais artificiais permite aumentar a eficiência e flexibilidade destas estruturas de processamento na solução de problemas com alto grau de complexidade.

Três arquiteturas de rede neural para treinamento supervisionado foram investigadas mais detalhadamente, resultando:

- um modelo paramétrico estático;
- um modelo não-paramétrico estático;
- um modelo paramétrico dinâmico.

Estes modelos devem ser ajustados de forma a otimizar critérios de desempenho adequadamente definidos para cada problema de aplicação.

O modelo paramétrico estático corresponde a uma rede neural não-recorrente com uma camada intermediária, estruturado na forma de combinações lineares de funções básicas definidas previamente. O processo de construção desta estrutura paramétrica vai atuar na translação, rotação e escalamento destas funções básicas, empregando técnicas não-lineares de ajuste de parâmetros. Isto diferencia este modelo conexionista de outros métodos utilizando funções básicas, como séries de Fourier e séries de polinômios ortogonais, que requerem ajuste linear de parâmetros na forma de soluções fechadas, mas exigem um número maior de funções básicas.

O modelo não-paramétrico estático também corresponde a uma rede neural não-recorrente com uma camada intermediária, gerada a partir de um algoritmo iterativo desenvolvido para resolver o problema de otimização do referido critério de desempenho. A partir de um conjunto inicial de informações disponíveis, este algoritmo fornece funções de ativação para os neurônios da camada intermediária e uma política de composição das ativações na produção de estruturas de processamento capazes de utilizar os recursos computacionais disponíveis de forma ótima.

Já o modelo paramétrico dinâmico corresponde a uma rede neural multicamada com recorrência externa. Apesar de não garantir a otimização de recursos computacionais, é demonstrado que este tipo de rede neural recorrente, quando adequadamente dimensionado, representa uma solução eficaz para o problema de geração de modelos dinâmicos não-lineares genéricos. Contudo, ao mesmo tempo em que as não-linearidades presentes e a realimentação de informação garantem capacidade de representação universal para esta estrutura conexionista, acabam também por dificultar sobremaneira a análise e interpretação dos resultados. Por ser um modelo dinâmico de tempo discreto, esta rede neural recorrente permite tratar o tempo pelo seu efeito real no processamento, conduzindo a abordagens eficazes para problemas de identificação de sistemas dinâmicos, etapa fundamental em problemas de previsão de séries temporais e controle de processos.

Enquanto as redes neurais não-recorrentes produzem modelos na forma de equações algébricas, as redes neurais recorrentes correspondem a um sistema de equações a diferenças, denominado sistema de funções iterativas. No entanto, em termos de aplicação, tanto a rede neural não-recorrente como a recorrente podem ser vistas como uma composição de um sistema de processamento de sinais (arquitetura da rede neural) e um sistema de adaptação (processo de treinamento supervisionado).

8.1 Contribuições

Em linhas gerais, as principais contribuições deste estudo são as seguintes:

- dedução dos modelos de redes neurais artificiais a partir de teoria de aproximação de funções, teoria de regularização e métodos de redução de dimensionalidade;
- aplicação de conceitos de geometria analítica no aperfeiçoamento de algoritmos construtivos para modelos de projeção, os quais apresentam propriedades de convergência bem definidas e têm redes neurais com uma camada intermediária como caso particular;

- comparação de modelos paramétricos e não-paramétricos via resultados do teorema de Stone-Weierstrass e teorema de Kolmogorov;
- formalização de redes neurais recorrentes de tempo discreto como modelos dinâmicos não-lineares genéricos;
- uso de conceitos de análise de sensibilidade no desenvolvimento de processos de ajuste utilizando informação exata de 2ª ordem para redes neurais multicamadas com recorrência externa, tendo redes neurais multicamadas não-recorrentes como caso particular;
- formulação básica de problemas de aproximação de funções e identificação de sistemas dinâmicos com o objetivo de apresentar a solução via inclusão direta dos modelos de redes neurais artificiais desenvolvidos.

8.2 Extensões

Como tópicos relacionados ao tema deste trabalho, mas não abordados ou apenas referenciados e, portanto, passíveis de um estudo mais aprofundado, é possível citar:

- desenvolvimento de métodos construtivos por composição híbrida, ou seja, envolvendo composição aditiva e multiplicativa;
- ampliação do escopo do método construtivo para a geração de outros tipos de arquiteturas;
- verificação da influência representada pela presença de erro também junto aos dados de entrada;
- uso de resultados da teoria de sistemas não-lineares em um estudo mais aprofundado de propriedades de estabilidade e tipos de comportamento dinâmico que podem ser apresentados pelas redes neurais recorrentes em estudo;
- introdução de técnicas de treinamento não-supervisionado junto aos processos de definição dos modelos de redes neurais, com o objetivo de ampliar o campo de aplicação, incluindo problemas onde a informação disponível é, por si só, insuficiente para permitir a aplicação exclusiva de treinamento supervisionado.
- aplicação a problemas originados a partir de dados reais;
- comparação mais abrangente, em termos de desempenho computacional, entre os modelos de redes neurais artificiais implementados.

Apêndices

Apêndice A

Conceitos básicos de álgebra linear e notação

As definições e axiomas apresentados a seguir representam uma coletânea de conceitos básicos de álgebra linear, fundamentais para a formulação de boa parte dos resultados apresentados neste estudo. A inclusão destes conceitos na forma de apêndice se justifica pela sua utilização generalizada na apresentação dos resultados principais deste trabalho. Sendo assim, no texto da tese, assume-se familiaridade com os tópicos aqui referenciados. Criam-se, portanto, condições para que o texto se ocupe apenas de tópicos mais avançados.

Na apresentação que se segue, os resultados são apenas referenciados, sendo que discussões mais detalhadas podem ser encontradas na literatura (CHEN, 1984; LUENBERGER, 1969; STRANG, 1988).

A.1 Escalar

Uma variável que assume valores no eixo dos números reais é denominada escalar. Os escalares são descritos por letras minúsculas do alfabeto romano expressas em itálico. O conjunto de todos os escalares reais é representado por $\Re$.

- O módulo de um escalar real x é dado na forma: $|x| = \begin{cases} x & \text{se } x \geq 0 \\ -x & \text{se } x < 0 \end{cases}$

A.2 Vetor

Um arranjo ordenado de n escalares $x_i \in \Re$ ($i=1,2,\dots,n$) é denominado vetor de dimensão n . Os vetores são descritos por letras minúsculas do alfabeto romano expressas em negrito, e assumem a forma de coluna ou linha:

$$\mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} \quad \text{ou} \quad \mathbf{x} = [x_1 \ x_2 \ \cdots \ x_n].$$

Para coerência de notação, salvo menção contrária, o texto desta dissertação considera os vetores na forma de coluna. O conjunto de todos os vetores de dimensão n com elementos reais é representado por $\mathfrak{R}^n$.

- Um escalar é um vetor de dimensão 1.
- Vetor $\mathbf{0}_n$: é o vetor nulo de dimensão n , com todos os elementos iguais a zero. O subscrito n é suprimido quando não há margem à dúvida.
- Vetor $\mathbf{1}_n$: é o vetor de dimensão n com todos os elementos iguais a 1.
- Vetor $\mathbf{e}_i$: é o vetor normal unitário de dimensão n (a dimensão deve ser indicada pelo contexto) com todos os elementos iguais a 0, exceto o i -ésimo elemento que é igual a 1. Neste caso, $1 \leq i \leq n$.

A.3 Matriz

Um arranjo ordenado de $m \cdot n$ escalares x_{ij} ($i=1,2,\dots,m; j=1,2,\dots,n$) é denominado matriz de dimensão $m \times n$. As matrizes são descritas por letras maiúsculas do alfabeto romano expressas em negrito, e assumem a forma:

$$\mathbf{X} = \begin{bmatrix} x_{11} & x_{12} & \cdots & x_{1n} \\ x_{21} & x_{22} & \cdots & x_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ x_{m1} & x_{m2} & \cdots & x_{mn} \end{bmatrix}.$$

O conjunto de todas as matrizes $m \times n$ com elementos reais é representado por $\mathfrak{R}^{m \times n}$.

- As colunas da matriz $\mathbf{X}$ são vetores-coluna descritos por $\mathbf{x}_i = \begin{bmatrix} x_{1i} \\ x_{2i} \\ \vdots \\ x_{mi} \end{bmatrix}$, $i=1,\dots,n$.
- As linhas da matriz $\mathbf{X}$ são vetores-linha descritos por $\mathbf{x}_{(j)} = [x_{j1} \ x_{j2} \ \cdots \ x_{jn}]$, $j=1,\dots,m$.
- Um vetor é uma matriz com número unitário de linhas e/ou colunas.

A.4 Conjuntos e operações com conjuntos

Um conjunto pode ser definido como uma agregação de objetos. Os conjuntos são descritos por letras maiúsculas do alfabeto romano expressas em itálico. Por conveniência de notação, alguns conjuntos especiais são descritos por símbolos específicos. Exemplos:

- $\mathbb{N}$: conjunto dos números naturais
- $\mathbb{R}$: conjunto dos números reais
- $\mathbb{C}$: conjunto dos números complexos

O estado lógico ou associação de um elemento x a um conjunto X qualquer é representado por

$$x \in X: x \text{ pertence a } X$$

$$x \notin X: x \text{ não pertence a } X$$

Um conjunto pode ser especificado listando-se seus elementos entre colchetes

$$X = \{x_1, x_2, \dots, x_n\}$$

ou evidenciando uma ou mais propriedades comuns aos seus elementos

$$X_2 = \{x \in X_1 \text{ tal que } P(x) \text{ é verdade}\} \text{ ou } X_2 = \{x \in X_1: P(x)\}$$

As principais operações entre conjuntos são:

- União: $X_1 \cup X_2 = \{x: x \in X_1 \text{ ou } x \in X_2\}$;
 - Interseção: $X_1 \cap X_2 = \{x: x \in X_1 \text{ e } x \in X_2\}$;
- $X_1 \cap X_2 = \emptyset$ (conjunto vazio) se X_1 e X_2 são conjuntos disjuntos.

O complemento de um conjunto X é representado por $\overline{X}$ e é definido na forma:

$$\overline{X} = \{x: x \notin X\}.$$

S é um subconjunto de X , se $x \in S$ implica $x \in X$. Neste caso, diz-se que S está contido em X ($S \subset X$) ou que X contém S ($X \supset S$). Se $S \subset X$ e S não é igual a X , então S é um subconjunto próprio de X .

A.5 Conjuntos especiais de números reais

Se a e b são números reais ($a, b \in \mathbb{R}$), define-se:

$$[a, b] = \{x: a \leq x \leq b\}$$

$$(a, b] = \{x: a < x \leq b\}$$

$$[a, b) = \{x: a \leq x < b\}$$

$$(a, b) = \{x: a < x < b\}$$

Se X é um conjunto de números reais, então o menor limitante superior de X , dado por

$$\bar{x} = \sup x = \sup\{x: x \in X\},$$

é o supremo de X , e o maior limitante inferior de X , dado por

$$\underline{x} = \inf x = \inf\{x: x \in X\},$$

é o ínfimo de X . Se $\bar{x} = +\infty$ então X não é limitado superiormente. De forma análoga, se $\underline{x} = -\infty$ então X não é limitado inferiormente.

A.6 Espaço vetorial linear

Um espaço vetorial linear $(X, \mathfrak{S})$ consiste de um conjunto X de vetores e de um campo $\mathfrak{S}$, sobre os quais estão definidas duas operações:

- **Adição (+)**: mapeamento $X \times X \rightarrow X$ tal que, $\forall \mathbf{x}, \mathbf{y} \in X$, $(\mathbf{x}, \mathbf{y}) \rightarrow (\mathbf{x} + \mathbf{y})$ produz $\mathbf{x} + \mathbf{y} \in X$.
- **Multiplicação por escalar ($\cdot$)**: mapeamento $\mathfrak{R} \times X \rightarrow X$ tal que, $\forall \mathbf{x} \in X$ e $\forall a \in \mathfrak{S}$, $(a, \mathbf{x}) \rightarrow (a \cdot \mathbf{x})$ produz $a \cdot \mathbf{x} \in X$.

Uma série de axiomas podem ainda ser deduzidos para o espaço $(X, \mathfrak{S})$:

- | | |
|---|-------------------------------------|
| (1) $\mathbf{x} + \mathbf{y} = \mathbf{y} + \mathbf{x}$ | (comutatividade na adição) |
| (2) $(\mathbf{x} + \mathbf{y}) + \mathbf{z} = \mathbf{x} + (\mathbf{y} + \mathbf{z})$ | (associatividade na adição) |
| (3) $a \cdot (\mathbf{x} + \mathbf{y}) = a \cdot \mathbf{x} + a \cdot \mathbf{y}$ | (distributividade na adição) |
| (4) $\mathbf{x} + \mathbf{0} = \mathbf{x}$ | (elemento nulo) |
| (5) $(ab) \cdot \mathbf{x} = a \cdot (b \cdot \mathbf{x})$ | (associatividade na multiplicação) |
| (6) $(a + b) \cdot \mathbf{x} = a \cdot \mathbf{x} + b \cdot \mathbf{x}$ | (distributividade na multiplicação) |
| (7) $1 \cdot \mathbf{x} = \mathbf{x}$ | (elemento neutro) |

Exemplos: $(\mathfrak{R}, \mathfrak{R}), (\mathfrak{R}^n, \mathfrak{R}), (\mathfrak{R}^{m \times n}, \mathfrak{R})$, onde n e m são inteiros positivos.

A.7 Subespaço vetorial linear

Seja $(X, \mathfrak{S})$ um espaço vetorial linear e S um subconjunto de X . Diz-se então que $(S, \mathfrak{S})$ é um subespaço linear de $(X, \mathfrak{S})$ se S forma um espaço linear sobre $\mathfrak{S}$ através das mesmas operações definidas sobre $(X, \mathfrak{S})$.

A.8 Variedade linear

A translação de um subespaço é chamada de variedade linear. Uma variedade linear V pode ser descrita na forma:

$$V = \mathbf{x} + S,$$

onde $(S, \mathfrak{S})$ é um subespaço linear de $(X, \mathfrak{S})$ e $\mathbf{x}$ é qualquer elemento de X .

A.9 Conjunto convexo

Diz-se que um conjunto X em um espaço vetorial linear é convexo se dados quaisquer elementos $\mathbf{x}_1, \mathbf{x}_2 \in X$, todos os elementos da forma $a\mathbf{x}_1 + (1-a)\mathbf{x}_2$, com $0 \leq a \leq 1$ também estão em X . São conjuntos convexos:

- o conjunto vazio, por definição;
- subespaços;
- variedades lineares;
- interseção de conjuntos convexos.

A.10 Combinação linear e combinação convexa

Seja $S = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n\}$ um subconjunto de um espaço vetorial linear $(X, \mathfrak{S})$. Combinações lineares de elementos de S são formadas através de

$$a_1\mathbf{x}_1 + a_2\mathbf{x}_2 + \dots + a_n\mathbf{x}_n,$$

onde $a_1, a_2, \dots, a_n \in \mathfrak{S}$.

O conjunto de todas as combinações lineares de elementos de S é chamado de subespaço gerado e é representado por $[S]$. $[S]$ é um subespaço linear de $(X, \mathfrak{S})$.

Se os escalares $a_1, a_2, \dots, a_n \in \mathfrak{S}$ são tais que $a_i \geq 0$ ($i=1,2,\dots,n$) e $\sum_{i=1}^n a_i = 1$, então a combinação linear é chamada combinação convexa dos elementos $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in X$.

A.11 Dependência linear

Diz-se que um vetor $\mathbf{x}$ é linearmente dependente em relação a um conjunto de vetores S se $\mathbf{x}$ pode ser expresso como uma combinação linear de elementos de S , ou seja, se $\mathbf{x} \in [S]$. Caso contrário, diz-se que $\mathbf{x}$ é linearmente independente em relação ao conjunto S . Um conjunto de vetores $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ é linearmente independente se e somente se a equação

$$\sum_{i=1}^n a_i \mathbf{x}_i = \mathbf{0} \text{ tem como única solução } a_i = 0, i = 1, 2, \dots, n.$$

A.12 Base e dimensão de um espaço vetorial linear

Um conjunto S de vetores linearmente independentes forma uma base para $(X, \mathfrak{S})$ se S gera $(X, \mathfrak{S})$. Se $S = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n\}$ é uma base para $(X, \mathfrak{S})$, então diz-se que o espaço vetorial $(X, \mathfrak{S})$ é de dimensão n (ou n -dimensional). Qualquer outra base para o mesmo espaço $(X, \mathfrak{S})$ possui o mesmo número n de elementos.

Qualquer $\mathbf{y} \in X$ pode ser expresso na forma:

$$\mathbf{y} = b_1 \mathbf{x}_1 + b_2 \mathbf{x}_2 + \dots + b_n \mathbf{x}_n = [\mathbf{x}_1 \ \mathbf{x}_2 \ \dots \ \mathbf{x}_n] \cdot \begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_n \end{bmatrix},$$

onde $\mathbf{b} = [b_1 \ b_2 \ \dots \ b_n]^T$ é a representação de $\mathbf{y}$ na base S (o superscrito T representa a operação de transposição).

A representação de um vetor numa determinada base é única.

Se n é finito, diz-se que o espaço vetorial $(X, \mathfrak{S})$ é de dimensão finita. Para n infinito, $(X, \mathfrak{S})$ é dito ser de dimensão infinita.

A.13 Produtos entre vetores

Produto interno

O produto interno é uma função $f_{PI}: \mathfrak{R}^n \times \mathfrak{R}^n \rightarrow \mathfrak{R}$ que associa a cada par de vetores $\mathbf{x}, \mathbf{y} \in \mathfrak{R}^n$ um escalar dado por:

$$\langle \mathbf{x}, \mathbf{y} \rangle = \mathbf{x}^T \mathbf{y} = x_1 y_1 + x_2 y_2 + \dots + x_n y_n,$$

onde $\mathbf{x} = [x_1 \ x_2 \ \dots \ x_n]^T$ e $\mathbf{y} = [y_1 \ y_2 \ \dots \ y_n]^T$.

Produto externo

O produto externo é uma função $f_{PE}: \mathfrak{R}^m \times \mathfrak{R}^n \rightarrow \mathfrak{R}^{m \times n}$ que associa a cada par de vetores $\mathbf{x} \in \mathfrak{R}^m$, $\mathbf{y} \in \mathfrak{R}^n$ uma matriz de dimensão $m \times n$ definida na forma:

$$\langle \mathbf{x}, \mathbf{y} \rangle = \mathbf{xy}^T = \begin{bmatrix} x_1 y_1 & x_1 y_2 & \dots & x_1 y_n \\ x_2 y_1 & \ddots & & \vdots \\ \vdots & & & \\ x_m y_1 & \dots & & x_m y_n \end{bmatrix}$$

A.14 Norma, semi-norma e quase-norma

A norma é uma função $f_N: \mathfrak{R}^n \rightarrow \mathfrak{R}$ que associa a cada vetor $\mathbf{x} \in \mathfrak{R}^n$ um escalar representado por $\|\mathbf{x}\|$. A norma satisfaz os seguintes axiomas:

- $\|\mathbf{x}\| \geq 0, \forall \mathbf{x} \in X; \|\mathbf{x}\| = 0 \Leftrightarrow \mathbf{x} = \mathbf{0};$
- $\|\mathbf{x} + \mathbf{y}\| \leq \|\mathbf{x}\| + \|\mathbf{y}\|, \forall \mathbf{x}, \mathbf{y} \in X$ (desigualdade triangular);
- $\|a\mathbf{x}\| = |a|\|\mathbf{x}\|, \forall \mathbf{x} \in X, \forall a \in \mathfrak{I}.$

Dentre as funções f_N que atendem a estes axiomas, é possível destacar:

- $\|\mathbf{x}\|_p = \left(\sum_{i=1}^n |x_i|^p \right)^{1/p}$ (norma p);
- Para $p = 1$ tem-se $\|\mathbf{x}\|_1 = \sum_{i=1}^n |x_i|;$
- Para $p = 2$ tem-se $\|\mathbf{x}\|_2 = \langle \mathbf{x}, \mathbf{x} \rangle^{1/2} = (\mathbf{x}^T \mathbf{x})^{1/2} = \left(\sum_{i=1}^n |x_i|^2 \right)^{1/2}$ (norma euclidiana);
- Para $p = \infty$ tem-se $\|\mathbf{x}\|_\infty = \max_i |x_i|.$

A semi-norma é uma função $f_{SN}: \mathfrak{R}^n \rightarrow \mathfrak{R}$ que satisfaz todas as propriedades de norma, com exceção do primeiro axioma. O subespaço linear $X_0 \subset \mathfrak{R}^n$ onde $\|\mathbf{x}\| = 0$ é denominado espaço nulo da semi-norma.

A quase-norma é uma função $f_{QN}: \mathfrak{R}^n \rightarrow \mathfrak{R}$ que satisfaz todas as propriedades de norma, com exceção do segundo axioma (desigualdade triangular), o qual assume a forma:

- $\|\mathbf{x} + \mathbf{y}\| \leq b(\|\mathbf{x}\| + \|\mathbf{y}\|), \forall \mathbf{x}, \mathbf{y} \in X, \text{ com } b \in \mathfrak{R}.$

Os conceitos de norma, semi-norma e quase-norma são fundamentais para a definição de propriedades topológicas.

A.15 Vetores ortogonais e ortonormais

Dois vetores $\mathbf{x}, \mathbf{y} \in \mathfrak{R}^n$ são ortogonais entre si, $\mathbf{x} \perp \mathbf{y}$, se o seu produto interno se anula:

$$\mathbf{x} \perp \mathbf{y} \Rightarrow \langle \mathbf{x}, \mathbf{y} \rangle = 0.$$

Além disso, se $\langle \mathbf{x}, \mathbf{x} \rangle = \langle \mathbf{y}, \mathbf{y} \rangle = 1$, então os vetores $\mathbf{x}, \mathbf{y} \in \mathfrak{R}^n$ são ortonormais entre si.

Um vetor $\mathbf{x}$ é ortogonal a um conjunto S ($\mathbf{x} \perp S$) se $\langle \mathbf{x}, \mathbf{s} \rangle = 0, \forall \mathbf{s} \in S$.

A.16 Relação entre produto interno e norma euclidiana

A desigualdade de Cauchy-Schwartz-Buniakowsky

$$|\langle \mathbf{x}, \mathbf{y} \rangle|^2 \leq \langle \mathbf{x}, \mathbf{x} \rangle \langle \mathbf{y}, \mathbf{y} \rangle \Rightarrow |\langle \mathbf{x}, \mathbf{y} \rangle| \leq \|\mathbf{x}\|_2 \cdot \|\mathbf{y}\|_2$$

estabelece uma importante relação entre produto interno e norma euclidiana. A igualdade é válida se e somente se os vetores $\mathbf{x}$ e $\mathbf{y}$ estiverem alinhados.

