

SQLskills Immersion Event

IEPTO2: Performance Tuning and Optimization

Whiteboard Drawings

Kimberly L. Tripp
Kimberly@SQLskills.com



**NOTE: Includes whiteboard drawings from prior IE courses
+ drawings done the week of September 19, 2016 in Bellevue, WA – IEPTO2 Training**

These next 5 slides come from this Pluralsight course.

Module 4: Statement Caching

[Here's a link to the course.](#)

SQL Server: Optimizing Ad Hoc Statement Performance

Module 4: Statement Caching

Kimberly L. Tripp

Kimberly@SQLskills.com

<http://www.SQLskills.com/blogs/Kimberly>



pluralsight 
hardcore developer training

Stored Procedure Plans and Caching

- A plan is generated when no plan already exists in cache
- Plans are never saved on disk and may not persist within the cache
 - Some operations completely remove / evict the plan from the cache
 - **Server-level:** Server restart | DBCC FREEPROCCACHE | some RECONFIGURE changes
 - See recordings starting with **Side Effect: Plan Cache Flush** for version-specific details and demo
 - **Database-level:** DBCC FLUSHPROCINDB (undocumented) | sp_dbcmptlevel
 - **Procedure-level:** sp_recompile or they're aged out through non-use
 - Some operations cause the plan to be invalidated (but it still remains an object in the cache); these can be tracked
 - Schema of base object changes (ALTER TABLE / ALTER <object>)
 - Including when indexes are added to the base object
 - Statistics of base objects change
 - See recordings starting with **Plan Invalidation** for version-specific details and demos
- When a plan exists in cache, all subsequent executions use that plan

Side Effect: Plan Cache Flush (1)

- **In SQL Server 2005 (only), numerous operations cause the entire plan cache to be cleared / flushed (all plans evicted from cache)**
 - Database-level operations that cause an entire plan cache flush include:
 - When a database auto-closes (from AUTO_CLOSE option)
 - When a database is detached
 - When a database snapshot is dropped
 - When a database state is changed (OFFLINE, ONLINE, READ_ONLY, READ_WRITE)
 - When a database backup is restored
 - ...
 - **There are quite a few others;** see KB article 917828 for complete list
- **In all versions, numerous server-level operations that cause the entire plan cache to be cleared include:**
 - max server memory (MB)
 - min server memory (MB)
 - **... There are quite a few others;** more information coming up

Side Effect: Plan Cache Flush (2)

- **From SQL Server 2005 SP2 onward, messages written to the error log show you that the cache was flushed:**
 - SQL Server has encountered 4 occurrence(s) of cachestore flush for the 'Object Plans' cachestore (part of plan cache) due to some database maintenance or reconfigure operations.
 - SQL Server has encountered 4 occurrence(s) of cachestore flush for the 'SQL Plans' cachestore (part of plan cache) due to some database maintenance or reconfigure operations.
 - SQL Server has encountered 4 occurrence(s) of cachestore flush for the 'Bound Trees' cachestore (part of plan cache) due to some database maintenance or reconfigure operations.

Side Effect: Plan Cache Flush (3)

- From SQL Server 2008 onward, many database-level operations that caused the entire plan cache to be flushed no longer do so
 - Instead, they only cause the plans for that database to be flushed
- Database-level operations that cause the database plan cache to be flushed include:
 - DBCC FLUSHPROCINDB (dbid)
 - When a database is detached / dropped / restored
 - When a database (or filegroup with the database) changes between READ_ONLY and READ_WRITE
 - When a database changes between OFFLINE and ONLINE, or auto closes
 - There are others... (for example, changing recovery model ...)

NOTE: There are many places where a few of these are documented incorrectly. Some aren't detailed at all and others say that many of these operations **STILL** cause the entire plan cache to be cleared.

Side Effect: Plan Cache Flush (4)

- Many server-level configuration changes cause the entire plan cache to be cleared (at the time of RECONFIGURE, not sp_configure)
- Configuration options that could cause the entire cache to be cleared:

<i>Some Versions</i>	<i>Most Versions</i>	<i>SQL Server 2014</i>
cost threshold for parallelism	cross db ownership chaining	access check cache bucket count
max text repl size (B)	index create memory (KB)	access check cache quota
query governor cost limit	max degree of parallelism	clr enabled
remote query timeout (s)	min or max server memory (MB)	max worker threads
	min memory per query (KB)	
	query wait (s)	
	user options	

- Key point: be careful making database-level or server-level configuration changes during production hours
 - Test them in development first
 - sp_configure, then RECONFIGURE – check the cache messages
 - Database changes – check [sys].[dm_exec_procedure_cache]

