

第3章 前馈神经网络

人工神经网络（Artificial Neural Network, ANN）是指一系列受生物学和神经科学启发的数学模型。这些模型主要是通过对人脑的神经元网络进行抽象，构建人工神经元，并按照一定的拓扑结构来建立人工神经元之间的连接，进而模拟生物神经网络。在人工智能领域，人工神经网络也常常简称为神经网络（Neural Network, NN）或神经模型（neural model）。

神经网络最早是作为一种主要的连接主义模型。20 世纪 80 年代中后期，最流行的一种连接主义模型是分布式并行处理（Parallel Distributed Processing, PDP）模型，其有如下 3 个主要特性。

- （1）信息表示是分布式的（非局部的）。
- （2）记忆和知识是存储在单元之间的连接上。
- （3）通过逐渐改变单元之间的连接强度来学习新的知识。

连接主义的神经网络有着多种多样的网络结构以及学习方法，虽然早期模型强调模型的生物学合理性（biological plausibility），但后期更关注对某种特定认知能力的模拟，如物体识别、语言理解等。尤其在引入误差反向传播来改进其学习能力之后，神经网络也越来越多地应用在各种机器学习任务上。随着训练数据的增多以及（并行）计算能力的增强，神经网络在很多机器学习任务上已经取得了很大的突破，特别是在语音、图像等感知信号的处理上，神经网络表现出了卓越的学习能力。

在本章中，主要关注采用误差反向传播来进行学习的神经网络，即作为一种机器学习模型的神经网络。从机器学习的角度来看，神经网络一般可以看作一个非线性模型，其基本组成单元为具有非线性激活函数的神经元，通过大量神经元之间的连接，使神经网络成为一种高度非线性的模型。神经元之间的连接权重就是需要学习的参数，可以在机器学习的框架下通过梯度下降方法来进行学习。

3.1 神 经 元

人工神经元（artificial neuron），简称神经元（neuron），是构成神经网络的基本单元，主要是模拟生物神经元的结构和特性，接收一组输入信号并产生输出。生物学家在 20 世纪初就发现了生物神经元的结构。一个生物神经元通常具有多个树突和一条轴突。树突用来接收信息，轴突用来发送信息。当神经元所获得的输入信号的积累超过某个阈值时，它就处于兴奋状态，产生电脉冲。轴突尾端有许多末梢可以给其他神经元的树突产生连接（突



触)，并将电脉冲信号传递给其他神经元。1943 年，心理学家 McCulloch 和数学家 Pitts 根据生物神经元的结构，提出了一种非常简单的神经元模型——MP 神经元。现代神经网络中的神经元和 MP 神经元的结构并无太多变化。不同的是，MP 神经元中的激活函数 f 为 0 或 1 的阶跃函数，而现代神经元中的激活函数通常要求是连续可导的函数。假设一个神经元接收 D 个输入 $x_1, x_2, \dots, x_D$ ，令向量 $\mathbf{x} = [x_1, x_2, \dots, x_D]$ 来表示这组输入，并用净输入（net input） $\mathbf{z} \in \mathbb{R}$ 表示一个神经元所获得的输入信号 $\mathbf{x}$ 的加权和。

$$\begin{aligned} z &= \sum_{d=1}^D w_d x_d + b \\ &= \mathbf{w}^T \mathbf{x} + b \end{aligned} \quad (3-1)$$

其中， $\mathbf{w} = [w_1, w_2, \dots, w_D] \in \mathbb{R}^D$ 是 D 维的权重向量， $b \in \mathbb{R}$ 是偏置。净输入 $\mathbf{z}$ 在经过一个非线性函数 $f(\cdot)$ 后，得到神经元的活性值（activation） $\mathbf{a}$ ：

$$\mathbf{a} = f(\mathbf{z}) \quad (3-2)$$

其中，非线性函数 $f(\cdot)$ 称为激活函数（activation function）。

图 3-1 给出了一个典型的神经元结构示例。

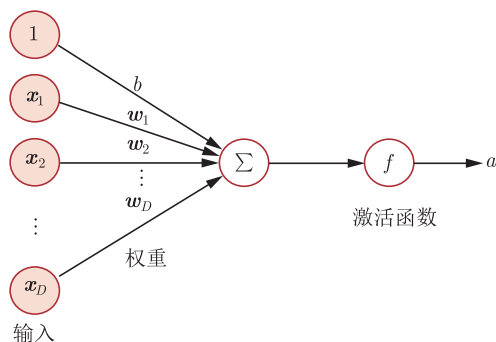


图 3-1 典型的神经元结构示例

激活函数在神经元中是非常重要的。为了增强网络的表示能力和学习能力，激活函数需要具备以下几点特性。

(1) 连续并可导（允许少数点上不可导）的非线性函数。可导的激活函数可以直接利用数值优化的方法来学习网络参数。

(2) 激活函数及其导函数要尽可能简单，有利于提高网络计算效率。

(3) 激活函数的导函数的值域要在一个合适的区间内，不能太大也不能太小，否则会影响训练的效率和稳定性。

下面介绍几种在神经网络中常用的激活函数。

3.1.1 Sigmoid 型函数

Sigmoid 型函数是指一类 S 型曲线函数，为两端饱和函数。常用的 Sigmoid 型函数有 Logistic 函数和 Tanh 函数。



1. Logistic 函数

Logistic 函数的定义为

$$\sigma(x) = \frac{1}{1 + \exp(-x)} \quad (3-3)$$

Logistic 函数可以看成是一个“挤压”函数，把一个实数域的输入“挤压”到 $(0, 1)$ 。当输入值在 0 附近时，Sigmoid 型函数近似为线性函数；当输入值靠近两端时，对输入进行抑制。输入越小，越接近于 0；输入越大，越接近于 1。这样的特点也和生物神经元类似，对一些输入会产生兴奋（输出为 1），对另一些输入产生抑制（输出为 0）。和感知器使用的阶跃激活函数相比，Logistic 函数是连续可导的，其数学性质更好。

因为 Logistic 函数的性质，使得装备了 Logistic 激活函数的神经元具有如下两点性质。

- (1) 其输出直接可以看作概率分布，使神经网络可以更好地和统计学习模型进行结合。
- (2) 可以看作一个软性门（soft gate），用来控制其他神经元输出信息的数量。

2. Tanh 函数

Tanh 函数也是一种 Sigmoid 型函数，其定义为

$$\tanh(x) = \frac{\exp(x) - \exp(-x)}{\exp(x) + \exp(-x)} \quad (3-4)$$

Tanh 函数可以看作放大并平移的 Logistic 函数，其值域是 $(-1, 1)$ 。

$$\tanh(x) = 2\sigma(2x) - 1 \quad (3-5)$$

图 3-2 给出了 Logistic 函数和 Tanh 函数的形状。Tanh 函数的输出是零中心化的（zero-centered），而 Logistic 函数的输出恒大于 0。非零中心化的输出会使其后一层的神经元的输入发生偏置偏移（bias shift），并进一步使得梯度下降的收敛速度变慢。

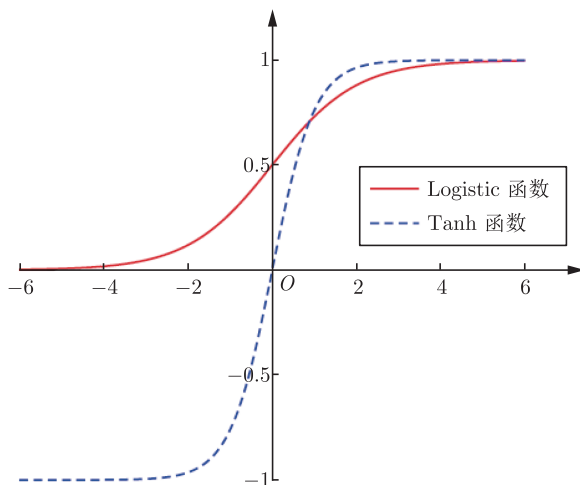


图 3-2 Logistic 函数和 Tanh 函数



3. Hard-Logistic 函数和 Hard-Tanh 函数

Logistic 函数和 Tanh 函数都是 Sigmoid 型函数，具有饱和性，但是计算开销较大。因为这两个函数都是在中间（0 附近）近似线性，两端饱和。因此，这两个函数可以通过分段函数来近似。