A.17 Projeção de um vetor em uma determinada direção

Dado um espaço vetorial linear X , seja $\mathbf{y} \in X$ um vetor que fornece uma determinada direção. A projeção de qualquer vetor $\mathbf{x} \in X$ na direção de $\mathbf{y}$ é dada na forma:

$$\text{proj}_{\mathbf{y}}(\mathbf{x}) = \frac{\langle \mathbf{x}, \mathbf{y} \rangle}{\langle \mathbf{y}, \mathbf{y} \rangle^{1/2}} \cdot \frac{\mathbf{y}}{\langle \mathbf{y}, \mathbf{y} \rangle^{1/2}} = \frac{\mathbf{x}^T \mathbf{y}}{\mathbf{y}^T \mathbf{y}} \cdot \mathbf{y}$$

A.18 Vetores ortonormais gerados a partir de vetores linearmente independentes

Apesar de todo conjunto de vetores ortonormais não-nulos ser linearmente independente, nem todo conjunto de vetores linearmente independentes é ortonormal, mas pode passar por um processo de ortonormalização, como segue.

Dado um conjunto de m vetores n -dimensionais ($m \leq n$) linearmente independentes $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m\}$, é sempre possível estabelecer uma combinação linear adequada destes vetores que produza m vetores n -dimensionais mutuamente ortogonais $\{\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_m\}$ que geram o mesmo espaço. Além disso, se os vetores $\mathbf{u}_i$ ($i=1, \dots, m$) apresentarem norma unitária, eles são mutuamente ortonormais.

Um conjunto de vetores ortonormais $\{\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_m\}$ pode ser obtido a partir de um conjunto de vetores linearmente independentes $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m\}$ através do processo de ortogonalização de Gram-Schmidt, o qual pode ser dividido em duas etapas:

Etapa (1) $\mathbf{y}_1 = \mathbf{x}_1$

$$\mathbf{y}_i = \mathbf{x}_i - \sum_{j=1}^{i-1} \frac{\langle \mathbf{x}_i, \mathbf{y}_j \rangle}{\langle \mathbf{y}_j, \mathbf{y}_j \rangle} \mathbf{y}_j, \quad i = 2, \dots, m$$

Etapa (2) $\mathbf{u}_i = \frac{\mathbf{y}_i}{\langle \mathbf{y}_i, \mathbf{y}_i \rangle^{1/2}}, \quad i = 1, \dots, m$

Com isso tem-se

$$\mathbf{u}_1 = a_{11}\mathbf{x}_1$$

$$\mathbf{u}_2 = a_{21}\mathbf{x}_1 + a_{22}\mathbf{x}_2$$

$$\mathbf{u}_3 = a_{31}\mathbf{x}_1 + a_{32}\mathbf{x}_2 + a_{33}\mathbf{x}_3$$

$$\vdots \quad \vdots \quad \vdots \quad \ddots$$

com $a_{ii} > 0$ ($i=1, \dots, m$). Certamente existem outros processos de ortonormalização mais gerais, que não impõem qualquer tipo de restrição aos coeficientes a_{ij} ($i \geq j; i, j=1, \dots, m$).

A.19 Algumas propriedades de matrizes

Simetria

Uma matriz $\mathbf{A} = \{a_{ij}\}$ de dimensão $n \times n$ é dita ser simétrica se $a_{ij} = a_{ji}$, $i, j=1, 2, \dots, n$.

Inversão

Dada uma matriz $\mathbf{A}$ de dimensão $n \times n$, se existe uma matriz $\mathbf{B}$ de mesma dimensão tal que $\mathbf{AB} = \mathbf{BA} = \mathbf{I}$ então $\mathbf{B}$ é única e é denominada a matriz inversa de $\mathbf{A}$, ou seja, $\mathbf{B} = \mathbf{A}^{-1}$.

Pseudo-inversão

Dada uma matriz $\mathbf{A}$ de dimensão $m \times n$, sempre existe uma matriz $\mathbf{B}$ satisfazendo as seguintes igualdades:

- $\mathbf{ABA} = \mathbf{A}$;
- $\mathbf{BAB} = \mathbf{B}$;
- $(\mathbf{AB})^T = \mathbf{AB}$;
- $(\mathbf{BA})^T = \mathbf{BA}$;

A matriz $\mathbf{B}$, geralmente denominada $\mathbf{A}^+$, é única e é conhecida como a inversa de Moore-Penrose ou a pseudo-inversa da matriz $\mathbf{A}$. A matriz $\mathbf{A}^+$ pode ser diretamente determinada através da decomposição de $\mathbf{A}$ em valores singulares.

Se a matriz $\mathbf{A}$ é de posto completo, ou seja, se $\rho(\mathbf{A}) = \min(m,n)$ (veja definição mais adiante), a pseudo-inversa é dada na forma:

- se $m < n$, $\rho(\mathbf{A}) = m$ e $\mathbf{A}^+ = \mathbf{A}^T(\mathbf{A}\mathbf{A}^T)^{-1}$;
- se $m > n$, $\rho(\mathbf{A}) = n$ e $\mathbf{A}^+ = (\mathbf{A}^T\mathbf{A})^{-1}\mathbf{A}^T$;
- se $m = n$, $\rho(\mathbf{A}) = m = n$ e $\mathbf{A}^+ = \mathbf{A}^{-1}$.

Portanto, se $\mathbf{A}$ é uma matriz inversível (matriz quadrada e de posto completo), a pseudo-inversa é igual à inversa.

Cofator

Dada uma matriz $\mathbf{A}$ de dimensão $n \times n$, o cofator do elemento a_{ij} ($i, j = 1, 2, \dots, n$) é dado na forma:

$$c_{ij} = (-1)^{i+j} m_{ij},$$

onde m_{ij} é o determinante da matriz formada eliminando-se a i -ésima linha e a j -ésima coluna da matriz $\mathbf{A}$.

Determinante

Dada uma matriz $\mathbf{A}$ de dimensão $n \times n$, o determinante de $\mathbf{A}$ é dado na forma:

$$|\mathbf{A}| = \det(\mathbf{A}) = \begin{cases} \sum_{i=1}^n a_{ij} c_{ij}, & \text{para qualquer } i \\ \text{ou} \\ \sum_{j=1}^n a_{ij} c_{ij}, & \text{para qualquer } j \end{cases}$$

onde c_{ij} é o cofator do elemento a_{ij} .

Seja $\mathbf{A} = [\mathbf{a}_1 \ \dots \ \mathbf{a}_j \ \dots \ \mathbf{a}_n]$ ($1 \leq j \leq n$). O determinante de $\mathbf{A}$ possui as seguintes propriedades:

- Invariância: $\det([\mathbf{a}_1 \ \dots \ \mathbf{a}_j \ \dots \ \mathbf{a}_n]) = \det([\mathbf{a}_1 \ \dots \ \mathbf{a}_j + \mathbf{a}_k \ \dots \ \mathbf{a}_n])$, $j \neq k$, $1 \leq j, k \leq n$;
- Homogeneidade: $\det([\mathbf{a}_1 \ \dots \ b\mathbf{a}_j \ \dots \ \mathbf{a}_n]) = b \cdot \det([\mathbf{a}_1 \ \dots \ \mathbf{a}_j \ \dots \ \mathbf{a}_n])$

Com isso, se $\mathbf{B}$ é uma matriz obtida de $\mathbf{A}$ pela troca de duas de suas colunas, então $\det(\mathbf{B}) = -\det(\mathbf{A})$.

Traço

Dada uma matriz $\mathbf{A}$ de dimensão $n \times n$, o traço de $\mathbf{A}$, representado por $\text{tr}(\mathbf{A})$, é a soma dos elementos da diagonal de $\mathbf{A}$, ou seja:

$$\text{tr}(\mathbf{A}) = \sum_{i=1}^n a_{ii}$$

Adjunta

Dada uma matriz $\mathbf{A}$ de dimensão $n \times n$, a adjunta de $\mathbf{A}$, representada por $\text{adj}(\mathbf{A})$, é dada na forma:

$$\text{adj}(\mathbf{A}) = \{ a'_{ij} \}$$

onde $a'_{ij} = c_{ji}$, o cofator do elemento a_{ji} .

São válidas as seguintes igualdades: $\begin{cases} \mathbf{A} \cdot \text{adj}(\mathbf{A}) = \det(\mathbf{A}) \cdot \mathbf{I} \\ \text{adj}(\mathbf{A}) \cdot \mathbf{A} = \det(\mathbf{A}) \cdot \mathbf{I} \end{cases} \Rightarrow \mathbf{A}^{-1} = \frac{\text{adj}(\mathbf{A})}{\det(\mathbf{A})}$

Singularidade

Diz-se que uma matriz $\mathbf{A}$ de dimensão $n \times n$ é singular se $\det(\mathbf{A}) = 0$. Caso contrário, diz-se que $\mathbf{A}$ é não-singular. Como $\mathbf{A}^{-1} = \text{adj}(\mathbf{A}) / \det(\mathbf{A})$, $\mathbf{A}$ admite inversa se e somente se $\det(\mathbf{A}) \neq 0$, ou seja, se $\mathbf{A}$ é não-singular.

Range e posto

O range de uma matriz $\mathbf{A}$ de dimensão $m \times n$ é um subespaço do $\mathfrak{R}^m$ definido na forma:

$$R(\mathbf{A}) = \{ \mathbf{y} \in \mathfrak{R}^m : \mathbf{y} = \mathbf{A}\mathbf{x}, \text{ para algum } \mathbf{x} \in \mathfrak{R}^n \}.$$

Portanto, $R(\mathbf{A})$ é dado pelo conjunto de todas as possíveis combinações lineares das colunas de $\mathbf{A}$. A dimensão de $R(\mathbf{A})$ é denominada posto da matriz $\mathbf{A}$ e é representado por $\rho(\mathbf{A})$. O posto é o número de colunas linearmente independentes de $\mathbf{A}$.

Como $\rho(\mathbf{A}) = \rho(\mathbf{A}^T)$, o posto pode ser também considerado como sendo o número de linhas linearmente independentes de $\mathbf{A}$. Se $n \geq m$ e $\rho(\mathbf{A}) = m$, então $R(\mathbf{A}) \equiv \mathfrak{R}^m$, donde se conclui que

$$\dim\{R(\mathbf{A})\} = \rho(\mathbf{A}) = \rho(\mathbf{A}^T) \leq \min(m, n).$$

Nulidade

O nulo de uma matriz $\mathbf{A}$ de dimensão $m \times n$ é um subespaço do $\mathfrak{R}^n$ definido na forma:

$$N(\mathbf{A}) \triangleq \{\mathbf{x} \in \mathfrak{R}^n: \mathbf{Ax} = \mathbf{0}\}$$

A dimensão de $N(\mathbf{A})$ é denominada nulidade da matriz $\mathbf{A}$ e é representada por $v(\mathbf{A})$. A nulidade e o posto estão relacionados por:

$$v(\mathbf{A}) + \rho(\mathbf{A}) = n$$

São propriedades importantes: $N(\mathbf{A}) \perp R(\mathbf{A}^T)$ no $\mathfrak{R}^n$ $N(\mathbf{A}^T) \perp R(\mathbf{A})$ no $\mathfrak{R}^m$

Autovalor e autovetor

Seja uma matriz $\mathbf{A}$ de dimensão $n \times n$. Diz-se que um escalar $\lambda \in \mathbb{C}$ (conjunto dos números complexos) é um autovalor de $\mathbf{A}$ se existe um vetor não-nulo $\mathbf{x} \in \mathbb{C}^n$, chamado de autovetor associado a λ , tal que

$$\mathbf{Ax} = \lambda \mathbf{x}.$$

- $\mathbf{Ax} = \lambda \mathbf{x}$ pode ser reescrito como $(\lambda \mathbf{I} - \mathbf{A})\mathbf{x} = \mathbf{0}$;
- $\exists \mathbf{x} \in \mathbb{C}^n, \mathbf{x} \neq \mathbf{0}$ tal que $(\lambda \mathbf{I} - \mathbf{A})\mathbf{x} = \mathbf{0}$ se e somente se $\det(\lambda \mathbf{I} - \mathbf{A}) = 0$;
- $\Delta(\lambda) \triangleq \det(\lambda \mathbf{I} - \mathbf{A})$ é o polinômio característico de $\mathbf{A}$;
- Como o grau de $\Delta(\lambda)$ é n , a matriz $\mathbf{A}$ possui n autovalores.

Positividade

Uma matriz $\mathbf{A}$ de dimensão $n \times n$ é dita ser definida positiva se e somente se

$$\langle \mathbf{x}, \mathbf{Ax} \rangle = \mathbf{x}^T \mathbf{Ax} > 0, \forall \mathbf{x} \in \mathfrak{R}^n, \mathbf{x} \neq \mathbf{0},$$

ou, de forma equivalente, se todos os seus autovalores forem positivos.

Uma matriz $\mathbf{A}$ de dimensão $n \times n$ é dita ser semi-definida positiva se e somente se

$$\langle \mathbf{x}, \mathbf{Ax} \rangle = \mathbf{x}^T \mathbf{Ax} \geq 0, \forall \mathbf{x} \in \mathfrak{R}^n,$$

ou, de forma equivalente, se todos os seus autovalores forem não-negativos.

Uma matriz $\mathbf{A}$ de dimensão $n \times n$ é dita ser definida negativa (semi-definida negativa) se $-\mathbf{A}$ é definida positiva (semi-definida positiva).

A.20 Teoremas envolvendo propriedades de matrizes

Teorema: Seja $\mathbf{A}$ uma matriz de dimensão $m \times n$ e considere o sistema de equações lineares $\mathbf{Ax} = \mathbf{y}$. É possível afirmar que:

- 1) Dados $\mathbf{A} \in \mathfrak{R}^{m \times n}$ e $\mathbf{y} \in \mathfrak{R}^m$, existe $\mathbf{x} \in \mathfrak{R}^n$ tal que $\mathbf{Ax} = \mathbf{y}$ se e somente se $\mathbf{y} \in R(\mathbf{A})$, ou seja, se e somente se $\mathbf{y}$ puder ser expresso como uma combinação linear das colunas de $\mathbf{A}$.

Nota: $\mathbf{y} \in R(\mathbf{A}) \Rightarrow \rho(\mathbf{A}) = \rho([\mathbf{A}; \mathbf{y}])$.

- 2) Dado $\mathbf{A} \in \mathfrak{R}^{m \times n}$, para todo $\mathbf{y} \in \mathfrak{R}^m$ existe $\mathbf{x} \in \mathfrak{R}^n$ tal que $\mathbf{Ax} = \mathbf{y}$ se e somente se $R(\mathbf{A}) \equiv \mathfrak{R}^m$.

Nota: $R(\mathbf{A}) \equiv \mathfrak{R}^m \Rightarrow \rho(\mathbf{A}) = m$

Teorema: Seja $\mathbf{A}$ uma matriz de dimensão $m \times n$ tal que $\rho(\mathbf{A}) = m$. Então a matriz $\mathbf{Q} = \mathbf{AA}^T \in \mathfrak{R}^m$ é definida positiva.

Teorema: Seja $\mathbf{A}$ uma matriz de dimensão $n \times n$. Então os autovetores associados a autovalores distintos são linearmente independentes.

Teorema: Seja $\mathbf{A}$ uma matriz simétrica de dimensão $n \times n$. Então os autovalores de $\mathbf{A}$ são reais e autovetores associados a autovalores distintos são ortogonais.

Teorema: Toda matriz $\mathbf{A}$ de dimensão $m \times n$ que tenha posto unitário pode ser fatorada na forma $\mathbf{A} = \mathbf{uv}^T$, para algum $\mathbf{u} \in \mathfrak{R}^m$ e $\mathbf{v} \in \mathfrak{R}^n$.

Apêndice B

Espaços normados e espaços lineares de funções¹

A definição de espaços vetoriais lineares apresentada no Apêndice A estabelece sete propriedades algébricas de elementos do espaço, enunciadas na forma de axiomas. Com base apenas nestas propriedades algébricas de adição e multiplicação por escalar, os espaços vetoriais lineares não possuem estrutura suficiente para sua utilização em análise funcional.

Realizado por Fréchet, Hausdorff, Riesz, Hilbert e Banach no início deste século, o desenvolvimento de conceitos de análise funcional, particularmente o estudo de espaços topológicos, métricos e normados, forneceu subsídios valiosos para o estabelecimento da teoria de aproximação. Estes espaços permitem medir o erro de aproximação e conduzem a operadores que podem ser utilizados diretamente em processos de aproximação.

Este apêndice apresenta a definição de uma série de espaços, menos pela importância de suas propriedades intrínsecas e mais pelo papel que exercem no texto principal desta tese.

B.1 Espaço métrico

Um espaço métrico X é um espaço linear constituído por uma coleção de elementos e uma medida de distância $d(\mathbf{x}_1, \mathbf{x}_2)$ definida para todo par ordenado de elementos $\mathbf{x}_1, \mathbf{x}_2 \in X$. A função $d(\cdot)$ satisfaz as seguintes propriedades:

- $d(\mathbf{x}_1, \mathbf{x}_2) \geq 0$;
- $d(\mathbf{x}_1, \mathbf{x}_2) = 0$ se e somente se $\mathbf{x}_1 = \mathbf{x}_2$;
- $d(\mathbf{x}_1, \mathbf{x}_2) = d(\mathbf{x}_2, \mathbf{x}_1)$;
- $d(\mathbf{x}_1, \mathbf{x}_2) + d(\mathbf{x}_2, \mathbf{x}_3) \geq d(\mathbf{x}_1, \mathbf{x}_3)$

¹ Este apêndice contém uma reprodução parcial de conceitos apresentados por DAVIS (1975), LUENBERGER (1969) e SCHUMAKER (1981), empregando a mesma notação utilizada por SCHUMAKER (1981).

B.2 Espaço linear normado

Um espaço linear normado é um espaço linear X em que está definido um funcional $\|\cdot\|$ que mapeia cada elemento $\mathbf{x} \in X$ em um escalar real $\|\mathbf{x}\|$, denominado norma de $\mathbf{x}$. Conforme descrito no apêndice A, a norma satisfaz os seguintes axiomas:

- $\|\mathbf{x}\| \geq 0$ para todo $\mathbf{x} \in X$, $\|\mathbf{x}\| = 0$ se e somente se $\mathbf{x} = \mathbf{0}$ (definibilidade positiva);
- $\|a\mathbf{x}\| = |a|\|\mathbf{x}\|$ para todo escalar a e todo $\mathbf{x} \in X$ (homogeneidade);
- $\|\mathbf{x}_1 + \mathbf{x}_2\| \leq \|\mathbf{x}_1\| + \|\mathbf{x}_2\|$ para todo $\mathbf{x}_1, \mathbf{x}_2 \in X$ (desigualdade triangular).

Observe que o conceito de espaço métrico é mais primitivo que o de espaço normado, pois um espaço normado é um espaço métrico com

$$d(\mathbf{x}_1, \mathbf{x}_2) = \|\mathbf{x}_1 - \mathbf{x}_2\|.$$

B.3 Espaço de Banach

Espaços normados em que toda sequência de Cauchy (veja apêndice C) é convergente são importantes em análise funcional. Nestes espaços é possível identificar seqüências convergentes sem identificar explicitamente seus limites.

Um espaço linear normado X é completo se toda sequência de Cauchy em X tem um limite em X . Espaços lineares normados completos são denominados espaços de Banach.

Teorema: Em um espaço de Banach, um subconjunto é completo se e somente se ele é fechado. Além disso, em um espaço linear normado, qualquer subespaço de dimensão finita é completo.

Em problemas de otimização que buscam um conjunto de valores ótimos para um vetor de parâmetros que maximize (minimize) um dado critério, a principal vantagem de espaços de Banach está na garantia de convergência de processos iterativos de busca. Estes processos fornecem, em seqüência, valores para o vetor de parâmetros que superam os valores anteriores no atendimento do critério. O vetor ótimo desejado é, dessa forma, o limite da seqüência.

B.4 Espaço com produto interno

Um espaço linear X é denominado espaço com produto interno, ou espaço pré-Hilbert se, para todo par de elementos $\mathbf{x}_1, \mathbf{x}_2 \in X$, está definido um número real denominado produto interno (veja apêndice A) e representado por $\langle \mathbf{x}_1, \mathbf{x}_2 \rangle$, com as seguintes propriedades:

- $\langle \mathbf{x}_1 + \mathbf{x}_2, \mathbf{x}_3 \rangle = \langle \mathbf{x}_1, \mathbf{x}_3 \rangle + \langle \mathbf{x}_2, \mathbf{x}_3 \rangle$ (linearidade);
- $\langle \mathbf{x}_1, \mathbf{x}_2 \rangle = \langle \mathbf{x}_2, \mathbf{x}_1 \rangle$ (simetria);
- $\langle a\mathbf{x}_1, \mathbf{x}_2 \rangle = a\langle \mathbf{x}_1, \mathbf{x}_2 \rangle, \quad a \in \mathfrak{R}$ (homogeneidade);
- $\langle \mathbf{x}_1, \mathbf{x}_1 \rangle \geq 0, \quad \langle \mathbf{x}_1, \mathbf{x}_1 \rangle = 0$ se e somente se $\mathbf{x}_1 = 0$ (positividade);

B.5 Espaço de Hilbert

Se X é um espaço com produto interno, ou seja, um espaço pré-Hilbert, a equação

$$\|\mathbf{x}\| = \sqrt{\langle \mathbf{x}, \mathbf{x} \rangle}$$

define uma norma em X , e X torna-se um espaço normado.

O espaço de Hilbert é um espaço de Banach (espaço normado e completo) equipado com um produto interno que induz a norma associada ao espaço. Em outras palavras, o espaço de Hilbert é um espaço pré-Hilbert completo.

Ortogonalidade, ortonormalidade, expansão em funções ortonormais (série de Fourier generalizada) e minimização no sentido de quadrados mínimos são conceitos que assumem uma definição natural em espaços de Hilbert.

B.6 Espaço linear de funções

Uma função real pode ser considerada como um elemento de um espaço linear. Seja X um conjunto de pontos em um espaço multidimensional. Seja F o conjunto de todas as funções reais com domínio X . Para todo $f, g \in F$ e $a \in \mathfrak{R}$, definindo-se a soma

$$(f + g)(\mathbf{x}) = f(\mathbf{x}) + g(\mathbf{x}), \quad \mathbf{x} \in X,$$

o produto escalar

$$(af)(\mathbf{x}) = af(\mathbf{x}), \quad \mathbf{x} \in X,$$

e o vetor nulo como sendo a função de F que se anula identicamente, F é um espaço linear. Se X contém mais que um número finito de pontos, então F tem dimensão infinita.

Apesar de serem extensíveis ao caso multidimensional, os espaços lineares de funções apresentados a seguir se restringem a funções reais monovariáveis $f: [a, b] \rightarrow \mathfrak{R}$ definidas em um intervalo finito $[a, b]$, com $a, b \in \mathfrak{R}$.

