

数据库技术系列

**DATABASE
PROFESSIONALS
LIBRARY**

SQL

Server 2000

提 高

北京希望电子出版社 总策划
吴 豪 编 著

红 旗 出 版 社



北京希望电子出版社
Beijing Hope Electronic Press
www.bhp.com.cn

数据库技术系列

DATABASE
PROFESSIONAL S
LIBRARY

SQL Server 2000 提高

北京希望电子出版社 总策划
吴 豪 编 著

红旗出版社



北京希望电子出版社
Beijing Hope Electronic Press
www.bhp.com.cn

图书在版编目 (CIP) 数据

SQL Server 2000 提高 / 吴豪编著.

—北京: 红旗出版社, 2005.4

ISBN 7-5051-1041-1

I. S... II. 吴... III. 关系数据库—数据库管理系统,
SQL Server 2000 IV. TP311.138

中国版本图书馆 CIP 数据核字 (2004) 第 102821 号

内 容 简 介

本书在《SQL Server 2000 基础》上, 进一步介绍了 SQL Server 2000 数据库管理与开发的高级知识。

本书由 18 章和两个附录组成。主要内容包括范式化设计、SQL Server 存储与索引结构、事务与锁、高级查询、XML 简介、XML 集成、全文检索、英文查询、分布式查询与事务、数据转换服务、高级 DTS、大容量复制程序 (bcp)、高级设计、分析服务、数据库安全性、性能测试、管理概览等内容。将各个知识点与具体实例结合, 以印证知识使用的方法。另外, 在附录部分提供了 SQL Server 2000 的系统和操作函数。

本书适合于 SQL Server 数据库管理员和高级开发者, 重点面向开发者。结合《SQL Server 2000 基础》, 对初学者的继续提高也有很大帮助。

本书部分素材请从 www.b-xr.com 下载。

需要本书或技术支持的读者, 请与北京中关村 083 信箱 (邮编: 100080) 发行部联系, 电话:
010-82702660, 82702658, 62978181 (总机) 转 103 或 238 传真: 010-82702698 E-mail: tbd@bhp.com.cn。

系 列 名 数据库技术系列

书 名 SQL Server 2000 提高

编 著 吴 豪

责 任 编 辑 但明天 雷 铎

出 版 发 行 红旗出版社 北京希望电子出版社

地 址 红旗出版社 北京市沙滩北街 2 号 (100727) 电话: (010) 64037138
北京希望电子出版社 北京市海淀区上地三街 9 号金隅嘉华大厦 C 座 610

经 销 各地新华书店 软件连锁店

排 版 希望图书输出中心 王春桥

印 刷 北京媛明印刷厂

版 次 / 印 次 2005 年 4 月第 1 版 2005 年 4 月第 1 次印刷

开 本 / 印 张 787×1092 1/16 37.875 印张 890.7 千字

印 数 1-5000 册

书 号 ISBN 7-5051-1041-1

定 价 52.00 元

前 言

SQL Server 2000 在 SQL Server 6.5 和 SQL Server 7.0 的基础上有了飞跃的发展,增加了不少功能:XML 集成、分布式分区视图、索引视图、INSTEAD OF 触发器、用户自定义函数、两阶段数据抽取——用于数据转换服务(DTS)、数据挖掘等。作为目前应用广泛的一种数据库,它不仅采用了合理的体系结构,而且利用了优越的可视化开发环境,加上其优秀的编辑界面,使其很多性能又远远超过了 Oracle 及其他数据库,从而为广大用户所青睐。目前有不少应用软件利用了 SQL Server 2000 作为后台存储数据库,如部分医院信息管理系统,银行、证券管理系统等。为了使越来越多的用户逐步深入掌握 SQL Server 2000 这种优秀的数据库,笔者编写了数据库技术系列这套丛书,与大家分享。

笔者是一位具有多年数据库研发经验的资深系统分析员和软件工程师,用过 SQL Server 和 Oracle 等大型数据库,开发过医院信息管理系统和远程医疗视频会议等大型管理系统。在学习和研究这些数据库的过程中,笔者深深地感受到学习过程的艰辛,为把推敲再三的心得殷实详尽地记录下来,为后学者铺出一条广阔的道路,笔者与同事一起系统地编写了这套丛书。

这是一套 SQL Server 进阶的系列丛书,更是一组经验的结晶,可满足不同层次读者的切实需要。本书适合于 SQL Server 数据库设计领域的初学者和高级开发者,但重点面向开发者。这意味着,为了简明扼要,那些适合于数据库管理员而不适合于开发者的内容有时一带而过,甚至避而不谈,但对开发者有影响或在开发过程中需要考虑的管理因素,会在部分章节对相关内容进行简要阐述。对于初学者或具有一定关系数据库经验而不熟悉 SQL Server 的人来说,使用这套丛书可以迅速利用 SQL Server 进行数据库开发,设计出自己需要的数据库结构;对于具有一定基础的中级水平的开发人员,可以利用本书提供的内容,融会贯通、灵活运用,并利用本丛书提供的最新特点和使用方法,帮助自己提高开发和运用工具的水平;对于高级开发人员,这是一套不可多得的参考书,书中讨论的内容表明:本版软件包含了更多内容,SQL Server 将以一个全新的面孔展现出精彩的世界。

