

涵盖世界知名IT公司技术面试的程序设计问题及其解题思路

供不同复杂度的解决方法，兼顾自学教材和面试辅导的不同需求

数据结构与算法 经典问题解析

(原书第2版)

[印度] 纳拉辛哈·卡鲁曼希 (Narasimha Karumanchi) 著

沈华 李兵兵 杜江毅 张明武 译



DATA STRUCTURES
AND ALGORITHMS MADE EASY
SECOND EDITION



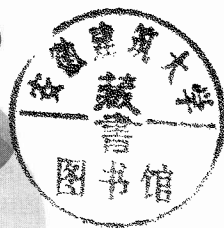
机械工业出版社
China Machine Press

数据结构与算法 经典问题解析

(原书第2版)

[印度] 纳拉辛哈·卡鲁曼希 (Narasimha Karumanchi) 著

沈华 李兵兵 杜江毅 张明武 译



DATA STRUCTURES
AND ALGORITHMS MADE EASY
SECOND EDITION



机械工业出版社
China Machine Press

图书在版编目 (CIP) 数据

数据结构与算法经典问题解析 (原书第 2 版) / (印) 纳拉辛哈·卡鲁曼希 (Narasimha Karumanchi) 著; 沈华等译. —北京: 机械工业出版社, 2018.10

书名原文: Data Structures and Algorithms Made Easy, Second Edition

ISBN 978-7-111-61241-4

I. 数… II. ①纳… ②沈… III. ①数据结构 ②算法分析 IV. ① TP311.12 ② TP312

中国版本图书馆 CIP 数据核字 (2018) 第 241550 号

本书版权登记号: 图字 01-2016-0678

Translation from the English language edition:

Data Structures and Algorithms Made Easy, Second Edition, by Narasimha Karumanchi (ISBN: 9780615459813).

Copyright © 2014 by CareerMonk. com.

All Rights reserved. No part of this book may be reproduced or transmitted in any form or by any means, electronic or mechanic, including photocopying, recording, or by any information storage retrieval system, without permission of Brainy Software, Inc.

Chinese simplified language edition published by China Machine Press.

Copyright © 2019 by China Machine Press.

本书中文简体字版由 CareerMonk Publications 授权机械工业出版社独家出版。未经出版者书面许可, 不得以任何方式复制或抄袭本书内容。

本书以问题驱动的方式系统地介绍了什么是数据结构, 什么是算法, 如何评价一个算法的好坏, 以及各种常用的数据结构和常用的算法设计技巧。全书共 21 章, 其覆盖的知识面很广泛, 大致包括以下内容: 基础知识 (第 1 ~ 2 章, 主要涉及数据结构的概念、算法的概念、算法分析的基本方法、递归等方面的内容)、常见的几种数据结构 (第 3 ~ 9 章)、几种重要的算法 (第 10 ~ 14 章)、常用的算法设计技巧 (第 15 ~ 19 章)、基本的计算复杂性理论 (第 20 章) 和其他主题 (第 21 章)。

本书可以作为高等院校计算机科学与技术及相关专业本科生的教材, 也可以作为相关专业人员的参考资料。

出版发行: 机械工业出版社 (北京市西城区百万庄大街 22 号 邮政编码: 100037)

责任编辑: 朱秀英

责任校对: 殷虹

印刷: 北京诚信伟业印刷有限公司

版次: 2019 年 1 月第 1 版第 1 次印刷

开本: 186mm × 240mm 1/16

印张: 30.75

书号: ISBN 978-7-111-61241-4

定价: 79.00 元

凡购本书, 如有缺页、倒页、脱页, 由本社发行部调换

客服热线: (010) 88378991 88361066

投稿热线: (010) 88379604

购书热线: (010) 68326294 88379649 68995259

读者信箱: hzsj@hzbook.com

版权所有·侵权必究

封底无防伪标均为盗版

本书法律顾问: 北京大成律师事务所 韩光/邹晓东

数据结构和算法是计算机科学的核心内容之一，是计算机专业学生的重要专业课程，对培养计算机专业设计与创新型人才起着关键作用。它具有概念多、内容广、抽象性强、需要一定的数学基础等特点，学习难度较大，对刚刚开始学习的学生来说更是如此。

在翻译 Narasimha Karumanchi 博士撰写的这本数据结构和算法入门书的过程中，我们发现该书有以下特点：

第一，语言通俗，与现实中的问题紧密相连，并且没有使用非常复杂的数学知识，非常适合数据结构和算法的初学者。

第二，对于经典的应用问题，常常给出基于不同数据结构的多种算法的求解方案，让读者充分了解和感受相同问题的不同算法在时间开销、空间开销和编程难度上的不同。

第三，习题丰富。本书的每一章都有包含大量习题的问题集，其中的每道题，作者均结合本章知识点或综合本章和其他章的知识点给出了分析和求解方案。

非常有幸得到机械工业出版社的信任参与了该书的翻译工作。由于水平有限，译文中表达不准确之处在所难免，还请业内专家和广大读者包涵，在此也敬请各位专家、学者和读者批评指正。欢迎各位读者将发现的错误或提出的任何意见和建议发送到邮箱 nancy78733@126.com。