B.6.1 Espaço de funções mensuráveis

O espaço de funções mensuráveis em $[a,b]$ é definido na forma:

$$M[a,b] = \{f: f \text{ é mensurável em } [a,b]\}.$$

B.6.2 Espaço de Lebesgue

Dado um escalar real $p > 0$, o espaço de Lebesgue é definido como segue:

$$L_p[a,b] = \left\{f: f \in M[a,b] \text{ e } \|f\|_{L_p[a,b]} < \infty\right\},$$

onde $\| \cdot \|_{L_p[a,b]}$ é a norma de Lebesgue definida na forma:

$$\|f\|_{L_p[a,b]} = \left[\int_a^b |f(x)|^p dx \right]^{1/p}, \quad p > 0.$$

Note que $\|f\|_{L_p[a,b]} = 0$ implica $f = 0$ em *quase* todo o intervalo $[a,b]$. Ou seja, é possível que exista um número finito de valores de $x \in [a,b]$ para os quais $f(x) \neq 0$. De forma equivalente, dado $f \in L_p[a,b]$, se $g \in L_p[a,b]$ representa funções tais que

$$\|f - g\|_{L_p[a,b]} = 0, \quad p > 0,$$

então, para $p \rightarrow \infty$, tem-se que

$$\|f\|_{L_\infty[a,b]} = \inf_g \sup_{x \in [a,b]} |g(x)|$$

Quando não houver necessidade de distinção entre funções que são idênticas em quase todo o intervalo $[a,b]$, o espaço de Lebesgue é um espaço linear normado. Neste caso, $L_p[a,b]$ apresenta as seguintes propriedades:

- (1) $L_p[a,b]$ é um espaço de Banach;
- (2) Se $f \in L_1[a,b]$, então $g(x) = \int_a^x f(x)dx$ é contínua;
- (3) Se $f \in L_p[a,b]$ com $p \geq 1$, é possível encontrar uma função uniformemente contínua $g(x)$ tal que $\int_a^b |f(x) - g(x)|^p dx \leq \varepsilon$ para $\varepsilon > 0$ arbitrário;
- (4) Como $[a,b]$ é um intervalo finito, então, para $1 \leq q \leq p$ é válida a seguinte relação:
 $L_q[a,b] \subseteq L_p[a,b]$.

B.6.3 Espaço de funções limitadas

O espaço de funções limitadas em $[a,b]$ é definido na forma:

$$B[a,b] = \{f: \exists m < \infty \text{ tal que } |f(x)| \leq m \text{ para todo } x \in [a,b]\}$$

$B[a,b]$ é um espaço linear normado quando associado à norma $\|f\| = \sup_{x \in [a,b]} |f(x)|$.

Para $f \in M[a,b]$ e com $p > 0$, é válida a seguinte relação: $B[a,b] \subseteq L_p[a,b]$.

B.6.4 Espaço de funções contínuas e uniformemente contínuas

O espaço de funções contínuas em $[a,b]$ é definido na forma:

$$C[a,b] = \{f: \text{dado } \varepsilon > 0, \exists \delta > 0 \text{ tal que, } \forall x_0 \in [a,b], |f(x) - f(x_0)| < \varepsilon \text{ sempre que } |x - x_0| < \delta, x \in [a,b]\},$$

onde δ depende de f , x_0 e ε . Se é possível encontrar um δ que dependa somente de f e ε , então o espaço é dito ser uniformemente contínuo. As noções de continuidade e continuidade uniforme coincidem em funções definidas em intervalos compactos.

Quando associado à norma $\|f\| = \max_{x \in [a,b]} |f(x)|$, o espaço $C[a,b]$ é um espaço de Banach.

Observe que se o intervalo fosse aberto, não seria possível utilizar o conceito de máximo, devendo-se recorrer ao conceito de supremo.

Considerando-se um outro tipo de norma como, por exemplo,

$$\|f\| = \int_a^b |f(x)| dx,$$

tem-se um espaço linear normado diferente de $C[a,b]$. Esta norma se aplica somente se $|f(x)|$ é integrável em $[a,b]$.

É válida a seguinte relação: $C[a,b] \subseteq B[a,b]$.

B.6.5 Espaço de funções absolutamente contínuas

O espaço de funções absolutamente contínuas em $[a,b]$ é definido como segue:

$$AC[a,b] = \left\{ f: \begin{array}{l} \text{para qualquer } \varepsilon > 0, \text{ existe } \delta > 0 \text{ tal que } \sum_{i=1}^n |f(\bar{t}_i) - f(\underline{t}_i)| < \varepsilon \text{ para} \\ \text{todo } n \text{ e todo } a \leq \underline{t}_1 \leq \bar{t}_1 \leq \underline{t}_2 \leq \bar{t}_2 \leq \dots \leq \underline{t}_n \leq \bar{t}_n \leq b \text{ com } \sum_{i=1}^n |\bar{t}_i - \underline{t}_i| < \delta \end{array} \right\}$$

É válida a seguinte relação: $AC[a,b] \subseteq C[a,b]$.

B.6.6 Espaço de Lipschitz

O espaço de funções que satisfazem uma condição de Lipschitz de ordem α em $[a, b]$ é definido na forma:

$$Lip^\alpha[a, b] = \left\{ f: \exists D < \infty \text{ tal que } |f(x_1) - f(x_2)| \leq D|x_1 - x_2|^\alpha \quad \forall x_1, x_2 \in [a, b] \right\}.$$

Quando for necessário explicitar a constante D , é utilizada a notação $Lip_D^\alpha[a, b]$.

São válidas as seguintes relações:

- $Lip^\beta[a, b] \subseteq Lip^\alpha[a, b]$, para $\beta \geq \alpha$;
- $Lip^\alpha[a, b] \subseteq L_\infty[a, b]$;
- $Lip^\alpha[a, b] \subseteq AC[a, b] \subseteq C[a, b]$.

B.6.7 Espaço de funções continuamente diferenciáveis até ordem r

O espaço de funções continuamente diferenciáveis até ordem r ($r > 0$) em $[a, b]$ é um subespaço de $C[a, b]$ definido na forma:

$$C^r[a, b] = \left\{ f: \frac{d^r f(x)}{dx^r} \in C[a, b] \right\},$$

Note que, se f é apenas diferenciável até ordem r ($r > 0$) em $[a, b]$, então $f \in C^{r-1}[a, b]$.

Associado a este espaço é possível definir a semi-norma $\|f^{(r)}\| = \max_{x \in [a, b]} \left| \frac{d^r f(x)}{dx^r} \right|$, a qual

tem como espaço nulo o conjunto de todos os polinômios com grau menor que r .

Uma norma associada a este espaço é dada na forma:

$$\|f\| = \max_{x \in [a, b]} |f(x)| + \max_{x \in [a, b]} \left| \frac{d^r f(x)}{dx^r} \right|.$$

São válidas as seguintes relações: $C^r[a, b] \subseteq Lip^1[a, b]$, $r > 0$;

$$\cdots \subseteq C^2[a, b] \subseteq C^1[a, b] \subseteq AC[a, b] \subseteq C[a, b].$$

B.6.8 Espaço de Sobolev

O espaço de Sobolev é um subespaço do espaço de Lebesgue, formado por funções mais suaves. Seja $L_p[a,b]$ o espaço de Lebesgue. Dado qualquer inteiro positivo $r \geq 1$, o correspondente espaço de Sobolev é definido na forma:

$$L_p^r[a,b] = \left\{ f: \frac{d^{r-1}f(x)}{dx^{r-1}} \in AC[a,b] \text{ e } \frac{d^r f(x)}{dx^r} \in L_p[a,b] \right\}.$$

$L_p^r[a,b]$ é um espaço linear normado quando associado à norma $\|f\|_{L_p^r[a,b]} = \sum_{j=0}^r \left\| \frac{d^j f}{dx^j} \right\|_{L_p[a,b]}$.

Para todo escalar real $1 \leq p \leq \infty$ e qualquer inteiro positivo r , são válidas as seguintes relações:

- $L_p^q[a,b] \subseteq L_p^r[a,b]$, para $q \geq r$;
- $C^r[a,b] \subseteq L_\infty^r[a,b] \subseteq L_p^r[a,b] \subseteq L_1^r[a,b] \subseteq C^{r-1}[a,b]$.

B.6.9 Espaço de funções analíticas

Com $f \in C^\infty[a,b]$ e $x_0 \in [a,b]$ é possível formar a expansão de Taylor de $f(x)$ na forma:

$$f(x) \approx \sum_{k=0}^{\infty} \frac{f^{(k)}(x_0)}{k!} (x - x_0)^k.$$

Para um dado $x \in [a,b]$, esta série pode ou não convergir. Se ela converge, ela pode ou não convergir para $f(x)$.

Se a série converge para $f(x)$ no intervalo $[a,b]$, diz-se que $f(x)$ é uma função analítica em $[a,b]$. O espaço de funções analíticas em $[a,b]$ é, portanto, definido na forma:

$$An[a,b] = \left\{ f: f(x) = \sum_{k=0}^{\infty} \frac{f^{(k)}(x_0)}{k!} (x - x_0)^k \text{ para todo } x, x_0 \in [a,b] \right\}.$$

Funções analíticas podem ser completamente caracterizadas pelo crescimento de suas derivadas, resultando em uma definição alternativa na forma:

$$An[a,b] = \left\{ f: \text{para todo } x \in [a,b], \exists r > 0 \text{ tal que } \left| \frac{d^n f(x)}{dx^n} \right| \leq n! r^n, n = 0, 1, \dots \right\}.$$

É válida a seguinte relação: $An[a,b] \subseteq C^\infty[a,b]$.

B.6.10 Espaço de funções íntegras

O espaço de funções íntegras em $[a, b]$ é definido na forma:

$$Itg[a, b] = \left\{ f: \text{para todo } x \in [a, b], f(x) = \sum_{k=0}^{\infty} c_k x^k \text{ com } \lim_{k \rightarrow \infty} |c_k|^{1/k} = 0 \right\}.$$

É válida a seguinte relação: $Itg[a, b] \subseteq An[a, b]$.

B.6.11 Espaço de polinômios de ordem até M

O espaço de polinômios de ordem até M (M finito) com coeficientes reais é definido no intervalo $[a, b]$ na forma:

$$P_M[a, b] = \left\{ p_M(x) = \sum_{i=0}^M c_i x^i, \quad c_0, \dots, c_M \in \mathfrak{R} \text{ e } x \in [a, b] \right\}.$$

Dado qualquer número real d , uma base para $P_M[a, b]$ pode ser definida na forma:

$$[P_M[a, b]] = \{1, x-d, (x-d)^2, \dots, (x-d)^M\}.$$

A derivada e a integral indefinida de um polinômio são também polinômios. De fato, se

$$p_M(x) = \sum_{i=0}^M c_i (x-d)^i, \text{ então}$$

$$\frac{dp_M(x)}{dx} = \sum_{i=1}^M i c_i (x-d)^{i-1} \quad \text{e} \quad \int_d^x p_M(t) dt = \sum_{i=0}^M \frac{c_i}{i+1} (x-d)^{i+1}.$$

Seja $a \leq x_1 < x_2 < \dots < x_{M+1} \leq b$. $P_M[a, b]$ é um espaço linear normado quando associado à norma $\|p\| = \max_{1 \leq j \leq (M+1)} |p(x_j)|$.

É válida a seguinte relação: $P_M[a, b] \subseteq Itg[a, b]$.

B.6.12 Espaço de funções constantes

O espaço de funções constantes em $[a, b]$ é definido na forma:

$$Cnt[a, b] = \{f: \text{para todo } x \in [a, b], f(x) = c, \text{ com } c \in \mathfrak{R} \text{ uma constante}\}.$$

Se $f \in Cnt[a, b]$, então $f \in Lip^\alpha[a, b]$ para todo $\alpha > 0$.

É válida a seguinte relação: $Cnt[a, b] \subseteq P_M[a, b]$.

Apêndice C

Conceitos básicos de análise funcional e topológica

Em 1930, Dunham Jackson concluiu um livro abordando conceitos de teoria de aproximação, com um capítulo dedicado ao estudo da geometria de espaços funcionais. A partir de então, tem aumentado o interesse no estudo e aplicação da teoria de análise funcional e topológica.

Os resultados apresentados a seguir tratam preferencialmente funções reais multivariáveis, definidas em regiões compactas (fechadas e finitas).

C.1 Conjunto aberto e conjunto fechado

Definição: Seja P um subconjunto de um espaço vetorial linear. Um ponto $\mathbf{p} \in P$ é um ponto interior de P se existe $\varepsilon > 0$ tal que todos os vetores $\mathbf{x}$ que satisfazem $\|\mathbf{x} - \mathbf{p}\| < \varepsilon$ são também elementos de P . A coleção de todos os pontos interiores de P é chamada de interior de P e é representada por $\hat{P}$. Introduzindo a notação de esfera aberta centrada em $\mathbf{p}$ na forma $S(\mathbf{x}, \varepsilon) = \{\mathbf{y}: \|\mathbf{y} - \mathbf{p}\| < \varepsilon\}$, então $\mathbf{p}$ é um ponto interior de P se para algum $\varepsilon > 0$ existe uma esfera $S(\mathbf{x}, \varepsilon)$ contida em P .

Definição: Diz-se que um conjunto P é aberto se $P = \hat{P}$. O conjunto vazio é um conjunto aberto, pois seu interior é um conjunto vazio.

Definição: Seja P um subconjunto de um espaço vetorial linear X . Um ponto $\mathbf{x} \in X$ é dito ser um ponto de fecho do conjunto P se, dado $\varepsilon > 0$, existe um ponto $\mathbf{p} \in P$ satisfazendo $\|\mathbf{x} - \mathbf{p}\| < \varepsilon$. O conjunto de todos os pontos de fecho de P é chamado de fecho de P e é representado por $\bar{P}$. Um ponto $\mathbf{x}$ é um ponto de fecho de P se toda esfera centrada em $\mathbf{x}$ contém um ponto de P .

Definição: Diz-se que um conjunto P é fechado se $P = \overline{P}$. Em geral, se P não for fechado, $P \subset \overline{P}$.

Definição: Seja P um subconjunto de um espaço vetorial linear X . Um ponto $\mathbf{x} \in X$ é dito ser um ponto de fronteira do conjunto P se, dado $\varepsilon > 0$, a esfera $S(\mathbf{x}, \varepsilon)$ contém ao menos um elemento de P e um elemento fora de P . O conjunto de todos os pontos de fronteira de P é chamado de fronteira de P e é representado por ∂P .

Algumas relações envolvendo $\hat{P}$, $\overline{P}$ e ∂P :

$$\overline{P} = P \cup \partial P \quad \hat{P} = \overline{P} - \partial P \quad P \neq \hat{P} \cup \partial P$$

C.2 Distância entre um ponto e um espaço vetorial linear normado

Considere um espaço vetorial linear normado $X \subset \mathfrak{R}^n$ e um ponto $\mathbf{x} \in \mathfrak{R}^n$. A distância entre o ponto $\mathbf{x}$ e o espaço X é dada por

$$\text{dist}(\mathbf{x}, X) = \inf_{\mathbf{y} \in X} \|\mathbf{x} - \mathbf{y}\|,$$

onde $\|\cdot\|$ é a norma definida no espaço X . Observe que, se $\mathbf{x} \in X$, então $\text{dist}(\mathbf{x}, X) = 0$.

C.3 Distância entre dois espaços vetoriais lineares normados

Considere espaços vetoriais lineares normados $X, Y \subset \mathfrak{R}^n$. A distância entre os espaços X e Y é dada por

$$\text{dist}(X, Y) = \inf_{\mathbf{x} \in X} \text{dist}(\mathbf{x}, Y) = \inf_{\substack{\mathbf{x} \in X \\ \mathbf{y} \in Y}} \|\mathbf{x} - \mathbf{y}\|,$$

onde $\|\cdot\|$ é a norma definida nos dois espaços. Observe que, se os dois espaços se interceptam, então $\text{dist}(X, Y) = 0$.

C.4 Seqüência de números reais

Uma seqüência $a_1, a_2, \dots, a_n, \dots$ de números reais é representada como $\{a_i\}_{i=1}^{\infty}$ ou simplesmente $\{a_i\}$.

C.5 Seqüência de vetores

Uma seqüência $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n, \dots$ de vetores em X é representada como $\{\mathbf{x}_i\}_{i=1}^\infty$ ou simplesmente $\{\mathbf{x}_i\}$.

Um vetor $\mathbf{x}$ é um limite da seqüência $\{\mathbf{x}_i\}$ de vetores em X se

$$\forall \varepsilon > 0, \exists j \text{ finito tal que } \|\mathbf{x}_i - \mathbf{x}\| < \varepsilon, \forall i \geq j$$

Uma seqüência $\{\mathbf{x}_i\}$ de vetores em X converge para um vetor $\mathbf{x}$ ($\mathbf{x}_i \rightarrow \mathbf{x}$) se a seqüência $\{\|\mathbf{x}_i - \mathbf{x}\|\}$ de números reais converge para 0. Neste caso, o limite $\mathbf{x}$ da seqüência $\{\mathbf{x}_i\}$ é único.

Além disso, $\mathbf{x}_i \rightarrow \mathbf{x}$ implica $\|\mathbf{x}_i\| \rightarrow \|\mathbf{x}\|$. Note que a condição $\mathbf{x} \in X$ pode não ser atendida.

O símbolo $\|\cdot\|$ denota uma norma (veja Apêndice A).

C.6 Seqüência de Cauchy

Uma seqüência $\{\mathbf{x}_i\}$ de vetores em X é uma seqüência de Cauchy se $\{\|\mathbf{x}_i - \mathbf{x}_j\|\} \rightarrow 0$ quando $i, j \rightarrow \infty$, ou seja,

$$\forall \varepsilon > 0, \exists k \text{ finito tal que } \|\mathbf{x}_i - \mathbf{x}_j\| < \varepsilon, \forall i, j \geq k$$

- Toda seqüência convergente é uma seqüência de Cauchy.
- Um espaço X em que toda seqüência de Cauchy $\{\mathbf{x}_i\}$ é convergente, isto é, tem um limite $\mathbf{x} \in X$, é dito ser completo.
- Espaços de dimensão finita são completos.

C.7 Transformação

Sejam $(X, \mathfrak{S})$ e $(Y, \mathfrak{S})$ espaços vetoriais lineares e S um subconjunto de X . A regra que associa cada elemento $\mathbf{x} \in S$ a um elemento $\mathbf{y} \in Y$ é chamada de transformação e é representada como $\mathbf{T}: S \subset X \rightarrow Y$, de forma que $\mathbf{y} = \mathbf{T}(\mathbf{x})$.

Uma transformação $\mathbf{T}: S \subset X \rightarrow Y$ é linear se

$$\mathbf{T}(a_1\mathbf{x}_1 + a_2\mathbf{x}_2) = a_1\mathbf{T}(\mathbf{x}_1) + a_2\mathbf{T}(\mathbf{x}_2)$$

para quaisquer $a_1, a_2 \in \mathfrak{S}$ e $\mathbf{x}_1, \mathbf{x}_2 \in S$. Neste caso, se $\dim(X) = n$ e $\dim(Y) = m$, então a transformação linear $\mathbf{T}$ pode ser interpretada como uma matriz de dimensão $m \times n$ que mapeia elementos do $\mathfrak{R}^n$ no $\mathfrak{R}^m$, ou seja, $\mathbf{y} = \mathbf{T}\mathbf{x}$.

Uma transformação $\mathbf{T}: S \subset X \rightarrow Y$ é contínua em $\mathbf{x}_0 \in S$ se

$$\forall \varepsilon > 0, \exists \eta > 0 \text{ tal que } \|\mathbf{x} - \mathbf{x}_0\| < \eta \text{ implica } \|\mathbf{T}(\mathbf{x}) - \mathbf{T}(\mathbf{x}_0)\| < \varepsilon.$$

Se $\mathbf{T}$ é contínua sobre S então diz-se simplesmente que $\mathbf{T}$ é contínua.

C.8 Funcional

Uma transformação de um espaço vetorial $(X, \mathfrak{S})$ para o espaço formado por escalares reais (ou complexos) é chamada de funcional.

Exemplo: $f(\mathbf{x}) = \|\mathbf{x}\|$, $\mathbf{x} \in X \subset \mathfrak{R}^n$.

Um funcional $f: X \subset \mathfrak{R}^n \rightarrow \mathfrak{R}$ é linear se

$$f(a_1\mathbf{x}_1 + a_2\mathbf{x}_2) = a_1f(\mathbf{x}_1) + a_2f(\mathbf{x}_2)$$

para quaisquer $a_1, a_2 \in \mathfrak{S}$ e $\mathbf{x}_1, \mathbf{x}_2 \in X$.

Um funcional $f: X \subset \mathfrak{R}^n \rightarrow \mathfrak{R}$ é contínuo em $\mathbf{x}_0 \in X$ se

$$\forall \varepsilon > 0, \exists \eta > 0 \text{ tal que } \|\mathbf{x} - \mathbf{x}_0\| < \eta \text{ implica } \|f(\mathbf{x}) - f(\mathbf{x}_0)\| < \varepsilon.$$

Se f é contínuo sobre X então diz-se simplesmente que f é contínuo.

No $(\mathfrak{R}^n, \mathfrak{R})$, todo funcional linear pode ser expresso na forma:

$$f(\mathbf{x}) = \langle \mathbf{c}, \mathbf{x} \rangle = \sum_{i=1}^n c_i x_i,$$

onde $c_i \in \mathfrak{R}$ ($i=1, \dots, n$).

C.9 Hiperplano

Considere um vetor constante não-nulo $\mathbf{a}$ de um espaço vetorial linear X . A equação $\langle \mathbf{a}, \mathbf{x} \rangle = 0$, $\mathbf{x} \in X$, define um subespaço H_1 de X dado por:

$$H_1 = \{\mathbf{x}; \langle \mathbf{a}, \mathbf{x} \rangle = 0, \mathbf{x} \in X\}.$$

Para qualquer escalar b , $\langle \mathbf{a}, \mathbf{x} \rangle = b$ define uma variedade linear. Dado qualquer vetor constante $\mathbf{c} \in X$ tal que $\langle \mathbf{a}, \mathbf{c} \rangle = b$, a variedade linear definida por

$$H_2 = \{\mathbf{x}; \langle \mathbf{a}, \mathbf{x} \rangle = b, \mathbf{x} \in X\}$$

pode ser considerada como uma translação de H_1 por $\mathbf{c}$. H_1 e H_2 são denominados hiperplanos.

Hiperplano é a maior variedade linear própria contida em X , ou seja, se V é uma variedade linear contendo um hiperplano H em um espaço vetorial linear X , então, ou $V = X$, ou $V = H$.

Dado que $f(\mathbf{x}) = \langle \mathbf{a}, \mathbf{x} \rangle$ ($\mathbf{a}$ constante) é um funcional linear, o teorema a seguir relaciona funcionais lineares com hiperplanos.

Teorema: H é um hiperplano de um espaço vetorial linear X se e somente se existe um funcional linear $f: X \rightarrow \mathfrak{R}$ e um escalar constante $b \in \mathfrak{R}$ tal que H pode ser expresso na forma:

$$H = \{\mathbf{x} \in X: f(\mathbf{x}) = b\}.$$

C.10 Diferencial

Seja $\mathbf{x}_0 \in X \subset \mathfrak{R}^n$, $\mathbf{d} \in X$ um vetor arbitrário e $f: X \subset \mathfrak{R}^n \rightarrow \mathfrak{R}$. Se o limite

$$Df(\mathbf{x}_0, \mathbf{d}) = \lim_{a \rightarrow 0} \frac{1}{a} [f(\mathbf{x}_0 + a\mathbf{d}) - f(\mathbf{x}_0)]$$

existe, então $Df(\mathbf{x}_0, \mathbf{d})$ é chamado de diferencial de f em $\mathbf{x}_0$ com incremento $\mathbf{d}$. Se o limite existe para todo $\mathbf{d} \in X$, diz-se que o funcional f é diferenciável em $\mathbf{x}_0$. Note que $Df(\mathbf{x}_0, \mathbf{d})$ é uma generalização do conceito de derivada direcional.

