

第3章

航空通信编解码技术

无线通信系统的技术性能主要从通信的数量和质量两方面度量,一般数量指标用有效性度量,质量指标用可靠性度量。前者主要与信源统计特性有关,后者则主要取决于信道的统计特性。

信源发出的信息通过无线信道进行传输。要实现信源信息的高速率、高可靠地通过无线信道传送给信宿,需要解决两个问题:一是在不失真或允许一定失真的条件下,如何用尽可能少的符号传送信源信息,提高通信的有效性;二是在信道存在噪声和干扰的情况下,如何增强信号的抗噪声和抗干扰能力,提高通信的可靠性。这两方面的问题分别是信源编码和信道编码的主要任务。

本章首先介绍信源和信源编码的概念、分类以及常见的信源编码方法,然后介绍信道编码原理、分类及航空无线通信中分组码、卷积码、Turbo 码、LDPC 码等常用的信道编码技术。



视频

3.1 信源与信源编码

3.1.1 信源及其分类

信源就是信息的来源,可以是人、机器、自然界的物体等。实际信源是由最基本的单个消息(符号)组合而成的,可抽象概括为离散信源和连续信源两大类,文字、电报以及各类数据属于离散信源,而未经数字化的话音、图像则属于连续信源。

离散信源又可分为无记忆信源和有记忆信源。当序列信源中的各个消息互相统计独立时称为离散无记忆信源,当序列信源中各个消息前后有关联时称为离散有记忆信源。

若任意两个不同时刻信源发出的符号的概率分布完全相同,则称为离散平稳信源;否则,称为离散非平稳信源。

根据信源符号的统计特性可分为统计特性已知和未知的信源。实际上初始信源的统计特性(符号出现的概率)不可能总是统计好的,所以对于需要知道统计特性的编码方法,编码前应该先估算信源的统计特性。

根据信源的不同分类,实际信源可构成不同的组合信源,常见的组合信源是统计特性已知的离散平稳无记忆信源。

信源输出的平均信息量可定义为信息熵,为单个消息产生的自信息量的概率统计平均值,即

$$H(X) = E\{I[P(x_i)]\} = E[-\log_2 P(x_i)] = - \sum_{i=1}^n P(x_i) \log_2 P(x_i) \quad (3-1)$$

式中: n 为信源产生消息的可能种类数; $P(x_i)$ 为各信源符号出现的概率。

3.1.2 信源编码及其分类

从信息论观点看,实际的信源若经过信息处理会存在大量的统计冗余成分,这些冗余信息完全没有必要通过信道传送给接收端。信源编码的任务是在分析信源统计特

性的基础上设法通过信源的压缩编码去掉这些统计冗余成分,从而提高信息传输的有效性。

信源编码是将信源的原始符号序列按一定的数学规则映射(变换)成码元序列的过程。信源编码的一般模型如图 3-1 所示,输入信源符号 $S = \{s_1, s_2, \dots, s_q\}$,同时存在另一符号 $X = \{x_1, x_2, \dots, x_r\}$,一般元素 x_j 比较适合信道传输,称为码元。编码器是将信源符号集合中的符号 $s_i (i=1, 2, \dots, q)$ 变换成由 $x_j (j=1, 2, \dots, r)$ 组成的长度为 L_i 的一一对应的码符号序列,称为码字。



图 3-1 信源编码器的一般模型

对不同类型的信源,是否存在各自最佳的信源编码,通常由各自的信源编码定理给出。

编码效率是信源编码的重要性能指标,其值越大,则编码的数据压缩能力越强。编码效率定义为

$$\eta = \frac{H(X)}{\bar{L}} \quad (3-2)$$

式中: $H(X)$ 为信源平均信息量(信源熵); $\bar{L}$ 为输出编码的平均码长,且有

$$\bar{L} = - \sum_{i=1}^n P(x_i) l_i \quad (3-3)$$

信源编码可以有以下三种类型:

1. 无失真编码和限失真编码

根据对信源编码的要求是无失真地恢复出原始信源的输出符号还是可允许一定程度上的失真,可将信源编码分为无失真信源编码和限失真信源编码。无失真信源编码能够无失真地恢复信源的数据信息,但大多情况下允许有一定程度的失真。一般离散信源可做到无失真编码,而连续信源则只能采用限失真编码。对连续信源进行编码时,可以通过对信源进行抽样和量化转化为离散信源后再进行编码,所以信源编码主要是对离散信源进行编码。

2. 等长编码和变长编码

根据信源编码输出码长的特点,可将信源编码分为等长编码和变长编码。等长编码对于信源不同的输出符号,码字的长度总是相同的,而变长编码产生的码长不完全相同。变长编码的基本思想是对于给定的信源,当信源符号不是等概率分布时,为了提高编码效率,给概率大的符号分配较短码字,概率小的符号分配较长码字,从而使得平均码长尽可能短。在序列长度 N 不是很大时,变长编码往往可实现高效地无失真信源编码。表 3-1 所示的“编码 1”为等长码,“编码 2”为变长码。

表 3-1 等长码和变长码

信源符号 x_i	符号出现概率 $P(x_i)$	编码 1	编码 2
x_1	$P(x_1)$	00	0
x_2	$P(x_2)$	01	01
x_3	$P(x_3)$	10	001
x_4	$P(x_4)$	11	101

3. 话音编码和数据压缩编码

根据信源信息的业务类型不同,可将信源编码分为话音编码和数据压缩编码。话音编码是把模拟话音信号变成数字话音信号,以便在信道中进行传输的过程,通常可分为波形编码、参量编码和混合编码三类。数据压缩编码是按照一定的编码机制对数据进行重新组织,力求用较少的数据表示信息,以减少数据冗余和存储的空间。根据数据压缩方法的具体原理不同,主要的数据压缩编码技术有预测编码、变换编码、统计编码、模型编码等。



视频

3.2 话音编码

话音是最传统的通信业务,早期的航空通信主要业务是话音。话音编码包含模拟话音信号数字化和对数字化后的信号进行编码两个过程,其目的是在保持一定的算法复杂度和通信时延的前提下,占用尽可能少的通信容量,传送尽可能高质量的话音。

话音编码技术可分为波形编码、参量编码和混合编码三大类。

3.2.1 波形编码

波形编码是将模拟话音信号经过采样、量化、编码后直接变换成数字码流的过程,其基本思想是尽可能保持话音波形不失真。由于在设计上波形编码基本上是与信号源分离的,因此对各种各样的信号进行编码均可以达到良好的效果。波形编码的特点是编码速率高,一般为 $16 \sim 64 \text{ kb/s}$,译码后话音质量较高、时延小。波形编码可以分为时域波形编码和频域波形编码,时域波形编码如 PCM、自适应差分脉冲编码调制(ADPCM)、增量调制(DM)等,频域波形编码如子带编码(Sub-Band Coding, SBC)、自适应变换编码(Adaptive Transform Coding, ATC)等。波形编码在传统的公共电话交换网(Public Switching Telephone Network, PSTN)中应用广泛,但由于需要占用较宽频带,故不适合用于无线通信系统。

1. 时域编码

最常用的时域编码是脉冲编码调制(Pulse Coding Modulation, PCM)编码,其基本原理是将时域的模拟波形信号经过取样、量化、编码而形成数字话音信号。

话音信号的带限特性($300 \sim 3400 \text{ Hz}$)使人们对话音信号进行抽样成为可能。为了保证数字话音信号解码后的高保真度,取样速率应满足奈奎斯特抽样定理,即对于一个频带限制在 $0 \sim f_m$ 内的低通信号,若用 $f_s \geq 2f_m$ 的抽样频率对其进行抽样,则抽样过程中不会丢失信号。在普通话音数字化过程中,话音信号经过滤波,将频带限制在 $300 \sim$

3400Hz 内,用 8kHz 的采样频率对其进行抽样。抽样后的信号在时间上离散了,但幅度值仍然连续。

用有限个电平表示模拟话音信号抽样值的过程称为量化。量化的过程是将信号幅度划分为若干区间,落入同一区间的抽样值量化为同一个幅度值,目的是将波形的幅度值离散化。量化过程中抽样值与量化幅度值之间的差值称为量化误差,量化误差将导致接收端出现量化噪声。为使量化误差尽可能小,量化分层数要大(量化间隔小),但是这将导致编码速率增大,从而使传输带宽增加。实际中,常利用话音信号幅度分布的不均匀性,采用非均匀量化方式(如 13 折线的方法)来折中解决这一问题。

量化后的有限个幅度值用二进制码表示的过程称为编码。编码后的数字信号通过数字通信系统进行传输。

PCM 系统虽然能获得较好的话音质量,但是传输时占用较宽的带宽。在 PCM 的基础上,人们研究出各种改进形式,如增量调制、差分脉码调制(Differential PCM, DPCM)和 ADPCM 等,这些调制方式在“通信原理”课程中已有详细阐述。

2. 频域编码

频域编码又可分为以下两种类型:

(1) 子带编码:在话音信号频域上,利用带通滤波器将话音频带分成若干子带,然后分别进行抽样、量化、编码。因为取样速率降低(相当于若干单边带信号),速率为 9.6~32kb/s,在该范围内话音质量可与同比特率的 ADPCM 质量相当。

(2) 自适应变换编码(Adaptive Transform Coding, ATC):将话音在时间上分段,对每段进行抽样(一般每时段信号有 64~512 个样点),再经数字正交变换转至频域(时域到频域变换),取相应各组频域系数,然后对系数进行量化、编码和传输。接收端则进行相反的处理,以恢复时域信号。这里时域/频域变换一般用离散余弦变换。ATC 速率为 12~16kb/s。在收端将解码得到的每组系数经频域/时域反变换变成时段信号,再将各时段信号连成话音。

3.2.2 参量编码

参量编码是基于人类话音的发声机理,提取表征话音的特征参数,对特征参数进行编码并发送,接收端根据这些参数还原话音的一种方法。参量编码器又称声源编码器或声码器。这种编码器不是跟踪话音信号的波形,而是提取产生话音信号的特征参数。其原理是把人的发音器官看成是一个滤波器,它由来自声带振动的脉冲来激励,滤波器由咽喉、舌头和嘴组成。这个“滤波器”和“激励脉冲串”是不断变化的,但是由于发音器官的延迟特性,可以认为在 10~30ms 发音器官没有变化。因此可以把较短时间段(如 20ms)内相应的滤波器参数指定下来,并提取出这段时间的激励源脉冲串。把滤波器参数和激励源参数一起发送出去,就能代表这段时间(20ms)的话音特性。不同的 20ms 时间段的话音有不同的特征参数。接收端根据收到的话音参数,重建话音。由于接收到的话音是“合成”的话音,有时很难分辨是谁的声音,即话音的自然度下降。

参量编码只传送话音的特征参量,可实现低速率的话音编码,码率一般为 1.2~

4.8kb/s, 占用带宽较小, 较适合于无线通信系统; 缺点是接收端合成的话音虽有一定的可懂度, 自然度却下降很多, 话音质量只能达到中等水平, 时延相对也较大, 不能满足商用话音通信的要求。线性预测编码 (Linear Predictive Coding, LPC) 是一种典型的参量编码。

线性预测编码提供一组话音信号模型参数, 该参数较精确地表征了话音信号的频谱幅度。线性预测由过去的样本值来预测或估计当前信号的结束值, 该值为线性预测值。信号值与线性预测值之差称为线性预测误差。设计一个预测误差滤波器, 使得在某个预定的准则下误差最小, 这个过程就称作线性预测分析 (LPC)。

在人们发声时气流都要通过声管 (喉管、口腔、牙齿、嘴唇等), 声管形状的变化 (口腔大小、伸缩、舌唇位置和形状的变化) 造成不同的气流冲击而形成声音。人类发出的声音包括清音和浊音。发清音时, 声带不振动, 无基音成分 (平坦频谱, 类似白噪声); 发浊音时, 声带振动, 含基音成分 (振动的基本频率)。

滤波器的激励可以是基本频率上的一个脉冲, 也可以是随机白噪声, 这取决于激励信号是清音还是浊音。预测原理与 ADPCM 中的原理相似。LPC 系统传输的只是误差信号 (预测波形与实际波形之间的差别) 中有选择的特征值, 而不是传输误差信号的量化值。

接收端利用收到的预测系数来设计合成滤波器。实际上, 许多线性预测编码器传送的滤波器参数值已经表达了误差信号, 可以直接被接收端合成。

可用图 3-2 的模型表示话音的生成。话音激励源激励一个参数变化的信道来产生话音信号。这里, 以具有一定周期的脉冲源来表示浊音的激励, 以具有平坦的噪声源表示清音的激励, 而声道的变化则以时变参数 (发声随时间变化) 的滤波器来模拟。当不同的激励源加于不同参数的滤波器后, 其输出即为话音信号。