SQLskills Immersion Event

IEPTO2: Performance Tuning and Optimization

Module 9: Statement Execution, Stored Procedures, and the Plan Cache

Kimberly L. Tripp

Kimberly@SQLskills.com



Sch_M on table

Sp_recompile
tname

NOLOCK longrunning

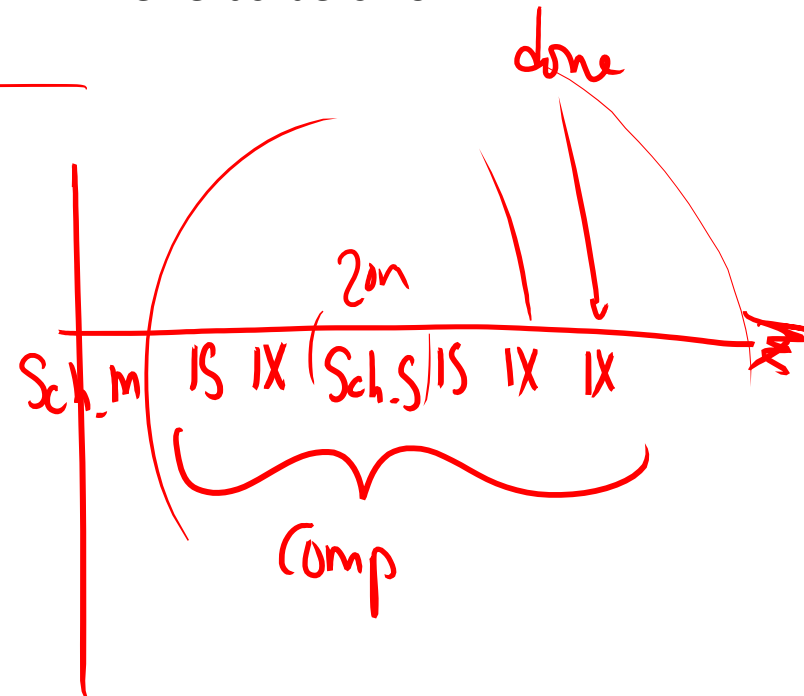
Sch_S

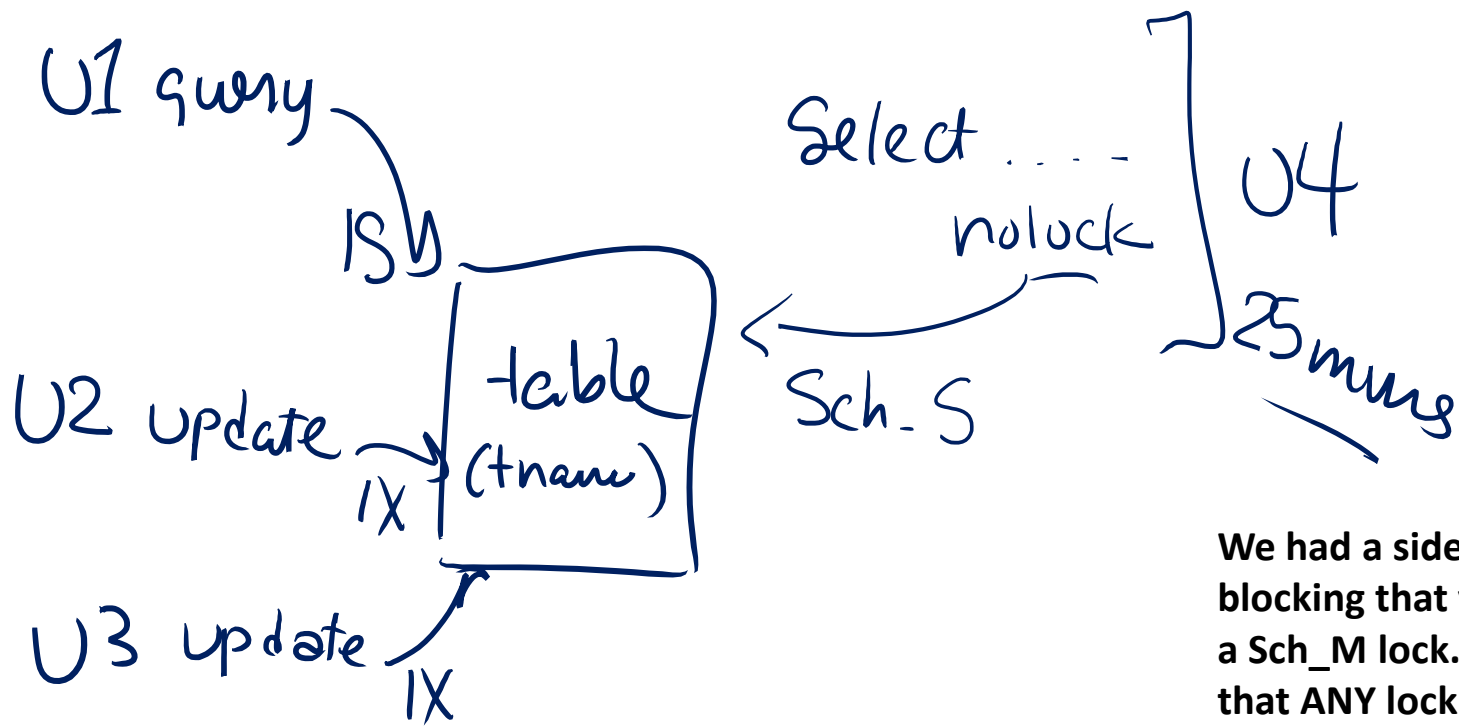
20 min



We had a side discussion on the blocking that you could have with a Sch_M lock. The key point is that ANY lock blocks a SCH_M lock but a SCH_S held from a NOLOCK query that's extremely long running – will cause nasty secondary blocking while the SCH_M waits for ALL locks to drain off (where the SCH_S lock is probably the longest running)

The next slide is from IE1.



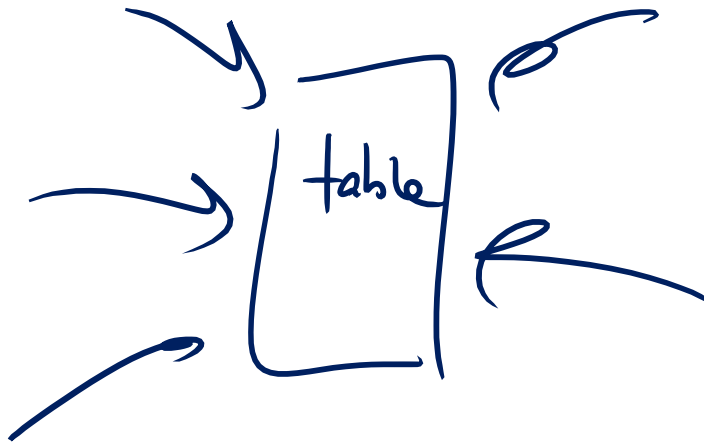


We had a side discussion on the blocking that you could have with a Sch_M lock. The key point is that ANY lock blocks a SCH_M lock but a SCH_S held from a NOLOCK query that's extremely long running – will cause nasty secondary blocking while the SCH_M waits for ALL locks to drain off (where the SCH_S lock is probably the longest running) The next slide talks about my script (PreventLongBlockingChains.sql). The following slide is a graphic from IEPTO1.

sp_recompile tname
WANTS Sch_M

Every other request WANTS

Be sure to check out the script:
PreventLongBlockingChains.sql



WANT sp_recompile frame

any long running statements?
(configurable)
30 sec

if not

then get in queue

if I can't get lock in 5 sec

Get out

Nasty Blocking CHAINS

Relaxed FIFO isn't Always Enough

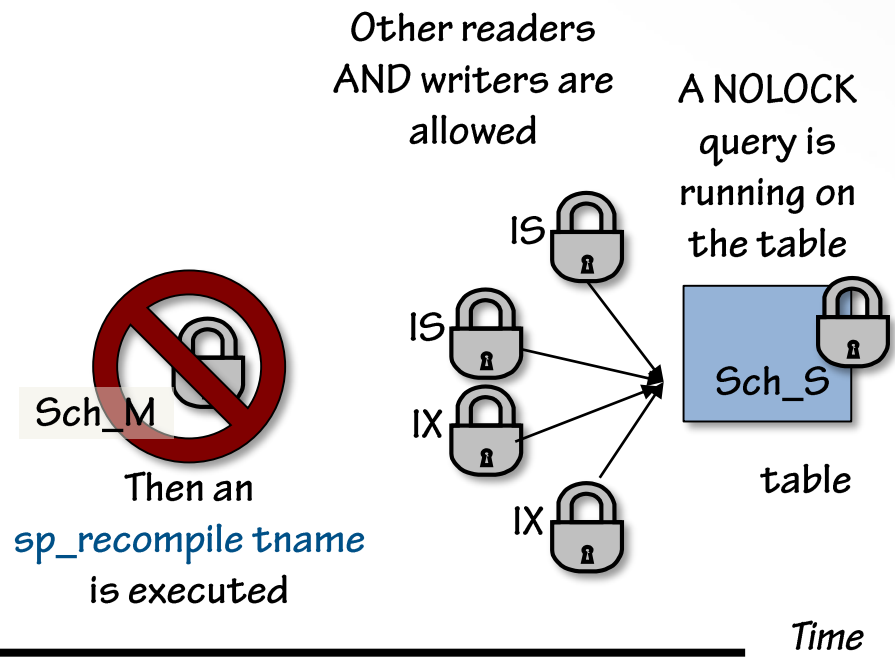
We had a side discussion on the blocking that you could have with a Sch_M lock. This slide is from IEPTO1.

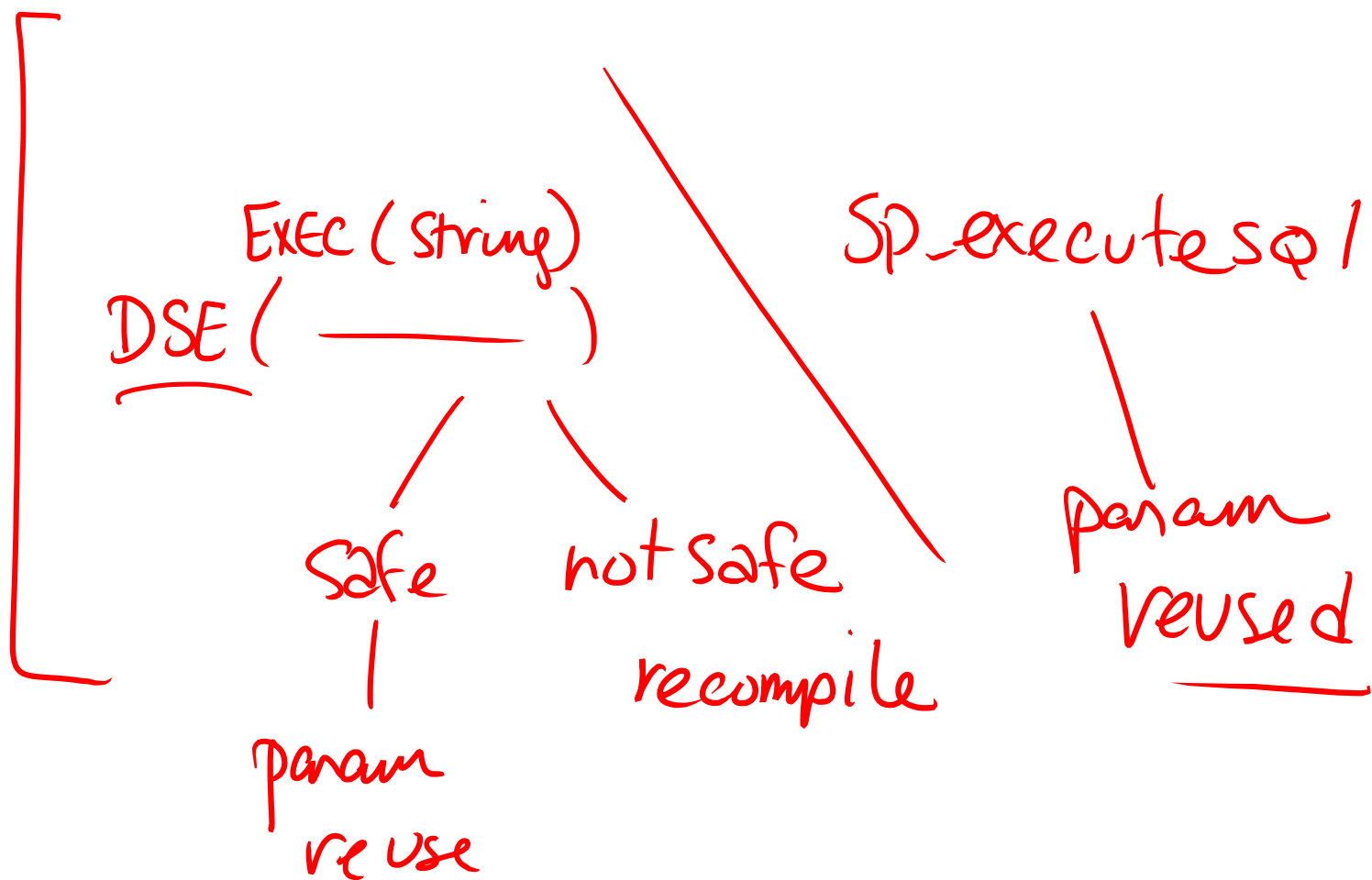
- Imagine someone's running a NOLOCK query against a very large table (the query takes 4 minutes to run [ok, not that nasty])
- This query's running off-hours but the standard is to use NOLOCK
- An **sp_recompile tname** is requested

Unfortunately now...
everybody waits.....

