



辽宁省“十二五”普通高等教育本科省级规划教材

普通高校本科计算机专业特色教材精选·算法与程序设计

数据结构（C语言版）（第3版）

秦玉平 马靖善 主编

清华大学出版社

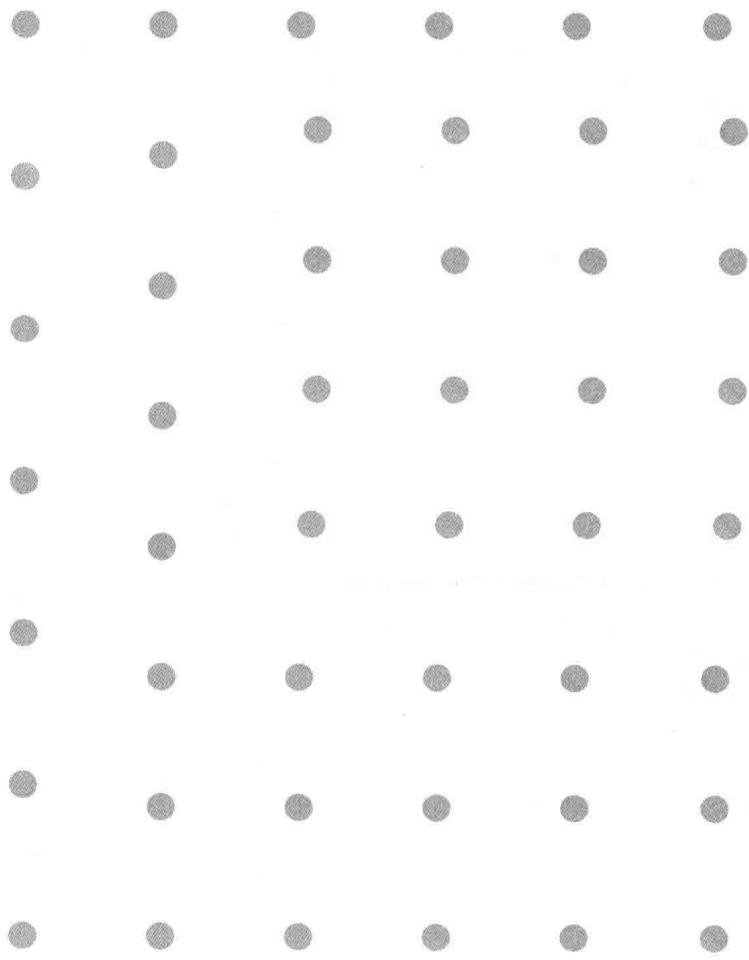


辽宁省“十二五”普通高等教育本科省级规划教材

普通高校本科计算机专业特色教材精选·算法与程序设计

数据结构 (C语言版) (第3版)

秦玉平 马靖善 主编



清华大学出版社
北京

内 容 简 介

数据结构是计算机及其相关专业的核心课程,是计算机程序设计的基础,也是程序员考试和许多高校研究生入学考试的必考科目。

全书共 10 章。第 1 章是数据结构的概述;第 2~10 章分别讨论了顺序表、链表、栈、队列、串、数组、广义表、树、二叉树、图、查找、内部排序、外部排序和文件等基本类型的数据结构。本书中的算法都已经过调试,不必修改就能在 Turbo C 2.0 系统下正常运行。