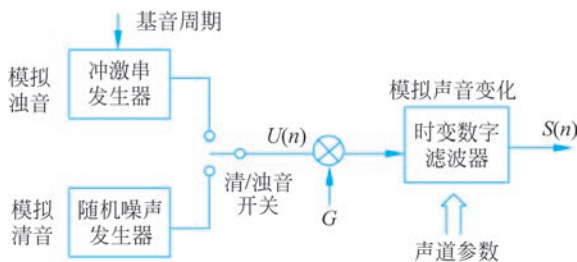


图 3-2 话音生成模型

不同的话音取决于清音和浊音的区分, 清音可用白噪声模拟, 浊音取决于脉冲源的周期和滤波器的参数值。把声音的参数 (如清/浊音判定、浊音周期、滤波器参数) 及时分析出来并编成二进制码进行发送, 收方收到后, 按参数调整收方模型, 即可重建话音。

3.2.3 混合编码

混合编码器是介于波形编码和参量编码之间的一种编码方法, 它综合了波形编码和参量编码的长处, 兼有参量编码低速率和波形编码高质量的优点。实现混合编码的基本

思想是以参量编码原理,特别是以 LPC 原理为基础,保留参量编码低速率的优点,并适当地吸收波形编码中能反映波形个性特征的因素,重点改善话音自然度性能。

混合编码的编码速率一般为 $4\sim 16\text{kb/s}$,话音质量可达到商用话音通信要求的标准。混合编码是无线通信广泛应用的话音编码技术,如规则脉冲激励长时预测线性预测(Regular Pulse Excited-Long Time Prediction, RPE-LTP)编码、码激励线性预测(Code Excited Linear Prediction, CELP)编码、矢量和激励线性预测(Vector Sum Excited Linear Prediction, VSELP)编码等,这些编码方法在现代移动通信领域获得了广泛应用。

3.3 数据压缩编码



视频

现代航空通信,通信业务已扩展到包括话音、数据和图像等在内的多媒体业务,因此信源编码不仅包含话音编码,还包括各类图像压缩编码、多媒体数据压缩编码等方面的内容。多媒体信息所包含的信息量比较大,如彩色图像和视频信息,尤其是视频信息(电视、电影等),在相同条件下要比话音的数据量大 1000 倍以上。将这些数据量大的信息在有限的空间进行存储和传输就需要用到数据压缩技术。

数据压缩编码是按照一定的编码机制对数据进行重新组织,力求用较少的数据来表示信息,以减少数据的冗余和存储的空间。根据解码后数据与原始数据是否完全一致,数据压缩编码技术分为无损压缩和有损压缩两类。无损压缩是指解码的数据与原始数据严格相同,压缩比为 $2:1\sim 5:1$,如霍夫曼编码、算术编码等。有损压缩是指还原的数据与原始数据存在一定的误差,例如对于有损图像压缩,压缩比可以从几倍到上百倍。

根据数据压缩方法的具体原理不同,数据压缩编码技术主要有预测编码、统计编码、变换编码等。

3.3.1 预测编码

预测编码是根据空间中相邻数据存在的相关性,利用前面一个或多个数据来预测下一个数据,然后对实际值和预测值的差值(预测误差)进行编码。常用的预测编码有 DPCM 和 ADPCM。预测编码较适合于声音、图像数据的压缩。

预测编码的基本思想不是直接对信源符号进行编码,而是通过充分利用信源符号间的统计相关性,先将信源输出符号序列进行预测变换,再对信源输出与被预测值的差值进行编码。正是由于信源符号之间存在相关性,所以才能使预测成为可能。如果信源的相关性很强,则采用实际值与预测值的差值进行预测编码可获得较高的压缩率。对于独立信源,预测就没有可能。预测的理论基础就是估计理论,即用实验数据组成一个统计量作为某一物理量的估值或预测值。

预测编码的基本原理如图 3-3 所示,编码端由一个符号编码器和一个预测器组成,译码端由一个符号译码器和一个预测器组成。

一个典型的差值预测编码的例子是 DPCM。DPCM 是采用样本与样本之间存在的信息冗余进行编码的一种数据压缩技术。其基本思想是:根据过去的样本去估算下一个样本信号的幅度大小(这个值称为预测值),然后对实际信号值与预测值之差进行量化编

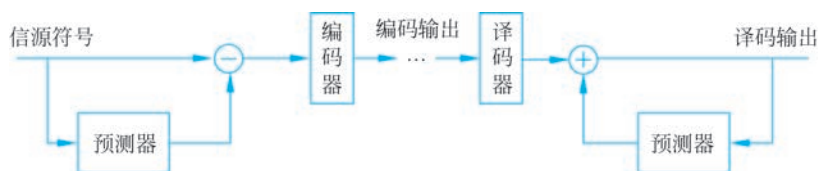


图 3-3 预测编码的基本原理

码,从而减少了表示每个样本信号的位数。

图 3-4 是 DPCM 系统原理框图。图中, $s(n)$ 为发送端输入端抽样信号, $s'_r(n)$ 为接收端重建信号,作为预测器下一个信号估计值的输入信号。发送端 $s_p(n)$ 为预测话音信号, $d(n)$ 为输入信号与预测信号 $s_p(n)$ 的差值即预测误差信号, $d_q(n)$ 为量化后的差值。DPCM 系统实际上就是对这个差值进行量化编码,用来补偿过去编码中产生的量化误差。DPCM 系统是一个反馈系统,采用这种结构可以避免量化误差的积累。 $c(n)$ 为 $d_q(n)$ 经编码后输出的码字。

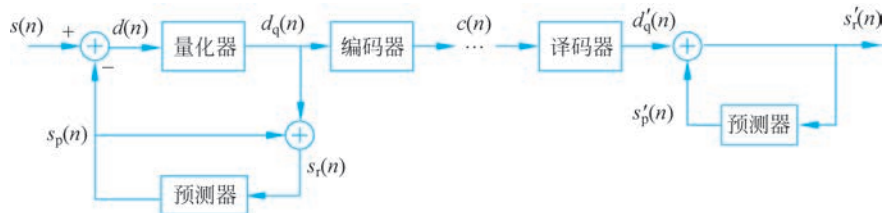


图 3-4 DPCM 系统原理框图

接收端 $d'_q(n)$ 为译码器恢复的差值,译码器中的预测值 $s'_p(n)$ 与发送端编码器中的预测值 $s_p(n)$ 完全相同。因此在无传输误码的情况下,接收端译码器输出的重建信号 $s'_r(n)$ 与编码器的 $s_r(n)$ 完全相同。

自适应差分脉冲编码调制利用自适应的思想改变量化阶的大小,即使用小的量化阶来编码小的差值,使用大的大量化阶来编码大的差值。

增量调制是差分脉码调制的一个重要特例,也称 1 位量化的差值编码。在增量调制系统中,其抽样频率要比 PCM 高得多,称为过抽样。增量调制可根据量化阶距的取值是否为常数,分为线性增量调制(Linear DM, LDM)和自适应增量调制(Adaptive DM, ADM)等。

当量化阶距为常数时,系统估值或者系统恢复值以线性变化的趋势跟踪输入信号,这种线性增量调制的优点是系统简单,容易实现,采用 1 位编码,收发不需要同步;但由于其量化阶距的动态范围很小,其传输质量不高,故发展受到限制。

改进量化阶距动态范围的方法很多,其基本原理是采用自适应方法使量化阶距的大小随输入信号的统计特性变化而跟踪变化。若量化阶距能随信号瞬时压扩,则此增量调制称为自适应增量调制。这是一种自动调节量阶的增量调制,它采用量阶自动调节的办法来适应信号斜率的变化,以避免斜率过载的影响。如果信号斜率增大,则量化阶距也相应增大;输入信号斜率减小,则量化阶距也减小。所以,ADM 具有动态范围大的优点。若量化阶距随音节时间间隔(5~20ms)中信号平均斜率变化,则成为连续可变斜率

增量(Continuous Variable Slope Delta, CVSD)调制。由于这种方法中信号斜率是根据码流中连“1”或连“0”的个数来检测的,所以又称为音节压扩的自适应增量调制,即数字压扩增量调制。

3.3.2 统计编码

统计编码是一种给已知统计信息的符号分配代码的数据无损压缩方法。其基本思想是:根据信息熵原理,将出现概率大的符号用短码来表示,将出现概率小的符号用长码来表示。常用的统计编码方法有香农编码、霍夫曼编码及算术编码等。

1. 香农编码

香农第一定理指出了信源与平均码长的关系,即 $-\log_2 P(x_i) \leq l_i < -\log_2 P(x_i) + 1$,同时也指出了可以通过编码使平均码长达到极限值。按照上述不等式选择的码长构成的码称为香农码。香农编码采用信源符号的累积概率分布函数分配码字。

设某离散信源为

$$\begin{bmatrix} X \\ P(x_i) \end{bmatrix} = \begin{bmatrix} x_1 & x_2 & \cdots & x_n \\ P(x_1) & P(x_2) & \cdots & P(x_n) \end{bmatrix}$$

香农编码的步骤如下:

- (1) 将各信源符号按概率递减的方式进行排列: $P(x_1) \geq P(x_2) \geq \cdots \geq P(x_n)$ 。
- (2) 按香农不等式计算出每个信源符号的码长: $-\log_2 P(x_i) \leq l_i < -\log_2 P(x_i) + 1$ 。
- (3) 计算第 i 个符号的累加概率: $P_i = \sum_{k=1}^{i-1} P(x_k)$ ($i=2, 3, \cdots, n$), $P_1=0$ 。
- (4) 将累加概率用二进制数表示。
- (5) 取 P_i 对应二进制的小数点后 l_i 位作为信源第 i 个符号的二进制码字。

【例 3-1】 设有离散信源的概率分布如下:

$$\begin{bmatrix} X \\ P(X) \end{bmatrix} = \begin{bmatrix} x_1 & x_2 & x_3 & x_4 & x_5 & x_6 & x_7 \\ 0.20 & 0.19 & 0.18 & 0.17 & 0.15 & 0.10 & 0.01 \end{bmatrix}$$

对这一信源进行香农编码,并计算平均码长及编码效率。

解: 香农编码的计算过程见表 3-2。可以得到

信源的熵为

$$H(X) = - \sum_{i=1}^7 P(x_i) \log_2 P(x_i) = 2.16 (\text{比特 / 符号})$$

码的平均长度为

$$\bar{L} = - \sum_{i=1}^7 P(x_i) l_i = 3.14$$

编码效率为

$$\eta = \frac{H(X)}{\bar{L}} = \frac{2.16}{3.14} = 83.1\%$$

表 3-2 香农编码过程

信源符号 x_i	符号概率 $P(x_i)$	累加概率 $P_i = \sum_{k=1}^{i-1} P(x_k)$	每个符号信息量 $-\log_2 P(x_i)$	每个符号 码长	累加概率对应的 二进制码	符号对 应码字
x_1	0.20	0	2.34	3	0.0000000	000
x_2	0.19	0.20	2.41	3	0.0011001	001
x_3	0.18	0.39	2.48	3	0.0110001	011
x_4	0.17	0.57	2.56	3	0.1001000	100
x_5	0.15	0.74	2.74	3	0.1011110	101
x_6	0.10	0.89	3.34	4	0.1110001	1110
x_7	0.01	0.99	6.66	7	0.1111110	1111110

如果一组编码中所有码字都不相同,则称为非奇异码;否则,称为奇异码。非奇异码一定是唯一可译码。可以看出,香农编码所得的码字没有完全相同的,是一种非奇异码。由于香农码的编码效率不高,且冗余度较大,所以它不是最佳码,实用性也受到较大限制。

2. 霍夫曼编码

霍夫曼编码是一种效率较高的变长、无失真信源编码方法。二进制霍夫曼编码的具体步骤如下:

- (1) 将信源消息符号 X 按概率大小自上而下排序。
- (2) 从最小两个概率开始编码,并赋予一定规则,比如上支路为“0”,下支路为“1”,且该规则在整个编码过程中保持不变。
- (3) 将已编码的两个支路概率合并,并重新排序、编码。
- (4) 重复步骤(3),直至合并概率归 1 时为止。
- (5) 从概率归一端沿树图路线逆行至对应消息和概率,并将沿线已编的“0”与“1”编为一组,即为该消息(符号)的编码。例如, x_1 编为“10”, x_7 编为“0111”。

【例 3-2】 设有一个离散简单消息(符号)信源如下:

$$\begin{bmatrix} X \\ P(x_i) \end{bmatrix} = \begin{bmatrix} x_1 & x_2 & x_3 & x_4 & x_5 & x_6 & x_7 \\ 0.20 & 0.19 & 0.18 & 0.17 & 0.15 & 0.10 & 0.01 \end{bmatrix}$$