C.11 Derivada parcial e derivada total

Sejam $\mathbf{e}_i$ ($i=1, \dots, n$) os vetores n -dimensionais normais definidos na forma:

$$\mathbf{e}_i = [0 \ \dots \ 0 \ \underset{\substack{\uparrow \\ i\text{-ésima posição}}}{1} \ 0 \ \dots \ 0]^T$$

então, se o funcional $f: X \subset \mathfrak{R}^n \rightarrow \mathfrak{R}$ for diferenciável em $\mathbf{x}_0$,

$$Df(\mathbf{x}_0, \mathbf{e}_i) = \left. \frac{\partial f}{\partial x_i} \right|_{\mathbf{x}=\mathbf{x}_0} = \frac{\partial f(\mathbf{x}_0)}{\partial x_i} \quad (i=1, \dots, n)$$

é denominada derivada parcial de f em relação a x_i no ponto $\mathbf{x}_0$.

Se $n=1$ ($\mathbf{x} = x$), a derivada parcial no ponto $\mathbf{x}_0$ é também a derivada total dada na forma:

$$\frac{df(x_0)}{dx} = f'(x_0) = \lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0}.$$

Sejam os funcionais $f, g: X \subset \Re^n \rightarrow \Re$ diferenciáveis até ordem p no ponto $\mathbf{x} \in X$. Então a q -ésima derivada parcial ($q \leq p$) de $(f \cdot g)(\mathbf{x})$ em relação a x_i ($i=1, \dots, n$) é dada pela fórmula de Leibniz

$$\frac{\partial^q (f \cdot g)(\mathbf{x})}{\partial x_i^q} = \sum_{r \leq q} \frac{q!}{r!(q-r)!} \cdot \frac{\partial^r f(\mathbf{x})}{\partial x_i^r} \cdot \frac{\partial^{q-r} g(\mathbf{x})}{\partial x_i^{q-r}}.$$

C.12 Convergência ponto-a-ponto

Suponha que $\{f_n\}$, $n = 1, 2, \dots$, é uma seqüência de funcionais definidos em X , e suponha que a seqüência de números $\{f_n(\mathbf{x})\}$ converge para todo $\mathbf{x} \in X$. Então, é possível definir um funcional f na forma:

$$f(\mathbf{x}) = \lim_{n \rightarrow \infty} f_n(\mathbf{x}), \quad \mathbf{x} \in X. \quad (\text{C.1})$$

Neste caso, diz-se que $\{f_n\}$ converge ponto-a-ponto para f em X . De forma análoga, se $\sum f_n(\mathbf{x})$ converge para todo $\mathbf{x} \in X$, então é possível definir o funcional

$$f(\mathbf{x}) = \sum_{n=1}^{\infty} f_n(\mathbf{x}), \quad \mathbf{x} \in X, \quad (\text{C.2})$$

denominado a soma da série $\sum f_n(\mathbf{x})$.

A convergência ponto-a-ponto, no entanto, não é suficientemente rigorosa de forma a servir como ponto de partida para a obtenção de resultados matematicamente expressivos. Propriedades importantes como continuidade, diferenciabilidade e integrabilidade dos funcionais f_n , $n=1, 2, \dots$, não são necessariamente preservadas em f . Por exemplo, se os funcionais f_n são contínuos em X , ou seja,

$$\lim_{\mathbf{t} \rightarrow \mathbf{x}} f_n(\mathbf{t}) = f_n(\mathbf{x}), \quad \forall \mathbf{t}, \mathbf{x} \in X, \quad n = 1, 2, \dots,$$

o mesmo não se pode afirmar a respeito de f dado pelas operações de limite (C.1) e (C.2) (RUDIN, 1976).

C.13 Convergência uniforme

Uma seqüência de funcionais $\{f_n\}$, $n = 1, 2, \dots$, converge uniformemente em X para um funcional f se para cada $\varepsilon > 0$ existe um inteiro positivo N tal que $n \geq N$ implica

$$|f_n(\mathbf{x}) - f(\mathbf{x})| \leq \varepsilon, \quad \forall \mathbf{x} \in X. \tag{C.3}$$

Convergência uniforme é mais rigorosa que convergência ponto-a-ponto, pois toda seqüência que converge uniformemente é convergente ponto-a-ponto. De fato:

- se $\{f_n\}$ converge ponto-a-ponto em X , então existe um funcional f tal que, para $\varepsilon > 0$ qualquer e para todo $\mathbf{x} \in X$, existe um inteiro positivo, dependente de ε e $\mathbf{x}$, tal que a desigualdade (C.3) é válida se $n \geq N$.
- se $\{f_n\}$ converge uniformemente em X , então existe um funcional f tal que, para $\varepsilon > 0$ qualquer e para todo $\mathbf{x} \in X$, existe um inteiro positivo, dependente apenas de ε , tal que a desigualdade (C.3) é válida se $n \geq N$.

De forma análoga, a série $\sum f_n(\mathbf{x})$ converge uniformemente em X se a seqüência $\{s_n\}$ de somas parciais definidas por

$$\sum_{i=1}^n f_i(\mathbf{x}) = s_n(\mathbf{x})$$

converge uniformemente em X .

C.14 Critério de Cauchy para convergência uniforme

Teorema: A seqüência de funcionais $\{f_n\}$, definida em X , converge uniformemente em X se e somente se, para todo $\varepsilon > 0$ e $\mathbf{x} \in X$, existe um inteiro positivo N tal que

$$|f_n(\mathbf{x}) - f_m(\mathbf{x})| \leq \varepsilon,$$

onde $m \geq N$, $n \geq N$, com $f_n(\mathbf{x})$ e $f_m(\mathbf{x})$ assumindo valores finitos.

Prova: Veja RUDIN (1976).

C.15 Convergência uniforme e continuidade

Teorema: Suponha que a seqüência de funcionais $\{f_n\}$ converge uniformemente para f em X . Seja $\mathbf{x}$ um ponto limite de X e suponha que

$$\lim_{t \rightarrow x} f_n(t) = A_n, \quad n = 1, 2, \dots$$

Então a sequência $\{A_n\}$ converge, e

$$\lim_{t \rightarrow x} f(t) = \lim_{n \rightarrow \infty} A_n$$

Sendo assim, a conclusão é que

$$\lim_{t \rightarrow x} \lim_{n \rightarrow \infty} f_n(t) = \lim_{n \rightarrow \infty} \lim_{t \rightarrow x} f_n(t)$$

Prova: Veja RUDIN (1976).

Como um resultado imediato tem-se:

Corolário: Se $\{f_n\}$ é uma sequência de funcionais contínuos em X e $\{f_n\}$ converge uniformemente para f em X , então f é contínuo em X . O inverso não é necessariamente verdadeiro, ou seja, uma sequência de funcionais contínuos pode convergir para um funcional contínuo sem que a convergência seja uniforme.

C.16 Convergência uniforme e diferenciação

A convergência uniforme da sequência $\{f_n\}$ não permite extrair qualquer conclusão a respeito da sequência $\{f'_n\}$. Para poder afirmar que $\{f'_n\}$ converge uniformemente para f' sempre que $\{f_n\}$ converge uniformemente para f , algumas hipóteses adicionais são necessárias.

Teorema: Seja $\{f_n\}$ uma sequência de funcionais diferenciáveis em X ($f_n \in C[X]$) e tal que, para algum ponto $x_0 \in X$, $\{f_n(x_0)\}$ converge. Se $\{f'_n\}$ converge uniformemente em X , então $\{f_n\}$ converge uniformemente para um funcional f em X , e

$$f'(x) = \lim_{n \rightarrow \infty} f'_n(x), \quad x \in X.$$

Prova: Veja RUDIN (1976).

C.17 Funcionais contínuos e uniformemente contínuos

Definição: Uma família F de funcionais f definidos em um conjunto X de um espaço métrico X_M é dita ser contínua em $x_0 \in X$ se dado $\varepsilon > 0$ existe um $\delta > 0$ tal que

$$|f(x) - f(x_0)| \leq \varepsilon, \quad \forall f \in F, \quad \forall x \in X \text{ com } \|x - x_0\| < \delta,$$

onde $\|\cdot\|$ representa uma métrica de X_M . Se $f \in F$ é contínua para todo $\mathbf{x}_0 \in X$, então diz-se que $f \in C[X_M]$.

Observe que δ depende de F , $\mathbf{x}_0$ e ε . Para δ independente de $\mathbf{x}_0$, tem-se:

Definição: Uma família F de funcionais f definidos em um conjunto X de um espaço métrico X_M é dita ser uniformemente contínua em X se dado $\varepsilon > 0$ existe um $\delta > 0$ tal que

$$|f(\mathbf{x}) - f(\mathbf{y})| \leq \varepsilon, \forall f \in F, \forall \mathbf{x}, \mathbf{y} \in X \text{ com } \|\mathbf{x} - \mathbf{y}\| < \delta,$$

onde $\|\cdot\|$ representa uma métrica de X_M .

Exemplo: O funcional $f(x) = 1/(1 + x^2)$ é uniformemente contínuo em todo o eixo real pois

$$|f(x) - f(y)| = \left| \frac{1}{1 + x^2} - \frac{1}{1 + y^2} \right| = \frac{|y^2 - x^2|}{(1 + x^2)(1 + y^2)} \leq |y - x| \frac{|x| + |y|}{(1 + x^2)(1 + y^2)} \leq |y - x|.$$

Assim,

$$|f(x) - f(y)| \leq \varepsilon \text{ sempre que } |y - x| \leq \varepsilon = \delta.$$

Teorema: Se uma família F de funcionais f é contínua em um conjunto compacto X de um espaço métrico X_M , então F também é uniformemente contínua em X .

C.18 Módulo de continuidade

Determinar o grau de suavidade de uma função por diferenciabilidade resulta muitas vezes em critérios imprecisos para efeito de aproximação. Por exemplo, no caso de funcionais uniformemente contínuos, é muitas vezes necessário saber qual o melhor ε que atende às condições da definição para um dado δ . Uma alternativa para estes casos é o módulo de continuidade de f em $[a, b]$, definido na forma:

Definição: Dada uma função $f: [a, b] \rightarrow \Re$, para quaisquer $x_1, x_2 \in [a, b]$ tal que $|x_1 - x_2| \leq \delta$, o módulo de continuidade de f em $[a, b]$ é definido na forma:

$$\mu_{[a,b]}^f(\delta) = \sup_{x_1, x_2 \in [a,b]} |f(x_1) - f(x_2)|.$$

Note que, para $f \in Lip_D^\alpha[a, b]$, tem-se $\mu_{[a,b]}^f(\delta) \leq D\delta^\alpha$. Para $f \in C[a, b]$ tem-se:

Teorema: Para $f \in C[a,b]$, o módulo de continuidade de f em $[a,b]$ tem as seguintes propriedades:

- $\mu_{[a,b]}^f(0) = 0$;
- Se $0 < \delta_1 < \delta_2$ então $\mu_{[a,b]}^f(\delta_1) \leq \mu_{[a,b]}^f(\delta_2)$;
- $\mu_{[a,b]}^f(\delta_1 + \delta_2) \leq \mu_{[a,b]}^f(\delta_1) + \mu_{[a,b]}^f(\delta_2)$;
- $\mu_{[a,b]}^f(n\delta) \leq n\mu_{[a,b]}^f(\delta)$;
- $\mu_{[a,b]}^f(\delta) \in C[0, b-a]$.

Prova: Veja DAVIS (1975).

C.19 Série de Taylor

Considere $f: X \subset \mathbb{R}^n \rightarrow \mathbb{R}$ e $\mathbf{x}_0 \in X$. Se f é continuamente diferenciável até 2ª ordem ($f \in C^2[X]$), então a expansão em série de Taylor até 2ª ordem de $f(\mathbf{x})$ em torno de $\mathbf{x}_0$, para $\mathbf{x} \in X$ um vetor arbitrário, é dada na forma:

$$f(\mathbf{x}) = f(\mathbf{x}_0) + \sum_{i=1}^n [Df(\mathbf{x}_0, \mathbf{y}^i)(x_i - x_{0i})] + \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \left\{ (x_i - x_{0i}) Df[Df(\mathbf{x}_0, \mathbf{y}^i), \mathbf{y}^j] (x_j - x_{0j}) \right\} + O(2)$$

$$f(\mathbf{x}) = f(\mathbf{x}_0) + \sum_{i=1}^n \left[\frac{\partial f(\mathbf{x}_0)}{\partial x_i} (x_i - x_{0i}) \right] + \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \left[(x_i - x_{0i}) \frac{\partial^2 f(\mathbf{x}_0)}{\partial x_i \partial x_j} (x_j - x_{0j}) \right] + O(2)$$

$$f(\mathbf{x}) = f(\mathbf{x}_0) + \nabla f(\mathbf{x}_0)^T (\mathbf{x} - \mathbf{x}_0) + \frac{1}{2} (\mathbf{x} - \mathbf{x}_0)^T \mathbf{F}(\mathbf{x}_0) (\mathbf{x} - \mathbf{x}_0) + O(2)$$

onde $\nabla f(\mathbf{x}_0)$ é o vetor gradiente de f no ponto $\mathbf{x} = \mathbf{x}_0$, $\mathbf{F}(\mathbf{x}_0)$ é a matriz hessiana de f no ponto $\mathbf{x} = \mathbf{x}_0$ e $O(2)$ representa termos de ordem superior a 2. Sendo assim,

$$\nabla f(\mathbf{x}_0) = \begin{bmatrix} \frac{\partial f(\mathbf{x}_0)}{\partial x_1} \\ \frac{\partial f(\mathbf{x}_0)}{\partial x_2} \\ \vdots \\ \frac{\partial f(\mathbf{x}_0)}{\partial x_n} \end{bmatrix} \quad \text{e} \quad \mathbf{F}(\mathbf{x}_0) = \begin{bmatrix} \frac{\partial^2 f(\mathbf{x}_0)}{\partial x_1 \partial x_1} & \frac{\partial^2 f(\mathbf{x}_0)}{\partial x_1 \partial x_2} & \dots & \frac{\partial^2 f(\mathbf{x}_0)}{\partial x_1 \partial x_n} \\ \frac{\partial^2 f(\mathbf{x}_0)}{\partial x_2 \partial x_1} & \frac{\partial^2 f(\mathbf{x}_0)}{\partial x_2 \partial x_2} & \dots & \frac{\partial^2 f(\mathbf{x}_0)}{\partial x_2 \partial x_n} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2 f(\mathbf{x}_0)}{\partial x_n \partial x_1} & \frac{\partial^2 f(\mathbf{x}_0)}{\partial x_n \partial x_2} & \dots & \frac{\partial^2 f(\mathbf{x}_0)}{\partial x_n \partial x_n} \end{bmatrix}.$$

C.20 Teorema diferencial do valor médio

Teorema: Seja $f: X \subset \mathbb{R}^n \rightarrow \mathbb{R}$ um funcional continuamente diferenciável em X ($f \in C^1[X]$).

Então existe $a \in [0, 1]$ tal que

$$f(\mathbf{x}) = f(\mathbf{x}_0) + \nabla f[a\mathbf{x}_0 + (1-a)\mathbf{x}]^T (\mathbf{x} - \mathbf{x}_0), \quad \forall \mathbf{x}, \mathbf{x}_0 \in X.$$

Se f é continuamente diferenciável até 2ª ordem ($f \in C^2[X]$),

$$f(\mathbf{x}) = f(\mathbf{x}_0) + \nabla f(\mathbf{x}_0)^T (\mathbf{x} - \mathbf{x}_0) + \frac{1}{2} (\mathbf{x} - \mathbf{x}_0)^T \mathbf{F}[a\mathbf{x}_0 + (1-a)\mathbf{x}] (\mathbf{x} - \mathbf{x}_0), \quad \forall \mathbf{x}, \mathbf{x}_0 \in X.$$

C.21 Funcional convexo e quase-convexo

Seja X um conjunto convexo não-vazio. Um funcional $f: X \subset \mathbb{R}^n \rightarrow \mathbb{R}$ é convexo se

$$f(a\mathbf{x}_0 + (1-a)\mathbf{x}) \leq af(\mathbf{x}_0) + (1-a)f(\mathbf{x}), \quad \forall \mathbf{x}, \mathbf{x}_0 \in X, \quad \forall 0 \leq a \leq 1.$$

Se as desigualdades forem estritas, diz-se que f é estritamente convexo.

Um funcional $f: X \subset \mathbb{R}^n \rightarrow \mathbb{R}$ definido sobre um subconjunto convexo é (estritamente) côncavo se $-f$ for (estritamente) convexo.

Se $f: X \subset \mathbb{R}^n \rightarrow \mathbb{R}$ for diferenciável em X ($f \in C^1[X]$), então f é convexo se e somente se

$$f(\mathbf{x}) \geq f(\mathbf{x}_0) + \nabla f(\mathbf{x}_0)^T (\mathbf{x} - \mathbf{x}_0), \quad \forall \mathbf{x}, \mathbf{x}_0 \in X.$$

Se $f: X \subset \mathbb{R}^n \rightarrow \mathbb{R}$ for duas vezes diferenciável em X ($f \in C^2[X]$), então f é convexo se e somente se a matriz hessiana $\mathbf{F}$ da função f é semi-definida positiva sobre X . Se $\mathbf{F}$ é definida positiva sobre X , então f é estritamente convexa.

Um funcional $f: X \subset \mathbb{R}^n \rightarrow \mathbb{R}$ é quase-convexo se

$$f(a\mathbf{x}_0 + (1-a)\mathbf{x}) \leq \max\{f(\mathbf{x}_0), f(\mathbf{x})\}, \quad \forall \mathbf{x}, \mathbf{x}_0 \in X, \quad \forall 0 \leq a \leq 1.$$

Apêndice D

Métodos de otimização paramétrica não-linear irrestrita

O problema de aproximação de uma função $g(\cdot): X \subset \mathbb{R}^m \rightarrow \mathbb{R}^r$ por uma função paramétrica $\hat{g}(\cdot, \theta): X \times \mathbb{R}^P \rightarrow \mathbb{R}^r$, onde $\theta \in \mathbb{R}^P$ (P finito) é um vetor de parâmetros, pode ser dividido em dois sub-problemas básicos:

- (1) Representação paramétrica: que classe de funções $g(\cdot)$ podem ser aproximadas por que classes de funções de aproximação $\hat{g}(\cdot, \theta)$?
- (2) Otimização paramétrica: a partir do conjunto de dados de aproximação, e fixada a função de aproximação $\hat{g}(\cdot, \theta)$, como encontrar um valor ótimo para o vetor de parâmetros $\theta \in \mathbb{R}^P$?

O problema geral de aproximação paramétrica pode, então, ser formalmente apresentado como segue:

Problema geral de aproximação paramétrica: Considere a função $g: X \subset \mathbb{R}^m \rightarrow \mathbb{R}^r$, que mapeia pontos de um subespaço compacto $X \subset \mathbb{R}^m$ em pontos de um subespaço compacto $g[X] \subset \mathbb{R}^r$. Com base nos pares de vetores de entrada-saída $\{(\mathbf{x}_l, \mathbf{s}_l)\}_{l=1}^N$ amostrados a partir do mapeamento determinístico definido pela função g na forma:

$$\mathbf{s}_l = g(\mathbf{x}_l) + \varepsilon_l, \quad l = 1, \dots, N,$$

e dada a função de aproximação $\hat{g}: X \times \mathbb{R}^P \rightarrow \mathbb{R}^r$, determine o vetor de parâmetros $\theta^* \in \mathbb{R}^P$ tal que

$$\text{dist}(g(\cdot), \hat{g}(\cdot, \theta^*)) \leq \text{dist}(g(\cdot), \hat{g}(\cdot, \theta)), \quad \text{para todo } \theta \in \mathbb{R}^P,$$

onde o operador $\text{dist}(\cdot, \cdot)$ mede a distância entre as duas funções em todo o espaço X . O vetor ε_l expressa o erro no processo de amostragem, sendo assumido ser de média zero e variância fixa. A solução deste problema, se existir, é denominada a melhor aproximação e depende diretamente da classe de funções para a qual $\hat{g}$ pertence, conforme discutido no capítulo 3.

D.1. Avaliação do nível de aproximação paramétrica

Em aproximação paramétrica utilizando um número finito de dados amostrados, a distância entre a função a ser aproximada e sua aproximação é uma função apenas do vetor de parâmetros $\theta \in \mathfrak{R}^P$. Tomando a norma euclidiana como a medida de distância, produz-se a seguinte expressão:

$$J(\theta) = \frac{1}{N} \sum_{l=1}^N (s_l - \hat{g}(\mathbf{x}_l, \theta))^2. \quad (\text{D.1})$$

O funcional $J: \mathfrak{R}^P \rightarrow \mathfrak{R}$ é denominado superfície de erro do problema de aproximação, pois pode ser interpretado como uma hipersuperfície localizada acima do espaço de parâmetros $\mathfrak{R}^P$, sendo que para cada ponto $\theta \in \mathfrak{R}^P$ corresponde uma altura $J(\theta)$.

Portanto, dada a superfície de erro, solução do problema de representação paramétrica, a solução do problema de otimização paramétrica é o vetor $\theta^* \in \mathfrak{R}^P$ que minimiza $J(\theta)$, ou seja,

$$\theta^* = \arg \min_{\theta \in \mathfrak{R}^P} J(\theta). \quad (\text{D.2})$$

São verdadeiras as seguintes afirmações:

- dado o conjunto de aproximação $\{(\mathbf{x}_l, s_l)\}_{l=1}^N$, o vetor de parâmetros $\theta = \theta^*$ fornece a melhor função de aproximação possível com base na representação paramétrica $\hat{g}(\cdot, \theta)$ e na medida de distância dada pela equação (D.1);
- se a superfície de erro for contínua e diferenciável em relação ao vetor de parâmetros, então os mais eficientes métodos de otimização irrestrita (os parâmetros podem assumir qualquer valor real) podem ser aplicados para minimizar $J(\theta)$;

D.2. Modelos de aproximação paramétrica

Dentre os modelos de aproximação paramétrica mais comuns é possível citar:

- linear (ANDERSON, 1984):

$$\hat{g}(\mathbf{x}, \theta) = \mathbf{x}^T \theta$$

- parcialmente linear (HECKMAN, 1986):

$$\hat{g}(\mathbf{x}, \theta) = \mathbf{x}_1^T \theta_1 + f(\mathbf{x}_2, \theta_2), \quad \mathbf{x} = \begin{bmatrix} \mathbf{x}_1 \\ \dots \\ \mathbf{x}_2 \end{bmatrix}, \quad \theta = \begin{bmatrix} \theta_1 \\ \dots \\ \theta_2 \end{bmatrix}$$

- não-linear por partes (ROBINSON, 1988):

$$\hat{g}(\mathbf{x}, \theta) = f_1(\mathbf{x}_1, \theta_1) + f_2(\mathbf{x}_2, \theta_2), \quad \mathbf{x} = \begin{bmatrix} \mathbf{x}_1 \\ \cdots \\ \mathbf{x}_2 \end{bmatrix}, \quad \theta = \begin{bmatrix} \theta_1 \\ \cdots \\ \theta_2 \end{bmatrix}$$

- aditivo (HASTIE & TIBSHIRANI, 1986; STONE, 1985):

$$\hat{g}(\mathbf{x}, \theta) = f_1(x_1, \theta_1) + \cdots + f_m(x_m, \theta_m), \quad \mathbf{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_m \end{bmatrix}, \quad \theta = \begin{bmatrix} \theta_1 \\ \cdots \\ \theta_m \end{bmatrix}$$

- expansão unidirecional (HÄRDLE & STOKER, 1989):

$$\hat{g}(\mathbf{x}, \theta) = f(\mathbf{x}^T \theta_1, \theta_2), \quad \theta = \begin{bmatrix} \theta_1 \\ \cdots \\ \theta_2 \end{bmatrix}$$

- redução de dimensionalidade (LI, 1991):

$$\hat{g}(\mathbf{x}, \theta) = f(\mathbf{x}^T \theta_1, \dots, \mathbf{x}^T \theta_j, \theta_{j+1}), \quad \theta = \begin{bmatrix} \theta_1 \\ \cdots \\ \vdots \\ \cdots \\ \theta_{j+1} \end{bmatrix}, \quad j \leq m;$$

- expansão unidirecional aditiva (FRIEDMAN & STUETZLE, 1981; HUBER, 1985):

$$\hat{g}(\mathbf{x}, \theta) = f_1(\mathbf{x}^T \theta_1, \theta_2) + \cdots + f_n(\mathbf{x}^T \theta_{2n-1}, \theta_{2n}), \quad \theta = \begin{bmatrix} \theta_1 \\ \cdots \\ \vdots \\ \cdots \\ \theta_{2n} \end{bmatrix}$$

- expansão unidirecional multiplicativa (FRIEDMAN *et al.*, 1984):

$$\hat{g}(\mathbf{x}, \theta) = f_1(\mathbf{x}^T \theta_1, \theta_2) * \cdots * f_n(\mathbf{x}^T \theta_{2n-1}, \theta_{2n}), \quad \theta = \begin{bmatrix} \theta_1 \\ \cdots \\ \vdots \\ \cdots \\ \theta_{2n} \end{bmatrix}$$

Dado o problema de aproximação, uma questão natural que surge é como escolher a função de aproximação que melhor se adapta, ou seja, que tipo de restrição paramétrica deve ser imposta ao modelo de aproximação. SAMAROV (1993) apresenta algumas alternativas com base nos dados de aproximação $\{(\mathbf{x}_l, s_l)\}_{l=1}^N$ e em estimativas de derivadas médias (HÄRDLE & STOKER, 1989).