最后，我要感谢湖北工业大学数据安全与隐私保护技术研究所、湖北工业大学工业大数据协同创新中心的诸位同事、同学的建议和帮助。在翻译过程中，张明武教授提出了很多有益的意见，在此表示诚挚的感谢。非常感谢机械工业出版社朱劼编辑对我的包容、理解、帮助和支持，也非常感谢机械工业出版社其他工作人员的辛勤付出。

沈 华

2018 年 3 月于武汉南湖

我知道很多人不会阅读前言。但我强烈建议你至少浏览一遍本书的前言，因为你会发现一些与众不同的东西。

本书不是展示数据结构和算法的定理和证明，而是另辟蹊径，通过改进不同复杂度问题的解决方案来讲解数据结构和算法的相关内容（对于每一个问题，都有复杂度由大到小的多种解决方案）。这些方法基本上涵盖了所有可能的问题解决之道。这种讲解模式可以让你在遇到新问题时，获得思考所有可能解决方案的思维方式。本书也可以让你在准备面试、考试和升学面试时收益良多。

作为一名求职者，如果透彻理解了本书的内容，那么在面对面试官时将会从容不迫，甚至感到绰绰有余。作为一名教师，如果通读了本书，我相信你自然会提升教学质量，最终会让你的学生对选择了计算机科学/信息技术专业而感到庆幸。

对工程专业的本科生和研究生来说，本书在他们的学术准备中非常有用。本书每章都把重点放在问题和对问题的分析上，而不是单纯的理论。每章都是先介绍必需的基础理论，然后是问题集。本书大约有 700 个算法问题，并且给出了解决方法。

如果你是一名正在为参加计算机科学/信息技术竞赛考试做准备的学生，你会发现这本书涵盖了你所需要的所有主题，并且给出了详细的阐述。在写这本书的时候，我的关注点就已经放在了帮助参加这些考试的学生上。

对许多问题，本书给出了多种不同复杂度级别的解决方案。我们以蛮力法解决方案开始，慢慢地向可能的最优方案靠近。对每个问题，我们将设法弄清楚算法花费的时间和占用的内存空间。

本书建议的阅读方式是，先至少通读一遍以获得对所有主题的全面了解，随后读者可以按需跳到感兴趣的章节。尽管我们做了足够多的校验，书中难免还会存在一些不足。我们将在 www.careerMonk.com 网站上更新发现的任何错误。请读者关注该网站，后续任何错误修正、新问题和解决方案都会在该网站上更新。此外，请将你的宝贵意见发到下面的邮箱：Info@CareerMonk.com。

献上最美好的祝福。我相信你会发现此书物有所值。

Narasimha Karumanchi

译者序
前言

第 1 章 绪论 1

1.1 变量	1
1.2 数据类型	1
1.3 数据结构	2
1.4 抽象数据类型	2
1.5 什么是算法	3
1.6 为什么需要分析算法	3
1.7 算法分析的目的	3
1.8 什么是运行时间分析	3
1.9 如何比较算法	4
1.10 什么是增长率	4
1.11 常用的增长率	4
1.12 算法分析的类型	5
1.13 渐近符号	5
1.14 O 符号	6
1.15 Ω 符号	7
1.16 Θ 符号	8
1.17 为什么称为渐近分析	9
1.18 渐近分析的准则	9
1.19 渐近符号的性质	11
1.20 常用的对数公式和求和公式 ..	11
1.21 分治法的主定理	11
1.22 与分治法主定理相关的问题 ..	12
1.23 减治递推的主定理	13
1.24 减治主定理的另一种形式	13
1.25 猜测与确认的方法	13
1.26 平摊分析	15
1.27 关于算法分析的问题集	15

第 2 章 递归与回溯 28

2.1 引言	28
--------------	----

2.2 什么是递归	28
2.3 为什么需要递归	28
2.4 递归函数的格式	28
2.5 递归与内存(图形化演示)	29
2.6 递归与迭代	30
2.7 递归的要点	30
2.8 递归算法举例	30
2.9 关于递归的问题集	31
2.10 什么是回溯	32
2.11 回溯算法举例	32
2.12 关于回溯的问题集	32

第 3 章 链表 35

3.1 什么是链表	35
3.2 链表的抽象数据类型	35
3.3 为什么需要链表	35
3.4 数组回顾	35
3.5 链表与数组、动态数组的 比较	37
3.6 单链表	37
3.7 双链表	43
3.8 循环链表	48
3.9 一种存储高效的双链表	54
3.10 松散链表	55
3.11 跳表	61
3.12 关于链表的问题集	64