Relaxed FIFO only let's through the folks that are compatible with ALL pending locks (not just those that are currently held).

This can create terribly long blocking chains.





Using EXEC (string) [DSE] and/or sp_ExecuteSql [ES] inside of stored procedures

DSE and ES are not the same.

DSE is unknown until runtime. At that point the statement is treated as though it's adhoc. Because most statements are **not** going to be safe, the statement will end up being recompiled (and optimized) for each execution.

ES is forced parameterization. When a statement is stable you can use this to reduce compilation / CPU costs.

OSFA Dialog (An attempt at *One Size Fits All* Programming)

fast
reason
Oh no...
OMG

where (c1 = @p1
or @p1 is null)
AND (c2 = @p2
or @p2 is null)
AND

Cust ID
LName FName
SSN
City State

Multipurpose screens/dialog boxes that lead to multipurpose procedures

They looked/sounded good at the time?

But, they're often VERY hard to optimize with where clauses that look like:

```
WHERE (Column1 = @Value1 OR @Value1 IS NULL)  
AND (Column2 = @Value2 OR @Value2 IS NULL)
```

...

Check out the demo scripts in order to fully see how to better optimize this – using `sp_executesql` (for combinations that are stable) and `sp_executesql WITH RECOMPILE` for combinations that could generate different plans (e.g. unstable plans).

brisque

Plan

P—

plan

plan

Opt
ad
hoc

Stub

Stub

Stub

too
much
to
look for
Plan
Stub

limited

at 1^{or} 2
GB

Clear

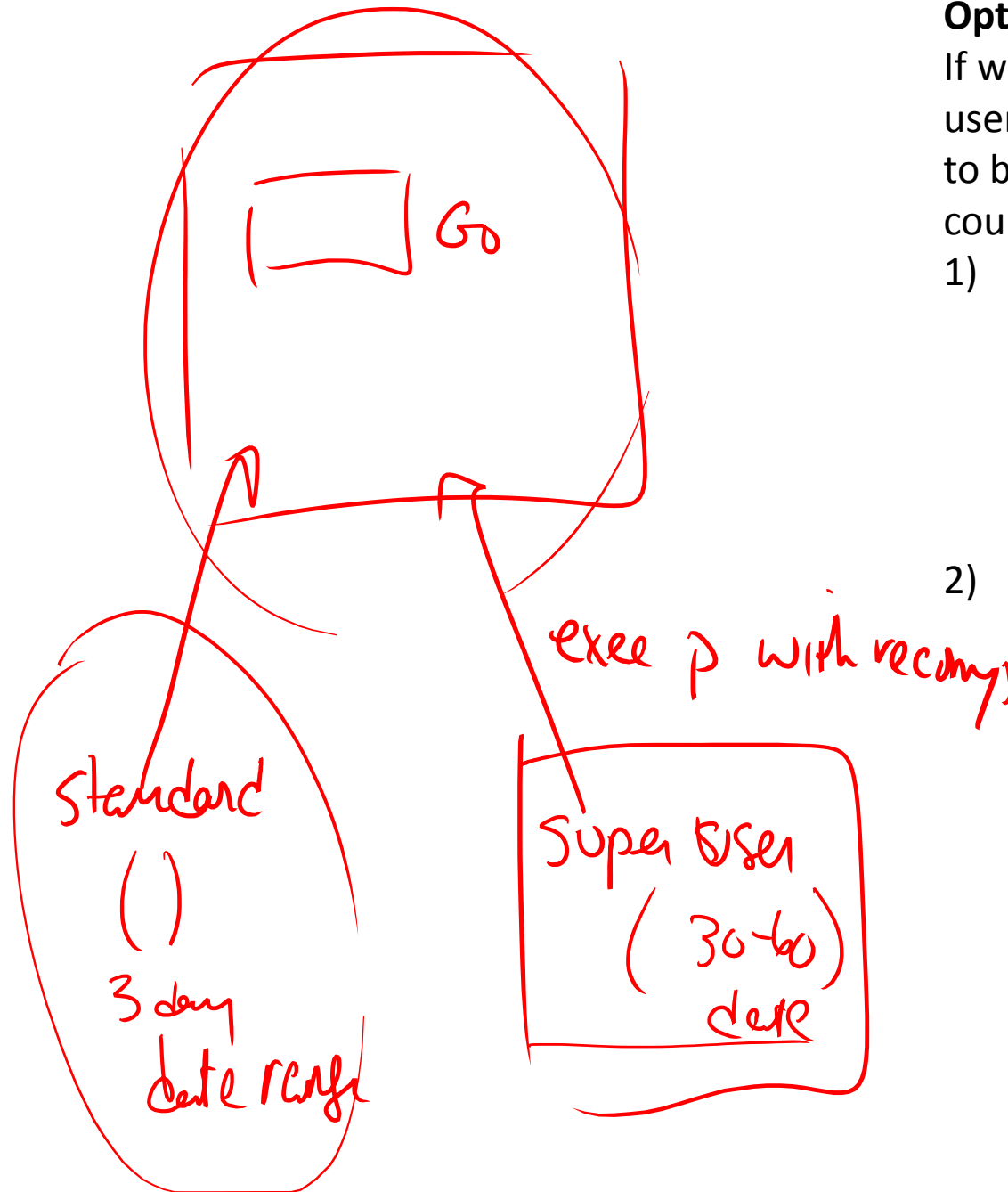
If you turn on “optimize for ad hoc workloads” then the number of “stubs” that can fit in cache is much higher.

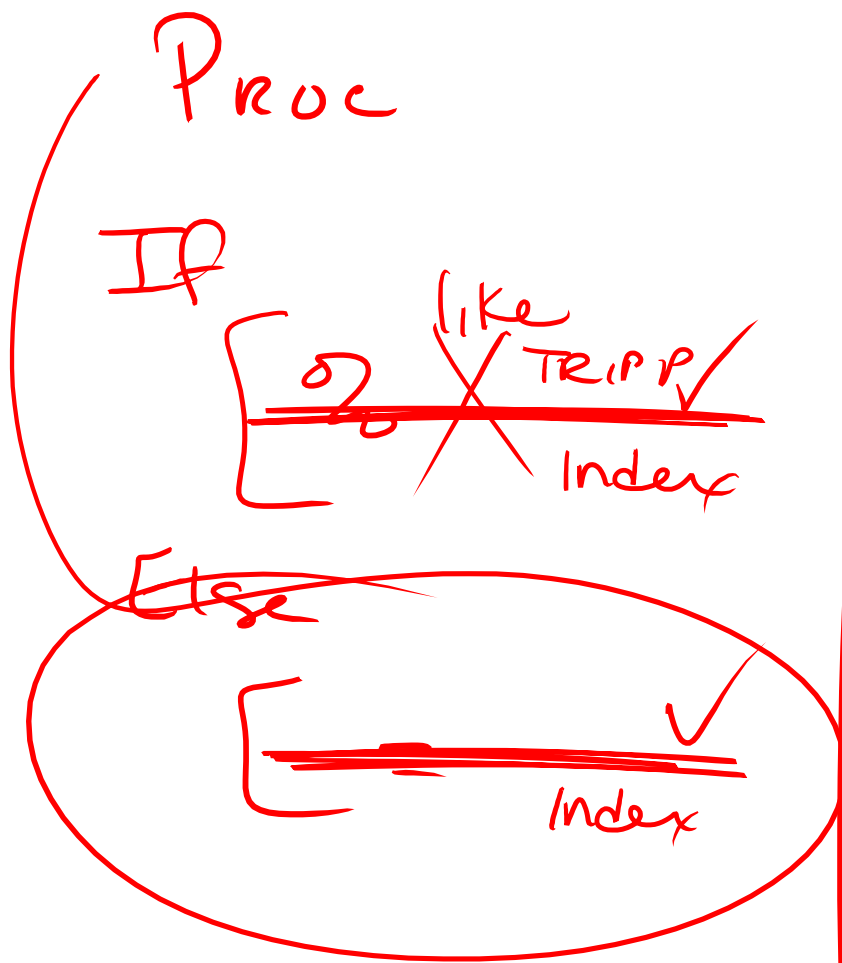
At that point, searching to find whether or not the statement has executed is harder... this is why you want to make sure that you don’t let this cache get really large (clearing it at 1 or 2GB).