É evidente que as funções de aproximação mais simples são casos particulares das funções mais complexas. Particularmente, redes neurais artificiais são capazes de representar todos os modelos apresentados acima. Observe também que boa parte dos modelos apresentados realiza a aproximação com base em uma soma ponderada das variáveis de entrada, importante passo no sentido de reduzir a dimensionalidade do problema de aproximação.

D.3. Algoritmos de otimização paramétrica utilizando diferenciação

Para a maioria dos modelos de aproximação paramétrica $\hat{g}(\cdot, \theta)$, o problema de otimização apresentado na equação (D.2) tem a desvantagem de ser não-linear e não-convexo, mas as vantagens de ser irrestrito e permitir a aplicação de conceitos de cálculo variacional na obtenção da solução θ^* . Estas características impedem a existência de uma solução analítica, mas permitem obter processos iterativos de solução, a partir de uma condição inicial θ_0 , na forma:

$$\theta_{i+1} = \theta_i + \alpha_i \mathbf{d}_i, i \geq 0 \quad (\text{D.3})$$

onde $\theta_i \in \mathcal{R}^p$ é o vetor de parâmetros, $\alpha_i \in \mathcal{R}^+$ é um escalar que define o passo de ajuste e $\mathbf{d}_i \in \mathcal{R}^p$ é a direção de ajuste, todos definidos na iteração i . Na obtenção do passo e da direção de ajuste do processo iterativo dado pela equação (D.3), pode-se empregar os mais eficientes algoritmos de otimização utilizando diferenciação.

Uma vez definida a estrutura paramétrica $\hat{g}(\cdot, \theta)$, o esquema básico dos algoritmos a serem apresentados contém necessariamente as seguintes etapas (o operador $\|\cdot\|$ corresponde à norma euclidiana):

- Defina uma condição inicial θ_0 para o vetor de parâmetros;
- Defina um escalar $\varepsilon > 0$ arbitrariamente próximo de zero;
- Faça $i = 0$;
- Enquanto $\|\nabla J(\theta_i)\| > \varepsilon$ faça:

calcule α_i e $\mathbf{d}_i$ tal que $J(\theta_{i+1}) < J(\theta_i)$ para $\theta_{i+1} = \theta_i + \alpha_i \mathbf{d}_i$;

$i = i+1$.

Procedimentos de validação cruzada para garantia de bom desempenho em termos de generalização podem ser prontamente incorporados.

D.3.1. O método do gradiente

Dentre os métodos que utilizam diferenciação, o método do gradiente é o método mais simples de obtenção da direção $\mathbf{d}_i$, pois utiliza apenas informações de 1ª ordem. Na i -ésima iteração, a direção $\mathbf{d}_i$ é definida como a direção de maior decrescimento da função J .

Teorema D.1: A direção $\mathbf{d} = -\frac{\nabla J(\theta)}{\|\nabla J(\theta)\|}$ é a direção de maior decrescimento da função J .

Prova: Para $J: \Re^P \rightarrow \Re$ e $\theta, \mathbf{d} \in \Re^P$, com $\|\mathbf{d}\| \leq 1$, a derivada de J na direção $\mathbf{d}$ é dada por:

$$D(J, \mathbf{d}) = \lim_{\lambda \rightarrow 0} \frac{J(\theta + \lambda \mathbf{d}) - J(\theta)}{\lambda} = \nabla J(\theta)^T \mathbf{d}.$$

Logo, para provar o teorema basta mostrar que $\mathbf{d} = -\frac{\nabla J(\theta)}{\|\nabla J(\theta)\|}$ minimiza $D(J, \mathbf{d})$.

Seja γ o ângulo entre os vetores $\nabla J(\theta)$ e $\mathbf{d}$, então

$$\nabla J(\theta)^T \mathbf{d} = \|\nabla J(\theta)\| \cdot \|\mathbf{d}\| \cdot \cos \gamma \geq -\|\nabla J(\theta)\| \cdot \|\mathbf{d}\| \geq -\|\nabla J(\theta)\|.$$

Para $\mathbf{d} = -\frac{\nabla J(\theta)}{\|\nabla J(\theta)\|}$, como

$$\nabla J(\theta)^T \mathbf{d} = -\|\nabla J(\theta)\|,$$

conclui-se que $\mathbf{d} = -\frac{\nabla J(\theta)}{\|\nabla J(\theta)\|}$ minimiza $D(J, \mathbf{d})$.

A lei de ajuste do método do gradiente é, então, dada por:

$$\theta_{i+1} = \theta_i - \alpha_i \frac{\nabla J(\theta_i)}{\|\nabla J(\theta_i)\|}. \quad (\text{D.4})$$

Como a direção $\mathbf{d}_i$ escolhida é a de maior decrescimento da função $J(\theta)$ no ponto $\theta = \theta_i$, poder-se-ia imaginar que qualquer valor positivo para α_i fornece uma direção minimizante. Isto só é verdade se $J(\theta)$ é uma função linear, ou seja, com derivadas de ordem maior que 1 nulas. Como este não é o caso em estudo, tem-se o seguinte problema de busca unidirecional:

$$\min_{\alpha_i > 0} J(\theta_{i+1}), \text{ com } \theta_{i+1} \text{ dado pela equação (D.4)}. \quad (\text{D.5})$$

Como este problema deve ser resolvido a cada iteração i e levando-se em conta o elevado custo computacional vinculado à sua solução ($J(\theta)$ deve ser avaliado em vários pontos), é adotado aqui um procedimento alternativo. Em lugar do problema (D.5), resolve-se um outro problema de busca unidirecional mais simples:

Encontre $\alpha_i > 0$ tal que $J(\theta_{i+1}) < J(\theta_i)$.

Este problema pode ser resolvido adotando-se os seguintes passos a cada iteração:

- $\begin{cases} \alpha_i = 1 & \text{se } i = 0 \\ \alpha_i = 2\alpha_{i-1} & \text{se } i > 0 \end{cases};$
- $\theta_{i+1} = \theta_i - \alpha_i \frac{\nabla J(\theta_i)}{\|\nabla J(\theta_i)\|};$
- Enquanto $J(\theta_{i+1}) \geq J(\theta_i)$ faça:

$$\alpha_i = \frac{1}{2} \alpha_i;$$

$$\theta_{i+1} = \theta_i - \alpha_i \frac{\nabla J(\theta_i)}{\|\nabla J(\theta_i)\|}.$$

Sempre que $J(\theta)$ tiver pelo menos um mínimo, o método do gradiente associado a este procedimento de busca unidirecional vai seguramente fornecer uma solução θ^* , mínimo local do problema (D.2). No entanto, um problema prático importante está no número de iterações necessário para atingir esta solução. Utilizando apenas informações de 1ª ordem, o método do gradiente é sabido ser pouco eficiente, apresentando convergência lenta, principalmente nas vizinhanças de θ^* (BAZARAA & SHETTY, 1979; LUENBERGER, 1989). A introdução de um termo de momento junto à equação (D.4), na forma:

$$\theta_{i+1} = \theta_i - \alpha_i \frac{\nabla J(\theta_i)}{\|\nabla J(\theta_i)\|} + \beta_i (\theta_i - \theta_{i-1}) \quad (\text{D.6})$$

tem a vantagem de permitir acelerar a convergência e até superar mínimos locais, mas a desvantagem de criar mais um parâmetro a ser arbitrado (β_i). A definição de valores adequados para o par (α_i, β_i) não é uma tarefa simples, pois agora a lei de ajuste passa a ser uma equação a diferenças de 2ª ordem.

D.3.2. O método de Newton

O método de Newton resolve iterativamente o problema de otimização (D.2) aproximando a função $J(\theta)$ por uma função quadrática, a qual é minimizada exatamente. Sendo assim, em uma vizinhança de θ_i , a função $J(\theta)$ pode ser aproximada pela sua expansão em série de Taylor até 2ª ordem, dada por:

$$J_{seg}(\theta) = J(\theta_i) + \nabla J(\theta_i)^T (\theta - \theta_i) + (\theta - \theta_i)^T \nabla^2 J(\theta_i) (\theta - \theta_i) \quad (D.7)$$

onde $\nabla J(\theta_i)$ é o vetor gradiente e $\nabla^2 J(\theta_i)$ é a matriz hessiana de $J(\theta)$, ambos calculados no ponto $\theta = \theta_i$. Observe que a matriz $\nabla^2 J(\theta_i)$ é sempre simétrica.

O vetor θ_{i+1} é, então, a solução que minimiza exatamente $J_{seg}(\theta)$ dado pela equação (D.7), ou seja,

$$\frac{\partial J_{seg}(\theta_{i+1})}{\partial \theta_{i+1}} = 0, \quad (D.8)$$

produzindo

$$\theta_{i+1} = \theta_i - [\nabla^2 J(\theta_i)]^{-1} \nabla J(\theta_i). \quad (D.9)$$

Utilizando o mesmo raciocínio empregado no caso do método do gradiente, como a função $J(\theta)$ não é necessariamente quadrática, a minimização de sua aproximação quadrática $J_{seg}(\theta)$ pode não fornecer um solução θ_{i+1} tal que $J(\theta_{i+1}) < J(\theta_i)$. Por outro lado, é sabido que $J_{seg}(\theta)$ é tão mais próximo de $J(\theta)$ quanto mais próximos estiverem θ de θ_i (BAZARAA & SHETTY, 1979; LUENBERGER, 1989). Sendo assim, a modificação da lei de ajuste (D.9) para

$$\theta_{i+1} = \theta_i - \alpha_i [\nabla^2 J(\theta_i)]^{-1} \nabla J(\theta_i) \quad (D.10)$$

garante a existência de $0 < \alpha_i \leq 1$ tal que $J(\theta_{i+1}) < J(\theta_i)$. Um procedimento de busca unidirecional pode, então, ser aplicado a cada iteração, na forma:

- Defina r tal que $0 < r < 1$;
- $\alpha_i = 1$;
- $\theta_{i+1} = \theta_i - \alpha_i [\nabla^2 J(\theta_i)]^{-1} \nabla J(\theta_i)$;

- Enquanto $J(\theta_{i+1}) \geq J(\theta_i)$ faça:

$$\alpha_i = r\alpha_i;$$

$$\theta_{i+1} = \theta_i - \alpha_i [\nabla^2 J(\theta_i)]^{-1} \nabla J(\theta_i).$$

Nesta forma, o método de Newton ainda não apresenta qualquer garantia de convergência, pois nada pode ser afirmado sobre o sinal da matriz hessiana, que deve ser definida positiva. Portanto, é necessário testar a positividade de $\nabla^2 J(\theta_i)$ a cada iteração, e na eventualidade de se obter $\nabla^2 J(\theta_i) \leq 0$, deve-se aplicar um procedimento de positivação da matriz. Os dois resultados a seguir auxiliam nesta tarefa.

Teorema D.2: Uma matriz simétrica é definida positiva se seus autovalores forem todos positivos.

Prova: Veja STRANG (1988).

Teorema D.3: Se os autovalores da matriz simétrica $\mathbf{A}$ de dimensão $P \times P$ são $\{\lambda_1, \dots, \lambda_P\}$, então os autovalores da matriz $\mathbf{A} + \gamma \mathbf{I}$ são $\{\lambda_1 + \gamma, \dots, \lambda_P + \gamma\}$.

Prova: Veja STRANG (1988).

Destes resultados segue o seguinte corolário:

Corolário D.1: Dado que se conhece o menor autovalor $\lambda_{\min}$ de uma matriz simétrica $\mathbf{A}$, é sempre possível obter uma matriz definida positiva $\mathbf{B}$ a partir de $\mathbf{A}$ na forma:

- se $\lambda_{\min} > 0$ tome $\mathbf{B} = \mathbf{A}$;
- se $\lambda_{\min} \leq 0$ tome $\mathbf{B} = \mathbf{A} + (\varepsilon - \lambda_{\min}) \mathbf{I}$, com $\varepsilon > 0$ arbitrário.

A lei de ajuste do método de Newton assume, então, a forma:

$$\theta_{i+1} = \theta_i - \alpha_i \mathbf{M}_i^{-1} \nabla J(\theta_i), \quad (\text{D.11})$$

onde $\mathbf{M}_i$ é dado por

$$\begin{cases} \mathbf{M}_i = \nabla^2 J(\theta_i) & \text{se } \lambda_{\min}^{[i]} > 0 \\ \mathbf{M}_i = \nabla^2 J(\theta_i) + (\varepsilon - \lambda_{\min}^{[i]}) \mathbf{I} & \text{se } \lambda_{\min}^{[i]} \leq 0 \end{cases}, \quad (\text{D.12})$$

com $\lambda_{\min}^{[i]}$ o autovalor mínimo de $\mathbf{M}_i$.

O método de Newton modificado dado pela equação (D.11) vai convergir para um mínimo local sempre que a função $J(\theta)$ tiver mínimos. Observe que todos os resultados obtidos acima são imediatamente aplicáveis a problemas de maximização, desde que se troque o sinal da matriz hessiana, que agora deve ser definida negativa.

A figura D.1 apresenta alguns exemplos da aplicação do método de Newton modificado na maximização de uma função monovariável $J(\theta)$, demonstrando que o ajuste do método de Newton é sempre no sentido de atender o critério. Os pontos '●' correspondem a $(\theta_i, J(\theta_i))$, enquanto que os pontos '○' correspondem a $(\theta_{i+1}, J(\theta_{i+1}))$. A função $J(\theta)$ é representada pela curva de traço cheio, enquanto que $J_{seg}(\theta)$ (ou sua versão modificada) é representada pela curva tracejada.

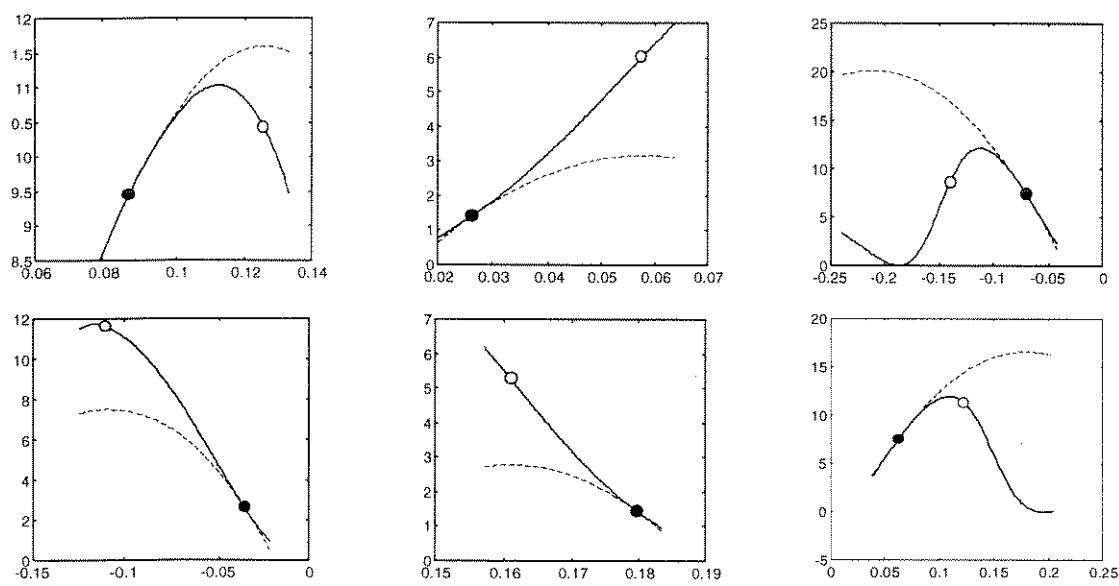


Figura D.1 - Exemplos da aplicação do método de Newton modificado

Observe que nos gráficos da coluna central foi necessário negatizar a hessiana, enquanto que nos gráficos da última coluna houve a necessidade de se utilizar $\alpha_i < 1$. Situações em que estes dois procedimentos sejam necessários simultaneamente também podem ocorrer.

É importante mencionar ainda que:

- em lugar de positivar $\mathbf{M}_i$ na equação (D.12) (negatizar em problemas de maximização), poder-se-ia, em princípio, utilizar qualquer outra matriz definida positiva de dimensões apropriadas. A razão para optar pelo processo de positivação da hessiana é a analogia com o tradicional método de Levenberg-Marquardt (DENNIS & SCHNABEL, 1983), que pode ser

interpretado como uma combinação entre a lei de ajuste do método do gradiente e a lei de ajuste do método de Newton.

- Apesar de apresentar uma taxa de convergência bem maior que o método do gradiente, o método de Newton é excessivamente custoso e numericamente instável do ponto de vista computacional, pois exige a inversão, a análise espectral e armazenagem de uma matriz que, apesar de simétrica, é da ordem do número de parâmetros P a serem ajustados. Em se tratando de redes neurais paramétricas, P pode ser da ordem de centenas a milhares de unidades.

D.3.3. O método do gradiente conjugado

O método do gradiente conjugado é um procedimento intermediário entre o método do gradiente e o método de Newton. Ele é projetado para exigir menos cálculos que o método de Newton e apresentar taxas de convergência maiores que as do método do gradiente. Tal como o método de Newton, este procedimento iterativo de otimização é projetado para obter soluções exatas para problemas quadráticos em um número finito de iterações.

D.3.3.1. Motivação para a origem do método

A lei de adaptação dos processos iterativos em estudo assume a forma da equação (D.3), sendo que, havendo convergência, a solução ótima $\theta^* \in \mathfrak{R}^P$ pode ser expressa por:

$$\theta^* = \theta_0 + \alpha_0 \mathbf{d}_0 + \alpha_1 \mathbf{d}_1 + \dots$$

Sendo um vetor no espaço $\mathfrak{R}^P$, θ^* pode sempre ser expresso na forma:

$$\theta^* = \alpha_0 \mathbf{d}_0 + \alpha_1 \mathbf{d}_1 + \dots + \alpha_{P-1} \mathbf{d}_{P-1} = \sum_{i=0}^{P-1} \alpha_i \mathbf{d}_i ,$$

desde que o conjunto $\{\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_{P-1}\}$ forme uma base de $\mathfrak{R}^P$ e $\alpha = [\alpha_0 \dots \alpha_{P-1}]^T$ seja a representação de θ^* nesta base (CHEN, 1984). Isto implica que dois procedimentos alternativos podem ser utilizados para obter θ^* em até P iterações:

- tome vetores $\mathbf{d}_i \in \mathfrak{R}^P$ ($i=0, \dots, P-1$) linearmente independentes e encontre a representação de θ^* na base $\{\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_{P-1}\}$;
- tome escalares $\alpha_i \in \mathfrak{R}$ ($i=0, \dots, P-1$) não todos nulos e encontre a base $\{\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_{P-1}\}$ que tenha $\alpha = [\alpha_0 \dots \alpha_{P-1}]^T$ como a representação de θ^* .

O primeiro procedimento é certamente mais adequado, principalmente para valores elevados de P , sendo o princípio de operação do método do gradiente conjugado.

D.3.3.2. Método das direções conjugadas

Esta seção apresenta uma forma de obtenção de $\alpha_0^*, \alpha_1^*, \dots, \alpha_{n-1}^*$ tal que

$$\theta^* = \sum_{i=0}^{P-1} \alpha_i^* \mathbf{d}_i. \quad (\text{D.13})$$

Para tanto, são utilizados os resultados a seguir.

Definição D.1: Dada uma matriz simétrica $\mathbf{A}$ de dimensão $P \times P$, as direções $\mathbf{d}_i \in \mathbb{R}^P$, $i=0, \dots, P-1$, são ditas serem $\mathbf{A}$ -conjugadas (ou $\mathbf{A}$ -ortogonais, ou então conjugadas em relação à matriz $\mathbf{A}$) se

$$\mathbf{d}_j^T \mathbf{A} \mathbf{d}_i = 0, \text{ para } i \neq j \text{ e com } i, j = 0, \dots, P-1.$$

Observe que o conceito de conjugação em relação a uma matriz $\mathbf{A}$ é uma generalização do conceito de ortogonalidade entre vetores, o qual é definido tomando-se $\mathbf{A} = \mathbf{I}$, onde $\mathbf{I}$ é a matriz identidade.

Teorema D.4: Para uma matriz $\mathbf{A}$ simétrica e definida positiva, direções $\mathbf{A}$ -conjugadas são necessariamente linearmente independentes.

Prova: Veja Hestenes (1980).

Corolário D.2: Se $\mathbf{A}$ é uma matriz simétrica, definida positiva e de dimensão $P \times P$, o conjunto de P direções $\mathbf{A}$ -conjugadas forma uma base do $\mathbb{R}^P$.

Com isso, os coeficientes α_j^* , $j=0, \dots, P-1$, podem ser determinados adotando-se o seguinte procedimento:

- dada uma matriz $\mathbf{A}$ simétrica, definida positiva e de dimensão $P \times P$, multiplicando-se a equação (D.13) à esquerda por $\mathbf{d}_j^T \mathbf{A}$, com $0 \leq j \leq P-1$, resulta:

$$\mathbf{d}_j^T \mathbf{A} \theta^* = \sum_{i=0}^{P-1} \alpha_i^* \mathbf{d}_j^T \mathbf{A} \mathbf{d}_i, \quad j = 0, \dots, P-1$$

- escolhendo as direções $\mathbf{d}_i \in \Re^P$, $i=0,\dots,P-1$, como sendo $\mathbf{A}$ -conjugadas, é possível aplicar os resultados apresentados acima para obter:

$$\alpha_j^* = \frac{\mathbf{d}_j^T \mathbf{A} \theta^*}{\mathbf{d}_j^T \mathbf{A} \mathbf{d}_j}, \quad j = 0, \dots, P-1. \quad (\text{D.14})$$

Esta expressão para α_j^* , $j=0,\dots,P-1$, ainda não representa uma solução viável para se obter θ^* , pois os próprios coeficientes são fornecidos em função de θ^* . Para eliminar θ^* da expressão (D.14) duas hipóteses adicionais devem ser consideradas:

- assuma que o problema de otimização é quadrático, ou seja,

$$J(\theta) = \frac{1}{2} \theta^T \mathbf{Q} \theta - \mathbf{b}^T \theta, \quad (\text{D.15})$$

então, no ponto de mínimo θ^* , deve valer:

$$\nabla J(\theta^*) = 0 \Rightarrow \mathbf{Q} \theta^* - \mathbf{b} = 0 \Rightarrow \mathbf{Q} \theta^* = \mathbf{b}; \quad (\text{D.16})$$

- assuma que $\mathbf{A} = \mathbf{Q}$.

