

SQL Server 200x Indexing Best Practices

Kimberly L. Tripp
SQLskills.com

Email: Kimberly@SQLskills.com

Blog: <http://www.SQLskills.com/Blogs/Kimberly>
<http://www.SQLskills.com>



Speaker – Kimberly L. Tripp

- Independent Consultant/Trainer/Speaker/Writer
- Founder, **SY Solutions, Inc.** www.SQLskills.com
 - Email: Kimberly@SQLskills.com
 - *Become a subscriber on SQLskills.com and learn about new resources which can improve your productivity and server performance!*
- Microsoft Regional Director <http://msdn.microsoft.com/isv/rd/>
- SQL Server MVP <http://mvp.support.microsoft.com/>
- Author for some SQL Server 2005 Whitepapers on MSDN (links from home page on www.SQLskills.com)
- SQL Server 2005 Launch Content Manager – Data Platform Track Sessions, Demos and Cross-training
- Writer/Editor for SQL Magazine www.sqlmag.com

SQL
skills.com

immerse yourself in sql server

@Copyright 2005, Kimberly L. Tripp
Immersion Events – Intense, Focused, Real-world Training!

Overview

- Index Usage and Effective Queries
- Index Concepts
- Table Structures: Internals
- Clustered Tables: The Pros and Cons
- Non-clustered Indexes: The Pros and Cons
- Finding the Right Balance
- Adding Non-Clustered Indexes for Better Performance
- Keeping Performance Optimal

Index Usage Conceptual

- Allows faster access to data, improving:
 - Lookup/query time
 - Insert/Update time – record location defined
 - Delete time – record location also defined
- Con: overhead for modifications
- Index Strategy Concepts
 - Fewer indexes are better than lots of indexes
 - Wider indexes have more uses
 - Can be used for point queries
 - Can be used by NARROW low selectivity queries
- Write Effective Queries to better take advantage of indexes...

Accessing Data: Limit Data

- Subset of columns = Projection
 - Do not use * (unless against view)
 - Optimizer has more chances for optimizing query when result set is NARROW (only the required columns)
- Subset of rows = Selection
 - Use positive search arguments
 - Isolate the column to one side of the expression
 - USE: `MonthlySalary > value/12` (constant, seekable)
 - DO NOT USE: `MonthlySalary * 12 > value` (must scan)
 - Be cautious with LEADING wildcards
 - USE: `LastName LIKE 'S%'`
 - Try not to just append `%val%` to every application
- Consider using Views, Stored Procedures and Functions to limit the columns/rows

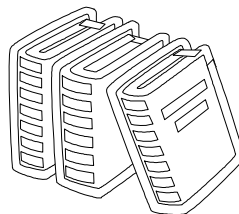
Index Concepts Book analogy

- Think of a book—with indexes in the back
- The book has one form of logical ordering
- For references, you use the indexes in the back...to find the data in which you are interested you look up the key
- When you find the key, you must lookup the data based on its location... i.e., a “bookmark” lookup
- The bookmark always depends on the (book) content order

Index – Species
Common Name

Index – Animal by Type,
Name *Bird, Mammal, Reptile,*
etc...

Index – Animal by
Country, Name



Index – Animals by Habitat,
Name *Air, Land, Water*

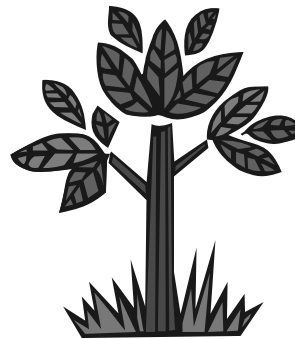
Index – Species Scientific Name

Index – Animal by
Continent, Country, Name

Index Concepts

Tree analogy

- If a tree were data and you were looking for leaves with a certain property, you would have two options to find that data:
- 1) Touch every leaf – interrogating each one to determine if they held that property...SCAN
- 2) If those leaves (which had that property) were grouped such that you could start at the root, move to the branch and then directly to those leaves...SEEK