本书可作为计算机及相关专业的教材,也可作为自学或各种计算机培训班的教材。

本书封面贴有清华大学出版社防伪标签,无标签者不得销售。

版权所有,侵权必究。侵权举报电话:010-62782989 13701121933

图书在版编目(CIP)数据

数据结构: C 语言版/秦玉平,马靖善主编. —3 版. —北京:清华大学出版社,2015

普通高校本科计算机专业特色教材精选·算法与程序设计

ISBN 978-7-302-40251-0

I. ①数… II. ①秦… ②马… III. ①数据结构—高等学校—教材 ②C 语言—程序设计—高等学校—教材 IV. ①TP311.12 ②TP312

中国版本图书馆 CIP 数据核字(2015)第 106329 号

责任编辑:焦 虹 战晓雷

封面设计:傅瑞学

责任校对:白 蕾

责任印制:杨 艳

出版发行:清华大学出版社

网 址: <http://www.tup.com.cn>, <http://www.wqbook.com>

地 址:北京清华大学学研大厦 A 座 邮 编:100084

社 总 机:010-62770175 邮 购:010-62786544

投稿与读者服务:010-62776969, c-service@tup.tsinghua.edu.cn

质量反馈:010-62772015, zhiliang@tup.tsinghua.edu.cn

课件下载: <http://www.tup.com.cn>, 010-62795954

印 刷 者:北京富博印刷有限公司

装 订 者:北京市密云县京文制本装订厂

经 销:全国新华书店

开 本:185mm×260mm 印 张:17.75 字 数:400 千字

版 次:2005 年 10 月第 1 版 2015 年 8 月第 3 版 印 次:2015 年 8 月第 1 次印刷

印 数:1~2000

定 价:34.50 元

产品编号:064712-01

第3版 前言

PREFACE

数据结构是计算机专业的一门重要的专业必修课,是绝大多数高校招收计算机专业及相关专业硕士研究生的必考科目之一。

本课程主要研究数据在计算机中的存储和操作。它涉及一系列较为实用的算法,这些算法在实际的程序设计中是非常有用的。但这门课程内容丰富、学习量大,其算法又十分抽象。经过多年的教学实践,总结出一些该课程的课程特点和教学方法,为此,我们编写了这部教材,以满足广大同学的要求和计算机教学的需要。

本书在《数据结构(C语言版)》(第2版)的基础上,秉承注重实用和便于学习的宗旨,删除了动态存储管理一章的内容,调整了部分章节的内容和结构,优化了算法的代码并给出了详细注释,在文字表述和习题设置等方面又进行了反复的推敲并做了相应修改。

本书共10章,第1章为概述,主要介绍数据结构的简单发展史、基本概念和算法的描述与分析方法;第2章为线性表,主要介绍顺序表和各种链表的存储表示与实现;第3章为特殊线性表,主要介绍栈、队列和串的存储表示与实现;第4章为数组和广义表,主要介绍数组和广义表的存储表示与实现;第5章为树和二叉树,主要介绍二叉树的基本知识、性质、存储、遍历及其应用;第6章为图,主要讨论图的基本概念、存储、遍历及其应用;第7章为查找,主要介绍静态查找、动态查找和散列表;第8章为内部排序,分别介绍几种常用的排序算法及性能;第9章为外部排序,主要研究在内存和外存之间如何调动和组织数据进行排序;第10章为文件,主要讨论文件在外存储器中的表示方法和各种运算的实现方法。

本书的算法都用C语言函数实现,无须任何修改即可被其他函数调用。本书结构合理,内容紧凑,知识连贯,逻辑性强。为了使读者更好地掌握所学内容,各章末配有精选的大量习题,可使读者快速熟悉和掌握所学的知识。本书既适合作计算机专业的本、专科教材,也适合用作与计算机学科相关的其他专业的教材。

本书第1和第5章由马靖善编写;第2~4章由秦玉平编写;第6章由

周连秋编写;第7章由王丽君编写;第8章由冯佳昕编写;第9章由沈泽刚编写;第10章由李春杰编写。全部书稿由范立南和秦玉平审校,所有算法由马靖善和王丽君调试。

在本书编写过程中,编者参考了大量有关数据结构的书籍和资料,在此对这些参考文献的作者表示感谢。

由于编者水平有限,书中难免存在错误和不当之处,恳请广大读者批评指正,以便再版时改进。

注:目录中带**的章节可选讲;带*的章节为专科选学。

本书受2014年辽宁省普通高等学校本科教育教学改革研究项目(UPRP20140661)资助。

编 者

2015年5月

目 录

CONTENTS

第 1 章 概述	1
1.1 数据结构的发展	1
1.2 基本概念	2
1.3 算法描述与分析	5
习题 1	10
第 2 章 线性表	13
2.1 线性表的定义及基本操作	13
2.1.1 线性表的基本概念	13
2.1.2 线性表的基本操作	14
2.2 顺序表	15
2.2.1 顺序表的定义	15
2.2.2 基本操作在顺序表上的实现	16
2.3 链表	19
2.3.1 单链表的表示和实现	19
2.3.2 双链表的表示和实现	27
2.3.3 循环链表的表示和实现	32
**2.3.4 静态链表的表示与实现	39
习题 2	44
第 3 章 特殊线性表	49
3.1 栈	49
3.1.1 栈的定义及其基本操作	49
3.1.2 顺序栈的表示和实现	50
* 3.1.3 链栈的表示与实现	55