本丛书目的在于使读者阅读以后能在前进的道路上掌握更高级的技术,假设你已经是一个有经验的开发者(并不一定在数据库方面)。为了彻底理解本书,虽然需要懂得编程的基础知识,如变量、数据类型、程序化设计等,但并不要求以前就看过有关查询的内容。

从头到尾通读全书,就可以很好地掌握本书介绍的相关内容。本书侧重于帮助读者利用 SQL Server 2000 成功地开发应用程序,并为备战微软认证考试的人提供帮助,因此笔者会简要介绍一些力求更完美的内容。

这套书不是万能工具书,但它能够提供有关 SQL Server 2000 开发的许多信息,如 DTS、英文查询和联机分析处理(OLAP)等。

内容结构

本套丛书分为《SQL Server 2000 基础》和《SQL Server 2000 提高》两本(以下简称《基础》和《提高》)。

为了使读者迅速利用 SQL Server 进行数据库开发,设计出自己需要的数据库结构,笔

者从数据库设计的角度出发,先介绍如何使用 SQL Server 2000 必备的工具,然后讨论如何进行数据库设计和开发应用,最后就高级数据库设计中可能遇到的各种问题做详细剖析。在《基础》中,笔者介绍了 SQL Server 2000 的历史背景、数据对象、常用工具、Transact-SQL、联接、创建和修改表、约束、视图、脚本和批处理、存储过程、用户自定义函数、触发器及游标等内容。在《提高》中,笔者先讨论了如何进行范式化设计,并就数据库设计中不得不讨论的索引结构、事务和锁展开探讨,再结合目前广为应用的 XML 技术,探讨了如何将 SQL Server 2000 与 XML 结合起来,并就如何利用 XML 技术进行高级查询、全文检索、英文查询做了深入介绍,同时就如何设计合理的分布式分区视图展开描述,最后就数据转换服务和数据库之间的数据复制等内容做了详细讨论。

运行环境

本书讨论的所有内容几乎都有例子,而且列出了相应的代码。为了能够更好地运行本书提供的示例,需要配备:

- SQL Server 2000
- Windows NT 4.0 和 SP5 补丁,最好使用 Windows 2000 (Windows 9X 也可以工作,但在某些方面会失败)

约 定

本书采用许多不同文本方式和显示方式,用于帮助区分不同内容。其中:

- 对于需要特别注意和需要提示的内容,相应的文字使用楷体,并在段落前标以“注意:”或方框字样。例如:

注意: Personal 版本虽然能与其他工具协调使用,但建议不要这样使用,因为微软不支持这种用法。

- 新代码或重要代码的格式设为:

```
SELECT CustomerID, ContactName, Phone  
FROM Customers
```

- 以前见过的代码或者无须讨论的代码格式为:

```
SELECT ProductName FROM Products
```

- 背景信息、附加信息和参考内容的字体格式如下:

20 世纪 80 年代,美国国家标准化组织 (American National Standards Institute, ANSI) 制订了 SQL 语言规范, ANSI-SQL 诞生了,这是 RDBMS 计算领域的关键时刻。

由于本书篇幅有限,加之时间仓促,难免有疏漏、不足之处,诚请广大读者提出宝贵意见并批评、指正。

目 录

第1章 范式化设计.....	1	2.4 选择创建索引的时机.....	50
1.1 表.....	1	2.4.1 可选择性.....	50
1.2 实现数据范式化.....	1	2.4.2 成本.....	51
1.2.1 准备工作.....	2	2.4.3 选择聚集索引.....	51
1.2.2 第一范式.....	4	2.4.4 列顺序问题.....	53
1.2.3 第二范式.....	6	2.4.5 删除索引.....	53
1.2.4 第三范式.....	7	2.4.6 利用索引调试向导.....	53
1.2.5 其他范式.....	7	2.5 维护索引.....	54
1.3 关系.....	8	2.5.1 碎片.....	55
1.3.1 1-1 关系.....	8	2.5.2 定义碎片与页拆分的可能性.....	55
1.3.2 1-1 或 1-多.....	9	2.6 小结.....	59
1.3.3 多对多.....	10	第3章 事务和锁.....	60
1.4 关系图.....	12	3.1 事务.....	60
1.5 降低范式化.....	21	3.1.1 BEGIN TRAN.....	61
1.6 超出范式化.....	21	3.1.2 COMMIT TRAN.....	61
1.6.1 保持简化.....	21	3.1.3 ROLLBACK TRAN.....	61
1.6.2 选择数据类型.....	22	3.1.4 SAVE TRAN.....	61
1.6.3 有关存储的错误.....	22	3.2 SQL Server 日志的工作方式.....	66
1.7 示例.....	22	3.2.1 失败和恢复.....	68
1.7.1 创建数据库.....	23	3.2.2 隐式事务.....	69
1.7.2 添加关系图和初始表.....	23	3.3 锁和并发性.....	70
1.7.3 添加关系.....	26	3.3.1 锁可以阻止的问题.....	71
1.7.4 添加一些约束.....	27	3.3.2 可以加锁的资源.....	74
1.8 小结.....	28	3.3.3 锁增加和锁对性能的影响.....	74
第2章 SQL Server 存储和索引结构.....	29	3.3.4 锁模式.....	75
2.1 SQL Server 存储历史.....	29	3.3.5 锁的兼容性.....	77
2.1.1 SQL Server 存储.....	30	3.3.6 控制特定锁的类型.....	77
2.1.2 在 6.5 版和以前存储.....	32	3.4 设置隔离层.....	83
2.1.3 7.0 及以后的 SQL Server 存储.....	34	3.5 处理死锁.....	85
2.2 索引定义.....	37	3.5.1 怎样指出存在死锁.....	85
2.2.1 B-树.....	38	3.5.2 怎样选择死锁牺牲品.....	86
2.2.2 在 SQL Server 中访问数据.....	40	3.5.3 避免死锁.....	86
2.2.3 索引类型和索引遍历.....	41	3.6 小结.....	88
2.3 创建和删除索引.....	46	第4章 高级查询.....	89
2.3.1 CREATE INDEX 语句.....	46	4.1 什么是子查询.....	89
2.3.2 创建约束时暗含的索引.....	49	4.2 建立嵌套子查询.....	90
		4.2.1 使用单值 SELECT 语句嵌套	

