

TURING

计算机程序设计竞赛权威指导书

程序设计中实用的 数据结构

王建德 吴永辉 编著



人民邮电出版社
POSTS & TELECOM PRESS

The background of the cover features a complex, abstract design. It consists of numerous small, dark dots of varying sizes, some of which are arranged in a grid-like pattern. Overlaid on these dots are several thin, curved lines that sweep across the page, creating a sense of movement and depth. The overall effect is a modern, technical aesthetic that suggests data and computation.

TURING

程序设计中实用的 数据结构

王建德 吴永辉 编著

人民邮电出版社
北京

绪 言

最近几年，我们长期从事程序设计研究和竞赛辅导，并在此基础上精心编写了 ACM/ICPC 竞赛的系列教程，其中包括《新编实用算法分析与程序设计》（人民邮电出版社 2008 年出版）、《程序设计中常用的计算思维方式》（人民铁道出版社 2009 年出版）、《程序设计中常用的解题策略》（人民邮电出版社 2009 年出版）和《程序设计中实用的数据结构》（即本书）。按照“数据结构 + 算法 = 程序”的定义，《新编实用算法分析与程序设计》和《程序设计中实用的数据结构》属于程序设计的知识范畴，而《程序设计中常用的计算思维方式》和《程序设计中常用的解题策略》分别从思维方式和行为特征这两个角度给出解决问题的基本方法，属于基于算法和数据结构知识的编程能力范畴。上述所有著作贴近实际，结合实例进行阐释，使读者对什么是程序设计，为什么要学习程序设计，怎样进行程序设计等诸多问题产生浓厚的兴趣和探索意识。因此可以说，这 4 部著作从知识、能力和情感体验 3 个维度，建立了程序设计学习的基本框架。

数据结构是计算机软件专业的核心课程，也是编程人员必学的知识。现在关于数据结构的书很多，章节结构大致雷同，具体描述也大同小异。为了避免本书落入“千篇一律”的俗套，我们按照数据结构知识的分类，将全书分为 3 篇。

上篇“讨论线性表”根据线性结构存储方式的不同，讲解了顺序存储结构和链式存储结构，并根据元素存取方式的不同，讲解了栈和队列两种特殊的线性结构，另外还介绍了线性结构的如下 3 个重要的应用。

(1) 排序与顺序统计。探讨了贪心策略中的序关系和分治策略在排序中的应用，并在排序知识的基础上阐释了顺序统计的有关方法。

(2) 散列表。介绍了如何使用散列函数确定关键字存储位置，并介绍了在多个关键字的运算结果相同的情况下，如何采用链式存储结构的分离链接法和采用顺序存储结构的开放定址法解决“冲突”。

(3) 后缀数组。介绍了一种专用于字符串处理的重要工具。

中篇“讨论树型问题”详细阐释了什么是树，树有哪些类型，如何通过不同的遍历规则将非线性结构的树转化为线性序列，程序中可采用哪些数据类型存储树，这些数据类型与树操作的效率之间有什么样的关系，怎样通过树形式的并查集实现集合间的合并，以及在哪个集合中查找元素。这一篇着重阐释了结构简单、操作方便且便于二分处理数据的二叉树，介绍了普通有序树转化为等价二叉树的基本方法和如下 4 种特殊类型的二叉树。

- (1) 用二叉堆实现元素入队和按照优先级出队的优先队列。
- (2) 在数据通信等方面有广泛用途的最优二叉树。
- (3) 用于区间处理和数据统计的线段树。
- (4) 用于高效查找数据的 3 种二叉查找树, 即二叉排序树、静态二叉排序树和平衡树。

下篇“讨论图型问题”详尽地阐释了图的基本概念和简单分类, 介绍了图的 4 种存储方式(相邻矩阵、邻接表、关联矩阵和邻接压缩表)和 2 种遍历规则(宽度优先遍历和深度优先遍历), 深入讨论了图的连通性问题, 分别给出了在有向图中计算强连通分量和传递闭包以及在不向图中计算割点、桥和双连通分量的基本思路, 并重点讲解了计算最小生成树和最短路径的策略。二分图和引入流量因素的单源单汇简单有向图(网络流图)是两类重要的图, 我们给出了求解二分图最大匹配以及在加权二分图中计算权和最大匹配的基本方法, 并分别叙述了应对不同网络流问题的算法。

