

Oracle Endeca Commerce

MDEX Engine Performance Tuning Guide

Version 6.4.0 • December 2012



Contents

Preface.....	9
About this guide.....	9
Who should use this guide.....	9
Conventions used in this guide.....	9
Contacting Oracle Support.....	10
 Chapter 1: Before You Begin.....	 11
About the Dgraph.....	11
Important concepts.....	11
 Chapter 2: System Characteristics and Hardware.....	 13
MDEX Engine architecture and performance.....	13
Storage considerations.....	15
Locally attached RAID storage (RAID 5/6, RAID 10, or RAID 0).....	15
SAN-backed network-attached storage.....	16
Memory considerations.....	16
About Dgraph process memory usage.....	16
Memory usage recommendations for optimizing performance.....	17
Dgraph virtual memory vs. RAM: use cases.....	17
Solutions for memory-based Dgraph performance problems.....	18
About the Dgraph cache.....	19
About the File System Cache.....	20
Cache-tuning recommendations for optimizing performance.....	20
Warming the Dgraph cache after an update.....	22
The Dgraph cache and its impact on virtual process size.....	22
Estimating the MDEX Engine RAM requirements.....	23
Network considerations.....	25
Dgidx performance recommendations.....	25
Operating system considerations.....	26
Windows 2008 performance considerations.....	26
VMware performance considerations.....	26
Linux considerations.....	28
Load balancer considerations.....	29
Load balancing and session affinity.....	30
High availability considerations.....	30
 Chapter 3: Using Multithreaded Mode.....	 33
About multithreaded mode.....	33
Benefits of multithreaded MDEX Engine.....	33
The MDEX Engine threading pool.....	34
Configuring the number of MDEX Engine threads.....	35
When to increase the number of threads.....	35
Multithreaded MDEX Engine performance.....	36
Recommended threading strategies and OS platform.....	36
 Chapter 4: Diagnosing Dgraph Problems.....	 39
Information you need.....	39
System state characteristics.....	39
Performance tools overview.....	40
Dgraph performance issues.....	42
Improving the speed of Dgraph startup.....	42
Tips for troubleshooting long processing time.....	42
Warming performance vs. steady state performance.....	44
About planning for peak Dgraph load.....	44
About tuning the number of threads.....	44
Multithreaded Dgraphs on machines with multithreaded processors.....	44

Multiple Dgraphs on one machine vs. multithreaded Dgraphs.....	45
Disk access recommendations for optimizing performance.....	45
CPU recommendations for optimizing performance.....	46
I/O recommendations for optimizing performance.....	46
Identifying problems with resource usage by the application.....	46
Coding practices for the front-end application.....	47
Web application ephemeral port contention.....	47
Recommendations for identifying network problems.....	47
Troubleshooting connection errors.....	49
Next steps.....	50

Chapter 5: Dgraph Analysis and Tuning.....51

Feature performance overview.....	51
Endeca record configuration.....	51
Record select.....	51
Aggregated records.....	52
Dimensions and dimension values.....	52
Hidden dimensions.....	53
Dimensions and dimension values with high record coverage.....	53
Flat dimension hierarchy.....	53
Displaying multiselect dimensions.....	53
Multi-assign dimensions.....	53
Displaying refinement dimension values.....	54
Dynamic statistics on dimension values.....	54
Aggregated refinement counts.....	55
Dynamic refinement ranking and performance.....	55
Disabled refinements.....	55
Displaying dimension value properties.....	56
Collapsible dimension values.....	56
Mapping source properties.....	56
Indexing all properties with Dgidx.....	57
Record sorting and filtering.....	57
Sorting records by dimension or property.....	57
Geospatial sorting and filtering.....	58
Range filters.....	58
Record filters.....	58
Optimizing URL record filters that use complex logic.....	59
EQL expressions and Record Relationship Navigation.....	60
When to use EQL-based filters vs. other filter types.....	60
Performance impact of EQL-based filters.....	62
Performance impact of RRN.....	63
Tips for troubleshooting EQL filters.....	65
Typical causes of EQL filter errors.....	65
Snippeting.....	66
Spelling auto-correction and Did You Mean.....	66
Spelling auto-correction.....	66
Did You Mean.....	67
Stemming and thesaurus.....	68
Guidelines for thesaurus development.....	68
Record, phrase, and dimension search.....	69
Record search.....	69
Boolean search.....	69
Phrase search.....	70
Wildcard search.....	70
Dimension search.....	72
Precedence rules.....	73
About precedence rules.....	73
Relevance ranking.....	74
Minimizing the performance impact of relevance ranking.....	74
Dynamic business rules.....	75
Analytics performance considerations.....	75

Appendix A: The MDEX Engine Request Log.....77

About the MDEX Engine request log.....	77
--	----

Request log file format.....	77
Extracting information from request logs.....	81
Storing logs on a separate physical drive.....	82
Request log rolling.....	82

Appendix B: The MDEX Engine Parameter Listing.....85

Understanding the URL parameter mapping.....	85
Mappings between request log and UriENQuery URL parameters.....	85
List of request log parameters.....	88
Example: interpreting error log messages	88
Description of query types.....	89
allbins.....	89
allgroups.....	90
analytics.....	90
attrs.....	91
autophrase.....	91
autophrasedwim.....	92
compound.....	92
dym.....	92
filter.....	93
format.....	93
group.....	94
groupby.....	94
id.....	95
ignore.....	95
irversion.....	96
keyprops.....	96
lang.....	97
log.....	97
merchdebug.....	97
merchpreviewtime.....	98
merchrulefilter.....	98
model.....	98
nbins.....	99
nbulkbins.....	99
node.....	100
num.....	101
offset.....	101
op.....	102
opts.....	102
pred.....	103
pretendtime.....	103
profiles.....	104
rank.....	104
refinement.....	105
relrank.....	105
select.....	105
sort.....	106
structured.....	106
terms.....	107

Appendix C: Creating Eneperf input files with the Request Log Parser.109

Installation location.....	109
Log format requirements.....	109
Invoking the Request Log Parser.....	109
Example output from the Request Log Parser.....	111
Using the Request Log Parser with Eneperf.....	112

Appendix D: Using the Eneperf Tool.....113

About Eneperf.....	113
Using Eneperf.....	113
Required settings.....	114
Optional settings.....	117
Example of Eneperf output.....	121

About the format of logs for use with Eneperf.....	123
The Request Log Parser.....	123
Recommendations for generating a representative log for Eneperf.....	123
Running Eneperf in two-stream mode: regular logs and logs with updates.....	124
Converting an MDEX Engine request log file for Eneperf.....	126
Performance testing .NET 2.0 applications that contain long or complex queries.....	126
Creating a log file by hand using substitute search terms.....	126
Debugging Eneperf.....	127
Appendix E: Using the Request Log Analyzer.....	129
About the Request Log Analyzer.....	129
Installation location.....	129
Log format requirements.....	129
Invoking the Request Log Analyzer.....	130
Show flags.....	131
Threshold flags	131
Ignore flags.....	132
Timeframe flags.....	133
Interpreting reports.....	134
Statistics.....	135
Common metrics.....	135
Hourly results.....	136
Longest-running requests by round-trip response time	137
Longest-running requests by engine-only processing time.....	137
Query types.....	137
Extended query types.....	139
Response codes.....	140
Request profiling	141
Response profiling.....	142
Peak performance	143
Threading and queueing information	143
Summary information	145
Appendix F: MDEX Engine Statistics and Auditing.....	147
About the MDEX Engine Statistics page.....	147
Viewing the MDEX Engine Statistics page.....	147
Sections of the MDEX Engine Statistics page.....	148
The Performance Summary tab.....	148
The General Information tab.....	148
The Index Preparation tab.....	149
The Cache tab.....	149
The Details tab.....	150
About the Endeca MDEX Engine Auditing page.....	152
Viewing the MDEX Engine Auditing page.....	153
Audit persistence file details.....	153
Sections of the MDEX Engine Auditing page	153
The Audit Stats tab.....	154
The General Information tab.....	155
Appendix G: Useful Third-Party Tools.....	157
Cross-platform tools.....	157
Solaris and Linux tools.....	157
Solaris-specific tools.....	158
Linux-specific tools.....	159
Windows tools.....	159
Appendix H: Tuning the Network Performance.....	161
Tuning network performance on Windows.....	161
Tuning network performance on Solaris.....	162
Configuring the FIN_WAIT_2 timeout interval.....	162
Configuring FIN_WAIT_2 timeout on Linux.....	163
Configuring FIN_WAIT_2 timeout on Solaris.....	163
Configuring FIN_WAIT_2 timeout on Windows.....	163

Copyright and disclaimer

Copyright © 2003, 2012, Oracle and/or its affiliates. All rights reserved.

Oracle and Java are registered trademarks of Oracle and/or its affiliates. Other names may be trademarks of their respective owners.

This software and related documentation are provided under a license agreement containing restrictions on use and disclosure and are protected by intellectual property laws. Except as expressly permitted in your license agreement or allowed by law, you may not use, copy, reproduce, translate, broadcast, modify, license, transmit, distribute, exhibit, perform, publish or display any part, in any form, or by any means. Reverse engineering, disassembly, or decompilation of this software, unless required by law for interoperability, is prohibited.

The information contained herein is subject to change without notice and is not warranted to be error-free. If you find any errors, please report them to us in writing.

If this is software or related documentation that is delivered to the U.S. Government or anyone licensing it on behalf of the U.S. Government, the following notice is applicable:

U.S. GOVERNMENT END USERS: Oracle programs, including any operating system, integrated software, any programs installed on the hardware, and/or documentation, delivered to U.S. Government end users are "commercial computer software" pursuant to the applicable Federal Acquisition Regulation and agency-specific supplemental regulations. As such, use, duplication, disclosure, modification, and adaptation of the programs, including any operating system, integrated software, any programs installed on the hardware, and/or documentation, shall be subject to license terms and license restrictions applicable to the programs. No other rights are granted to the U.S. Government.

This software or hardware is developed for general use in a variety of information management applications. It is not developed or intended for use in any inherently dangerous applications, including applications that may create a risk of personal injury. If you use this software or hardware in dangerous applications, then you shall be responsible to take all appropriate fail-safe, backup, redundancy, and other measures to ensure its safe use. Oracle Corporation and its affiliates disclaim any liability for any damages caused by use of this software or hardware in dangerous applications.

This software or hardware and documentation may provide access to or information on content, products and services from third parties. Oracle Corporation and its affiliates are not responsible for and expressly disclaim all warranties of any kind with respect to third-party content, products, and services. Oracle Corporation and its affiliates will not be responsible for any loss, costs, or damages incurred due to your access to or use of third-party content, products, or services.

Preface

The Oracle Endeca Commerce solution enables your company to deliver a personalized, consistent customer buying experience across all channels — online, in-store, mobile, or social. Whenever and wherever customers engage with your business, the Oracle Endeca Commerce solution delivers, analyzes, and targets just the right content to just the right customer to encourage clicks and drive business results.

Oracle Endeca Commerce is the most effective way for your customers to dynamically explore your storefront and find relevant and desired items quickly. An industry-leading faceted search and Guided Navigation solution, Oracle Endeca Commerce enables businesses to help guide and influence customers in each step of their search experience. At the core of Oracle Endeca Commerce is the MDEX Engine™, a hybrid search-analytical database specifically designed for high-performance exploration and discovery. The Endeca Content Acquisition System provides a set of extensible mechanisms to bring both structured data and unstructured content into the MDEX Engine from a variety of source systems. Endeca Assembler dynamically assembles content from any resource and seamlessly combines it with results from the MDEX Engine.

Oracle Endeca Experience Manager is a single, flexible solution that enables you to create, deliver, and manage content-rich, cross-channel customer experiences. It also enables non-technical business users to deliver targeted, user-centric online experiences in a scalable way — creating always-relevant customer interactions that increase conversion rates and accelerate cross-channel sales. Non-technical users can control how, where, when, and what type of content is presented in response to any search, category selection, or facet refinement.

These components — along with additional modules for SEO, Social, and Mobile channel support — make up the core of Oracle Endeca Experience Manager, a customer experience management platform focused on delivering the most relevant, targeted, and optimized experience for every customer, at every step, across all customer touch points.

About this guide

This guide describes how to diagnose and tune Dgidx and the Dgraph to provide optimal performance. It also includes hardware provisioning recommendations as well as storage, memory, and network support recommendations.

Who should use this guide

This guide is intended for system administrators and developers responsible for the performance of an Endeca implementation.

Conventions used in this guide

This guide uses the following typographical conventions:

Code examples, inline references to code elements, file names, and user input are set in monospace font. In the case of long lines of code, or when inline monospace text occurs at the end of a line, the following symbol is used to show that the content continues on to the next line: ↵

When copying and pasting such examples, ensure that any occurrences of the symbol and the corresponding line break are deleted and any remaining space is closed up.

Contacting Oracle Support

Oracle Support provides registered users with important information regarding Oracle Endeca software, implementation questions, product and solution help, as well as overall news and updates.

You can contact Oracle Support through Oracle's Support portal, My Oracle Support at <https://support.oracle.com>.

Chapter 1

Before You Begin

This section provides background information you should know before you begin to diagnose performance problems in your Endeca implementation.

About the Dgraph

A typical Endeca implementation includes one or more Dgraphs. The Dgraph is the name of the process for the MDEX Engine, which is the query engine that provides the backbone for all Endeca solutions. The Dgraph uses proprietary data structures and algorithms that allow it to provide real-time responses to client requests. You can use a single Dgraph or a set of load-balanced Dgraphs. Because the Dgraph is key to every Endeca implementation, its performance is critical.

Important concepts

There is a small set of terms and concepts you should be familiar with as you read this guide.

The following terms are used to discuss the performance of the MDEX Engine:

- *Throughput* is the number of requests processed by the MDEX Engine per unit of time. In this guide, unless otherwise specified, it is expressed as query operations per second (ops/sec). Throughput is measured with the performance tool Eneperf using an MDEX Engine request log.
- *Dgraph sustained throughput* is the measure of query capacity, that is, the maximum number of requests that can be consistently processed by the MDEX Engine per second.
- *Latency* is how fast the MDEX Engine responds to queries, or the time it takes for a query to be returned by the Engine, typically in milliseconds.
- *Maximum latency* is the maximum time it takes for the longest query to be returned by the MDEX Engine.



Note: Though latency and throughput are related, they are not directly derivable from one another. The inverse of the average latency is a lower bound on the maximum throughput. For example, if the average latency for a shopper in a supermarket checkout line is 5 minutes, we know that the checkout throughput of the store must be at least 0.2 shoppers per minute. In addition, latency and throughput are tied together by concurrency. Using the same example, the real maximum throughput may be 10 shoppers per minute because there are many checkout lanes.

- An *operation* is defined as a single request to the MDEX Engine.

Such a request may have one of the following types:

- Navigation (possibly including record search, analytics, and so on)
 - Dimension search
 - Record search
 - Aggregated record
 - Administration (such as a Web Service invocation for administrative purposes, statistics, configuration update, partial update, and so on)
- *Memory bandwidth* is the rate at which data can be read from or stored in memory by a processor. It is measured in bytes per second. In relation to MDEX Engine performance, you may be interested in the memory bandwidth that a system can sustain while running a Dgraph or multiple Dgraphs.
 - The *virtual process size (or address space)* for the Dgraph is the total amount of virtual memory allocated by the operating system to the MDEX Engine process at any point in time. This includes the Dgraph code, the MDEX Engine data as represented on disk, the Dgraph cache and any temporary work space.
 - *Resident set size (RSS)* is the amount of physical memory currently allocated and used by the MDEX Engine process. As the MDEX Engine process runs, the active executable code and data are brought into RAM, becoming part of the RSS for the MDEX Engine.

You can view the resident set size of a process on Linux by using `ps -o pid,ucomm, rss` commands, `ucomm`, or by using the `top` program which reports the RSS size.

- The *working set size (WSS)* of the MDEX Engine process is the amount of physical memory needed for those parts of the process that have been most recently and frequently accessed. In other words, the Dgraph WSS is the amount of memory a Dgraph process is consuming now and that is needed to avoid paging.

The WSS of the Dgraph process directly affects RAM usage. As the working set increases, the Dgraph process memory demand increases. With a larger WSS, a process needs more memory to run with acceptable performance.

You cannot measure the WSS, but you can make assumptions about it when you measure the resident set size and observe performance; performance tends to degrade if the RSS cannot equal the WSS.

- The *Dgraph cache* is an area of memory set aside for dynamically saving the partial and complete results of processing queries.
- *Warming* is the process during which the MDEX Engine performance gradually increases to a steady state. A gradual increase in performance takes place either as the MDEX Engine starts up and processes queries or following a partial update.
- *Utilization* is the percentage of the total capacity of a resource that is actually being used.
- *The number of concurrent users* is the number of site users engaging the MDEX Engine at any given time. When planning for Dgraph capacity based on the number of concurrent users, take into account the consideration that multiple users on a site do not continually issue queries. Between queries, a user has some "think time" before issuing another query.

System Characteristics and Hardware

This section provides recommendations for hardware used for an Endeca implementation and discusses typical hardware-based issues that affect performance of the MDEX Engine.