查询	90	6.2 HTTP 访问	158
4.2.2 使用子查询返回多个值的 嵌套查询	91	6.2.1 设置 HTTP 访问	159
4.2.3 ANY、SOME 和 ALL 操作符	94	6.2.2 基于 URL 的查询	164
4.3 关联子查询	95	6.2.3 使用模板	165
4.3.1 关联子查询是怎样工作的	95	6.2.4 POST	168
4.3.2 WHERE 子句中的关联子 查询	95	6.2.5 XPath	172
4.3.3 SELECT 清单中的关联子 查询	97	6.2.6 设计结果的样式	175
4.4 衍生表	98	6.2.7 更新程序简介	178
4.5 EXISTS 操作符	100	6.3 使用 IcommandStream 设置 XML 命令	178
4.6 混合数据类型 CAST 和 CONVERT	103	6.4 XML 相关资源	200
4.7 性能比较	105	6.5 小结	202
4.8 小结	107	第 7 章 全文检索	203
第 5 章 XML 简介	108	7.1 全文检索体系结构	204
5.1 XML 基础	108	7.2 设置全文索引和分类	205
5.1.1 XML 文档的组成	110	7.2.1 授予数据库全文搜索能力	205
5.1.2 良构 XML	114	7.2.2 创建全文目录	206
5.1.3 XML 实例	116	7.2.3 给独立的表启用全文检索 功能	207
5.1.4 确定元素和属性	119	7.2.4 索引组装	210
5.2 命名空间	120	7.3 全文查询语法	211
5.3 DTD 和架构	122	7.3.1 CONTAINS	212
5.3.1 DTD	123	7.3.2 FREETEXT	213
5.3.2 XML 架构	125	7.3.3 CONTAINSTABLE	213
5.3.3 DTD/架构和性能总结	125	7.3.4 FREETEXTTABLE	214
5.4 转换——XSLT	126	7.3.5 近似 (Proximity)	216
5.5 小结	131	7.3.6 前缀条件	217
第 6 章 XML 集成	132	7.3.7 权重	217
6.1 FOR XML 子句	133	7.3.8 词尾变化	218
6.1.1 RAW	134	7.3.9 对等级的简单总结	218
6.1.2 AUTO	135	7.4 无效单词	219
6.1.3 EXPLICIT	136	7.5 所支持的语言	219
6.1.4 OPENXML	146	7.6 sp_fulltext_service	220
6.1.5 XML 语法规则	154	7.7 实现全文本检索	221
6.1.6 元素的语法	156	7.8 小结	222
6.1.7 注释的语法	156	第 8 章 英文查询	223
6.1.8 CDATA 的语法	156	8.1 什么是英文查询	224
6.1.9 Namespaces 的语法	157	8.2 英文查询的优点和局限性	225
6.1.10 entity 的语法	157	8.3 英文查询应用程序的体系结构	225
		8.4 创建英文查询工程	226