Com isso, a equação (D.14) pode ser reescrita na forma:

$$\alpha_j^* = \frac{\mathbf{d}_j^T \mathbf{b}}{\mathbf{d}_j^T \mathbf{Q} \mathbf{d}_j}, \quad j = 0, \dots, P-1, \quad (\text{D.17})$$

produzindo

$$\theta^* = \sum_{j=0}^{P-1} \frac{\mathbf{d}_j^T \mathbf{b}}{\mathbf{d}_j^T \mathbf{Q} \mathbf{d}_j} \mathbf{d}_j. \quad (\text{D.18})$$

Logicamente, esta solução é válida apenas no caso de a função objetivo $J(\theta)$ ser quadrática, condição para que a matriz $\mathbf{Q}$ e o vetor $\mathbf{b}$ sejam constantes em todo o espaço $\Re^P$.

Como esta solução vai ser obtida iterativamente a partir de uma condição inicial θ_0 , é possível obter um resultado mais interessante, do ponto de vista prático, que a expressão (D.18), por permitir a substituição do vetor $\mathbf{b}$. Para tanto, assumindo solução iterativa, a solução ótima θ^* é expressa na forma

$$\theta^* = \theta_0 + \alpha_0^* \mathbf{d}_0 + \alpha_1^* \mathbf{d}_1 + \dots + \alpha_{P-1}^* \mathbf{d}_{P-1}. \quad (\text{D.19})$$

Com isso, os coeficientes α_j^* , $j=0,\dots,P-1$, são dados por:

$$\alpha_j^* = \frac{\mathbf{d}_j^T \mathbf{Q}(\theta^* - \theta_0)}{\mathbf{d}_j^T \mathbf{Q} \mathbf{d}_j}, j=1, \dots, P-1. \quad (\text{D.20})$$

Na j -ésima iteração ($0 < j < P-1$), θ_j assume a forma:

$$\theta_j = \theta_0 + \alpha_0^* \mathbf{d}_0 + \dots + \alpha_{j-1}^* \mathbf{d}_{j-1}, \quad (\text{D.21})$$

permitindo obter

$$\mathbf{d}_j^T \mathbf{Q} \theta_0 = \mathbf{d}_j^T \mathbf{Q} \theta_j. \quad (\text{D.22})$$

Substituindo a equação (D.22) na equação (D.20) resulta:

$$\alpha_j^* = \frac{\mathbf{d}_j^T \mathbf{Q}(\theta^* - \theta_j)}{\mathbf{d}_j^T \mathbf{Q} \mathbf{d}_j}, j=1, \dots, P-1. \quad (\text{D.23})$$

Com base na equação (D.16), e sabendo-se que

$$\nabla J(\theta_j) = \mathbf{Q} \theta_j - \mathbf{b},$$

a equação (D.23) pode ser colocada na forma:

$$\alpha_j^* = -\frac{\mathbf{d}_j^T \nabla J(\theta_j)}{\mathbf{d}_j^T \mathbf{Q} \mathbf{d}_j}, j=1, \dots, P-1. \quad (\text{D.24})$$

ou seja, a lei de ajuste do método das direções conjugadas é dada por:

$$\theta_{i+1} = \theta_i - \frac{\mathbf{d}_i^T \nabla J(\theta_i)}{\mathbf{d}_i^T \mathbf{Q} \mathbf{d}_i} \mathbf{d}_i, \quad (\text{D.25})$$

produzindo a solução ótima em P iterações, se o problema for quadrático.

D.3.3.3. Método do Gradiente Conjugado

Para a aplicação da lei de ajuste dada pela equação (D.25), é ainda necessário obter as direções $\mathbf{Q}$ -conjugadas $\mathbf{d}_i \in \mathbb{R}^P$, $i=0, \dots, P-1$. Para tanto, uma alternativa é tomá-las na forma (BAZARAA & SHETTY, 1979; HESTENES, 1980; LUENBERGER, 1984):

$$\begin{cases} \mathbf{d}_0 = -\nabla J(\theta_0) \\ \mathbf{d}_{i+1} = -\nabla J(\theta_{i+1}) + \beta_i \mathbf{d}_i \text{ para } i \geq 0 \end{cases} \quad \text{com} \quad \beta_i = \frac{\nabla J(\theta_{i+1})^T \mathbf{Q} \mathbf{d}_i}{\mathbf{d}_i^T \mathbf{Q} \mathbf{d}_i}. \quad (\text{D.26})$$

As equações (D.25) e (D.26) compõem o método do gradiente conjugado para a obtenção iterativa da solução do problema (D.2). Vale mais uma vez salientar que este método é especialmente desenvolvido para o caso de a função $J(\theta)$ ser quadrática, quando são válidas as seguintes propriedades (BAZARAA & SHETTY, 1979; HESTENES, 1980; LUENBERGER, 1984):

- o número de iterações necessário para se obter θ^* é finito e limitado pelo número de autovalores distintos da matriz hessiana $\mathbf{Q}$;
- enquanto ainda não se atingiu a solução, o gradiente $\nabla J(\theta_i)$ tem módulo maior que zero e é ortogonal a todas as direções $\mathbf{d}_j$ tal que $j < i$;
- cada iteração do método do gradiente conjugado é pelo menos tão eficiente quanto uma iteração do método do gradiente a partir do mesmo ponto;
- cada ciclo de P iterações é equivalente a uma iteração do método de Newton;
- não existe a necessidade de se inverter a matriz hessiana;
- ao contrário do método de Newton modificado e do método do gradiente, não existe a necessidade de se realizar qualquer tipo de busca unidirecional;
- o custo computacional de cada iteração do método do gradiente conjugado é comparável ao custo de uma iteração do método do gradiente;
- a propriedade essencial dos vetores $\mathbf{d}_i \in \mathbb{R}^P$, $i=0, \dots, P-1$, está em suas direções e não em seu módulo;
- a vantagem de se calcular as direções conjugadas com base no gradiente está na produção de um processo iterativo com convergência mais uniforme.

D.3.3.4. Aplicação a problemas não-quadráticos

Algumas modificações devem ser introduzidas junto ao método do gradiente conjugado para possibilitar sua aplicação a problemas não-quadráticos. A principal delas é a substituição da matriz $\mathbf{Q}$ pela matriz hessiana calculada no ponto θ_i , $\nabla^2 J(\theta_i)$. Com isso, o algoritmo do gradiente conjugado para problemas não-quadráticos assume a forma:

Passo 0: Atribua um valor inicial $\theta_0 \in \mathbb{R}^P$ para o vetor de parâmetros e um valor arbitrariamente pequeno para a constante $\varepsilon > 0$;

Passo 1: Calcule $\nabla J(\theta_0)$ e faça $\mathbf{d}_0 = -\nabla J(\theta_0)$;

Passo 2: Para $i=0, \dots, P-1$ faça:

(a) Calcule $\nabla^2 J(\theta_i)$;

(b) Faça $\theta_{i+1} = \theta_i + \alpha_i \mathbf{d}_i$, com $\alpha_i = -\frac{\mathbf{d}_i^T \nabla J(\theta_i)}{\mathbf{d}_i^T \nabla^2 J(\theta_i) \mathbf{d}_i}$;

(c) Calcule $\nabla J(\theta_{i+1})$;

(d) Se $i < P-1$, faça $\mathbf{d}_{i+1} = -\nabla J(\theta_{i+1}) + \beta_i \mathbf{d}_i$, onde $\beta_i = \frac{\nabla J(\theta_{i+1})^T \nabla^2 J(\theta_i) \mathbf{d}_i}{\mathbf{d}_i^T \nabla^2 J(\theta_i) \mathbf{d}_i}$;

Passo 3: Se $\|\nabla J(\theta_{i+1})\| > \varepsilon$, faça $\theta_0 = \theta_P$ e retorne ao Passo 1;

Dentre as vantagens do algoritmo do gradiente conjugado é possível destacar:

- ao contrário dos algoritmos desenvolvidos para os métodos do gradiente e de Newton, não é necessário qualquer tipo de busca unidirecional;
- não é preciso inverter a matriz hessiana $\nabla^2 J(\cdot)$ em nenhum momento.

São consideradas desvantagens do algoritmo do gradiente conjugado:

- devido à necessidade de se calcular $\nabla^2 J(\cdot)$ a cada iteração, o algoritmo não pode ser considerado globalmente convergente, pois não existe garantia de a matriz $\nabla^2 J(\cdot)$ ser definida positiva em todo o espaço $\mathbb{R}^P$;
- como o ajuste do vetor de parâmetros θ toma informações até 2ª ordem de $J(\theta)$, não existe garantia de que o passo do algoritmo do gradiente conjugado seja minimizante.
- principalmente para valores elevados de P , calcular e armazenar a matriz hessiana $\nabla^2 J(\cdot)$ a cada iteração representa um custo computacional acentuado.

Estes três problemas são solucionados nas próximas três seções.

D.3.3.5. A positividade da hessiana

Problemas de otimização não-quadráticos não apresentam a propriedade de convexidade. Sendo assim, a ausência de convergência global para problemas não-convexos é uma característica comum a todos os algoritmos iterativos que utilizam informação local, incluindo o algoritmo do gradiente conjugado.

Dado um ponto θ do espaço de parâmetros, a informação *local* corresponde ao conhecimento dos valores de derivadas de ordem *finita* da função objetivo $J(\cdot)$ no ponto θ , ou seja, só é possível prever o comportamento *global* de $J(\cdot)$, a partir da informação em um ponto θ , se *todas* as derivadas de $J(\cdot)$ neste ponto forem conhecidas exatamente. Logo, os *únicos* métodos de otimização passíveis de implementação prática são aqueles que utilizam informação local.

A forma de garantir convergência para um mínimo local de $J(\cdot)$, sempre que mínimos de $J(\cdot)$ existirem, é assegurar a positividade da matriz hessiana $\nabla^2 J(\cdot)$ a cada iteração. Com o seu método do gradiente conjugado escalonado, MOLLER (1993) apresenta um procedimento eficiente para garantir a positividade de $\nabla^2 J(\cdot)$ sem precisar recorrer a procedimentos de análise espectral ou métodos de busca. O ponto de partida é um resultado de FLETCHER (1987) baseado no método de Levenberg-Marquardt.

D.3.3.6. A garantia de ajustes minimizantes

O procedimento para garantia de que o ajuste de parâmetros a cada iteração é sempre no sentido de decrescimento da função $J(\cdot)$ pode seguir o procedimento de busca unidirecional já adotado para os métodos do gradiente e de Newton, ou então utilizar a função proposta por MOLLER (1993), dada na forma:

$$\Delta_i = \frac{J(\theta_i) - J(\theta_i + \alpha_i \mathbf{d}_i)}{J(\theta_i) - J_{seg}(\theta_i + \alpha_i \mathbf{d}_i)}, \quad (\text{D.27})$$

que permite comparar a função original $J(\cdot)$ e sua aproximação de 2ª ordem $J_{seg}(\cdot)$. Quanto mais próximas estiverem estas duas funções, mais confiável é o passo do algoritmo do gradiente conjugado.

D.3.3.7. Cálculo exato da informação de 2ª ordem

O elevado custo computacional associado ao cálculo e armazenagem da hessiana $\nabla^2 J(\cdot)$ a toda iteração pode ser substancialmente reduzido aplicando-se os resultados de PEARLMUTTER (1994). Além de permitir o cálculo exato da informação de 2ª ordem, o custo computacional associado é da mesma ordem que o exigido para o cálculo da informação de 1ª ordem.

Utilizando um operador diferencial, é possível calcular exatamente o produto entre a matriz $\nabla^2 J(\cdot)$ e qualquer vetor desejado, sem a necessidade de se calcular e armazenar $\nabla^2 J(\cdot)$. Este resultado é de grande valia para o método do gradiente conjugado, em que a matriz hessiana $\nabla^2 J(\cdot)$ aparece invariavelmente multiplicada por um vetor.

Expandindo o vetor gradiente $\nabla J(\cdot)$ em torno de um ponto $\theta \in \mathbb{R}^p$ resulta:

$$\nabla J(\theta + \Delta\theta) = \nabla J(\theta) + \nabla^2 J(\theta)\Delta\theta + O(\|\Delta\theta\|^2) \quad (D.28)$$

onde $\Delta\theta$ representa uma pequena perturbação. Escolhendo $\Delta\theta = a\mathbf{v}$, com a uma constante positiva próxima de zero e $\mathbf{v} \in \mathbb{R}^p$ um vetor arbitrário, é possível calcular $\nabla^2 J(\theta)\mathbf{v}$ como segue:

$$\nabla^2 J(\theta)\mathbf{v} = \frac{1}{a} [\nabla J(\theta + a\mathbf{v}) - \nabla J(\theta) + O(a^2)] = \frac{\nabla J(\theta + a\mathbf{v}) - \nabla J(\theta)}{a} + O(a). \quad (D.29)$$

Tomando o limite quando $a \rightarrow 0$,

$$\nabla^2 J(\theta)\mathbf{v} = \lim_{a \rightarrow 0} \frac{\nabla J(\theta + a\mathbf{v}) - \nabla J(\theta)}{a} = \left. \frac{\partial}{\partial a} \nabla J(\theta + a\mathbf{v}) \right|_{a=0}. \quad (D.30)$$

Portanto, definindo o operador diferencial

$$\Psi_{\mathbf{v}} \{f(\theta)\} = \left. \frac{\partial}{\partial a} f(\theta + a\mathbf{v}) \right|_{a=0}, \quad (D.31)$$

é possível aplicá-lo a todas as operações realizadas para se obter o gradiente, produzindo

$$\Psi_{\mathbf{v}} \{\nabla J(\theta)\} = \nabla^2 J(\theta)\mathbf{v} \quad \text{e} \quad \Psi_{\mathbf{v}} \{\theta\} = \mathbf{v}. \quad (D.32)$$

Por ser um operador diferencial, $\Psi_{\mathbf{v}}\{\cdot\}$ obedece as regras usuais de diferenciação (RUDIN, 1976).

No capítulo 6, este operador é empregado na geração da informação de 2ª ordem a ser utilizada para ajuste de parâmetros de redes neurais multicamadas recorrentes via método do gradiente conjugado. Para tanto, é necessário aplicar $\Psi_{\mathbf{v}}\{\cdot\}$ a toda operação realizada para obter o vetor gradiente, isto é, a todos os passos do método de retropropagação dinâmica.

Bibliografia

Bibliografia

- ACKLEY, D., HINTON, G., SEJNOWSKI, T. A learning algorithm for Boltzmann machines. *Cognitive Science*, vol. 9, no. 1, pp. 147-169, 1985.
- ADAMS, R.A. **Sobolev spaces**. New York: Academic Press, 1975.
- ANDERSON, T.W. Estimating Linear Statistical Relationships. *The Annals of Statistics*, vol. 12, no. 1, pp. 1-45, 1984.
- ANDREWS, D.F. Plots of high-dimensional data. *Biometrics*, vol. 28, no. 1, pp. 125-136, 1972.
- ANDREWS, D.F., GNANADESIKAN, R., WARNER, J.L. Transformations of multivariate data. *Biometrics*, vol. 27, no. 4, pp. 825-840, 1971.
- ASH, T. Dynamic node creation in backpropagation neural networks. *Connection Science*, vol. 1, no. 4, pp. 365-375, 1989.
- ÄSTROM, K.J. Theory and Applications of Adaptive Control - A Survey. *Automatica*, vol. 19, no. 5, pp. 471-486, 1983.
- ÄSTROM, K.J., WITTENMARK, B. **Adaptive Control**. Menlo Park, CA: Addison-Wesley Publishing Company, 1989.
- ÄSTROM, K.J., WITTENMARK, B. **Computer-controlled systems - Theory and Design**. Englewood Cliffs (NJ): Prentice-Hall, 1990.
- ATKINSON, K. **An Introduction to Numerical Analysis**. New York: John Wiley & Sons, 1989.
- BABA, N., MOGAMI, Y., KOHZAKI, M., SHIRAISHI, Y., YOSHIDA, Y. A Hybrid Algorithm for Finding the Global Minimum of Error Function of Neural Networks and Its Applications. *Neural Networks*, vol. 7, no. 8, pp. 1253-1265, 1994.
- BALAKRISHNAN, K., HONAVAR, V. Evolutionary Design of Neural Architectures - A Preliminary Taxonomy and Guide to Literature. *Technical Report CS TR #95-01*, Department of Computer Science, Iowa State University, USA, 1995.
- BALDI, P.F., HORNIK, K. Learning in Linear Neural Networks: A Survey. *IEEE Transactions on Neural Networks*, vol. 6, no. 4, pp. 837-858, 1995.

- BÄRMANN, F., BIEGLER-KÖNIG, F. On a Class of Efficient Learning Algorithms for Neural Networks. *Neural Networks*, vol. 5, no. 1, pp. 139-144, 1992.
- BARNESLEY, M. **Fractals Everywhere**. San Diego, CA: Academic Press, 1988.
- BARRON, A.R. Universal approximation bounds for superpositions of a sigmoidal function. *IEEE Transactions on Information Theory*, vol. 39, no. 3, pp. 930-945, 1993.
- BARRON, A.R., BARRON, R.L. Statistical Learning Networks: A Unifying View. in E. Wegman (ed.) **Computing Science and Statistics. Proceedings of 20th Symposium of the Interface**, American Statistical Association, Washington, D.C., pp. 192-203, 1988.
- BARTLETT, M.S. Chance or Chaos? *Journal of the Royal Statistical Society A*, vol. 153, Part 3, pp. 321-347, 1990.
- BARTO, A.R. Connectionist learning for control: An overview. in T. Miller, R.S. Sutton, P.J. Werbos (eds.) **Neural Networks for Control**. Cambridge: M.I.T. Press, 1990.
- BATES, D.M., LINDSTROM, M.J., WAHBA, G., YANDELL, B.S. GCVPACK - Routines for Generalized Cross Validation. *Communications in Statistics - Simulation and Computation*, vol. 16, no. 1, pp. 263-297, 1987.
- BATTITI, R. First- and Second-Order Methods for Learning: Between Steepest Descent and Newton's Method. *Neural Computation*, vol. 4, no. 2, pp. 141-166, 1992.
- BAUER, H.-U., GEISEL, T. Nonlinear dynamics of feedback multilayer perceptrons. *Physical Review A*, vol. 42, no. 4, pp. 2401-2409, 1990.
- BAUM, E.B., HAUSSLER, D. What Size Net Gives Valid Generalization? in D.S. Touretzky (ed.) **Advances in Neural Information Processing Systems 1**, pp. 81-90, San Mateo, CA: Morgan Kaufmann Publishers, 1989.
- BAZARAA, M.S., SHETTY, C.M. **Nonlinear Programming - Theory and Algorithms**. New York: John Wiley & Sons, 1979.
- BELLMAN, R.E. **Adaptive Control Processes**. Princeton University Press, Princeton, NJ, 1961.
- BERNSTEIN, S. **Leçons sur les Propriétés Extremales**. Paris, 1926.
- BIEGLER-KÖNIG, F., BÄRMANN, F. A Learning Algorithm for Multilayered Neural Networks Based on Linear Least Squares Problems. *Neural Networks*, vol. 6, no. 1, pp. 127-131, 1993.
- BILLINGS, S. Identification of nonlinear systems - a survey. *IEE Proceedings*, vol. 126, no. 6, pp. 272-305, 1980.

- BINDER, K., YOUNG, A.P. Spin Glasses: Experimental Facts, Theoretical Concepts, and Open Questions. *Reviews of Modern Physics*, vol. 58, pp. 801-976, 1986.
- BISHOP, C. Exact Calculation of the Hessian Matrix for the Multilayer Perceptron. *Neural Computation*, vol. 4, no. 4, pp. 494-501, 1992.
- BLACK, H.S. Inventing the negative feedback amplifier. *IEEE Spectrum*, pp. 54-60, 1977.
- BOOKSTEIN, F.L. Discussion. in Huber, P.J. Projection pursuit. *The Annals of Statistics*, vol. 13, no. 2, pp. 481-484, 1985.
- BOX, G.E.P., JENKINS, G.M. **Time series analysis: Forecasting and control**. San Francisco (CA): Holden-Day, 1970.
- BREIMAN, L., FRIEDMAN, J.H. Estimating Optimal Transformations for Multiple Regression and Correlation (with discussion). *Journal of the American Statistical Association (JASA)*, vol. 80, no. 391, pp. 580-619, 1985.
- BREIMAN, L., FRIEDMAN, J.H. Predicting multivariate responses in multiple linear regression. *Journal of the Royal Statistical Society B*, 1996 (como citado em MALTHOUSE, 1995).
- BREIMAN, L., FRIEDMAN, J.H., OLSHEN, R.A., STONE, C.J. **Classification and Regression Trees**. Pacific Grove: Wadsworth, 1984.
- BRONSON, R. **Theory and Problems of Matrix Operations**. Schaum's Outline Series, New York: McGraw-Hill, 1989.
- BUCK, R.C. Linear Spaces and Approximation Theory. in R.E. Langer (ed.) **On Numerical Approximation**, pp. 11-23, Madison, 1959.
- BUJA, A., STUETZLE, W. Discussion. in Huber, P.J. Projection pursuit. *The Annals of Statistics*, vol. 13, no. 2, pp. 484-490, 1985.
- CARDALIAGUET, P., EUVRARD, G. Approximation of a Function and its Derivative with a Neural Network. *Neural Networks*, vol. 5, no. 2, pp. 207-220, 1992.
- CASDAGLI, M. Nonlinear Prediction of Chaotic Time Series. *Physica D*, vol. 35, no. 3, pp. 335-356, 1989.
- CHAKRABORTY, K., MEHROTRA, K., MOHAN, C.K., RANKA, S. Forecasting the Behavior of Multivariate Time Series Using Neural Networks. *Neural Networks*, vol. 5, no. 6, pp. 961-970, 1992.
- CHEN, C.-T. **Linear System Theory and Design**. New York: Holt, Rinehart and Winston, 1984.
- CHEN, H. Estimation of a projection-pursuit type regression model. *The Annals of Statistics*, vol. 19, no. 1, pp. 142-157, 1991.