在上述 3 篇的论述中, 我们尽力凸显如下特色。

(1) 与时俱进。尽可能将近年来数据结构研究上涌现的新知识、新成果纳入本书, 并且以贴近实际、深入浅出的形式展现给读者。

(2) 注重实际应用能力。通过实例教会读者, 怎样寻找问题中各对象之间的关系, 确定数学模型所使用的逻辑结构; 怎样实现对各个对象的操作, 即确定数据所采用的存储结构; 怎样融数据类型与定义在该类型上的运算于一体, 为面向对象的程序设计方法奠定基础。

(3) 掌握并联两种数据结构的映射方法。教会读者在运算对象需要同时采用两种数据结构时, 怎样将两种数据结构的元素一一对应起来, 使得当一种数据结构中的关键元素发生变化时, 另一种数据结构的相应元素也随之变化。

(4) 辩证地看问题。使读者既能了解线性结构这种基本数据结构的局限性, 又可以掌握采用基本数据结构实现复杂非线性结构的基本方法, 还进一步探索了用基本数据结构取代非线性结构的简化途径。如果说基本数据结构到非线性结构是一种提高, 那么从非线性结构回到基本数据结构却不是一种简单的回归, 而是一种螺旋式发展。对于基本数据结构的再次运用, 是基于对整个数据结构知识的一种深刻领悟, 在进一步理解和运用基本数据结构的过程中, 可使我们的知识与思想得到升华。

(5) 精益求精。随着试题难度和问题规模的不断提升, 对算法时空效率的要求愈来愈高, 数据结构对算法效率的制约力也愈来愈大。因此, 本书旨在引导读者充分考虑数据结构对程序效率所产生的影响, 在选择数据结构时, 尽可能采用“扬长避短”的原则并结合“取长补短”的方法, 更加有效地优化算法。

为了使选用本教程的辅导教师容易教, 参与竞赛培训的学生容易学, 读者在解题查阅时容易懂, 本书在表述上采取了如下 4 条措施。

- (1) 对于每个比较难懂的概念, 都举实例加以说明; 对于每个比较抽象的定理, 都提供具体的应用例证帮助理解; 对于每个经典的算法, 都有清晰的程序流程给予示范。
- (2) 大量运用图表, 使得概念、定理和算法的由来更具体、形象和直观。
- (3) 定理证明尽量采用初等数学知识, 公式推导尽量浅显和详细。

(4) 为了使程序样例不受单一程序设计语言的限制, 不让某种特定语言的特殊性掩盖算法的本质内容, 书中采用了一种结构清晰、移植性强且贴近自然语言的类程序设计语言。无论你具备 Pascal、C、true.bascal 等哪一种语言的基础, 都可以比较轻松地读懂其语义。本书之所以未给出可直接编译运行的程序代码, 除了考虑到语言的通用性外, 还有一个重要的原因, 就是想给读者留出动手编程的空间, 避免无意义的“依葫芦画瓢”。

最后, 由衷感谢复旦大学 ACM 队的舒天民同学。他花费了数月时间, 字斟句酌地查找书中的瑕疵, 反复测试所有的程序示例, 并编程解答了其中部分试题。当我们向他表示感谢时, 他却说, 这本书使他学到了许多过去闻所未闻的数据结构新知识, 编程解题的视野更加开阔了。当然我们自知, 这本书不可能是完美的, 更不可能涵盖全部的数据结构知识和应用, 但在数据结构知识的开发上, 可为初学者或浅尝过这门课程的读者提供十分新颖和有用的信息。但愿你能够像舒天民同学一样, 通过本书的学习拓展新知, 获取灵感, 开发睿智, 提升能力。

王建德 吴永辉

2011/1/10

目 录

上篇 讨论线性表

第 1 章 数组	3
1.1 数组的基本概念	3
1.1.1 数组是一种顺序存储结构	3
1.1.2 数组是程序设计中使用时频率最高的数据类型	4
1.2 优化数组的存储方式	6
1.2.1 规则矩阵的压缩存储	6
1.2.2 稀疏矩阵的压缩存储	7
1.2.3 01 矩阵的压缩存储	11
1.3 排序与顺序统计	14
1.3.1 排序的基本概念	14
1.3.2 计数排序与贪心策略	15
1.3.3 采用“二分”策略的排序方法	21
1.3.4 顺序统计的基本方法	28
第 2 章 链式存储结构	34
2.1 链表的基本概念	34
2.1.1 单链表	34
2.1.2 循环链表	35
2.1.3 双向链表	35
2.2 链表的基本运算	35
2.2.1 构建单链表	36
2.2.2 插入操作	36
2.2.3 删除操作	36
2.2.4 读取操作	36
2.3 链表的应用	37
第 3 章 两种存取方式特殊的线性表	43
3.1 “后进先出”的栈	43