8.4.1 使用 SQL 工程向导	228	10.2 使用导入/导出向导	294
8.4.2 基本英文查询	232	10.2.1 Execute SQL	299
8.4.3 提炼和扩展实体和关系	235	10.2.2 工作流	300
8.4.4 提炼和扩展字段实体	236	10.2.3 连接对象	300
8.4.5 定义附加的短语	238	10.2.4 转换	300
8.4.6 保存查询和回归测试	239	10.3 创建一个简单的转换包	301
8.4.7 定义表之间的关系	240	10.4 保存包	311
8.4.8 基于多个联接的关系	242	10.5 使用 DTS 代替 bcp	312
8.4.9 表和字段实体之间的短语和关系	243	10.6 从命令行运行 dtstrun	314
8.4.10 定义子集短语	244	10.7 小结	316
8.4.11 定义形容词短语	246	第 11 章 高级 DTS	317
8.4.12 定义同义词	249	11.1 DTS 对象模型介绍	317
8.4.13 字典条目	249	11.2 Package 对象	319
8.5 重新测试应用程序	251	11.3 动态包	321
8.6 分发英文查询应用程序	252	11.4 动态属性任务	323
8.6.1 工程文件	252	11.5 多阶段数据抽取	327
8.6.2 建立英文查询应用程序	255	11.5.1 启用多阶段数据提取	328
8.6.3 使用英文查询应用程序	255	11.5.2 多阶段数据抽取的例子	329
8.6.4 分发	263	11.6 在 DTS 中优化数据加载	335
8.7 小结	263	11.7 小结	336
第 9 章 分布式查询和事务	265	第 12 章 大容量复制程序 (bcp)	337
9.1 分布式事务	265	12.1 bcp 的作用	337
9.1.1 准备阶段	266	12.2 bcp 导入	340
9.1.2 提交阶段	266	12.2.1 数据导入的例子	341
9.2 分布式查询	267	12.2.2 记录日志与不记录日志	344
9.2.1 创建链接服务器	267	12.3 格式化文件	344
9.2.2 使用链接服务器	269	12.3.1 不匹配的字段顺序	347
9.2.3 在链接服务器上执行 sprocs	273	12.3.2 使用格式化文件	347
9.2.4 从远程服务器上收集元数据	274	12.3.3 最大化导入性能	348
9.2.5 创建和使用通道查询	277	12.4 bcp 导出	348
9.2.6 在远程数据源上使用特别查询	278	12.5 大容量插入	349
9.2.7 其他分布式查询的注意事项	279	12.6 小结	350
9.3 联合服务器 (分布式分区视图)	281	第 13 章 复制	351
9.4 小结	288	13.1 支持分布式数据	352
第 10 章 数据转换服务	289	13.2 架构复制时的考虑	352
10.1 DTS 包编辑器	289	13.2.1 独立性	352
10.1.1 连接 (connection)	289	13.2.2 延迟性	352
10.1.2 任务	291	13.2.3 数据一致性	353
10.1.3 工作流	293	13.2.4 架构一致性	353
		13.2.5 其他考虑	353

13.3	发布比喻	353
13.3.1	订阅刊物	354
13.3.2	订阅服务器类型	355
13.3.3	过滤数据	355
13.4	复制类型	355
13.4.1	快照复制	356
13.4.2	合并复制	358
13.4.3	事务复制	360
13.4.4	立即更新的订阅服务器	362
13.4.5	混合复制类型	363
13.5	复制模型假设	363
13.5.1	标准模型	363
13.5.2	混合模型	365
13.5.3	附加说明	366
13.6	实现示例	366
13.6.1	会诊专家申请方案	367
13.6.2	交互	367
13.7	计划复制	368
13.7.1	数据考虑	368
13.7.2	复制类型	369
13.8	复制向导	369
13.9	启用发布和分发	369
13.9.1	启动向导	370
13.9.2	默认配置	371
13.9.3	配置之后	373
13.9.4	禁用发布	374
13.9.5	Transact-SQL 过程	374
13.10	事务/快照发布刊物	375
13.10.1	创建和管理发布向导	375
13.10.2	配置之后	380
13.10.3	Transact-SQL 过程	381
13.11	合并发布刊物	382
13.11.1	创建和管理发布刊物向导	382
13.11.2	配置之后	386
13.11.3	Transact-SQL 存储过程	387
13.12	强制订阅刊物	388
13.12.1	强制订阅刊物向导	388
13.12.2	关于合并订阅	390
13.12.3	Transact-SQL 过程	390
13.13	请求订阅刊物	390

13.13.1	请求订阅刊物向导	391
13.13.2	Transact-SQL 过程	392
13.14	管理复制	392
13.14.1	复制脚本	392
13.14.2	支持各种复制	392
13.14.3	发布到 Internet	393
13.14.4	复制和架构变化	394
13.14.5	复制监视器	396
13.14.6	后续	397
13.15	小结	398
第 14 章	高级设计	399
14.1	讨论更多的关系图和关系	399
14.1.1	一组关系类型	400
14.1.2	实体盒	400
14.1.3	关系线	401
14.1.4	终止符	401
14.2	逻辑与物理设计	402
14.2.1	逻辑模型的目的	402
14.2.2	逻辑模型的组成	403
14.3	处理基于文件的信息	404
14.4	子类	406
14.4.1	子类的类型	406
14.4.2	实现子类	407
14.4.3	子类的物理实现	408
14.4.4	把子类添加到扩展属性中	409
14.5	数据库重用	409
14.5.1	候选的重用数据库	410
14.5.2	怎样分解	410
14.5.3	高代价的重用性	411
14.6	分区实现扩展性	411
14.7	小结	412
第 15 章	分析服务	413
15.1	终端用户的要求	413
15.1.1	联机事务处理 (OLTP)	413
15.1.2	联机分析处理 (OLAP)	414
15.1.3	数据挖掘	415
15.1.4	OLTP 或 OLAP	415
15.2	维度数据库	416
15.2.1	事实数据表	417
15.2.2	维度表	417

