

搜索引擎

——原理、技术与系统

(第二版)

李晓明 闫宏飞 王继民 著



科学出版社

搜索引擎

——原理、技术与系统

(第二版)

李晓明 闫宏飞 王继民 著

科学出版社
北 京

内 容 简 介

本书系统介绍了互联网搜索引擎的工作原理、实现技术及系统构建方案。全书分三篇共 13 章。上篇介绍搜索引擎的基本原理和技术,讲述一个小型简单搜索引擎实现的具体细节;中篇详细讨论了大规模分布式搜索引擎系统的设计要点及其关键技术;下篇结合“中国 Web 信息博物馆”和“中国互联网数字资源财富库藏”的实践经验,介绍了构建大规模 Web 历史网页和非网页仓储系统的技术和方法,以及中文网页的自动分类与聚类、开放域问题系统的构建等。

本书层次分明,由浅入深,上篇和中篇涉及内容提供了源代码下载地址;既有深入的理论分析,也有大量的实验数据和程序,具有学习和实用双重意义。

本书可作为高等院校计算机科学与技术、软件工程、信息管理 with 信息系统、电子商务等专业的研究生或高年级本科生的教学参考书和技术资料;对广大从事网络技术、Web 站点管理、数字图书馆、Web 挖掘等研究和应用开发的科技人员有很高的参考价值;书中提供了大量源代码,除了用于构建搜索引擎之外,对于学习编程,提高编程技巧,以及实现一个大规模应用开发也有一定的参考价值。

图书在版编目(CIP)数据

搜索引擎:原理、技术与系统/李晓明,闫宏飞,王继民著.—2 版.—北京:科学出版社,2012

ISBN 978-7-03-034258-4

I. ①搜… II. ①李…②闫…③王… III. ①互联网络-情报检索-高等学校-教材 IV. ①G354.4

中国版本图书馆 CIP 数据核字(2012)第 090075 号

责任编辑:匡 敏 刘兰霞/责任校对:包志虹
责任印制:阎 磊/封面设计:迷底书装

科 学 出 版 社 出 版

北京东黄城根北街16号

邮政编码:100717

<http://www.sciencep.com>

北京市农林印务有限公司印刷

科学出版社发行 各地新华书店经销

*

2005 年 4 月 第 一 版 开本:B5(720×1000)

2012 年 5 月 第 二 版 印张:21 3/4

2012 年 5 月第八次印刷 字数:451 000

定价:48.00 元

(如有印装质量问题,我社负责调换)

第二版前言

2005年4月,在华夏英才基金的支持下,本书的第一版问世。当时,互联网搜索引擎虽然在网民中已经不是一个陌生的概念,但是很少见到针对实际系统的构建,比较系统地介绍搜索引擎原理与实现技术的书。因此本书第一版的出版应该说比较及时,对不少年轻计算机技术人员起到了启蒙的作用。的确,在过去这些年里,我们多次收到读者的反馈,告知那本书对他们职业生涯的作用。最初的读者中,现在有许多都是搜索公司的骨干了。截止到2011年5月,《搜索引擎——原理、技术与系统》已经印刷了七次。一本比较专业的书,能得到那么多读者的喜欢,我们甚感欣慰。

同时我们也看到,随着互联网长盛不衰的发展,搜索的重要性日益突出;而且我们还看到一个现象,那就是相当一批有实力的互联网公司都有进入搜索领域的战略和行动,并不因为有了百度和谷歌就放弃那个市场。这种现象造成了搜索技术人才的很大缺口,尤其是有一定实际开发搜索引擎经验的人才缺口。很多高校看到这种需求,纷纷开设相关课程。许多学者看到这种需求,也纷纷写出各具特色的教材和专著。例如,郭军的《Web 搜索》,董守斌和袁华的《网络信息检索》,以及刘奕群、马少平、洪涛和刘子正的《搜索引擎技术基础》。许多出版社也看到了这种需求,引进了大量相关的外版书籍。