Index Usage

Physical

- “Data” is stored in Index
 - Leaf level of an index (regardless of index type) stores relevant data for every row of the table
 - Clustered Index Leaf Level = Data
 - Non-clustered Index Leaf Level
 - Without additional “include” columns
 - » Stores the NC Index KEY + the CL Key (or RID of Heap)
 - With additional “include” columns
 - » Stores the NC Index KEY + additionally included columns + the CL Key (or RID of Heap)
 - Non-leaf level(s) are for navigation
- Indexes can be used to answer a query OR to help “point” to full data row to answer query

INCLUDE Non-Key Columns New in SQL Server 2005

- Leaf level of index can include non-key columns
- Index key limited to 900 bytes/16 columns – this is to keep tree structure and non-leaf levels optimal/small
- Allows more covering indexes
- Allows LOB data types (use sparingly)
- Also, consider Indexed Views for more interesting/wider index options

Table Structure

- HEAP: A table without a clustered index (collection of unordered pages)
- Clustered Table: A table with a clustered index (defined and ordered sets)
- Non-clustered Indexes *do not* affect the base table's structure
- However, Non-clustered Indexes are affected by whether or not the table is Clustered...

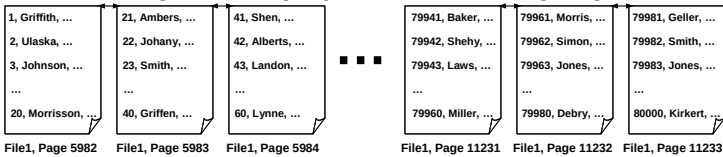
Hint: The non-clustered index dependency on the clustered index should impact your choice for the clustering key!

Table Structure

Clustered table

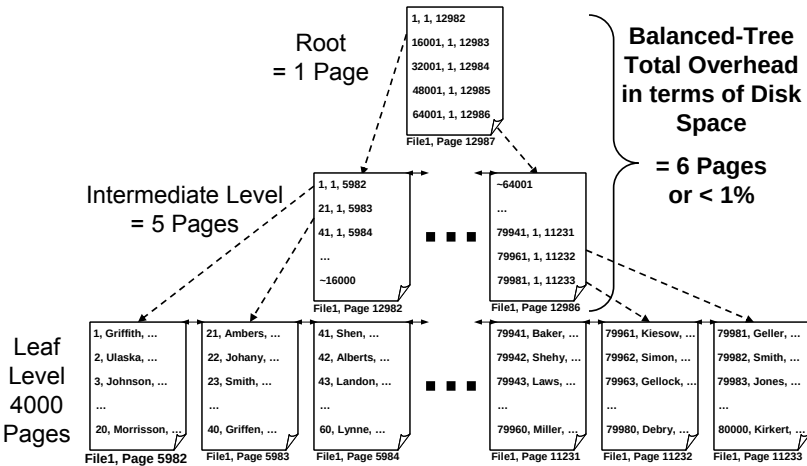
- Clustered Index defines order – applied at CREATION
- Expensive—in time and space—to build (an additional 1-2 times the index size for (re)build)
- Table is a doubly-linked list – order maintained LOGICALLY
- *Might* be expensive to maintain

4000 Pages of Employees in Clustering Key Order



Clustered EmployeeID

Consider a clustered index on an identity column (even better if this is the table's Primary Key)