15.2.3	星型和雪花型架构	417	16.5	应用程序角色	468
15.2.4	数据多维集	418	16.5.1	创建应用程序角色	468
15.3	OLAP 存储类型	418	16.5.2	添加应用程序角色许可权	469
15.3.1	MOLAP	419	16.5.3	使用应用程序角色	469
15.3.2	ROLAP	419	16.5.4	删除应用程序角色	470
15.3.3	HOLAP	419	16.6	XML 的安全性	470
15.4	数据仓库的概念	419	16.7	更高级的安全性	470
15.4.1	数据仓库特性	420	16.7.1	处理 Guest 账户	470
15.4.2	数据集市	421	16.7.2	TCP/IP 端口设置	471
15.5	数据转换服务	421	16.7.3	不要使用 sa 账户	472
15.5.1	数据校验	421	16.7.4	围绕 xp_cmdshell	472
15.5.2	数据清理	422	16.7.5	视图、存储过程和 UDF	472
15.5.3	数据移植	422	16.8	小结	473
15.5.4	数据转换	422	第 17 章	性能调试	474
15.5.5	DTS 组件	423	17.1	调试时机	474
15.6	元数据和知识库	423	17.2	索引选择	475
15.7	数据挖掘模型	423	17.2.1	全覆盖索引	476
15.8	数据挖掘算法	424	17.2.2	索引调试向导	476
15.8.1	决策树	424	17.2.3	索引调试向导的“清闲”	477
15.8.2	聚集	425	17.3	客户与服务器端的处理	477
15.9	分析管理器	426	17.4	策略性地降低范式化要求	478
15.10	建立数据挖掘模型	438	17.5	程序维护	478
15.11	小结	441	17.6	很好地组织 Sprocs	479
第 16 章	数据库安全性	443	17.6.1	使事务尽可能短	479
16.1	安全性基础	445	17.6.2	使用限制最少的事务 隔离层	479
16.1.1	一个人、一个帐号和口令	445	17.6.3	执行多种解决方案	479
16.1.2	口令有效期	446	17.6.4	避免使用游标	480
16.1.3	口令长度和组成	448	17.7	使用临时表	480
16.1.4	登录次数	449	17.8	细节考虑	480
16.1.5	用户和口令信息存储	449	17.9	硬件考虑	481
16.2	安全性选项	450	17.9.1	服务器的独占使用	481
16.2.1	SQL Server 安全性	451	17.9.2	I/O 与 CPU 速度	482
16.2.2	NT 集成安全性	454	17.9.3	OLTP 与 OLAP	485
16.3	用户权限	456	17.9.4	在线或离线	485
16.3.1	对特定数据库授予访问权	456	17.9.5	宕机的危险	486
16.3.2	授予对象许可权	457	17.9.6	丢失数据	486
16.3.3	用户权限和语句级许可权	462	17.9.7	性能不是全部	486
16.4	服务器和数据库角色	463	17.9.8	硬盘支持	487
16.4.1	服务器角色	464	17.9.9	理想系统	487
16.4.2	数据库角色	465			

17.10 问题纷争的解决方法.....	487	18.5 数据库复制.....	523
17.10.1 显示计划和统计信息.....	488	18.5.1 复制数据库向导.....	523
17.10.2 数据库一致性检验程序.....	492	18.5.2 备份和恢复.....	525
17.10.3 查询控制器.....	492	18.5.3 连接/分开.....	525
17.10.4 SQL Server Profiler.....	493	18.6 索引重建.....	526
17.10.5 性能监视器 (Perfmon) ...	496	18.7 归档数据.....	526
17.11 小结.....	497	18.8 小结.....	527
第 18 章 管理概览.....	498	附录 A 系统函数.....	528
18.1 调度表作业.....	498	附录 B 函数列表.....	536
18.1.1 创建操作员.....	499	B.1 聚集函数.....	536
18.1.2 创建作业和任务.....	502	B.2 游标函数.....	538
18.2 备份和恢复操作.....	513	B.3 日期和时间函数.....	539
18.2.1 备份介质.....	514	B.4 数学函数.....	541
18.2.2 备份.....	515	B.5 元数据函数.....	544
18.2.3 恢复数据.....	517	B.6 行集函数.....	552
18.3 自动响应警报.....	520	B.7 安全性函数.....	554
18.3.1 在 EM 中创建响应警报.....	520	B.8 字符串函数.....	556
18.3.2 在 T-SQL 中创建响应警报.....	521	B.9 系统函数.....	559
18.4 全文目录操纵.....	522	B.10 文本和图像函数.....	568
18.4.1 备份与恢复.....	522	B.11 小结.....	568
18.4.2 安排注入的调度.....	523		