那么,既然市场上已经有了这么多关于搜索引擎技术的书,为什么还要出版本书呢?其实这个想法我们4年前就有了,由于客观原因而耽误下来。最根本的原因是,搜索引擎技术的发展和搜索引擎技术前沿认识的深入,使得我们感到原来书中的有些内容不再重要,而有些新的内容应该包括进去。

2003年秋,在编写本书第一版的时候,主要工作基础是“天网搜索”,它曾经是中国最好的,是我们引以为自豪的搜索引擎。围绕“天网搜索”的开发,北京大学网络实验室培养了一批优秀的学生。本书第一版的内容大都是那些同学实际工作经验的总结,因此一方面总体看的确比较实用,但另一方面某些部分也不够成熟和深入。同时,过去7年来,北京大学网络实验室在搜索引擎技术方面的研究工作也有了深入进展,尤其在搜索评测与高性能索引结构方面取得了一批前沿成果,这些都是在修订中要考虑的内容。

本书保留了第一版上篇的大部分内容,即搜索引擎的基本原理,过去这么些年并没有什么变化;删除了第一版中的第九,第十二和十三章,增加了第十,第十一和十三章,分别介绍基于搜索引擎技术开发并从2002年一直运行至今的“中国 Web 信息博物馆”、“中国数字财富库藏”及开放域问答系统。同时,较大幅度修订了第一版中的部分小节内容。总的来看,第二版中约45%的内容是新的,且总篇幅比第一版增加

约 30%。

鉴于我们在第一版中的一个特色——详细介绍了一个小型搜索引擎(TSE)并提供源代码,引起了许多读者的兴趣,纷纷下载和来邮件咨询,在此高兴地告诉读者:北京大学网络实验室将开放天网搜索系统的所有源代码。用该源代码构建的系统能搜集和处理上亿量级的网页,其体现的技术与本书中的几章内容相对应。

还有两个原因促使我们完成第二版的修订。一是 2011 年我们与百度合作承担一个国家项目“基于框计算的新一代搜索引擎与浏览器”,尽管规定的任务中没有要求,但我们认为能有这么一本书在项目完成之际面世,也是一件令人高兴的事情。二是 2003 年我们发起了“全国搜索引擎与网上信息挖掘学术研讨会”,首届在北京大学举办后,每年在全国各地由不同的高校轮流举办,很多人在会议组织工作中付出了巨大努力,他们是华南理工大学的董守斌、清华大学的李星、山东大学的马军、海南大学的雷景生、江西师范大学的王明文、大连理工大学的林鸿飞、西华大学的杜亚军,河北大学的袁方。今年已是第十届了,又回到北京大学来举办,本书的出版既是对本届会议的一个献礼,也是对过去 10 年来曾经为“全国搜索引擎与网上信息挖掘学术研讨会”做过贡献的所有朋友,以及参加会议的所有人员的感谢。

只要互联网还在发展,Web 还在发展,搜索引擎技术就会不断发展。我们相信《搜索引擎——原理、技术与系统(第二版)》能向读者更深入地展示搜索技术视野,同时也有助于启迪创新的搜索技术。同时我们愿意与读者交流,并欢迎读者指出书中难免的疏漏(lxm@pku.edu.cn; yhf@net.pku.edu.cn; wjm@pku.edu.cn; http://sewm.pku.edu.cn/book/)。

作 者

2012 年 3 月于北京大学燕园

第一版前言

随着互联网的不断发展和日益普及,网上的信息量在爆炸性增长,全球 Web 页面的数目已经超过 40 亿,中国的网页数目估计也超过了 3 亿。目前人们从网上获得信息的主要工具是浏览器,而通过浏览器得到信息通常有三种方式:第一,直接向浏览器输入一个关心的网址(URL),如 <http://net.pku.edu.cn>,浏览器返回所请求的网页,根据该网页内容及其包含的超链接文字(anchor text)的引导,获得自己需要的内容;第二,登录到某个知名门户网站,如 <http://www.yahoo.com>,根据该网站提供的分类目录和相关链接,逐步“冲浪”浏览,寻找自己感兴趣的東西;第三,登录到某个搜索引擎网站,如 <http://e.pku.edu.cn>,输入代表自己所关心信息的关键词或者短语,依据返回的相关信息列表、摘要和超链接引导,试探寻找自己需要的内容。