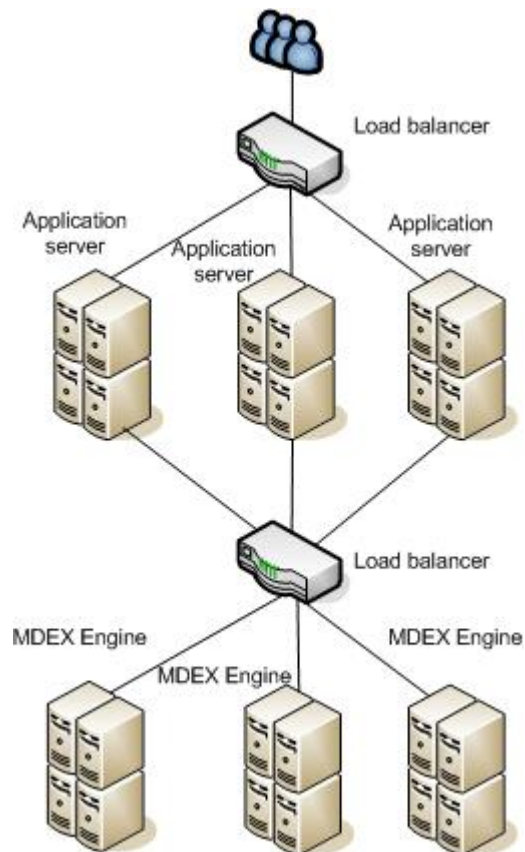
MDEX Engine architecture and performance

The MDEX Engine is optimized for performance. This section reviews those characteristics of the Engine that have a direct impact on its performance.

Hardware architecture diagram

The following diagram represents a typical MDEX Engine deployment architecture. It shows a set of application servers and MDEX Engines, each with a dedicated hardware load balancer. The Information Transformation Layer (ITL) that supplies data to the MDEX Engine index is not shown.

In this diagram, a load balancer directs query requests to one of the MDEX Engines. If you are using servers with dual-core or quad-core processors, multiple multithreaded MDEX Engines can be configured on the same machine, with two or more threads configured for each MDEX Engine.



Resource utilization

The MDEX Engine stores index structures in system memory to provide rapid access during query execution. Less frequently accessed structures and record data are stored on disk; these are only pulled into RAM as needed.

Storage locality

The data and indexes are stored in memory and on disk in a manner that provides optimal locality for common access patterns. When queries have to access disk to retrieve information, they find all the data required with the minimum number of seek operations. This decreases the cumulative disk access seek times thereby decreasing the time needed for query processing and increasing query throughput.

Unified Dgraph cache

The MDEX Engine has a unified dynamic cache where it stores intermediate results and index structures for future processing. When similar requests are made to the Engine with slight changes (example: sorting by price, then ranking, then popularity), the Engine stores intermediate results in the cache. This allows for the optimal reuse of data previously retrieved from slower sources, such as disk. The cache is dynamically managed by the MDEX Engine to keep the optimal data cached for the current query patterns.

Stateless architecture combined with load balancing

The Endeca implementation has a stateless server architecture. Query processing does not require any state information about prior queries from this client or other clients. Because of this, when multiple

identical MDEX Engines are placed in parallel behind a load balancer, the response will be identical regardless of which server receives the request. The throughput of such a system is equal to the throughput of a single server times the number of parallel servers.

Multithreaded mode

The MDEX Engine always runs in multithreaded mode with the total number of threads set to 1 by default. Oracle recommends to increase this number to maximize the use of system resources. On processors that are multithreaded or multicore, multiple query threads can use a single processor at the same time.

64-bit architecture

The MDEX Engine utilizes 64-bit operating systems and processors, and can store and access larger volumes of data with scale. The MDEX Engine can utilize as much physical memory as can be placed in a server. Running in the 64-bit environment, the MDEX Engine can service many memory-intensive requests simultaneously without the risk of running out of memory address space. This, combined with a large Dgraph cache (1GB), provides a significant performance benefit.

Storage considerations

Oracle recommends using one of two storage approaches with Oracle Endeca Guided Search implementations -- RAID or SAN-backed network-attached storage (if using RAID is not possible).

Locally attached RAID storage (RAID 5/6, RAID 10, or RAID 0)

For RAID disks, use these Endeca recommendations.

Storage availability after disk failure is usually a requirement for your RAID configuration. In this case, you may opt for either a read/write balanced configuration or a more purely read-oriented configuration.

For most implementations, a configuration that balances the demands of disk read and write activities is the best choice.

- RAID 5/6. For some implementations, disk read speed is paramount and write speed is much less important to performance. For example, suppose the baseline index is never modified by partial updates, and new baseline indexes are moved into production only infrequently. In these implementations, a RAID 5 (or RAID 6) configuration improves availability with the least cost in spindles.
- RAID 10 (also known as RAID 1+0) is an excellent choice for devices that are partitioned across a disk array of four or more spindles. RAID 10 provides the performance benefits of striping and the redundancy of mirroring.
- RAID 0. The RAID 0 configuration is useful when storage availability after disk failure is not a concern. This is because both read and write activities are parallelized across all available spindles to decrease access latency and increase read and write throughput.

In any RAID configuration, high rotational speeds (such as 15k RPM or 10k RPM) are very beneficial to performance. Performance-oriented RAID controller features, such as battery-backed write caching, or a large cache size within the RAID controller, are also very beneficial to performance.

SAN-backed network-attached storage

As an alternative to using RAID disks, you can also use SAN-backed storage with a Fibre Channel backplane network from the MDEX Engine server to the SAN.

A storage area network (SAN) is a network to which remote storage devices are attached, usually accessible by a single machine in a one-to-one relationship. The storage devices appear to the operating system as locally attached to the server, as opposed to network- attached disks.



Note: To enable a successful Endeca implementation, ensure that the SAN is properly configured. It is also preferable that the MDEX Engine has dedicated access to its own SAN disk arrays.

In Endeca implementations, a SAN is in many cases faster and easier to work with than local storage. SAN-backed storage provides the following benefits:

- Faster promotion of index images from staging to production
- Faster backup of index images in production
- Faster copying of data from staging to production server
- Simpler backups of Endeca index files due to built-in functions for backups and snapshots in SAN



Note: Network-attached storage with NFS is known to cause performance issues and is not recommended in Endeca implementations.

Memory considerations

This section discusses the relationship between the amount of RAM, the Dgraph process virtual memory usage, the Dgraph cache, the working set size (WSS), and the resident set size (RSS) for the Dgraph process and their impact on performance.

In general, storing information on disk, instead of in memory, increases disk activity, which slows down the server. Although all the information the MDEX Engine may need is stored on disk, the running MDEX Engine aims to store in memory as many of its structures it currently needs as possible.

The decisions on what to keep in memory at any given time are based on which parts of the Dgraph are most frequently used. This affects the resident set size and the working set size of the running Dgraph, which, as they increase, lead to the increase of RAM being consumed.

Related Links

[Dgraph virtual memory vs. RAM: use cases](#) on page 17

While the amount of virtual memory consumed by the Dgraph process may grow and even exceed RAM at times, it is important for performance reasons that the working set size of the Dgraph process does not exceed RAM.

About Dgraph process memory usage

The Dgraph performs best when the working set of its process fits in RAM without swapping memory pages to disk.

The working set of the Dgraph process is a collection of pages in the virtual address space of the process that is resident in physical memory. The pages in the working set have been most recently and frequently referenced. In other words, the Dgraph working set is the amount of memory a Dgraph process is consuming now. This is the amount of memory that is needed to avoid paging.

In general, depending on the query load, the virtual memory process size of the Dgraph fluctuates. In some cases, it can exceed physical memory to a degree without affecting performance.

The section “*Dgraph virtual memory vs. RAM: use cases*” illustrates these statements.

Many factors affect the amount of memory needed by the Dgraph process. The number of records in the source data and their complexity are the most obvious factors, but the use of almost any feature will cause some increase in RAM usage.

The amount of memory needed for the Dgraph process also depends on other aspects of the query mix, such as which of the items that typically constitute Guided Navigation are being used and requested (records, dimensions, refinements, or other), and their particular usage in the query mix.

Memory usage recommendations for optimizing performance

Use the following recommendations to measure memory and optimize its usage for best performance.

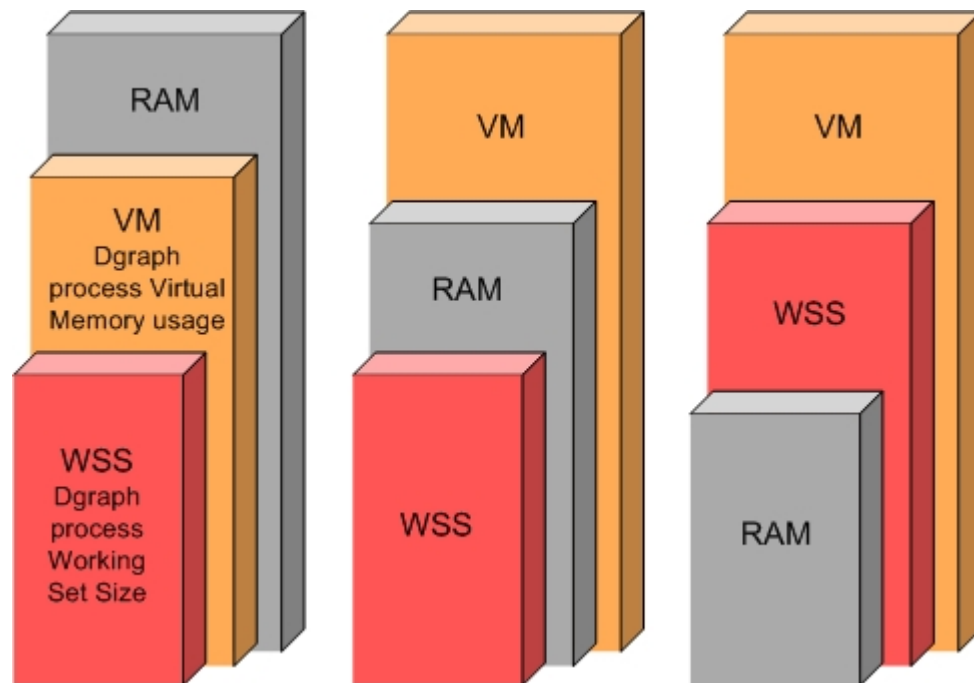
- Periodically measure the virtual memory process size of the MDEX Engine and its resident set size. The goal for these tests is to check whether the working set size (WSS) of the MDEX Engine starts to significantly exceed physical memory (it may exceed physical memory to a degree). The WSS cannot be computed, although it is always less than or equal to the amount of virtual process size for the MDEX Engine.
- Determine the WSS experimentally: if you notice that increasing RSS (by adding RAM or subtracting competing processes) improves performance of the MDEX Engine, this means that the WSS was previously larger than the RSS. This was likely the cause of the performance degradation.
- If the size of the WSS grows too close to the amount of RAM, or starts to exceed it, paging to disk begins and you will notice rapid decreases in performance.

The most noticeable symptom of paging is a large increase in Dgraph query latency. You can directly measure the amount of paging by using tools. For a list of some commonly used tools, see “*Useful Third-Party Tools*” in this guide.

Dgraph virtual memory vs. RAM: use cases

While the amount of virtual memory consumed by the Dgraph process may grow and even exceed RAM at times, it is important for performance reasons that the working set size of the Dgraph process does not exceed RAM.

The following diagram illustrates this relationship:



In this diagram:

- RAM is the amount of physical memory
- VM is the Dgraph process virtual memory usage
- WSS is the Dgraph process working set size

The diagram illustrates three distinct use cases:

- **Typical operation with normal memory saturation.** The graph on the left side illustrates the case where the amount of virtual memory used by the Dgraph process completely fits into RAM and thus the working set size of the Dgraph process also fits into RAM. This is a standard situation under which the Dgraph maintains its optimal performance.
- **Typical operation in an out-of-memory situation.** The graph in the middle illustrates the case where, while the amount of virtual memory exceeds RAM, the working set size of the Dgraph process fits into RAM. In this case, the Dgraph also maintains its optimal performance.
- **Potentially I/O bound operation with poor performance where WSS starts to exceed RAM.** The graph on the right side illustrates a situation that you should avoid. In this case, both the amount of virtual memory consumed by the Dgraph and the working set size of the Dgraph exceed RAM. Two situations are possible in this scenario that are of particular interest to you: the WSS can start to exceed RAM mildly or significantly. Subsequently, the degradation in I/O performance can also be mild or significant. Identify the level of I/O performance that is acceptable to your implementation. Depending on the acceptable I/O performance, you can decide whether you need to address the situation with WSS exceeding RAM. In general, if WSS starts to considerably exceed RAM, this causes Dgraph performance to drop dramatically.

Solutions for memory-based Dgraph performance problems

Use the following tips when addressing paging or out-of-memory situations with the Dgraph process.

- Add more RAM to the server hosting a single Dgraph or multiple Dgraphs. This is the simplest solution to paging issues with the Dgraph. If multiple Dgraphs are sharing a machine, you can

spread them out over a larger number of machines, thus giving each Dgraph a larger share of RAM. This solution has limits based on your hardware capabilities.

In addition, you can take a conservative approach, and add additional RAM in cases where the Dgraph memory consumption (WSS) approaches the amount of RAM available for the Dgraph, but does not exceed it yet. In such cases, while additional RAM may not be necessary to create an environment free of I/O contention, it provides a buffer and ensures that memory is available when needed.

- Defragment the file system periodically. Note that this may alleviate some performance problems.
- Consider tuning the `read_ahead_kb` kernel parameter on Linux. For example, a large data scale implementation that is operating out of memory can be a candidate for tuning this parameter.
- Explore how you use features such as wildcard search, multi-assign for dimensions, and others.

Related Links

[Dgraph Analysis and Tuning](#) on page 51

This section describes Dgraph performance tuning tips feature by feature. Features are not presented in order of severity of system impact.

[Memory usage recommendations for optimizing performance](#) on page 17

Use the following recommendations to measure memory and optimize its usage for best performance.

[Solutions for memory-based Dgraph performance problems](#) on page 18

Use the following tips when addressing paging or out-of-memory situations with the Dgraph process.

[Cache-tuning recommendations for optimizing performance](#) on page 20

There is no hard and fast rule for how to best allocate memory between internal Dgraph cache and FS cache, if you do not have enough memory to maximize both. For example, if you are operating at large data scale, it is likely that you will not have enough memory to maximize both the FS cache and the Dgraph cache. In practice it is easier to determine the right answer experimentally.

[Tuning the `read\_ahead\_kb` kernel parameter](#) on page 28

Oracle recommends setting the `read_ahead_kb` kernel parameter to 64 kilobytes on all Linux machines (RHEL 5). This setting controls how much extra data the operating system reads from disk when performing I/O operations.

About the Dgraph cache

The MDEX Engine cache (or the Dgraph cache) is a storage area in memory that the Dgraph uses to dynamically save potentially useful data structures, such as partial and complete results of processing queries.

Since the Dgraph has direct access to the structures it needs, it does not need to repeat the computational work previously done. The structures that are chosen for storing enable the Dgraph to answer queries faster by using fewer server resources.

The Dgraph cache is unified and adaptive:

- The cache is unified in that all sorts of data structures go into the same cache. The whole Dgraph has one cache, and all the threads share it.
- The cache is adaptive in that it uses a dynamic mechanism for evicting those data structures that it finds no longer useful. Its eviction algorithm attaches value to each cache object based on the empirical evidence of how useful the object actually is to the Engine, based on your current data,

and responding to your visitors' current queries. When this information changes, the Dgraph cache detects the change and adjusts, but you do not have to retune it.

The default Dgraph cache size (specified by the `--cmem` flag) is 1024MB (1GB).

The Dgraph cache improves both throughput and latency by taking advantage of similarities between processed queries. When a query is processed, the Dgraph checks to see whether processing time can be saved by looking up the results of some or all of the query computation from an earlier query.

The Dgraph cache is used to dynamically cache query results as well as partial or intermediate results. For example, if you perform a text search query the result will be stored, if it was not already, in the cache. If you then refine the results by selecting a dimension value, your original text search query is augmented with a refinement. It is likely that the Dgraph can take advantage of the cached text search result from your original query and avoid recomputing that result. If the navigation refinement result is also in the cache, the Engine does not need to do that work either.

To a large extent, the contents of the Dgraph cache are self-adjusting: what information is saved there and how long it is kept is decided automatically.

However, when deploying a Dgraph you need to decide how much memory to allocate for the Dgraph cache.

Allocating more memory to the cache improves performance by increasing the amount of information that can be stored in it. Thus, this information does not have to be recomputed.

Your MDEX Engine is well-tuned only when the Dgraph cache and the file system cache are well-balanced; therefore you need to understand them both.

About the File System Cache

The file system (FS) cache is a mechanism that the operating system is using to speed up disk read and write operations.

FS caching is very beneficial for the MDEX Engine, and it is important to tune the file system cache and the Dgraph cache on the server that runs the Dgraph.