对它进行霍夫曼编码,并计算平均码长及编码效率。

解: 首先根据信源符号概率大小排队并按图 3-5 的形式进行霍夫曼编码。

信源的熵为

$$H(X) = - \sum_{i=1}^7 P(x_i) \log_2 P(x_i) = 2.16 (\text{比特/符号})$$

码的平均长度为

$$\bar{L} = - \sum_{i=1}^7 P(x_i) l_i = 2.72$$

编码效率为

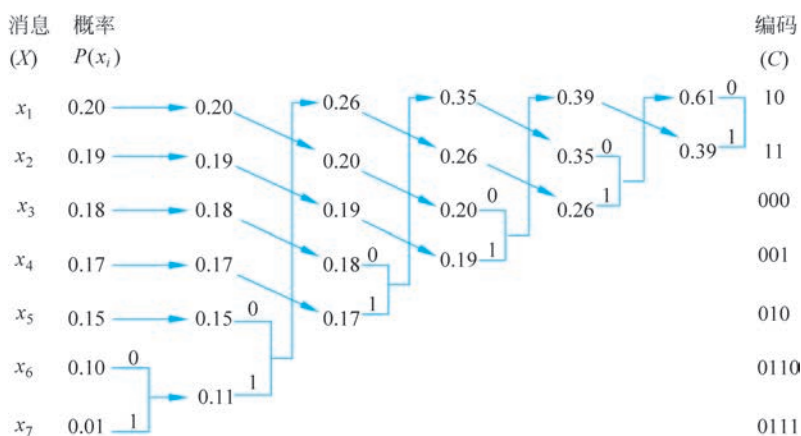


图 3-5 霍夫曼编码

$$\eta = \frac{H(X)}{\bar{L}} = \frac{2.16}{2.72} = 96.3\%$$

与香农编码相比,霍夫曼编码的平均码长更小,信息传输速率更大,编码效率更高。二进制霍夫曼编码可以推广到 m 进制霍夫曼编码,所不同的是每次把 m 个概率最小的符号分别用 $0, 1, \dots, m-1$ 等码元表示,再将它们缩减为信源符号,其余步骤与二进制霍夫曼编码相同。

霍夫曼编码有如下特点:

- (1) 概率大的符号安排在离根节点较近的终端节点,从而保证大概率符号对应于短码,小概率符号对应于长码。
- (2) 每次缩减信源的最长两个码字具有相同码长。
- (3) 每次缩减信源的最后两个码字总是最后一位码元不同,前面各位码元相同。
- (4) 编码不具有唯一性(开始编码时,赋“0”或“1”是任意的;信源缩减时,概率相同的信源放置次序是任意的)。

香农编码与霍夫曼编码之间的区别:两者都考虑了信源的统计特性,使经常出现的信源符号对应较短的码字,进一步缩短信源的平均码长,从而实现了信源的压缩;香农码编码结果唯一,但在很多情况下编码效率不是很高,霍夫曼码的编码方法不唯一;霍夫曼码对信源的统计特性没有特殊要求,编码效率比较高,对编码设备的要求也比较简单,因此综合性能优于香农编码;霍夫曼编码通常适用于多元信源,对于二元信源,必须采用合并符号的方法才能得到较高的编码效率。

霍夫曼编码被称为最优的变长信源编码,但这一最佳性能建立在稳定、确知的概率统计特性基础上,一旦统计特性不稳定或发生变化或不完全确知,变长编码将失去统计匹配的前提,其性能必然引起恶化,实际信源往往不可能提供很稳定、确知的概率特性。于是人们开始研究比较稳健、适应性比较强的准最佳信源编码,算术编码就是最出色的一个编码。算术编码理论上性能虽然比霍夫曼码稍有逊色,但是其实际性能往往优于霍夫曼编码,是发展迅速的一种实用化的无失真离散信源编码。

香农编码和霍夫曼编码主要针对无记忆信源进行编码,当信源有记忆时,编码效率不高。可以将霍夫曼编码进行推广,即将霍夫曼编码中对单个消息(符号)的统计匹配编码推广至信源中 0 序列和 1 序列的消息序列进行统计匹配。具体地说,就是将符号序列中“连 0 串”或“连 1 串”(游程)映射成游程长度和对应符号序列的位置的标志序列,然后对游程序列按霍夫曼编码方法进行编码。这种编码称为游程编码。游程编码效率较高,主要应用于黑、白二值灰度的文件传真及图像编码中。

3. 算术编码

香农编码和霍夫曼编码都是建立在信源符号与码字一一对应的基础上的,这种编码方法通常称为块编码或分组码。算术编码是一种非分组编码方法,是一种从整体符号序列出发,采用递推形式进行编码的方法,信源符号和码字间不再是一一对应的关系。

算术编码的基本思想是:从整体符号序列出发,将各信源序列的概率映射到 $[0,1)$ 区间上,使每个符号序列对应于区间内的一点,也就是一个二进制的小数。这些点把 $[0,1)$ 区间分成许多小段,每段的长度等于某一信源序列的概率,在段内取一个二进制小数,其长度可与该序列的概率匹配,达到高效率编码的目的。这种方法与香农编码有些类似,只是它们考虑的信源对象有所不同,在香农编码中考虑的是单个信源符号,而在算术编码中考虑的是整个信源符号序列。

信源符号集 $A = \{a_1, a_2, \dots, a_n\}$, 信源序列 $\alpha = \{a_{i1}, a_{i2}, \dots, a_{il}, \dots, a_{iL}\}, a_{il} \in A$, 则总共有 n^L 种可能的序列。由于考虑的是整个符号序列,因而整页纸上的信息也许就是一个序列,所以长度 L 一般都很很大。在实际中很难得到对应信源序列的概率,一般从已知的信源符号概率 $P = \{p_1, p_2, \dots, p_n\}$ 递推得到。

定义各符号的累积概率 $P_i = \sum_{l=1}^{i-1} p_l$, 则可得 $P_1 = 0, P_2 = p_1, P_3 = p_1 + p_2 = P_2 + p_2, \dots, P_n = P_{n-1} + p_{n-1}$ 。图 3-6 为算术编码累积概率示意图。

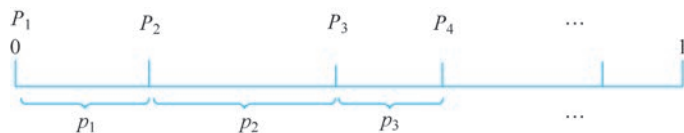


图 3-6 算术编码累积概率示意图

不同的信源符号有不同的概率区间,它们互不重叠,因此可以用这个小区间中的任意一点的取值作为该信源符号的代码。需确定的是这个代码所对应的长度,并使这个长度与信源符号的概率相匹配。对于整个信源符号序列而言,要把一个算术码字赋给它,则必须确定这个算术码字所对应的位于 $[0,1)$ 区间内的实数区间,即由整个信源符号序列的概率本身确定 0 和 1 之间的一个实数区间。随着符号序列中的符号数量的增加,用来代表它的区间会减小,而用来表达区间所需的信息比特数会增大。

每个符号序列中随着符号数量的增加,即信源符号的不断输入,用于代表符号序列概率的区间将随之减小,区间减小的过程如图 3-7 所示。

区间宽度的递推公式如下:

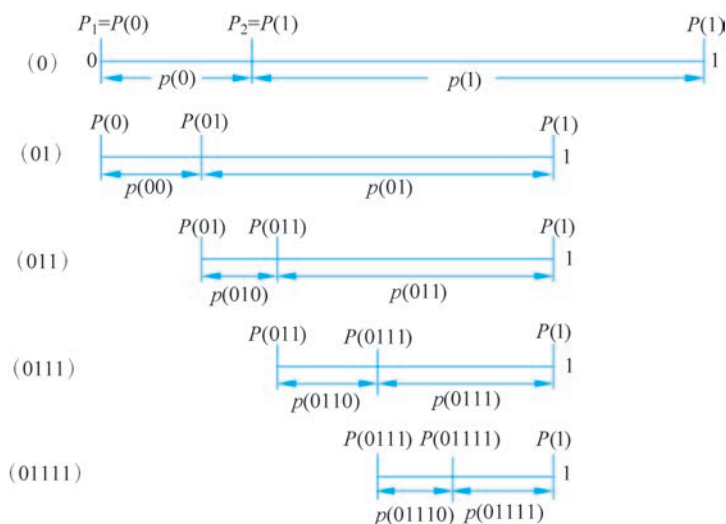


图 3-7 算术编码区间减小的过程示意图

$$A(ar) = A(\alpha)p(r) = p(\alpha)p(r) = p(ar), \quad r=0,1 \quad (3-4)$$

累积概率的递推公式为

$$P(ar) = P(\alpha) + P(\alpha)P(r) = p(ar), \quad r=0,1 \quad (3-5)$$

式中: $P(\alpha)$ 为信源符号序列的累积概率; $p(\alpha)$ 为信源符号序列的联合概率。 $P(0)=0$, $P(1)=p(0)$ 。

取该小区间内的一点代表这个信源符号序列,那么选取此点的方法可以有多种,实际中常取小区间的下界值。对信源符号序列的编码方法也可有多种,下面介绍常用的一种算术编码方法。

将信源符号序列 α 的累积概率值写成二进制小数: $P(\alpha)=0.c_1c_2\cdots c_L, c_i \in \{0,1\}$, 取小数点后 L 位,若后面有尾数,就进位到第 L 位,并使 L 满足

$$L = \left\lceil \log_2 \frac{1}{p(\alpha)} \right\rceil \quad (3-6)$$

式中: $\lceil x \rceil$ 表示大于或等于 x 的最小整数。这样就得到信源符号序列所对应的一个算术编码 $c_1c_2\cdots c_L$ 。

【例 3.3】 设二元无记忆信源 $X=\{0,1\}$, 其中 $p(0)=\frac{1}{4}, p(1)=\frac{3}{4}$, 试对二元序列 $\alpha=11111100$ 进行算术编码。

解: 根据算术编码的方法,先计算信源符号序列 $\alpha=11111100$ 的联合概率

$$p(\alpha) = p(11111100) = p(1)^6 p(0)^2 = (3/4)^6 (1/4)^2 = 0.01112366$$

信源符号序列的算术码字长度为

$$L = \left\lceil \log_2 \frac{1}{p(\alpha)} \right\rceil = \lceil 6.49 \rceil = 7$$

再计算信源符号序列的累积概率。按递推公式有

$$p(\alpha) = p(11111100) = p(0) + p(10) + p(110) + p(1110) + p(11110) + p(111110)$$

$$\begin{aligned}
&= (1/4) + (3/4)(1/4) + (3/4)^2(1/4) + (3/4)^3(1/4) + (3/4)^4(1/4) + \\
&\quad (3/4)^5(1/4) \\
&= 0.82202148 = (0.110100.001)_2
\end{aligned}$$

累积概率值 $p(\alpha)$ 即为输入符号序列“11111100”区间的下界值。取 $p(\alpha)$ 二进制表示小数点后 L 位, 得到信源符号序列的算术码字为 1101010。

平均码长为

$$\bar{L} = \frac{7}{8}$$

编码效率为

$$\eta = \frac{H(X)}{\bar{L}} = \frac{0.811}{7/8} = 92.7\%$$

因为这里需编码的序列长为 8 位, 故一共要将半开区间 $[0, 1)$ 分成 256 个小区间, 以对应任一个可能的序列, 由于任意一个码字必在某个特定的区间, 所以其解码具有唯一性。

从性能上看, 算术编码具有很多优点, 它所需的参数少、编码效率高、编译码简单, 在实际实现时, 常用自适应算术编码对输入的信源序列自适应地估计其概率分布。算术编码在图像压缩标准(如 JPEG)中应用广泛。

3.3.3 变换编码

在许多情况下, 信源输出符号之间具有很强的相关性, 若按照霍夫曼编码、算术编码等编码算法实现高效数据压缩, 则需要使用扩展信源进行编码, 因此需要知道序列间的联合概率分布。由于编码算法本身就比较复杂, 因此当信源符号数量较多时, 如果采用较长的序列进行编码, 就会进一步增加编码系统的复杂度。为了提高编码效率, 同时降低编码复杂度, 可以对信源输出的数据进行变换, 在变换域上解除相关性, 即采用变换编码方式将数据之间的相关冗余变换为系数之间的统计冗余, 从而降低系数编码的复杂度, 提高编码效率。

变换编码的基本原理是将信号从一种信号空间变换到另一种有利于压缩编码的信号空间(如信号由时域变换到频域), 然后进行编码。在时域空间上具有强相关性的数据信息, 反映在频域某些特定的区域内能量常常被集中在一起, 只需将注意力放在相对小的区域上, 从而实现压缩。一般的变换编码的系统框图如图 3-8 所示。变换编码常用于图像的压缩编码, 图像变换域编码通常将空间域相关的像素点通过某种变换关系映射到另一个频域上, 使变换后的系数之间的相关性降低, 在变换后的频域上满足所有系数相互独立, 能量集中于少数几个系数上, 且这些系数集中在一个最小的区域内, 然后对这些变换系数进行量化、编码、传输。在接收端, 则先将收到的信号进行解码操作, 再进行反映射变换, 以再现原始信号。