Options for different executions

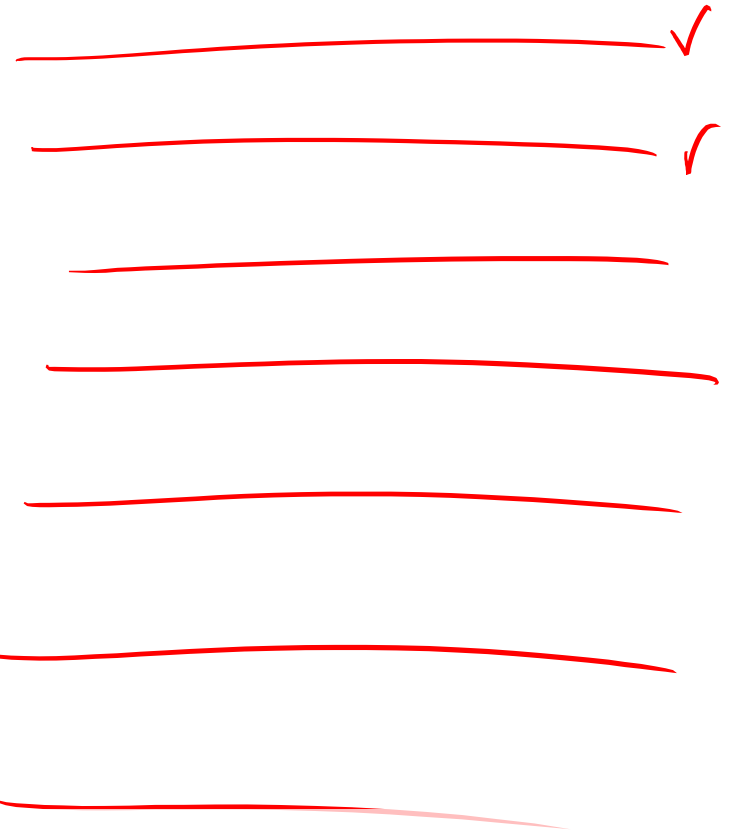
If we largely execute with “standard” users and we absolutely want THAT plan to be in cache and to be re-used then we could:

- 1) Use option (optimize FOR...) and use a literal that will generate a plan that's great for the “typical” scenario. However, the performance won't be great for the atypical (super user) scenario.
- 2) We could use EXEC for the typical scenario and EXEC proc WITH RECOMPILE for the super user scenario. The benefit of this is:
 - 1) We get a cached plan for the typical user (no CPU/plan stability)
 - 2) We get a specific plan for the atypical user (the super user) as opposed to a possibly inefficient plan (so, we have to compile but we're only compiling this one type of plan)





Proc



Conditional logic and modularization

This is another way of looking at the same thing from the prior diagram. Really, don't think of the "logic" in your procedure; think only of the procedure as a list of statements (regardless of conditional logic). SQL Server tries to optimize ANY statement that's "*optimizable*" with the parameters that were supplied (regardless of those specific parameters apply for *that* execution). This is another case where you can end up with an inefficient plan because of parameter sniffing.

IF

WILDCARD

Select... Like TS

INDEX

ELSE

EQUALITY

Select... = %e%

INDEX

1ST EXEC
%e%

1ST EXEC
Tripp

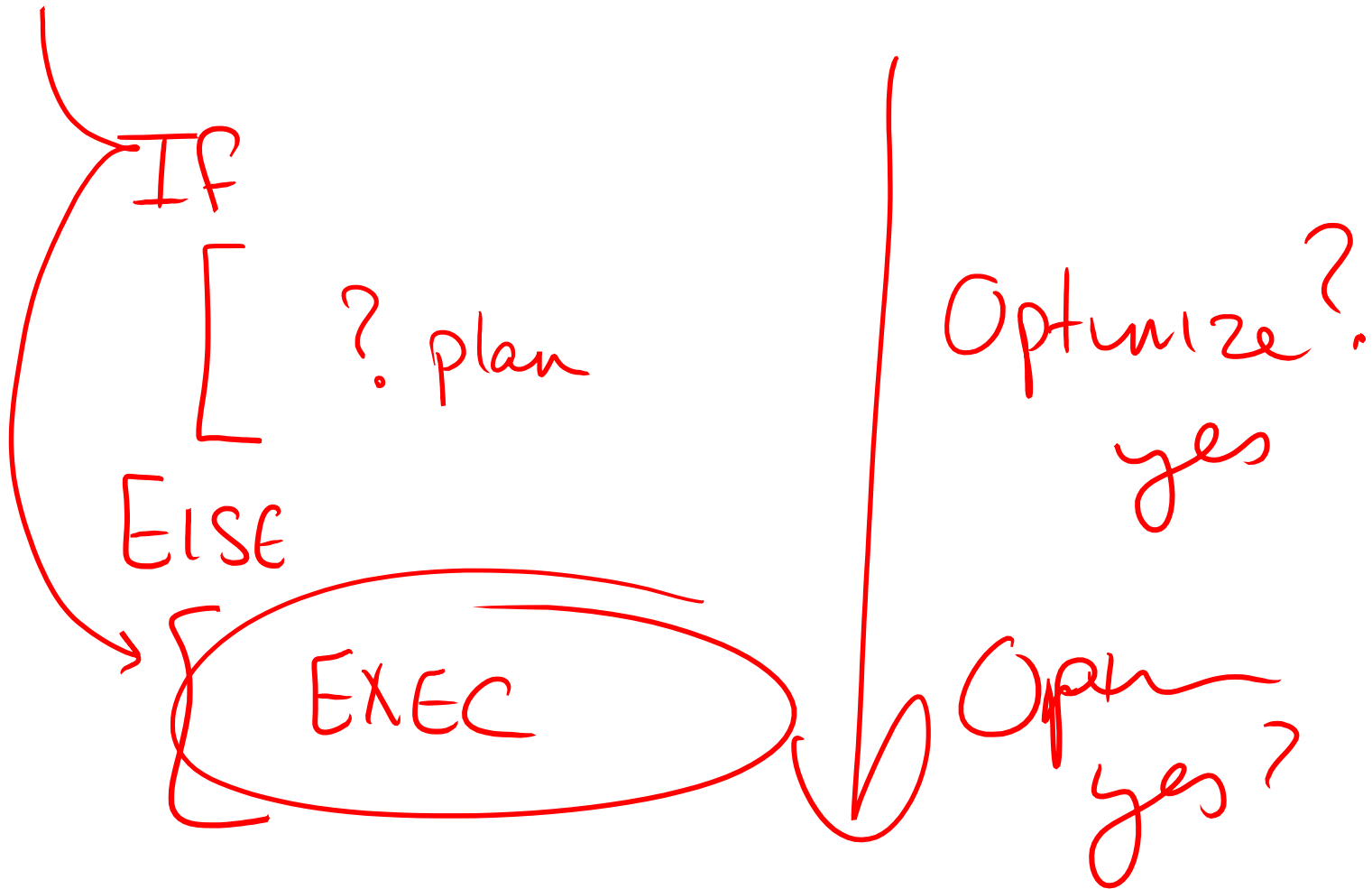
Conditional Logic example from the demo

When we ran our procedure to get members tied to the input of the name supplied – the conditional logic would execute only ONE branch... but, SQL Server would have already had to optimize BOTH branches

If the first execution was EXEC GetMembers '%e%' then the wildcard statement can be optimized using a table scan. The equality case will know that '%e%' (as a LITERAL name) is HIGHLY selective so there, they'll use an index...

However, if the first execution happens to be 'Tripp' then BOTH statements will be using an index and this will make the first statement inefficient.

The point, you can't control who gets there first and you never want statements to be optimized for parameters that weren't intended for them. **SOLUTION? Modularize your code!**



Conditional logic and modularization

SQL Server will never step into a sub-procedure unless it's executed and in this case it will only be executed with exactly the right parameters.

CR PROC — } Sniffed
 (param1
pa 2
 —)

as

Declare var — unknown why?
 we don't know
 until EXECUTION

Select .. — Param?
 Select ... — param?