For example, consider read acceleration since it is the aspect of the FS cache that matters most when tuning the MDEX Engine for performance. The FS cache speeds up reads by holding recently accessed information in RAM (on the theory that your process will need this data again), and by proactively reading ahead beyond the area recently accessed, and holding that information in RAM too (on the theory that your process will likely ask for that data next).

Related Links

[Tuning the `read\_ahead\_kb` kernel parameter](#) on page 28

Oracle recommends setting the `read_ahead_kb` kernel parameter to 64 kilobytes on all Linux machines (RHEL 5). This setting controls how much extra data the operating system reads from disk when performing I/O operations.

Cache-tuning recommendations for optimizing performance

There is no hard and fast rule for how to best allocate memory between internal Dgraph cache and FS cache, if you do not have enough memory to maximize both. For example, if you are operating at large data scale, it is likely that you will not have enough memory to maximize both the FS cache and the Dgraph cache. In practice it is easier to determine the right answer experimentally.

Use the following practices for optimizing the Dgraph and the file system caches for best performance:

- Examine the Cache tab of the MDEX Engine Stats page, especially if you need to tune the cache. In particular, pay attention to these columns in the Cache tab:
 - "Number of rejections". Examining this column is useful if you want to see whether you need to increase the amount of disk space used for the MDEX cache. Counts greater than zero in the "Number of rejections" column indicate that the cache is undersized and you may want to increase it.
 - "Number of reinsertions". Examining this column is useful if you want to examine your queries for similarities and improve performance by considering the redesign of the front-end application. Large counts in the "Number of reinsertions" column indicate that simultaneous queries are computing the same values, and it may be possible to improve performance by sequencing queries, if the application design permits.
 - "Total reinsertion time". Examining this column is useful for quantifying the overall performance impact of queries that contribute to the "Number of reinsertions" column. This column represents the aggregated time that has been spent calculating identical results in parallel with other queries. This is the amount of compute time that potentially can be saved by sequencing queries in a re-design of the front-end application.
- Experiment and increase the size of the Dgraph cache as your hardware allows. However, do not set the Dgraph cache to use all the free memory available on your server, since you also want to allocate memory for the file system cache and query working memory.

Use the Dgraph - -cmem flag to experimentally tune the Dgraph cache. It specifies the size of the cache in megabytes of RAM, and is the major mechanism for tuning the Dgraph cache. By default, if - -cmem is not specified, the size of the cache is 1024MB (1GB) for the Dgraph.

If you want the Engine to perform better, and you have physical memory to spare, you can increase the cache size and see if it works. Once the MDEX Engine obtains extra memory for its cache, the cache algorithm identifies the best strategy for storing the most useful data structures and eviction of those structures that are less likely to be needed frequently.

- For a specific MDEX Engine on any server, find the optimal performance point experimentally:

Gradually increase the size of the Dgraph cache until it no longer improves performance. When you notice that performance stops improving and starts degrading, you will know you have gone too far.

Back off the Dgraph cache setting by a fair amount (such as 500MB). The right answer depends on both raw data size and some subtle characteristics of the workload (such as, how much disk-backed information the average query needs, and how similar or different queries are from each other).
- Review your query mix to see if it exhibits a high degree of similarity between queries (either because of a highly constrained user interface or a highly homogeneous user base). This is one of the cases where performance improvements from a larger Dgraph cache may not be noticeable. If all your queries are similar, a large Dgraph cache is unlikely to be valuable.
- Find the right balance between the Dgraph cache and the FS cache. When tuning the size of the Dgraph cache, ensure that you do not accidentally displace the amount of memory allocated to the FS cache.

In general, the Dgraph cache may contain a slightly larger number of objects useful to the Dgraph compared with the FS cache. This is often beneficial for the Dgraph performance. However, this causes a significant performance degradation when information that is not in the FS cache is needed. This is because real disk access (not just access to the FS cache reads from RAM) will be needed more often, and disk reads are quite slow relative to the reads from the FS cache.

- Be aware of the paging situation when you experimentally determine the best strategy for allocating RAM to the Dgraph internal cache and the file system cache.

If you increase the Dgraph cache size in large increments between experiments, you may create a configuration where the Dgraph process memory (including the Dgraph cache) does not fit into physical RAM. In this situation not only there is not enough room for the FS cache, but the Dgraph process starts paging and performance degrades significantly.

- As your hardware permits, experiment with increasing the FS cache, along with the Dgraph cache. In general, performance gains from using the FS cache vary depending on what processes you are running and what they are doing with the disk.

For information on the file system caching mechanism, refer to the online sources of information that are specific to your operating system and the file system that you use.

Warming the Dgraph cache after an update

You can improve the performance of a Dgraph by warming the internal Dgraph cache after an update has been processed. To warm the cache, you specify the `--warmupseconds <seconds>` flag to the Dgraph.

The `--warmupseconds` flag instructs the Dgraph to store a sample of client queries, and after an update is processed, the Dgraph internally re-runs a sample of those queries to warm the cache before processing external queries using the updated data. The warming queries run for the `<seconds>` value that you specify.

During the period of time that the cache is warming, the Dgraph continues to serve incoming queries against the current data set. The Dgraph begins using the updated data as soon as the warming period is over. In other words, the amount of time needed to process an update is increased by the number of seconds you specify for `--warmupseconds`.

In the Dgraph request log, all warming query URLs are tagged with the additional parameter `&log=warming` as a convenience to identify the log entries produced by warming queries.

The Dgraph cache and its impact on virtual process size

The amount of memory allocated to the Dgraph cache directly affects the virtual process size of the Dgraph. An example in this topic shows how to adjust the Dgraph cache.

Furthermore, since the cache is accessed frequently, the amount of virtual memory allocated to it affects the working set size of the Dgraph. This may cause virtual memory paging, which can adversely affect throughput and especially the maximum latency. Whether this is a problem depends on your deployment scenario.

Example: Adjusting the Dgraph cache

Consider a scenario where a single Dgraph runs on a machine with 8GB of physical memory:

- If the virtual process size of the Dgraph is 6GB with a default (1GB) Dgraph cache, and the machine is not being used for any other processes, it makes sense to experiment with increasing the Dgraph cache size to 2.5GB to improve performance. The resulting 8.5GB virtual process size will not cause undue memory pressure.
- If the virtual process size of the Dgraph is 9GB, this exceeds the amount of RAM (8GB) and creates significant memory pressure. However, it may still make sense to increase the Dgraph cache size above the default, if the increase is not aggressive. Although in such a situation, increasing the cache size further will slow down those queries that are not assisted by the Dgraph cache, that may be acceptable if the effect of speeding up queries by providing a larger cache is greater than the effect of slowing down queries by causing virtual memory paging.

To make the right trade-off in this situation, increase the cache size while watching throughput, average latency, and maximum latency. At some point you will see that throughput is improving but average latency has gotten worse. Whether you are willing to trade latency degradation for throughput improvement will depend on the specific performance numbers, on your application, and on the expectations of your users.

Estimating the MDEX Engine RAM requirements

This topic provides recommendations for estimating the requirements for physical memory for an Endeca 6.1.x system given the anticipated growth of your data set.

The size of the Dgraph process is impacted by:

- The size of the Dgraph index generations in memory
- the size of the precomputed sorts in memory (if precomputed sorts are used)
- the size of the Dgraph cache
- Other factors, such as the size of the in-flight data

Each of these areas is discussed below in a separate section.

Impact of the MDEX Engine cache on WSS

Use `--cmem` to identify (or change) the Dgraph cache, and take it into account when estimating the projected amount of RAM needed for the MDEX Engine operations in view of the projected growth of the data set.

Impact of partial updates on WSS

Partial updates can have a significant impact on RSS and WSS. The precise details of the Endeca generation merging strategy are complex and proprietary. However, the rough pattern of memory usage that you can expect to see from a Dgraph running with partial updates is as follows:

- Expect a jump in address space usage each time a partial update is applied. The size of the jump depends on the size of the update. Each partial update causes one or more index generation files to be created.
- When merges of partial update generations occur, the MDEX Engine allocates space for a new generation file and merges two or more existing generations into that new generation file. This allocation causes a spike in the address space usage. Since some of the merged operations may cancel each other (for example, adding a record in generation file N is canceled by the deletion of that record in generation file N+1), the new total generation size may be smaller after the partial merge.
- Additionally, when full generation merges occur, all existing generations are merged into a single new generation file. Since the new generation file is roughly the same size as the sum of all pre-existing generation files (minus any canceled operations), the WSS roughly doubles during this period.
- While a full merge may cause the WSS to increase significantly, the effects on WSS are muted by the paging behavior of the operating system. Based on Endeca's recommendations, it is unnecessary for an MDEX Engine server to have a quantity of RAM equal to twice the generation file sizes when partial or full merge is occurring.
- It is fairly easy to detect the occurrence of a full merge. Watch the generations directory, found in `<dgidx_output>/<dataset_prefix>_indexes/generations`, and notice when the number of generation files drops to 1.

- During this testing, push enough of partial updates through the system to trigger the full merge. This will provide you with a good enough estimate of how much RAM you need for handling partial updates.



Note: Beginning with version 6.1.4 of the MDEX Engine, you can set the partial updates merges to use a balanced or aggressive merge strategy. For details on the merge policy, see the *MDEX Engine Partial Updates Guide*.

Impact of sorting strategies on WSS

When measuring WSS, account for the sorting strategies used by the MDEX Engine. To ensure that you measure the full "eventual" WSS of the Dgraph in 6.1.x, include a wide range of queries in your testing logs, ensuring that a portion of your queries utilizes sorting strategies, including precomputed sorts.



Note: You can confirm whether your sorting queries utilize precomputed sort by checking whether any of your properties is configured in Developer Studio so that it can be used for record sort, or by checking the <RECORD_SORT_CONFIG> element in your application's XML configuration files. This element lists properties that are configured to use precomputed sort. Precomputed sort techniques may be used by the MDEX Engine in the default sort queries. Therefore, to verify whether any of your sorting queries use precomputed sort, you can check the **Index Preparation Tab** of the Stats page that contains **Precomputed Sorts** statistics. This metric displays how much time the Dgraph has spent computing sorts, including computing sorts and incremental sort updates.

Impact of in-flight data on WSS

In addition to the types of impact that are already listed in this topic, other factors, such as in-flight processing and data can have an effect on WSS. These factors cannot be measured directly, but you should be aware of their effect.

Recommendations for estimating projected RAM requirements



Important: Use the following recommendations with the understanding that estimating RSS and WSS depends to a large degree on the operating system processing, and is also highly dependent on the context of your Endeca implementation. The size of the Dgraph process is affected by several aspects, such as the size of the index in memory, cache, and other computations. These artifacts, in turn, depend on the features you are using, such as types of updates you run (partial or baseline), or whether the application relies on precomputed sorts. To summarize, while estimating requirements for physical memory, use these recommendations in the context of your own implementation, to account for variability in the RSS size due to these factors.

To estimate projected requirements for physical memory for an Endeca 6.1.x system, use the following recommendations:

- Measure RSS. Perform evaluations of your average resident set size for your indexes, and peak resident set size, while noting the record set size on disk. For example, you may find it useful to identify various ratios between average record size on disk, average resident set size of your indexes, and peak resident set size. For testing these numbers, employ tests with varying levels of request activity sent to the Dgraph. For example, send a considerable number of requests to the 6.1.x MDEX Engine with periodic cache flushes to force the Dgraph to go to memory or disk as needed to fulfill some of the requests (this is true if you replay request logs for your test).

- If your implementation uses partial updates, account for this fact in your MDEX Engine testing. Include in your tests large enough files that contain records which will be updated through partial updates. For more information, see the section in this topic on Impact of partial updates on RSS.
- Similarly, account for the size of the Dgraph cache, for sorting queries that utilize precomputed sorts, and for the size of in-flight data (see sections in this topic on each of these aspects of the RSS).
- Identify the ratio of the RSS to on-disk representation of the record set, and confirm that with different tests this ratio remains the same.
- Based on these evaluations, draw conclusions and identify the following numbers:
 - The average on-disk record set size and the largest on-disk record set size.
 - The peak resident set size observed with the current record set.



Note: If you are not using partial updates, this number could be roughly equivalent to the on-disk representation of the MDEX Engine data plus the size of the cache for each of your Dgraphs, the size of the in-flight processing and data, and the fact whether precomputed sort is being used. If you are using partial updates, see a section in this topic for their impact on WSS and RSS.

- Using these recommendations, you can identify the following numbers for the MDEX Engine 6.1.x:
 - The average on disk record size that is used for your number of records.
 - The peak resident set size (RSS) of the Dgraph.
 - The peak virtual memory usage.
- Predict the growth of the RSS that you will need. You can do so based on the projected growth of the on-disk representation of the data set and the numbers that you obtain for the peak resident size, peak virtual memory usage and their ratios to your data set size.

Once you predict the growth of the resident set size, you can estimate memory requirements for your Endeca implementation. This will make it possible to provision enough hardware to support the MDEX Engines with the projected data set growth.

Network considerations

Oracle recommends that you use 100Mbit or Gigabit Ethernet. Also, make sure that all NICs in your implementation use the same duplex setting. The full-duplex setting is highly recommended.

Dgidx performance recommendations

This topic provides information about performance considerations for Dgidx.

RAM and disk swap size recommendations

Although this guide deals with MDEX Engine performance, since the Dgidx program is involved in the indexing process, it is important to plan for adequate Dgidx performance as well. It is especially important to plan for Dgidx performance if you have a large data set.

Oracle recommends provisioning your hardware for running Dgidx using these estimates:

- Plan to run Dgidx with the provisioned amount of RAM that is equal to the size of the finished index size, that is the size of the data/dgidx_output directory after a successful Dgidx run.
- Increase the amount of swap space size to at least the amount of RAM provisioned on your system.

Troubleshooting tips for Dgidx

if a record takes longer than 60 seconds to process by Dgidx, Dgidx prints out a warning enabling you to identify and fix the record. This information can be useful to you if you need to identify a record with extremely large numbers of property assignments. This may occur as a result of an issue with the ETL process. After you identify the record, you can review it to decide whether all of its assignments are required by the application.

Operating system considerations

This section discusses various tuning changes on the Operating system level that you can perform on the server running the MDEX Engine to optimize its performance.

Windows 2008 performance considerations

If you experience poor performance on an Intel Xeon processor-based servers running Windows Server 2008, Oracle recommends changing the default BIOS setting for power management from "Dynamic" mode to "Static High Performance" mode.

The BIOS has a mode setting that controls the power regulator. In the default "Dynamic" mode, the system attempts to balance high performance with power savings. Setting the regulator to "Static High Performance" mode forces the system to always favor performance.

This issue has been observed only on some Xeon-based servers.

VMware performance considerations

This topic discusses performance expectations of MDEX Engine deployments on VMware (all supported versions) and provides recommendations for such deployments.

Virtualizing Endeca deployments on VMware is motivated by cost management reduction that is typically associated with server consolidation, as well as by human cost reduction associated with simplified server administration and maintenance.

Supported guest operating systems

See the "Supported operating systems" section of the *Endeca MDEX Engine Installation Guide* for supported guest operating systems.

Configuration guidelines

Oracle recommends using the following guidelines for MDEX Engine deployments on VMware:

- Configure four VCPUs on a virtual machine.
- Specify four threads for each Dgraph. Overall, the number of threads should not exceed the number of VCPUs.
- Allocate a single Dgraph per virtual machine. Endeca does not recommend running more than one MDEX Engine per virtual machine.

Performance expectations

Overall, for server-level performance, the average and sustained throughput decrease in a VM environment, while the latency and the warmup time increase.

If you consider deploying an MDEX Engine with the Dgraph that is configured with four threads and where the MDEX Engine is assumed to be utilized at full capacity, expect a 10-30% performance overhead with a VMware-based deployment compared with a non-VM deployment. The indexing performance is also expected to be in the range of 10-30% overhead above the non-VM deployment. In some deployments, depending on your hardware, storage and implementation strategy, performance overhead can be up to 50%.

These performance expectations manifest in the decrease in sustained throughput, increase in average latency, increase in the amount of time it takes the Dgraph to reach 80% of its expected level of throughput, and increase in the latency of the longest query (99% of queries perform better than this query).

Additional performance recommendations

Performance risk associated with virtualizing the MDEX Engine is directly related to the performance and scalability requirements of your application. While Oracle recommends virtualization, customers interested in virtualizing HPC (high-performance computing) applications should analyze the risk associated with such projects and seek IT support with strong virtualization skills and experience. Endeca believes that virtualization of the MDEX Engine on VMware is most appropriate at smaller data scale.