变换编码与预测编码的区别: 预测编码是在空间域(或时间域)内进行的, 而变换编码则是在变换域(或频率域)内进行的。



图 3-8 一般的变换编码系统框图

变换编码中的关键技术在于正交变换,常用的变换方法有离散傅里叶变换(DFT)、离散余弦变换(Discrete Cosine Transformation, DCT)、沃尔什-哈达玛变换(Walsh-Hadamard Transformation, WHT)和小波变换(Wavelet Transformation, WT)等。不同的变换会有不同的图像压缩效果,在语音、图像编码等应用中,由于离散余弦变换算法简单、有效,因此得到了广泛应用。

3.4 信道编码概述



视频

3.4.1 差错控制方式

信道编码是提高信息传输可靠性的有效手段,它是在信源编码后的数字序列上增加一些冗余信息,增加信息传输的抗干扰能力,以抵抗各种噪声和衰落的影响。在实际信道中传输数字信号时,由于信道特性不理想及加性噪声的影响,接收端接收到的数字信号不可避免地会产生错码,影响通信质量,信道编码通过加入冗余码来减少误码,增加了信息的冗余度。通常冗余度是特定的、有规律的,故可利用其在接收端进行检错和纠错,因此信道编码也称为差错控制编码。

常用的差错控制方式有前向纠错(Forward Error Correction, FEC)、检错重发(Automatic Repeat reQuest, ARQ)、混合纠错(Hybrid Error Correction, HEC)和信息反馈(Information Repeat reQuest, IRQ),如图 3-9 所示。

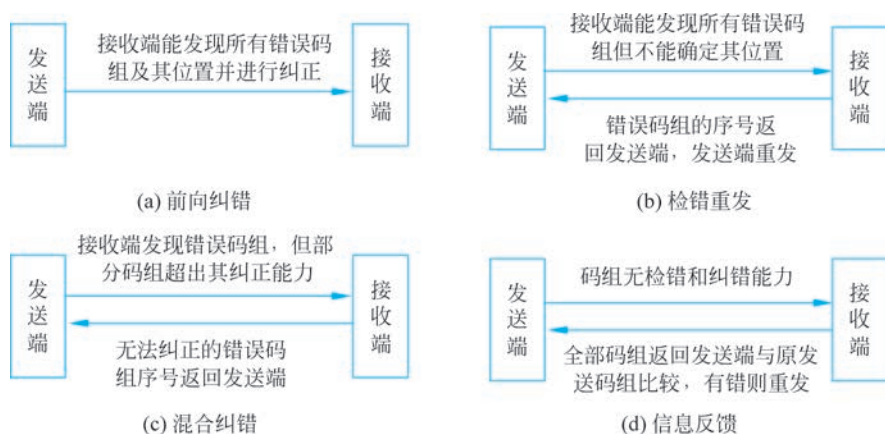


图 3-9 差错控制方式

3.4.2 信道编码基本原理

信道编码是用来改善无线信道通信可靠性的一种信号处理技术,基本思想是在被传送的信息中加入一些多余的码元(监督码元),在收和发之间建立某种检验关系,以达到检验信息码元在传输中是否出现差错的目的。

一般来说,信源发出的信息均可用二进制信号表示。例如,要传送的消息为 A 和 B,可用 0 代表 A,1 代表 B。若传输过程中出错,接收端无法发现。如果在 0 和 1 后面加一位监督位,如 00 代表 A,11 代表 B。经过传输后,接收端接收到码组的组合有 00、01、10、11,00 和 11 是许用码组,01 和 10 是禁用码组。如果接收端收到 01 或 10,译码器就认为是误码,但不能判断是哪一位出错,因而不能自动纠错。如果传输过程中两位码都发生错误,如由代表 B 的码组 11 错成 00,译码器就判为 A,从而造成差错。按这种 1 位信息位加 1 位监督位的编码方式,只能发现 1 位错误。但是,如果按一定数学关系增加监督位,就可能使接收到的码组不但有检错的功能,而且能知道出错的位置,从而加以纠正。纠错编码的实质就是在传输的信息码元之外加入监督码元,使码元之间具有某种相关性。

在一个码组集合中,码组之间的最小码距 $d_{\min}$ 是衡量该码组检错和纠错能力的重要依据,也是编码的重要参数。 $d_{\min}$ 越大,则从一个码字错成另一个码字的可能性就越小,其检纠错能力也就越强。

常用编码效率和编码增益衡量信道编码性能。

1. 编码效率

码字中的信息位数为 k ,监督位数为 r ,码长 $n=k+r$ 。编码效率为信息码位数 k 在码长 n 中所占的比例,表达式为

$$R = \frac{k}{n} = \frac{n-r}{n} = 1 - \frac{r}{n} \quad (3-7)$$

在信道编码过程中,监督位越多,纠错能力就越强,但编码效率就越低。编码效率表示传输信息的有效性,它说明了信道的利用效率,是衡量信道编码性能的主要指标。

2. 编码增益

在数字通信中,信噪比通常用 E_b/n_0 表示,其中 E_b 为信号的比特能量, n_0 为噪声单边功率谱密度。编码增益为

$$G = \left(\frac{E_b}{n_0} \right)_{\text{UC}} - \left(\frac{E_b}{n_0} \right)_{\text{C}} \quad (\text{dB}) \quad (3-8)$$

式中: $\left(\frac{E_b}{n_0} \right)_{\text{UC}}$ 和 $\left(\frac{E_b}{n_0} \right)_{\text{C}}$ 分别为未编码和编码后所需要的信噪比。

3.4.3 信道编码分类

信道编码有多种实现方式,主要的分类如下:

按照信道编码的不同功能可分为检错码和纠错码。检错码仅能检测误码,例如,在

计算机串行通信中常用到的奇偶校验码等；纠错码具有检错能力的同时还能纠正误码，当发现不可纠正的错误时可以发出出错指示。

按照信息码元与监督码元之间的检验关系可分为线性码和非线性码。若信息码元与监督码元之间的关系为线性关系，则称为线性码；反之，则为非线性码。

按照信息码元和监督码元之间的约束方式不同可分为分组码和卷积码。分组码中，编码后的码元序列每 n 位分为一组，其中 k 位信息码元， r 位监督码元， $r = n - k$ ，监督码元仅与本码字的信息码元有关；而在卷积码中，码组中的监督码元不但与本组信息码元有关，而且与前面码组的信息码元有约束关系，就像链条那样一环扣一环，因此卷积码也称连环码或链码。

按照信息码元在编码后是否保持原来的形式可分为系统码和非系统码。系统码中，编码后的信息码元保持不变，而非系统码中的信息码元发生变化。大多情况下，系统码的性能大体上与非系统码相同，但是非系统码的译码较复杂，因此系统码得到了广泛应用。

按照纠正错误的类型不同可分为纠正随机错误码和纠正突发错误码。纠正随机错误码主要用于发生零星独立错误的信道；纠正突发错误码用于以突发错误（成串出现的错误）为主的信道。

按照信道编码所采用的数学方法不同可分为代数码、几何码和算术码。其中代数码是目前发展最为完善的编码，线性码就是代数码的一个重要分支。

下面结合典型航空通信系统的应用，介绍几种常用的信道编码技术。

3.5 线性分组码



视频

3.5.1 线性分组码基本概念

分组码是一组固定长度的码字，可表示为 (n, k) 。在分组码中，监督码元被加到信息码元之后，形成新的码字。在编码时， k 个信息码元被编为 n 位码字长度，而 $r = n - k$ 个监督码元的作用就是实现检错与纠错，当分组码的信息位和监督位之间为线性关系时，这个分组码就称为线性分组码。奇偶校验码就是一种只有一位监督码元的线性分组码。由于线性分组码一般是按照代数规律构造的，故又称代数编码。

线性分组码主要性质：具有封闭性，即码组中任意两个码字的模 2 和仍是该码组中的码字；码组间的最小码距等于非零码的最小汉明重量；全零码必属于线性分组码。

线性分组码的编码规则取决于监督码元和信息码元之间的数学关系（也称监督关系）。对于 (n, k) 线性分组码，有 r 位监督码元就可以对应 r 个监督方程式，如果以矩阵形式表示，就可以得到一个监督矩阵 H 。从 r 个监督方程式中可以得到 r 个校正子，根据校正子可以确定错误图样， r 个校正子可以用来指示 $2^r - 1$ 个错误图样。假设只有 1 位错误，那就可供收端指示 $2^r - 1$ 个错误位置。对于码组 (n, k) ，若希望用 r 个监督位构造出 r 个监督关系式来指示一位错码的 n 种可能，则要求

$$2^r - 1 \geq n \quad \text{或} \quad 2^r \geq k + r + 1 \quad (3-9)$$

典型的线性分组码有奇偶校验码、汉明码、循环码等。

3.5.2 汉明码

纠正 1 位错误的线性分组码称为汉明码。二进制汉明码可以表示为

$$(n, k) = (2^m - 1, 2^m - 1 - m) \quad (3-10)$$

汉明码的特点是：码长 $n = 2^m - 1$ ，信息位数 $k = 2^m - 1 - m$ ，监督位数 $r = m$ ，最小码距 $d_{\min} = 3$ ，纠错能力 $t = 1$ 。如果要产生一个系统的汉明码，可以将矩阵 H 转换成典型监督矩阵，得到相应的生成矩阵 G 。当 r 为 3、4、5 时，线性分组码 $(7, 4)$ ， $(15, 11)$ ， $(31, 26)$ ，…都是汉明码。汉明码的编码效率 $R = \frac{k}{n} = \frac{n-r}{n} = 1 - \frac{r}{n}$ ，当 n 较大时，编码效率接近于 1，所以汉明码是一种高效率的纠错码。

3.5.3 循环码

循环码(Cyclic Redundancy Code, CRC)是一种特殊的线性分组码，具有许多特殊的代数性质。循环码的编码是根据给定的 (n, k) 值选择一个生成多项式 $g(x)$ ，即从 $x^n + 1$ 的因子中选择一个最高次幂为 $r = n - k$ 次多项式作为 $g(x)$ ，根据所有码多项式 $T(x)$ 都可以被 $g(x)$ 整除这一原则，对给定的信息码元进行编码。编码具体步骤如下：

- (1) 根据给定的信息码元，写出信息码多项式 $m(x)$ ；
- (2) 用 x^{n-k} 乘 $m(x)$ ，得到 $x^{n-k}m(x)$ ；
- (3) 用生成多项式 $g(x)$ 除 $x^{n-k}m(x)$ ，得到商和余式，即

$$\frac{x^{n-k}m(x)}{g(x)} = Q(x) + \frac{r(x)}{g(x)}$$

- (4) 编出码组，即联合 $x^{n-k}m(x)$ 和余式 $r(x)$ 得到码多项式

$$T(x) = x^{n-k}m(x) + r(x)$$

例如，信息码为 1011001，则 $m(x) = x^6 + x^4 + x^3 + 1$ ，假设生成多项式为 $g(x) = x^4 + x^3 + 1$ ，则 $x^4m(x) = x^{10} + x^8 + x^7 + x^4$ ，采用多项式除法，得到余式 $r(x) = x^3 + x$ ，从而可知监督码元应为 1010，编出的码组为 1011001 1010。

3.5.4 BCH 码

BCH 码是由 Bose、Chaudhuri 和 Hocquenghem 共同提出，并因此而得名。BCH 码是一种能纠正多位错误的循环码，它可以根据纠错能力的要求直接确定码的构造，其参数可以在大范围内变化，选用灵活，适用性强，是一类广泛应用的差错控制编码。

BCH 码分为本原 BCH 码和非本原 BCH 码两类。最常见的 BCH 码是本原 BCH 码。

本原 BCH 码具有如下特点：

- (1) 码长 $n = 2^m - 1$ ，其中 m 为大于或等于 3 的整数。
- (2) 生成多项式 $g(x)$ 中含有最高次为 m 的本原多项式。

非本原 BCH 码具有如下特点:

(1) 码长 n 是 $2^m - 1$ 的一个因子,其中 m 为大于或等于 3 的整数。

(2) 它的生成多项式 $g(x)$ 中不含有最高次为 m 的本原多项式。

BCH 码可以纠正或检出 t 位错误,并可以提供灵活的参量选择,码长可达到上百位,因此它是目前同样码长和码率的所有分组码中的最优码。然而,求 BCH 码的生成多项式是一项非常繁琐的工作。前人已经将 BCH 码生成多项式的研究结果列成表,通过查表可以得到不同 n, k, t 取值情况下 BCH 码的生成多项式 $g(x)$,工程设计中查表可直接获得。表 3-3 和表 3-4 示例了部分本原 BCH 码和非本原 BCH 码的生成多项式,表中生成多项式系数用八进制数字表示。