3.2 队列	58
3.2.1 队列的定义及其基本操作	58
3.2.2 顺序队列的表示和实现	59
3.2.3 链队列的表示和实现	64
3.3 串	66
3.3.1 串的定义及其基本操作	66
3.3.2 顺序串的表示与实现	67
* 3.3.3 链串的表示和实现	72
** 3.3.4 串的模式匹配	78
习题 3	84
第 4 章 数组和广义表	87
4.1 数组	87
4.1.1 数组的定义及基本操作	87
4.1.2 数组存储结构	88
* 4.1.3 矩阵的压缩存储	89
* 4.2 广义表	103
4.2.1 广义表的定义和基本操作	103
4.2.2 广义表的存储	104
习题 4	109
第 5 章 树和二叉树	113
5.1 树的定义和基本操作	113
5.1.1 树的定义和基本术语	113
5.1.2 树的基本操作	114
5.2 二叉树的定义和性质	115
5.2.1 二叉树的定义	115
5.2.2 二叉树的性质与结论	115
5.3 二叉树的存储	118
5.3.1 二叉树的顺序存储结构	118
5.3.2 二叉树的链式存储结构	122
5.4 二叉树的遍历及应用	123
5.4.1 二叉树的遍历	123
5.4.2 二叉树递归遍历应用举例	127
* 5.4.3 二叉树的非递归遍历	130
* 5.5 线索二叉树	132
5.5.1 线索二叉树的定义	132
5.5.2 线索化处理算法	133

5.6 树和森林	137
5.6.1 树的存储结构	137
5.6.2 树、森林与二叉树之间的转换	140
5.6.3 树和森林的遍历	141
5.7 赫夫曼树及其应用	142
5.7.1 赫夫曼树	142
5.7.2 赫夫曼编码	145
习题 5	148
第 6 章 图	153
6.1 图的基本概念	153
6.2 图的存储	156
6.2.1 邻接矩阵	156
6.2.2 邻接表与逆邻接表	158
**6.2.3 十字链表	161
**6.2.4 邻接多重表	162
6.3 图的遍历	163
6.3.1 深度优先搜索及其生成树	163
6.3.2 广度优先搜索及其生成树	164
6.4 最小生成树	166
6.4.1 Kruskal 算法	166
6.4.2 Prim 算法	167
6.5 图的应用	168
6.5.1 拓扑排序	169
6.5.2 关键路径	170
*6.5.3 最短路径	173
习题 6	175
第 7 章 查找	177
7.1 查找的基本概念	177
7.2 静态查找表	178
7.2.1 顺序查找	179
7.2.2 二分查找	181
*7.2.3 分块查找	184
7.3 动态查找表	186
7.3.1 二叉排序树	186
7.3.2 平衡二叉树	193
**7.3.3 B 树与 B+ 树	199
**7.3.4 键树	201

7.4 散列表	203
7.4.1 散列表的定义	203
7.4.2 散列函数的构造方法	204
7.4.3 处理冲突的方法	206
* 7.4.4 散列表的查找与分析	212
习题 7	213
第 8 章 内部排序	217
8.1 排序的基本概念	217
8.2 插入排序	220
8.3 交换排序	228
8.4 选择排序	232
8.5 归并排序	239
8.6 计数排序	242
8.7 基数排序	243
8.8 各种排序方法的综合比较	246
习题 8	247
第 9 章 外部排序	251
9.1 外存储器简介	251
9.2 外部排序的方法	252
9.3 多路归并排序	254
9.4 置换-选择排序	256
9.5 最佳归并树	258
习题 9	259
第 10 章 文件	261
10.1 文件的基本概念	261
10.2 顺序文件	263
10.3 索引文件	265
10.3.1 ISAM 文件	266
10.3.2 VSAM 文件	267
10.4 散列文件	268
10.5 多关键字文件	270
10.5.1 多重表文件	270
10.5.2 倒排文件	270
习题 10	271
参考文献	273

第 1 章

概 述

CHAPTER

1.1 数据结构的发展

计算机从开始发展到现在已经深入到社会的各个领域,应用越来越广泛。人们已经离不开计算机了,计算机的应用与普及对人类的生产和生活产生了巨大的影响。在计算机上开发的各种软硬件产品层出不穷。就编程而言,一是要提高程序的执行效率,二是要想办法降低存储空间需求。对同一个问题,每个人的理解不完全一样,编程的方法和手段也不尽相同,但是,追求完美是人们的共识。那么,如何才能做到这一点呢?为了编写出“好”的程序,必须分析待处理对象的特性以及各个对象之间的关系。这就是“数据结构”这门学科形成和发展的背景。