Parameters are “sniffed”

SQL Server uses parameters to optimize. Sniffing implies that SQL Server uses the histogram to evaluate the data selectivity. For the FIRST execution, parameter sniffing is great. It's the subsequent executions that can suffer from parameter sniffing (and therefore end up with PSP = parameter sniffing problems).

Variables are “unknown”

Variables are only known at runtime as the statements execute and the variables are assigned. As a result, they are unknown during optimization. Without an actual value, SQL Server cannot use the histogram. Instead SQL Server uses the density vector for the “average” numbers of rows that meet that criteria.

PROC

param list

@p1	value
@p2	value

SQL → uses a variable unknown
D.V.

SQL → uses a literal Known → special features

SQL → uses a parameter — histogram sniffing
(no special features)

If

[SQL → uses param — hist. sniffing

Else

[SQL → uses param — hist sniffing

SQL Server optimizes everything that's optimizable (sp?) ☺

Again, parameters are known. Parameters can be sniffed. SQL Server will optimize every statement that it can — regardless of what actually ends up executing. At the time of optimization they do not know what will or will not execute and they do not know the value of a variable.

DEFAULT
Proc

DSE — protects against SQLinj
execute under context
CALLER

Dynamic String Execution and SQL Injection

To prevent SQL Injection, SQL Server requires that strings execute under the context of the caller.

Dynamic String Execution and Execute As

I have a love/hate relationship with EXECUTE AS. On the one hand, you can give someone rights to something that you wouldn't otherwise be able to give...

However, be careful giving elevated privs around DSE. Generally, my recommendation around DSE is to only create a login-less user with very low privs.

fname

sproc

EXEC(~~~~~)
Delete ~ ~ fname

Executes As?

CALLER

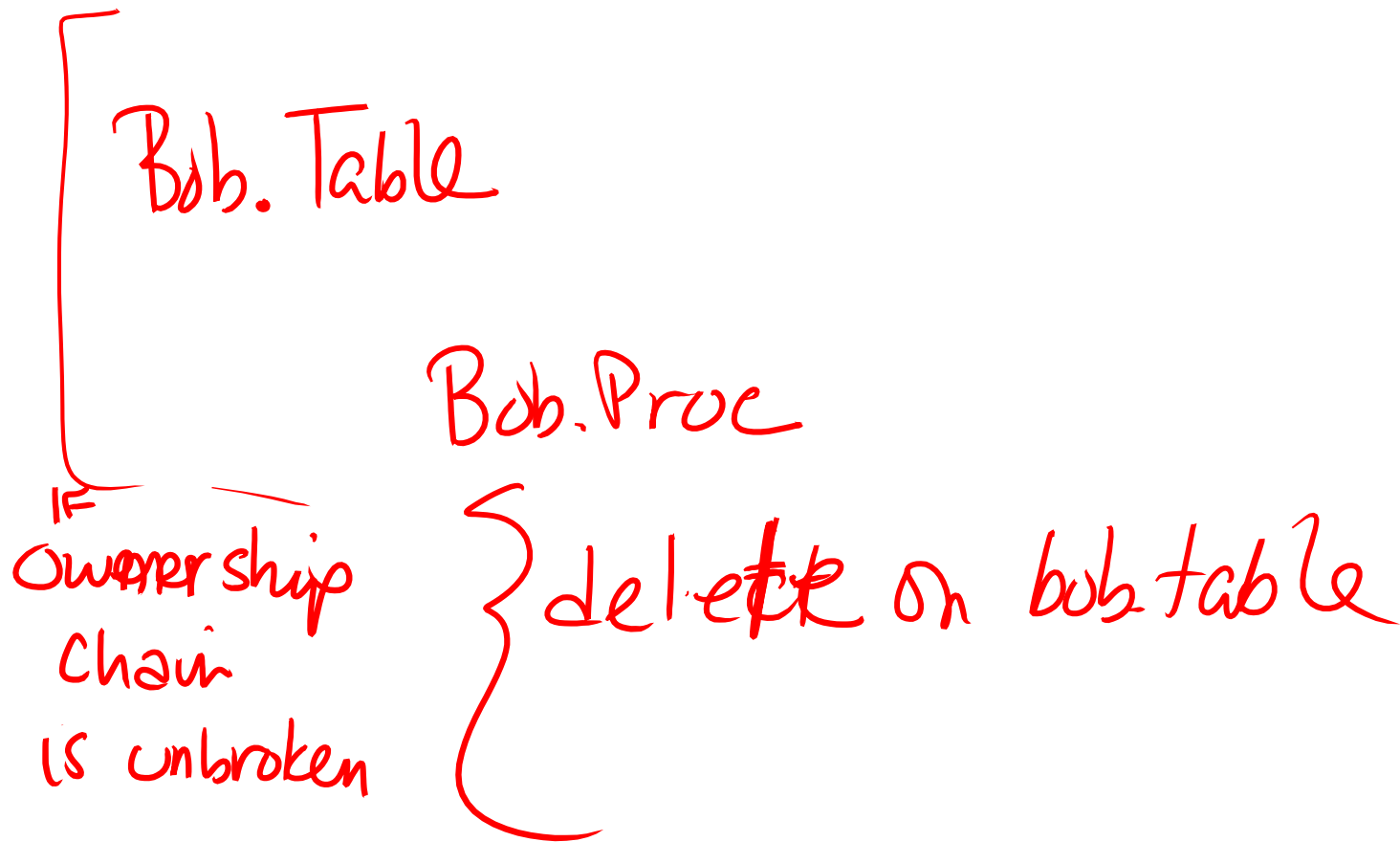
USER

~~(higher)~~

low

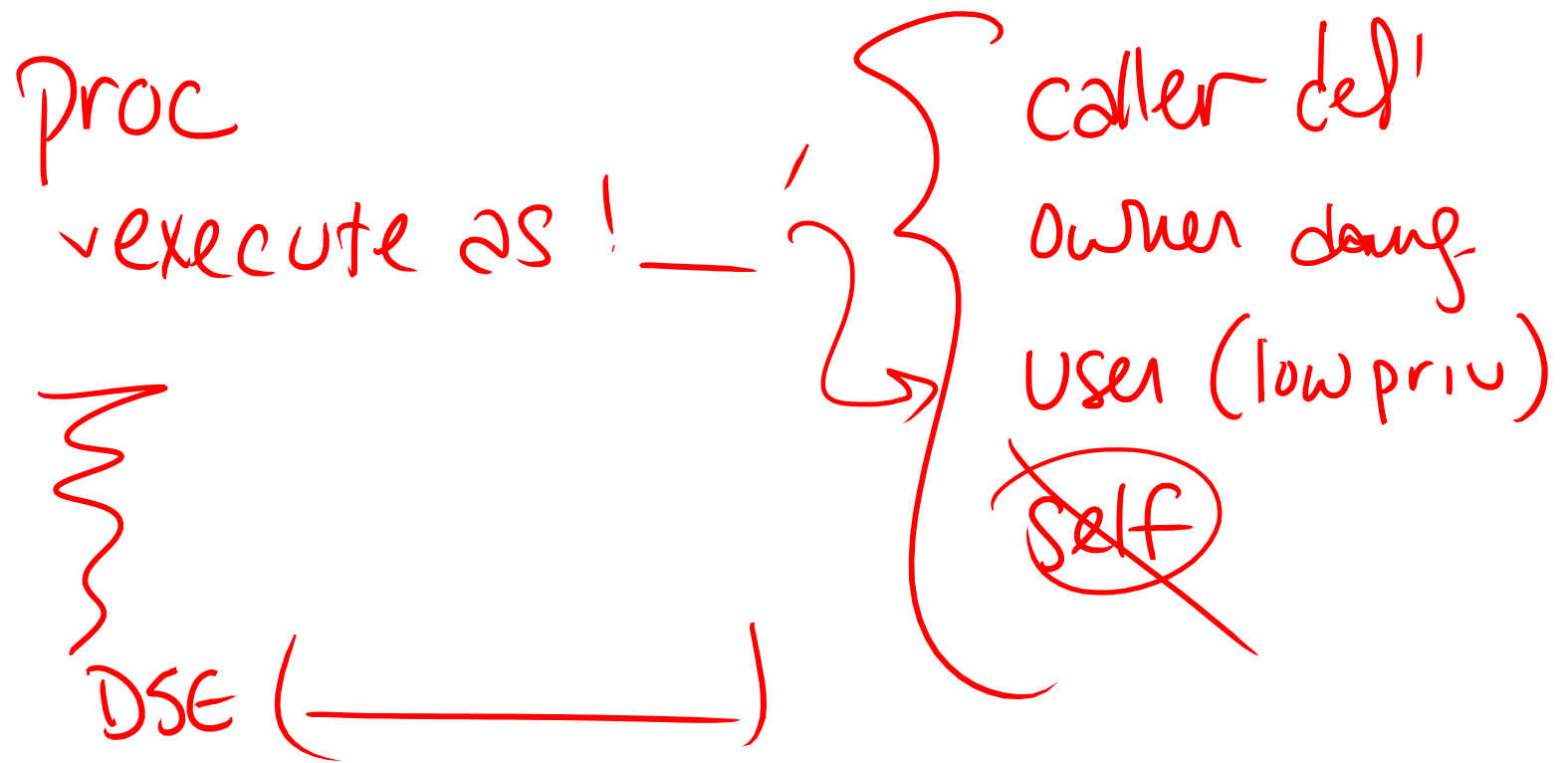
~~owner~~

~~self~~ stupid



Without Dynamic String Execution permissions flow through the ownership chain

The idea is that when the ownership chain is unbroken then direct permissions (for example, to do a delete) are not necessary. As an owner you are giving rights to a process; this procedure can be protected with more business rules/logic and even reduced/isolated to specific rows. The end user does NOT need direct delete permissions to be able to execute the procedure as long as they've been given execute rights on the sproc.