以 Logistic 函数 $\sigma(x)$ 为例，其导数为 $\sigma'(x) = \sigma(x)(1 - \sigma(x))$ 。Logistic 函数在 0 附近的一阶泰勒展开（Taylor expansion）为

$$\begin{aligned} g_l(x) &\approx \sigma(0) + x \times \sigma'(0) \\ &= 0.25x + 0.5 \end{aligned} \quad (3-6)$$

这样 Logistic 函数可以用分段函数 $\text{hard-logistic}(x)$ 来近似。

$$\begin{aligned} \text{hard-logistic}(x) &= \begin{cases} 1 & g_l(x) \geq 1 \\ g_l & 0 < g_l(x) < 1 \\ 0 & g_l(x) \leq 0 \end{cases} \\ &= \max(\min(g_l(x), 1), 0) \\ &= \max(\min(0.25x + 0.5, 1), 0) \end{aligned} \quad (3-7)$$

同样，Tanh 函数在 0 附近的一阶泰勒展开为

$$\begin{aligned} g_t(x) &\approx \tanh(0) + x \times \tanh'(0) \\ &= x \end{aligned} \quad (3-8)$$

Tanh 函数也可以用分段函数 $\text{hard-tanh}(x)$ 来近似。

$$\begin{aligned} \text{hard-tanh}(x) &= \max(\min(g_t(x), 1), -1) \\ &= \max(\min(x, 1), -1) \end{aligned} \quad (3-9)$$

图 3-3 给出了 Hard-Logistic 函数和 Hard-Tanh 函数的形状。

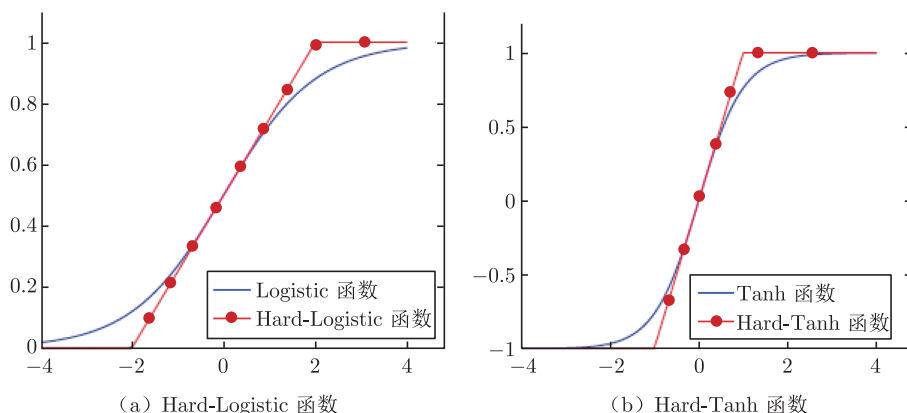


图 3-3 Hard-Sigmoid 型激活函数



3.1.2 ReLU 函数

ReLU (Rectified Linear Unit, 修正线性单元), 也叫作 Rectifier 函数, 是目前深度神经网络中经常使用的激活函数。ReLU 实际上是一个斜坡 (ramp) 函数, 定义为

$$\begin{aligned}\text{ReLU}(x) &= \begin{cases} x & x \geq 0 \\ 0 & x < 0 \end{cases} \\ &= \max(0, x)\end{aligned}\quad (3-10)$$

ReLU 函数的优点如下。① 采用 ReLU 的神经元只需要进行加、乘和比较的操作, 计算上更加高效。② ReLU 函数被认为具有生物学合理性 (biological plausibility), 如单侧抑制、宽兴奋边界 (即兴奋程度可以非常高)。在生物神经网络中, 同时处于兴奋状态的神经元非常稀疏。人脑中在同一时刻大概只有 1%~4% 的神经元处于活跃状态。Sigmoid 型激活函数会导致一个非稀疏的神经网络, 而 ReLU 却具有很好的稀疏性, 大约 50% 的神经元会处于激活状态。③ 在优化方面, 相比于 Sigmoid 型函数的两端饱和, ReLU 函数为左饱和函数, 且在 $x > 0$ 时导数为 1, 在一定程度上缓解了神经网络的梯度消失问题, 加速梯度下降的收敛速度。

ReLU 函数的缺点如下。① ReLU 函数的输出是非零中心化的, 给后一层的神经网络引入偏置偏移, 会影响梯度下降的效率。② ReLU 神经元在训练时比较容易“死亡”。在训练时, 如果参数在一次不恰当的更新后, 第一个隐藏层中的某个 ReLU 神经元在所有的训练数据上都不能被激活, 那么这个神经元自身参数的梯度永远都是 0, 在以后的训练过程中永远不能被激活。这种现象称为死亡 ReLU 问题 (dying ReLU problem), 并且有可能会发生在其他隐藏层。

在实际使用中, 为了避免上述缺点, 有几种 ReLU 的变种可能会被广泛使用。

1. 带泄露的 ReLU

带泄露的 ReLU (leaky ReLU) 在输入 $x < 0$ 时, 保持一个很小的梯度 γ 。当神经元非激活时也能有一个非零的梯度可以更新参数, 避免永远不能被激活。带泄露的 ReLU 的定义如下:

$$\begin{aligned}\text{LeakyReLU}(x) &= \begin{cases} x & x > 0 \\ \gamma x & x \leq 0 \end{cases} \\ &= \max(0, x) + \gamma \min(0, x)\end{aligned}\quad (3-11)$$

其中, γ 是一个很小的常数, 如 0.01。当 $\gamma < 1$ 时, 带泄露的 ReLU 可以写为

$$\text{LeakyReLU}(x) = \max(x, \gamma x) \quad (3-12)$$

这相当于一个比较简单的 maxout 单元。



2. 带参数的 ReLU

带参数的 ReLU (Parametric ReLU, PReLU) 引入一个可学习的参数, 不同神经元可以有不同的参数。对于第 i 个神经元, 其 PReLU 的定义为

$$\begin{aligned} \text{PReLU}_i(x) &= \begin{cases} x & x > 0 \\ \gamma_i x & x \leq 0 \end{cases} \\ &= \max(0, x) + \gamma_i \min(0, x) \end{aligned} \quad (3-13)$$

其中, γ 为 $x \leq 0$ 时函数的斜率。因此, PReLU 是非饱和函数。如果 $\gamma = 0$, 那么 PReLU 就退化为 ReLU。如果 γ 为一个很小的常数, 则 PReLU 可以看作带泄露的 ReLU。PReLU 可以允许不同神经元具有不同的参数, 也可以一组神经元共享一个参数。

3. ELU 函数

ELU (Exponential Linear Unit, 指数线性单元) 是一个近似的零中心化的非线性函数, 其定义为

$$\begin{aligned} \text{ELU}(x) &= \begin{cases} x & x > 0 \\ \gamma(\exp(x) - 1) & x \leq 0 \end{cases} \\ &= \max(0, x) + \min(0, \gamma(\exp(x) - 1)) \end{aligned} \quad (3-14)$$