第 4 章 栈 87

4.1 什么是栈	87
4.2 如何使用栈	87
4.3 栈的抽象数据类型	87
4.4 栈的应用	88
4.5 栈的实现	88
4.6 栈实现的比较	94

4.7 关于栈的问题集	94	第8章 不相交集	226
第5章 队列	114	8.1 引言	226
5.1 什么是队列	114	8.2 等价关系和等价类	226
5.2 如何使用队列	114	8.3 不相交集的抽象数据类型	227
5.3 队列的抽象数据类型	114	8.4 不相交集的应用	227
5.4 操作异常	115	8.5 不相交集实现的折中方案	227
5.5 队列的应用	115	8.6 快速查找 Fast FIND 的实现 (Quick FIND)	227
5.6 队列的实现	115	8.7 快速合并 Fast UNION 的实现 (Quick UNION)	228
5.7 关于队列的问题集	121	8.8 快速合并 Fast UNION 的实现 (Slow FIND)	228
第6章 树	127	8.9 快速合并 Fast UNION 的实现 (Quick FIND)	231
6.1 什么是树	127	8.10 小结	234
6.2 相关术语	127	8.11 关于不相交集的问题集	234
6.3 二叉树	128	第9章 图算法	235
6.4 几种特殊的二叉树	128	9.1 引言	235
6.5 二叉树的性质	129	9.2 相关术语	235
6.6 二叉树的遍历	131	9.3 图的应用	238
6.7 一般的树(N 叉树)	153	9.4 图的表示	238
6.8 线索二叉树的遍历(与栈/队列 无关的遍历)	159	9.5 图的遍历	242
6.9 表达式树	166	9.6 拓扑排序	249
6.10 XOR 树	168	9.7 最短路径算法	250
6.11 二叉搜索树	169	9.8 最小生成树	256
6.12 平衡二叉搜索树	184	9.9 关于图算法的问题集	259
6.13 AVL 树	184	第10章 排序	280
6.14 其他形式的树	200	10.1 什么是排序	280
第7章 优先队列和堆	204	10.2 为什么需要排序	280
7.1 什么是优先队列	204	10.3 排序算法的分类	280
7.2 优先队列的抽象数据类型	204	10.4 其他分类方式	281
7.3 优先队列的应用	205	10.5 冒泡排序	281
7.4 优先队列的实现	205	10.6 选择排序	282
7.5 堆和二项堆	206	10.7 插入排序	283
7.6 二项堆	207	10.8 希尔排序	285
7.7 堆排序	213	10.9 归并排序	287
7.8 关于优先队列(堆)的问题集	214		

10.10	堆排序	289	第 14 章	散列法	346
10.11	快速排序	289	14.1	什么是散列法	346
10.12	树排序	292	14.2	为什么需要散列法	346
10.13	排序算法的比较	292	14.3	散列表的抽象数据类型	346
10.14	线性排序算法	292	14.4	理解散列法	346
10.15	计数排序	293	14.5	散列法的构成要素	347
10.16	桶排序(或箱排序)	293	14.6	散列表	348
10.17	基数排序	294	14.7	散列函数	348
10.18	拓扑排序	295	14.8	装填因子	348
10.19	外部排序	295	14.9	冲突	349
10.20	关于排序的问题集	296	14.10	解决冲突的方法	349
第 11 章	搜索	306	14.11	拉链法	349
11.1	什么是搜索	306	14.12	开放地址法	349
11.2	为什么需要搜索	306	14.13	冲突解决方法的比较	351
11.3	搜索的类型	306	14.14	如何得到复杂度为 $O(1)$ 的 散列法	352
11.4	无序线性搜索	306	14.15	散列技术	352
11.5	排序/有序线性搜索	307	14.16	不合适构造散列表的 问题	352
11.6	二分搜索	307	14.17	Bloom 过滤器	353
11.7	基本搜索算法的比较	308	14.18	关于散列法的问题集	355
11.8	符号表和散列	308	第 15 章	字符串算法	366
11.9	字符串搜索算法	308	15.1	引言	366
11.10	关于搜索的问题集	308	15.2	字符串匹配算法	366
第 12 章	选择算法(中位数)	333	15.3	蛮力法	366
12.1	什么是选择算法	333	15.4	Robin-Karp 字符串匹配 算法	367
12.2	基于排序的选择	333	15.5	基于有限自动机的字符串 匹配	368
12.3	基于划分的选择算法	333	15.6	KMP 算法	370
12.4	线性选择算法—— Median of Median 算法	333	15.7	Boyce-Moore 算法	373
12.5	按序寻找第 k 小元素	333	15.8	字符串的存储结构	373
12.6	关于选择算法的问题集	334	15.9	字符串的散列表	374
第 13 章	符号表	343	15.10	字符串的二叉搜索树	374
13.1	引言	343	15.11	字典树	374
13.2	什么是符号表	343	15.12	三叉搜索树	377
13.3	符号表的实现	343	15.13	二叉搜索树、字典树和 三叉搜索树的比较	380
13.4	符号表实现的比较	344			