数据结构是一门综合性的学科,它是在 20 世纪 60 年代末期开始形成和发展起来的。1968 年,在美国一些大学的计算机系开始开设了这门课程。当时,数据结构几乎和图论,特别是表、树等理论为同义语。随后,“数据结构”这个概念被扩充到包括网络、集合代数论、格、关系等方面,从而变成了现在称为“离散数学”的内容。1968 年,美国学者唐·欧·克努特(Knuth D. E)教授开创了数据结构的最初体系。他所著的《计算机程序设计技巧》第一卷《基本算法》是一本较系统地阐述数据的逻辑结构和存储结构及其操作的著作。从 20 世纪 60 年代末到 70 年代初,出现了大型程序,软件也相对独立,结构程序设计成为程序设计方法学中的主要内容。人们越来越重视数据结构,认为程序设计的实质是对确定的问题选择一种好的结构及合适的算法。从 20 世纪七八十年代开始,各种版本的数据结构方面的著作和教材相继问世。

目前,“数据结构”已不仅是各个高校计算机专业的核心课程,也是其他相关专业的必修课程之一。它是集数学、计算机硬件和软件于一身的综合学科。它以数学方法为基础,研究如何更加合理有效地组织数据,编制出高质量的程序。“数据结构”作为一门较新的学科还在不断地发展中,也在不断地受到专业人士的关注。

很多计算机工作者认为,程序设计的实质就是通过分析问题,确定数学模型和算法,然后再选择一个好的数据结构,即

程序 = 算法 + 数据结构

因此,建议读者在学习“数据结构”这门课程时应做到以下4点:

- (1) 牢记典型算法的基本思想和求解步骤。
- (2) 注意各种结构之间的相互联系和区别。
- (3) 按照算法编制程序,上机调试。
- (4) 分析遇到的问题,并研究解决问题的方法。

学习这门课程,要有前驱课程的准备。本教材中的算法都是用C语言编写的,不用做改动就能上机运行,所用的编程环境为Turbo C。建议读者在学习过程中一定要牢牢掌握数据结构中的一些基本概念和典型算法的基本思想,并注重上机实验,加深对所学算法的理解和掌握,同时也有助于提高C语言编程能力。

1.2 基本概念

要学好数据结构这门课程,必须要明确各种概念及相互之间的联系。本节只介绍一些主要的基本概念,其他的相关术语将在后续章节中介绍。

1. 数据

数据(data)是对客观事物的符号表示。在计算机学科中,数据是指所有能输入到计算机中,并能被计算机程序所处理的符号的总称。因此,除了日常所习惯的符号,比如数字构成的整数和实数、字母构成的串之外,还有标点符号、键盘符号,甚至图形、图像、声音等也都是数据的表示形式。

2. 数据元素和数据项

数据元素(data element)是描述数据的基本单位。**数据项(data item)**是描述数据的最小单位。在计算机中表示数据时,都是以一个数据元素为单位。比如,一个整数表示的一个数据元素、一条记录表示的一个数据元素等。而用一条记录表示一个数据元素时,这条记录中一般还会有多个描述记录属性的分项,称为数据项。比如,描述一辆自行车的记录中可以包括车型、颜色、出厂日期、材质等分项,描述一个班级的记录可以包括班级名、人数、男女生比例、教室、班委会成员和团支部成员等小项。数据项是具有独立含义的最小标识单位。

3. 数据对象

数据对象(data object)是性质相同的数据元素的集合。数据是非常广泛的一个概念,是用来描述千变万化的客观世界的。从中取出一部分,而这部分元素都有共同的性质,这些数据元素就可以组成一个数据对象,实质上,数据对象是数据的一个子集。比如,整数集、字符集、由记录组成的文件等。

4. 数据结构

数据结构(data structure)由数据和结构两部分组成。其中,数据部分是指数据元素的集合;结构就是关系,结构部分就是指数据元素之间关系的集合。所以,笼统地说数据

结构是指数据元素的集合及数据元素之间关系的集合。概括地讲,数据结构是指相互之间有一种或多种特定关系的数据元素的集合。在计算机上要处理数据,就要保存数据及它们之间的关系。在这里,关系就是结构。通常,数据之间有如下4种关系。