Oracle recommends the following practices to ensure adequate performance on VMware:

- Implement vendor best practices for tuning performance of network and storage in a VM environment. For example, be aware of the limitation of four virtual CPUs per virtual machine.
- Be aware of the virtualization performance tax. The performance overhead, or "tax", of virtualizing the MDEX Engine varies by data set and by performance metric. When a deployment is properly configured and sized, the performance overhead is generally about 10%-30%. Endeca expects that the virtualization performance tax will exceed the range of 10%-30% and may reach up to 50% in the following situations:
 - Improperly configured or improperly sized deployments. Adequate memory allocation is especially important. Plan for additional memory and storage requirements due to index replication.
 - Write-heavy workloads. In particular, the following Endeca configurations are susceptible: (1) deployments where Dgidx and Forge are used heavily, and (2) Dgraphs under extensive and sustained partial update load.
- Rely on a robust deployment architecture. Most of the initial performance problems associated with deploying VMware occur due to mis-configurations or inadequate system resources.
- The approach to disk storage can be a significant factor in performance. Both locally-attached storage and network-attached storage solutions are supported. To ensure adequate performance, pay special attention to testing and tuning the bandwidth and latency of your storage solution with VMware. Consult with the documentation for your storage manufacturer for information on tuning your storage configuration for VMware.
- Expect that lower throughput will lead to longer warmup periods.
- Plan for lower ratio of query threads to update threads for applications leveraging frequent partial updates. Frequent partial updates are recommended in such implementations because each Dgraph is limited to four threads by the virtual machine limit of four virtual CPUs. On non-VM platforms, a Dgraph can be configured with significantly more threads, improving the ratio of query threads to update threads during partial update processing.

Linux considerations

This section lists recommended tuning changes on RHEL 4 and RHEL 5 configurations for the MDEX Engine.

About the `read_ahead_kb` kernel parameter

Starting with the MDEX Engine version 6.0, the MDEX Engine takes advantage of the `readahead` function.

Readahead is a technique employed by the Linux kernel that can improve file reading performance. If the kernel assumes that a particular file is being read sequentially, it attempts to read subsequent blocks from the file into memory before the application requests them. Setting the `readahead` can speed up the system's throughput, since the reading application does not have to wait as long for its subsequent requests, since they are served from cache in RAM, not from disk. However, in some cases the `readahead` setting generates unnecessary I/O operations and occupies memory pages which are needed for some other purpose. Therefore, tuning `readahead` for best performance is recommended.

You can tune `readahead` for optimum performance based on the settings recommended by Endeca.

Tuning the `read_ahead_kb` kernel parameter

Oracle recommends setting the `read_ahead_kb` kernel parameter to 64 kilobytes on all Linux machines (RHEL 5). This setting controls how much extra data the operating system reads from disk when performing I/O operations.

Reducing this value from the default typically increases sustained throughput for the MDEX Engine while also increasing its warmup time. Warmup is defined as initial performance of the MDEX Engine after startup (throughput and query latency), until the sustained level of performance is reached. Therefore, if you decide to tune this parameter, choose a value to balance these concerns.

Reducing `read_ahead_kb` has a noticeable effect and increases throughput for the MDEX Engine only in cases where a large data set may not fit into the MDEX Engine memory.

In cases when the index fits into memory, reducing `read_ahead_kb` from its default has no noticeable effect on the MDEX Engine performance.

When operating the MDEX Engine on a large data set that is running out of memory, consider adding more memory in addition to tuning `read_ahead_kb` to improve performance.

Setting `read_ahead_kb` to 64 kilobytes is a reasonable choice for most applications running on Linux.

To tune the `read_ahead_kb` kernel parameter on RHEL 5:

Add a command to `/etc/rc.local` as root:

```
echo 64 > /sys/block/sda/queue/read_ahead_kb
```

where `sda` is the name of the disk device for the MDEX Engine, and 64 is the number of kilobytes for the new `read_ahead_kb` setting.

Changing the I/O scheduler on RHEL 5

Oracle recommends changing the default I/O scheduler that the Linux kernel uses from CFQ to DEADLINE.

This dramatically speeds up performance of Endeca applications with large data sets in cases where both the amount of physical memory available to the MDEX Engine and disk I/O are limited. This recommendation applies to Endeca implementations on both RAID disk arrays and individual disks.

To adjust the I/O scheduler on a device:

1. Add a command similar to the following to `/etc/rc.local` as root:

```
echo deadline > /sys/block/sda/queue/scheduler
```

where `sda` is the name of the block device where the Dgraph input resides on your system. This changes the scheduler to DEADLINE.

2. Use performance tools to validate the results.

Disabling the swap token timeout on RHEL 5

Oracle recommends disabling the swap token timeout by setting it to zero. The swap token is a mechanism in Linux that allows some processes to make progress when the total working set size of all processes exceeds the size of physical RAM.

In situations when only one process is active, and the virtual memory size of that process gets close to, or exceeds the size of the available RAM, enabling the swap token negatively affects performance. In the context of the Dgraph, this can happen if the physical server is dedicated exclusively to running the MDEX Engine, and the index size is close to, or exceeds the size of the available RAM.

Oracle recommends disabling the swap token for those MDEX Engine configurations running on Linux that serve large data sets and are memory- and disk-bound.

If you choose not to disable the swap token, and experience erratic Dgraph performance, you may wish to examine the system to determine whether the swap token is causing problems. The swap token can cause "direct steal" operations.

To measure "direct steal" operations, check the contents of `/proc/vmstat`, adding `pgsteal_dma32` and `pgsteal_normal` values and subtracting `kswapd_steal`.



Note: Oracle recommends that you disable the swap token explicitly for the MDEX Engine disk devices even though you can obtain a patch for the Linux kernel that disables it.

To disable the swap token timeout on RHEL 5:

As part of the boot process, add one of the following options to your `/etc/rc.local` file as root:

```
sysctl -w vm.swap_token_timeout=0
```

or

```
echo 0 > /proc/sys/vm/swap_token_timeout
```

Or, add `vm.swap_token_timeout = 0` to `/etc/sysctl.conf`.

Load balancer considerations

For all deployment architectures, Oracle recommends the following load balancing practices.

- Use load balancers with the MDEX Engine to increase throughput and ensure availability in the event of hardware failure. Oracle recommends including two hardware-based load-balancing

switches configured redundantly in your configuration. Having two load balancers ensures their availability in the event of a load balancer hardware failure.

- Use the "least connections" model as the best routing algorithm for balancing traffic to the Dgraphs. The "round robin" model can have negative consequences, especially when occasional long-running queries are possible and the site is operating near its maximum traffic load.
- Ensure that return traffic from Dgraphs to the client tier is directly transmitted, and does not pass back through the load balancer hardware.
- Use scripting for load balancers. For example, you can use `http://[host]:[port]/admin?op=ping` on the load balancer to check whether the Dgraph process is running on this port. If it is not running, the load balancer fails over to another port, and directs queries to the MDEX Engine that is currently available.

Load balancing and session affinity

In a load balancing situation, consider enabling session affinity on the application server that directs server requests to the load balanced Dgraphs.

Session affinity, also known as "sticky sessions", is the function of the load balancer that directs subsequent requests from each unique session to the same Dgraph in the load balancer pool. Implementing session affinity makes the utilization of the Dgraph cache more effective, which improves performance of Dgraph access and the application server.

To facilitate session affinity, your application code can call `ENEQuery.setQueryInfo()` to create an `ENEQueryInfo` object. In this object, you set query-specific information in name/value pairs (such as the session ID and query ID) for the MDEX Engine to log.

Alternatively, you can also set this information by calling `HttpENEConnection.addHTTPHeader()` and specifying a name/value pair.

In either approach, the Web application sends the name/value pairs to the MDEX Engine. However, the `setQueryInfo()` method adds the name/value pairs to the query object itself; while the `addHTTPHeader()` method adds the name/value pairs to the header of the HTTP GET request.

In cases where long URLs interact poorly with a load balancer, you may need to force a POST request. You can force a POST request by calling `HttpENEConnection.setMaxUrl()` and specifying an upper limit on the length of the URL. Any URLs longer than the specified value are sent to the MDEX Engine using a POST request. You can also call `setMaxUrl()` and specify a value of 0 to force a POST request for all queries regardless of URL length.

Remember that application code automatically sends a query using a POST if the URL becomes too long to send using a GET request. The `setMaxUrl()` provides a way to force the request type if necessary.

Session affinity increases the latency overhead of the load balancer. Therefore, Oracle recommends testing the load balanced environment for performance optimization. This helps to determine whether the benefit of increased leverage from the Dgraph cache exceeds the cost of increased latency in the load balancer.

High availability considerations

Oracle recommends the following practices to ensure high performance and high availability.

- Use the Dgraph in multithreaded mode and experiment with increasing the number of threads. By default, the Dgraph runs in multithreaded mode with the number of threads set to one. It can be configured to run with a larger number of threads.
- Protect your configuration from hardware failures:

Use redundant disk drives with RAID (RAID 0, RAID 0+1, RAID 5), or SAN.

Utilize device redundancy for servers, load balancer devices and routers.

Use a second data center (hot or cold standby) to protect from site failures, such as power and network outages or enterprise data center failures.

- Protect your configuration from software failures:

If you have not done this already, upgrade to 64-bit systems to avoid “out of address space” failures. (Starting with its version 6.*, MDEX Engine installations are only supported on 64-bit systems.)

Use respawning monitors to protect against unexpected fatal process errors.

Watch out for paging with process memory usage.

Periodically examine your application for slow queries, or massive responses (too many results returned not all of which may be needed by the users).

Chapter 3

Using Multithreaded Mode

This section discusses MDEX Engine performance in multithreaded mode.

About multithreaded mode

The MDEX Engine always runs in multithreaded mode with the default number of threads set to 1. The multithreaded mode cannot be disabled.

The MDEX Engine always starts with a pool of threads that you can control with the `--threads` flag. These threads include query processing and partial update processing threads and additional threads that support query and update processing.

Each thread acts like an independent MDEX Engine, processing client requests one at a time and performing other tasks that support these requests, such as sorting and background index merging. It is important that the threads share data, memory, and the server network port. Essentially, this allows a multithreaded MDEX Engine with N threads to appear as a single MDEX Engine process that can work on N queries at a time. Each of the independent threads can run on independent CPUs (or cores), allowing a single multithreaded MDEX Engine to make use of multi-processor hardware.

Multiple threads can also share a processor, especially a multi-core processor, allowing an MDEX Engine running on a single-processor host to remain responsive as long-running queries are handled.

Benefits of multithreaded MDEX Engine

The MDEX Engine normally runs in multithreaded mode with the default number of threads set to 1. For many applications, Oracle recommends running the MDEX Engine with the number of threads greater than 1. These applications have the following characteristics.

- **Large index files on disk.** Only one set of index files is required for the multithreaded MDEX Engine. Thus, in addition to reduced hardware costs, the multithreaded approach reduces the hardware hosting disk space required.
- **Long-running queries.** For applications that rely on commonly used MDEX Engine features, the vast majority of queries complete in a fraction of a second. This allows the MDEX Engine to remain responsive at all times. However, many applications make use of more advanced features (such as computing complex aggregate Analytics queries) and can encounter longer running queries. For such applications, multithreaded mode allows the MDEX Engine to remain responsive while working on long-running queries.

- **Simplified system management and network architecture.** Enabling multithreading and tuning it is much simpler than adding new distinct servers that will run additional MDEX Engines, which includes reconfiguring the file system, adding load balancers and other infrastructure changes.
- **Applications with high throughput requirements with limited hardware resources.** The most efficient way to achieve simultaneous high throughput is to add MDEX Engines and run multiple MDEX Engines on distinct servers. But, when hardware resources are limited, running a multithreaded MDEX Engine on the same server requires fewer hardware resources than multiple distinct Engines, because all threads in the multithreaded MDEX Engine share resources.

The MDEX Engine relies on in-memory index structures to provide sub-second responses to complex queries. As the scale of application data increases, so does the memory required to host a single instance of the MDEX Engine.

Multithreaded execution mode enables more efficient utilization of RAM through SMP (Symmetric Multi-Processing) configurations. For example, if your current data scale requires 4GB of RAM, and query throughput requires four CPUs, multithreaded execution allows the site to be hosted on a single quad-processor machine with 5-6GB of RAM, rather than using more costly options, such as four single-processor machines, each with 4GB of RAM, or a 16GB machine with four Dgraphs on it.

- **Applications that heavily use the MDEX Engine dynamic cache.** Such applications cause a multithreaded MDEX Engine (with threads greater than 1) to perform better than multiple singlethreaded MDEX Engines because all threads in a multithreaded Engine share the same dynamic cache. This is especially true when that cache is cleared frequently due to restarts or partial updates, or when the cache is typically under heavy eviction pressure.

The MDEX Engine threading pool

The MDEX Engine consistently manages all processor-intensive tasks related to query processing using its preconfigured threading pool.

The `--threads` flag reflects the total number of threads in the MDEX Engine threading pool.

You define the number of threads in the threading pool at MDEX Engine startup, based on the setting for the `--threads` flag.

Recall that the recommended number of threads for the MDEX Engine is typically equal to the number of cores on the MDEX Engine server. By managing the threading pool, the MDEX Engine lets you more accurately limit the available computation resources to each core. This ensures that the system resources are used effectively for the highly prioritized tasks in the MDEX Engine all of which support query processing and high performance.

The threading pool manages the following MDEX Engine tasks:

- Query processing tasks
- Update and administrative operations
- All tasks that support query processing in the MDEX Engine. The MDEX Engine allocates these tasks for threads in the threading pool. The tasks include all high-priority, CPU-intensive, frequently performed operations the MDEX Engine runs in production. For example, they include precomputed sorting, background merging of index generations, and operations that support high performance of updates, among others.

Other MDEX Engine operations that do not have a significant impact on CPU usage are not managed by the threading pool.



Note: If you use operating system commands such as `top` to examine the number of threads used by the MDEX Engine server, you may see a number that is larger than the number you specify with the `--threads` flag. This is because in addition to this number of threads, the MDEX Engine may use additional threads for other tasks. These additional threads support tasks that are run infrequently, are less-CPU intensive, and do not affect overall MDEX Engine performance. You cannot control these additional threads.

Configuring the number of MDEX Engine threads

For most applications, Oracle recommends experimenting and increasing the number of threads.

By default, the MDEX Engine runs in multithreaded mode with the number of threads set to 1.

To increase the number of threads:

Specify it for the `--threads` flag when starting the MDEX Engine (Dgraph).
For example: `--threads 4`

This starts the MDEX Engine in multithreaded mode with four threads that are used for query processing and other MDEX Engine tasks that support query processing.

When to increase the number of threads

Oracle recommends using a higher setting for threads than in previous releases. Increasing the number of threads allows the MDEX Engine to handle more queries simultaneously.

Use the following recommendations:

- If you are using an application with a low throughput without long-running queries, this implementation can run in a singlethreaded mode in which one thread is used to process all query requests to the MDEX Engine. The same thread is used for other query-related processes of the MDEX Engine.
- If you are using a single MDEX Engine server with one thread, it is worth increasing the number of threads to improve performance.

A simple recommendation is to configure at least one thread per core. Higher ratios may generate more throughput, but due to the potential impact on latencies, Oracle recommends running further testing to find the thread count most beneficial to the needs of a specific application.

If increasing the number of threads stops improving query performance, this is an inflection point at which you can start considering the need to switch to a configuration with more Dgraphs.

A typical estimate that you can use to start testing with the increased number of threads is about 1 thread per core. For example:

- On a standard processor, enable 1 thread per processor
- On a dual-core processor, enable 2 threads per processor
- On a quad-core processor, enable 4 threads per processor

Multithreaded MDEX Engine performance

The performance of an MDEX Engine process is a function of a number of factors.

These factors include:

- Base, single-threaded performance, given the application data and query profile
- Number of processors on the host system
- Query characteristics
- Host operating system

Generally, on a host system with N CPUs or cores, where one single-threaded MDEX Engine can serve K operations/seconds of query load, N or more independent MDEX Engine processes will serve somewhat less than N times K, commonly in the 80-90% utilization range. In other words, given the base single-instance performance of K, the expected N-processor performance is given by

$$U \times K \times N \text{ where } (0.8 \leq U \leq 0.9)$$

The expected performance for one multithreaded MDEX Engine with more than one thread is similar, but generally somewhat less. In this case, the expected performance is given by the above formula, except with utilization in the 65% to 85% range ($0.65 \leq U \leq 0.85$). However, less RAM is required for running one multithreaded MDEX Engine with threads more than one compared with running separate single-threaded MDEX Engines.

For example, if one single-threaded MDEX Engine provides 20 ops/sec on a given load, running two MDEX Engines on a dual processor may provide around 36 ops/sec (U=90%, K=20, N=2). Running the same application with an MDEX Engine with threads more than one may provide 32 ops/sec (U=80%, K=20, N=2).

Similarly, if a single MDEX Engine requires 16GB of RAM, two Engines will require 32GB. Whereas a single MDEX Engine with more than one processing thread will only require slightly more than 16GB of RAM.

To summarize, Oracle recommends that you run a single MDEX Engine with the number of threads set to more than one, as opposed to multiple MDEX Engines. (Running multiple MDEX Engines introduces implementation complexity and also requires a load balancer.)

Recommended threading strategies and OS platform

The size of the thread pool and the host operating system impact performance and processor utilization.

In general, Oracle recommends using one thread per processor or core for good performance in most cases. The actual optimal number of threads for a given application depends on many factors, and is best determined through experimental performance measurements using expected query load on production data.