15.14	后缀树	380	18.10	关于分治法的问题集	409
15.15	关于字符串的问题集	384			
第 16 章	算法设计技巧	393	第 19 章	动态规划	422
16.1	引言	393	19.1	引言	422
16.2	算法的分类	393	19.2	什么是动态规划策略	422
16.3	基于实现方法的算法分类	393	19.3	动态规划策略的性质	422
16.4	基于设计方法的算法分类	394	19.4	动态规划可以求解所有的问题吗	422
16.5	其他方式的算法分类	395	19.5	动态规划方法	422
第 17 章	贪婪算法	396	19.6	动态规划算法举例	423
17.1	引言	396	19.7	理解动态规划	423
17.2	贪婪策略	396	19.8	最长公共子序列	426
17.3	贪婪算法的要素	396	19.9	关于动态规划的问题集	429
17.4	贪婪技术总是奏效吗	396	第 20 章	复杂性类	461
17.5	贪婪方法的优缺点	396	20.1	引言	461
17.6	贪婪技术的应用	397	20.2	多项式/指数时间	461
17.7	理解贪婪技术	397	20.3	什么是判定性问题	461
17.8	关于贪婪算法的问题集	400	20.4	判定性过程	462
第 18 章	分治算法	407	20.5	什么是复杂性类	462
18.1	引言	407	20.6	复杂性类的类型	462
18.2	什么是分治策略	407	20.7	归约	464
18.3	分治法总是奏效吗	407	20.8	关于复杂性类的问题集	466
18.4	分治法原理的图形化演示	407	第 21 章	其他主题	469
18.5	理解分治法	408	21.1	引言	469
18.6	分治法的优点	408	21.2	位编程的技巧	469
18.7	分治法的缺点	409	21.3	其他编程问题	474
18.8	主定理	409	参考文献	481	
18.9	分治法的应用	409			

绪 论

本章的目标是说明算法分析的重要性、算法分析需要用到的各种符号、这些符号之间的关系和运用算法分析的方法去解决尽可能多的问题。首先，让我们关注算法的基本要素、算法分析的重要性，然后，再逐渐介绍上述提到的其他主题。学完本章，对任何给定的算法(特别是递归函数)，读者应该能够给出它的复杂度。

1.1 变量

在给出变量的定义之前，让我们将它与古老的数学等式联系起来。从孩提时代开始，我们已经求解过许多数学等式。作为一个例子，考虑下面的等式：

$$x^2 + 2y - 2 = 1$$

我们不需要担心这个等式是否有使用价值。我们需要理解的关键问题是，这个等式中的一些名字(x 和 y)具有值(数据)。这意味着，这些名字(x 和 y)是用来表示数据的占位符。类似地，在计算机科学编程中，我们需要某些东西来代表数据，变量(variable)就扮演了这个角色。

1.2 数据类型

在上述数学等式中，变量 x 和 y 可以取任意值，如(10, 20)范围内的整数、(0.23, 5.5)范围内的实数，或者仅仅取值为 0 和 1。因此，为了求解这个等式，我们必须确定变量 x 和 y 取什么类型的值。在计算机科学编程中，数据类型(data type)就是用来实现这个目的的。在编程语言中，数据类型是预先定义了值的数据的集合，如整型、浮点型、字符型、字符串等。

计算机内存中存储的都是 0、1 代码串。如果我们想对某个问题进行编程求解，那么提供以 0、1 代码串表示的解决方案是非常困难的。为了解决这个困难，编程语言和编译器向我们提供了数据类型。例如，整型数据占 2 字节(这个值取决于具体的编译器)，浮点型数据占 4 字节，等等。这就是说，在内存中 2 字节(16 位)被合并了，称为一个整数。同样，4 字节(32 位)被合并了，称为一个浮点数。数据类型降低了编程的难度。在编程语言中，有两类数据类型：

- 系统定义的数据类型(也称为基本数据类型)
- 用户定义的数据类型(也称为用户自定义类型)

系统定义的数据类型(基本数据类型)

系统定义的数据类型被称为基本数据类型。许多编程语言提供以下基本数据类型：int、float、char、double、bool，等等。基本数据类型的数据在内存中占据的位数由编程语言、编译器和操作系统确定。即使是相同的基本数据类型，在不同的编程语言中它的

大小(即所占字节数)也可能不同。因为数据类型所有可能的取值(即取值范围)取决于它在内存中占据的位数,所以在不同编程语言中相同数据类型的取值范围也可能不同。

例如, `int` 类型数据的大小可能是 2 字节或者 4 字节。如果它的大小是 2 字节(16 位),那么它的取值范围是 $(-32\,768, +32\,767)$ (即 $(-2^{15}, 2^{15}-1)$); 如果它的大小是 4 字节(32 位),那么它的取值范围是 $(-2\,147\,483\,648, 2\,147\,483\,647)$ (即 $(-2^{31}, 2^{31}-1)$)。其他数据类型也有相同的情况。

用户定义的数据类型(用户自定义类型)