Dynamic String Execution, EXECUTE AS and SQL Injection

People complained that granting permissions directly to the user was a problem... so, they wanted an alternative. Enter: EXECUTE AS.

While it solves the problem of permissions, it adds more potential for SQLInjection.

So..... Check out these posts:

[Little Bobby Tables, SQL Injection and EXECUTE AS](#)

["EXECUTE AS" and an important update your DDL Triggers \(for auditing or prevention\)](#)

✓ fn
✓ (fn + ln) ?

Stable plans vs. unstable plans – how do you know?

- You could test different combinations using execute with recompile.
- You might know because of selectivity?

* fn + ln + mno ← stable/cons
✓ ln

* ln + mno ←
* mno ←
* mno + fn ←
} ← stable/cons

Options that better balancing recompiling too much vs. not recompiling enough

This is from the multipurpose procedure demo. There are 3 parameters and 7 possible combinations that can be submitted. Instead of adding `OPTION (RECOMPILE)` to the statement (which doesn't always work [remember, it has a very checkered past]) you can build the statement dynamically and then programmatically (by setting a `RECOMPILE` flag to 0 or 1) decide whether or not `OPTION(RECOMPILE)` should be added to the dynamically constructed statement. Then, you can use `sp_executesql` to place ALL seven statements in cache. Those without the `OPTION(RECOMPILE)` clause will be cached. Those with `OPTION (RECOMPILE)` will get a plan on each execution.

15 executions?

15 plans?

2 plans

n 6/5

no good /
Single plan

OPTION
(RECOMPILE)

DSE?

8/7

IF Subproc
Else Subproc

14

(1) Sep proc
recompile only work if
you did not
OPTION opt for

How many plans COULD be generated?

When you're testing your procedures you might have 15 test cases. Executing each of the procedures (with the 15 different combinations of parameters) WITH RECOMPILE. Then, review the number of different plans

If there are 15 different plans for 15 different executions – Use OPTION (RECOMPILE)

If there are 2 primary plans that are fairly split – see if you can create subprocedures and then control the behavior that way (especially if they either don't need to recompile OR only one needs to recompile).

If there are only 2 plans where one is 14 of the executions (and, is the norm) then you might use OPTION (OPTIMIZE FOR ...) but it depends on how bad the performance is the for 15th. Can you control that through [other] parameters or programmatically?

① recompile always

- more CPU

+ Plan spec. to those parameters

Default

PSP

+ less CPU

- possibly horrible plans

- not consistent (who gets there first)

UNKNOWN (Avg)

- Bad plan? for atypical values

How many plans COULD be generated?

This is another way to describe what was on the prior page...

These are the pros/cons of the different procedure (and statements within) strategies:

The default

- Prone to parameter sniffing problems
- Possibly HORRIBLE plans for subsequent executions
- No guarantee of what plan actually gets into cache
- + Less CPU (not always though – you do save compilation time but if the execution is really expensive then you might lose all of the gains of having compiled/saved a plan)

Forcing a recompile for EVERY execution

- More CPU
- + Plan specific to those parameters

Optimize for UNKNOWN

- Guarantees an AVERAGE plan
- + Works well when the data is evenly distributed (but, the problems are less likely as well – but, if plans were getting into cache from atypical values then this might solve the problem)