3.1.1 栈的基本运算	43
3.1.2 栈的应用	44
3.2 “先进先出”的队列	52
3.2.1 队列的基本运算	52
3.2.2 队列的应用	54
第 4 章 散列技术	59
4.1 散列表	59
4.2 散列函数的设计	59
4.3 消除冲突的基本方法	61
4.3.1 使用开放寻址法消除冲突	61
4.3.2 使用分离链接法消除冲突	67
第 5 章 后缀数组	71
5.1 后缀数组的基本概念	71
5.2 采用倍增算法求解 <i>rank</i> 数组	73
5.3 利用 <i>rank</i> 数组计算最长公共前缀	74
5.3.1 计算最长公共前缀是一个典型的 RMQ 问题	75
5.3.2 计算最长公共前缀的基本方法	75
5.4 后缀数组的应用	78
5.4.1 利用后缀数组处理单个字符串	78
5.4.2 两个字符串的公共子串问题	87
5.4.3 多个字符串共享子串的问题	90
上篇小结	97

中篇 讨论树型问题

第 6 章 树的基本概念和遍历规则	101
6.1 树的递归定义	101
6.2 节点的分类	101

6.3 有关度的定义	101	压缩	137
6.4 树的深度(高度)	102	9.3 合并两个元素所在的集合	138
6.5 森林	102	第 10 章 堆	143
6.6 有序树和无序树	102	10.1 二叉堆的概念	143
6.7 树的表示方法	103	10.2 在插入或删除节点时维护堆性质	144
6.8 树的遍历规则	103	10.2.1 插入节点	144
6.8.1 先根次序遍历树	103	10.2.2 删除最小值元素	144
6.8.2 后根次序遍历树	104	10.3 建堆	145
第 7 章 树的存储结构	105	10.4 堆排序	146
7.1 采用数组存储入边信息	105	第 11 章 最优二叉树	154
7.1.1 存储无权树的入边信息	105	11.1 最优二叉树的基本概念	154
7.1.2 存储加权树的入边信息	106	11.2 构造最优二叉树	155
7.2 采用数组存储所有儿子的地址信息	108	第 12 章 线段树	160
7.2.1 采用整数存储儿子的数组下标	108	12.1 线段树的基本概念	160
7.2.2 采用指针存储儿子的地址	109	12.1.1 用于区间运算的线段树	160
7.3 采用邻接表存储出边信息	110	12.1.2 用于数据统计的线段树	161
7.3.1 采用数组存储方式的邻接表	110	12.1.3 线段树的数据结构	162
7.3.2 采用单链表存储方式的邻接表	114	12.2 线段树的基本操作	162
7.4 无根树的一般存储方式	116	12.2.1 建立线段树	162
第 8 章 二叉树	123	12.2.2 在区间内插入线段或数据	163
8.1 二叉树的基本概念和存储结构	123	12.2.3 删除区间内的线段或数据	164
8.1.1 二叉树的基本概念	123	12.2.4 计算区间内的线段或数据 状态	165
8.1.2 二叉树的存储结构	125	12.3 线段树在静态统计问题上的应用	165
8.2 将普通有序树和森林转换成对应的 二叉树	128	12.4 线段树在动态统计问题上的应用	168
8.2.1 将普通有序树转换成对应的 二叉树	128	第 13 章 二叉查找树	176
8.2.2 将普通有序树组成的森林转 换成对应的二叉树	129	13.1 二叉排序树	176
8.3 二叉树的遍历	130	13.1.1 二叉排序树的基本概念	176
8.3.1 前序遍历	130	13.1.2 二叉排序树的基本操作	177
8.3.2 中序遍历	130	13.2 静态二叉排序树	180
8.3.3 后序遍历	131	13.2.1 静态二叉排序树的特征	180
8.3.4 由两种遍历确定二叉树结构	133	13.2.2 静态二叉排序树的构造方法	180
第 9 章 并查集	135	13.2.3 在静态二叉排序树上进行数 据统计	181
9.1 并查集的基本概念	135	13.3 子树大小平衡树(SBT)	186
9.2 查找元素所在树的根节点并进行路径 压缩	137	13.3.1 SBT 的性质	186