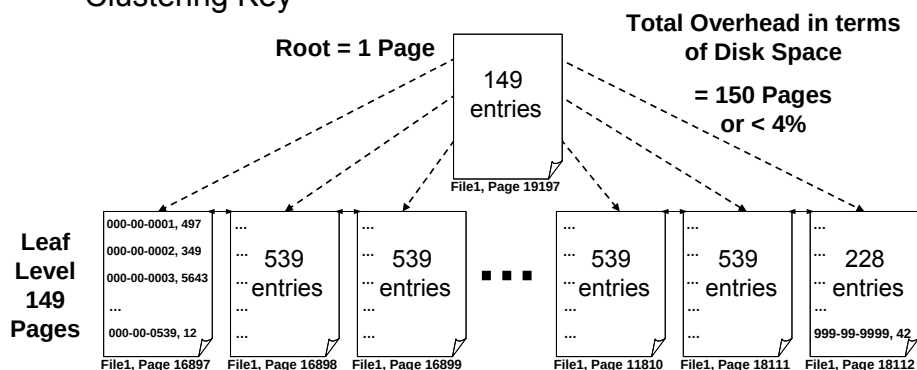


Non-Clustered Indexes Physical

- Depend on whether the table is a Heap or Clustered
- Clustered Table
 - Rows use the Clustering Key
 - No additional OH to add this column – may use actual data to define the clustering key
 - Minimizes NC Index manipulation if rows move (the NC still points to the CL Key)
 - Clustering key should be static and narrow
- Heap
 - Generally, not recommended

Non-Clustered Index Unique Key SSN

- Leaf level contains the non-clustered key(s) to define the order
- Also includes either the Heap's Fixed RID or the Table's Clustering Key



Clustered Index Criteria

- Unique
 - Yes – No overhead, data takes care of this criteria
 - NO – SQL Server must “uniquify” the rows on INSERT. This costs time and space. In SQL Server 2000, all uniquifiers are regenerated when the CL index is rebuild (fixed in 2005).
- Narrow
 - Yes – Keeps the NC indexes narrow
 - NO – Possibly wastes space
- Static
 - Yes – Improves Performance
 - NO – Costly to maintain during updates to the key especially if row movement and/or splits

In fact, an identity column that's ever increasing is ideal...

Clustering on an Identity

- **Naturally Unique**
(although not guaranteed – should combine with key constraint)
- **Naturally Static**
(although should be enforced through permissions and/or trigger)
- **Naturally Narrow**
(only numeric values possible, whole numbers with scale = 0)
- **Naturally creates a hot spot...**
 - Needed pages for INSERT already in cache
 - Minimizes cache requirements
 - Helps reduce fragmentation due to INSERTs
 - Helps improve availability by naturally needing less defrag

Clustering Key Choices

- Identity column
 - Order Date, identity
 - Not date alone as not unique
 - GUID
 - ✗ Populated by client-side call to .NET client GUID generator – NOT as the CL key but may be the Primary Key
 - ✗ Populated by server-side newid() function
 - ✓ Populated by NEW server-side newsequentialid() function (Creates an ever-increasing GUID)
Or on 2000 – populated with Gert's xp_guid
- Key points: Narrow, static, unique, ever-increasing*

Optimal Table Structures (1 of 3)

Poor Clustered Index Choice

- Worst base table structure is a poor choice for Clustered Index, for example: LastName
 - Fragmentation ⇒ poor performance
 - Non-unique ⇒ poor performance
 - Volatile ⇒ poor performance
 - Most duplicated value
 - Could cause record relocation and therefore fragmentation
 - Wide ⇒ poor performance
 - Wastes space/time – wide keys take longer to maintain/insert.
 - The wider the key the wider the non-clustered indexes – often with little benefit (although debatable as the non-clustered indexes will cover more queries)

TIP: Don't do this! ☺

Optimal Table Structures (2 of 3)

Heap is Better

- Better base table structure is a HEAP
 - No Fragmentation ⇨ forwarding rows
 - Wastes space in non-clustered indexes by using fixed-RID
 - Scans (due to locking/consistency requirements) can be more expensive!
 - Optimized for space over speed
 - Inserts are slower re: IAM and PFS lookup (reduces the “Swiss cheese problem”)
 - Updates can be slower with wide modification – same reason that inserts are slower... record relocation

TIP: Create HEAPs for staging tables as an interim table for high performance data loading!