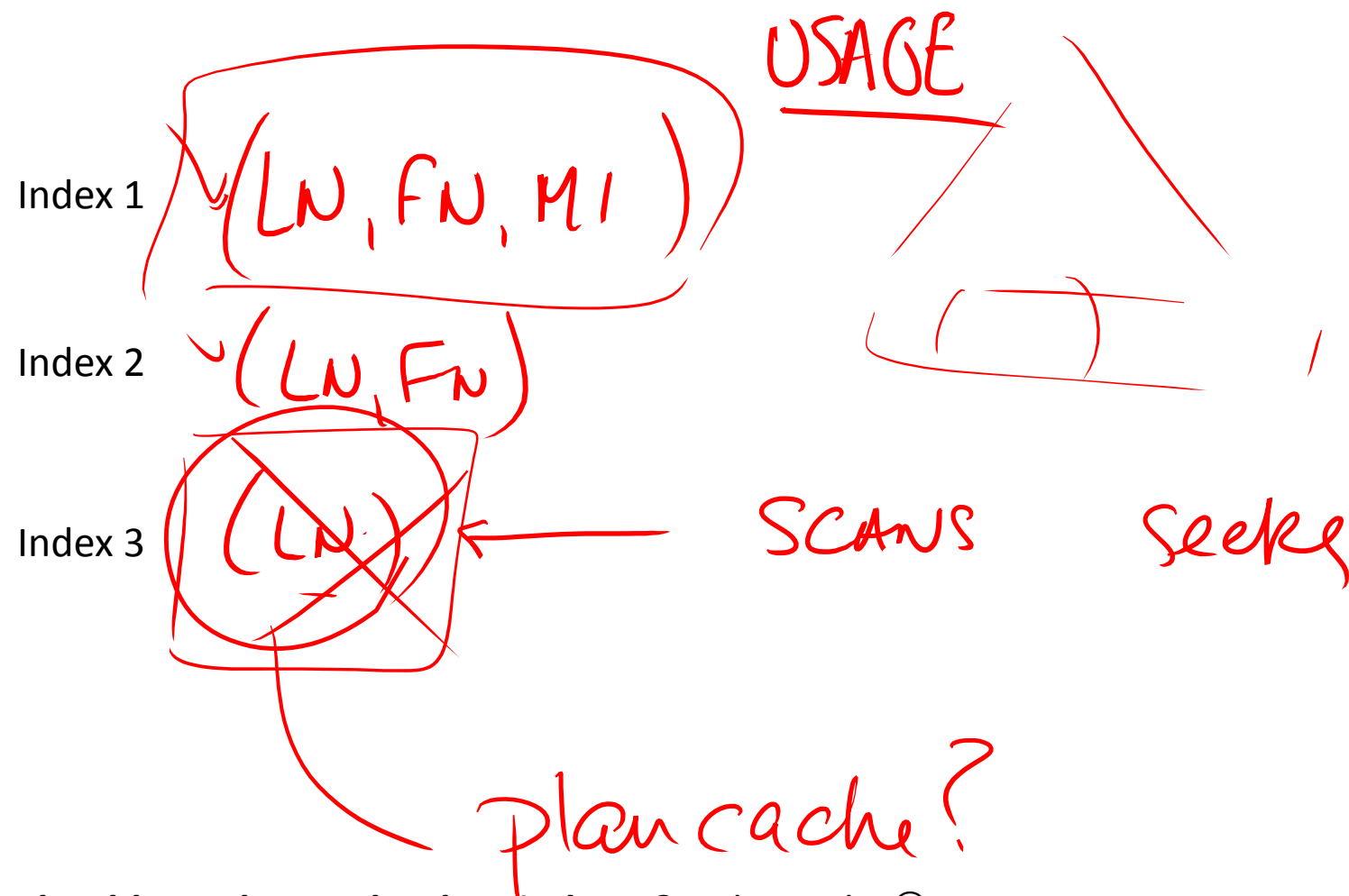
SQLskills Immersion Event

IEPTO2: Performance Tuning and Optimization

Module 10: Index Analysis

Kimberly L. Tripp
Kimberly@SQLskills.com

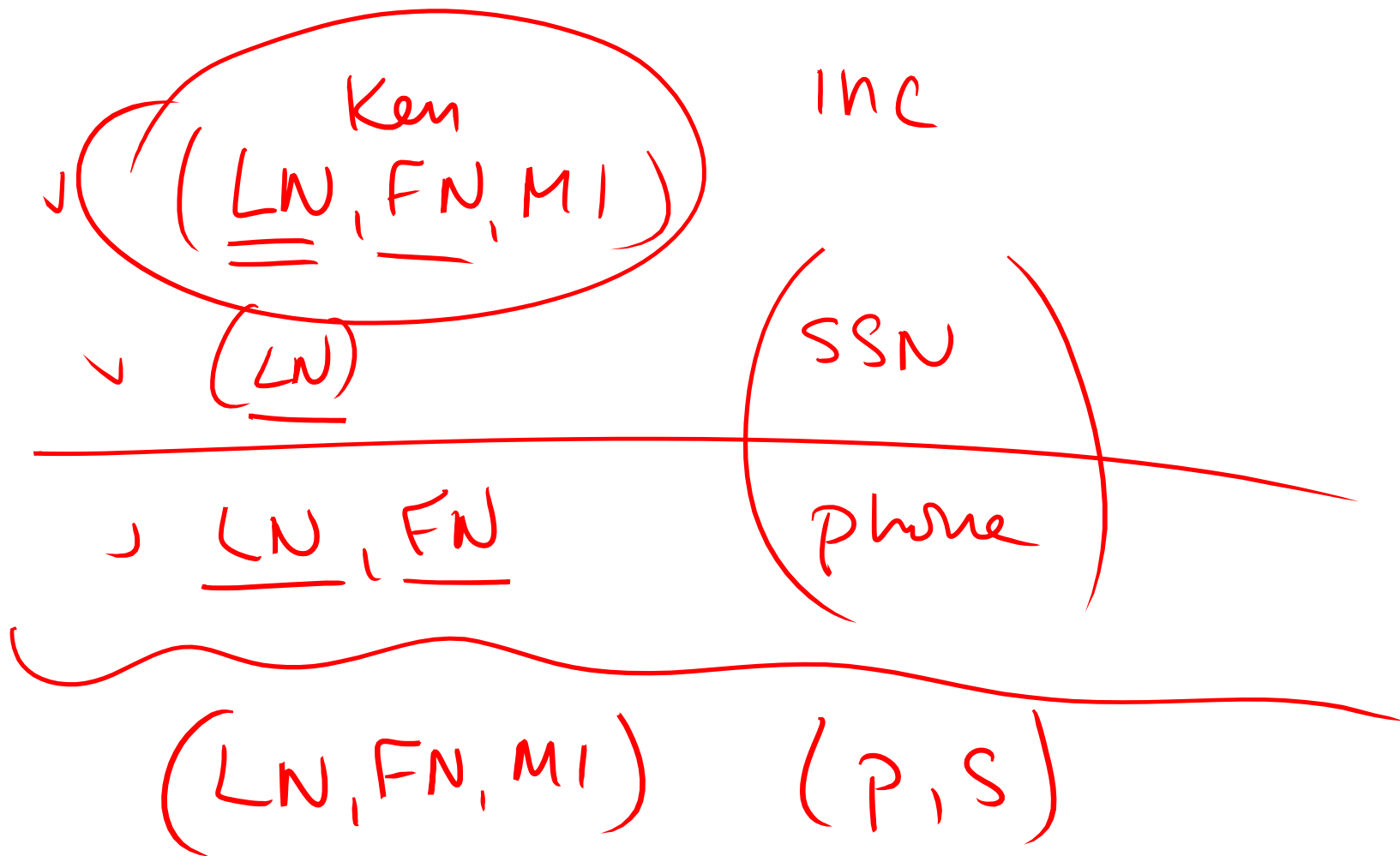




Should you drop redundant indexes? It depends. 😊

It's TRUE that a query that uses a left-based subset of one index (let's say that a query uses index 2) that it COULD use a wider version of the same (left-based) index. It could use Index 1 instead of 2.

The problem is that there could be specific (and very frequent) uses that warrant the smaller version. Questions to ask: Is the smaller version used by scans? Are they executed VERY frequently? How critical / important are they?



Index consolidation

No production environment can support THE BEST index for **every** query as you'd have too many indexes. To be really successful with indexing – you need to find a balance. This includes consolidating some of your index recommendations to have one index handle multiple purposes!

Key Inc

C1 = C2 ~~C3~~

C1, C2, C3 no inc

Key = navigation

Key must be pres.

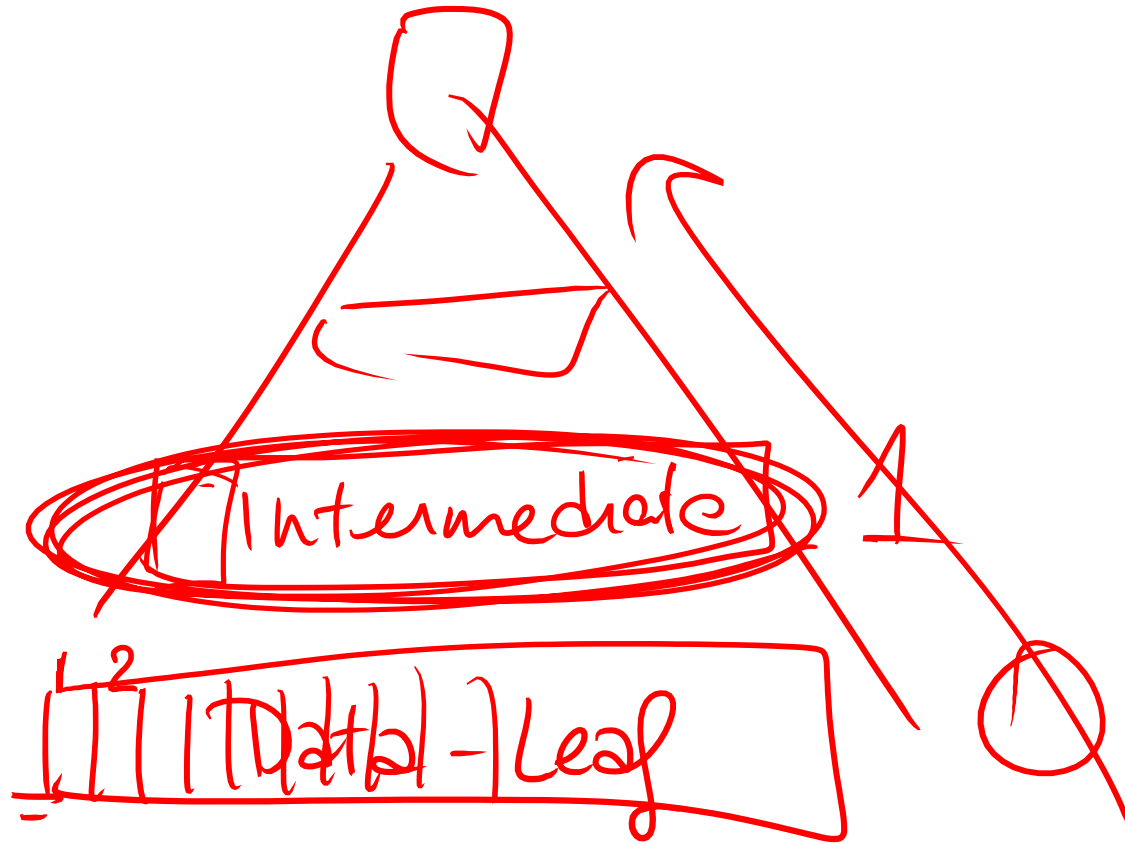
Key Inc

<u>LN, FN</u>	
<u>LN</u>	
<u>LN, FN, MI</u>	
LN	SSN
<u>LN, FN, MI (Inc SSN)</u>	

Consolidation?

You can CONSIDER removing indexes that are left-based subsets of others (in terms of the key). However, be sure to preserve the key as that's used for navigation/seeking. The columns that are included can be combined in any order.