其中, $\gamma \geq 0$ 是一个超参数, 决定 $x \leq 0$ 时的饱和曲线, 并调整输出均值在 0 附近。

4. Softplus 函数

Softplus 函数可以看作 Rectifier 函数的平滑版本, 其定义为

$$\text{Softplus}(x) = \log(1 + \exp(x)) \quad (3-15)$$

Softplus 函数的导数恰好是 Logistic 函数。Softplus 函数虽然也具有单侧抑制、宽兴奋边界的特性, 却没有稀疏激活性。

图 3-4 给出了 ReLU、Leaky ReLU、ELU 以及 Softplus 函数的示例。

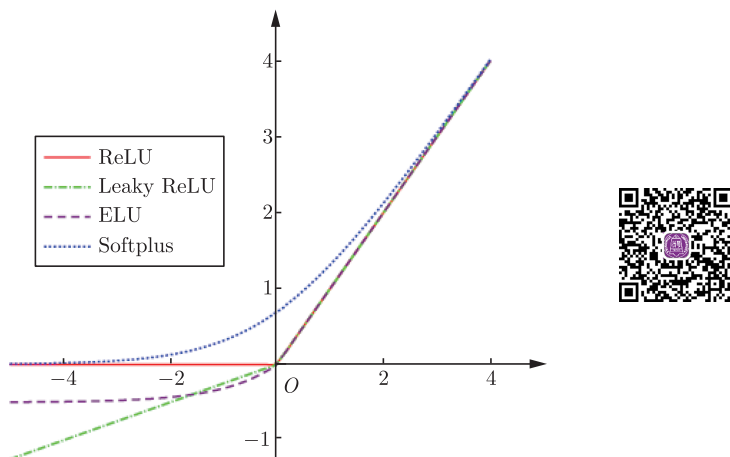


图 3-4 ReLU、Leaky ReLU、ELU 以及 Softplus 函数示例



3.1.3 Swish 函数

Swish 函数是一种自门控 (self-gated) 激活函数, 定义为

$$\text{swish}(x) = x\sigma(\beta x) \quad (3-16)$$

其中, $\sigma(\cdot)$ 为 Logistic 函数, β 为可学习的参数或一个固定超参数。 $\sigma(\cdot) \in (0, 1)$ 可以看作一种软性的门控机制。当 $\sigma(\beta x)$ 接近于 1 时, 门处于“开”状态, 激活函数的输出近似于 x 本身; 当 $\sigma(\beta x)$ 接近于 0 时, 门的状态为“关”, 激活函数的输出近似于 0。

图 3-5 给出了 Swish 函数的示例。当 $\beta = 0$ 时, Swish 函数变成线性函数 $x/2$ 。当 $\beta = 1$ 时, Swish 函数在 $x > 0$ 时近似线性, 在 $x < 0$ 时近似饱和, 同时具有一定的非单调性。当 $\beta \rightarrow +\infty$ 时, $\sigma(\beta x)$ 趋向于离散的 0-1 函数, Swish 函数近似为 ReLU 函数。因此, Swish 函数可以看作线性函数和 ReLU 函数之间的非线性插值函数, 其程度由参数 β 控制。

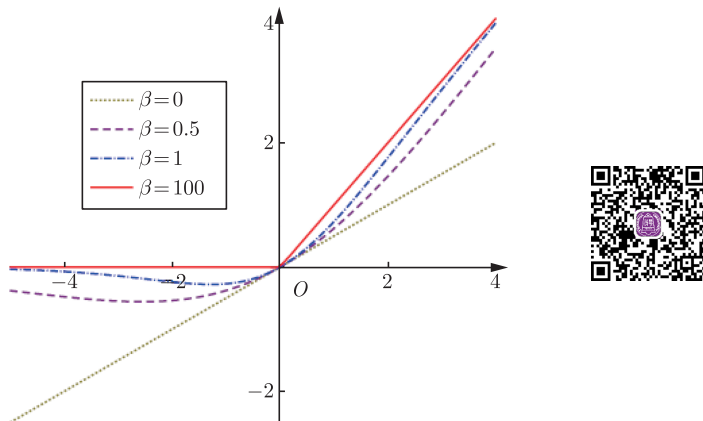


图 3-5 Swish 函数示例

3.1.4 GELU 函数

GELU (Gaussian Error Linear Unit, 高斯误差线性单元) 是一种通过门控机制来调整其输出值的激活函数, 和 Swish 函数比较类似。

$$\text{GELU}(x) = xP(X \leq x) \quad (3-17)$$

其中, $P(X \leq x)$ 是高斯分布 $N(\mu, \sigma^2)$ 的累积分布函数, μ, σ 为超参数, 一般设 $\mu = 0, \sigma = 1$ 即可。由于高斯分布的累积分布函数为 S 型函数, 因此 GELU 函数可以用 Tanh 函数或 Logistic 函数来近似:

$$\text{GELU}(x) \approx 0.5x \left(1 + \tanh \left(\sqrt{\frac{2}{\pi}} (x + 0.044715x^3) \right) \right) \quad \text{或} \quad (3-18)$$

$$\text{GELU}(x) \approx x\sigma(1.702x)$$

当使用 Logistic 函数近似时, GELU 相当于一种特殊的 Swish 函数。



3.1.5 Maxout 单元

Maxout 单元是一种分段线性函数。Sigmoid 型函数、ReLU 等激活函数的输入是神经元的净输入 z ，是一个标量。而 Maxout 单元的输入是上一层神经元的全部原始输出，是一个向量 $\mathbf{x} = [x_1, x_2, \dots, x_D]$ 。

每个 Maxout 单元有 K 个权重向量 $\mathbf{w}_k \in \mathbb{R}^D$ 和偏置 $b_k (1 \leq k \leq K)$ 。对于输入 $\mathbf{x}$ ，可以得到 K 个净输入, $1 \leq k \leq K$ 。

$$z_k = \mathbf{w}_k^T \mathbf{x} + b_k \quad (3-19)$$

其中, $\mathbf{w}_k = [w_{k,1}, w_{k,2}, \dots, w_{k,D}]^T$ 为第 k 个权重向量。Maxout 单元的非线性函数定义为

$$\text{maxout}(\mathbf{x}) = \max(z_k) \quad (3-20)$$

Maxout 单元不单是净输入到输出之间的非线性映射，而且是整体学习输入到输出之间的非线性映射关系。Maxout 激活函数可以看作任意凸函数的分段线性近似，并且在有限的点上是不可微的。

3.2 网络结构

生物神经细胞的功能比较简单，而人工神经元只是生物神经细胞的理想化和简单实现，功能更加简单。要想模拟人脑的能力，单一的神经元是远远不够的，需要通过很多神经元一起协作来实现复杂的功能。这样通过一定的连接方式或信息传递方式进行协作的神经元可以看作一个网络，就是神经网络。到目前为止，研究者已经发明了各种各样的神经网络结构，常用的神经网络结构有如下 3 种。

3.2.1 前馈网络

前馈网络中各神经元按接收信息的先后分为不同的组，每一组可以看作一个神经层，每一层中的神经元接收前一层神经元的输出，并输出到下一层神经元。整个网络中的信息是朝一个方向传播，没有反向的信息传播，可以用一个有向无环路图表示。前馈网络包括全连接前馈网络和卷积神经网络等。

前馈网络可以看作一个函数，通过简单、非线性函数的多次复合，实现输入空间到输出空间的复杂映射。这种网络结构简单，易于实现。

3.2.2 记忆网络

记忆网络，也称为反馈网络，网络中的神经元不但可以接收其他神经元的信息，也可以接收自己的历史信息。和前馈网络相比，记忆网络中的神经元具有记忆功能，在不同的时刻具有不同的状态。记忆神经网络中的信息传播可以是单向或双向传播，因此，可以用



一个有向循环图或无向图来表示。记忆网络包括循环神经网络、Hopfield 网络、玻尔兹曼机、受限玻尔兹曼机等。

记忆网络可以看作一个程序，具有更强的计算和记忆能力。

为了增强记忆网络的记忆容量，可以引入外部记忆单元和读写机制，用来保存一些网络的中间状态，称为记忆增强神经网络（Memory Augmented Neural Network, MANN），如神经图灵机和记忆网络等。

3.2.3 图网络

前馈网络和记忆网络的输入都可以表示为向量或向量序列。但实际应用中很多数据是图结构的数据，如知识图谱、社交网络、分子（molecular）网络等。前馈网络和记忆网络很难处理图结构的数据。