(1) **集合结构**: 结构中的数据元素除了“同属于一个集合”的关系外,再无其他关系。

(2) **线性结构**: 结构中的数据元素之间存在着“一对一”的邻接关系。比如,生产过程中的流水作业、体育比赛中的接力赛跑等。

(3) **树状结构**: 结构中的数据元素之间存在着“一对多”的关系。比如,多米诺骨牌、政府组织机构等。

(4) **图状结构或网状结构**: 结构中的数据元素之间存在着“多对多”的关系。比如,城市交通图、电话网等。

图1.1是上述4种结构的关系图。其中,树状结构和图状结构又称为非线性结构。由于集合中数据元素之间关系是非常松散的,因此,常用其他几种结构来描述集合。数据结构依据抽象描述方式和机内存储形式可分为逻辑结构和物理结构。

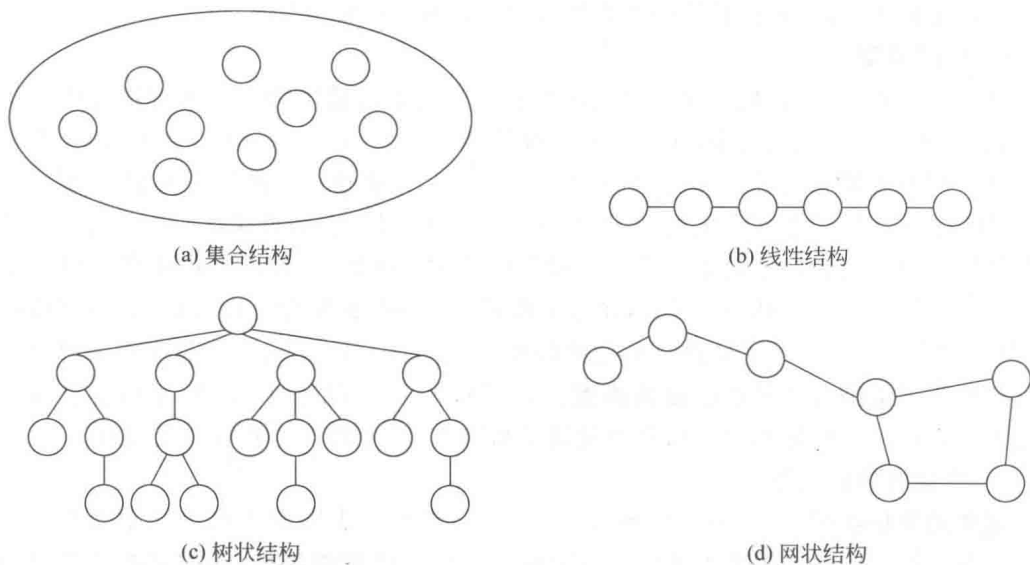


图 1.1 4 种基本结构关系图

(1) **逻辑结构**(logical structure)是以抽象的数学模型来描述数据结构中数据元素之间的逻辑关系。通常用二元组来描述这种关系:

$$\text{Data_Structure} = (D, R)$$

其中, D 是数据元素的有限集, R 是 D 上的关系的有限集。例如,复数的逻辑结构可以用如下的二元组来描述:

$$\text{Complex} = (D, R)$$

其中, $D = \{x | x \text{ 为实数}\}$, $R = \{ \langle x, y \rangle | x, y \in D, x \text{ 为实部}, y \text{ 为虚部} \}$ 。其中, $\langle x, y \rangle$ 表示的是一个有序偶,即 x 和 y 有顺序关系, $\langle x, y \rangle$ 不等价于 $\langle y, x \rangle$ 。常用 (x, y) 表示无序偶,即 x 和 y 无顺序关系, (x, y) 等价于 (y, x) 。若在 D 中任取两个实数,比如5和

3, 则通过关系 $\langle 5, 3 \rangle$ 可以表示一个复数 $5 + 3i$, 而通过关系 $\langle 3, 5 \rangle$ 可以表示一个复数 $3 + 5i$ 。