See presentation “High Performance Data Loading” on www.SQLDev.Net for more tips!

Optimal Table Structures (3 of 3)

The RIGHT Clustered Index

- Best base table structure is the RIGHT Clustered Index
 - Cluster on identity column or add one for clustering
 - Cluster on composite key (date, identity) for tables that also have this pattern
 - Key Criteria: Static, Narrow and Unique
 - Secondary Criteria: Effective Hot Spot(s) using Identity or composite key to create multiple (but not random) hot spots – consider partitioning for large tables/high volume OLTP
- ***TIP: If this doesn’t exist consider adding a surrogate column solely to cluster on it!***
- ***TIP: Find the RIGHT Clustered index for YOUR environment***

Clustering for Performance

- Choose Clustering Key for internal benefits which lead to:
 - Better insert/update performance re: less splits
 - Better Availability as full table maintenance may not be needed
- Reliance on Nonclustered indexes is greater:
 - Faster access to narrow/low selectivity range queries
 - NC Indexes are used in numerous non-obvious ways (multiple NC Indexes can be joined to cover a query)
 - More flexible in the definition (i.e. Indexed Views can include computations, substrings, etc. and with INCLUDE...)
 - Non-clustered indexes are easier/faster to rebuild when they become fragmented (only Shared Table Lock)
 - Non-clustered indexes are easier to keep less fragmented – i.e. rebuilds with fillfactor has greater benefit since the index row is narrower

Finding the Right Balance

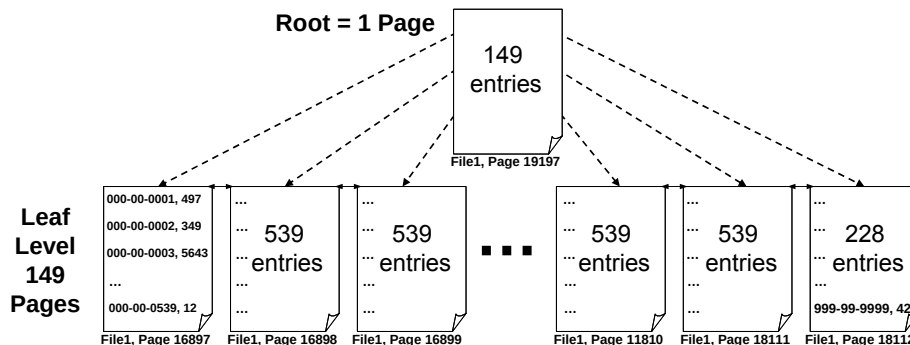
- Start with a minimal number of indexes
 - Clustered Index
 - Primary Key (nonclustered, if not clustered)
 - Unique Keys
- Manually index foreign keys
 - Non-unique indexes
 - Speed up join performance
- Use Database Tuning Advisor
- Manually index based on either:
 - Specific query tuning where DTA didn't help
 - Query frequency

Non-Clustered Indexes

- Internally, non-clustered indexes always store an entry for EVERY row of the table – effectively, they are a “mini” clustered table (of only the index columns)
- If the columns requested in the query are IN the index (regardless of order) then the index **COVERs** the query
- If queries are narrow this is more likely to happen
- Real-world queries...for example, accessing a large number of tables in a join are typically narrow (the join condition, maybe a search argument and possibly one or two columns in the select list)
- Is it likely that you would ever execute:

```
SELECT * FROM 8-table-join
```

What kind of structure is this?



- A non-clustered index on SSN on a table that has a clustered index on EmpID (as we've seen), *or*
- A clustered index on SSN on a table that has ONLY two columns, SSN and EmpID

NOTE: The structure of a non-clustered index is EXACTLY the same as a clustered index – but only on the columns that make up the NC index (with the RID or CL Key added to “point to the data row”).