如果系统定义的数据类型不能满足应用需求,那么一些编程语言允许用户定义自己的数据类型,这样的数据类型被称为用户定义的数据类型(或用户自定义类型)。C/C++ 语言中的结构(structure)和 Java 中的类(class)都是用户自定义类型的典型例子。例如,在下面的代码段中,我们运用结构体合并多个系统定义的数据类型生成了一个名为“newType”的用户自定义类型。这在处理计算机内存方面提供了更多的灵活性和舒适性。

```
struct newType {  
    int data1;  
    float data2;  
    ...  
    char data;  
};
```

1.3 数据结构

基于上面的讨论可知,一旦有了变量中的数据,我们就需要一些使用这些数据解决问题的机制。在计算机中,数据结构(data structure)是以高效使用数据为目标的存储和组织数据的特定方式。一种数据结构就是组织和存储数据的一种特殊格式。常见的数据结构包括数组、文件、链表、栈、队列、树、图等。

从组织元素的角度,数据结构被分为以下两类:

- 1) 线性数据结构:该数据结构中的元素按顺序依次进行访问,但并非一定按顺序依次进行存储。例如链表、栈和队列。
- 2) 非线性数据结构:该数据结构中的元素以非线性的方式进行存储/访问。例如树和图。

1.4 抽象数据类型

给出抽象数据类型(ADT)的定义之前,让我们从另外一个角度去认识系统定义的数据类型。众所周知,默认情况下,所有基本数据类型(如 `int`、`float` 等)均支持诸如加法、减法这样的基本操作。系统提供这些基本数据类型的操作的实现。因此,对用户自定义数据类型而言,我们也需要定义它能够完成的操作。想使用用户自定义数据类型去解决实际问题,我们还需要给出这些操作的实现。这就意味着,通常用户自定义数据类型的定义包括对其操作的定义。

为了简化求解问题的过程,我们将数据结构和它的操作封装在一起称为抽象数据类型(Abstract Data Type, ADT)。一个抽象数据类型由两个部分构成:

- 1) 数据的声明
- 2) 操作的声明

常用的 ADT 包括：链表、栈、队列、优先队列、二叉树、字典、不相交集(合并和查找)、散列表、图等。例如，栈使用 LIFO(Last-In-First-Out)机制存储数据到栈中。最后入栈的元素最先出栈。栈的基本操作包括：栈的创建、入栈(即压入一个元素到栈顶)、出栈(即弹出栈顶元素)、访问当前的栈顶元素、求栈的长度(即栈中元素的个数)，等等。

在定义 ADT 时，我们无须关心其实现细节。只有当我们真正需要使用 ADT 时，才去考虑它们的实现细节。不同的 ADT 适用于不同的应用，适用于特定任务的 ADT 则是非常专业化的。通读完本书，读者会认识很多的 ADT，而且能够根据需要解决的应用问题选择合适的数据结构。

1.5 什么是算法

让我们考虑准备煎蛋的问题。为了煎蛋，我们会按以下步骤进行准备：

- 1) 准备煎锅；
- 2) 准备油；
 - a. 有油吗？
 - i. 如果有，那么将油倒入锅内；
 - ii. 如果没有，那么我们需要去买油吗？
 1. 如果需要，那么我们去买油；
 2. 如果不需要，那么煎蛋终止；
- 3) 打开炉子；
-

对于一个给定问题(准备煎一个蛋)，我们所做的是，给出这个问题的一步一步的求解过程。算法的定义是：

算法是求解给定问题的一系列指令。

注意：我们不需要证明算法的每一个步骤。

1.6 为什么需要分析算法

为了实现从城市 A 移动到城市 B，我们有很多可选的方式：搭飞机、坐巴士、乘火车、骑自行车。根据可行性和便利性，我们会从中选择适合自己的一种。类似地，在计算机科学中，求解同一个问题有多种可选的算法(举个例子，排序问题有很多求解算法，如插入排序法、选择排序法、快速排序法，等等)。算法分析可以帮助我们判断哪个算法在时间和空间开销方面是高效的。

1.7 算法分析的目的

算法分析(analysis of algorithm)的目的是，主要从运行时间方面比较算法(或解决方案)，并且也从其他方面(如内存、开发人员的工作量等)对算法进行比较。

1.8 什么是运行时间分析

运行时间分析用来确定随着问题规模(即输入大小)的增大算法处理时间是如何增加的。输入大小是输入的元素个数。根据问题的类型，输入类型可以不同。下面是常见的

几种输入类型：

- 数组的大小
- 多项式的次数
- 矩阵的元素个数
- 输入的二进制表示的位数
- 图的顶点个数和边的条数

1.9 如何比较算法

为了比较算法，我们定义了如下几个客观度量(objective measure)：

执行时间。不是好的度量，因为执行时间依赖于具体的计算机。

执行的语句数。也不是好的度量，因为语句的个数依赖于具体的编程语言和程序员的个人风格。

理想的客观度量是什么？假设我们将给定算法的运行时间表示为关于输入大小 n 的函数(例如 $f(n)$)，那么比较算法就是比较与运行时间相关的不同函数。这种比较与具体的机器时间、编程风格等无关。