13.3.2 旋转	186
13.3.3 动态维护 SBT 的平衡特性	191
13.3.4 SBT 的基本操作	196
中篇小结	205

下篇 讨论图型问题

第 14 章 图的基本概念及其存储结构	209
14.1 图的基本概念	209
14.2 图的简单分类	211
14.2.1 无向图和有向图	212
14.2.2 无权图和加权图	212
14.2.3 稀疏图和稠密图	212
14.2.4 完全图和补图	212
14.2.5 树和森林	213
14.2.6 图的生成树和生成森林	213
14.2.7 平面图	214
14.2.8 二分图	214
14.2.9 相交图和区间图	214
14.3 图的存储结构	215
14.3.1 存储节点间相邻关系的相 邻矩阵	215
14.3.2 存储边信息的 3 种数据 结构	217
第 15 章 图的遍历及其应用	220
15.1 广度优先遍历 (BFS 算法)	220
15.1.1 BFS 算法的基本概念	220
15.1.2 BFS 算法的应用	222
15.2 深度优先遍历 (DFS 算法)	233
15.2.1 DFS 的基本概念	233
15.2.2 在 DFS 遍历过程中记录节点 颜色变化的时间	239
15.2.3 根据节点颜色对边进行分类	240
15.2.4 分析 DFS 森林的结构	242
15.2.5 使用 DFS 算法进行拓扑 排序	244
15.2.6 使用 DFS 算法计算欧拉 回路	248

第 16 章 有向图的强连通分量和传递 闭包	255
16.1 判定仙人掌图	255
16.2 计算强连通分量	259
16.3 传递闭包的应用	266
第 17 章 无向图的连通性分析	271
17.1 计算节点的 low 函数	271
17.2 计算连通图的割点和桥	272
17.2.1 计算连通图的割点	272
17.2.2 计算连通图的桥	273
17.3 计算双连通子图	275
17.4 分析连通图的连通程度	276
17.4.1 连通图的顶连通度	277
17.4.2 连通图的边连通度	278
第 18 章 最小生成树	279
18.1 基本概念	279
18.2 最小生成树的应用价值	279
18.3 最小生成树的计算策略	281
18.4 计算最小生成树的两种算法	281
18.4.1 Kruskal 算法	282
18.4.2 Prim 算法	283
18.5 最小生成树算法的应用实例	285
第 19 章 加权图的单源最短路径问题	293
19.1 基本概念	293
19.1.1 单源算法是高效解决所有最 短路径问题的基础	293
19.1.2 负权回路影响单源最短路径 的计算	294
19.1.3 松弛技术是单源算法的核心	295
19.2 求解单源最短路径问题的 3 种算法	296
19.2.1 Dijkstra 算法	296
19.2.2 Bellman-Ford 算法	298
19.2.3 SPFA 算法	299
19.3 单源最短路径算法的应用实例	301
第 20 章 二分图的匹配问题	310
20.1 基本概念	310

20.1.1	图的匹配概念	310	21.3.1	Dinic 算法的基本思路	334
20.1.2	二分图的概念和判定方法	311	21.3.2	Dinic 算法的基本流程	335
20.2	计算无权二分图的最大匹配	315	21.4	利用最大流最小割定理解题	339
20.2.1	匈牙利算法的基本思路	316	21.4.1	割的概念	339
20.2.2	匈牙利算法的基本流程	316	21.4.2	最小割的计算方法和应用实例	340
20.2.3	匈牙利算法的应用实例	317	21.5	计算多源多汇网络的可行流	346
20.3	计算带权二分图的最佳匹配	320	21.6	网络增加容量下界因素后的流量计算问题	348
20.3.1	最佳匹配的概念	320	21.6.1	求容量有上下界的网络的最大流	348
20.3.2	KM 算法的基本思路	321	21.6.2	求容量有上下界的网络的最小流	353
20.3.3	KM 算法的基本流程和应用实例	323	21.7	网络增加费用因素后的流量计算问题	358
第 21 章	最大流问题	327	21.7.1	计算最小费用最大流	359
21.1	基本概念	327	21.7.2	计算容量有上下界的网络的最小费用最小流	364
21.2	在可增广路径的基础上计算最大流	329	下篇小结		370
21.2.1	可增广路径的基本概念	329			
21.2.2	基于最大流定理上的最大流算法	334			
21.3	按层次计算最大流的 Dinic 算法	334			

Part 1

上 篇

讨论线性表

本篇内容

- 第1章 数组
- 第2章 链式存储结构
- 第3章 两种存取方式特殊的线性表
- 第4章 散列技术
- 第5章 后缀数组

线性结构是由有限个数据元素组成的有序集合，其中每个数据元素都有一个数据项或者多个数据项。这种数据结构是最简单、最常用的，其特征有以下几点。

- **均匀性。**同一线性表中各数据元素的数据类型一致且数据项数相同。例如，字符串是一个线性结构，该结构中每一个数据元素都为单个字符；又如，学生成绩表中的每个数据元素都含学生姓名、学号、若干学科成绩等数据项，因此也是一个线性结构。
- **有序性。**表中数据元素之间的相对位置是线性的，即存在唯一的“第一个”和“最后一个”数据元素。除第一个和最后一个外，其他元素前后均只有一个数据元素（直接前趋和直接后继）。例如，字符串中的字符和学生成绩表中的前后元素间就存在着这种“一一对应”的关系。