图网络是定义在图结构数据上的神经网络。图中每个节点都由一个或一组神经元构成。节点之间的连接可以是有向的，也可以是无向的。每个节点可以收到来自相邻节点或自身的信息。图网络是前馈网络和记忆网络的泛化，包含很多不同的实现方式，如图卷积网络（Graph Convolutional Network, GCN）、图注意力网络（Graph Attention Network, GAT）、消息传递神经网络（Message Passing Neural Network, MPNN）等。

图 3-6 给出了前馈网络、记忆网络和图网络的网络结构示例，其中圆形节点表示一个神经元，方形节点表示一组神经元。

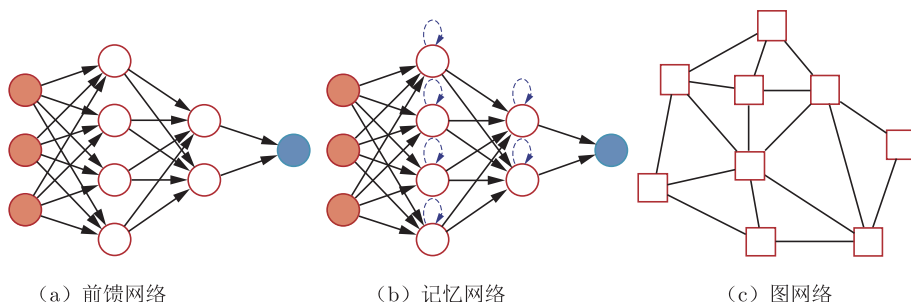


图 3-6 3 种不同的网络结构示例

3.3 前馈神经网络

给定一组神经元，就可以将神经元作为节点来构建一个网络。不同的神经网络模型有着不同网络连接的拓扑结构。一种比较直接的拓扑结构是前馈网络。前馈神经网络（Feed-forward Neural Network, FNN）是最早发明的简单人工神经网络。前馈神经网络经常被称为多层感知器（Multi-Layer Perceptron, MLP）。但多层感知器的叫法并不是十分合理，因为前馈神经网络其实是由多层的 Logistic 回归模型（连续的非线性函数）组成，而不是由多层的感知器（不连续的非线性函数）组成。

在前馈神经网络中，各神经元分别属于不同的层。每一层的神经元可以接收前一层神



经元的信号，并产生信号输出到下一层。第 0 层称为输入层，最后一层称为输出层，其他中间层称为隐藏层。整个网络中无反馈，信号从输入层向输出层单向传播，可用一个有向无环图表示。

图 3-7 给出了多层前馈神经网络的示例。

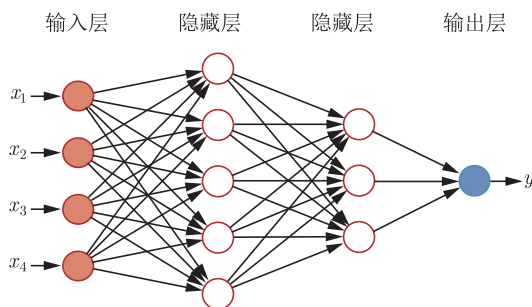


图 3-7 多层前馈神经网络

表 3-1 给出了描述前馈神经网络的记号。

表 3-1 前馈神经网络的记号

记 号	含 义
L	神经网络的层数
M_l	第 l 层神经元的个数
$f_l(\cdot)$	第 l 层神经元的激活函数
$\mathbf{W}^{(l)} \in \mathbb{R}^{M_l \times M_{l-1}}$	第 $l-1$ 层到第 l 层的权重矩阵
$\mathbf{b}^{(l)} \in \mathbb{R}^{M_l}$	第 $l-1$ 层到第 l 层的偏置
$\mathbf{z}^{(l)} \in \mathbb{R}^{M_l}$	第 l 层神经元的净输入（净活性值）
$\mathbf{a}^{(l)} \in \mathbb{R}^{M_l}$	第 l 层神经元的输出（活性值）

令 $\mathbf{a}^{(0)} = \mathbf{x}$ ，前馈神经网络通过不断迭代下面的公式进行信息传播：

$$\mathbf{z}^{(l)} = \mathbf{W}^{(l)} \mathbf{a}^{(l-1)} + \mathbf{b}^{(l)} \quad (3-21)$$

$$\mathbf{a}^{(l)} = f_l(\mathbf{z}^{(l)}) \quad (3-22)$$

首先根据第 $l-1$ 层神经元的活性值（activation） $\mathbf{a}^{(l-1)}$ 计算出第 l 层神经元的净活性值（net activation） $\mathbf{z}^{(l)}$ ，然后经过一个激活函数得到第 l 层神经元的活性值。因此，可以把每个神经层看作一个仿射变换（affine transformation）和一个非线性变换。

式 (3-21) 和式 (3-22) 可以合并写为

$$\mathbf{z}^{(l)} = \mathbf{W}^{(l)} f_{l-1}(\mathbf{z}^{(l-1)}) + \mathbf{b}^{(l)} \quad (3-23)$$

或者

$$\mathbf{a}^{(l)} = f_l(\mathbf{W}^{(l)} \mathbf{a}^{(l-1)} + \mathbf{b}^{(l)}) \quad (3-24)$$



这样，前馈神经网络可以通过逐层的信息传递，得到网络最后的输出 $a^{(L)}$ 。整个网络可以看作一个复合函数 $\phi(\mathbf{x}; \mathbf{W}, b)$ ，将向量 $\mathbf{x}$ 作为第 1 层的输入 $a^{(0)}$ ，将第 L 层的输出 $a^{(L)}$ 作为整个函数的输出。

$$\mathbf{x} = a^{(0)} \rightarrow \mathbf{z}^{(1)} \rightarrow a^{(1)} \rightarrow \mathbf{z}^{(2)} \rightarrow \dots \rightarrow a^{(L-1)} \rightarrow \mathbf{z}^{(L)} \rightarrow a^{(L)} = \phi(\mathbf{x}; \mathbf{W}, b) \quad (3-25)$$

其中， $\mathbf{W}, b$ 表示网络中所有层的连接权重和偏置。

3.3.1 通用近似定理

前馈神经网络具有很强的拟合能力，常见的连续非线性函数都可以用前馈神经网络来近似。

定理 3.1 通用近似定理 (Universal Approximation Theorem): 令 $\phi(\cdot)$ 是一个非常数、有界、单调递增的连续函数， $\mathcal{J}_D$ 是一个 D 维的单位超立方体 $[0, 1]^D$ ， $C(\mathcal{J}_D)$ 是定义在 $\mathcal{J}_D$ 上的连续函数集合。对于任意给定的一个函数 $f \in C(\mathcal{J}_D)$ ，存在一个整数 M ，和一组实数 $V_m, b_m \in \mathbb{R}$ 以及实数向量 $\mathbf{w}_m \in \mathbb{R}^D$ ， $m = 1, 2, \dots, M$ ，可以定义函数

$$F(\mathbf{x}) = \sum_{m=1}^M v_m \phi(\mathbf{w}_m^T \mathbf{x} + b_m) \quad (3-26)$$

作为函数 f 的近似实现，即

$$|F(\mathbf{x}) - f(\mathbf{x})| < \epsilon, \forall \mathbf{x} \in \mathcal{J}_D \quad (3-27)$$

其中， $\epsilon > 0$ 是一个很小的正数。

通用近似定理在实数空间 $\mathbb{R}^D$ 中的有界闭集上依然成立。

根据通用近似定理，对于具有线性输出层和至少一个使用“挤压”性质的激活函数的隐藏层组成的前馈神经网络，只要其隐藏层神经元的数量足够，它可以以任意精度来近似任何一个定义在实数空间 $\mathbb{R}^D$ 中的有界闭集函数。所谓“挤压”性质的函数是指像 Sigmoid 函数的有界函数，但神经网络的通用近似性质被证明对于其他类型的激活函数（如 ReLU）也是适用的。