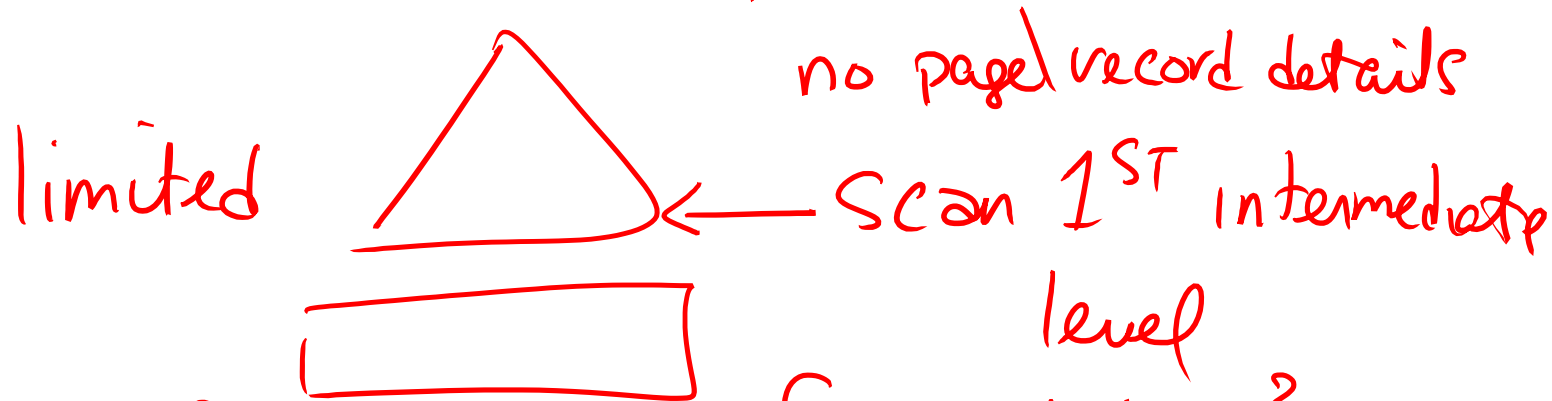
47



Index Health – Checking for fragmentation

To get the value: **avg_fragmentation_in_percent** for **sys.dm_db_index_physical_stats**, SQL Server uses a scan of the first intermediate level (index level 1). All of the modes (limited, sampled, and detailed), use this same method. Then, SAMPLED and DETAILED do additional work (on the leaf level of the index) to give you things like average page density, etc. If your analysis/rebuild scripts only use **avg_fragmentation_in_percent** to determine fragmentation then you should only do a LIMITED mode scan. SAMPLED and DETAILED offer more page-specific details (page density, forwarding pointers, etc.) and if you're NOT using those during your automated maintenance process (and most people aren't) don't waste time!

Sys. dm-db-index-physical-staff



Good for avg-fragmentation - in percent fragmentation?

sampled } Avg page density, page fullness
detailed } row size

→ 1/100 of leaf / detailed reads all pages
all levels

FN like K%

rno > 6

mno < 5000

fn, mno, rno

Katie 6 4

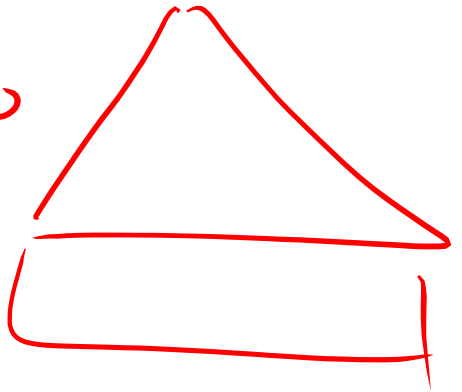
Katie 7892 6

Kiera

Kiera

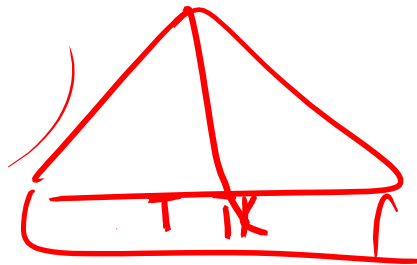
Kimberly

Kimberly

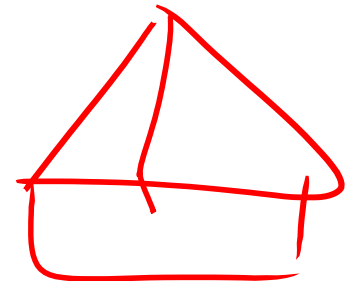


FN = Kimberly

LN = TREPP



LN, FN



FN, LN

Adding more indexes

Here we were discussing the order of columns. If there are range-based predicates then it doesn't usually matter which column is second. You tend to want to put the most selective first (in terms of the PREDICATES supplied). If all of the predicates are equality-based then it doesn't really matter. You're best ordering them by the LEFT-based combination that's used most commonly.

KEY

Density_Vector

c1

c2

c3

c4

c2

c4

c2

c3

c3

c4

c2

c3

c4

Multi-column, column-level statistics

Only DTA recommends multi-column, column-level statistics...

If you are about to trust (and create) a “green hint” recommendation (which comes from the Missing Index DMVs) then you might want to first run the query through DTA. If you end up creating the index that DTA recommends then you should also create the recommended stats (for that same table). These statistics can be helpful to the optimizer to better understand non-left-based subsets of the index for other possible uses.

<u>Step</u>			DIST.	AVG
50	}			
100		57		
200		<u>250K</u>	5	<u>50K</u>
		<u>29472</u>		

This was around our discussion about the limitation of SQL Server's knowledge of the data with gaps in the step values.

If the step describes the range from 101 to 200 (not including those values)

- There are 250K rows over that range
- There are 5 distinct values

The average will show 50K and EVERY value supplied between 101 and 199 will get an estimate of 50K. There's no way for SQL Server to know WHICH values actually exist.

Since we discussed columnstore indexes a bit, I thought I'd throw in some of the general indexing strategy slides delivered as part of my SQLintersection "Indexing Strategies" session (delivered in Orlando in April 2016).

SQLintersection

SQL Server Indexing Strategies

Kimberly L. Tripp

President / Founder, SQLskills.com

Kimberly@SQLskills.com

@KimberlyLTripp



SQL
intersection



Biggest Concern is SQL Server Version

- **SQL Server 2008 is the lowest (IMO) version for large tables, performance, scalability**
 - Added data compression (row and page compression)
 - Added filtered indexes / filtered statistics
 - Fixed fast-switching for partition-aligned, indexed views
- **SQL Server 2012 adds read-only, nonclustered columnstore indexes**
 - Some frustrating limitations but still amazing performance when possible and/or workarounds used (see this WIKI for SQL Server 2012 workarounds: <http://bit.ly/1eHVV00>)
- **SQL Server 2014 adds updateable, clustered columnstore indexes**
 - Many of the most frustrating limitations with CS fixed – for example, UNION ALL supports batch mode (which means you can use these with partitioned views)
 - Added “incremental statistics” to help reduce when to rebuild as well as time to rebuild
- **SQL Server 2016 takes columnstore indexes even further with updateable nonclustered columnstore indexes and row-based nonclustered indexes on clustered columnstore indexes!**

SQL Server 2008 / R2 and SQL Server 2012

- **If you're not on SQL Server 2012 – your options are limited to row-based indexes**
 - Skip 2012 on your migration path and go directly to 2014 or 2016*
- **If you're on SQL Server 2012 AND you have DSS / RDW workloads with partitioning (partitioned views or partitioned tables):**
 - Create and test a nonclustered columnstore index on your large fact tables
 - They're super easy to create
 - You can have only one
 - There's no column order
 - It's highly compressed so it doesn't take a lot of space
 - You might get HUGE gains for large scan / aggregate queries
 - If you find you're not getting batch mode processing
 - Check out Eric Hanson's WIKI and workarounds: <http://bit.ly/ZP63Lr>
 - Consider upgrading to SQL Server 2014 (lots of batch mode fixes)

What Type of Workload is Running?

- **OLTP – Online transaction processing**

- Priority is toward modifications
- Lots of point queries (highly-selective queries of a small number of rows)

Search for a sale, search for a customer, lookup a product, ...

- **Dedicated Decision Support System / Relational Data Warehouse**

- Priority is toward large-scale aggregates
- Very large percentage or even the entire dataset – is being evaluated often

- **Hybrid**

- OLTP might be the priority and some point query activity
- Some range-based queries because “management” wants real-time analysis

*(NOTE: We're talking indexes in this session but in the hybrid environment you should really be running with versioning enabled! Check out the data option: `read_committed_snapshot`. For more info, see resources. However, **clustered** columnstore indexes do not work yet in versioned databases until SQL Server 2016.)*

General Indexing Strategies Based on Workload

- **OLTP – Online transaction processing**
 - Traditional row-based clustered and nonclustered indexes
- **Dedicated Decision Support System / Relational Data Warehouse**
 - Prior to SQL Server 2012
 - Traditional row-based clustered and nonclustered indexes
 - SQL Server 2012+:
 - Consider adding a read-only, nonclustered, columnstore index for partitioned objects leveraging partition switching as additional data is added
 - SQL Server 2014+:
 - Use the SQL Server 2012+ strategy above
 - Or, consider the new updateable clustered columnstore index
- **Hybrid**
 - Most likely to use traditional row-based clustered and nonclustered indexes and possibly nonclustered columnstore indexes if you've partitioned your data in the hybrid scenario