本篇根据线性结构存储方式的不同，讲解了以下两种存储结构。

- (1) 顺序存储结构，即所谓的数组。
- (2) 链式存储结构，简称链表，包括单链表、双链表或循环链表。

根据元素存取方式的不同，本篇还讲解了以下两种特殊的线性结构。

- (1) 栈，即“先进后出”的线性表。
- (2) 队列，即“先进先出”的线性表。

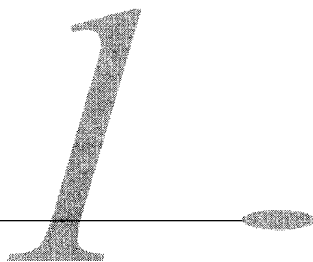
这两种数据结构通常用数组实现。另外，本篇还介绍了线性结构的如下3个重要的应用。

- (1) 排序与顺序统计。本篇介绍了两种特别重要的排序方法：计数排序和快速排序，探讨了贪心策略中的序关系和分治策略在排序中的应用，并在排序知识的基础上阐释了顺序统计的有

关方法。

(2) 散列表。一种对关键字做某种运算后直接确定其存储位置，并可在常数平均时间内执行插入、删除和查找的技术。对于多个关键字运算结果相同的情况，即产生了所谓的“冲突”，本篇介绍了两种方法解决“冲突”：采用链式存储结构的分离链接法和采用顺序存储结构的开放定址法。

(3) 后缀数组。一种专用于处理字符串这种特殊类型数组的重要工具。



数组是程序设计中使用的最多的数据结构，无论是线性结构的还是非线性结构的数据对象，都可用数组存储。

本章将揭示数组的物理地址与逻辑地址之间的线性对应关系，介绍使用数组存储各类数据对象的一般方法，并针对 3 种特殊的二维矩阵（规则矩阵、稀疏矩阵和 01 矩阵）提出压缩存储的优化方式。

在数组的应用中，我们顺便介绍了计数排序、采用二分策略的快速排序和合并排序，以及包含排序过程的贪心策略，阐释了顺序统计的基本方法。这些算法一般都是以数组为其存储结构的。

1.1 数组的基本概念

我们称有限个同类型数据元素的序列为数组，它是一种定长的线性表。若数据元素不再有分量，则称该序列为一维数组；若数据元素为数组，则称该序列为多维数组。例如，下列代码中， A 为一维数组，表长为 $d-c+1$ ，元素类型为 atype ； B 为二维数组，表长为 $(d_1-c_1+1)(d_2-c_2+1)$ ，元素类型为 btype 。

```
var
  A:array[c..d] of atype;
  B:array[c1..d1,c2..d2] of btype;
```

一维数组和多维数组的元素类型和元素个数需在程序的说明部分预先定义，编译系统按照数组说明预留静态数据区，执行过程中不能增减其内存空间。

1.1.1 数组是一种顺序存储结构

在数组存储结构中，元素依次用一组地址连续的存储单元存储，每个元素在序列中的逻辑位置由数组下标指明，映射了元素之间在物理地址上的相邻关系，因此数组亦称顺序存储结构。例如，一维数组 a 的长度为 maxlen ， $a[i]$ 的内存地址为 $\text{loc}(a[i])$ ($1 \leq i \leq \text{maxlen}$)，元素的存储单位长度为 la 。数组元素 $a[1]..a[n]$ 的存储情况如图 1-1 所示。

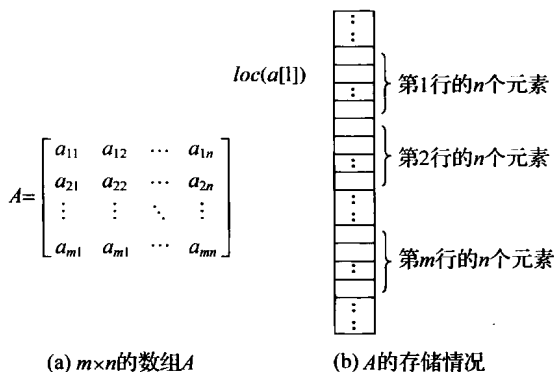
显然， $a[i]$ 的内存地址为 $\text{loc}(a[i]) = \text{loc}(a[1]) + (i-1)la$ ，其中 $\text{loc}(a[1])$ 为存储区的首地址， $(i-1)la$ 为偏移量。一维数组按照下标递增的顺序访问表中元素，即 $a[1] \rightarrow a[2] \rightarrow \cdots \rightarrow a[n]$ 。