这三种方式各有特点,各有自己最适合自己的应用场合。第一种方式的应用是最有针对性的,例如,要了解北京大学计算机系网络与分布式系统实验室在做些什么工作,从某个渠道得知该实验室的网址为 <http://net.pku.edu.cn>,于是直接用它驱动浏览器就是最有效的方式。第二种方式的应用类似于读报,用户不一定有明确的目的,只是想看看网上有什么有意思的消息;当然这其中也可能是关心某种主题,如体育比赛、家庭生活等。第三种方式适用于用户大致知道自己要关心的内容,如“国有股减持”,但不清楚哪里能够找到相关信息(即不知道哪些 URL 能给出这样的信息);在这种场合,搜索引擎能够为用户提供一个相关内容的网址及其摘要的列表,由用户一个个试探看是否为自己需要的。现在的搜索引擎技术已经能做到在多数情况下满足用户的这种需要。CNNIC 的信息统计指出,目前搜索引擎已经成为继电子邮件之后人们用得最多的网上信息服务系统。

同时,随着网上信息资源规模的增长,尤其是其内容总体和我们社会的演化发生着越来越密切的联系,研究网上存在的海量信息逐渐成为许多学科关注的一个方向。为此,不少研究人员也有采样搜集特定内容、一定数量网页的需要。

本书以我们设计、实现并维护、运行北大“天网”搜索引擎的实际经验,介绍大规模搜索引擎的工作原理和实现技术。我们要向读者揭示,为什么向搜索引擎输入一个关键词或者短语,就能够在几秒钟内得到那么多相关的文档及其摘要,而点击其中的链接就能够被引导到文档的全文,且其中相当一部分可能正是用户需要的。

我们按照上、中、下三篇展开相关的内容。上篇讲搜索引擎的基本工作原理和技术,要解决的是为什么搜索引擎能提供如此庞大的信息查找服务这一问题,以及它在功能上有什么本质的局限性。这一篇的内容包括网页的搜集过程,网页信息的提取、组织方式和索引结构,查询提交和响应的过程以及结果产生,等等。这其中,虽然我们假定读者熟悉 URL、HTML、HTTP、CGI、MIME 等基本概念,但在上下文中也给

予了必要的介绍,力图保持行文的流畅性。这一部分内容对于需要构建小规模搜索引擎的研究人员会有直接的参考价值。

中篇讨论和大规模实用搜索引擎有关的技术问题。所谓大规模在这里指至少维护超过 1000 万的网页信息,提供相关的查询服务。所涉及的内容包括并行分布处理技术的应用,数据局部性的开发,缓存技术的应用以及搜集的网页在提供服务之前的预处理问题和高效倒排文件的建立技术等。这一部分的讨论有比较强的计算机系统结构的风格,我们将向读者展示计算机系统结构课程中的那些概念是如何生动地体现在一个实际应用系统中的。这一部分内容对构建大规模数字图书馆的技术人员也应该有帮助。

下篇介绍挑战性更强一些的内容。一般地讲,前面所述可以称为是“通用搜索引擎”,为最广泛的人群提供信息查询服务是它的基本宗旨。这意味着它的应用模式必须尽量简单,即关键词或查询短语的提交和匹配响应。尽管这已经可以解决许多问题了,但对有些重要的信息需求依然显得力不从心。例如,一个人可能会关心最近半年来网上出现了哪些关于他(她)的信息,一个企业可能要关心它做了一次大规模促销活动后一个月内网上有什么反响,一个政府机构可能会关心在一项政策法规颁布后的网上舆论。面向主题和个性化的信息查询服务就是我们试图描述的一种基本途径。这一部分内容更多地和网上中文信息处理技术有关。更准确地讲,我们要介绍网络与并行分布处理技术和中文处理技术的结合,从而实现大规模、高性能、高质量、有针对性的网上信息查询服务。这一部分内容反过来可能对从事中文信息处理的研究人员有启发作用。