通用近似定理只是说明了神经网络的计算能力可以去近似一个给定的连续函数，但并没有给出如何找到这样一个网络，以及是否是最优的。此外，当应用到机器学习时，真实的映射函数并不知道，一般是通过经验风险最小化和正则化来进行参数学习。因为神经网络的强大能力，反而容易在训练集上过拟合。

3.3.2 应用到机器学习

根据通用近似定理，神经网络在某种程度上可以作为一个“万能”函数来使用，可以用来进行复杂的特征转换，或逼近一个复杂的条件分布。

在机器学习中，输入样本的特征对分类器的影响很大。以监督学习为例，好的特征可以极大提高分类器的性能。因此，要取得好的分类效果，需要将样本的原始特征向量 $\mathbf{x}$ 转换到更有效的特征向量 $\phi(\mathbf{x})$ ，这个过程叫作特征抽取。



多层前馈神经网络可以看作一个非线性复合函数 $\phi: \mathbb{R}^D \rightarrow \mathbb{R}^{D'}$ ，将输入 $\mathbf{x} \in \mathbb{R}^D$ 映射到输出 $\phi(\mathbf{x}) \in \mathbb{R}^{D'}$ 。因此，多层前馈神经网络可以看成是一种特征转换方法，其输出 $\phi(\mathbf{x})$ 作为分类器的输入进行分类。

给定一个训练样本 $(\mathbf{x})$ ，先利用多层前馈神经网络将 $\mathbf{x}$ 映射到 $\phi(\mathbf{x})$ ，然后再将 $\phi(\mathbf{x})$ 输入到分类器 $g(\cdot)$ ，即

$$\hat{y} = g(\phi(\mathbf{x}); \theta) \quad (3-28)$$

其中， $g(\cdot)$ 为线性或非线性的分类器， θ 为分类器 $g(\cdot)$ 的参数， $\hat{y}$ 为分类器的输出。特别地，如果分类器 $g(\cdot)$ 为 Logistic 回归分类器或 Softmax 回归分类器，那么 $g(\cdot)$ 可以看成是网络的最后一层，即神经网络直接输出不同类别的条件概率 $p(y|\mathbf{x})$ 。

对于二分类问题 $y \in \{0, 1\}$ ，若采用 Logistic 回归，那么 Logistic 回归分类器可以看成神经网络的最后一层。也就是说，网络的最后一层只用一个神经元，并且其激活函数为 Logistic 函数。网络的输出可以直接作为类别 $y = 1$ 的条件概率，即

$$p(y = 1|\mathbf{x}) = \mathbf{a}^{(L)} \quad (3-29)$$

其中， $\mathbf{a}^{(L)} \in \mathbb{R}$ 为第 L 层神经元的活性值。

对于多分类问题 $y \in \{1, 2, \dots, C\}$ ，如果使用 Softmax 回归分类器，相当于网络最后一层设置 C 个神经元，其激活函数为 Softmax 函数。网络最后一层（第 L 层）的输出可以作为每个类的条件概率，即

$$\hat{y} = \text{softmax}(\mathbf{z}^{(L)}) \quad (3-30)$$

其中， $\mathbf{z}^{(L)} \in \mathbb{R}^C$ 为第 L 层神经元的净输入； $\hat{y} \in \mathbb{R}^C$ 为第 L 层神经元的活性值，每一维分别表示不同类别标签的预测条件概率。

3.3.3 参数学习

如果采用交叉熵损失函数，对于样本 $(\mathbf{x}, y)$ ，其损失函数为

$$\mathcal{L}(y, \hat{y}) = -y^T \log \hat{y} \quad (3-31)$$

给定训练集为 $D = \{(\mathbf{x}^{(n)}, y^{(n)})\}_{n=1}^N$ ，将每个样本 $\mathbf{x}^{(n)}$ 输入给前馈神经网络，得到网络输出为 $\hat{y}^{(n)}$ ，其在数据集 D 上的结构化风险函数为

$$R(\mathbf{W}, \mathbf{b}) = \frac{1}{N} \sum_{n=1}^N \mathcal{L}(y^{(n)}, \hat{y}^{(n)}) + \frac{1}{2} \lambda \|\mathbf{W}\|_F^2 \quad (3-32)$$

其中， $\mathbf{W}$ 和 $\mathbf{b}$ 分别表示网络中所有的权重矩阵和偏置向量； $\|\mathbf{W}\|_F^2$ 是正则化项，用来防止过拟合； $\lambda > 0$ 为超参数。 λ 越大， $\mathbf{W}$ 越接近于 0。这里的 $\|\mathbf{W}\|_F^2$ 一般使用 Frobenius 范数：

$$\|\mathbf{W}\|_F^2 = \sum_{l=1}^L \sum_{i=1}^{M_l} \sum_{j=1}^{M_{l-1}} \left(w_{ij}^{(l)}\right)^2 \quad (3-33)$$



有了学习准则和训练样本，网络参数可以通过梯度下降法来进行学习。在梯度下降法的每次迭代中，第 l 层的参数 $\mathbf{W}^{(l)}$ 和 $\mathbf{b}^{(l)}$ 参数更新方式为

$$\begin{aligned}
 \mathbf{W}^{(l)} &\leftarrow \mathbf{W}^{(l)} - \alpha \frac{\partial R(\mathbf{W}, \mathbf{b})}{\partial \mathbf{W}^{(l)}} \\
 &= \mathbf{W}^{(l)} - \alpha \left(\frac{1}{N} \sum_{n=1}^N \left(\frac{\partial \mathcal{L}(y^{(n)}, \hat{y}^{(n)})}{\partial \mathbf{W}^{(l)}} \right) + \lambda \mathbf{W}^{(l)} \right) \\
 \mathbf{b}^{(l)} &\leftarrow \mathbf{b}^{(l)} - \alpha \frac{\partial R(\mathbf{W}, \mathbf{b})}{\partial \mathbf{b}^{(l)}} \\
 &= \mathbf{b}^{(l)} - \alpha \left(\frac{1}{N} \sum_{n=1}^N \frac{\partial \mathcal{L}(y^{(n)}, \hat{y}^{(n)})}{\partial \mathbf{b}^{(l)}} \right)
 \end{aligned} \tag{3-34}$$

其中， α 为学习率。

梯度下降法需要计算损失函数对参数的偏导数，通过链式法则逐一对每个参数进行求偏导比较低效，在神经网络的训练中经常使用反向传播算法来高效地计算梯度。

3.4 反向传播算法

假设采用随机梯度下降进行神经网络参数学习，给定一个样本 $(\mathbf{x}, \mathbf{y})$ ，将其输入神经网络模型中，得到网络输出为 $\hat{\mathbf{y}}$ 。假设损失函数为 $\mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})$ ，要进行参数学习就需要计算损失函数关于每个参数的导数。

对第 l 层中的参数 $\mathbf{W}^{(l)}$ 和 $\mathbf{b}^{(l)}$ 计算偏导数。因为 $\frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{W}^{(l)}}$ 的计算涉及向量对矩阵的微分，十分烦琐，因此先计算 $\mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})$ 关于参数矩阵中每个元素的偏导数 $\frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{w}_{ij}^{(l)}}$ 。根据链式法则：

$$\frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{w}_{ij}^{(l)}} = \frac{\partial z^{(l)}}{\partial \mathbf{w}_{ij}^{(l)}} \frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial z^{(l)}} \tag{3-35}$$

$$\frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{b}^{(l)}} = \frac{\partial z^{(l)}}{\partial \mathbf{b}^{(l)}} \frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial z^{(l)}} \tag{3-36}$$

式 (3-35) 和式 (3-36) 中的第二项都是目标函数关于第 l 层的神经元 $z^{(l)}$ 的偏导数，称为误差项，可以一次计算得到。这样只需要计算 3 个偏导数，分别为 $\frac{\partial z^{(l)}}{\partial \mathbf{w}_{ij}^{(l)}}$ 、 $\frac{\partial z^{(l)}}{\partial \mathbf{b}^{(l)}}$ 和 $\frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial z^{(l)}}$ 。