存储地址	数组元素	元素在序列中的逻辑位置
$loc(a[1])$	$a[1]$	1
$loc(a[1]) + la$	$a[2]$	2
...		...
$loc(a[1]) + (i-1)la$	$a[i]$	i
...		...
$loc(a[1]) + (n-1)la$	$a[n]$	n
$loc(a[1]) + nla$		
...		
$loc(a[1]) + (maxlen-1)la$		

} 空闲

图 1-1 一维数组 $a[1..n]$ 的存储情况

二维数组一般是按照以行为主的顺序存储在计算机的连续存储空间中的,即先存储数组的第1行再存储数组的第2行,依次类推,每一行的元素按从左到右的顺序存储。例如,图1-2a所示为一个 $m \times n$ 的数组 A ,图1-2b所示为数组 A 在计算机中的存储形式。

图 1-2 数组 A 及其在计算机中的存储形式

设 $loc(a[i, j])$ 为元素 $a[i, j]$ 的内存地址 ($1 \leq i \leq m, 1 \leq j \leq n$), 存储单位长度为 la , 则

$$loc(a[i, j]) = loc(a[1, 1]) + [(i-1)n + j - 1] la$$

其中 $loc(a[1, 1])$ 为存储区的首地址, $[(i-1)n + j - 1]la$ 为偏移量。二维数组按照先行后列的顺序访问表中元素, 即 $a[1, 1] \rightarrow a[1, 2] \rightarrow \cdots \rightarrow a[1, n] \rightarrow \cdots \rightarrow a[i-1, n] \rightarrow a[i, 1] \rightarrow \cdots \rightarrow a[m, n-1] \rightarrow a[m, n]$ 。

在数组中, 元素的下标间接反映了元素的存储地址, 而计算机内存是随机存取的装置, 在数组中存取一个元素只要通过下标计算它的存储地址就行了, 因此在数组中存取任意一个元素的时间都为 $O(1)$ 。从这个意义上讲, 数组的存储结构是一种随机存取的结构。

1.1.2 数组是程序设计中频率最高的数据类型

数组在程序中几乎“无处不在”, 无论是线性结构的还是非线性结构的数据对象, 都可用数

组存储。例如,可以使用数组构造链表中的数据元素和指针。图 1-3a 所示为一个存储动态集合 $\{1, 4, 9, 16\}$ 的双链表 L , 表头为 $head[L]$, 数据元素含关键字字段 key 、前驱指针字段 $prev$ 和后继指针字段 $next$, 首元素的 $prev$ 值和末元素的 $next$ 值为 nil (用 “/” 表示)。双链表 L 也可以用记录类型数组 L 存储, 其中 $L[i].key$ 存储节点 i 的关键字, $L[i].prev$ 和 $L[i].next$ 分别存储节点 i 的前驱指针和后继指针; 或者分别用 3 个一维数组 $key[]$ 、 $prev[]$ 和 $next[]$ 存储链表 L 中每个节点的关键字字段、前驱的数组下标和后继的数组下标。图 1-3b 给出了用数组 L 实现图 1-3a 中所示双链表 L 的情形, 表头的下标为 7。有关链表的知识将在第 2 章中详细介绍。

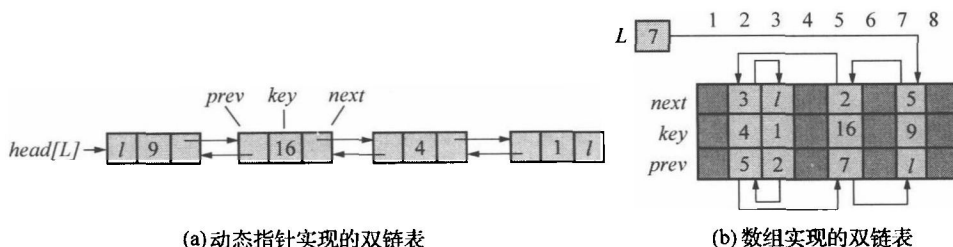


图 1-3 动态指针实现的双链表与数组实现的双链表

数组亦可存储非线性结构的数据对象。例如, 一棵含 n 个节点的完全二叉树可用一维数组 $A[1..n]$ 存储, 其中节点 $A[1]$ 为根, 节点 $A[i]$ ($2 \leq i \leq \lfloor \frac{n}{2} \rfloor$) 的父节点为 $A[\lfloor \frac{i}{2} \rfloor]$, 左儿子为 $A[2i]$, 右儿子为 $A[2i+1]$, 节点 $A[\lfloor \frac{n}{2} \rfloor + 1]$ 至 $A[n]$ 为叶节点。图 1-4 给出了一棵含 14 个节点的完全二叉树。