(2) **物理结构**(physical structure) 又称存储结构, 是数据结构在计算机内的存储表示, 也称内存映象。一个存储在内存中的数据元素又称为**结点**, 数据元素中的每个数据项又称为**域**。因此, 数据元素或结点可以看成是数据元素在计算机中的映象。数据元素可以存放到内存的某个单元, 那么如何存储数据元素之间的关系呢? 在计算机内部主要是用顺序存储和链式存储这两种结构来表示数据元素之间的逻辑关系。顺序存储结构的特点是用物理地址相邻接来表示数据元素在逻辑上的相邻关系; 链式存储结构的特点是逻辑上相邻接的数据元素在存储地址上不一定相邻接, 数据元素逻辑上的相邻关系是通过指针来描述的, 常用它来描述树状结构和图状结构在计算机内的存储。除这两种存储结构外, 在计算机内还有索引存储结构和散列存储结构等, 其实质都是通过顺序存储结构和链式存储结构复合而成。

逻辑结构和物理结构是描述数据结构的密不可分的两个方面。任何一个算法的设计取决于选定的逻辑结构, 而算法的实现则取决于依托的存储结构。

5. 数据类型

在客观世界中, 任何数据元素都应该有它自身的取值范围和它所允许进行的运算操作。比如, 车是一个数据对象, 车族中有各种各样的汽车、火车、三轮车、自行车等, 如果给车安装一对翅膀那就变成了飞机, 超出了车的范围。车能进行的操作是安装、行驶、载人、运货、比赛等, 如果让车到天上去飞, 那就太难为它了。因此, **数据类型**(data type) 就是一个值的集合和定义在这个值集上的一组操作的总称。例如, C 语言中的基本整型(signed int), 它的值集是 $-32\ 768 \sim 32\ 767$, 在这个值集上能进行的操作有加、减、乘、除和取余数等, 而在实型(float)上就不能进行取余数的操作。按值的不同特性, 数据类型又可分为不可分解的**原子类型**及可分解的**结构类型**。比如 C 语言中的整型、实型、字符型就属于原子类型, 而数组、结构体和共用体类型就属于结构类型, 可由其他类型构造得到。

6. 抽象的数据类型

抽象的数据类型(Abstract Data Type, ADT) 是指一个数学模型以及定义在该模型上的一组操作。抽象数据类型的定义仅取决于它的一组逻辑特性, 而与其在计算机内部如何表示和实现无关, 即不论其内部结构如何变化, 只要它的数学特性不变, 都不影响其外部的使用。抽象的数据类型可以分为如下 3 种类型:

(1) **原子类型**(atomic data type), 其值是不可分的。

(2) **固定聚合类型**(fixed-aggregate data type), 其值由确定数目的成分按某种结构组成。

(3) **可变聚合类型**(variable-aggregate data type), 其值由不确定数目的成分构成。

一个抽象的数据类型的软件模块通常包含定义、表示和实现 3 个部分。

7. 多型数据类型

多型数据类型(polymorphic data type) 是指其值的成分不确定的数据类型。

1.3 算法描述与分析

要解决实际问题,就是要找出解决问题的方法。要用计算机解决实际问题,就要先给出解决问题的算法,再依据算法编制程序完成要求。所谓**算法** (algorithm)就是对求解问题步骤的一种描述,也称为**算法设计**。描述算法的工具有许多,比如自然语言、框图、计算机语言程序、伪代码、类计算机语言等。本书主要采用C语言函数来描述算法,但在书写算法时可能会忽略一些细节,如变量定义、交换两个变量的值等。读者只要补充这些细节,再添加主函数就可以上机调试算法程序了。

1. 算法的5个重要特性

(1) **有穷性**: 一个算法必须总是(对任何合法的输入值)在执行有穷步之后结束,且每一步都可在有穷的时间内完成。这也是算法与程序的最主要区别,程序可以无限地循环下去,如操作系统的监控程序,在机器启动后就一直在监测着操作者的鼠标动作和输入的命令。

(2) **确定性**: 算法中的每一条指令都必须有明确的含义,不应产生二义性。并且,在任何条件下,算法只有唯一的一条执行路径,即对于相同的输入只能得到相同的输出。

(3) **可行性**: 一个算法是可以被执行的,即算法中的每个操作都可以通过已经实现的基本运算执行有限次来完成。

(4) **有输入**: 根据实际问题需要,一个算法在执行时可能要接收外部数据,也可能无须外部输入。所以,一个算法应有零个或多个输入,这取决于算法本身要实现的功能。