1.10 什么是增长率

算法的运行时间被表示为关于输入的函数，它随着输入规模增大的增长速度称为增长率(rate of growth)。假设你到一家车行去买一辆汽车和一辆自行车。如果你的朋友恰巧在车行碰到你并且询问你在买什么，那么通常会回答在买汽车。为什么会这样？这是因为一辆汽车的花费远远大于一辆自行车的费用(即买自行车的费用被忽略在买汽车的费用中了)。

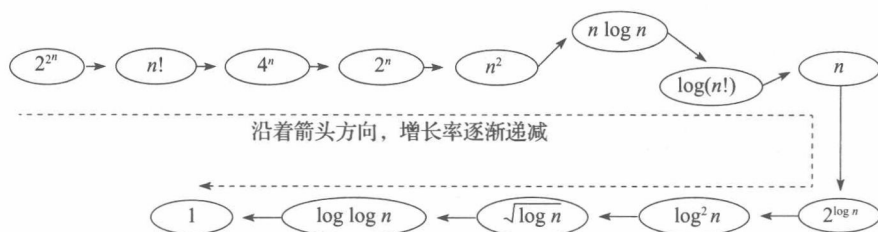
$$\begin{aligned}\text{Total Cost} &= \text{cost_of_car} + \text{cost_of_cycle} \\ \text{Total Cost} &\approx \text{cost_of_car}(\text{近似})\end{aligned}$$

我们可以用函数来刻画上述例子中买汽车和自行车的开销。对给定的函数可以忽略其低阶项，因为低阶项对整个函数的贡献相对来说微不足道(当输入规模 n 很大时)。在下面的例子中， n^4 、 $2n^2$ 、 $100n$ 和 500 是某个函数的各项成本组成，可以将函数近似处理为 n^4 ，因为 n^4 是函数中增长速度最快的项。

$$n^4 + 2n^2 + 100n + 500 \approx n^4$$

1.11 常用的增长率

下图展示了不同增长率之间的关系。



下表列出了后续章节会涉及的增长率。

时间复杂度	名称	例子
1	常数阶	在链表首部插入一个元素
$\log n$	对数阶	在一个有序序列中查找元素
n	线性阶	在一个无序序列中查找元素
$n \log n$	线性对数阶	运用分治归并排序(即自顶向下的归并排序)对 n 个元素进行排序
n^2	平方阶	求图中两顶点之间的最短路径
n^3	立方阶	矩阵相乘
2^n	指数阶	汉诺塔问题

1.12 算法分析的类型

为了分析给定算法，我们需要知道算法在哪些输入下花费的时间较短(即表现良好)和在哪些输入下花费的时间较长。我们已经知道算法的时间开销可以表示成表达式的形式。这就意味着，我们可以用多个表达式来表达不同情况下算法的时间开销：需要较少时间的情况和需要较多时间的情况。

通常，第一种情况被称为算法的最好情况(best case)，第二种情况被称为算法的最坏情况(worst case)。为了分析算法，我们需要某种文法的基础和渐近分析/符号的基础。这里介绍三种类型的算法分析：

● 最坏情况

- 定义了算法需要最多时间的输入。
- 该输入是使算法运行最慢的输入。

● 最好情况

- 定义了算法需要最少时间的输入。
- 该输入是使算法运行最快的输入。

● 平均情况

- 提供了关于算法运行时间的预测。
- 假设输入是随机的。

$$\text{下界} \leq \text{平均时间} \leq \text{上界}$$

对一个给定的算法，我们可以用表达式来分别表示它在最好情况、最坏情况和平均情况下的时间开销。例如，令 $f(n)$ 是表示某算法时间开销的函数。

$$f(n) = n^2 + 500, \text{最坏情况}$$

$$f(n) = n + 100n + 500, \text{最好情况}$$

平均情况也是类似的，表达式定义了算法需要花费平均运行时间(或内存空间)的那些输入。

1.13 渐近符号

现在我们已经得到了在最好情况、平均情况和最坏情况下算法时间开销(或空间开销)的表达式。下面我们需要做的是，确定这三种情况下算法开销的上界和下界。为此，

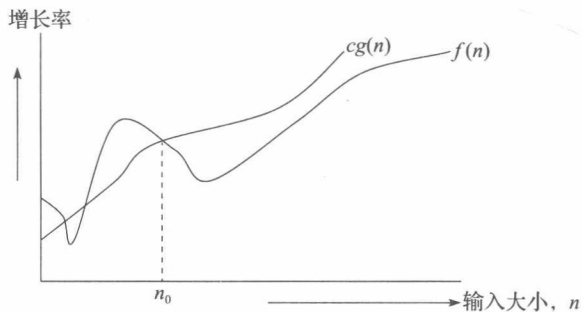
我们需要用某种文法来描述这些界限。为了方便后面几个小节对这个主题的讨论，我们假设某算法的(时间或空间)开销为 $f(n)$ 。