表 3-3 部分本原 BCH 码的生成多项式

n	k	t	生成多项式 $g(x)$ 的系数(八进制)
7	4	1	13
15	11	1	23
	7	2	721
	5	3	2467
31	21	2	3551
	16	3	107657
	11	5	5423325
63	51	2	12471
	45	3	1701317
	39	4	166623567
	30	6	157464165547
127	113	2	41567
	106	3	11554743
255	239	2	267543
	231	3	156720665
511	493	2	1132353

表 3-4 部分非本原 BCH 码的生成多项式

n	k	t	生成多项式 $g(x)$ 的系数(八进制)
17	9	2	727
21	12	2	1663
23	12	3	5343
33	22	2	5145
33	12	4	3777
41	21	4	6647133
47	24	5	43073357
65	3	2	10761
65	40	4	354303067
73	46	4	1717773537

例如,表 3-3 中, $n=15, k=11, t=1$ 的 BCH 码生成多项式为 $g(x)=(23)_8$,将八进

制数表示成二进制数得到 $g(x) = (23)_8 = (010011)_2$, 它对应了生成多项式各项的系数, 即 $g(x) = x^4 + x + 1$ 。

表 3-4 中, $(23, 12)$ 是一个特殊的非本原 BCH 码, 称为格雷码。码距为 7, 可以纠正 3 个随机错误, 其生成多项式为 $g(x) = (5343)_8 = (101011100011)_2$, 相应的生成多项式为 $g(x) = x^{11} + x^9 + x^7 + x^6 + x^5 + x + 1$ 。实际应用中, BCH 码的码长都是奇数, 有时为了得到偶数码长, 可将 BCH 码的生成多项式乘以一个因子 $x + 1$, 它相当于在原 BCH 码上增加了一个校验位, 从而得到更强的纠错能力。

3.5.5 RS 码

RS(Reed-Solomon)码是一类具有很强纠错能力的多进制 BCH 码, 能同时纠正随机差错和突发差错, 是 20 世纪以来信道编码技术中应用最为广泛的一种码型。RS 码具有最大的汉明距离, 对于纠正突发错误比较有效, 与其他类型的纠错码相比, 在冗余符号相同的情况下, RS 码的纠错能力最强。RS 码特别适合用于存在突发错误的信道如移动通信衰落信道和多进制调制的场合, 它也被广泛应用于数字电视、数字音频、数字图像等工程实践, 以及卫星通信、深空探测、航空通信等数字通信系统中。在战术数据链 Link-16 中也采用了 RS 纠错编码的方法。

在 (n, k) RS 码中, 输入信号每 kn 位分成一组, 每组包括 k 个符号, 每个符号由 m 位组成。一般 RS 码常用 $m = 8$ 位, 8 位 RS 码具有很大的应用价值。可以纠正 t 个符号错误的 RS 码参量如下:

- (1) 码长: $n = 2^m - 1$ 个符号或 $m(2^m - 1)$ 位。
- (2) 信息段: k 个符号或 km 位。
- (3) 监督段: $n - k = 2t$ 个符号或 $m(n - k)$ 位。
- (4) 最小汉明距离: $d_{\min} = 2t + 1$ 个符号或 $m(2t + 1)$ 位。

例如, $(255, 223)$ RS 码表示码块长度共有 255 个符号, 其中信息段有 223 个符号, 监督段有 32 个检验符号。在这个由 255 个符号组成的码块中, 可以纠正在这个码块中出现的 16 个分散的或者 16 个连续的错误符号。



视频

3.6 卷积码

3.6.1 卷积码基本概念

分组码是把 k 位的序列编成 n 位的码组, 每个码组的 $n - k$ 个校验位仅与本组码的 k 个信息位有关, 而与其他码组无关。为了达到一定的纠错能力和编码效率, 分组码的码组长度一般比较大, 编译码时必须把整个信息码组存储起来, 而由此产生的译码时延随 n 的增大而增加。与分组码不同, 卷积码的 n 位编码, 不仅与当前段的 k 个信息位有关, 还与前面 m 段的信息位有关, 整个编码过程可以看成是输入信息序列与由移位寄存器和模 2 和连接方式所决定的另一个序列的卷积, 卷积码也由此得名。

卷积码常表示为 (n, k, m) , n 表示卷积码编码器输出端码元数, k 表示编码器输入端信息位, m 表示编码器中寄存器的节数, 卷积码的 k 和 n 通常很小。从编码器的输入端

来看,卷积码仍以 k 位数据为一组,分组输入。从输出端来看,卷积码是非分组的,它输出的 n 位码元不仅与当前输入的 k 个信息位有关,还与前面 m 段的输入信息位有关,所以卷积码属于有记忆的非分组码。卷积码为有记忆编码,其记忆(或称约束)度为 $N = m + 1$,编码过程中互相关联的码元个数为 nN , nN 称为编码的约束长度。卷积码的纠错能力随 N 的增大而提高,而差错率随 N 的增大呈指数下降。在编码器复杂性相同的情况下,卷积码的性能优于分组码,但卷积码没有分组码那样的严密的数学分析,需要通过计算机进行优码的搜索。卷积码的码长较小,因此适合于串行形式传输且时延较小。

3.6.2 卷积码的编译码原理

1. 卷积码编码器

卷积码编码器的一般结构如图 3-10 所示,它包括: n 段输入移位寄存器,每段移位寄存器有 k 级存储单元,共 nk 位寄存器;一组 n 个模 2 加法器, n 等于卷积编码器输出位数,每个模 2 加法器连接到一些移位寄存器的输出端,数目可以不同,连接的移位寄存器也可以不同; n 位输出移位寄存器, n 个模 2 加法器与 n 位输出移位寄存器一一对应链接,模 2 加法器的运算结果即为卷积编码输出,每输入 k 位,得到 n 位输出。

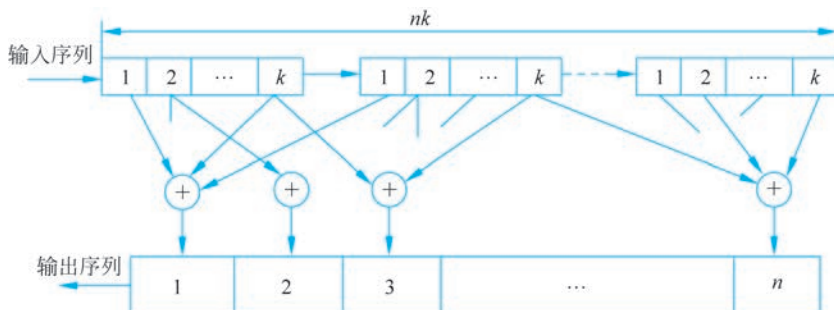


图 3-10 卷积编码器的一般结构

图 3-11 为 $(2,1,2)$ 卷积编码器,图中 $n=2, k=1, m=2$,它由移位寄存器、模 2 加法器和开关电路组成。每输入一位信息经编码器产生两位输出。假设移位寄存器的起始状态为全零,即 $S_1 S_2 S_3$ 为 000, S_1 为当前输入数据,而移位寄存器状态 $S_2 S_3$ 存储之前的数据,输出码字由下式确定。

$$\begin{cases} C_1 = S_1 \oplus S_2 \oplus S_3 \\ C_2 = S_1 \oplus S_3 \end{cases} \quad (3-11)$$

各移位寄存器初始状态为全零。当第一个输入位为 0 时,输出位为 00; 输入位为 1, 输出位为 11; 当第二个输入位进入时,第一位右移一位,此时输出位受到当前输入位和前面输入位的影响; 第三个输入位到来时,第一、二个输入位分别右移一位,此时输出位受到当前输入位和前面两个输入位的影响; 第四个输入位到来时,第一个输入位移出寄存器而消失,此时输出位受到当前输入位和前面两个输入位的影响。

对于图 3-11,假设输入的信息序列为 11010,为了使信息全部通过移位寄存器,必须在信息码元后面加 3 个 0,即输入序列为 11010000。表 3-5 给出了 $(2,1,2)$ 卷积编码器在移位寄存器的起始状态为全零、编码器输入序列为 11010000 时,对应的编码器输出移位

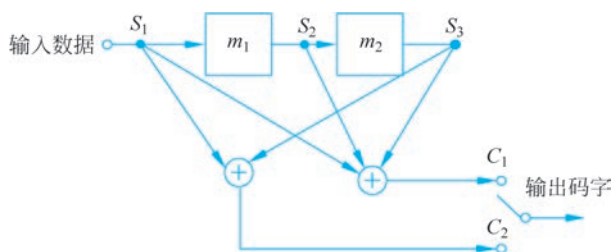


图 3-11 (2,1,2)卷积编码器

寄存器的状态变化过程,其中输入移位寄存器 $S_3 S_2$ 的四种状态 00、01、10、11 分别用 a 、 b 、 c 、 d 表示。(2,1,2)卷积码的编码约束度为 3,约束长度为 6。

表 3-5 (2,1,2)卷积编码器的状态变化过程

S_1	1	1	0	1	0	0	0	0
$S_3 S_2$	00	01	11	10	01	10	00	00
$C_1 C_2$	11	01	01	00	10	11	00	00
状态	a	b	d	c	b	c	a	a

2. 卷积码的描述

描述卷积编码器方法有图解法和解析法。图解法包括树状图、状态图、网格图；解析法包括矩阵形式、生成多项式形式。解析法较为抽象,本节仅介绍图解法。

1) 树状图

(2,1,2)卷积码编码器中随着输入序列的进入,编码器移位的过程和各种可能的输出序列可用图 3-12 所示的树状图来表示。从节点 a 开始,由 a 出发有两条路径, $a=0$ 取

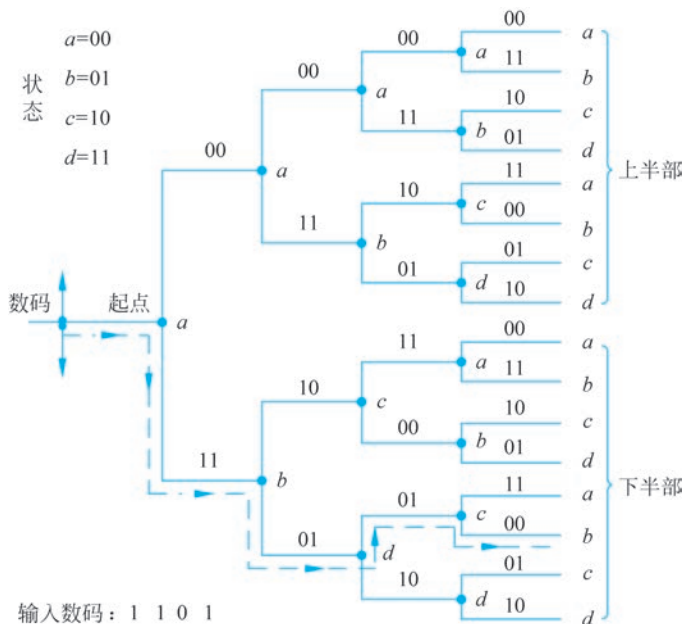


图 3-12 (2,1,2)卷积码树状图

上支路, $a=1$ 取下支路, 输入第二位时, 移位寄存器右移一位, 上支路的状态为 00, 下支路状态变为 01。新的比特位到来时, 随着移位寄存器状态和输入位的不同, 树状图分为 4 个支路, 以此类推, 可得整个树状图。输入不同的信息序列, 编码器走的路径就不同, 从而得到输出不同的码序列。例如, 输入数据为 1101 时, 其路径如图中虚线所示, 输出码序列为 11010100, 与表 3-5 结果相同。

2) 状态图

将已到达稳定状态的一节网络取出, 当前状态与下一状态合并, 得到图 3-13 所示的状态图。图中, 有 a 、 b 、 c 、 d 共 4 个节点, 同样分别表示 S_2S_3 的 4 种可能状态: 00、01、10、11。每个节点有两条线离开该节点, 实线表示输入数据为 0, 虚线表示输入数据为 1, 两个闭合圆圈分别表示“ $a \rightarrow a$ ”和“ $d \rightarrow d$ ”的状态转移, 线旁的数字即为输出码字。

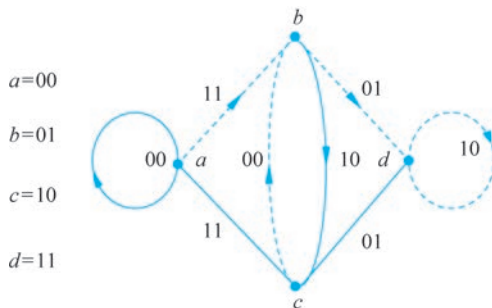


图 3-13 (2,1,2) 卷积码的状态图

3) 网格图

将树状图进行适当的变形可以得到一种更为紧凑的图解标志方法, 即网格图法, 如图 3-14 所示。在网格图中, 将树状图中具有相同状态的节点(状态)合并在一起, 码树上的上分支(对应输入 0)用实线表示, 下分支(对应输入 1)用虚线表示。分支上标注的码元为对应的输出, 自上而下 4 行节点分别表示 a 、 b 、 c 、 d 四种状态。