- CHEN, S., BILLINGS, A. Neural networks for nonlinear dynamic system modelling and identification. *International Journal of Control*, vol. 56, no. 2, pp. 319-346, 1992.
- CHEN, T., CHEN, H., LIU, R. Approximation Capability in $C(\overline{\mathbb{R}}^n)$ by Multilayer Feedforward Networks and Related Problems. *IEEE Transactions on Neural Networks*, vol. 6, no. 1, pp. 25-30, 1995.
- CHEN, A.M., LU, H., HECHT-NIELSEN, R. On the Geometry of Feedforward Neural Network Error Surfaces. *Neural Computation*, vol. 5, no. 6, pp. 910-927, 1993.
- CHENEY, E.W. **Introduction to Approximation Theory**. New York: McGraw-Hill, 1966.
- CHENG, B., TITTERINGTON, D.M. Neural Networks: A Review from a Statistical Perspective, *Statistical Science*, vol. 9, no. 1, pp. 2-54, 1994.
- CHENTOUF, R., JUTTEN, C. Incremental Learning with a Stopping Criterion: experimental results. in J. Mira, F. Sandoval (eds.) IWANN'95: From Natural to Artificial Neural Computation. *Lecture Notes in Computer Science 930*, Springer-Verlag, pp. 519-526, 1995.
- CLEEREMANS, A., SERVAN-SCHREIBER, D., MCCLELLAND, J. Finite state automata and simple recurrent networks. *Neural Computation*, vol. 1, no. 3, pp. 372-381, 1989.
- COETZEE, F.M., STONICK, V.L. Topology and Geometry of Single Hidden Layer Network, Least Squares Weight Solutions. *Neural Computation*, vol. 7, no. 4, pp. 672-705, 1995.
- COHN, D., TESAURO, G. How Tight Are the Vapnik-Chervonenkis Bounds? *Neural Computation*, vol. 4, no. 2, pp. 249-269, 1992.
- CONNOR, J.T., ATLAS, L.E., MARTIN, D.R. Recurrent Networks and NARMA Modeling. in J.E. Moody, S.J. Hanson, R.P. Lippmann (eds.) **Advances in Neural Information Processing Systems 4**, pp. 301-308, San Mateo, CA: Morgan Kaufmann Publishers, 1992.
- CONNOR, J.T., MARTIN, D.R., ATLAS, L.E. Recurrent Neural Networks and Robust Time Series Prediction. *IEEE Transactions on Neural Networks*, vol. 5, no. 2, pp. 240-254, 1994.
- COTTER, N.E. The Stone-Weierstrass Theorem and Its Application to Neural Networks. *IEEE Transactions on Neural Networks*, vol. 1, no. 4, pp. 290-295, 1990.
- CRAVEN, P., WAHBA, G. Smoothing Noisy Data with Spline Functions: Estimating the Correct Degree of Smoothing by the Method of Generalized Cross-Validation. *Numerische Mathematik*, vol. 31, Fasc. 4, pp. 377-403, 1979.

- CRUZ JR., J.B., PERKINS, W.R. A New Approach to the Sensitivity Problem in Multivariable Feedback System Design. *IEEE Transactions on Automatic Control*, vol. 9, no. 3, pp. 216-223, 1964.
- CYBENKO, G. Approximation by superposition of sigmoidal functions. *Mathematics of Control, Signals and Systems*, vol. 2, no. 4, pp. 303-314, 1989.
- DAHMEN, W., MICCHELLI, C.A. Some remarks on ridge functions. *Approximation Theory and its Applications*, vol. 3, nos. 2-3, pp. 139-143, 1987.
- DAVIS, P.J. **Interpolation and Approximation**. Brown University, Dover Publications, 1975.
- DE BOOR, C. **A Practical Guide to Splines**. New York: Springer-Verlag, 1978
- DENNIS, J., SCHNABEL, R. **Numerical Methods for Unconstrained Optimization and Nonlinear Equations**. Englewood Cliffs (NJ): Prentice-Hall, 1983.
- DIACONIS, P. Discussion. in Huber, P.J. Projection pursuit. *The Annals of Statistics*, vol. 13, no. 2, pp. 494-496, 1985.
- DIACONIS, P., FRIEDMAN, J.H. Asymptotics of Graphical Projection Pursuit. *The Annals of Statistics*, vol. 12, no. 3, pp. 793-815, 1984.
- DIACONIS, P., SHAHSHAHANI, M. On Nonlinear Functions of Linear Combinations. *SIAM Journal on Scientific and Statistical Computing*, vol. 5, no. 1, pp. 175-191, 1984.
- DINGANKAR, A.T. **On applications of approximation theory to identification, control and classification**. Ph.D. Thesis, University of Texas, Austin, USA, 1995.
- DONOHO, D.L., JOHNSTONE, I.M. Projection-based approximation and a duality with kernel methods. *The Annals of Statistics*, vol. 17, no. 1, pp. 58-106, 1989.
- DOYLE, J.C., GLOVER, K., KHARGONEKAR, P.P., FRANCIS, B.A. State-Space Solutions to Standard H_2 and H_∞ Control Problems. *IEEE Transactions on Automatic Control*, vol. 34, no. 8, pp. 831-847, 1989.
- ELMAN, J. L. Finding structure in time. *Cognitive Science*, vol. 14, no. 2, pp. 179-211, 1990.
- EUBANK, R.L. **Spline Smoothing and Nonparametric Regression**. New York: Marcel Dekker, 1988.
- FAHLMAN, S.E. The Recurrent Cascade-Correlation Architecture. in R.P. Lippmann, J.E. Moody, D.S. Touretzky (eds.) **Advances in Neural Information Processing Systems 3**, pp. 190-196, San Mateo, CA: Morgan Kaufmann Publishers, 1991.

- FAHLMAN, S.E., LEBIERE, C. The Cascade-Correlation Learning Architecture. in D.S. Touretzky (ed.) **Advances in Neural Information Processing Systems 2**, pp. 524-532, San Mateo, CA: Morgan Kaufmann Publishers, 1990.
- FEJÉR, L. Über die Lage der Nullstellen von Polynomen, die aus Minimumforderungen gewisser Art entspringen. *Math. Ann.*, vol. 85, pp. 41-48, 1922.
- FLETCHER, R. **Practical Methods of Optimization**. New York: John Wiley & Sons, 1987.
- FREAN, M. The upstart algorithm: A method for constructing and training feedforward neural networks. *Neural Computation*, vol. 2, no. 2, pp. 198-209, 1990.
- FRIEDMAN, J.H. SMART User's Guide. *Report LCM001*, Department of Statistics, Stanford University, 1984.
- FRIEDMAN, J.H. Discussion. in Huber, P.J. Projection pursuit (with Discussion). *The Annals of Statistics*, vol. 13, no. 2, pp. 475-481, 1985.
- FRIEDMAN, J.H. Exploratory Projection Pursuit. *Journal of the American Statistical Association (JASA)*, vol. 82, no. 397, pp. 249-266, 1987.
- FRIEDMAN, J.H. Adaptive Spline Networks. in R.P. Lippmann, J.E. Moody, D.S. Touretzky (eds.) **Advances in Neural Information Processing Systems 3**, pp. 675-683, San Mateo, CA: Morgan Kaufmann Publishers, 1991a.
- FRIEDMAN, J.H. Multivariate adaptive regression splines (with Discussion). *The Annals of Statistics*, vol. 19, no. 1, pp. 1-141, 1991b.
- FRIEDMAN, J.H. An overview of predictive learning and function approximation. in J.H. Friedman, H. Wechsler (eds.) *From Statistics to Neural Networks: Theory and Pattern Recognition Applications. Proceedings of the NATO/ASI Workshop*, Springer-Verlag, 1994.
- FRIEDMAN, J.H., STUETZLE, W. Projection Pursuit Regression. *Journal of the American Statistical Association (JASA)*, vol. 76, no. 376, pp. 817-823, 1981.
- FRIEDMAN, J.H., STUETZLE, W., SCHROEDER, A. Projection Pursuit Density Estimation. *Journal of the American Statistical Association (JASA)*, vol. 79, no. 387, pp. 599-608, 1984.
- FRIEDMAN, J.H., TUKEY, J. A Projection Pursuit Algorithm for Exploratory Data Analysis. *IEEE Transactions on Computers*, vol. 23, no. 9, pp. 881-890, 1974.
- FUJITA, O. Optimization of the Hidden Unit Function in Feedforward Neural Networks. *Neural Networks*, vol. 5, no. 5, pp. 755-764, 1992.

- FUNAHASHI, K.-I. On the approximate realization of continuous mappings by neural networks. *Neural Networks*, vol. 2, no. 3, pp. 183-192, 1989.
- FUNAHASHI, K.-I., NAKAMURA, Y. Approximation of dynamical systems by continuous time recurrent neural networks. *Neural Networks*, vol. 6, no. 5, pp. 801-806, 1993.
- GALLANT, A.R., WHITE, H. On Learning the Derivatives of an Unknown Mapping With Multilayer Feedforward Networks. *Neural Networks*, vol. 5, no. 1, pp. 129-138, 1992.
- GEMAN, S., BIENENSTOCK, E., DOURSAT, R. Neural Networks and the Bias/Variance Dilemma. *Neural Computation*, vol. 4, no. 1, pp. 1-58, 1992.
- GHOSH, J., TUMER, K. Structural adaptation and generalization in supervised feed-forward networks. *Journal of Artificial Neural Networks*, vol. 1, no. 4, pp. 431-458, 1994.
- GILES, C., MILLER, C., CHEN, D., CHEN, H., SUN, G., LEE, Y. Learning and extracting finite state automata with second-order recurrent neural networks. *Neural Computation*, vol. 4, no. 3, pp. 393-405, 1992.
- GIROSI, F. Some extensions of radial basis functions and their applications in artificial intelligence. *Computers and Mathematics with Applications*, vol. 24, no. 12, pp. 61-80, 1992.
- GIROSI, F., JONES, M., POGGIO, T. Regularization Theory and Neural Networks Architectures. *Neural Computation*, vol. 7, no. 2, pp. 219-269, 1995.
- GIROSI, F., POGGIO, T. Representation Properties of Networks: Kolmogorov's Theorem Is Irrelevant. *Neural Computation*, vol. 1, no. 4, pp. 465-469, 1989.
- GIROSI, F., POGGIO, T., CAPRILE, B. Extensions of a Theory of Networks for Approximation and Learning: Outliers and Negative Examples. in R.P. Lippmann, J.E. Moody, D.S. Touretzky (eds.) **Advances in Neural Information Processing Systems 3**, pp. 750-756, San Mateo, CA: Morgan Kaufmann Publishers, 1991.
- GOLUB, G., REINSCH, C. Singular value decomposition and least squares solutions. *Numerische Mathematik*, vol. 14, no. 5, pp. 403-420, 1970.
- GOVER, J.C. Discussion. in Jones, M.C., Sibson, R. What is Projection Pursuit? *Journal of the Royal Statistical Society A*, vol. 150, Part 1, pp. 19-21, 1987.
- HABER, R., UNBEHAUEN, H. Structure Identification of Nonlinear Dynamic Systems - A Survey on Input / Output Approaches. *Automatica*, vol. 26, no. 4, pp. 651-677, 1990.
- HALL, P. On Polynomial-based Projection Indices for Exploratory Projection Pursuit. *The Annals of Statistics*, vol. 17, no. 2, pp. 589-605, 1989a.

- HALL, P. On Projection Pursuit Regression. *The Annals of Statistics*, vol. 17, no. 2, pp. 573-588, 1989b.
- HALL, P., LI, K.-C. On almost linearity of low dimensional projections from high dimensional data. *The Annals of Statistics*, vol. 21, no. 2, pp. 867-889, 1993.
- HANSON, S.J. Meiosis Networks. in D.S. Touretzky (ed.) **Advances in Neural Information Processing Systems 2**, pp. 533-541, San Mateo, CA: Morgan Kaufmann Publishers, 1990.
- HÄRDLE, W., STOKER, T. Investigating Smooth Multiple Regression by the Method of Average Derivatives. *Journal of the American Statistical Association (JASA)*, vol. 84, no. 408, pp. 986-995, 1989.
- HARTMAN, E.J., KEELER, J.D., KOWALSKI, J.M. Layered neural networks with gaussian hidden units as universal approximators. *Neural Computation*, vol. 2, no. 2, pp. 210-215, 1990.
- HARVEY, A.C. **The econometric analysis of time series**. Cambridge: M.I.T. Press, 1990.
- HASSIBI, B., STORK, D.G. Second Order Derivatives for Network Pruning: Optimal Brain Surgeon. in S.J. Hanson, J.D. Cowan, C.L. Giles (eds.) **Advances in Neural Information Processing Systems 5**, pp. 164-171, San Mateo, CA: Morgan Kaufmann Publishers, 1993.
- HASTIE, T., TIBSHIRANI, R. Generalized Linear Models (with Discussion). *Statistical Science*, vol. 1, no. 3, pp. 297-318, 1986.
- HECHT-NIELSEN, R. Kolmogorov's mapping neural network existence theorem. *Proceedings of the IEEE International Conference on Neural Networks*, vol. III, pp. 11-13, 1987.
- HECHT-NIELSEN, R. **Neurocomputing**. New York: Addison-Wesley Publishing Company, 1990.
- HECKMAN, N. Spline Smoothing in Partly Linear Models. *Journal of the Royal Statistical Society B*, vol. 48, no. 2, pp. 244-248, 1986.
- HEERMANN, P.D. **Neural Network Techniques for Stable Learning Control of Nonlinear Systems**. Ph.D. Thesis, University of Texas, Austin, 1992.
- HERTZ, J., KROGH, A., PALMER, R.G. **Introduction to the Theory of Neural Computation**. New-York: Addison-Wesley Publishing Company, 1991.
- HESTENES, M.R. **Conjugate Direction Methods in Optimization**. New York: Springer-Verlag, 1980.

- HILBERT, D. Mathematische Probleme. *Nachr. Akad. Wiss. Göttingen*, pp. 253-297, 1900; *Gesammelte Abhandlungen*, Bd. 3, pp. 290-329, 1935.
- HIROSE, Y., YAMASHITA, K., HIJIIYA, S. Back-propagation algorithm which varies the number of hidden units. *Neural Networks*, vol. 4, no. 1, pp. 61-66, 1991.
- HOPFIELD, J. Neural networks and physical systems with emergent collective computational abilities. *Proceedings of the National Academy of Sciences* 79, p. 2554-2558, 1982.
- HORNIK, K. Approximation capabilities of multilayer feedforward networks. *Neural Networks*, vol. 4, no. 2, pp. 251-257, 1991.
- HORNIK, K., STINCHCOMBE, M., WHITE, H. Multi-layer feedforward networks are universal approximators. *Neural Networks*, vol. 2, no. 5, pp. 359-366, 1989.
- HORNIK, K., STINCHCOMBE, M., WHITE, H. Universal approximation of an unknown function and its derivatives using multilayer feedforward networks. *Neural Networks*, vol. 3, no. 5, pp. 551-560, 1990.
- HORNIK, K., STINCHCOMBE, M., WHITE, H., AUER, P. Degree of Approximation Results for Feedforward Networks Approximating Unknown Mappings and Their Derivatives. *Neural Computation*, vol. 6, no. 6, pp. 1262-1275, 1994.
- HSING, T., CARROLL, R.J. An Asymptotic for Sliced Inverse Regression. *The Annals of Statistics*, vol. 20, no. 2, pp. 1040-1061, 1992.
- HUBER, P.J. Projection pursuit (with Discussion). *The Annals of Statistics*, vol. 13, no. 2, pp. 435-475, 1985.
- HUNT, K.J., SBARBARO, D., ZBIKOWSKI, R., GAWTHROP, P.J. Neural Networks for Control Systems - A Survey. *Automatica*, vol. 28, no. 6, pp. 1083-1112, 1992.
- HWANG, J.-N., LAY, S.R., MAECHLER, M., MARTIN, R.D., SCHIMERT, J. Regression modeling in back-propagation and projection pursuit learning. *IEEE Transactions on Neural Networks*, vol. 5, no. 3, pp. 342-353, 1994.
- HWANG, J.-N., LI, H., MAECHLER, M., MARTIN, R.D., SCHIMERT, J. A Comparison of Projection Pursuit and Neural Network Regression Modeling. in J.E. Moody, S.J. Hanson, R.P. Lippmann (eds.) **Advances in Neural Information Processing Systems 4**, pp. 1159-1166, San Mateo, CA: Morgan Kaufmann Publishers, 1992.
- IGLESIAS, P.A., GLOVER, K. State-space approach to discrete-time H_∞ control. *International Journal of Control*, vol. 54, no. 5, pp. 1031-1073, 1991.

- INTRATOR, N. Combining Exploratory Projection Pursuit and Projection Pursuit Regression with Application to Neural Networks. *Neural Computation*, vol. 5, no. 3, pp. 443-455, 1993a.
- INTRATOR, N. On the Use of Projection Pursuit Constraints for Training Neural Networks. in S.J. Hanson, J.D. Cowan, C.L. Giles (eds.) **Advances in Neural Information Processing Systems 5**, pp. 3-10, San Mateo, CA: Morgan Kaufmann Publishers, 1993b.
- INTRATOR, N. On the combination of supervised and unsupervised learning. *Physica A*, vol. 200, nos. 1-4, pp. 655-661, 1993c.
- ISIDORI, A. **Nonlinear Control Systems - An Introduction**. Berlin: Springer-Verlag, 1989.
- ITO, Y. Approximation Capability of Layered Neural Networks with Sigmoid Units on Two Layers. *Neural Computation*, vol. 6, no. 6, pp. 1233-1243, 1994.
- JOLLIFFE, I.T. **Principal Component Analysis**. Springer Series in Statistics, New York: Springer-Verlag, 1986.
- JONES, L.K. On a conjecture of Huber concerning the convergence of projection pursuit regression. *The Annals of Statistics*, vol. 15, no. 2, pp. 880-882, 1987.
- JONES, L.K. Constructive Approximations for Neural Networks by Sigmoidal Functions. *Proceedings of the IEEE*, vol. 78, no. 10, 1990.
- JONES, L.K. A Simple Lemma on Greedy Approximation in Hilbert Space and Convergence Rates for Projection Pursuit Regression and Neural Network Training. *The Annals of Statistics*, vol. 20, no. 1, pp. 608-613, 1992.
- JONES, M.C., SIBSON, R. What is Projection Pursuit? (with Discussion). *Journal of the Royal Statistical Society A*, vol. 150, Part 1, pp. 1-36, 1987.
- JORDAN, M.I. **The Learning of Representations for Sequential Performance**. Ph.D. Thesis, University of California, San Diego, 1985.
- KAHANE, J.-P. Sur le théorème de superposition de Kolmogorov. *Journal of Approximation Theory*, vol. 13, no. 3, pp. 229-234, 1975.
- KAILATH, T. **Linear Systems**. Englewood Cliffs (NJ): Prentice-Hall, 1980.
- KARNIN, E.D. A simple procedure for pruning back-propagation trained neural networks. *IEEE Transactions on Neural Networks*, vol. 1, no. 2, pp. 239-242, 1990.
- KIRKPATRICK, S., GELATT JR., C.D., VECCHI, M.P. Optimization by Simulated Annealing. *Science*, vol. 220, no. 4598, pp. 671-680, 1983.

- KOLEN, J.F. **Exploring the Computational Capabilities of Recurrent Neural Networks.** Ph.D. Thesis, The Ohio State University, 1994.
- KOLEN, J.F., POLLACK, J.B. Back Propagation is Sensitive to Initial Conditions. *in* R.P. Lippmann, J.E. Moody, D.S. Touretzky (eds.) **Advances in Neural Information Processing Systems 3**, pp. 860-867, San Mateo, CA: Morgan Kaufmann Publishers, 1991.
- KOLMOGOROV, A.N. On the Representation of Continuous Functions of Many Variables by Superposition of Continuous Functions of One Variable and Addition. *Dokl. Akad. Nauk. SSSR*, vol. 114, pp. 953-956, 1957.
- KRUSKAL, J.B. Toward a practical method which helps uncover the structure of a set of multivariate observations by finding the linear transformation which optimizes a new 'index of condensation'. *in* R.C. Milton, J.A. Nelder (eds.) **Statistical Computation**, pp. 427-440, New York: Academic Press, 1969.
- KUAN, C.-M., HORNIK, K., WHITE, H. A Convergence Result for Learning in Recurrent Neural Networks. *Neural Computation*, vol. 6, no. 3, pp. 420-440, 1994.
- KURKOVÁ, V. Kolmogorov's Theorem Is Relevant. *Neural Computation*, vol. 3, no. 4, pp. 617-622, 1991.
- KURKOVÁ, V. Kolmogorov's Theorem and Multilayer Neural Networks. *Neural Networks*, vol. 5, no. 3, pp. 501-506, 1992.
- KWOK, T.-Y., YEUNG, D.-Y. Constructive Feedforward Neural Networks for Regression Problems: A Survey. *Technical Report HKUST-CS95-43*, Department of Computer Science, Hong Kong University of Science and Technology, Hong Kong, 1995.
- LAPEDES, A.S., FARBER, R.M. Nonlinear Signal Processing Using Neural Networks: Prediction and System Modeling. *Technical Report LA-UR-87-2662*, Los Alamos National Laboratory, 1987.
- LAY, S.-R., HWANG, J.-N., YOU, S.-S. Extensions to Projection Pursuit Learning Networks with Parametric Smoothers. *Proceedings of the IEEE International Conference on Neural Networks (ICNN'94)*, vol. 3, pp. 1325-1330, 1994.
- LE CUN, Y., DENKER, J.S., SOLLA, S.A. Optimal Brain Damage. *in* D.S. Touretzky (ed.) **Advances in Neural Information Processing Systems 2**, pp. 598-605, San Mateo, CA: Morgan Kaufmann Publishers, 1990.

- LEE, S., KIL, R.M. A Gaussian potential function network with hierarchically self-organizing learning. *Neural Networks*, vol. 4, no. 2, pp. 207-224, 1991.
- LEE GILES, C., CHEN, D., SUN, G.-Z., CHEN, H.-H., LEE, Y.-C., GOUDREAU, M.W. Constructive Learning of Recurrent Neural Networks: Limitations of Recurrent Cascade Correlation and a Simple Solution. *IEEE Transactions on Neural Networks*, vol. 6, no. 4, pp. 829-836, 1995.
- LEVIN, A.U., NARENDRA, K.S. Identification Using Feedforward Networks. *Neural Computation*, vol. 7, no. 2, pp. 349-357, 1995.
- LEVIS, A.H., MARCUS, S.I., PERKINS, W.R., KOKOTOVIC, P., ATHANS, M., BROCKETT, R.W., WILLSKY, A.S. Challenges to Control: A Collective View - Report of the Workshop Held at the University of Santa Clara on September 18-19, 1986. *IEEE Transactions on Automatic Control*, vol. 32, no. 4, pp. 275-285, 1987.
- LI, K.-C. Sliced Inverse Regression for Dimension Reduction. *Journal of the American Statistical Association (JASA)*, vol. 86, no. 414, pp. 316-342, 1991.
- LI, K.-C. On Principal Hessian Directions for Data Visualization and Dimension Reduction: Another Application of Stein's Lemma. *Journal of the American Statistical Association (JASA)*, vol. 87, no. 420, pp. 1025-1039, 1992.
- LIN, J.-N., UNBEHAUEN, R. On the Realization of a Kolmogorov Network. *Neural Computation*, vol. 5, no. 1, pp. 18-20, 1993.
- LJUNG, L. **System Identification: Theory for the User**. New Jersey: Prentice-Hall, Inc., 1987.
- LOGAN, B.F. The Uncertainty Principle in Reconstructing Functions from Projections. *Duke Mathematical Journal*, vol. 42, no. 4, pp. 661-706, 1975.
- LOGAN, B.F., SHEPP, L.A. Optimal Reconstruction of a Functions from its Projections. *Duke Mathematical Journal*, vol. 42, no. 4, pp. 645-659, 1975.
- LORENTZ, G.G. **Approximation of Functions**. New York: Holt, Reinhart and Winston, 1966.
- LUENBERGER, D.G. **Optimization by Vector Space Methods**. New York: John Wiley & Sons, 1969.
- LUENBERGER, D.G. **Linear and Nonlinear Programming**. Reading, Massachusetts: Addison-Wesley Publishing Company, 1989.
- LYTTKENS, E. Regression aspects of canonical correlations. *Journal of Multivariate Analysis*, vol. 2, no. 4, pp. 418-439, 1972.