本书的内容是集体智慧的结晶,主要概括了北京大学计算机科学技术系网络与分布式系统实验室自 1996 年以来的研究成果。其中许多段落直接来自同学的博士和硕士论文,他们是雷鸣、赵江华、冯是聪、单松巍、谢正茂、彭波、张志刚、龚笔宏、孟涛、胥红英等。署名作者的主要工作是把这些内容系统化,使其表述的风格统一。我们特别感谢陈葆珏教授,是她在北京大学计算机系开创了搜索引擎这一研究方向,从而使我们能在其后发扬光大,还要感谢刘建国和王建勇,是他们分别带领攻关队伍,实现了天网 1.0 和天网 2.0 版本。感谢黄蕊为本书进行的文字校对。最后,我们要感谢国家九五攻关计划、973 计划和 985 计划的支持,是它们的不断支持使我们得以将天网不断推上新的台阶,实现“让天网和中国网上信息资源规模同步成长”的理想。

作 者

2004 年 5 月于北京大学燕园

目 录

第二版前言

第一版前言

第一章 引论	1
第一节 搜索引擎的概念	2
第二节 搜索引擎的发展历史	3
第三节 一些著名的搜索引擎	6
第四节 小结	11

上篇 Web 搜索引擎基本原理和技术

第二章 Web 搜索引擎工作原理和体系结构	15
第一节 基本要求	15
第二节 网页搜集	16
第三节 预处理	18
第四节 查询服务	20
第五节 体系结构	23
第六节 小结	25
第三章 Web 信息的搜集	26
第一节 概述	26
一、超文本传输协议	26
二、一个小型搜索引擎系统	27
第二节 网页搜集	30
一、定义 URL 类和 Page 类	31
二、与服务器建立连接	35
三、发送请求和接收数据	37
四、网页信息存储的天网格式	38
第三节 多道搜集程序并行工作	40
一、多线程并发工作	41
二、控制对一个站点并发搜集线程的数目	42
第四节 如何避免网页的重复搜集	43
一、记录未访问、已访问 URL 和网页内容摘要信息	43
二、域名与 IP 的对应问题	43

第五节 搜集信息的类型	45
第六节 小结	46
第四章 对搜集信息的预处理	47
第一节 索引网页库	47
第二节 网页编码识别	50
一、基本而重要的概念	50
二、常用字符编码	52
三、常用字符编码算法	55
四、字符的输入和显示	57
五、编码识别	58
第三节 中文自动分词	60
第四节 分析网页和建立倒排文件	64
第五节 小结	67
第五章 信息查询服务	68
第一节 检索的定义	68
第二节 查询服务的实现	69
一、结果集合的形成	69
二、查询结果显示	70
第三节 小结	71

中篇 对质量和性能的追求

第六章 可扩展搜集子系统	75
第一节 天网系统概述和集中式搜集系统结构	75
一、天网系统结构	75
二、集中式搜集系统	76
第二节 利用并行处理技术高效搜集网页的一种方案	82
一、节点间 URL 的划分策略	82
二、关于性能的讨论	85
三、性能测试和评价	87
四、系统的动态可配置性设计	90
第三节 天网分布式搜集系统	92
第四节 对 Deep Web 的认识	93
一、Deep Web 的成因	93
二、搜索 Deep Web 的方法	96
第五节 小结	98
第七章 网页净化与消重	100

第一节 网页净化与元数据提取	100
一、DocView 模型	102
二、网页的表示	103
三、提取 DocView 模型要素的方法	108
四、模型应用及实验研究	112
第二节 网页消重算法	115
一、消重算法	116
二、算法评测	118
第三节 小结	121
第八章 高性能检索子系统	122
第一节 检索系统基本技术	122
一、系统设计与结构	122
二、索引创建	125
三、检索过程	127
第二节 适于查询的网页索引结构	129
一、倒排索引结构	129
二、平面位置索引	131
第三节 倒排索引压缩	135
一、倒排索引压缩技术	136
二、词典与倒排表的压缩	142
第四节 索引剪枝	150
一、静态索引剪枝方法	151
二、动态索引剪枝方法	153
第五节 混合索引技术	168
一、混合索引的原理	169
二、混合索引的实现	171
第六节 倒排文件缓存机制	173
一、倒排文件缓存	174
二、负载特性	176
三、缓存策略的选择	178
第七节 小结	178
第九章 相关排序与系统质量评估	180
第一节 传统 IR 的相关排序技术	180
第二节 链接分析与相关排序	182
一、链接分析	182
二、Web 查询模式下的新信息	184