1.14 O 符号

O 符号给出了函数的紧致上界。它通常被描述为 $f(n) = O(g(n))$ ，表示对于一个很大的 n ， $f(n)$ 的上界是 $g(n)$ 。例如，如果 $f(n) = n^4 + 100n^2 + 10n + 50$ ，那么 $g(n) = n^4$ 。这意味着，对于一个很大的 n ， $g(n)$ 给出了 $f(n)$ 的最大增长率。

下面让我们更加详细地了解 O 符号。O 符号定义为： $O(g(n)) = \{f(n) : \text{存在一个正常数 } c \text{ 和正整数 } n_0, \text{ 使得对所有 } n \geq n_0 \text{ 有 } 0 \leq f(n) \leq cg(n) \text{ 成立}\}$ 。 $g(n)$ 是 $f(n)$ 的渐近紧致上界。我们的目标是给出不小于算法增长率 $f(n)$ 的最小增长率 $g(n)$ 。

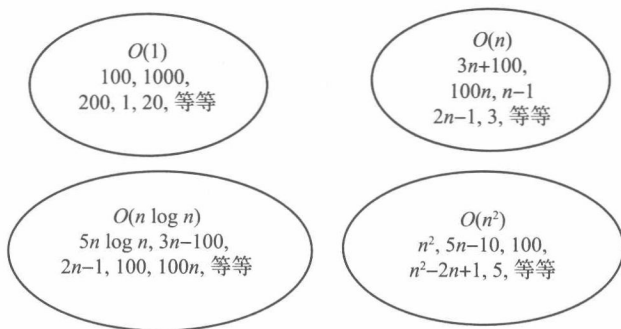
因为 n 较小时的增长率是没有太大意义的，所以通常我们会忽略掉它。在下图中， n_0 是我们考虑某算法增长率的分界点。可见， n 较小时(即小于 n_0)的增长率可能与 n 较大时(即大于 n_0)的增长率不同。



O 符号的图形化演示

$O(g(n))$ 是一个函数的集合，这个集合中的任意一个函数具有比 $g(n)$ 小或相同阶的增长率。例如， $O(n^2)$ 包括 $O(1)$ 、 $O(n)$ 、 $O(n \log n)$ ，等等。

注意：我们仅仅进行 n 取大值时的算法分析。也就是说，我们不考虑小于 n_0 的算法增长率。



关于 O 符号的例子

例 1 求 $f(n) = 3n + 8$ 的上界。

求解：对所有 $n \geq 8$ 有 $3n+8 \leq 4n$ ，所以 $3n+8 = O(n)$ ，其中 $c=4$ ， $n_0=8$ 。

例 2 求 $f(n)=n^2+1$ 的上界。

求解：对所有 $n \geq 1$ 有 $n^2+1 \leq 2n^2$ ，所以 $n^2+1 = O(n^2)$ ，其中 $c=2$ ， $n_0=1$ 。

例 3 求 $f(n)=n^4+100n^2+50$ 的上界。

求解：对所有 $n \geq 11$ 有 $n^4+100n^2+50 \leq 2n^4$ ，所以 $n^4+100n^2+50 = O(n^4)$ ，其中 $c=2$ ， $n_0=11$ 。

例 4 求 $f(n)=2n^3-2n^2$ 的上界。

求解：对所有 $n \geq 1$ 有 $2n^3-2n^2 \leq 2n^3$ ，所以 $2n^3-2n^2 = O(2n^3)$ ，其中 $c=2$ ， $n_0=1$ 。

例 5 求 $f(n)=n$ 的上界。

求解：对所有 $n \geq 1$ 有 $n \leq n$ ，所以 $n = O(n)$ ，其中 $c=1$ ， $n_0=1$ 。

例 6 求 $f(n)=410$ 的上界。

求解：对所有 $n \geq 1$ 有 $410 \leq 410$ ，所以 $410 = O(1)$ ，其中 $c=1$ ， $n_0=1$ 。

不是唯一的解？

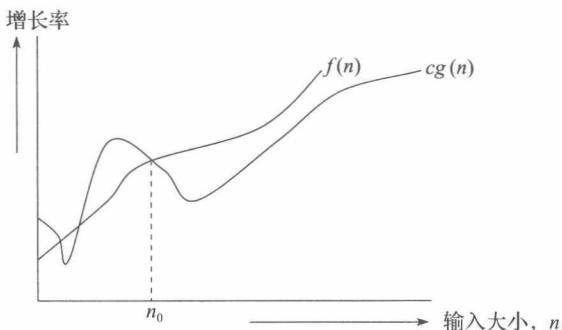
在给出渐近界限的时候， n_0 和 c 的取值并不唯一。例如， $100n+5 = O(n)$ 。求解过程中 n_0 和 c 有多个可能的取值。

求解 1：对所有 $n \geq 5$ 有 $100n+5 \leq 100n+n = 101n \leq 101n$ ，因此 $c=101$ ， $n_0=5$ 是一个可能的解。