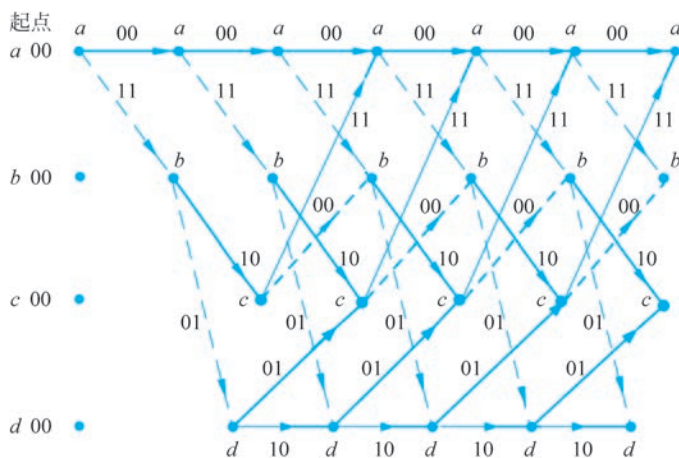


图 3-14 卷积码的网格图

3. 卷积码译码

卷积码译码可分为代数译码和概率译码两大类。代数译码是利用生成矩阵和监督矩阵来译码,最主要的方法大数逻辑译码,该方法利用编码本身的大数结构,而不考虑信道的统计,硬件实现简单,但性能较差。概率译码比较实用的有维特比译码和序列译码两种,其中维特比译码应用较多,这里只简单介绍维特比译码方法。

维特比译码是 1967 年由 Viterbi 提出的一种卷积码的译码算法,该算法在卫星通信中作为标准技术得到了广泛应用,算法的本质是利用最大似然准则进行译码,其基本思想是将接收码字序列与所有可能的发送码字序列进行比较,从中选择与接收序列汉明距离最小的发送序列(汉明距离最小相当于似然函数最大)作为译码输出。对于长度为 k 的二进制发送序列,需要对可能发送的 2^k 个不同的序列的 2^k 条路径似然函数累加进行比较,选取其中最大的一条。显然,译码算法的计算量将随着 k 的增加呈指数增加,实际中需要采取一定的措施简化。维特比译码使用网格图描述卷积码,每个可能的发送序列都与网格中的一条路径相对应。根据网格图的路径汇聚特性,如果在某个节点上发现某条路径已不可能与接收序列的距离最小,就放弃这条路径,然后在剩下的“幸存”路径中重新选择码路径,这样一直进行到最后第 L 级。由于这种方法过早丢弃了那些不可能的路径,因而减小了译码的工作量。

维特比译码的基本步骤:对于网格图第 i 级的每个节点,计算到该节点的所有路径的距离量度,即在前面 $i-1$ 级路径距离量度的基础上累加上第 i 条支路的距离量度,并从中选择距离量度最小的路径作为幸存路径,其他路径丢弃。上述译码过程可概括为“累加—比较—选择”(Accumulation-Compare-Selection, ACS)运算。

以图 3-11 所示的卷积编码器为例来说明维特比译码的过程。设发送信息序列为 11010,为了使全部信息位能通过编码器,在发送序列后面加上 3 个 0,即 11010000,可以得到如表 3-5 所示的计算结果,这时编码器输出的序列 1101010010110000,那么移位寄存器的状态转移路线为 $a \rightarrow b \rightarrow d \rightarrow c \rightarrow b \rightarrow c \rightarrow a \rightarrow a$,信息全部离开编码器,最后回到状态 a 。

假设接收序列为 0101011010010001,有 4 位码元发生误码,图 3-15 所示的维特比译码网格图说明了整个译码过程。由于该卷积码的约束长度为 6,故先选前 3 段接收序列 010101 作为标准。与到达第 3 级 4 个节点的 8 条路径进行对照,逐步算出每条路径与作为标准的接收序列 010101 之间的累计码距。

1) 到达第 3 级的情况

到达节点 a 的两条路径是 000000 与 111011,它们与 010101 之间的码距分别是 3 和 4;到达节点 b 的两条路径是 000011 和 111000,它们与 010101 之间的码距分别是 3 和 4;到达节点 c 的两条路径是 001110 和 110101,它们与 010101 之间的码距分别是 4 和 1;到达节点 d 的两条路径是 001101 和 110110,它们与 010101 之间的码距分别是 2 和 3。每个节点保留一条码距较小的路径作为幸存路径,它们分别是 000000、000011、110101 和 001101。这些路径就是如图 3-15 所示的到达第 3 级节点 a 、 b 、 c 、 d 的 4 条路径,累计码距分别由括号里的数字标出。

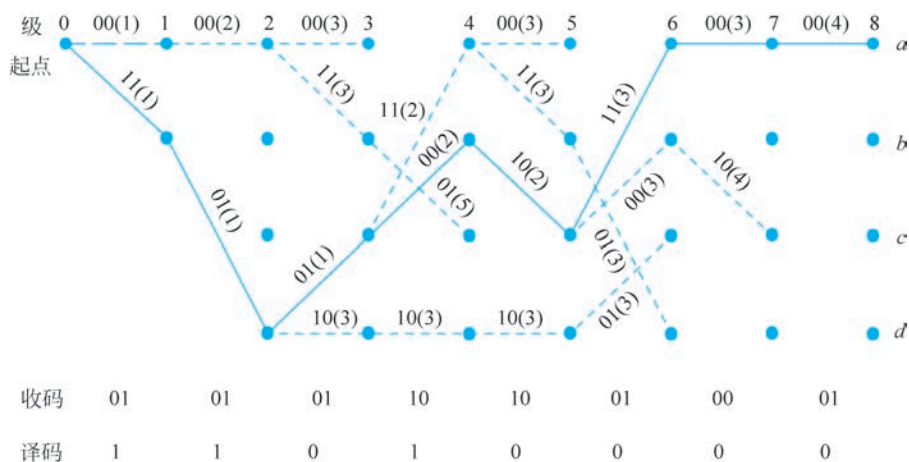


图 3-15 维特比译码网格图

2) 到达第 4 级的情况

节点 a 的两条路径是 00000000 和 11010111, 节点 b 的两条路径是 00000011 和 11010100, 节点 c 的两条路径是 00001110 和 00110101, 节点 d 的两条路径是 00110110 和 11011010。将它们与接收序列 01010110 对比求出累计码距, 每个节点仅留下一条码距小的路径作为幸存路径, 它们分别是 11010111、11010100、00001110 和 00110110。

3) 继续筛选幸存路径

逐步推进筛选幸存路径, 到第 7 级时, 只要选出到达节点 a 和 c 的两条路径即可, 因为到达终点第 8 级 a 只可能从第 7 级的节点 a 和 c 出发。最后得到到达终点 a 的一条幸存路径, 即为译码路径, 如图 3-15 中的实线所示。根据这条路径, 对照图 3-14 可知译码结果为 11010000, 与发送信息序列一致。

维特比译码需要在网格图的每一列节点处进行累计距离的“累加—比较—选择”运算, 并带有译码的回溯过程。译码器应由软、硬件联合组成, 译码器的复杂性也将随着状态数和约束长度的增加而上升。根据上面的讨论, 可以得出维特比译码器的基本结构, 如图 3-16 所示。图中, 网格图参量存储器用于寄存可能的发码路径参数, 输入变换和支路量度计算将输入信号变换为可用计算距离的量度, 并计算收码与可能的发码路径间距离, 然后由累加—比较—选择(ACS)电路完成支路量度的“累加—比较—选择”作用, 将计算结果存入路径量度存储器, 再进行最佳路径选择, 利用输出缓冲器完成回溯译码过程。

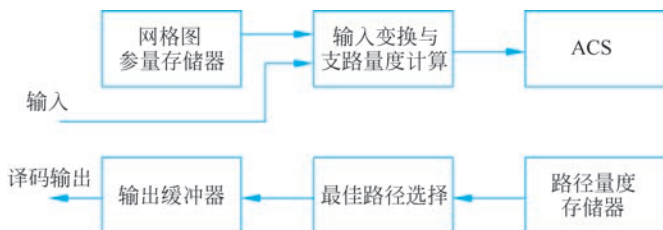


图 3-16 维特比译码器的基本结构



视频

3.7 高性能信道编码

3.7.1 Turbo 码

1. Turbo 码的概念

Turbo 码又称并行级联卷积码,它是由 Berrou 等于 1993 年在国际 ICC 会议上提出的一种新型的信道纠错编码。Turbo 码巧妙地将卷积码和随机交织器结合在一起,能够产生很长的码字并提供更好的传输性能,达到接近随机编码的目的。如果译码方式和参数选择得当,Turbo 码性能可以接近香农(Shannon)极限。因此,这一超乎寻常的优异性能引起了信息与编码理论界的轰动。

由于 Turbo 码的优异性能不是从理论研究角度给出的,而仅是计算机仿真的结果,因此 Turbo 码的理论基础还不完善。随着研究的不断深入,其性能也不断提高,并在强干扰环境下显示了其广阔的应用前景。目前,Turbo 码已经成为以大容量、高数据率和承载多媒体业务为目的的信道编码方案,例如移动卫星通信系统组织公布的 Inmarsat-F 系统中就用到了 Turbo 编码技术。此外,Turbo 码在高清晰度数字电视传输系统应用中能使大量的数字信号准确无误地传输,真正做到高清晰度的图像质量。不仅如此,迭代译码的思想已经作为“Turbo 原理”而广泛应用于均衡、调制、信道检测等领域。

香农极限定理指出,要想在一个存在噪声的确定带宽的信道中可靠地传送信号,可以加大信噪比或在信号编码中加入附加的纠错码;若要想通过提高信号编码效率来达到信道容量要求,就要使编码的长度尽可能长而且尽可能随机,但这在实际中计算量太大而不可能实现。Turbo 码的不同之处在于:它通过一个交织器,使之达到接近香农极限的性能。此外,它采用迭代译码策略来逼近最大似然译码,使得译码复杂性大大降低。

Turbo 码最初以并行级联卷积码(Parallel Concatenated Convolutional Code, PCCC)的形式出现,后来为了克服误码率的错误平层(误码率曲线随着信噪比增加在某一点斜率变缓),Benedetto 和 Divsalar 等提出了串行级联卷积码(Serial Concatenated Convolutional Code, SCCC),又称为串行级联 Turbo 码,后来又将 PCCC 与 SCCC 结合,设计了混合级联卷积码(Hybrid Concatenated Convolutional Code, HCCC)。鉴于理论分析不完善,实现过程复杂,此处仅对 PCCC 的编译码原理进行简单分析和介绍。

2. Turbo 码的编码原理

Turbo 编码器由两个或两个以上的简单分量编码器(也称子编码器)通过交织器并行级联在一起构成,如图 3-17 所示。图中两个分量编码器分别采用递归系统卷积(Recursive Systematic Convolutional, RSC)码编码器,且具有完全相同的结构。交织器是一个单输入单输出设备,由一定数量的存储单元构成, $M \times N$ 存储单元构成存储矩阵(其中 M 为存储矩阵的行数, N 为存储单元的列数),各个存储单元可用它在矩阵中所处的行数和列数来表示。交织器的输入和输出符号序列有相同的字符集,只是各符号在输入与输出序列中的排列顺序不同。编码器将输入数据的 n 位分为一组,由编码器 1 进行行编码,再经过交织器由编码器 2 进行列编码。对两路编码的校验位进行抽取,删除适

当码元以提高码率。卷积码码率较低,可进行抽取;分组码码率较高,可直接省略。最后进行复接,完成串并转换。

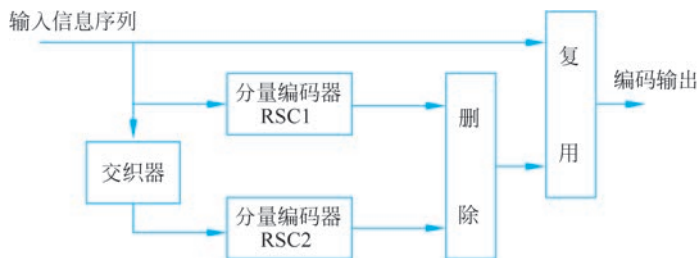


图 3-17 Turbo 码编码器结构

交织器的引入可以说是 Turbo 码的一大特色,信息位流顺序流入交织器,以某种方式乱序读出,或者以乱序形式读入,再以顺序形式读出。交织器的作用一般是对抗突发错误,更重要的是它可以改变码的重量分布,将原始信息序列打乱,使交织前后的信息序列的相关性减弱。交织长度越长,相邻反馈信号的相关性就越低,可更好地实现迭代译码。