下面分别计算这 3 个偏导数。

(1) 计算偏导数 $\frac{\partial z^{(l)}}{\partial \mathbf{w}_{ij}^{(l)}}$ 。



因为 $\mathbf{z}^{(l)} = \mathbf{W}^{(l)} \mathbf{a}^{(l-1)} + \mathbf{b}^{(l)}$, 偏导数

$$\begin{aligned} \frac{\partial \mathbf{z}^{(l)}}{\partial \mathbf{w}_{ij}^{(l)}} &= \left[\frac{\partial z_1^{(l)}}{\partial w_{ij}^{(l)}}, \dots, \frac{\partial z_i^{(l)}}{\partial w_{ij}^{(l)}}, \dots, \frac{\partial z_{M_l}^{(l)}}{\partial w_{ij}^{(l)}} \right] \\ &= \left[0, \dots, \frac{\partial \left(\mathbf{w}_{i:}^{(l)} \mathbf{a}^{(l-1)} + b_i^{(l)} \right)}{\partial w_{ij}^{(l)}}, \dots, 0 \right] \\ &= \left[0, \dots, a_j^{(l-1)}, \dots, 0 \right] \\ &\triangleq \mathbb{L}_i \left(a_j^{(l-1)} \right) \in \mathbb{R}^{1 \times M_l} \end{aligned} \quad (3-37)$$

其中, $\mathbf{w}_{i:}^{(l)}$ 为权重矩阵 $\mathbf{W}^{(l)}$ 的第 i 行, $\mathbb{L}_i \left(a_j^{(l-1)} \right)$ 表示第 i 个元素为 $a_j^{(l-1)}$, 其余为 0 的行向量。

(2) 计算偏导数 $\frac{\partial \mathbf{z}^{(l)}}{\partial \mathbf{b}^{(l)}}$ 。

因为 $\mathbf{z}^{(l)}$ 和 $\mathbf{b}^{(l)}$ 的函数关系为 $\mathbf{z}^{(l)} = \mathbf{W}^{(l)} \mathbf{a}^{(l-1)} + \mathbf{b}^{(l)}$, 因此偏导数

$$\frac{\partial \mathbf{z}^{(l)}}{\partial \mathbf{b}^{(l)}} = \mathbf{I}_{M_l} \in \mathbb{R}^{M_l \times M_l} \quad (3-38)$$

为 $M_l \times M_l$ 的单位矩阵。

(3) 计算偏导数 $\frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{z}^{(l)}}$ 。

偏导数 $\frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{z}^{(l)}}$ 表示第 l 层神经元对最终损失的影响, 也反映了最终损失对第 l 层神经元的敏感程度, 因此一般称为第 l 层神经元的误差项, 用 $\delta^{(l)}$ 来表示。

$$\delta^{(l)} \triangleq \frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{z}^{(l)}} \in \mathbb{R}^{M_l} \quad (3-39)$$

误差项 $\delta^{(l)}$ 也间接反映了不同神经元对网络能力的贡献程度, 从而较好地解决了贡献度分配问题 (Credit Assignment Problem, CAP)。

根据 $\mathbf{z}^{(l+1)} = \mathbf{W}^{(l+1)} \mathbf{a}^{(l)} + \mathbf{b}^{(l+1)}$, 有

$$\frac{\partial \mathbf{z}^{(l+1)}}{\partial \mathbf{a}^{(l)}} = \left(\mathbf{W}^{(l+1)} \right)^T \in \mathbb{R}^{M_l \times M_{l+1}} \quad (3-40)$$

根据 $\mathbf{a}^{(l)} = f_l(\mathbf{z}^{(l)})$, 其中 $f_l(\cdot)$ 为按位计算的函数, 因此有

$$\begin{aligned} \frac{\partial \mathbf{a}^{(l)}}{\partial \mathbf{z}^{(l)}} &= \frac{\partial f_l(\mathbf{z}^{(l)})}{\partial \mathbf{z}^{(l)}} \\ &= \text{diag} \left(f_l'(\mathbf{z}^{(l)}) \right) \in \mathbb{R}^{M_l \times M_l} \end{aligned} \quad (3-41)$$



因此, 根据链式法则, 第 1 层的误差项为

$$\begin{aligned}
 \delta^{(l)} &\triangleq \frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{z}^{(l)}} \\
 &= \frac{\partial \mathbf{a}^{(l)}}{\partial \mathbf{z}^{(l)}} \cdot \frac{\partial \mathbf{z}^{(l+1)}}{\partial \mathbf{a}^{(l)}} \cdot \frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{z}^{(l+1)}} \\
 \partial &= \text{diag}(f'_l(\mathbf{z}^{(l)})) \left(\mathbf{W}^{(l+1)} \right)^T \cdot \delta^{(l+1)} \\
 &= f'_l(\mathbf{z}^{(l)}) \odot \left(\left(\mathbf{W}^{(l+1)} \right)^T \delta^{(l+1)} \right) \in \mathbb{R}^{M_l}
 \end{aligned} \tag{3-42}$$

其中, $\odot$ 是向量的点积运算符, 表示每个元素相乘。

从式 (3-42) 可以看出, 第 1 层的误差项可以通过第 $l+1$ 层的误差项计算得到, 这就是误差的反向传播。反向传播算法的含义: 第 l 层的一个神经元的误差项 (或敏感性) 是所有与该神经元相连的第 $l+1$ 层的神经元的误差项的权重和。然后, 再乘以该神经元激活函数的梯度。

在计算出上面 3 个偏导数之后, 式 (3-35) 可以写为

$$\begin{aligned}
 \frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{w}_{ij}^{(l)}} &= \mathbb{L}_i(\mathbf{a}_j^{(l-1)}) \delta^{(l)} \\
 &= \left[0, \dots, \mathbf{a}_j^{(l-1)}, \dots, 0 \right] \left[\delta_1^{(l)}, \dots, \delta_i^{(l)}, \dots, \delta_{M_l}^{(l)} \right]^T \\
 &= \delta_i^{(l)} \mathbf{a}_j^{(l-1)}
 \end{aligned} \tag{3-43}$$

其中, $\delta_i^{(l)} \mathbf{a}_j^{(l-1)}$ 相当于向量 $\delta^{(l)}$ 和向量 $\mathbf{a}^{(l-1)}$ 的外积的第 i, j 个元素。式 (3-43) 可以进一步写为

$$\left[\frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{W}^{(l)}} \right]_{ij} = \left[\delta^{(l)} (\mathbf{a}^{(l-1)})^T \right]_{ij} \tag{3-44}$$

因此, $\mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})$ 关于第 l 层权重 $\mathbf{W}^{(l)}$ 的梯度为

$$\frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{W}^{(l)}} = \delta^{(l)} (\mathbf{a}^{(l-1)})^T \in \mathbb{R}^{M_l \times M_{l-1}} \tag{3-45}$$

同理, $\mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})$ 关于第 l 层偏置 $\mathbf{b}^{(l)}$ 的梯度为

$$\frac{\partial \mathcal{L}(\mathbf{y}, \hat{\mathbf{y}})}{\partial \mathbf{b}^{(l)}} = \delta^{(l)} \in \mathbb{R}^{M_l} \tag{3-46}$$

在计算出每一层的误差项之后, 就可以得到每一层参数的梯度。因此, 使用误差反向传播算法的前馈神经网络训练过程可以分为如下 3 步。

- (1) 前馈计算每一层的净输入 $\mathbf{z}^{(l)}$ 和激活值 $\mathbf{a}^{(l)}$, 直到最后一层。
- (2) 反向传播计算每一层的误差项 $\delta^{(l)}$ 。
- (3) 计算每一层参数的偏导数, 并更新参数。