第三节 相关排序的一种实现方案	188
一、形成网页中词项的基本权重	189
二、利用链接的结构	190
三、收集用户反馈信息	192
四、计算最终的权重	194
第四节 信息检索技术评估	195
一、信息检索技术评估指标	197
二、TREC 和 CWIRF 信息检索评估	206
三、搜索引擎技术评估	213
第五节 小结	217

下篇 Web 信息资源的组织与应用服务

第十章 大规模 Web 历史网页仓储系统的构建	221
第一节 国外 Web 历史网页保存现状	221
一、Internet Archive	222
二、PANDORA	222
三、其他相关 Web 保存项目	223
第二节 中国 Web 信息博物馆的系统设计	224
一、Web InfoMall 的设计目标	225
二、Web InfoMall 的体系结构	225
第三节 历史网页的存储	227
一、数据的组织	228
二、存储结构	229
三、数据管理与压缩	230
四、存储性能	232
第四节 数据访问	232
一、PageID 的索引	233
二、URL 的索引	233
三、数据服务	234
四、性能与优化	235
第五节 网页的格式保存	236
第六节 小结	236
第十一章 大规模 Web 非网页信息仓储系统的构建	238
第一节 网络资源库藏相关工作	238
一、Ibiblio	239
二、Internet Archive	240

三、Wikimedia	240
四、中国互联网数字资源财富库藏	241
第二节 CDAL 系统概况	242
第三节 CDAL 系统设计	244
一、系统体系结构	244
二、可扩展的存储组织方案	244
第四节 网络资源描述信息获取	246
一、Ontology 概述	247
二、描述信息获取机制	247
三、改进查询的方法	248
四、改进排序的方法	249
第五节 基于局部聚类思想的共现词汇算法	250
一、基本定义	251
二、FDC 共现词汇算法	251
第六节 小结	252
第十二章 中文网页自动分类与聚类	253
第一节 文档自动分类算法的类型	253
第二节 实现中文网页自动分类的一般过程	254
第三节 影响分类器性能的关键因素分析	256
一、实验设置	256
二、训练样本	258
三、特征选取	262
四、分类算法	265
五、截尾算法	270
六、中文网页分类器的设计方案	272
第四节 天网目录导航服务	272
一、问题的提出	272
二、天网目录导航服务的体系结构	273
三、天网目录的运行实例	274
第五节 文本聚类方法	275
一、文本聚类的一般过程	275
二、文本间相似性的度量	276
三、常用聚类算法	276
四、聚类结果的评估	279
五、搜索引擎返回结果的聚类	280
第六节 小结	281

第十三章 开放域问答系统	283
第一节 概述	283
一、问答系统的历史	283
二、著名开放域问答系统介绍	284
三、开放域问答系统的通用体系结构	285
第二节 问句的分析	287
一、问句中的指代消解	287
二、问句分类	288
三、问句主题提取	290
第三节 文档和段落检索	290
一、检索模型的选用	291
二、查询生成	291
三、查询结果排序	293
四、增强索引的功能	295
第四节 答案提取和验证模块	295
一、生成候选答案集合	295
二、答案提取	296
第五节 问答系统的改进方法	299
一、问答系统中外部资源的利用	299
二、寻找特殊类问题的解决方案	301
三、通过系综方法构建问答系统	302
第六节 问答系统的评测	303
一、TREC 问答系统评测	303
二、问答系统评测指标	304
第七节 实例:天网开放域问答系统	306
第八节 小结	308
参考文献	309
附录 术语	322