交织器的设计准则:最大程度地置乱原始数据排列顺序,避免置换前相距较近的数据在置换后仍相距较近,特别要避免置换前相邻数据在置换后再次相邻;尽可能避免与同一信息位直接相关的两个分量编码器中的校验位在复用时均被删除;对于不归零的编码器,交织器设计时要避免出现错误平层效应。在满足上述要求的交织器中再选择一个好的交织器,使码字之间的最小距离 $d_{\min}$ 尽可能大,而码重为 $d_{\min}$ 的码字数要尽可能少,以改善 Turbo 码在高信噪比时的性能。

3. Turbo 码的译码

Turbo 码的译码过程采用迭代反馈的方法,每次迭代采用的是软输入和软输出,并不断地将输出的码元反馈到输入,与输入进行比较以增加译码的正确性。Turbo 码译码器结构如图 3-18 所示,由两个软输入软、输出译码器 DEC1 和 DEC2 串行级联组成,交织器与编码器中所使用的交织器相同。Turbo 码的译码过程:译码器 DEC1 对分量编码器 RSC1 输出的校验码进行最佳译码,产生关于信息序列中每位的似然信息,并将其中的“新信息”经过交织送给 DEC2。译码器 DEC2 将此信息作为先验信息,对分量编码器 RSC2 输出的校验码进行最佳译码,产生关于交织后的信息序列中每位的似然信息,将其中的“新信息”经过解交织送给 DEC1,进行下一次译码。这样经过多次迭代,DEC1 和 DEC2 的新信息趋于稳定,似然比渐近值逼近于对整个码的最大似然译码,然后对此似然比进行判决,即可得到信息序列的每位的最佳估值序列。

Turbo 码译码器采用迭代译码,它是由 Bahl、Cocke、Jelinek 和 Raviv 提出来的,其策略是根据接收序列计算后验概率,为降低长码计算的复杂度,由两个分量码译码器分别计算后验概率。常见的 Turbo 码的译码算法有最大后验概率(MAP)算法、BCJR 算法、对数域的 LOG-MAP 算法及 MAX-LOG-MAP 算法、减少状态搜索的 M-BCJR 算法和 T-BCJR 算法等。

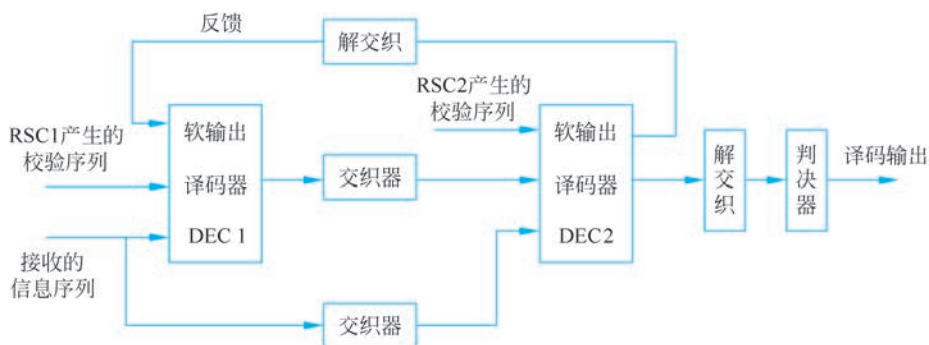


图 3-18 Turbo 码译码器结构

Turbo 码的优点是编码性能高,编码后的速率接近香农容量极限,同时编码的可靠性增强,但时延较大,因此 Turbo 码适用于对误块率要求严格,时延要求不严格的场合,如移动通信网络中的高速下载类业务。在 3G 移动通信、WiMAX、深空通信等许多国际标准中,Turbo 码作为首推的纠错编码。从目前的研究结果看,Turbo 码在学术界研究中的地位相当重要,Turbo 码与空时编码、网格编码调制(TCM)的结合、Turbo 码均衡技术、Turbo 码多用户检测技术以及在多输入多输出(Multiple Input and Multiple Output, MIMO)信道、协作通信中的应用均为当前的研究热点。

3.7.2 LDPC 码

1. LDPC 码的基本概念

虽然 Turbo 码获得了极大的成功,但是 Turbo 码也存在一些缺点,如译码复杂度相对较高、译码时延长等。美国麻省理工学院的 Gallager 于 1963 年提出一种具有稀疏奇偶校验矩阵的分组码——低密度奇偶校验(Low Density Parity Check, LDPC)码。LDPC 码具有以下优点:①码本身具有良好的内交织特性,抗突发差错能力强,从而避免了专门引入交织器所带来的时延;②具有更好的分组误码性能,错误平层大大降低;③描述简单,理论分析具有可验证性;④译码复杂度低于 Turbo 码,且可实现完全的并行操作,便于硬件实现;⑤吞吐量大,具有高速译码能力。

目前,LDPC 码被认为是迄今为止性能最好的码,是当今信道编码领域的最令人瞩目的研究热点,近几年国际上对 LDPC 码的理论研究、工程应用和超大规模集成电路(VLSI)实现方面的研究都已取得重要进展。目前 LDPC 编码技术在欧洲卫星广播系统 DVB-S3、深空通信、磁记录系统、4G 和 5G 移动通信等领域均有应用。另外 LDPC 与 MIMO、OFDM 等技术的结合也是当前热点研究的问题。

LDPC 码是一种线性分组码,由其校验矩阵 $\mathbf{H}$ 的稀疏性而得名,即校验矩阵中大部分元素为“0”,仅包含很少的非零元素。可以理解为 n 维线性空间上的 k 维线性子空间,也可以看作满足一系列约束方程的解向量的集合。码长为 n 、校验位长度为 r 、信息位长度 $k=n-r$ 的 LDPC 码 $\mathbf{C}$ 可以由校验矩阵 $\mathbf{H}$ 唯一确定。 $\mathbf{H}$ 的每一行对应一个校验方程,每一列对应码字的一位,满足如下方程的 n 维向量 $\mathbf{V}$:

$$\mathbf{H}\mathbf{V}^T = \mathbf{0} \quad (3-12)$$

即为码 $\mathbf{C}$ 的一个码字。 $\mathbf{H}$ 的每一行中非零元素的个数 ρ 称为 $\mathbf{H}$ 的行重,也称校验节点的“度”,代表参与该行对应校验方程的变量节点的个数;每一列中非零元素的个数 λ 称为 $\mathbf{H}$ 的列重,也称变量节点的“度”,代表该列对应码字比特参与校验方程的个数。

根据校验矩阵行列重的不同,LDPC 码可分为规则 LDPC 码和非规则 LDPC 码。若校验矩阵 $\mathbf{H}$ 每行的非零元素个数相同,且每列中非零元素个数也相同,则该 LDPC 称为规则 LDPC 码;否则,称为非规则 LDPC 码。

规则 LDPC 码是一个 m 行 n 列的稀疏矩阵 $\mathbf{H}_{m \times n}$ 的零空间, $\mathbf{H}$ 为 LDPC 码的校验矩阵。规则 LDPC 码的校验矩阵为稀疏矩阵,具有如下特性:

- (1) 所有行重都为一定值 ρ ;
- (2) 所有列重都为一定值 λ ;
- (3) 任意两行(列)中的“1”在共同位置最多只出现 1 次;
- (4) 行重和列重相对于码长来说都非常小。

规则 LDPC 码通常用 (n, λ, ρ) 表示。对于规则 LDPC 码,显然, $r\rho = n\lambda$ 成立。一个简单的规则 LDPC 码 $(10, 2, 4)$ 的校验矩阵如下:

$$\mathbf{H} = \begin{bmatrix} 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 1 & 0 & 0 & 1 & 1 & 0 \\ 0 & 0 & 1 & 0 & 0 & 1 & 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 1 & 1 \end{bmatrix} \quad (3-13)$$

式中的校验矩阵是一个系数矩阵,对应的 LDPC 码的码长 $n=10$,列重 $\lambda=2$,行重 $\rho=4$ 。

2. LDPC 码的 Tanner 图表示

LDPC 码的校验矩阵也可以利用图论中的二分图表示。图 3-19 给出了 $(10, 2, 24)$ 规则 LDPC 码的 Tanner 图。图的下方有 n 个节点,每个节点表示码字的信息位,对应于校验矩阵的各列;图的上方有 r 个节点,每个节点表示码字的一个校验集,称为校验节点,代表校验方程,对应于校验矩阵的各行;与校验矩阵中“1”的元素相对应的左右两节点之间存在连接边,这两条边两端的节点称为相邻节点,每个节点相连的边数称为该节点的度数。图 3-20 为 $(7, 3)$ 非规则 LDPC 码的 Tanner 图。

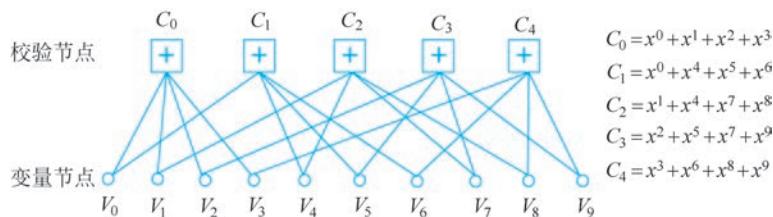


图 3-19 $(10, 2, 4)$ 规则 LDPC 码的 Tanner 图

对于规则 LDPC 码,其 Tanner 图中所有校验节点的度都相同(等于 ρ),所有变量节点的度也都相同(等于 λ)。而对于非规则 LDPC 码,其 Tanner 图中校验节点的度不一定

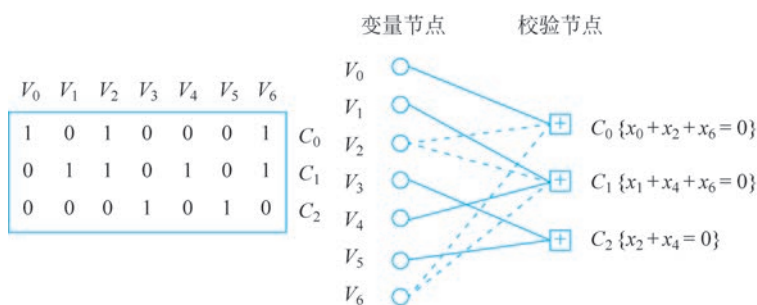


图 3-20 (7,3)非规则 LDPC 码的 Tanner 图

相同,变量节点的度也不一定相同。一般来说,非规则 LDPC 码的性能优于规则 LDPC 码。非规则 LDPC 码一般用其度分布序列 $\{\rho_i\}$ 和 $\{\lambda_i\}$ 或者多项式 $\rho(x)$ 和 $\lambda(x)$ 表示:

$$\rho(x) = \sum_{i=1}^{d_c} \rho_i x^{i-1}, \quad \lambda(x) = \sum_{j=1}^{d_v} \lambda_j x^{j-1} \quad (3-14)$$

式中: ρ_i 表示与度为 i 的校验节点相连的边的个数与总边数之比; λ_j 表示与度为 j 的变量节点相连的边的个数与总边数之比; d_c 和 d_v 分别是校验节点和变量节点的最大度数。

$$\text{显然, } \sum_{i=1}^{d_c} \rho_i = 1, \sum_{j=1}^{d_v} \lambda_j = 1.$$

一个 LDPC 码 C 有多个不同的校验矩阵,即不同的校验矩阵的零空间可能给出同一个码,因此码 C 具有多个 Tanner 图表示,尽管这些 Tanner 图对应同一个码,但基于不同 Tanner 图的迭代译码将会给出不同的译码性能,所以构造 LDPC 码的校验矩阵时需遵循一定的准则,以获得较好的纠错性能。

在 Tanner 图中,从一个节点出发,在经过不同节点后回到同一个节点所构成的图称为一个环。图 3-20 中,虚线部分构成了一个环。该环遍历 4 条边,称其为 4 环。在 Tanner 图中会存在许多这样的闭合环路,其中最短的闭合环路的长度称为该图的围长。因为 Tanner 图中任意两节点之间最多有一条边相连,因此图中的环的长度均为偶数且大于等于 4。由于 LDPC 码译码采用迭代译码,其算法的推导是基于在节点间传递的信息统计独立的假设的,当 LDPC 码校验矩阵对应的 Tanner 图中存在长度为 $2l$ 的环路时,这些消息只在前 l 轮迭代过程中满足独立性假设(译码的迭代过程包括一次比特节点更新和一次校验节点更新),因此小循环的存在会影响 LDPC 码的译码性能。在设计 LDPC 码的校验矩阵时,尤其要避免最短环 4 的出现,环 4 的出现将极大地影响 LDPC 码的性能。希望在 LDPC 码的 Tanner 图中大围长多,小围长少。