图 3-8 给出采用反向传播算法的随机梯度下降训练过程。



输入：训练集 $D = \{(\mathbf{x}^{(n)}, \mathbf{y}^{(n)})\}_{n=1}^N$ 验证集 $\mathcal{V}$, 学习率 α , 正则化系数 λ , 网络层数 L , 神经元数量 M_l , $1 \leq l \leq L$ 。

```
1 随机初始化:  $\mathbf{W}, \mathbf{b}$ ;  
2 repeat  
3   对训练集  $D$  中的样本随机重排序;  
4   for  $n=1, 2, \dots, N$  do  
5     从训练集  $D$  中选取样本  $(\mathbf{x}^{(n)}, \mathbf{y}^{(n)})$ ;  
6     前馈计算每一层的净输入  $\mathbf{z}^{(l)}$  和激活值  $\mathbf{a}^{(l)}$ , 直到最后一层;  
7     反向传播计算每一层的误差  $\delta^{(l)}$ ;  
8     // 计算每一层参数的导数  
9      $\forall l, \quad \frac{\partial \mathcal{L}(\mathbf{y}^{(n)}, \hat{\mathbf{y}}^{(n)})}{\partial \mathbf{W}^{(l)}} = \delta^{(l)}(\mathbf{a}^{(l-1)})^T$ ;  
10     $\forall l, \quad \frac{\partial \mathcal{L}(\mathbf{y}^{(n)}, \hat{\mathbf{y}}^{(n)})}{\partial \mathbf{b}^{(l)}} = \delta^{(l)}$ ;  
11    // 更新参数  
12     $\mathbf{W}^{(l)} \leftarrow \mathbf{W}^{(l)} - \alpha(\delta^{(l)}(\mathbf{a}^{(l-1)})^T + \lambda \mathbf{W}^{(l)})$ ;  
13     $\mathbf{b}^{(l)} \leftarrow \mathbf{b}^{(l)} - \alpha \delta^{(l)}$ ;  
14  end  
15 until 神经网络模型在验证集  $\mathcal{V}$  上的错误率不再下降;  
输出:  $\mathbf{W}, \mathbf{b}$ 
```

图 3-8 采用反向传播算法的随机梯度下降训练过程

3.5 自动梯度计算

神经网络的参数主要通过梯度下降来进行优化。当确定了风险函数以及网络结构后,就可以手动用链式法则来计算风险函数对每个参数的梯度,并用代码进行实现。但是手动求导并转换为计算机程序的过程非常琐碎并容易出错,导致实现神经网络变得十分低效。实际上,参数的梯度可以让计算机自动计算。目前,主流的深度学习框架都包含了自动梯度计算的功能,即可以只考虑网络结构并用代码实现,其梯度可以自动进行计算,无须人工干预,这样可以大幅提高开发效率。

自动梯度计算的方法可以分为以下 3 类:数值微分、符号微分和自动微分。

3.5.1 数值微分

数值微分 (numerical differentiation) 是用数值方法来计算函数 $f(x)$ 的导数。函数 $f(x)$ 的点 x 的导数定义为

$$f'(x) = \lim_{\Delta x \rightarrow 0} \frac{f(x + \Delta x) - f(x)}{\Delta x} \quad (3-47)$$

要计算函数 $f(x)$ 在点 x 的导数,可以对 x 加上一个很小的非零的扰动 Δx ,通过上述定义来直接计算函数 $f(x)$ 的梯度。数值微分方法非常容易实现,但找到一个合适的扰动 Δx 却十分困难。如果 Δx 过小,会引起数值计算问题,如舍入误差;如果 Δx 过大,会增加截断误差,使得导数计算不准确。因此,数值微分的实用性比较差。



在实际应用中,经常使用式 (3-48) 来计算梯度,可以减少截断误差。

$$f'(x) = \lim_{\Delta x \rightarrow 0} \frac{f(x + \Delta x) - f(x - \Delta x)}{2\Delta x} \quad (3-48)$$

数值微分的另外一个问题是计算复杂度。假设参数数量为 N , 则每个参数都需要单独施加扰动, 并计算梯度。假设每次正向传播的计算复杂度为 $O(N)$, 则计算数值微分的总体时间复杂度为 $O(N_2)$ 。

3.5.2 符号微分

符号微分 (symbolic differentiation) 是一种基于符号计算的自动求导方法。符号计算也叫作代数计算, 是指用计算机来处理带有变量的数学表达式。这里的变量被看作符号 (symbols), 一般不需要代入具体的值。符号计算的输入和输出都是数学表达式, 一般包括对数学表达式的化简、因式分解、微分、积分、解代数方程、求解常微分方程等运算。如数学表达式的化简:

$$\begin{aligned} \text{输入: } & 3x - x + 2x + 1 \\ \text{输出: } & 4x + 1 \end{aligned} \quad (3-49)$$

符号计算一般来讲是对输入的表达式, 通过迭代或递归使用一些事先定义的规则进行转换。当转换结果不能再继续使用变换规则时, 便停止计算。

符号微分可以在编译时就计算梯度的数学表示, 并进一步利用符号计算方法进行优化。此外, 符号计算的一个优点是符号计算和平台无关, 可以在 CPU 或 GPU 上运行。符号微分的不足之处如下。

- (1) 编译时间较长, 特别是对于循环, 需要很长时间进行编译。
- (2) 为了进行符号微分, 一般需要设计一种专门的语言来表示数学表达式, 并且要对变量 (符号) 进行预先声明。
- (3) 很难对程序进行调试。

3.5.3 自动微分

自动微分 (Automatic Differentiation, AD) 是一种可以对一个 (程序) 函数进行计算导数的方法。符号微分的处理对象是数学表达式, 而自动微分的处理对象是一个函数或一段程序。

自动微分的基本原理是所有的数值计算可以分解为一些基本操作, 包含 $+$ 、 $-$ 、 $\times$ 、 $/$ 和一些初等函数 $\exp$ 、 $\log$ 、 $\sin$ 、 $\cos$ 等, 然后利用链式法则自动计算一个复合函数的梯度。

为简单起见, 这里以一个神经网络中常见的复合函数的例子来说明自动微分的过程。令复合函数 $f(x; w, b)$ 为

$$f(x; w, b) = \frac{1}{\exp(-(wx + b)) + 1} \quad (3-50)$$

其中, x 为输入标量, w 和 b 分别为权重和偏置参数。



首先, 可以将复合函数 $f(x; w, b)$ 分解为一系列的基本操作, 并构成一个计算图 (computational graph)。计算图是数学运算的图形化表示。计算图中的每个非叶节点表示一个基本操作, 每个叶节点为一个输入变量或常量。图 3-9 给出了当 $x = 1, w = 0, b = 0$ 时复合函数 $f(x; w, b)$ 的计算图, 其中连边上的红色数字表示前向计算时复合函数中每个变量的实际取值。

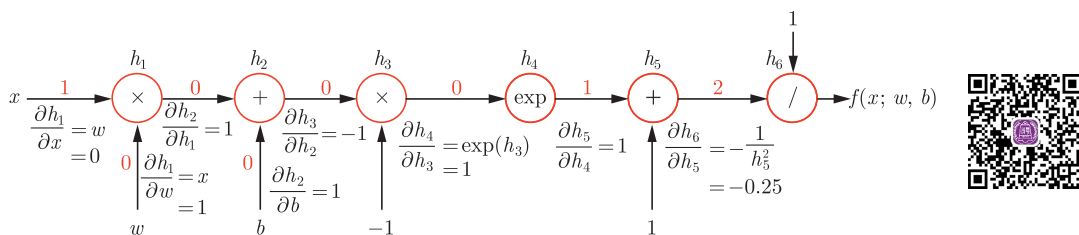


图 3-9 复合函数 $f(x; w, b)$ 的计算图

从图 3-9 上可以看出, 复合函数 $f(x; w, b)$ 由 6 个基本函数 $h_i (1 \leq i \leq 6)$ 组成。如表 3-2 所示, 每个基本函数的导数都十分简单, 可以通过规则来实现。