Seems Like...

- Non-clustered indexes are only useful for relatively selective queries...
- What about the following queries?

```
SELECT EmpID, SSN
FROM Employee
WHERE SSN between x and y
```

or

```
SELECT EmpID, SSN
FROM Employee
WHERE EmpID < 10000
```

Think about the access patterns:

- Table Scan (always an option)
- Index Seek w/partial scan
The nc index on SSN
"covers" the query

Think about the access patterns:

- SARG is on clustering key

The nc index on SSN
"covers" the query

Non-Clustered Indexes For optimal query performance

- Non-Clustered Indexes are BETTER for low selectivity/range queries because:
 - Non-clustered indexes ARE EXACTLY the same as a clustered index—but narrower
 - If your query only needs columns IN an index then the index is said to "cover" your query
 - No – this does *not* mean that you need to cover every query!!!
 - What it does mean: You can make one or two existing indexes slightly wider and you probably will not hurt INSERT, UPDATE, and DELETE performance and may RADICALLY improve SELECT performance
- Non-Clustered Indexes are narrower and easier to maintain (let ITW/DTA help...)
- Use INCLUDE to create narrower trees, more relevant leaf levels for covering!

Covering an Aggregate

- Increasing Gains with different types of indexes can always be achieved
- Each query can be isolated and tuned as if it were in a sandbox – *not* the goal
- Finding the right balance starts with:
 - Base Indexes (CL, PK, UK, FKs)
 - Capturing a workload, tuning with DTA
 - Prioritizing the queries that are still performing poorly and then choose the optimal level of gain by understanding the trade-offs

Indexing for Aggregations

- Two types of Aggregates:
Stream and Hash
- Try to Achieve Stream to Minimize Overhead in temp table creation
- Computation of the Aggregate Still Required
- Lots of Users, Contention and/or Minimal Cache can Aggravate the problem!

Aggregate Query

- Member has 10,000 Rows
- Charge has 1,600,000 Rows

```
SELECT c.member_no AS MemberNo,  
       sum(c.charge_amt) AS TotalSales  
FROM dbo.charge AS c  
GROUP BY c.member_no
```

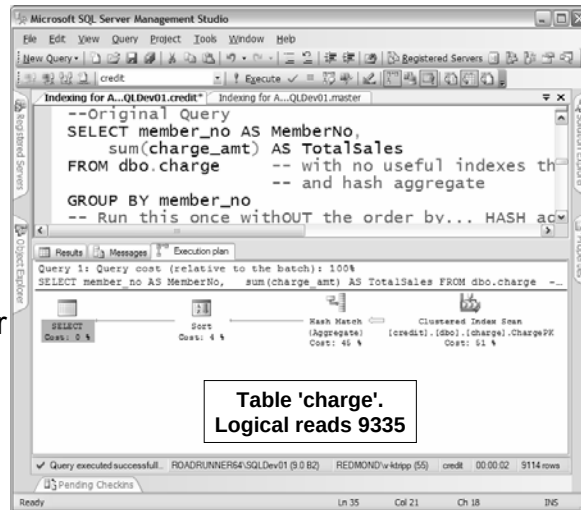
Aggregate Query (cont'd) Table scan + hash aggregate

```
SELECT c.member_no AS MemberNo,  
       sum(c.charge_amt) AS TotalSales  
FROM dbo.charge AS c  
GROUP BY c.member_no
```

- Table Scan of Charge Table
 - Largest structure to evaluate
 - Worst case scenario
- Worktable created to store intermediate aggregated results – OUT OF ORDER (HASH)
- Data Returned OUT OF ORDER – unless ORDER BY added
- Additional ORDER BY causes another step for SORT – sorting can be expensive!