3. LDPC 码的构造

LDPC 码是定义在稀疏奇偶校验矩阵上的一种线性分组码,LDPC 码的构造就是稀疏奇偶校验矩阵的构造。根据构造方式的不同,LDPC 码的构造大致可以分为随机构造法和结构化构造法。

随机构造法是一种非确定性的构造方法,是在一定的约束条件下通过计算机随机搜

索的方法构造校验矩阵,即在给定基本的构造参数 (n, λ, ρ) 基础上,通过向一个全零的稀疏校验矩阵中随机填充“1”,使填充进去的“1”的位置符合构造的约束条件。相对于随机构造方法构造出来的校验矩阵中非零元素的位置杂乱无序,结构化构造方法构造出来的校验矩阵的非零元素具有分块的特点。为了能够保证 LDPC 码易于硬件实现的特点,很多结构化构造方法都使用了基于准循环 LDPC(QC-LDPC)移位的方法。

4. LDPC 码的译码

Gallager 给出了硬判决算法和软判决算法两种迭代译码算法。硬判决算法操作简单,易于硬件实现,但是性能较差;软判决算法性能较好,但实现复杂度较高。

在硬判决算法方面,最基本的硬判决是 Gallager 提出的比特翻转(Bit-Flipping, BF)算法,其基本思想是传输序列的某一比特参与的校验方程错误越多,则认为此比特错误的概率越大,因此翻转该比特。BF 算法是简单的硬判决算法,仅需要逻辑运算,实现非常简单,但译码性能一般。为了改善硬判决译码的性能,Y. Kou 和 F. Guo 等提出了一种基于软信息的加权比特翻转(Weighted Bit-Flipping, WBF)算法以及可靠性加权比特翻转(Reliability Ratio based Weighted Bit-Flipping, RRWBF)算法。

在软判决算法方面,Gallager 提出了基于置信传播(Belief Propagation, BP)的迭代译码算法,这是一种软输入迭代概率译码算法。置信传播算法中传递的信息是前一级计算软判决输出的概率或似然比。BP 算法主要分为三个步骤。首先进行初始化,每个变量节点根据信道输出值计算其先验消息,然后进入消息迭代过程。每轮迭代中,每个校验节点根据收到的最新消息进行消息更新并将更新后的消息传递给与其相连的每个变量节点;然后每个变量节点利用初始消息及收到的最新校验节点消息进行消息更新,并将更新后的消息传递给与其相连的每个校验节点。每轮迭代后,每个变量节点进行译码判决,若得到的判决序列满足所有的校验方程则显示译码成功,否则进行新一轮迭代直至译码成功或者达到预先设定的最大迭代次数并显示译码失败。

不管是硬判决还是软判决,LDPC 码的译码均采用迭代算法。LDPC 码的译码既可以选择实现简单、性能稍差的 BF 硬判决算法,也可以选择实现复杂、性能优越的 BP 软判决算法,还可以选择趋于中间的各种改进型算法。

3.7.3 TCM 码

1. TCM 的基本概念

在传统的数字通信系统中,发送端的纠错编码和调制是分开进行的,接收端的译码和解调也是分开进行的。纠错编码需要冗余度,在码流中增加监督比特可达到检错或纠错的目的,但此时码流的比特速率增加,编码后信号速率增加,相应传输带宽需求增加。若系统带宽一定,则编码增益需要依靠降低信息传输速率来获得。若要求不降低信息传输速率,编码增益又需要依靠牺牲带宽来获得。若系统频带受限,则需要通过增加信号点集以降低因编码冗余引起的带宽增加。信号点集的增加使信号空间距离减小,因此造成误码性能变差,在系统功率受限情况下无法保证误码性能。若调制和编码仍按传统的相互独立方法进行设计,则无法得到令人满意的效果。

网格编码调制是由 Ungerboeck 在 1982 年提出的一种高级的编码调制方法,其本质是将编码和调制相结合,利用信号集空间的冗余度来进一步降低误码率。在加性高斯白噪声信道中,这样处理后使得决定系统性能的主要参数由卷积码的汉明距离转化为传输信号间的欧几里得距离。因此最佳网络编码的设计是基于欧几里得距离,这就要求必须将编码器和调制器当作一个统一的整体进行综合设计,使得编码器和调制器级联后产生的编码信号序列具有最大的欧几里得距离。从信号空间的角度看,这种最佳编码调制的设计实际上是一种对信号空间的最佳分割。

TCM 技术与常规的非编码多进制调制相比具有较大的编码增益且不降低频带利用率,所以特别适合限带信道的信号传输。网格编码有两个以下基本特征:

(1) 在信号空间的信号点数目比无编码的调制情况下对应的信号点数目要多(通常增加 1 倍),增加的信号点使编码有了冗余,而不牺牲带宽。

(2) 采用卷积码编码规则,在一系列信号点之间引入依赖关系,使得只有某些信号点图样是许可使用的信号序列,并可模型化为格状网络,因而称为“网格”编码调制。

2. TCM 的编解码原理

TCM 最优码是按照编码信号的网格图确定的,TCM 最优码方案通过一种特殊的信号映射可变成卷积码形式,这种映射的原理是将调制信号集合分割成子集,使得子集内信号间具有更大的欧几里得距离。因此,TCM 设计的一个主要目标是寻找与各种调制方式对应的卷积码,卷积码的每个分支与信号点映射后,使得每条信号路径之间有最大的欧几里得距离。根据这个目标,对于多电平/多相位的二维信号空间,把信号点集不断地扩大为 2、4、8 等多个子集,使它们之间的信号点的最小距离不断增大,这种映射关系称为集合分割映射。集合分割映射的每一次分割是将一个较大的信号子集分割成较小的两个子集,这样可得到一个表示集合分割的二叉树,每经过一次分割,子集数加倍,而子集内信号最小距离增加,一直分割到子集内只含欧几里得距离最大的两个点。

下面举例说明集合分割映射的具体实现。图 3-21 给出了 8PSK 信号的集合分割映射过程示意图。图中,所有 8 个信号点分布在一个圆周上,经连续 3 次划分后,分别产生 2、4、8 个子集,它们的共同点是最小码距逐次增大,即 $d_0 < d_1 < d_2$ 。设最上层星座图的圆的半径为 1,则该层信号的最小欧几里得距离 $d_0 = 2\sin(\pi/8) = 0.765$ 。第一级分割后,得到 2 个子集,每个子集相当于一个 QPSK,此时信号的最小欧几里得距离 $d_1 = 1.414$ 。第二级分割后,得到 4 个子集,每个子集相当于一个 2PSK,信号的最小欧几里得距离 $d_2 = 2$ 。从图 3-21 可以看出每次分割后,信号的欧几里得距离都得到增加。

根据集合分割的思想,可以设计比较简单而有效的 TCM 编码方案,TCM 编码调制器的系统结构如图 3-22 所示。设输入码字有 n 位,无编码调制时,二维信号空间中应有 2^n 个信号点与之对应。当采用编码调制时,为增加冗余度,有 2^{n+1} 个信号点对应。在图 3-21 中,可划分为 4 个子集,对应于码字的 1 位加到编码效率为 1/2 的卷积码输入端,输出 2 位,选择相应的子集,码字剩余的未编码数据比特确定信号与子集中信号点之间的映射关系。

TCM 系统使用冗余多进制调制与一个有限状态的网格编码器相结合,由编码器控

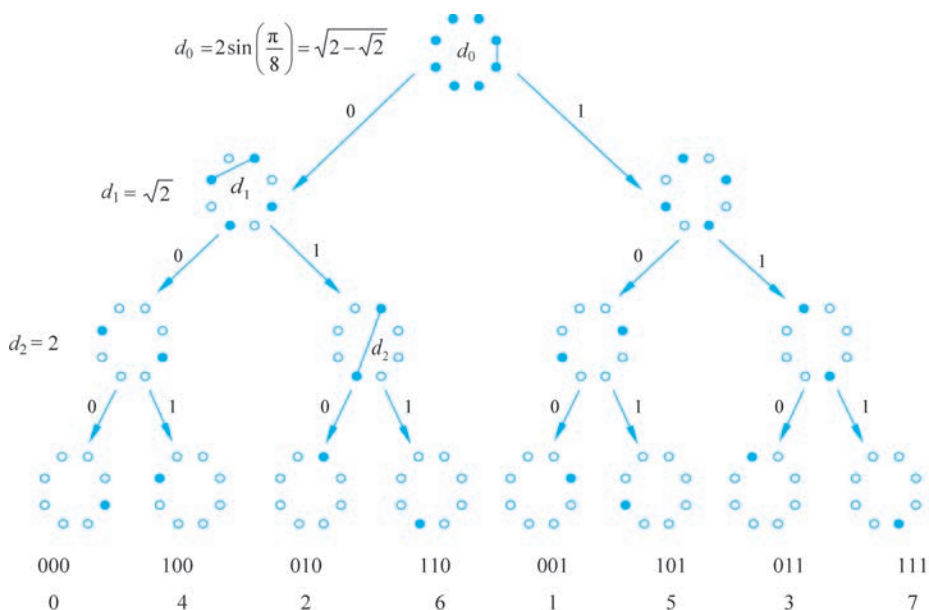


图 3-21 8PSK 信号的集合分割过程

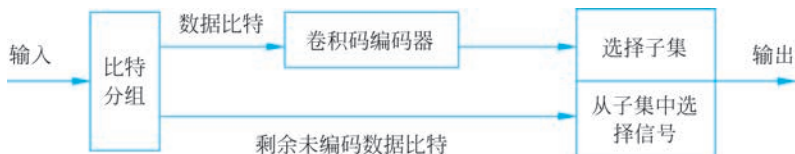


图 3-22 TCM 网格编码系统结构

制选择调制信号,以产生编码符号序列,编码器和调制器级联后产生的编码信号序列具有最大的欧几里得自由距离。当编码调制后的信号序列经过一个加性高斯白噪声信道后,在接收端可以采用维特比算法实现译码。用维特比算法寻找最佳格状路径,以最小码距为准则,执行最大似然序列检测解调出接收的信号序列。在格状图中,每一条支路对应一个子集,而不是一个信号点。检测的第一步是确定信号点,首先确定信号点所在的子集,在码间距离意义下,这个子集是最靠近接收信号点的子集。信号点确定后,它和接收点之间的平方欧几里得距离可用于以后的支路求解,并可用维特比算法继续求解。

TCM 技术自提出以后就得到了广泛的研究和应用。1984 年 Wei 针对信道中的各种干扰因素对相位的影响,提出了克服相位模糊的旋转不变码。从理论上分析,3dB 以内的编码增益可以通过增加网格码编码寄存器的状态数得到;但当状态数增加到一定程度后,编码增益的增加变得十分缓慢,而实现电路的复杂性却呈指数形式增长,实现电路的复杂性几乎抵消了由增加状态数而带来的编码增益。于是 Forney 等提出了带限信道上的多维 TCM 技术。由于网格编码调制可以得到具有最大欧几里得距离的码序列,因此在多进制调制场合获得了广泛应用,如计算机上用的调制解调器、卫星通信以及一些移动通信系统等。Turbo 码作为级联卷积码,也可以与调制技术联合设计,实现 Turbo-TCM。

习题

1. 话音编码和数据压缩编码的实现方法有哪些?
2. 简述波形编码、参量编码和混合编码各自的特点。
3. 设信源共有 8 个信源符号,其概率分布为

$$\begin{bmatrix} X \\ P(X) \end{bmatrix} = \begin{bmatrix} x_1 & x_2 & x_3 & x_4 & x_5 & x_6 & x_7 & x_8 \\ 0.37 & 0.16 & 0.14 & 0.13 & 0.07 & 0.06 & 0.04 & 0.03 \end{bmatrix}$$

试对该信源进行香农编码和霍夫曼编码,并计算各自的平均码长和编码效率。

4. 简述统计编码、预测编码和变换编码的基本原理及特点。
5. 简述差错控制方式的种类及各自的优缺点。
6. 已知(7,3)分组码的监督关系式为

$$\begin{cases} x_6 + x_3 + x_2 + x_1 = 0 \\ x_6 + x_2 + x_1 + x_0 = 0 \\ x_6 + x_5 + x_1 = 0 \\ x_6 + x_4 + x_0 = 0 \end{cases}$$

求其监督矩阵、生成矩阵、全部码字及纠错能力。

7. 已知(15,7)循环码,生成多项式 $g(x) = x^8 + x^7 + x^6 + x^4 + 1$ 。
(1)写出该循环码的生成矩阵;(2)若信息码组为 0011001,写出系统的输出码组。
8. 已知一个(3,1,2)卷积码 $g_1(x) = x^2 + x + 1, g_2(x) = x^2 + x + 1, g_3(x) = x^2 + 1$ 。
(1)画出该编码器的框图;(2)画出状态图、树图;(3)求该码的自由距离。
9. 简述卷积编码、Viterbi 译码的原理。
10. 简述 Turbo 码和 LDPC 码的基本原理和特点。