第 1 章 范式化设计

如果本书是你第一次真正涉足数据库，那么对迄今为止怎样和为什么要构造表可能仍然感到困惑。

本书讨论的内容基本都在 OLTP (Online Transaction Processing, 联机事务处理) 环境下进行，书中将讨论 OLTP 和其更具面向分析的 OLAP (Online Analytical Processing, 在线分析处理) 之间的差别。需要指出的是，在大部分例子中，会看到最常见的 OLTP 数据库表的优化设计，这种方式的数据库设计问题属于第三范式下数据库布局的范畴。

那么，什么是“范式”呢？简单地说，范式就是将数据分成逻辑、非重复的格式，并且容易组合成一个整体。除了范式化（它将数据变成标准格式）之外，我们还要验证 OLTP 和 OLAP 数据库的功能。此外，还要在本章介绍实现“约束”的许多例子。

本章内容对部分读者而言，可能是最难掌握的章节之一。建议仔细研读本章，充分掌握其中的细节，从而为今后设计自己的数据库打下坚实的基础。

1.1 表

在《SQL Server 2000 基础》（以下简称《基础》）中，已经对表进行了大量的阐述，这些内容似乎超出了基础知识的范围，但还是简要复习一下什么是表。

数据库表不是无关紧要的，而是数据库的核心对象。表是具有相同通用属性 (attributes) 的实例数据集合，这些数据实例组织为数据行 (row) 和列 (column)。一张表应该提供实际的数据集合，而且与其他表有关联 (relationships) 信息。各种实体 (表) 和关系 (它们怎样协同工作) 的描述通常指一个实体关系图表 (或者叫 ER 图)，有时将 ER 图写为 ERD (Entity-Relationship diagram)。

通过各种关系连接两张或多张表，可以根据需要在这些表的组合数据中临时创建其他表 (在《基础》的第 4 和第 5 章中已经对此进行过深入讨论)。关系实体集以成组的方式存放在数据库中。

1.2 实现数据范式化

范式是现代 OLTP 数据库设计的基础模型，它首先沿用了关系数据库的概念。范式和关系数据库的概念源自 E.F.Codd (IBM 公司) 1969 年的专著。当时，Codd 首先提出了数据库“由一系列可以返回表的非过程化操纵的无序表组成”的概念，这种概念在当时并没有引起计算机领域的研究人员的重视。在随后的几十年里，经过像 Sybase 和 Oracle 等公司的研究和开发，以及 ANSI 的不懈努力，最终形成了真正的关系数据库产品，从而使关系数据库最终变成当今占主流地位的数据库产品。

实现数据范式化的关键在于：

- 次序不是很重要。
- 一张基表应以非过程化方法与其他表相关（Codd 称表为关系）。
- 通过关联这些基表，应该能创建一个满足需要的虚拟表。

范式化是关系数据库设计的一个分支。

在编程过程中，范式化的概念已经成为被滥用的甚至是被误解的概念之一。人人都认为自己理解了这一概念，而且许多人至少在学术中做到了这一点。不幸的是，不少数据库设计者对此像走过场——在某种程度上，他们像“真正”的数据库设计师。实际上，他们也仅知道什么是范式。范式其实就是大型数据库设计的一张图纸，有时需要实现数据的范式，但有时需要再次故意使数据不满足范式的要求。即使在范式化处理过程中，通常也可有多种途径得到技术上的范式化数据库。

范式化是一种理论，而且仅是一种理论。一旦选择了是否采用范式化策略，就决定了你的数据库拥有的内容——满怀希望地期待它可能是你设计出的最好数据库。不要沉迷于书本知识，最好根据实际情况进行处理。这里介绍的内容只是告诉你相关的概念。为了得到最佳的解决方案，你需要掌握这些概念，但并不要求你按照笔者的想法去设计自己的数据库。下面开始讨论范式的形式和作用。

我们首先从 6 种范式形式说起。对那些已经接触过数据库和范式化设计的人来说，或许你对有 6 种范式感到惊奇。你可能听到完整的范式化设计通常归结为 3 种形式——难道不止这 3 种吗？本书重点讨论常用的 3 种形式，但为了后续的需要，还将对其他 3 种形式做简要介绍。

我们已经看到了怎样创建主键，并看到了在表中使用主键的原因，这就是：如果你只想返回一行，就需要能惟一定义该行。范式化的概念在很大程度上是围绕行的主键定义和主键所依赖的列展开的。在范式化设计中，最常见得说法是：