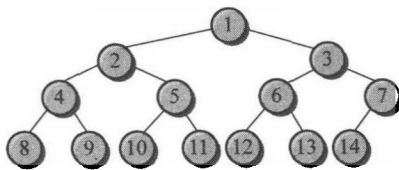


图 1-4 一棵含 14 个节点的完全二叉树的下标定义

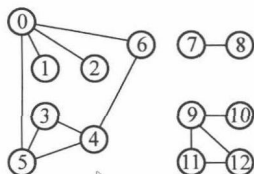
在完全二叉树中增删节点, 可以在一维数组的基础上“原地操作”, 不必另辟数组, 不过其操作必须维护完全二叉树的特征。有关这方面的内容请参阅第 10 章。

数组也可以用来存储图, 例如, 用二维数组 $A[1..n, 1..n]$ 存储图的相邻矩阵是最典型的应用。图 1-5a 给出了一个无向图 G , 其对应的相邻矩阵如图 1-5b 所示。

有关相邻矩阵的知识, 请参阅 14.3.1 节。

如果一个图是稠密图 ($|E|$ 比较接近 n^2) 或需要及时判别任一对节点间是否存在连接边, 通常采用相邻矩阵, 但是如果一个图为稀疏图 ($|E|$ 远小于 n^2) 则通常采用邻接表。邻接表由一个包含 n 个列表的一维数组 Adj 组成, 每个列表对应于图中的一个节点和一条含所有相邻点的链表, 即对于每一个节点 $u \in V$, $Adj[u]$ 链表包含了图中所有与 u 相邻的节点。例如, 图 1-6a 所示无向图

对应的邻接表如图 1-6b 所示。

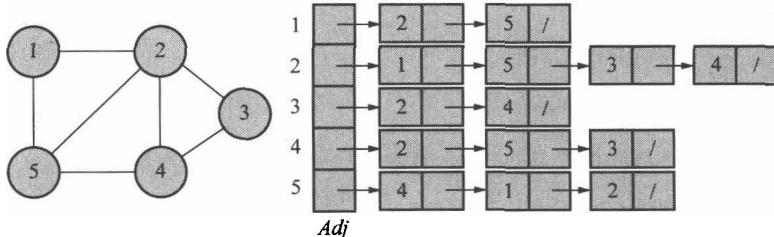


(a) 无向图

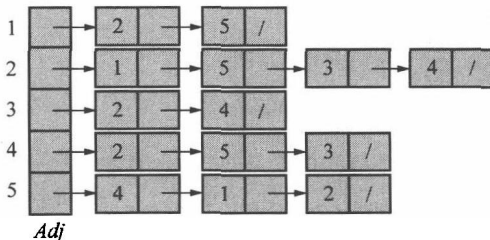
	0	1	2	3	4	5	6	7	8	9	10	11	12
0	0	1	1	0	0	1	1	0	0	0	0	0	0
1	1	0	0	0	0	0	0	0	0	0	0	0	0
2	1	0	0	0	0	0	0	0	0	0	0	0	0
3	0	0	0	0	1	1	0	0	0	0	0	0	0
4	0	0	0	1	0	1	1	0	0	0	0	0	0
5	1	0	0	1	1	0	0	0	0	0	0	0	0
6	1	0	0	0	1	0	0	0	0	0	0	0	0
7	0	0	0	0	0	0	0	0	1	0	0	0	0
8	0	0	0	0	0	0	0	1	0	0	0	0	0
9	0	0	0	0	0	0	0	0	0	0	1	1	1
10	0	0	0	0	0	0	0	0	0	1	0	0	0
11	0	0	0	0	0	0	0	0	1	0	0	1	0
12	0	0	0	0	0	0	0	0	0	1	0	1	0

(b) 对应的相邻矩阵

图 1-5 示例无向图及其对应的相邻矩阵



(a) 无向图



(b) 邻接表

图 1-6 示例无向图及其对应的邻接表

由图 1-6 可见，每个邻接表为一个单链表，当然这个单链表亦可用数组取而代之。设 $d[u]$ 为节点 u 的度， $Adj[u, 1..d[u]]$ 存储与节点 u 相邻的 $d[u]$ 个节点的序号。虽然这种方法的内存开销与相邻矩阵相同，但编程比单链表形式的邻接表方便一些，且运行速度比相邻矩阵要好一些。因为如果在相邻矩阵的基础上计算与 u 节点相连的所有边，则需要遍历 u 行的 n 个元素，而邻接表仅需遍历 u 行的前 $d[u]$ 个元素即可。有关相邻矩阵和邻接表的知识请参阅 14.3.2 节。