- MCCLELLAND, J.L., RUMELHART, D.E., editors. **Parallel Distributed Processing: Explorations in the Microstructure of Cognition**. Cambridge: M.I.T. Press, 1986.
- MADYCH, W.R., NELSON, S.A. Multivariate Interpolation and Conditionally Positive Definite Functions. *Approximation Theory and its Applications*, vol. 4, no. 4, pp. 77-89, 1988.
- MADYCH, W.R., NELSON, S.A. Multivariate Interpolation and Conditionally Positive Definite Functions II. *Mathematics of Computation*, vol. 54, no. 189, pp. 211-230, 1990.
- MAGNUS, W., OBERHETTINGER, F., SONI, R.P. **Formulas and Theorems for the Special Functions of Mathematical Physics**. Berlin: Springer-Verlag, 1966.
- MALTHOUSE, E.C. **Nonlinear Partial Least Squares**. Ph.D. Thesis, Northwestern University, Evanston, Illinois, 1995.
- MARCHAND, M., GOLEA, M., RUJÁN, P. A convergence theorem for sequential learning in two-layer perceptrons. *Europhysics Letters*, vol. 11, no. 6, pp. 487-492, 1990.
- MARDIA, K.V., KENT, J.T., BIBBY, J.M. **Multivariate Analysis**. London: Academic Press, 1979.
- MATLAB. **MATLAB User's Guide**. The MathWorks, Inc., Natick, MA, USA, 1992.
- MICCHELLI, C.A. Interpolation of Scattered Data: Distance Matrices and Conditionally Positive Definite Functions. *Constructive Approximation*, vol. 2, no. 1, pp. 11-22, 1986.
- MILLER, G.F., TODD, P.M., HEGDE, S.U. Designing neural networks using genetic algorithms. *Proceedings of the Third International Conference on Genetic Algorithms*, pp. 379-384, 1989.
- MOLLER, M.F. A Scaled Conjugate Gradient Algorithm for Fast Supervised Learning. *Neural Networks*, vol. 6, no. 4, pp. 525-533, 1993.
- MOODY, J.E. The *Effective* Number of Parameters: An Analysis of Generalization and Regularization in Nonlinear Learning Systems. in J.E. Moody, S.J. Hanson, R.P. Lippmann (eds.) **Advances in Neural Information Processing Systems 4**, pp. 847-854, San Mateo, CA: Morgan Kaufmann Publishers, 1992.
- MOODY, J.E. Prediction risk and architecture selection for neural networks. in V. Cherkassky, J.H. Friedman, H. Wechsler (eds.) *From Statistics to Neural Networks. Proceedings of the NATO/ASI Workshop*, pp. 143-156, Springer-Verlag, 1994.
- MORGAN, N., BOURLARD, H. Generalization and Parameter Estimation in Feedforward Nets: Some Experiments. in D.S. Touretzky (ed.) **Advances in Neural Information Processing Systems 2**, pp. 630-637, San Mateo, CA: Morgan Kaufmann Publishers, 1990.

- NABHAN, T.M., ZOMAYA, A.Y. Toward generating neural networks structures for function approximation. *Neural Networks*, vol. 7, no. 1, pp. 89-99, 1994.
- NARENDRA, K.S., MUKHOPADHYAY, S. Adaptive Control of Nonlinear Multivariable Systems Using Neural Networks. *Neural Networks*, vol. 7, no. 5, pp. 737-752, 1994.
- NARENDRA, K.S., PARTHASARATHY, K. Identification and control of dynamical systems using neural networks. *IEEE Transactions on Neural Networks*, vol. 1, no. 1, pp 4-27, 1990.
- NARENDRA, K.S., PARTHASARATHY, K. Gradient methods for the optimization of dynamical systems containing neural networks. *IEEE Transactions on Neural Networks*, vol. 2, no. 2, pp 252-262, 1991.
- NERRAND, O., ROUSSEL-RAGOT, P., PERSONNAZ, L., DREYFUS, G. Neural Networks and Nonlinear Adaptive Filtering: Unifying Concepts and New Algorithms. *Neural Computation*, vol. 5, no. 2, pp. 165-199, 1993.
- NIJMEIJER, H., VAN DER SCHAFT, A. **Nonlinear Dynamical Control Systems**. New York: Springer-Verlag, 1990.
- OPPENHEIM, A.V., SCHAFER, R.W. **Discrete-Time Signal Processing**. New Jersey: Prentice-Hall, Inc., 1989.
- OTT, E. **Chaos in Dynamical Systems**. New York: Cambridge University Press, 1993.
- PAREKH, R., YANG, J., HONAVAR, V. Constructive Neural Network Learning Algorithms for Multi-Category Pattern Classification. *ISU CS-TR 95-15*, Artificial Intelligence Research Group, Department of Computer Science, Iowa State University, USA, 1995.
- PARK, J., SANDBERG, I.W. Universal approximation using radial basis function networks. *Neural Computation*, vol. 3, no. 2, pp. 246-257, 1991.
- PEARLMUTTER, B.A. Fast Exact Multiplication by the Hessian. *Neural Computation*, vol. 6, no. 1, pp. 147-160, 1994.
- PHAM, D. T., LIU, X. Dynamic system modelling using partially recurrent neural networks. *Journal of Systems Engineering*, pp. 134-141, 1992.
- PINEDA, F.J. Recurrent Backpropagation and the Dynamical Approach to Adaptive Neural Computation. *Neural Computation*, vol. 1, no. 2, pp. 161-172, 1989.
- PISIER, G. Remarques sur un resultat non publié de B. Maurey. *Seminaire D'analyse Fonctionnelle 1980-1981*, Ecole Polytechnique, Centre de Mathematiques, Palaiseau, 1981.
- POGGIO, T. GIROSI, F. Networks for Approximation and Learning *Proceedings of the IEEE*, Special Issue on Neural Networks, vol. 78, no. 9, pp. 1481-1497, 1990a.

- POGGIO, T. GIROSI, F. Regularization algorithms for learning that are equivalent to multilayer networks. *Science*, vol. 247, no. 4945, pp. 978-982, 1990b.
- POPE, S.B., GADH, R. Fitting noisy data using cross-validated cubic smoothing splines. *Communications in Statistics - Simulation and Computation*, vol. 17, no. 2, pp. 349-376, 1988.
- PRIESTLEY, M.B. **Non-linear and Non-stationary Time Series Analysis**. San Diego, CA: Academic Press, 1988.
- RALSTON, A. **A First Course in Numerical Analysis**. London: McGraw-Hill, 1965.
- RAO, C.R., MITRA, S.K. **Generalized Inverse of Matrices and its Applications**. Toronto: John Wiley & Sons, 1971.
- RATKOWSKY, D.A. **Handbook of Nonlinear Regression Models**. Statistics: Textbooks and Monographs, vol. 107, New York: Marcel Dekker, 1989.
- REED, R. Pruning algorithms - a survey. *IEEE Transactions on Neural Networks*, vol. 4, no. 5, pp. 740-747, 1993.
- REINSCH, C.M. Smoothing by spline functions. *Numerische Mathematik*, vol. 10, no. 3, pp. 177-183, 1967.
- REINSCH, C.M. Smoothing by spline functions, II. *Numerische Mathematik*, vol. 16, no. 5, pp. 451-454, 1971.
- RIPLEY, B.D. Neural Networks and Related Methods for Classification. *Journal of the Royal Statistical Society B*, vol. 56, no. 3, pp. 409-456, 1994.
- ROBBINS, H., MONRO, S. A stochastic approximation method. *Annals of Mathematical Statistics*, vol. 22, pp. 400-407, 1951.
- ROBINSON, P.M. Root N -Consistent Semiparametric Regression. *Econometrica*, vol. 56, no. 4, pp. 931-954, 1988.
- ROOSEN, C.B., HASTIE, T.J. Automatic Smoothing Spline Projection Pursuit. *Journal of Computational and Graphical Statistics*, vol. 3, no. 3, pp. 235-248, 1994.
- RUDIN, W. **Principles of Mathematical Analysis**. London: McGraw-Hill, 1976.
- RUMELHART, D.E., HINTON, G.E., WILLIAMS, R.J. Learning internal representations by error propagation. in J.L. McClelland, D.E. Rumelhart (eds.) **Parallel Distributed Processing: Explorations in the Microstructure of Cognition**. vol. 1, pp. 318-362, Cambridge: M.I.T. Press, 1986.

- SAHA, A., WU, C.L., TANG, D.S. Approximation, dimension reduction, and nonconvex optimization using linear superpositions of Gaussians. *IEEE Transactions on Computers*, vol. 42, no. 10, pp. 1222-1233, 1993.
- SAMAROV, A.M. Exploring Regression Structure Using Nonparametric Functional Estimation. *Journal of the American Statistical Association (JASA)*, vol. 88, no. 423, pp. 836-847, 1993.
- SANGER, T.D. A tree-structured adaptive network for function approximation in high-dimensional spaces. *IEEE Transactions on Neural Networks*, vol. 2, no. 2, pp. 245-256, 1991.
- SCHOENBERG, I.J. Contributions to the problem of approximation of equidistant data by analytic functions, Part A: On the problem of smoothing of graduation, a first class of analytic approximation formulae. *Quart. Appl. Math.*, vol. 4, pp. 45-99, 1946a.
- SCHOENBERG, I.J. Contributions to the problem of approximation of equidistant data by analytic functions, Part B: On the problem of osculatory interpolation, a second class of analytic approximation formulae. *Quart. Appl. Math.*, vol. 4, pp. 112-141, 1946b.
- SCHOENBERG, I.J. Spline functions and the problem of graduation. *Proceedings of the National Academic of Science*, vol. 52, pp. 947-950, 1964.
- SCHUMAKER, L.L. **Spline Functions: Basic Theory**. New York: John Wiley & Sons, 1981.
- SETIONO, R., HUI, L.C.K. Use of a quasi-Newton method in a feedforward neural network construction algorithm. *IEEE Transactions on Neural Networks*, vol. 6, no. 1, pp. 273-277, 1995.
- SHEPP, L.A., KRUSKAL, J.B. Computerized tomography: the new medical x-ray technology. *The American Mathematical Monthly*, vol. 85, no. 6, pp. 420-439, 1978.
- SHIN, Y., GHOSH, J. Ridge polynomial networks. *IEEE Transactions on Neural Networks*, vol. 6, no. 2, pp. 610-622, 1995.
- SIEGELMANN, H.T., SONTAG, E.D. On the computational power of neural nets. *Fifth Workshop on Computational Learning Theory (COLT'92)*, pp. 440-449, 1992.
- SILVERMAN, B.W. Some Aspects of the Spline Smoothing Approach to Non-parametric Regression Curve Fitting (with Discussion). *Journal of the Royal Statistical Society B*, vol. 47, no. 1, pp. 1-52, 1985.
- SJÖBERG, J. **Non-Linear System Identification with Neural Networks**. Ph.D. Thesis, Department of Electrical Engineering, Linköping University, Sweden, 1995.

- SJÖBERG, J., ZHANG, Q., LJUNG, L., BENVENISTE, A., DELYON, B., GLORENNEC, P.-Y., HJALMARSSON, H., JUDITSKY, A. Nonlinear Black-box Modeling in System Identification: a Unified Overview. *Automatica*, vol. 31, no. 12, pp. 1691-1724, 1995.
- SKELTON, R.E. **Dynamic Systems Control**. New York: John Wiley & Sons, 1988.
- SONTAG, E. Some Topics in Neural Networks and Control. *Technical Report LS93-02*, Department of Mathematics, Rutgers University, 1993.
- SPRECHER, D.A. On the structure of continuous functions of several variables. *Transansactions of the American Mathematical Society*, vol. 115, pp. 340-355, 1965.
- SPRECHER, D.A. On the structure of representations of continuous functions of several variables as finite sums of continuous functions of one variable. *Proceedings of the American Mathematical Society*, vol. 17, no. 1, pp. 98-105, 1966.
- SPRECHER, D.A. A Universal Mapping for Kolmogorov's Superposition Theorem. *Neural Networks*, vol. 6, no. 8, pp. 1089-1094, 1993.
- STARK, J. Iterated function systems as neural networks. *Neural Networks*, vol. 4, no. 5, pp. 679-690, 1991.
- STEIN, C. Estimation of the Mean of a Multivariate Normal Distribution. *The Annals of Statistics*, vol. 9, no. 6, pp. 1135-1151, 1981.
- STEWART, I. **Does God Play Dice? The New Mathematics of Chaos**. London: Basil Blackwell, 1989.
- STONE, C.J. Consistent nonparametric regression (with Discussion). *The Annals of Statistics*, vol. 5, no. 4, pp. 595-645, 1977.
- STONE, C.J. Optimal Global Rates of Convergence for Nonparametric Regression. *The Annals of Statistics*, vol. 10, no. 4, pp. 1040-1053, 1982.
- STONE, C.J. Additive Regression and Other Nonparametric Models. *The Annals of Statistics*, vol. 13, no. 2, pp. 689-705, 1985.
- STONE, M.H. The generalized Weierstrass approximation theorem. *Mathematics Magazine*, vol. 21, pp. 167-184, 237-254, 1948.
- STRANG, G. **Linear Algebra and its Applications**. San Diego: Harcourt Brace Jovanovich Publishers, 1988.
- STYBLINSKI, M.A., TANG, T.-S. Experiments in Nonconvex Optimization: Stochastic Approximation with Function Smoothing and Simulated Annealing. *Neural Networks*, vol. 3, no. 4, pp. 467-483, 1990.

- SVARER, C. **Neural Networks for Signal Processing**. Ph.D. Thesis, Technical University of Denmark, 1994.
- SWITZER, P. Numerical Classification. In **Geostatistics**, New York: Plenum, 1970.
- TENG, C.C., WAH, B.W. An automated design system for finding the minimal configuration of a feedforward neural network. *Proceedings of the IEEE International Conference on Neural Networks*, vol. III, pp. 1295-1300, Orlando, Florida, USA, 1994.
- TENORIO, M.F., LEE, W.T. Self-organizing network for optimum supervised learning. *IEEE Transactions on Neural Networks*, vol. 1, no. 1, pp. 100-110, 1990.
- TIAO, G.C., TSAY, R.S. Model specification in multivariate time series. *Journal of the Royal Statistical Society B*, vol. 51, no. 2, pp. 157-213, 1989.
- TIKHONOV, A.N. Solution of incorrectly formulated problems and the regularization method. *Soviet Math. Dokl.*, vol. 4, pp. 1035-1038, 1963.
- TIKHONOV, A.N., ARSENIN, V.Y. **Solutions of Ill-posed Problems**. Washington, D.C.: W. H. Winston, 1977.
- TONG, H. Threshold models in non-linear time series analysis. *Lecture Notes in Statistics*, vol. 21, New York: Springer-Verlag, 1983.
- TUKEY, P.A., TUKEY, J.W. Preparation: prechosen sequences of views. in V. Barnett (ed.) **Interpreting Multivariate Data**. Chichester: John Wiley & Sons, pp. 189-213, 1981.
- TURING, A.M. On computable numbers, with an application to the entscheidungsproblem. *Proceedings of the London Mathematical Society*, vol. 2-42, pp. 230-265, 1936
- VAN DER SMAGT, P.P. Minimisation Methods for Training Feedforward Neural Networks. *Neural Networks*, vol. 7, no. 1, pp. 1-11, 1994.
- VAPNIK, V.N. Principles of Risk Minimization for Learning Theory. in J.E. Moody, S.J. Hanson, R.P. Lippmann (eds.) **Advances in Neural Information Processing Systems 4**, pp. 831-838, San Mateo, CA: Morgan Kaufmann Publishers, 1992.
- VAPNIK, V.N., CHERVONENKIS, A.Y. On the uniform convergence of relative frequencies of events to their probabilities. *Theory of Probability and Its Applications*, vol. 16, pp. 264-280, 1971.
- VERKOOIJEN, W., DANIELS, H. Connectionist projection pursuit regression. *Computational Economics*, vol. 7, pp. 155-161, 1994.
- VIDYASAGAR, M. **Nonlinear Systems Analysis**. New Jersey: Prentice-Hall, Inc., 1993.

- VITUSHKIN, A.G. On representation of functions by means of superpositions and related topics. *L'enseignement Mathématique*, II^e série, Tome XXIII, Fasc. 3-4, pp. 255-320, 1977.
- VON ZUBEN, F.J. **Redes Neurais Aplicadas ao Controle de Máquina de Indução**. Tese de Mestrado, Faculdade de Engenharia Elétrica, Unicamp, 1993.
- VON ZUBEN, F.J., NETTO, M.L.A. Exploring the Nonlinear Dynamic Behavior of Artificial Neural Networks. *Proceedings of the IEEE International Conference on Neural Networks*, vol. II, pp. 1000-1005, Orlando, Florida, USA, 1994a.
- VON ZUBEN, F.J., NETTO, M.L.A. Identificação de Sistemas Dinâmicos Não-lineares Utilizando Redes Neurais Recorrentes e Estudo do Comportamento de Convergência do Processo de Treinamento Associado. *Anais do 10^o Congresso Brasileiro de Automática / 6^o Congresso Latino Americano de Controle Automático*, Rio de Janeiro, RJ, Brasil, vol. 2, p. 1207-1212, 1994b.
- VON ZUBEN, F.J., NETTO, M.L.A. Speed Control of AC Drives Using Neural Networks. *Proceedings of the AMCA/IEEE International Workshop on Neural Networks Applied to Control and Image Processing (NNACIP'94)*, pp. 24-28, November 7-11, Mexico City, 1994c.
- VON ZUBEN, F.J., NETTO, M.L.A. Aprendizado construtivo para redes neurais com uma camada intermediária. *Anais do 2^o Simpósio Brasileiro de Automação Inteligente*, pp. 283-288, Curitiba, PR, 1995a.
- VON ZUBEN, F.J., NETTO, M.L.A. Second-order training for recurrent neural networks without teacher-forcing. *Proceedings of the IEEE International Conference on Neural Networks (ICNN'95)*, vol. 2, pp. 801-806, 1995b.
- VON ZUBEN, F.J., NETTO, M.L.A. Unit-growing learning optimizing the solvability condition for model-free regression. *Proceedings of the IEEE International Conference on Neural Networks (ICNN'95)*, vol. 2, pp. 795-800, 1995c.
- VON ZUBEN, F.J., NETTO, M.L.A. Implementação computacional de algoritmos de treinamento para redes neurais paramétricas e não-paramétricas. *Relatório Técnico RT-DCA 1/96*, Departamento de Engenharia da Computação e Automação Industrial, Faculdade de Engenharia Elétrica, Unicamp, 1996.
- VON ZUBEN, F.J., NETTO, M.L.A., BIM, E., SZAJNER, J. Adaptive Vector Control of a Three-Phase Induction Motor Using Neural Networks, *Proceedings of the IEEE International Conference on Neural Networks*, vol. VI, pp. 3750-3755, Orlando, Florida, USA, 1994a.

- VON ZUBEN, F.J., NETTO, M.L.A., BIM, E., SZAJNER, J. Aplicação de Técnicas Não-lineares baseadas em Redes Neurais no Controle Vetorial Indireto do Motor de Indução Trifásico com Rotor Tipo Gaiola de Esquilo. *Anais do 10º Congresso Brasileiro de Automática / 6º Congresso Latino Americano de Controle Automático*, Rio de Janeiro, RJ, Brasil, vol. 2, p. 930-935, 1994b.
- WAHBA, G. Smoothing Noisy Data with Spline Functions. *Numerische Mathematik*, vol. 24, no. 5, pp. 383-393, 1975.
- WAHBA, G. in Huber, P.J. Projection pursuit. *The Annals of Statistics*, vol. 13, no. 2, pp. 518-521, 1985.
- WAHBA, G. **Spline Models for Observational Data**. Society for Industrial and Applied Mathematics, 1990.
- WAHBA, G., WOLD, S. A completely automatic French curve: Fitting spline functions by cross-validation. *Communications in Statistics*, vol. 4, no. 1, pp. 1-17, 1975a.
- WAHBA, G., WOLD, S. Periodic splines for spectral density estimation: The use of cross-validation for determining the degree of smoothing. *Communications in Statistics*, vol. 4, no. 2, pp. 125-141, 1975b.
- WEIGEND, A.S., HUBERMAN, B.A., RUMELHART, D.E. Predicting the future: A connectionist approach. *International Journal of Neural Systems*, vol. 1, pp. 193-209, 1990.
- WEIGEND, A.S., RUMELHART, D.E., HUBERMAN, B.A. Generalization by Weight-Elimination with Application to Forecasting. in R.P. Lippmann, J.E. Moody, D.S. Touretzky (eds.) **Advances in Neural Information Processing Systems 3**, pp. 875-882, San Mateo, CA: Morgan Kaufmann Publishers, 1991.
- WEISS, G. Neural networks and evolutionary computation. Part 1: Hybrid approaches in artificial intelligence. *Proceedings of the IEEE International Conference on Evolutionary Computation*, vol. I, pp. 268-272, Orlando, Florida, USA, 1994.
- WERBOS, P.J. **Beyond regression: New tools for prediction and analysis in the behavioral sciences**. PhD thesis, Harvard University, Cambridge, Massachusetts, 1974.
- WERBOS, P.J. Backpropagation Through Time: What It Does and How to Do It. *Proceedings of the IEEE*, vol. 78, no. 10, pp. 1550-1560, 1990.
- WHITE, H. Learning in Artificial Neural Networks: A Statistical Perspective. *Neural Computation*, vol. 1, no. 4, pp. 425-464, 1989.

- WIDROW, B., STEARNS, S.D. **Adaptive Signal Processing**. Englewood Cliffs (NJ): Prentice-Hall, 1985.
- WIENER, N. **Cybernetics**. Cambridge: M.I.T. Press, 1948.
- WILLIAMS, R.J., ZIPSER, D. A Learning Algorithm for Continually Running Fully Recurrent Neural Networks. *Neural Computation*, vol. 1, no. 2, pp. 270-280, 1989a.
- WILLIAMS, R.J., ZIPSER, D. Experimental analysis of the real-time recurrent learning algorithm. *Connection Science*, vol. 1, no. 1, pp. 87-111, 1989b.
- YAO, X. A review of evolutionary artificial neural networks. *International Journal of Intelligent Systems*, vol. 8, pp. 539-567, 1993.
- ZBIKOWSKI, R.W. **Recurrent Neural Networks: Some Control Aspects**. Ph.D. Thesis, Faculty of Engineering, Glasgow University, 1994.
- ZHANG, B.T. An incremental learning algorithm that optimizes network size and sample size in one trial. *Proceedings of the IEEE International Conference on Neural Networks*, vol. I, pp. 215-220, Orlando, Florida, USA, 1994.
- ZHAO, Y., ATKESON, C.G. Some Approximation Properties of Projection Pursuit Learning Networks. in J.E. Moody, S.J. Hanson, R.P. Lippmann (eds.) **Advances in Neural Information Processing Systems 4**, pp. 936-943, San Mateo, CA: Morgan Kaufmann Publishers, 1992.