If high performance is required, enable more than one thread. Determine the optimal number of threads through load testing of different configurations.

As a starting point, enable the following number of threads:

- On a quad-core processor, enable 4 threads per processor
- On a hyperthreaded processor, enable 2 threads per processor
- On a standard processor, enable 1 thread per processor

For example, consider a server with two hyperthreaded processors and sufficient disk resources and RAM, on which a high-performance application will be deployed. The appropriate starting point for such an architecture would be one MDEX Engine running multithreaded with 4 threads.

Multithreaded MDEX Engine on Linux and Solaris

On Linux and Solaris, the MDEX Engine uses the POSIX Thread Library, `Pthreads`. You can examine the thread count using standard tools, such as `top`.

Multithreaded MDEX Engine on Windows

On Windows, the MDEX Engine uses native Windows threads. The thread count for an MDEX Engine can be examined in the Windows Task Manager in the Threads column.



Note: The number of threads listed may be greater than the value specified for the `--threads` flag; the additional threads that could be listed are those that are used infrequently by processes that are not CPU-intensive and represent internal maintenance tasks. All the CPU-intensive, query processing-related threads are controlled by the `--threads` flag.

Multithreaded MDEX Engine on VMware

On VMware, use the following configuration:

- Be aware of the limitation of four virtual CPUs per virtual machine.
- Specify four threads for each Dgraph. Overall, the number of threads should not exceed the number of VCPUs.

Chapter 4

Diagnosing Dgraph Problems

This section discusses techniques for determining the root cause of apparent poor MDEX Engine performance. It walks you through some example scenarios and points you in the appropriate direction, based upon the problems that may be present in your Endeca implementation.

Information you need

This section lists the information you need and the tools you can use to gather information in order to analyze and optimize the MDEX Engine performance.

Use the following sources of information:

- System state characteristics
- The MDEX Engine request log
- The Request Log Analyzer utility
- Eneperf

Sometimes poor application performance is the symptom of an operational problem (with the hardware, network, connections, or the application server). At other times, it may require you to review and revise the application coding, the Dgraph settings that were chosen previously and may need to be adjusted, or interactions between different features.

The first step in performance tuning is to find out what is causing the application to run more slowly than expected.

As you gather information about system performance, Oracle recommends that you note what steps you take and any changes you make to your environment, to ensure that you can analyze them or revert to your previous settings if needed.

When testing performance, make sure that the types of operations used to produce a load against the Dgraph are representative of an actual application usage scenario.

System state characteristics

The first clues to identifying the source of a performance problem are found in the system state. The following characteristics are easy to extract and may immediately indicate a direction in which to concentrate further investigation.

- The `Dgraph_input` directory.
- **Information about changes in the configuration.** This includes:

- Can the issue be replicated in the staging environment?
- Could the issue be caused by changes in network traffic or other network-related performance issues?
- Have there been any changes to the incoming data, pipeline or configuration files?
- **CPU utilization, disk I/O activity, and internal resource use.** This includes:
 - Physical number of CPUs available and the number of cores per CPU
 - The number of threads the Dgraph has been started with, and the total number of Dgraphs started on one machine
 - The type of disk I/O connection
 - CPU utilization statistics from the Dgraph host (especially when the performance problem is exhibited, if it is transient)
 - CPU utilization statistics from the front-end application host
 - disk I/O activity: processes other than the Dgraphs running on the machine that are not standard daemons or services (for example, a periodic backup process may interfere with disk access)
- **Memory utilization.** This includes:
 - Amount of allocated memory on the application server
 - Amount of physical memory (RAM) available on the Dgraph machine
 - Memory footprint of the Dgraph process. This includes the Dgraph cache (obtain it with `-c mem`), resident set size, and the amount of virtual memory available for the Dgraph process.
- **Storage capacity and configuration.** This includes:
 - Disk capacity in GB and disk rotation speed
 - Configuration and number of disks holding the index
 - Whether network-attached storage is used (SAN with Fibre Channel is recommended) versus local storage
 - Whether RAID configuration is used (the simultaneous use of two or more hard disk drives to achieve greater levels of performance)
 - If RAID is used, the configuration of the read-ahead policy for RAID. If the policy you have allows read-ahead, this lets the disk controller read additional data into the disk cache, which in turn increases the Dgraph performance.
 - Whether mirrored disks are used

This information defines the basic parameters for the performance problem. Typically, you base initial hypotheses on these findings, and confirm them with the next steps of the investigation.



Note: It is likely that you already have many of the tools you need to assess system state.

Related Links

[Useful Third-Party Tools](#) on page 157

This section lists some third-party tools that you may find useful during the Endeca performance monitoring process. The tools listed here are not supported by Endeca and are subject to change. In addition, these suggestions are not meant to overrule your choice of other tools.

Performance tools overview

You can use the following performance tools.

- The MDEX Engine Request Log

- The MDEX Engine Statistics page
- The MDEX Engine Auditing page
- The Request Log Analyzer
- Eneperf

The following sections describe these tools in detail.

Related Links

[The MDEX Engine Request Log](#) on page 77

This section describes the MDEX Engine (Dgraph) request log, which you can use to analyze Endeca application performance.

[MDEX Engine Statistics and Auditing](#) on page 147

The MDEX Engine Statistics page displays MDEX Engine (Dgraph) performance statistics. The MDEX Engine Auditing page tracks usage for licensing and performance purposes. This section describes these pages.

[Using the Eneperf Tool](#) on page 113

Eneperf is a performance testing tool that is included in your Endeca installation. This section describes how to use Eneperf.

The MDEX Engine request log

The MDEX Engine request log captures per-query metrics from a running Dgraph.

You can sort, filter, or otherwise manipulate the Dgraph request log to collect performance information. For example, you can sort the Dgraph request log based on query processing time to get the list of most expensive queries, or sort it on response duration to track latency trends.

The MDEX Engine Statistics page

The MDEX Engine Statistics page (also called the Dgraph Stats page) provides aggregated metrics since startup, and creates a detailed breakdown of what a running Dgraph is doing.

If performance is an issue, this page can help you to figure out which features are at fault.

Typically the feature in the Hot-spot Analysis section with the highest total is the best place to start your investigation. You can use the figures in the Dgraph Stats page to calculate useful metrics.

For example, to determine your application's network usage, you can multiply the number of ops/second by the average result page size.

The MDEX Engine Auditing page

The MDEX Engine Auditing page lets you view the aggregate MDEX Engine metrics over time and provides output of XML reports that track ongoing usage statistics.

These statistics persist through process restarts. This data can be used to verify compliance with licensing terms, and is also useful for tracking product usage. Each Dgraph in an implementation is audited separately.

The Request Log Analyzer

Use the Request Log Analyzer for processing request logs to analyze query load metrics for the MDEX Engine. The Request Log Analyzer reports actual performance, not the expected performance.

Use the Request Log Analyzer utility together with Eneperfto investigate whether you have performance under load.

Here are some of the ways you can use this utility:

- Isolate requests within a specific time range with the `--timelower` and `--timeupper` flags.
- Focus your attention on user-generated requests, by excluding admin, invalid, empty and error requests with the `--ignore` flag.
- Ensure that all statistics are logged. Request metrics in Endeca log reports do not correspond directly to query load metrics for the MDEX Engine. Differences in request metrics can arise from pages that issue multiple queries and from caching. For example, run the Request Log Analyzer with `--showAll` flag to ensure all statistics are logged:

```
reqloganalyzer --showAllGraph1.log > Graph1.stats
```

- Determine whether the performance bottleneck is caused by the Dgraph by comparing the statistics for “Engine-Only Processing Time” with “Round-Trip Response Time”.
- Show statistics based on threading with the `--showthreading` flag. This is useful when tuning your Dgraph threading configuration to increase the number of query threads.

Eneperft

Eneperft is a lightweight performance testing tool that is included in your Endeca installation. It makes Presentation API queries and XQuery-based queries against the Dgraph based on your Dgraph request logs and reveals how many operations per second the Dgraph responds with.

Dgraph performance issues

This section discusses locating and addressing Dgraph performance issues.

Improving the speed of Dgraph startup

Starting with the 6.1.x version of the MDEX Engine, Web services are loaded by default at startup. For this reason, Dgraph startup takes slightly longer than it did in the version 6.0.1. The Dgraph startup is typically faster than in Endeca IAP 5.1.

In most cases this increase in startup time is not an issue. However, if you find the startup time a problem and you are not planning to use Web services, you can turn off Web services and thus avoid the startup penalty. To do this, start the Dgraph with the `--disable_web_services` flag. (This flag is particularly useful during development, when you might be starting and stopping the Dgraph frequently.)

Tips for troubleshooting long processing time

You can use the Request Log Analyzer, installed with the MDEX Engine, to determine whether the performance bottleneck is caused by the Dgraph by comparing the statistics for “Engine-Only Processing Time” with “Round-Trip Response Time”.

If “Engine-Only Processing Time” as returned by the Request Log Analyzer tool is long, look further into specific query features to identify possible causes of the problem. This list identifies which problems you may want to isolate first:

- Is the long processing time for the Engine caused by limitations of hardware resources? Identify whether long query time is caused by CPU, memory, or disk I/O utilization.
- Is a high number of records being returned by the MDEX Engine? Identify how many records are being returned per query by looking for large nbins values in queries as reported by the Request Log Analyzer. This value indicates the maximum number of records that can be returned in the query. If this number is high, this can be expensive to compute and affects performance. Consider implementing paging control methods. For information on using paging control methods, see the *Endeca MDEX Engine Advanced Development Guide*.
- Are all dimension refinements (dimension values) exposed for navigation? That is, examine whether your queries are spending most of their time in refinement computation. Identify whether all dimension refinements are exposed by looking for allgroups=1 in the Dgraph request log (request URL parameter) or in Request Log Analyzer reports.

This setting corresponds to NavAllRefinements value of the ENEQuery method.

If the allgroups=1 setting is present in the URL parameter, review this configuration setting for your application to decide whether it is necessary. Exposing all refinements for navigation can decrease performance because the MDEX Engine has to examine each dimension value in the dimensions and determine whether or not that dimension value is a valid refinement given a current navigation state. Exposing all dimension refinements for navigation is not recommended.

For dimensions with many dimension values, Oracle recommends introducing a hierarchy (for example, a sift dimension hierarchy for automatically generated dimensions), so that the MDEX Engine has fewer dimension values to consider at one time.

- Are your longest queries similar? Check the longest queries for similarities, such as whether they all use the same search interface with relevance ranking, wildcard search, or record filters. See the sections in this guide about tuning performance of each of these features.
- Is record search being used? Identify whether a record search is being used by any queries by looking for "attrs=search_interface_name" in a query. This indicates that a record search is being used which means that possibly expensive relevance ranking modules can be contributing to high computation time.
- Which relevance ranking strategies are being used? Check the app_prefix.relrank_strategies.xml file for the presence of Exact, Phrase and Proximity ranking modules and test the same query with these modules removed.
- Is sorting enabled for properties or dimensions? Identify whether sorting with sort keys is enabled, for which properties and dimensions it is being used and whether it is needed. The first time a sort key is issued to a Dgraph after startup the key must be computed which can slow down performance. To isolate this problem, test the query in the staging environment by removing the sort key. If you confirm sort keys are the issue, consider using sort keys in a representative batch of queries used to warm up the Dgraph after startup. The sorts will become cached and these queries will be faster.



Note: Also, identify if sorting for properties and dimensions is necessary. In particular, it is not necessary to flag all sortable properties as sort keys in the project. This is often a performance problem itself.

Related Links

[The MDEX Engine Parameter Listing](#) on page 85

This section describes the parameters in the MDEX Engine request logs and provides mappings between the URL that is sent from the application to the Endeca Presentation API, and the URL that is sent from the API to the MDEX Engine.

[CPU recommendations for optimizing performance](#) on page 46

Use the following recommendations to optimize CPU performance.

[I/O recommendations for optimizing performance](#) on page 46

If you are testing the Dgraph maximum throughput using Eneperft with an adequate number of connections and the CPU is still not fully utilized, I/O could be a problem, especially if your application is search intensive but light on other features.

[Disk access recommendations for optimizing performance](#) on page 45

To optimize disk access performance, consider the following recommendations.

[Relevance ranking](#) on page 74

Relevance ranking can impose a significant computational cost in the context of affected search operations (that is, operations where relevance ranking is enabled).

Warming performance vs. steady state performance

When a Dgraph starts, its performance will gradually increase until it reaches a steady state. This process is known as Dgraph warming.

It is important to distinguish between the warming performance of the Dgraph and the steady state performance. Many of the techniques discussed in this guide address either one or the other, while others address both types of performance diagnostics and optimization.

The following considerations apply specifically to diagnosing and optimizing the warming performance of the Dgraph:

- Disk I/O problems can sometimes cause slow warming.
- It is helpful to run a Dgraph warming script at startup. For example, you can use a request log of characteristic queries played against the Dgraph to help warm it to a steady state.

About planning for peak Dgraph load

It is important that you plan your capacity to handle peak load. Sustained load above the projected peak load results in requests being queued for a long time. The system cannot keep up, and as a result, site performance (in particular latency) degrades.

About tuning the number of threads

Standard system diagnostic tools can tell you how busy CPUs on the machine are. If performance is poor and the CPUs are not very busy, try to increase the number of threads.

By default, starting with the MDEX Engine version 6.0, the Dgraph is running in multithreaded mode, with the `--threads` setting set to 1.

If increasing the number of threads does not help, one of the following is happening:

- You are using too many threads in one process. This is unlikely unless you exceed four threads, in which case consider using multiple Dgraphs.
- You have an I/O problem.
- There is an underlying network problem that needs to be investigated.

Multithreaded Dgraphs on machines with multithreaded processors

Processors with multithreading is a feature that allows a single microprocessor to act like two or more separate processors to the operating system and the application programs that use it.

Hyperthreading is a feature of Intel® Xeon® processors, as well as of Pentium 4® processors that support this technology.

Similarly, SPARC® Chip Multithreading (CMT) processors provide the technology for processor multithreading.

If your machine features hyperthreading or CMT, adding threads to your Dgraph can improve peak throughput by up to 30% per processor.

Multiple Dgraphs on one machine vs. multithreaded Dgraphs

You can run more than one Dgraph on a single machine, add additional threads to a single Dgraph, or run several Dgraphs with several threads enabled for each. Depending on your application, one choice might be better than the other.

The following use cases describe these choices:

- In most cases, the following recommendation applies: Dgraphs with a large memory footprint, especially in search-intensive applications, should be run in multithreaded mode with the number of threads greater than one for best performance.

For example, suppose you have a four-processor 16GB machine and a 3GB Dgraph. You could run four identical separate Dgraphs. A better alternative is to run one four-threaded Dgraph and thus reap the benefits of having more disk cache.

By running with more than one thread, I/O and computation can be overlapped. Although the time to process an individual request isn't improved (and can actually increase slightly due to contention for shared resources), overall throughput is significantly boosted.

- Likewise, in many cases it is appropriate to run two or more Dgraphs on one machine, each with several threads. Two four-threaded Dgraphs on one machine is an especially common configuration. The trade-off between thread contention and memory depends on the memory footprint that you estimate is needed for each Dgraph and the amount of memory available on the machine that will host multiple Dgraphs.

Disk access recommendations for optimizing performance

To optimize disk access performance, consider the following recommendations.

- Use a dedicated storage device with low latency and high IO ops/sec for all your Endeca indices and files. Locally-attached storage with a RAID controller is preferred. Only in cases where that is not possible, SAN using a Fibre Channel will typically provide strong performance assuming it has been configured correctly.
- If you are using an array controller, Oracle recommends using a striped disk configuration, such as RAID 5/6 or RAID 0+1 that enable you to avoid having redundant disks but ensures fault tolerance.
- Do not use disks with NFS, or other file system protocols. They are known to slow down performance.
- Ensure that the log files are saved locally. Turning off verbose mode, which prints information about each request to `stdout`, can sometimes help performance.
- Ensure that you have a fast disk subsystem and plenty of memory available for disk cache managed by the operating system, since the Dgraph keeps its various text search indices on disk, including search and navigation indexes.

CPU recommendations for optimizing performance

Use the following recommendations to optimize CPU performance.

- If the CPU is under-utilized, increase the number of threads for the Dgraph.
- If the CPU is over-utilized and you are not satisfied with throughput, investigate which activities make it busy. Add machines or make the queries less taxing by tuning individual features.

Related Links

[Dgraph Analysis and Tuning](#) on page 51

This section describes Dgraph performance tuning tips feature by feature. Features are not presented in order of severity of system impact.

I/O recommendations for optimizing performance

If you are testing the Dgraph maximum throughput using Eneperf with an adequate number of connections and the CPU is still not fully utilized, I/O could be a problem, especially if your application is search intensive but light on other features.

There is no absolute threshold that indicates that an application is I/O bound, but typical symptoms include very high numbers of I/O hits per second or KB per second. If I/O is below the specifications for the hardware, it is less likely to be a problem. In some cases, it is even possible to go beyond a device's theoretical maximum because of disk caching.