键，乃至整个键，什么也不是，就是一个键。

这是对第三种范式的最简短解释。如果有人说所有的列仅依赖于整个键（不多也不少），那么他所说的就是第三范式。

下面来看看各种标准范式，了解每种范式是怎样工作的。

1.2.1 准备工作

在设法把数据转换成第一范式化格式之前，还需要进行适当的准备。用关系数据库意义上的词来说，在将表考虑为实体之前，需要考虑以下几件事情：

- 表应该描述一个实体。
- 所有行必须惟一，而且必需有一个主键。
- 列和行的顺序不受影响。

这些内容有些非常明显，而有些则不明显，我们将在范式化设计过程中详细讨论这些内容。首先要定义实体，然后仔细检查并定义所有明显的实体。

面向对象的编程方法（OOP）通常会在对象模型中把最符合逻辑的上层实体比作对象。

下面考虑一个非常简单的模型——再次考虑库房管理模型。开始，我们不考虑各种变化，甚至不考虑拥有的列。相反，我们只考虑系统中将要定义的基本实体。

首先，考虑最基本的处理。在处理过程中，想为每一个原子体创建一个实体，从而能够维护其中的数据。

假设第一个要传递的实体是入库单。

如果NI是一位具有丰富范式化设计经验的系统设计师，那么你的第一印象可能是从开始就产生一些实体。但现在只采用这个实体，看看范式化处理进程如何给我们提供所需要的结果。

假定NI已经掌握了有关实体的概念，下一步要指出哪个列产生主键。记住，主键为每一行提供了惟一标识。

仔细看看列的清单并提供键的候选人（key candidates）。候选键的清单应该包括惟一标识实体中每一行的所有列。另外，规则可以指出哪一列能成为主键（这就是许多人在设计数据库时包含了相同基本信息的原因之一）。如果找不到候选键，就要采用组合的方法来形成一个候选键。表 1-1 是我们创建一张入库记录 IMaster。

表 1-1 入库记录表 IMaster

中文字段名	英文字段
入库编号	INo
入库日期	IDate
厂家编号	ISource
厂家名称	FacName
联系人	Connect
联系电话	Telephone
联系地址	Address
入库方式	IMode
入库设备	IstrumItem

因为这是一张入库表，一个入库编号就意味着一张入库单，所以要用 INo 作为主键。

现在已经拥有了基本实体，即入库记录表 IMaster，这个实体并不关心列的顺序（根据传统，表的组成通常将第一列作为主键，但从技术角度来讲，并不需要这样做）。在表的基本组成中，没有哪一列关心行的顺序。表面上，表描述的对象仅是一个实体。下面准备开始范式化进程（实际上，我们已经开始了）。

1.2.2 第一范式

第一范式 (The first normal form, 简称 1NF) 完全是定义消除数据重复并保证数据的原子性 (atomicity, 该数据是自包含而且是独立的)。它在高层创建主键, 把重复数据移到新表中, 再为这些表创建新键, 还把数据分成由列组成的单行数据。

在更传统的平面设计中, 经常会看见很多重复数据 (因为多种信息都存放在一列中), 这会带来很多问题:

- 磁盘存储非常昂贵, 多次存储数据意味着浪费空间。目前数据存储的成本已经不再昂贵, 所以这已经不是一个问题。
- 重复数据意味着要移动更多的数据, 从而导致大的 I/O 流量。这意味着性能受阻, 因为后台的大量数据块必须通过数据总线和网络移动, 即使当今的数据传输技术发展速度非常快, 在性能上也会有负面影响。
- 应该成为重复数据行之间的数据通常不一致, 并产生某种自相矛盾或缺乏完整性的数据。
- 如果想从组合数据列中查找信息列, 首先就得采用一种方式解析该列的数据 (这种工作相当慢)。
- 通过前一节的介绍, 我们知道该表中已经拥有很多列, 你还可以再加入一些列。

如果想了解入库单的所有内容, 最好是从一个地方查询出所有结果。

下面探究一下这到底意味着什么, 让我们看看表中的一些数据 (这里有意删减了一些列, 表 1-2)。

表 1-2 入库砍数据表

INo	IDate	ISource	FacName		IMode	IstrumItem
100	2004/01/06	0001	Linkstar		购买入库	卫星调制解调器 150 个, 单价 50000.00 元 串口服务器 150 个, 单价 3000.00 元 视频通信终端 200 套, 单价 60000.00 元
101	2004/01/06	0002	创佳公司		购买入库	LNB (低噪声放大器) 200 个, 单价 3000.00 元, 保质期 2 年 BUC (功率放大器) 200 个, 单价 20000.00 元, 保质期 2 年
102	2004/02/13	0001	Linkstar		赠送入库	① SUN 工作站一台, 估价 10 万元
103	2004/01/14	0004	新通公司		购买入库	① 网线 2 箱 (600 米/箱), 总价 1000.00 元

如果继续对该表实现范式化, 还需要处理一系列问题。虽然我们已经拥有了一个功能主键 (主键已经存在于关系系统中很久了), 但在第一范式中仍然要进行适当的处理。

- 存在数据重复组（入库来源），需要将该数据分成不同的表。
 - 入库设备（IstrumItem）列不包含本质上是原子的数据。
- 我们可以从该表中分离出几列，以形成新表 FacName 中（表 1-3 和表 1-4）。