图 表 目 录

图 1-1	2012 年 3 月在 Google 上检索“伊拉克战争”的结果	2
图 1-2	2012 年 3 月在 Open Directory 上检索“伊拉克战争”的结果	5
图 2-1	搜索引擎示意图	15
图 2-2	搜索引擎三段式工作流程	16
图 2-3	搜索引擎的体系结构	23
图 3-1	TSE 搜索引擎界面	28
图 3-2	TSE 查询结果页面	29
图 3-3	TSE 网页快照页面	29
图 3-4	TSE 系统结构	30
图 3-5	Web 信息的搜集	31
图 3-6	Sockets 和端口	35
图 3-7	通过 Socket 建立连接	36
图 4-1	网页预处理系统结构	47
图 4-2	原始网页库中的记录格式	48
图 4-3	索引网页库算法	49
图 4-4	字符的输入和显示流程	57
图 4-5	GB2312, Big5 和 GBK 字符编码分布	58
图 4-6	正向减字最大匹配算法流程	62
图 4-7	切词算法流程	63
图 4-8	分析网页与建立倒排文件流程	65
图 4-9	过滤网页中非正文信息算法	65
图 4-10	正向索引表记录格式	65
图 4-11	由正向索引建立反向索引	66
图 5-1	信息查询的系统结构	68
图 5-2	基本检索算法	69
图 5-3	动态摘要算法	71
图 5-4	用户查询日志的记录格式	71
图 6-1	天网系统概貌	76
图 6-2	搜集系统的主控结构	77
图 6-3	协调进程工作算法	84
图 6-4	分布式 Web 搜集系统结构	85

图 6-5	负载方差	88
图 6-6	并行搜集系统与集中式搜集系统的性能对比	89
图 6-7	分布式系统效率	89
图 6-8	URL 两阶段映射	91
图 6-9	天网分布式搜集系统 P_Arthur 体系结构	92
图 6-10	人才招聘网站首页	94
图 7-1	用 DocView 模型提取的网页要素	104
图 7-2	净化后的网页	104
图 7-3	HTML Tree 结构	105
图 7-4	内容块权值传递过程	107
图 7-5	有主题网页 DocView 模型生成过程	109
图 7-6	计算网页特征项权值的算法	109
图 7-7	正文段落识别过程	111
图 7-8	基于 anchor text 的超链选取算法	111
图 7-9	网页净化前后分类效果对比	113
图 7-10	查全率随选取关键词个数的变化	120
图 8-1	检索系统集成框架结构	124
图 8-2	天网 WWW 检索分布式系统构架	125
图 8-3	倒排索引结构示意图	129
图 8-4	按块组织的倒排链的结构	130
图 8-5	位置索引的结构	131
图 8-6	CLPS 结构示意图	135
图 8-7	倒排链中文档号之间的 d -gaps 分布图	146
图 8-8	不同文档号分配下平均每个查询对应文档号序列的压缩大小	146
图 8-9	不同压缩算法对文档号的解压速度	147
图 8-10	不同文档号分配下平均每个查询对应词频序列的压缩大小	147
图 8-11	不同压缩算法对词频的解压速度	148
图 8-12	平均每个查询对应的位置信息需要的存储空间	149
图 8-13	索引剪枝方法的分类	151
图 8-14	MAXSCORE 算法的示例	157
图 8-15	WAND 算法选择候选文档的过程	159
图 8-16	基于最大块索引的支点文档号的选择示例	161
图 8-17	Interval-Base 剪枝方法中文档子区间划分的示例	161
图 8-18	SAAT 方法处理查询处理模式及分数累加器数量的变化	164
图 8-19	当前支持高效 SR+IR 剪枝的索引结构	166
图 8-20	扩展词典树结构示例	172