To determine the level of I/O activity, use the following tools:

- On Solaris, run `iostat -2`
- On Linux, run `sar -b`
- On Windows, do the following:

On the **Task Manager**, open the **Processes** tab.

From the menus, select **View > Select Columns**.

Check **I/O Reads**, **I/O Read Bytes**, **I/O Writes**, and **I/O Write Bytes**. These options enable new columns in the **Processes** pane that provide similar information to `sar -b` on UNIX.

Identifying problems with resource usage by the application

Use the following recommendations to identify performance problems associated with resource usage.

- Isolate performance testing for those parts of the application that specifically use the Endeca MDEX Engine from testing for other parts of the application. In other words, measure the performance of those parts of the application that use the Endeca MDEX Engine separately from the performance of those parts that use other software that may cause performance problems, such as a relational database. For example, if the latency is high, consider testing the interaction of the application with the database, if you are using one.
- If you are sending a lot of requests to the front-end application and performance is slow but the MDEX Engine servers are idle, the front-end application and its resource usage is probably the issue. There are two possible fixes: you can reduce consumption of resources by the application by reviewing your coding practices for the front-end application, or add resources.

Coding practices for the front-end application

Reviewing your front-end application code can help reduce resource usage performance issues that affect it. Review your Web application to check for any of the following problems.

- Creating or discarding objects unnecessarily.
- Excessive looping, particularly over properties that are not going to be displayed.
- Creating too many variables.

Web application ephemeral port contention

Each client/server connection has a unique identifier (known as a quad) that includes an ephemeral port number that will later be reassigned. Each operating system has a range of numbers that it uses as ephemeral ports (for example, on Windows the range is 1024 through 4999).

The operating system allocates ephemeral ports when a new socket is set up.

If the range is relatively small and you are making several requests per page in parallel, you can run out of port numbers. At that point the ephemeral port numbers assigned by the operating system start colliding with ones already in use as they are recycled too quickly, and subsequent connections will be aborted.

To address this problem, try one of the following:

- Reduce the two-minute time interval that the system waits between a connection close and port reassignment. The minimum recommended time is 30 seconds.
- Change the ephemeral port range. The method varies depending on your operating system; however, details are easily obtained on the Web.

Recommendations for identifying network problems

Often the diagnosis of slow performance comes from a query load played against the front-end application. The front-end application, or the configuration of its application server, may be the reason for the poor performance.

Alternatively, the network may be the problem, although this is less likely.

To identify whether the network is a performance issue:

- Compare Eneperfb performance on the local host and a remote host. First, run Eneperfb against the Dgraph on the Dgraph machine. Next, run the same Eneperfb against the same Dgraph, but from the front-end machine (if possible), or somewhere on the other side of the network. If the difference is negligible, the network is not a problem. If Eneperfb across the network is slow, you need to consider both the network itself and the application configuration.
- Alternatively, you can run the Request Log Analyzer and compare the “Round-Trip Response Time” with the “Engine-Only Processing Time”. If “Round-Trip Response Time” is long but the “Engine-Only Processing Time” is short, this can indicate a network problem or a configuration of an application server for the front-end application.
- Measure network performance using Netperf, a freely available tool that can be used to measure bandwidth. Alternatively, you can FTP some large files across the network link. If these tools show poor throughput across the network, this can indicate a network hardware problem such as a failing network interface card (NIC) or cable.

- In addition, check Eneperft statistics, the Dgraph request logs, or the Dgraph Stats page to see how much data is being transmitted back from the Dgraph on an average request. Large average result page size can saturate the network.

If it seems as if your application is trying to move too much data, it is likely that you may need to change the configuration of your application. To determine if changes are needed, consider the following:

- Is all of the data actually being used by the application? In other words, does the MDEX Engine return record fields that are then ignored by the front-end application? This is an especially serious problem with large documents.
- Is your application returning unnecessary fields with the Select feature? (This is described in “Controlling Record Values with the Select Feature” in the *MDEX Engine Advanced Development Guide*.)
- Is your application returning navigation pages that are too large? (Navigation pages are result list pages, as opposed to record detail pages.) If the application returns a lot of detailed information in the result list pages, consider reserving the details for a click-through and reducing the size of the result list pages your application returns on initial requests.
- Is your application returning large numbers of records without using the bulk record API? (This is described in “Bulk Export of Records” in the *MDEX Engine Advanced Development Guide*.)
- Is the network saturated? Upgrade to Gigabit Ethernet and identify the transmission speed being used. Ensure there is ample network bandwidth between the front-end application and the Dgraph. To identify Gigabit Ethernet transmission speeds, work with your network administrator.
- What is the configuration of NIC cards? Ensure that NIC duplex settings match between the Dgraph host and the web application client host and that both are set to full duplex. A mismatch can cause latency issues.
- Could large response sizes returned by the Dgraph be saturating the network? Use the Request Log Analyzer analysis to confirm large response sizes returned by the Dgraph, which can be caused by the query features you use. The way certain features are used can cause slow processing time and also saturate the network.
- Do you have queries waiting in the Dgraph queue to be processed? Check “Threading/Queuing Information” summary in the Request Log Analyzer for the number of items experiencing queue issues and the number of HTTP Error request 408 timeouts. Review the Dgraph setting for the number of worker threads and consider increasing it, if it is set to 1. Queuing can also be caused by spikes in traffic.
- Does the front-end application process the responses returned by the Dgraph quickly enough? Check CPU, memory, and disk I/O utilization on the front-end application server. Ensure the application server does not need to be tuned and that large responses are not being returned by the Dgraph.

Related Links

[Useful Third-Party Tools](#) on page 157

This section lists some third-party tools that you may find useful during the Endeca performance monitoring process. The tools listed here are not supported by Endeca and are subject to change. In addition, these suggestions are not meant to overrule your choice of other tools.

[Tips for troubleshooting long processing time](#) on page 42

You can use the Request Log Analyzer, installed with the MDEX Engine, to determine whether the performance bottleneck is caused by the Dgraph by comparing the statistics for “Engine-Only Processing Time” with “Round-Trip Response Time”.

Troubleshooting connection errors

This topic discusses how to debug connection errors with ENEQuery exceptions.

Problem - The application server does not seem to connect to the Endeca server. The Endeca reference application has no difficulty connecting. A connection to the port works as confirmed by JUnit tests. A problem exists connecting to the server once all the reference application libraries are packaged into the EAR file that is run inside the WebSphere application server.

Solution - In general, the `HttpENEConnection.query ENEQuery` method is used to issue a query against the Dgraph. In the `HttpENEConnection.query` method in the Java version of the Endeca Presentation API, any connections problems are raised as an `ENEQueryException`. (There is an equivalent in .NET version of the Endeca Presentation API).

To diagnose a connection problem from an application server to an Endeca server, the following assumptions are made:

- The Java version of the Endeca Presentation API is being used.
- The connection from the application server to the MDEX Engine is running on HTTP, not HTTPS.
- The application server and the MDEX Engine on the Endeca server are configured on separate machines.

To troubleshoot the connection problem, do the following:

1. Verify from the application server machine that you can connect to the port on the Endeca server. Using telnet on Windows or Unix can help you determine if you can successfully make a connection:

```
telnet <hostname> <dgraph port>
```

- a) If you cannot establish a connection with telnet, check that the Dgraph process is running with the specified port. Check the Dgraph `stderr` log to confirm the Dgraph was able to successfully bind to the port and another process is not using the port. You can also verify the Endeca server machine is listening on a socket with the specified port using `netstat -a`. Check that a valid network route exists from the application server to the Endeca server. You can also use `ping`. Also, use `tracert` on Windows, `tracert` on Linux, or `traceroute` on Solaris. If no valid network paths exist, check with your network administrator to eliminate possible problems with a firewall or routing configuration.
- b) If you can obtain a connection from telnet, verify that the application server can talk to the Endeca server. Write a Java program with a `static void main` method to make a connection to the MDEX Engine on the Endeca server. Make sure the Endeca Navigation JAR file is included in your classpath. If this program makes a connection successfully, the problem should only occur within the application server.
2. Write a utility JSP page that connects to the MDEX Engine on the Endeca application server and place it on the application server to verify the connection. Alternatively, you can run the Reference Application on the application server.
3. If everything works correctly, to troubleshoot further check the application server configuration. For Websphere, do the following:
 - a) Check all log files in `IBM/WebSphere/AppServer/profiles/AppSrv01/logs/server1`.
 - b) Verify that the Reference application is correctly packaged as EAR file.
 - c) Make sure Websphere deployed the EAR file and the application is running in the WAS admin console.

Assuming that you have WAS 6.1, go to **Security > Secure Administration, application and infrastructure** and check whether Java 2 security is enabled. If it is enabled, make sure your `was.policy` file is saved in the META-INF directory.

Next steps

Your hardware needs should be based on the number of ops/second revealed by Eneperf testing. If you feel that the resulting hardware requirements are too great, the next thing to do is identify costly features in your front-end application and see what you can do about them.

Modifications you can make to your Dgraph settings in order to improve the performance of your Endeca application are discussed in the next chapter.

Chapter 5

Dgraph Analysis and Tuning

This section describes Dgraph performance tuning tips feature by feature. Features are not presented in order of severity of system impact.

Feature performance overview

Once you have determined that the Dgraph is the bottleneck using the techniques described in this guide, there are many things you can do to tune performance. In many cases, unnecessary complexity slows performance, so small changes can yield big returns.

It is best to begin making adjustments with a conservative strategy that you understand well. Do not modify too many features at once—it makes it difficult to assess the impact of any one change.

Details on tuning specific features can be found in the following sections. Where applicable, they discuss problematic feature interactions. Likewise, each section indicates whether the kind of data you are processing (for example, large text fields as opposed to many part numbers) significantly impacts a feature's performance.

This chapter calls out only those aspects of a feature that affect application performance. For more general information about implementing these features, see the *Platform Services Forge Guide*, *MDEX Engine Basic Development Guide* and the *MDEX Engine Advanced Development Guide*.

Endeca record configuration

This section discusses the performance implications of some aspects of Endeca record configuration.

Record select

The Select feature prevents the transfer of unneeded properties and dimension values when they are not used by the front-end Web application.

It therefore makes the application more efficient because the unneeded data does not take up network bandwidth and memory on the application server. This may be relevant if your logs are showing large result pages.

You set the selection list on the `ENEQuery.setSelection()` method (Java), or the `ENEQuery.Selection` property (.NET).

Aggregated records

Aggregated Endeca records are not necessarily an expensive feature in the MDEX Engine. However, use them only when necessary, because they add organizational and implementation complexity to the application (particularly if the rollup key is different from the display information).

Using aggregated records slows down the performance of sorting and paging.

Note also that dynamic statistics on regular and aggregated records (controlled with the `--stat-abins` Dgraph flag) are expensive computations for the Endeca MDEX Engine. See the topic in this section for more details.

Derived properties on aggregated records

Some overhead is introduced to calculate derived properties on aggregated records. In most cases this should be negligible. However, large numbers of derived properties and, more importantly, aggregated records with many member records may degrade performance.

The number of records returned with an aggregated record and performance

You can use the `Np` parameter to specify the number of records to be returned with an aggregated records. For example, `Np=1` means that a single representative record is returned with each aggregate record, and `Np=2` brings back all records.

Utilizing `Np=2` may adversely affect your performance, as it causes the MDEX Engine to serialize more records for each query. The degree to which performance is affected is proportional to the number of base records for each aggregate record that is returned.

In most cases, it is not recommended to bring back all records in each query and aggregate all records with `Np=2` as this computation could be expensive for the MDEX Engine to serialize the result. However, `Np=2` can be useful in some cases. The impact on performance is proportional to the number of records that will be returned as aggregates.

For example, if each aggregate record contains only 2 records, the record serialization time is only twice the time as it is for `Np=1`. If, however, each aggregated record has 100 records associated with it, it is 100 times more expensive to perform the record serialization for `Np=2` than for `Np=1`.

Record serialization time is typically only a large portion of the query processing time in very low latency applications or with very large numbers of returned records.

Note also that in many cases, a 100-fold increase in record serialization time is barely noticeable. You can examine the `Prefetching horizontal records` statistics in the `Hotspot Analysis` section of the Stats page to determine whether their performance issue is due to returning many records.

For example, if you have a very small data set with queries served almost entirely from the cache, where most of the computation done by the Dgraph for each query consists of assembling the records to be returned, the negative effect on performance is reflected in the `Prefetching horizontal records` statistics being very large in this case which indicates that `Np=2` should not be used.

Dimensions and dimension values

This section discusses tuning features related to dimensions and dimension values.

Hidden dimensions

You prevent a dimension from appearing in the navigation controls by designating it as a hidden dimension.

Hidden dimensions, like regular dimensions, are composed of dimension values that allow the user to refine a set of records. The difference between regular dimensions and hidden dimensions is that regular dimensions are returned for both navigation and record queries, while hidden dimensions are only returned for record queries and dimension search.

In cases where certain dimensions in an application are composed of many values, marking such dimensions as hidden improves Dgraph performance to the extent that queries on large dimensions are limited, reducing the processing cycles and amount of data the Dgraph must return.

Dimensions and dimension values with high record coverage

Consider a case where records have dimensions that have almost—but not quite—full coverage over the records. For example, 99% of the records have a dimension value for a Location dimension, but the remaining 1% do not.

While this factor does not affect performance significantly, you can add an “n/a” dimension value to fill the gap and make the dimension have 100% coverage, if you want to let users explicitly refine to records that do not have an assignment for that dimension.

Flat dimension hierarchy

In general, avoid using large, flat dimensions (that is, dimensions with thousands of dimension values at the same level of hierarchy).

This is doubly true if statistics are enabled for those dimensions. It is better to design dimensions that contain sensible levels of hierarchy.

For some applications with extremely large, non-hierarchical dimensions, larger values for `--esampmin` can meaningfully improve dynamic refinement ranking quality with minor performance cost.

Displaying multiselect dimensions

When making decisions about whether to configure a dimension as multiselect, keep in mind that users may take longer to refine the list of results, because the user can continue to refine a multiselect dimension until all leaf dimensions have been selected.

In particular, refinements for dimensions tagged as multiselect OR are expensive.

Multi-assign dimensions

A dimension is considered to be multi-assign if there exists a record which has more than one dimension value assigned to it from that dimension.

Making a dimension multi-assign can slow down refinement computation. To improve performance, you can use multi-assign only for those dimensions for which you need it, and avoid making dimensions multi-assign where it is not useful.

Displaying refinement dimension values

Run-time performance of the MDEX Engine is sometimes directly related to the number of refinement dimension values being computed for display. If any refinement dimension values are being computed by the MDEX Engine but not being displayed by the application, use the `Ne` parameter more strictly.

The worst-case scenario for run-time performance is having a data set with a large number of dimensions, each dimension containing a large number of refinement dimension values, and setting the `ENEQuery.setNavAllRefinements()` method (Java), or `ENEQuery.NavAllRefinements()` property (.NET) to `true`. This combination is slow to compute and creates a page with an overwhelming number of refinement choices for the user. Endeca does not recommend using this strategy.

In general, you may want to reconsider the number of refinements you display, as well as consider implementing precedence rules.

Related Links

[Precedence rules](#) on page 73

This section discusses precedence rules and explains their performance impact.

Dynamic statistics on dimension values

You should only enable a dimension for dynamic statistics if you intend to use the statistics in your Endeca-enabled Web application. Because the Dgraph performs additional computation for the statistics, there is a performance cost to enabling statistics that your application does not use.

Using dynamic refinement ranking can greatly speed up refinement computation by displaying only the top refinements for a dimension, rather than computing the exhaustive list of refinements.

To decide whether or not dynamic refinement count statistics are likely to be appropriate for a project, consider the following aspects of your configuration:

- The number of dimension value refinements per page, especially dimension values assigned to large numbers of records. The more refinements are returned on each page, the more counts that need to be computed, and the bigger the performance impact.

For example, if the data set has a large number of dimensions, and/or the application uses `ENEQuery.setNavAllRefinements(true)`, then the performance impact will be larger. This is especially true if many of the dimension values are assigned to large numbers of records. This frequently happens with hierarchical dimensions. For example, it is more expensive to count Red Wines than it is to count Merlots.

- The number of records in the data set. Data sets with large numbers of records will see a proportionally higher performance impact from record count statistics.
- The average number of results per query. Applications that tend to perform searches that match larger numbers of records will see proportionally higher impact from refinement count statistics.

As a simple rule, add up the counts for all of the refinements on the page. The performance impact of record count statistics grows proportionally with that sum over all refinements. All of the above considerations are aspects of the application that can make that sum larger, and increase your performance slowdown related to record counts.

You can speed up computation of dynamic statistics for refinements by doing the following:

- Set the following options in the STATS subelement in the `refinement_config.xml` file:
 - `RECORD_COUNT_DISABLE_THRESHOLD` specifies the maximum number of records in a result set above which the MDEX Engine does not compute or return any dynamic statistics for that query. This speeds up processing if you do not need the counts in this case.

- `MAX_RECORDS_COUNT` causes the MDEX Engine to stop computing dynamic statistics for a particular dimension value when it has reached the specified value. The count returned in this case is the minimum of the actual count and `MAX_RECORDS_COUNT`. Thus, you can set this parameter to a specific value if you do not need to know the count for a particular dimension value once it is sufficiently high.