Worst Case

- Clustered Index Scan
(table scan)
1,600,000 rows
- Hash Aggregate
yields 9,114
rows out of order
- Sort
only has to sort
9,114 rows
instead of
1,600,000 rows
- Return Data



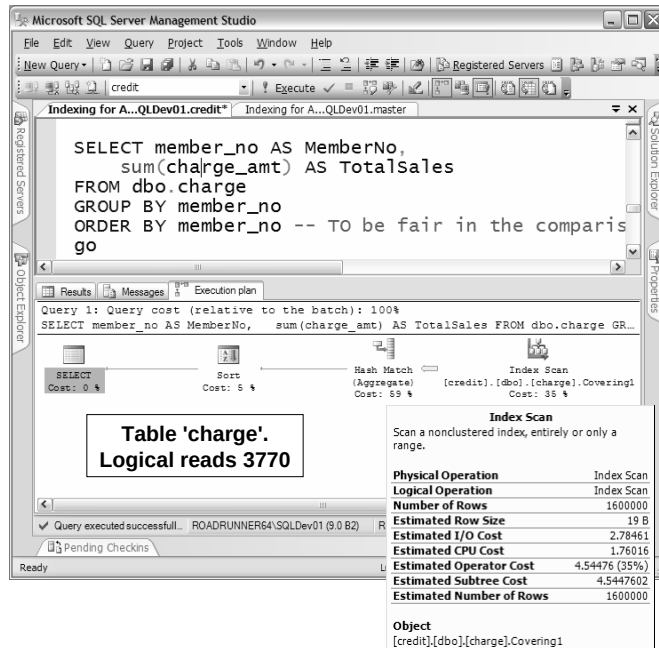
Aggregate Query Index scan + hash aggregate

```
SELECT c.member_no AS MemberNo,
       sum(c.charge_amt) AS TotalSales
FROM   dbo.charge AS c
GROUP BY c.member_no
```

- Out of Order Covering Index on Charge Table
 - Index Exists which is narrower than base table
 - Used instead of table – to cover the query
- Worktable still created to store intermediate aggregated results – OUT OF ORDER (HASH)
- Data Returned OUT OF ORDER – unless ORDER BY added
- Additional ORDER BY causes another step for SORT – sorting can be expensive!

Not as Bad

- COVERING Index Scan
1,600,000
narrower rows
- Hash Aggregate
yields 9,114
rows out of
order
- Sort
only has to sort
9,114 rows
instead of
1,600,000 rows
- Return Data



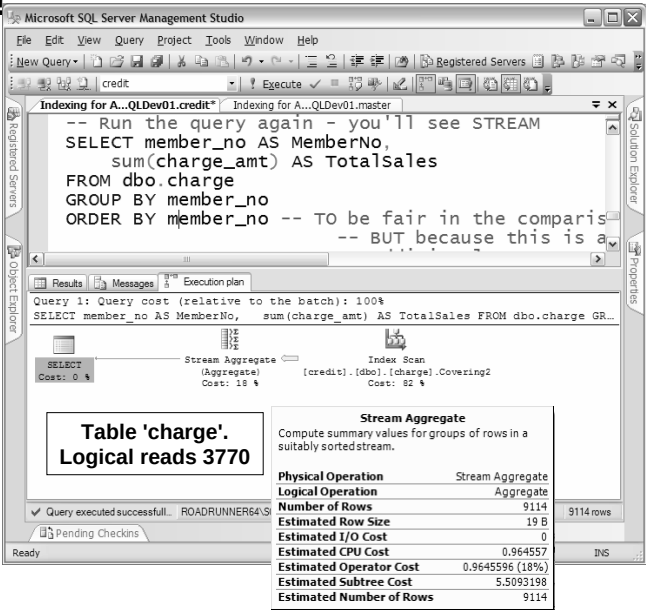
Aggregate Query Index scan + stream aggregate

```
SELECT c.member_no AS MemberNo,
       sum(c.charge_amt) AS TotalSales
FROM   dbo.charge AS c
GROUP BY c.member_no
```