图 8-21	扩展词典匹配查找算法	173
图 8-22	搜索引擎检索系统缓存结构	174
图 8-23	文档数据访问对象大小分布	176
图 8-24	I/O 与 PAGE 序列序号-频度分布	177
图 8-25	I/O 与 PAGE 序列时间间隔分布	177
图 8-26	I/O 和 PAGE 序列中唯一模式串	178
图 9-1	Inktomi 提供的几种搜索引擎技术的比较	185
图 9-2	词典在系统中的地位	186
图 9-3	新词学习	187
图 9-4	网页的互联结构示意	191
图 9-5	信息获取技术评估的“森林”	197
图 9-6	查准率和召回率基础定义图示	198
图 9-7	查准率和召回率例子	198
图 9-8	“省事的”11 点标准召回率例子	199
图 9-9	实践中召回率例子	200
图 9-10	实际中的 44 个查询词的评价统计表和 $P-R$ 图	202
图 9-11	测试集在检索评估中的角色	208
图 9-12	帮助判断相关结果页面的计算机辅助程序入口	211
图 9-13	帮助判断相关结果页面的计算机辅助程序操作界面	211
图 10-1	Web InfoMall 体系结构	226
图 10-2	网页数据的分割	229
图 10-3	Web InfoMall 的存储结构	230
图 10-4	网页的引用压缩示意图	232
图 11-1	CDAL 提供的资源访问方式	243
图 11-2	CDAL 系统结构图	245
图 11-3	基于 Ontology 的网络资源描述信息获取	248
图 11-4	概念的属性及其词汇扩展(以电影类资源为例)	249
图 11-5	获得描述信息的改进排序算法	250
图 11-6	网络资源描述信息展示	250
图 12-1	自动文档分类算法的分类	254
图 12-2	中文网页自动分类的一般过程	255
图 12-3	中文网页分类器的工作原理图	256
图 12-4	WebSmart——一个网页实例集搜集和整理工具	259
图 12-5	一种中文网页的分类体系	260
图 12-6	Macro- F_1 值随样本数的变化	261
图 12-7	Micro- F_1 值随样本数的变化	261

图 12-8	CHI、IG、DF、MI 的比较(Macro- F_1)	264
图 12-9	CHI、IG、DF、MI 的比较(Micro- F_1)	264
图 12-10	kNN 与 NB 分类结果的比较	267
图 12-11	k 的取值对分类器质量的影响(Marco- F_1)	268
图 12-12	k 的取值对分类器质量的影响(Micro- F_1)	268
图 12-13	兰式距离法与欧式距离法对 12 个不同类别的分类情况	269
图 12-14	基于层次模型的 kNN 与基本 kNN 的比较	270
图 12-15	RCut 和 SCut 截尾算法的比较	272
图 12-16	天网目录的体系结构	274
图 12-17	天网目录导航服务	274
图 12-18	文本聚类的一般过程	275
图 12-19	层次聚类实例	277
图 12-20	k -均值算法进行文本聚类的过程	278
图 12-21	搜索结果聚类系统 Carrot2	281
图 13-1	START 系统界面	285
图 13-2	Ask Jeeves 查询结果	285
图 13-3	问答系统的通用体系结构	287
图 13-4	天网开放域系统的体系结构	306
表 4-1	网页索引文件	49
表 4-2	URL 索引文件	50
表 6-1	SOIF 数据描述	78
表 6-2	SOIF 具体语法	80
表 6-3	参照序列,假设节点数为 2	87
表 7-1	类别编号对照表	113
表 7-2	消重实验结果	115
表 7-3	当 $N=10, \delta=0.01$ 时 5 种算法的查全率和准确率	119
表 7-4	考察 δ 的取值对算法 3 和 4 的影响	119
表 7-5	分段签名算法的时间复杂度及性能	120
表 7-6	基于关键词的各算法的时间复杂度及性能 ($N=10, \delta=0.01$)	121
表 8-1	MTF 对序列 $\langle 4, 4, 1, 4, 2 \rangle$ 进行转换的过程	142
表 8-2	对包含 100 万词条的词典使用不同编码所需要的空间	144
表 8-3	平均每个查询对应词频链的空间大小(文档号按 URL 序分配)	148
表 8-4	不同索引的组织结构及其支持的查询处理方式	155
表 8-5	数据集基本统计信息	176
表 9-1	新词学习对检索准确率的影响	188
表 9-2	影响权值的 HTML 标签	189