(5) **有输出**: 一个算法在执行完成后,一定要有一个或多个结果或结论。这就要求算法一定要有输出,这些输出是与输入有着某些联系的量。

2. 算法的评价

通常,解决同一个问题,不同的人有不同的想法,即使是同一个人,在不同的时间里可能对同一个问题的理解也不完全相同。而算法是依据个人的理解和想法人为设计出来的求解问题的步骤,不同的人,或同一个人不同的时间里设计出来的算法也不尽相同,那么哪种算法设计得好? 如何评价一个算法的好与不好? 通常,在算法设计时应该考虑如下5个方面。

(1) **正确性**: 这是算法设计的最基本要求,算法应该严格地按照特定的规格说明去设计,要能够解决给定的问题。但是,“正确”一词的含义在通常的用法中有很大的区别,大体上可分为以下4个层次:

- ① 依据算法所编制的程序中不含语法错误。
- ② 程序对于几组输入数据能够得到满足规格说明要求的结果。
- ③ 程序对于经过精心挑选较为苛刻的几组输入数据也能够得到令人满意的结果。
- ④ 程序对于所有符合要求的输入数据都能得到正确的输出。

对于大型软件需要进行专业测试,一般情况下,通常以第三个要求作为衡量算法正确性的标准。

(2) **可读性**: 设计算法的主要目的是解决实际问题,在设计实现一个项目时,往往不

是一个人去独立完成。如果别人看不懂你设计的算法,怎么去交流?又如何依据算法去编制程序?为了达到可读性的要求,在设计算法时,一般要使用有一定意义的标识符给变量、函数等起名,达到“见名知意”。再有,可以在算法的开头或指令的后面加注释,解释算法和指令的功能。

(3) **健壮性**:当输入不合法数据时,算法能作出相应的反应或进行适当的处理,避免带着非法数据执行,导致莫名其妙的结果。

(4) **高效率**:依据算法编制的程序运行速度较快。

(5) **低存储**:依据算法编制的程序运行时所需内存空间较少。

对于一个系统设计人员来说,前三项很容易实现。在使用电脑软件时,人们更加注重软件的运行速度,而后两点恰恰是影响速度的主要因素。

3. 时间复杂性

算法的执行时间需要通过依据算法所编制的程序在计算机上运行时所消耗的时间来度量,而度量一个程序执行的时间通常有两种方法。

(1) **事后统计法**:就是统计依据算法所编制的程序在计算机上运行时所消耗的时间。但是,同一个程序在不同类型的机器上运行所需的时间不一定相同,所以这种统计是片面的。

(2) **事先估算法**:就是根据每条指令的执行时间估算依据算法编制的程序在计算机上运行时所消耗的时间。但是,因为每种类型机器的指令集不同,执行的时间也不尽相同,这种方法也离不开具体的机器软硬件环境和设备。

显然,以具体的时间单位作为计算程序执行时的时间度量是不科学的。所以在计算算法的执行时间时,应该抛开具体机器软硬件环境和设备,而使用指令的执行次数作为时间单位更合理些。在算法中,可以使用基本语句的执行次数作为算法的时间度量单位。可以认为一个特定算法时间性能只依赖于问题的规模(通常用 n 来表示),或者说,它是关于问题规模 n 的一个函数 $f(n)$,当问题规模 n 趋近于无穷大时的时间量级就称为算法的渐近时间复杂性,简称**时间复杂性**或称**时间复杂度**。记作

$$T(n) = O(f(n))$$

即 $T(n)$ 是 $f(n)$ 的同阶无穷大。

【例 1.1】 分析如下程序段的时间性能。

```
s=0;
for(i=1;i<=n;i++)
    s=s+i;
```

分析:

s=0;	执行 1 次
i=1;	执行 1 次
i<=n;	执行 $n+1$ 次
s=s+i;	执行 n 次
i++;	执行 n 次

总的执行次数为 $3(n+1)$ 次, 因此, 该算法的时间复杂性为 $T(n) = O(3(n+1)) = O(n)$ 。

【例 1.2】 分析如下程序段的时间性能。

```
x=0; y=10;
while(y<100)
{ if(x==10)x=1;
  x++; y+=x;
}
```

分析: $x=0$; 和 $y=10$; 各执行 1 次; 循环部分语句的执行过程取决于 x 值和 y 值的变化, 参见表 1.1。

表 1.1 例 1.2 的时间性能统计表

x 值	0	1	2	3	4	5	6	7	8	9	10	2	3	4	5	6	7	8	总次数
y 值	10	11	13	16	20	25	31	38	46	55	65	67	70	74	79	85	92	100	
语句	当 x 值和 y 值变化时, 各语句执行情况统计(1: 执行, 0: 未执行)																		
$y < 100$	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	18
$x == 10$	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	0	17
$x = 1$	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	1
$x++$	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	0	17
$y += x$	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	0	17