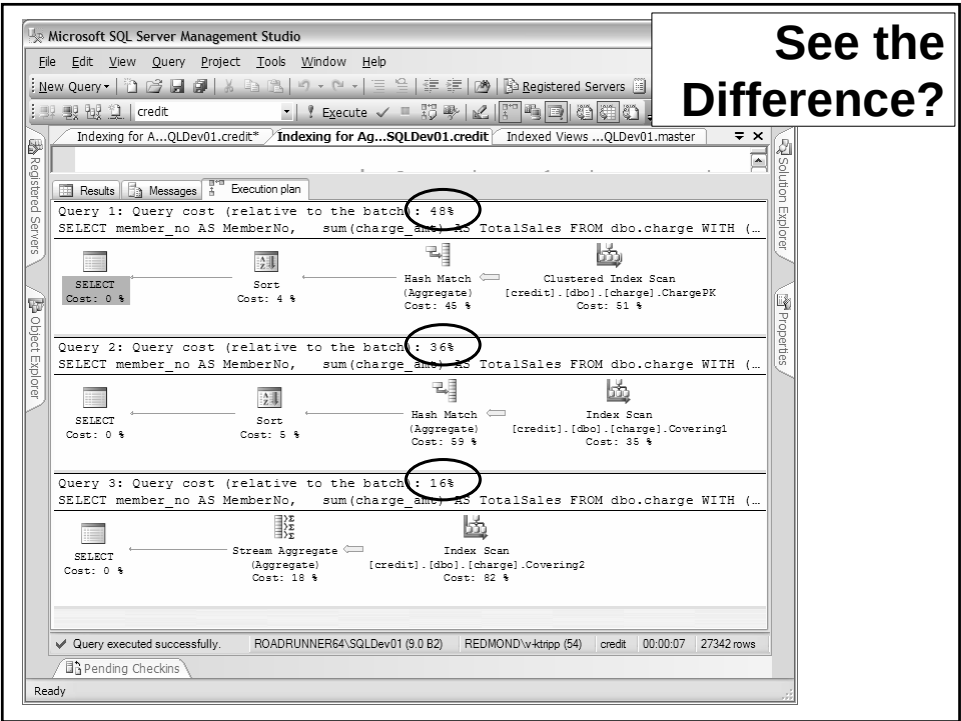
- Covering Index on Charge Table – In ORDER of GROUP BY Clause
 - Index Exists which is narrower than base table
 - Used instead of table – to cover the query
 - Covers the GROUP BY so data is grouped
- Less work to aggregate results IN ORDER
- Data Returned IN ORDER – unless ORDER BY or other joins added
- Adding an ORDER BY identical to the GROUP BY does NOT cause any additional step for sorting!

Much Better!

- COVERING Index Scan 1,600,000 narrower rows
- Stream Aggregate also yields 9,114 rows IN ORDER NO SORT REQUIRED
- Return Data



See the Difference?



Concerns

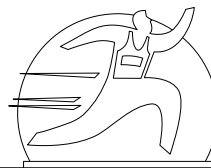
- More temp tables
- More contention in tempdb
- Larger tempdb required
- Performance Varies on each execution
- Aggregate needs to be computed

Is there a better way?
Indexed Views!

demo

The Punch Line?

What kinds of gains can you get?
Will it be worth it?



Aggregate Query Indexed view

Table 'SumOfAllCharges'. Logical reads 35

See the Difference?

Query with NO useful indexes

Query with covering Index in wrong order

Query with covering index in correct order

Query w/Indexed View and no computations!