Aggregated refinement counts

Dynamic statistics on regular and aggregated records are expensive computations for the Endeca MDEX Engine.

You should only enable a dimension for dynamic statistics if you intend to use the statistics in your Endeca-enabled Web application.

Similarly, you should only use the `--stat-abins` flag with the Dgraph to calculate aggregated record counts if you intend to use the statistics in your Endeca-enabled Web application. Because the Dgraph does additional computation for additional statistics, there is a performance cost for those that you are not using.

In applications where record counts or aggregated record counts are not used, these lookups are unnecessary. The MDEX Engine takes more time to return navigation objects for which the number of dimension values per record is high.

The `--stat-abins` flag for the Dgraph lets you calculate aggregated record counts beneath a given refinement. For more information on using this flag, see the *Endeca Basic Development Guide*.

Dynamic refinement ranking and performance

You can use `--esampmin` with the Dgraph, to specify the minimum number of records to sample during refinement computation. The default is 0.

For most applications, larger values reduce performance without improving dynamic refinement ranking quality. For some applications with extremely large, non-hierarchical dimensions (if they cannot be avoided), larger values for `--esampmin` can meaningfully improve dynamic refinement ranking quality with minor performance cost.

Disabled refinements

Performance impact from displaying disabled refinements falls into three categories. They are discussed in the order of importance.

- The cost of computation involved in determining the base and default navigation states.

The base and default navigation states are computed based on the top-level filters that may belong to these states. These filters are text searches, range, EQL and record filters and selections from dimensions. The types and numbers of these top-level filters in the base and default navigation states affect the MDEX Engine processing involved in computing the default navigation state. The more filters exist in the current navigation state, the more expensive is the task; some filters, such as EQL, are more expensive to take into account than others.

- The trade off between using dynamic refinement ranking and disabled refinements.

In general, these two features pursue the opposite goals in the user interface — dynamic ranking allows you to intelligently return less information to the users based on most popular dimension

values, whereas disabled refinements let you return more information to the users based on those refinements that are not available in the current navigation state but would have been available if some of the selections were not made by the users.

Therefore, carefully consider your choices for the user interface of your front-end application and decide for which of your refinements you would like to have one of these user experiences:

- Dynamically ranked refinements
- Disabled refinements

If, for example, for some dimensions you want to have only the most popular dimension values returned, you need dynamic ranking for those refinements. For it, you set the sampling size of records (with `--esampin`), which directly affects performance: the smaller the sampling, the quicker the computation. However, for those dimensions, the MDEX Engine then does not compute (and therefore, does not return) disabled refinements.

If, on the other hand, in your user experience you would like to show grayed out (disabled) refinements, and your performance allows it, you can decide to enable them, instead of dynamic ranking for those dimensions. This means that for those dimensions, you need to disable dynamic ranking. As a side effect, this involves a performance cost, since computing refinements without dynamic ranking is more expensive. In addition, with dynamic ranking disabled, the MDEX Engine will need to compute refinement counts for more dimension values.

- The cost of navigation queries.

Disabled refinements computation slightly increases the navigation portion of your query processing. This increase is roughly proportional to the number of dimensions for which you request the MDEX Engine to return disabled refinements.

Displaying dimension value properties

Dimension value properties (that is, key-value pairs that the Dgraph passes back along with a dimension value) could slightly increase the processing or querying time because additional data is moved through the system, but this effect is generally minimal.

If your Endeca application does complex formatting on the properties, this could slow down page loads. If the properties are used to add formatting HTML or perform other trivial operations, they have minimal impact on performance.

Collapsible dimension values

Collapsible dimension values have a negative impact on performance.

Mapping source properties

Automatically mapping source properties is a feature that, while it can be used in the staging environment to facilitate testing, is not recommended for using in the production environment.

The Property Mapper in Developer Studio allows you to automatically map source properties to Endeca properties or Endeca dimensions, if no mapping is found. (This feature is also known as Automapper). The option of the Property Mapper that lets you map source properties to Endeca properties or dimensions defines the setting that Forge uses to handle source properties that have neither explicit nor implicit mappings.

Use this option with caution because each source property that is mapped uses system resources. Ideally, you should only map source properties that you intend to use in your implementation. Many production-level implementations automatically pull and process new data when it is available. If this data has new source properties, they will be mapped and included in your MDEX Engine indices, which uses system resources unnecessarily. As a result, the Forge output is larger, the indexer is larger and the MDEX Engine has additional indices to process.

Indexing all properties with Dgidx

The `--nostrictattrs` flag for Dgidx allows you to index every property found on a record, including those properties that do not have corresponding property mapper settings. Using this flag may negatively affect performance of Dgidx and the MDEX Engine.

If a large number of unused properties are sent to Dgidx, they will get indexed and will consume system resources during the indexing process and at run-time. These properties can also affect performance of the front-end application API, because the amount of information communicated between the MDEX Engine and the API increases.

Record sorting and filtering

This section discusses the performance impact of record sorting and filtering.

Sorting records by dimension or property

Enabling dimensions and properties for sorting increases the size of the Dgraph process and may negatively affect partial update latency. The specific size of the increase is related to the number of records included in the data set.

Therefore, in Developer Studio, enable only those dimensions or properties for sorting which are specifically needed by an application. Sorting gets slower as the process size grows and paging gets deeper.

In general, the MDEX Engine explicitly uses precomputed sorts for properties that you specifically configure as sort keys in Developer Studio, using the **“Prepare sort offline”** option.

Sorting can be done on any property, whether configured for sort or not. Configuring for sort mainly controls the generation of a precomputed sort (an internal optimization done by the MDEX Engine), and secondarily enables the field to be returned in the API sort keys function. In cases where the precomputed sort is rarely or never used (such as when the number of search results is typically small), the memory can be saved.

If the Dgraph has to compute precomputed sort objects to answer queries, the precomputed sort process in the Dgraph can be time-consuming. As a side effect of this processing, if you issue the `admin?op=exit` command to shut down the Dgraph while the precomputed sort process is still running, the actual shutdown may be delayed from the time the command is issued. This delay occurs because the Dgraph shutdown process may still be waiting for the completion of its creating several precomputed sort objects.

Geospatial sorting and filtering

Geospatial sorting and filtering is a query-time operation. The computation time it requires increases as larger sets of records are sorted and filtered. For best performance, apply geospatial sorting and filtering once the set of records has been reduced by normal refinement or search.

To optimize performance of geofilters, consider using these recommendations:

- Examine the request log for the presence of long distance queries that contain a geofilter. If there is a noticeable percentage of such queries, remove the geofilter from them.

In other words, if a portion of your queries represents searches in which distance is very large and thus appears to be not an important factor in a query, remove the geofilter from such queries.

For example, for users searching for cars within a radius beyond 10,000 miles, remove the geofilter for those queries. Removing the geofilter does not affect the records returned, but cuts the MDEX Engine response times in half.

In general, when the MDEX Engine applies a geofilter, it first uses the area's bounding rectangle to reduce the number of records it has to consider, and then performs the computation on remaining records, to determine if the record falls within the specified radius. This computation is expensive. For queries containing a geofilter for very large distances, the bounding rectangle includes all records, which means that the MDEX Engine performs this expensive computation for each record.

- Restrict the number of records returned to speed up MDEX Engine performance.

Range filters

Range filters do not impact the amount of memory needed by the Dgraph. However, because the feature is evaluated entirely at request time, the Dgraph response times are directly related to the number of records being evaluated for a given range filter request.

You should test your application to ensure that the resulting performance is compatible with the requirements of the implementation.

Record filters

Record filters can impact the following areas.

- Spelling auto-correction and spelling Did You Mean. Record filters impose an extra performance cost on spelling auto-correction and spelling Did You Mean.
- Memory cost
- Expression evaluation
- Large OR filters ("part lists")
- Large scale negation
- Record filters with complex logic

Record filters: memory cost

The evaluation of record filter expressions is based on the same indexing technology that supports navigation queries in the Dgraph. Because of this, there is no additional memory or indexing cost associated with using navigation dimension values in record filters.

When using property values in record filter expressions, additional memory and indexing cost is incurred because String properties are not indexed for navigation by default.

In some cases, it may be worth replacing some of the filters with dimensions that have the same meaning. For example, if you notice that 20% of queries have a filter of "price > 0" on them, to improve performance, add a "has price?" dimension to your records instead of using a filter in this case.

Expression evaluation in record filters: impact on performance

Because expression evaluation is based on composition of indexed information, most expressions of moderate size (that is, tens of terms and operators) do not add significantly to request processing time. Furthermore, because the Dgraph caches the results of record filter operations, the costs of expression evaluation are typically only incurred on the first use of a filter during a navigation session. However, some expected uses of record filters have known performance bounds, which are described in the following sections.

Large OR filters ("part lists")

One common use of record filters is to specify lists of individual records to identify data subsets (for example, custom part lists for individual customers, culled from a superset of parts for all customers).

The total cost of processing records can be broken down into two main parts: the parsing cost and the evaluation cost. For large expressions such as "part lists", which are commonly stored as file-based filters, XML parsing performance dominates total processing cost.

XML parsing cost is linear in relation to the size of the filter expression, but incurs a much higher unit cost than actual expression evaluation. Though lightweight, expression evaluation exhibits non-linear slowdown as the size of the expression grows.

OR expressions with a small number of operands perform linearly in the number of results, even for large result sets. While the expression evaluation cost is reasonable into the low millions of records for large OR expressions, parsing costs relative to total query execution time can become too large, even for smaller numbers of records.

Part lists beyond approximately one hundred thousand records generally result in unacceptable performance (10 seconds or more load time, depending on hardware platform). Lists with over one million records can take a minute or more to load, depending on hardware. Because results are cached, load time is generally only an issue on the first use of a filter during a session. However, long load times can cause other Dgraph requests to be delayed and should generally be avoided.

Large-scale negation

In most common cases, where the NOT operator is used in conjunction with other positive expressions (that is, AND with a positive property value), the cost of negation does not add significantly to the cost of expression evaluation.

However, the costs associated with less typical, large-scale negation operations can be significant. For example, running top-level negation filtering, such as "NOT availability=FALSE" on a record set of several million records leads to lower throughput.

If possible, attempt to rephrase expressions to avoid the top-level use of NOT in Boolean expressions. For example, in the case where you want to list only available products, the expression "availability=TRUE" yields better performance than "NOT availability=FALSE".

Optimizing URL record filters that use complex logic

URL record filters with complex logic may cause an expected growth in memory usage for the MDEX Engine. You can create either a fast-running filter that heavily uses memory, or a slow-running filter

that uses minimum memory. This section explains the trade offs and recommends which filter logic you should use.

The filter syntax dictates the sequence in which queries are being run by the MDEX Engine.

Use these recommendations:

- If your goal is to run the record filter as quickly as possible, regardless of concerns for potential memory usage growth on the MDEX Engine server, use the query logic in your filter that is as flat as possible. In other words, use AND and OR operations directly on the records, and do not use nested operations.

For example, this filter lists several records directly without any nested operations. It maximizes query performance at the expense of memory usage:

```
Nr=OR(P_WineID:89955,P_WineID:73036,P_WineID:69087,P_WineID:69993,
P_WineID:60641,P_WineID:58831,P_WineID:44996,P_WineID:52212,
P_WineID:81192,P_WineID:75040,P_WineID:76632)
```

- If your goal is to run the record filter that minimizes memory usage by the MDEX Engine, each AND and OR statements should contain at most two direct records. Since in many cases you may need to include more than two records in your filters, you can nest AND and OR operations.

For example, this heavily nested filter minimizes memory usage at the expense of MDEX Engine query processing time:

```
Nr=OR(OR(OR(OR(OR(OR(OR(OR(OR(P_WineID:89955,P_WineID:73036),
P_WineID:69087),P_WineID:69993),P_WineID:60641),P_WineID:58831),
P_WineID:44996),P_WineID:52212),P_WineID:81192),P_WineID:75040),
P_WineID:76632)
```

To summarize, if the data set is large, the filter with flat query logic consumes more memory but runs faster than the filter with nested logic, which runs slower but consumes minimum memory.

If hardware limitations prevent you from accommodating the expected memory growth, change the logic of your existing URL record filter.

EQL expressions and Record Relationship Navigation

You can use Endeca Query Language (EQL) expressions for these purposes.

- To filter query results based on dimension values, individual property values, ranges of property values and search terms.
- To combine EQL expressions using Boolean logic.
- To enable an Endeca feature known as Record Relationship Navigation (RRN).

For more information on EQL and Record Relationship Navigation, see the *MDEX Engine Advanced Development Guide*.

When to use EQL-based filters vs. other filter types

You can use EQL expressions to express all of the filter capabilities that are also supported by range filters (Nf), text search filters (Ntt, Ntk, Ntx) and navigation refinements. This topic helps you decide which type of filters to use, EQL-based or regular.

In general, due to their Boolean logic capabilities, EQL expressions offer more flexibility than regular filters expressed through other `Ur1ENEQuery` parameters. However, EQL expressions have different performance characteristics, and demonstrate other effects that you should take into account when considering which type of filter to implement.

Consider the following characteristics when deciding which type of filters to use, EQL-based or regular:

- **Unless you need EQL filter functionality, use regular filters.**

In general, when it is possible to express a query using regular filters (range filters and other types), use those methods instead of EQL expressions, as they often provide better query performance. Use EQL expressions after you have evaluated using other features for expressing your query logic.

In particular:

- EQL-based filters may be slower than record filters (Nr).

Use record filters (Nr) for large filters. (Large filters are used to filter out lists of individual records that identify data subsets, for example custom part lists created for individual customers that are culled from a superset of parts for all customers.) Large filters are better expressed with file-based record filter expressions than with EQL expressions.

- EQL-based range filters are slower than range filters (Nf).

- **To utilize merchandising rules or other supplementary information generated by regular filters, use them alone or in combination with EQL filters.**

EQL-based filters do not trigger the same supplementary information as a similar refinement navigation or a text search filter. For example, a navigation refinement may trigger merchandising rules, but an EQL filter does not.

In cases when you want to take advantage of additional information, such as search reports, merchandising rules, DYM and - -whymatch, use either of the following solutions:

- Use regular filters.
- Use EQL expressions in conjunction with other query parameters (such as N, Ntt, and Nr filtering parameters, and Nf, NrK, Nrt, Nrr, and Nrm relevance ranking parameters).

EQL combined with these parameters provides such actions as triggering merchandising rules, sorting, search reports or relevance ranking.

For examples and information on the feature interaction possibilities, see the *Endeca Advanced Development Guide*.

- **To implement security, use record filters.**

Use record filters instead of EQL-based filters to implement security filtering, such as filtering based on user role or catalog type. Record filters (Nr) are useful also in cases when you want to use file-based filters. (File-based filters are the recommended method for filtering out large numbers of included or excluded records.)

- **To maximize the use of the Dgraph cache, use record filters.**

Use the Nr parameter instead of EQL for those parts of the filter that are static across many queries. This is because static parts of the filter are faster with Nr than with EQL, due to the maximized use of the filter cache.

EQL caches the results of the entire filter, as well as those of a few limited sub expressions. Record filters (Nr) also cache the full results of each filter. Thus, if some part of an EQL filter is static across

many queries and can be expressed in the language of the Nr parameter, it can be advantageous to use Nr for that part of the filter so as to maximize use of the cache.

- **For more flexibility, use filters in combination.**

Use EQL-based filters instead of record filters when you do not require security implications, or when you need more flexibility in expressing filter logic. In this case you may want to improve EQL filter performance by using record filters in conjunction with EQL-based filters, as explained in the next bullet.

- **To narrow down the set of records, use record filters first.**

Record filters act as pre-filters and narrow down the working set of records for future evaluation by the MDEX Engine. Other expressions in such a query operate only on records returned by a record filter. By comparison, EQL-based filters do not narrow down the working set of records in this way. This has performance implications.

When evaluating a query, the MDEX Engine first evaluates record filters of type Nr, and then all other filters.

Performance impact of EQL-based filters

Use the following recommendations to optimize query performance of EQL-based filters.

- To optimize the performance of EQL-based filters, use record filters in conjunction with EQL-based filters. Use record filters first (Nr) if you can, to narrow down the working set, and then use EQL logic to filter within the smaller working set of records.
- Monitor the size of the standard Dgraph request log file. EQL-based filters have verbose syntax. Since all Endeca queries are logged to the standard Dgraph request log, the size of EQL-based queries affects disk space due to the growing size of the Dgraph logs. As an alternative, consider using file-based record filters.
- Identify slow queries during testing. To determine whether an EQL-based filter is slowing down your navigation queries, set the EQL statistics logging in the Dgraph. For example:

```
--log_stats <file_name>
--log_stats_thresh N
```

This file contains timing for queries taking longer than the specified threshold. Oracle recommends setting a low threshold value during development, and a more conservative value for testing. Do not use statistics logging in production since the verbosity of the logs can cause heavy disk writes and consume available disk space. Look for nodes with large `self_time_ms` values to identify the total time, in milliseconds, spent in this query node and its descendants.