表 3-2 复合函数 $f(x; w, b)$ 的 6 个基本函数及其导数

函 数	导 数	偏 导 数
$h_1 = x \times w$	$\frac{\partial h_1}{\partial w} = x$	$\frac{\partial h_1}{\partial x} = w$
$h_2 = h_1 + b$	$\frac{\partial h_2}{\partial h_1} = 1$	$\frac{\partial h_2}{\partial b} = 1$
$h_3 = h_2 \times (-1)$	$\frac{\partial h_3}{\partial h_2} = -1$	
$h_4 = \exp(h_3)$	$\frac{\partial h_4}{\partial h_3} = \exp(h_3)$	
$h_5 = h_4 + 1$	$\frac{\partial h_5}{\partial h_4} = 1$	
$h_6 = 1/h_5$	$\frac{\partial h_6}{\partial h_5} = -\frac{1}{h_5^2}$	

整个复合函数 $f(x; w, b)$ 关于参数 w 和 b 的导数可以通过计算图上的节点 $f(x; w, b)$ 与参数 w 和 b 之间路径上所有的导数连乘来得到, 即

$$\frac{\partial f(x; w, b)}{\partial w} = \frac{\partial f(x; w, b)}{\partial h_6} \frac{\partial h_6}{\partial h_5} \frac{\partial h_5}{\partial h_4} \frac{\partial h_4}{\partial h_3} \frac{\partial h_3}{\partial h_2} \frac{\partial h_2}{\partial h_1} \frac{\partial h_1}{\partial w} \quad (3-51)$$

$$\frac{\partial f(x; w, b)}{\partial b} = \frac{\partial f(x; w, b)}{\partial h_6} \frac{\partial h_6}{\partial h_5} \frac{\partial h_5}{\partial h_4} \frac{\partial h_4}{\partial h_3} \frac{\partial h_3}{\partial h_2} \frac{\partial h_2}{\partial b} \quad (3-52)$$

以 $\frac{\partial f(x; w, b)}{\partial w}$ 为例, 当 $x = 1, w = 0, b = 0$ 时, 可以得到



$$\begin{aligned}
 \left. \frac{\partial f(x; w, b)}{\partial w} \right|_{x=1, w=0, b=0} &= \frac{\partial f(x; w, b)}{\partial h_6} \frac{\partial h_6}{\partial h_5} \frac{\partial h_5}{\partial h_4} \frac{\partial h_4}{\partial h_3} \frac{\partial h_3}{\partial h_2} \frac{\partial h_2}{\partial h_1} \frac{\partial h_1}{\partial w} \\
 &= 1 \times (-0.25) \times 1 \times 1 \times (-1) \times 1 \times 1 \\
 &= 0.25
 \end{aligned} \tag{3-53}$$

如果函数和参数之间有多条路径，可以将这多条路径上的导数再进行相加，得到最终的梯度。

按照计算导数的顺序，自动微分可以分为两种模式：前向模式和反向模式。

1. 前向模式

前向模式是按计算图中计算方向的相同方向来递归地计算梯度。以 $\frac{\partial f(x; w, b)}{\partial w}$ 为例，当 $x = 1, w = 0, b = 0$ 时，前向模式的累积计算顺序如下：

$$\frac{\partial h_1}{\partial w} = x = 1 \tag{3-54}$$

$$\frac{\partial h_2}{\partial w} = \frac{\partial h_2}{\partial h_1} \frac{\partial h_1}{\partial w} = 1 \times 1 = 1 \tag{3-55}$$

$$\frac{\partial h_3}{\partial w} = \frac{\partial h_3}{\partial h_2} \frac{\partial h_2}{\partial w} = -1 \times 1 \tag{3-56}$$

$$\frac{\partial h_4}{\partial w} = \frac{\partial h_4}{\partial h_3} \frac{\partial h_3}{\partial w} = 1 \times (-1) \tag{3-57}$$

$$\frac{\partial h_5}{\partial w} = \frac{\partial h_5}{\partial h_4} \frac{\partial h_4}{\partial w} = 1 \times (-1) \tag{3-58}$$

$$\frac{\partial h_6}{\partial w} = \frac{\partial h_6}{\partial h_5} \frac{\partial h_5}{\partial w} = -0.25 \times (-1) = 0.25 \tag{3-59}$$

$$\frac{\partial f(x; w, b)}{\partial w} = \frac{\partial f(x; w, b)}{\partial h_6} \frac{\partial h_6}{\partial w} = 1 \times 0.25 = 0.25 \tag{3-60}$$

2. 反向模式

反向模式是按计算图中计算方向的相反方向来递归地计算梯度。以 $\frac{\partial f(x; w, b)}{\partial w}$ 为例，当 $x = 1, w = 0, b = 0$ 时，反向模式的累积计算顺序如下：

$$\frac{\partial f(x; w, b)}{\partial h_6} = 1 \tag{3-61}$$

$$\frac{\partial f(x; w, b)}{\partial h_5} = \frac{\partial f(x; w, b)}{\partial h_6} \frac{\partial h_6}{\partial h_5} = 1 \times (-0.25) = -0.25 \tag{3-62}$$

$$\frac{\partial f(x; w, b)}{\partial h_4} = \frac{\partial f(x; w, b)}{\partial h_5} \frac{\partial h_5}{\partial h_4} = -0.25 \times 1 = -0.25 \tag{3-63}$$

$$\frac{\partial f(x; w, b)}{\partial h_3} = \frac{\partial f(x; w, b)}{\partial h_4} \frac{\partial h_4}{\partial h_3} = -0.25 \times 1 = -0.25 \tag{3-64}$$

$$\frac{\partial f(x; w, b)}{\partial h_2} = \frac{\partial f(x; w, b)}{\partial h_3} \frac{\partial h_3}{\partial h_2} = -0.25 \times (-1) = 0.25 \tag{3-65}$$



$$\frac{\partial f(x; w, b)}{\partial h_1} = \frac{\partial f(x; w, b)}{\partial h_2} \frac{\partial h_2}{\partial h_1} = 0.25 \times 1 = 0.25 \quad (3-66)$$

$$\frac{\partial f(x; w, b)}{\partial w} = \frac{\partial f(x; w, b)}{\partial h_1} \frac{\partial h_1}{\partial w} = 0.25 \times 1 = 0.25 \quad (3-67)$$

前向模式和反向模式可以看作应用链式法则的两种梯度累积方式。从反向模式的计算顺序可以看出，反向模式和反向传播的计算梯度的方式相同。

对于一般的函数形式 $f: \mathbb{R}^N \rightarrow \mathbb{R}^M$ ，前向模式需要对每一个输入变量都进行一次遍历，共需要 N 次。而反向模式需要对每一个输出都进行一次遍历，共需要 M 次。当 $N > M$ 时，反向模式更高效。在前馈神经网络的参数学习中，风险函数为 $f: \mathbb{R}^N \rightarrow \mathbb{R}^M$ ，输出为标量，因此采用反向模式为最有效的计算方式，只需要一次计算。

3. 静态计算图和动态计算图

计算图按构建方式可以分为静态计算图（static computational graph）和动态计算图（dynamic computational graph）。静态计算图是在编译时构建计算图，计算图构建好之后在程序运行时不能改变，而动态计算图是在程序运行时动态构建。两种构建方式各有优缺点。静态计算图在构建时可以进行优化，并行能力强，但灵活性比较差。动态计算图则不容易优化，当不同输入的网络结构不一致时，难以并行计算，但是灵活性比较高。

4. 符号微分和自动微分

符号微分和自动微分都利用计算图和链式法则自动求解导数。符号微分在编译阶段先构造一个复合函数的计算图，通过符号计算得到导数的表达式，还可以对导数表达式进行优化，在程序运行阶段才代入变量的具体数值来计算导数。而自动微分则无须事先编译，在程序运行阶段边计算边记录计算图，计算图上的局部梯度都直接代入数值进行计算，然后用前向或反向模式来计算最终的梯度。

图 3-10 给出了符号微分与自动微分的对比。



图 3-10 符号微分与自动微分的对比

3.6 优化问题

神经网络的参数学习比线性模型更加困难，主要原因有两点：非凸优化问题和梯度消失问题。