Indexed View for Aggregates

Key points

- TempDB access not necessary
- NO worktables are necessary
- Aggregated set should be small but not too small as to create a hot ROW spot of activity (which can create excessive blocking)
 - GROUP BY member_no – probably OK
 - GROUP BY state – too few rows in aggregate
 - GROUP BY country – *avoid* like the plague!
- Performance of data modification statements should be tested

INCLUDE Non-key Columns Compared to Indexed Views

- Indexed Views allow aggregates – adding interesting columns in the leaf level of an index offers “creative covering”
- With new INCLUDE clause, leaf level of index can include non-key columns
- Index key [has been since 7.0] limited to 900 bytes/16 columns – this is to keep tree structure and non-leaf levels optimal/small
- Allows more covering indexes
- Indexed View v. Include – depends on what needs to be in the leaf level

Finding the Right Balance

Index strategies: Summary

- Determine Primary Usage of Table – OLTP vs. OLAP vs. Combo? This determines Clustered Index
- Create Constraints – Primary Key and Alternate/Candidate Keys
- Manually Add Indexes to Foreign Key Constraints
- Capture a Workload(s) and Run through Index Tuning Wizard
- Add additional indexes to help improve SARGs, Joins, Aggregations

Are you done?

NO!

Review

- Index Usage and Effective Queries
- Index Concepts
- Table Structures: Internals
- Clustered Tables: The Pros and Cons
- Non-clustered Indexes: The Pros and Cons
- Finding the Right Balance
- Adding Non-clustered Indexes for Better Performance
- Keeping performance optimal

MSDN Webcast Series

<http://www.microsoft.com/events/series/msdnsqlserver2005.msp>

- Session 1: Interaction between data and log
- Session 2: Recovery Models
- Session 3: Table optimization strategies
- Session 4: Optimization through indexes
- Session 5: Optimization through maintenance
- Session 6: Isolation, locking and blocking
- Session 7: Optimizing procedural code
- Session 8: Partitioning...
- Session 9: Profiling for the unknown problems
- Session 10: Common Roadblocks, *A Series Wrapup*

Resources

Whitepapers written by Kimberly L. Tripp

- “SQL Server 2005 Snapshot Isolation”
On MSDN and www.SQLskills.com
- “SQL Server 2005 Partitioned Tables”
On MSDN and www.SQLskills.com
- Blogged this new one with a preview on SQLskills: The Database Administrator's Guide to the SQL Server Database Engine .NET Common Language Runtime Environment
www.sqlskills.com/blogs/kimberly/PermaLink.aspx?guid=385c0e05-497f-4c4c-b24a-155106a08959
- Another one in the works: SQL Server 2005 Management Tools

Finding the Right Balance (cont'd)

SQL Server 2000 resources with GREAT best practices – that STILL apply!

- Support Webcast: *Indexing for Performance: Finding the Right Balance*
(recorded 11 June 2004)
 - <http://msevents.microsoft.com/CUI/EventDetail.aspx?EventID=1032254503&Culture=en-US>
- Support Webcast: *Indexing for Performance: Proper Index Maintenance*
(recorded 19 July 2004)
 - <http://msevents.microsoft.com/CUI/EventDetail.aspx?EventID=1032256511&Culture=en-US>

Finding the Right Balance (cont'd)

SQL Server 2000 resources with GREAT best practices – that STILL apply!

- Support Webcast: *SQL Server 2000 Profiler: What's New and How to Effectively Use It*
<http://support.microsoft.com/default.aspx?scid=%2Fservicedesks%2Fwebcasts%2Fwc111400%2Fwcblurb111400%2Easp>
- Whitepaper: “Index Tuning Wizard for Microsoft SQL Server 2000”
 - <http://msdn.microsoft.com/library/en-us/dnsq12k/html/itwforsql.asp?frame=true>



Thank you!

Please take a moment to fill out your evaluation.

Kimberly L. Tripp

Consultant . Trainer . Writer . Speaker

email: Kimberly@SQLskills.com

**Make sure to register for special offers
and other helpful information and resources!**

www.SQLskills.com



immerse yourself in sql server

@Copyright 2005, Kimberly L. Tripp

Immersion Events – Intense, Focused, Real-world Training!