1.2 优化数组的存储方式

优化数组存储方式的目的是为了节省内存。下面我们分别针对二维数组中的以下 3 种特殊情况提出压缩存储的优化方式：

- 规则矩阵（元素分布有一定规律）；
- 稀疏矩阵（元素一半以上为无用 0）；
- 01 矩阵（元素值为 0 和 1）。

1.2.1 规则矩阵的压缩存储

规则矩阵是指矩阵中的非零元素都按照如下两种规律分布：

个 7×8 的稀疏矩阵中仅有 8 个非零元素，其余为零元素（图 1-8）。

显然，若这些零元素也存入计算机，则会浪费大量的存储空间。因此，在方便访问每个元素（包括零元素）的前提下，仅存储少量的非零元素，这就是稀疏矩阵的压缩存储。

与规则矩阵的压缩存储不同的是，稀疏矩阵中的非零元素没有规律，不能采用一个一维数组来依次存放，因此稀疏矩阵的压缩存储方式比规则矩阵稍复杂一些。设数组有 m 行 n 列，其中非零元素的个数为 t ，我们建立一个长度为 $t+1$ 的记录类型数组 d 。

$$\begin{bmatrix} 0 & 0 & 3 & 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 9 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 7 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 6 & 0 \\ 0 & 0 & 0 & 2 & 0 & 3 & 0 & 0 \\ 0 & 0 & 3 & 0 & 0 & 0 & 0 & 0 \end{bmatrix}$$

图 1-8 7×8 的稀疏矩阵

```

type                                /*d 数组的元素类型*/
node=record                         /*非零元素的类型*/
    i,j,v:integer;                 /*非零元素的行列序号为 i 和 j，值为 v*/
end;

dtype=array[1.. $\left\lfloor \frac{m \cdot n}{2} \right\rfloor$ ] of node; /*d 数组的元素类型定义*/

arr=array[1..m] of integer;         /*pos 和 num 序列的类型定义*/
arrtype= array[1..m,1..n] of integer; /*稀疏矩阵的类型定义*/

var
d:dtype;                          /*非零元素序列*/
pos,num:arr;                       /*稀疏矩阵第 k 行中非零元素的个数为 num[k]。若该行有非零元素的话，则第 1 个非零元素在 d 数组的序号为 pos[k]，即为 d[pos[k]]*/

```

数组 d 为压缩后的非零元素序列。其中 $k > 1$ 时， $d[k]$ 给出了第 $k-1$ 个非零元素 a_{ij} 的信息，即 $d[k].i = i$ ， $d[k].j = j$ ， $d[k].v = a_{ij}$ ； $k = 1$ 时， $d[1].i$ 和 $d[1].j$ 为数组的行和列数， $d[1].v$ 为非零元素的个数，即 $d[1].i = m$ ， $d[1].j = n$ ， $d[1].v = t$ 。

这样做的目的是为了确定唯一性，因为在稀疏矩阵的下方添加全 0 的若干行或在右方添加全 0 的若干列，得出的 d 序列也是相同的，而指明稀疏矩阵的规模和非零元素个数，则 d 对应的稀疏矩阵是唯一的。例如，用 $d[1..9]$ 存储图 1-8 所示的稀疏矩阵：

```

d[1].i=7  d[1].j=8  d[1].v=8
d[2].i=1  d[2].j=1  d[2].v=3
d[3].i=1  d[3].j=8  d[3].v=1
d[4].i=3  d[4].j=1  d[4].v=9
d[5].i=4  d[5].j=5  d[5].v=7
d[6].i=5  d[6].j=7  d[6].v=6
d[7].i=6  d[7].j=4  d[7].v=2
d[8].i=6  d[8].j=6  d[8].v=3
d[9].i=7  d[1].j=3  d[1].v=3

```

为了避免因访问稀疏矩阵中的某元素而遍历整个 d 序列的盲目性，附设了两个长度与稀疏矩阵行数相同的一维整数数组 pos 和 num ，其中 $num[k]$ 给出了稀疏矩阵第 k 行中非零元素的个数。若该行有非零元素的话，则 $pos[k]$ 给出了其中第 1 个非零元素在 d 数组中的序号，即 $d[pos[k]]$ 存储 k 行第 1 个非零元素的信息。显然