所以, 总的执行次数为 $2+18+17+1+17+17=72$ 次。这与问题规模 n 无关, 因此, 该算法的时间复杂性可以写成 $T(n) = O(1)$ 。

如果这样去计算算法的时间性能, 那就太不合适了。在分析算法时间性能时, 常用最基本语句的执行次数来估算。所谓最基本语句通常是指最深层循环体中的语句, 也是执行频度最高的语句。它的执行次数反映了整个算法的基本时间性能。如例 1.1 中的 $s = s+i$ 和 $i++$ 均被执行了 n 次, 所以 $T(n) = O(n)$; 例 1.2 中的 $x == 10$ 、 $x++$ 及 $y += x$ 都被执行了 17 次, 所以 $T(n) = O(1)$ 。

实际上, 算法的时间量级有多种形式, 见表 1.2, 其对应的函数曲线见图 1.2。

表 1.2 算法的时间量级分类表

名 称	时间复杂度 $T(n)$	说 明
常量阶	$O(1)$	与问题规模无关的算法
线性阶	$O(n)$	与问题规模相关的单重循环
平方阶	$O(n^2)$	与问题规模相关的二重循环
立方阶	$O(n^3)$	与问题规模相关的三重循环
指数阶	$O(e^n)$	较为复杂

续表

名 称	时间复杂度 $T(n)$	说 明
对数阶	$O(\log_2 n)$	折半查找算法
复合阶	如 $O(n \log_2 n)$	堆排序算法
其他	不确定	过于复杂

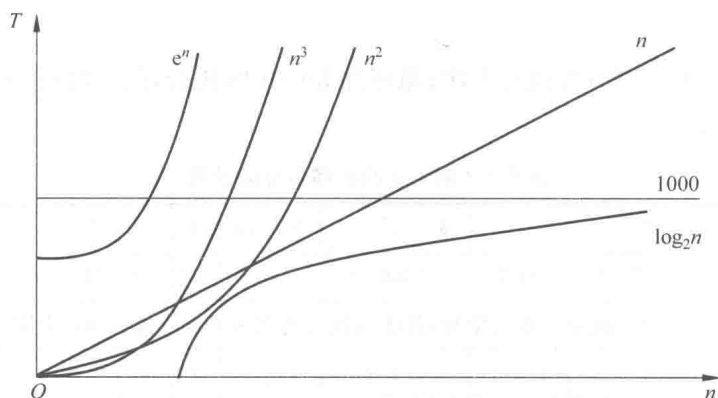


图 1.2 常见函数的增长率

【例 1.3】 讨论起泡排序算法的时间复杂性。

```

void bubblesort(int a[], int n)                                /* 0 单元作为临时存储空间 */
{ int i, j, change;                                           /* change 为交换标志 */
  for(i=change=1; change && i < n; i++)
    for(j=1, change=0; j <= n-i; j++)
      if(a[j] > a[j+1])
      { a[0]=a[j]; a[j]=a[j+1]; a[j+1]=a[0];
        change=1;                                           /* 有交换, 则进行下一趟比较 */
      }
}

```

这一算法的基本操作是内层循环的比较语句: $a[j] > a[j+1]$, 而该语句的执行次数既受到与问题规模 n 有关的循环控制变量 i 和 j 影响, 也受到是否进行下一次循环的条件变量 $change$ 的影响, 这与输入的数据有关。若输入的数据是已经有序的, 则外循环仅进行一次, 所以 $T(n) = O(n)$; 若输入的数据是杂乱无章的, 则外循环可能会执行 $n-1$ 次, 则比较语句的执行次数为 $(n-1) + (n-2) + \dots + (n-i) + \dots + 1 = n(n-1)/2$, 所以, $T(n) = O(n^2)$ 。

通过上例可以看出, 一个算法时间复杂性可能存在最好情况和最坏情况, 通常要以算法的平均时间复杂性来进行算法分析。但是算法的平均时间复杂性取决于各种数据出现的概率, 不好进行分析, 所以, 往往借助于最坏时间复杂性来进行算法分析与评价。