求解 2：对所有 $n \geq 1$ 有 $100n+5 \leq 100n+5n = 105n \leq 105n$ ，因此 $c=105$ ， $n_0=1$ 也是一个可能的解。

1.15 Ω 符号

类似于 O 符号的讨论， Ω 符号给出了函数的紧致下界。它通常被描述为 $f(n) = \Omega(g(n))$ ，表示对于一个很大的 n ， $f(n)$ 的紧致下界是 $g(n)$ 。例如，如果 $f(n)=100n^2+10n+50$ ，那么 $g(n)=\Omega(n^2)$ 。



Ω 符号可以定义为： $\Omega(g(n)) = \{f(n) : \text{存在一个正常数 } c \text{ 和正整数 } n_0, \text{ 使得对所有 } n \geq n_0 \text{ 有 } 0 \leq cg(n) \leq f(n) \text{ 成立}\}$ 。 $g(n)$ 是 $f(n)$ 的渐近紧致下界。我们的目标是给出不大于算法增长率 $f(n)$ 的最大增长率 $g(n)$ 。

关于 Ω 符号的例子

例 1 求 $f(n)=5n^2$ 的下界。

求解: $\exists c$ 和 n_0 使得 $0 \leq cn \leq 5n^2 \Rightarrow cn \leq 5n^2 \Rightarrow c=1, n_0=1$ 。

所以 $5n^2 = \Omega(n^2)$, 其中 $c=1, n_0=1$ 。

例2 证明 $f(n) = 100n + 5 \neq \Omega(n^2)$ 。

求解: $\exists c$ 和 n_0 使得 $0 \leq cn^2 \leq 100n + 5$, 又 $100n + 5 \leq 100n + 5n (\forall n \geq 1) = 105n$, 故 $cn^2 \leq 105n \Rightarrow n(cn - 105) \leq 0$;

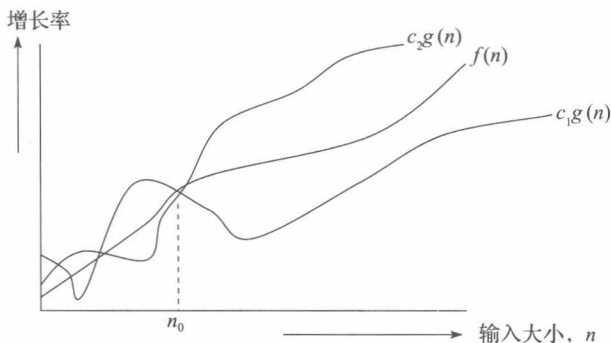
因为 n 是正整数 $\Rightarrow cn - 105 \leq 0 \Rightarrow n \leq 105/c \Rightarrow$ 矛盾: n 不可能比某一个常数小。

例3 $2n = \Omega(n), n^3 = \Omega(n^3), \log n = \Omega(\log n)$ 。

1.16 Θ 符号

Θ 符号可以确定函数(或算法)的上界和下界是否相同。算法的平均运行时间总是介于下界和上界之间的。如果上界(O)和下界(Ω)给出相同的结果, 那么 Θ 符号也会给出同样的结果。举个例子, 假设有 $f(n) = 10n + n$ 。那么, $f(n)$ 的紧致上界 $g(n)$ 是 $O(n)$, 并且最好情况下的增长率 $g(n) = O(n)$ 。

在这个例子中, 最好情况和最坏情况下的增长率相同。因此, 平均情况下也应该相同。对一个给定的函数(或算法), 如果 O 和 Ω 表示的增长率(即界限)不相同, 那么 Θ 表示的增长率就可能与它们不同。在这种情况下, 我们需要考虑针对函数(或算法)所有可能输入的时间复杂度, 然后求出它们的均值(例如, 快速排序的平均情况, 请参见第10章)。



现在给出 Θ 符号的定义: $\Theta(g(n)) = \{f(n) : \text{存在两个正常数 } c_1, c_2 \text{ 和正整数 } n_0 \text{ 使得对所有 } n \geq n_0 \text{ 有 } 0 \leq c_1 g(n) \leq f(n) \leq c_2 g(n) \text{ 成立}\}$ 。 $g(n)$ 是 $f(n)$ 的渐近紧致界。 $\Theta(g(n))$ 是所有增长率与 $g(n)$ 同阶的函数构成的集合。

关于 Θ 符号的例子

例1 求 $f(n) = \frac{n^2}{2} - \frac{n}{2}$ 的 Θ 界。

求解: 对所有 $n \geq 1$ 有 $\frac{n^2}{5} \leq \frac{n^2}{2} - \frac{n}{2} \leq n^2$, 所以 $\frac{n^2}{2} - \frac{n}{2} = \Theta(n^2)$, 其中 $c_1 = 1/5, c_2 = 1, n_0 = 1$ 。

例2 证明 $n \neq \Theta(n^2)$ 。

求解: $c_1 n^2 \leq n \leq c_2 n^2 \Rightarrow$ 仅当 $n \leq 1/c_1$ 时不等式成立, 所以 $n \neq \Theta(n^2)$ 。