表 1-3 分离入库单

ISource	FacName		IMode	IstrumItem
0001	Linkstar		购买入库	卫星调制解调器 150 个，单价 50000.00 元 串口服务器 150 个，单价 3000.00 元 视频通信终端 200 套，单价 60000.00 元
0002	创佳公司		购买入库	LNB（低噪声放大器）200 个，单价 3000.00 元，保质期 2 年 BUC（功率放大器）200 个，单价 20000.00 元，保质期 2 年
0001	Linkstar		赠送入库	① SUN 工作站一台，估价 10 万元
0004	新通公司		购买入库	① 网线 2 箱（600 米/箱），总价 1000.00 元

表 1-4 表 FacName

ISource	FacName	Connect	Telephone	Address
0001	Linkstar	卓娅	13006883113	北京市复兴路 33 号
0002	创佳公司	黄忠	010-82316945	北京市海淀区路 24 号
0004	新通公司	张剑锋	020-43561234	上海市华山路 30 号

在新表和旧表之间，需要注意：

- 在新表中必需有一个主键以确保每一行数据惟一。新表有两个候选键——厂家编号（ISource）和厂家名称（FacName）。厂家编号就是为了这种目的而创建的，它似乎仅是一个逻辑选择。
- 尽管从 IMaster 表中移出了有关厂家信息的数据，但仍然需要在新 IMaster 表中引用 FacName 表中的数据，这就是为什么在入库主记录表中仍然需要入库来源（主键）列的原因。建立引用时，将创建一个外部键约束，强迫所有入库单都拥有一个有效厂家编号。
- 对 Linkstar 有限公司而言，可以省略实例信息，这就是移动重复组中的数据的目的一只存储一次；这样既节省了空间又避免了冲突值。
- 只移动重复组数据。我们仍然可以看到多条记录的入库单日期相同，但它实际不在一个组中——它仅是该表中没有适当边界的相关随机数据。

前面已经处理了重复数据，下面准备讨论第一范式的第二种非常规行为。如果看一看 IstrumItem 列，就会看到实际有许多存在差别的数据，比如设备名称、设备规格和单价信息。

如果维持设备名称、规格和单价信息不变，那么它们就是原子信息，一旦把它们组合在一起，就不再是原子体了。

无论相信与否，有时事情就是以这种方式处理的。乍一看，似乎这是一件非常简单的事情。

下面继续分解——而且，要根据单价添加一些新信息。问题是，一旦将这些信息分解，主键就不再能惟一定义行信息了——行仍然惟一，但主键不再惟一。

现在，我们在表中添加流水号 IItemNo 作为主键的一部分，从而可以再次惟一定义行。因为每张入库单只要具备了入库单号和入库流水号，就能准确定位是哪一条记录。

注意：也可以不创建另一列，而将产品编号作为主键的一部分，但这样最终会导致在同一张订单中不能让同一种产品出现两次。我们将在第 2 章中简要讨论基于多列上的键——组合键。

此时已经满足了第一范式的标准：没有数据重复组，而且所有列都是原子体。但是还有一些列不得不存在重复数据的问题，随后将会对此进行讨论。

1.2.3 第二范式

范式化的下一个阶段是第二范式 (2NF)，第二范式进一步减少了重复数据的范围（不一定分组）。

第二范式有两个规则：

- 表必须满足第一范式的规则（范式化是一个建筑块的处理过程——如果没有准备好前两种范式，就不可能拥有第三范式）。
- 每个列必须依赖整个键。

该例子存在一个问题——实际上，在该区域中有一对规则。再看一看第一范式下的 IMaster 表，它的每一列都依赖整个键吗？难道真需要这部分键吗？

答案是既是也不是。有几列仅依赖于 INo 列，而不是 IItemNo 列。它们是：Idate, Isource, Imode, IAuthor, Ismanager 和 IHandler。无论该入库单上有多少条记录，这几列对整个入库单而言都一样。处理这些问题需要创建另一张表，这要涉及头表(header)与明细表(detail)的概念。

实际上，一个实体有时仍然需要分成两个实体（两张表）。头表在关系上是明细表的父表，它包含的信息只需要保存一次，而明细表存储的信息则可能存放于多个实体中。头表通常保存了原始表名，而明细表则以头表开头并添加了一些明细信息（如 IDetail）。头表的每一条记录通常至少有一条明细记录，也可能有多条明细记录，这是下一节要了解的一种对应关系（一对多的对应关系）。

为了讨论相关内容，我们需要再次对该表进行分离。首先考虑明细表，将它称为入库单明细表 (IDetail)：Ino (PK), IItemNo (PK), Item_code, Item_spec, units, batch_no, expire_date, firm_id, purchase_price, discount, retail_price, package_spec, quantity, package_units, sub_package_1, sub_package_units_1, sub_package_spec_1, sub_package_2, sub_package_units_2, sub_package_spec 2, invoice_no, invoice_date。

尽管可以把它当做两张表之间的新表，但还需要头表 Imaster，让它保存入库单名：INo (PK), Storage, import_date, supplier, account_receivable, account_payed, additional_fee, import_class, sub_storage, account_indicator, memos, operator。

这就是第二范式的例子，所有列都依赖于整个键，但是范式化设计过程还远未结束。