- To optimize EQL query performance, use EQL for queries based on property value instead of queries based on range. For example, if the application's price property contains only 0 or positive values, using an EQL expression to query for "not (price = 0)" provides a better query performance than using queries of type "price > 0". (This recommendation is true for regular range filters as well.)
- To speed up the MDEX Engine processing of queries, consider implementing the filtering logic in the Forge pipeline. For more complex range expressions, it is more efficient to implement the filtering logic in the Forge pipeline. Use expression logic in a record manipulator or Java manipulator to create a new property with a Boolean value.

For example, create an "onsale = true" property value if the record has "price > 0" and "price < listprice" properties, and then use the EQL expression to perform a query based on the property

value for the newly created property (that is, for “onsale = true”), rather than using EQL for computing range filter expressions on the original properties.

Performance impact of RRN

You can use EQL expressions for Record Relationship Navigation (RRN).



Important: You must configure the MDEX Engine in order to enable Record Relationship Navigation. This capability is an optional module that extends the MDEX Engine. Endeca customers who are entitled by their license to use Record Relationship Navigation can find instructions on the Endeca Support site. Contact your Endeca representative if you need to obtain a Record Relationship Navigation license.

Use the following recommendations to speed up the RRN queries:

- When writing RRN filters, take into account that RRN filter expressions work from the inside out. That is, the innermost, or most nested, expressions are evaluated by the MDEX Engine before the outer ones. The following example illustrates this bottom-up processing:

```
collection()/record
[
  author_bookref = collection()/record
  [
    book_year = "1843"
  ]
  /book_id
]
```

The MDEX Engine first finds the records that have the `book_year` property set to “1843”. Then it finds the list of all of the values in the `book_id` property for that set of records. Finally, it finds the set of records with the `author_bookref` property set to any of the values in that list.

- To speed up RRN queries, assign different property names for records representing different concepts. This is because RRN query performance depends on the number of records in the “nested” EQL query. Keep the number of records that match results for the innermost expression of the RRN filter relatively small. For example, in this query:

```
collection()/record[
  record_type = "Film"
  and
  endeca:matches(., "title", "Godfather")
  and
  actor_id = collection()/record[
    record_type = "Actor"
    and
    gender = "male"
    and
    nationality = "Italian"
  ]
]
```

the MDEX Engine uses its bottom-up query execution strategy in the following way:

It first evaluates the inner query and finds the set of records for which the `record_type` property has the value “Actor,” the `gender` property has the value “male,” and the `nationality` property has the value “Italian.”

It then creates a collection of all the values of the `id` property for this set of records.

Next, it iterates over the set of `/id` values to filter the set of `"Film"` records. Thus, if the size of the collection of `/id` values is really large, the iteration can be relatively slow.

In this example, if the number of film IDs that are returned from the innermost filter to the Actor filter is relatively small, the RRN filter that will evaluate these records will be fast; if the number of IDs returned is large, the RRN evaluation will be slow.

To generalize, when you know that the number of records that will have to be evaluated for a RRN filter is quite large (in this example, it is the number of Italian male actors), a query could be slow. To solve this problem, one solution is to use the user interface and force the users to narrow down the set of records early on in the navigation process.

If this is not a reasonable solution for your application, and you cannot guarantee that the user's navigation path will necessarily limit the set of records, you can narrow down this set by limiting the number of records that match in the innermost query, as shown in this example:

```
collection()/record[
  record_type = "Film"
  and
  endeca:matches(., "title", "Godfather")
  and
  actor_id = collection()/record[
    record_type = "Actor"
    and
    gender = "male"
    and
    nationality = "Italian"
    and
    film_id = collection()/record[
      record_type = "Film"
      and
      endeca:matches(., "title", "Godfather")
    ]/id
  ]/id
]
```

This method is mimicking a top-down execution of a query.

While building an application, test the performance of this inner query with EQL statistics logging to evaluate the time spent in it.

- To speed up RRN queries, assign different property names for different record types of the RRN `collection()/record` function.

For example, consider this generic RRN query:

```
collection()/record[propertyKey1 = recordPath/propertyKey2]
```

where:

`propertyKey1` is the NCName of an Endeca property on a record type to be filtered, such as record of type `Vineyard`. The resulting records will have this property.

`recordPath` is one or more of the `collection()/record` functions.

`propertyKey2` is the NCName of an Endeca property on another record type, such as record of type `Wine`, that will be compared to `propertyKey1`. Records that satisfy the comparison will be added by the MDEX Engine to the returned set of records.

In this example, instead of assigning the same value of `"ID"` for `propertyKey1` and `propertyKey2`, assign two different property names— `"wine_reference_ID"` on a record representing a vineyard, and `"wine_ID"` on a record representing a wine. As the number of records evaluated for

the RRN query increases, having the naming convention with different property names for different record types has a greater effect on performance.

When properties with the same name are assigned on each side of the RRN query, this negatively affects RRN query performance.

For more information about RRN, see the *Endeca Advanced Development Guide*.

Tips for troubleshooting EQL filters

To detect queries with errors in EQL, check the Dgraph standard error log located at `$ENDECA_PROJECT_DIR/logs/dgraphs/DgraphN/DgraphN.reqlog`.

We recommend using `tail -f` to follow the log during query development.

To troubleshoot EQL filters, use the following recommendations:

- **Watch for disk space limitations.** All Endeca queries are logged to the standard Dgraph request log. Be careful to monitor the size of this log file; there is a risk to run out of disk space due to large log files.
- **Watch for filter length limitations.** The Dgraph process has no limits on the length of a request. The Endeca APIs, however, may have limitations stemming from the programming languages in which they are implemented.
- **Detect slow EQL queries with a dedicated statistics log.** Use these Dgraph flags to enable a special EQL statistics log:
 - `--log_stats [path_to_file]`
 - `--log_stats_thresh N`

The log contains an execution plan, including timing, for queries taking longer than the specified threshold. To identify slow EQL queries, in the log, look for nodes with large `self_time_ms` values.

This statistics logging is turned off by default. Specifying a target for `--log_stats` implicitly turns it on.

Oracle recommends placing this log in the same directory as all other Dgraph logs, such as in: `$ENDECA_PROJECT_DIR/logs/dgraphs/DgraphN/DgraphN.eqlllog`.

You can specify values for the optional `--log_stats_thresh` argument either as seconds or milliseconds, such as 1s or 500. If unspecified, the default is 60 seconds. Oracle recommends setting a low threshold value during development and a more conservative value for testing to capture queries that take longer than the threshold. In general, do not use statistics logging in production, as additional logs can cause operational issues due to heavy disk usage and consumption of available disk space.

Typical causes of EQL filter errors

EQL filter errors are logged into the Dgraph standard error log. This topic lists the most frequent causes of EQL filter errors.

When errors with parsing or syntax occur in EQL filters, they are logged to the Dgraph standard error log located at `$ENDECA_PROJECT_DIR/logs/dgraphs/DgraphN/DgraphN.reqlog`.



Note: When EQL filter errors occur, the query returns zero results and no messages are included in the API response. Therefore, it is important to look into the Dgraph standard error log.

The top issues that may cause errors in the EQL filters are the following:

- **Missing brackets.** Make sure your expressions have matching brackets [] and parentheses ().
- **Case-sensitivity.** All fields and values are case-sensitive. This includes boolean operators which must be lower case.
- **Property is not indexed properly.** Ensure that you enable properties for record filters in Developer Studio. For any property enabled for record filtering, the Dgidx process creates an inverted index. If a property is not enabled, you may receive an error message like this: Property "p_name" is not invertible; comparison will fail.
- **Property or dimension is not an NCName.** For example, "Wine Type" is not correct, "Wine_Type" or "WineType" are correct.
- **Whitespace is present in values.** For example, this is applicable to property value filters: "Foo" != "Foo".

Snippeting

You can minimize the performance impact of snippeting by limiting the number of words in a property that the MDEX Engine evaluates to identify the snippet.

This approach is especially useful in cases where a snippet-enabled property stores large amounts of text. Provide the `--snip_cutoff <num words>` flag to the Dgraph to restrict the number of words that the MDEX Engine evaluates in a property. For example, `--snip_cutoff 300` evaluates the first 300 words of the property to identify the snippet.

If the `--snip_cutoff` Dgraph flag is not specified, or is specified without a value, the snippeting feature defaults to a cutoff value of 500 words.

Spelling auto-correction and Did You Mean

This section discusses tuning the spelling auto-correction and spelling Did You Mean features.

Spelling auto-correction

Spelling auto-correction performance is impacted by the size of the dictionary in use. Spell-corrected keyword searches with many words, in systems with very large dictionaries, can take a disproportionately long time to process relative to other Dgraph requests.

It is important to carefully analyze the performance of the system together with application requirements prior to production application deployment.

Performance of `admin?op=updateaspell`

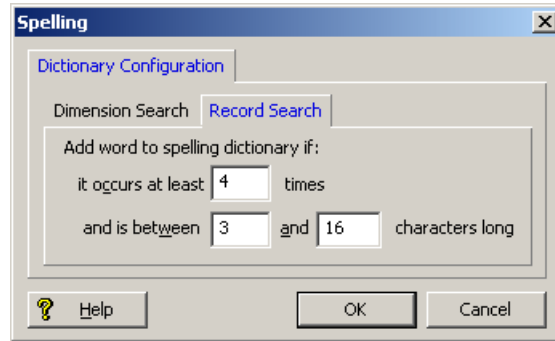
You can use the `admin?op=updateaspell` administrative query to make changes to the Aspell spelling dictionary without having to stop and restart the MDEX Engine. This administrative query causes the MDEX Engine to temporarily stop processing other regular queries, update the spelling dictionary and then resume its regular processing.

If the total amount of searchable text is large, this increases the latency of the `admin?op=updateaspell` operation, especially at large data scale.

Dictionary pruning

The performance of spelling correction in the Dgraph depends heavily on the size of the dictionary. An unnecessarily large dictionary can slow response times and provide less focused results.

Dictionary pruning techniques allow you to reduce the size of the dictionary without sacrificing much in the way of usefulness. To improve spelling correction performance, consider making the following adjustments in Developer Studio's Spelling editor:



- Set the minimum number of word occurrences to a number greater than one.

The first setting in the Spelling editor indicates the number of times a word must occur in the source data in order for it to be included in the dictionary. For record search, the default value is four, which means only words that appear four or more times are included in the dictionary.

- Set the minimum word length to a number greater than one.

The second setting in the Spelling editor specifies the minimum length (number of characters) of a word for inclusion in the dictionary. By default, words that are longer than three characters and shorter than sixteen characters are included.

While less dramatic than tuning the minimum word occurrences, adjusting the minimum word length can result in a cleaner, more useful dictionary.

Tuning word break analysis

Word-break analysis allows you to consider alternate queries computed by changing the word divisions in the user's query. The performance impact of word-break analysis can be considerable, depending on your data. Seemingly small deviations from default values, such as increasing the value of `--wb_maxbrks` from one to two or decreasing the value of `--wb_minbrklen` from two to one, can have a significant impact, because they greatly increase the workload on the MDEX Engine. Endeca suggests that you tune this feature carefully and test its impact thoroughly before exposing it in a production environment.

Did You Mean

Lowering the value for `--dym_hthresh` (a Dgraph spelling option) may improve the performance of Did You Mean.

The option `--dym_hthresh` indicates when spelling Did You Mean engages. The default is 20, meaning that spelling Did You Mean engages even if there are up to 20 results.

Depending upon your data, making Did You Mean suggestions at this point may be unnecessary or even overwhelming to your end users. Setting `--dym_hthresh` to 2 or 4 is often a better choice.

Stemming and thesaurus

Stemming and thesaurus equivalences generally introduce little memory overhead (beyond the amount of memory required to store the raw string forms of the equivalences).

In terms of online processing, both features expand the set of results for typical user queries.

While this generally slows search performance (search operations require an amount of time that grows linearly with the number of results), typically these additional results are a required part of the application behavior and cannot be avoided.

The overhead involved in matching the user query to thesaurus and stemming forms is generally low, but could slow performance in cases where a large thesaurus (tens of thousands of entries) is asked to process long search queries (dozens of terms).

Because matching for stemming entries is performed on a single-word basis, the cost for stemming-oriented query expansion does not grow with the size of the stemming database or with the length of the query. However, the stemming performance of a specific language is affected by the degree to which the language is inflected. For example, German nouns are much more inflected than English nouns.

Guidelines for thesaurus development

To avoid performance problems related to expensive and non-useful thesaurus search query expansions, consider the following thesaurus clean-up rules.

- Use `--thesaurus_cutoff <limit>` to set a limit on the number of words in a user's search query that are subject to thesaurus replacement. The default value of `<limit>` is 3. Up to 3 words in a user's search query can be replaced with thesaurus entries. If there are more terms in the query that match thesaurus entries, these terms are not replaced by thesaurus expansion. This option serves as a performance guard against very expensive thesaurus queries. Lower values improve thesaurus engine performance.
- Do not create a two-way thesaurus entry for a word with multiple meanings. For example, *khaki* can refer to a color as well as to a style of pants. If you create a two-way thesaurus entry for *khaki* = *pants*, then a user's search for *khaki towels* could return irrelevant results for pants.
- Do not create a two-way thesaurus entry between a general and several more-specific terms, such as *top* = *shirt* = *sweater* = *vest*. This increases the number of results the user has to go through while reducing the overall accuracy of the items returned.

In this instance, better results are attained by creating individual one-way thesaurus entries between the general term *top* and each of the more specific terms.

- Use care when creating thesaurus entries that include a term that is a substring of another term in the entry. Consider the following example with a two-way equivalency between *Adam and Eve* and *Eve*.

If users type *Eve*, they get results for *Eve* or (*Adam and Eve*) (that is, the same results they would have gotten for *Eve* without the thesaurus). If users type *Adam and Eve*, they get results for (*Adam and Eve*) or *Eve*, causing the Adam part of the query to be ignored.

There are times when this behavior might be desirable (such as in an equivalency between *George Washington* and *Washington*), but not always.

- Do not use stop words such as *and* or *the* in single-word thesaurus forms.

For example, if *the* has been configured as a stop word, thesaurus equivalency between *thee* and *the* is not useful.

You can use stop words in multi-word thesaurus forms, because multi-word thesaurus forms are handled as phrases. In phrases, a stop word is treated as a literal word and not a stop word.

- Avoid multi-word thesaurus forms where single-word forms are appropriate.

In particular, avoid multi-word forms that are not phrases that users are likely to type, or to which phrase expansion is likely to provide relevant additional results. For example, the two-way thesaurus entry *Aethelstan, King Of England (D. 939) = Athelstan, King Of England (D. 939)* should be replaced with the single-word form *Aethelstan = Athelstan*.

- Thesaurus forms should not use non-searchable characters. For example, the one-way thesaurus entry *Pikes Peak > Pike's Peak* should only be used if apostrophe (') is enabled as a search character.
- Use `--thesaurus_multiword_nostem` to specify that words in a multiple-word thesaurus form should be treated like phrases and should not be stemmed. This may increase performance for some query loads. Single-word terms will be subject to stemming regardless of whether this flag is specified.

This flag prevents the Dgraph from expanding multi-word thesaurus forms by stemming. Thesaurus entries continue to match any stemmed form in the query, but multi-word expansions only include explicitly listed forms. To get the multi-word stemmed thesaurus expansions, the various forms must be listed explicitly in the thesaurus.

Record, phrase, and dimension search

This section discusses the performance impact of various kinds of search.

Record search

Because record search is an indexed feature, each property enabled for record search increases the size of the Dgraph process. The specific size of the increase is related to the size of the unique word list generated by the specific property in the data set.

Therefore, only properties that are needed by an application for record searching should be configured as such.

Boolean search

The performance of Boolean search is a function of the number of terms and operators in the query and also the number of records associated with each term in the query.

As the number of records increases and as the number of terms and operators increase, queries become more expensive.

Proximity search impacts the system in various ways. The performance of proximity searches is as follows:

- Searches using the proximity operators will be slower than searches using the other Boolean operators.
- Proximity searches that operate on phrases will be slower than other proximity searches and slower than normal phrase searches.



Note: If you notice unexpected behavior while using Boolean search, use the Dgraph `-v` flag when starting the Dgraph. This flag prints detailed output to `stderr` describing the running Boolean query process.

Phrase search

The cost of phrase search operations depends mostly on how frequently the query words appear in the data and the number of words in the phrase. You can improve performance of phrase search by limiting the number of words in a phrase with the `--phrase_max <num>` flag for the Dgraph.

Searches for phrases containing relatively infrequent words (such as proper names) are generally very rapid.

You can use the `--phrase_max <num>` flag for the Dgraph to specify the maximum number of words in each phrase for text search. Using this flag improves performance of text search with phrases. The default number is 10. If the maximum number of words in a phrase is exceeded, the phrase is truncated to the maximum word count and a warning is logged.

Wildcard search

The MDEX Engine uses a mechanism for wildcard search that simplifies user configuration. In most cases, the size of the on-disk index is reduced considerably, and indexing performance is improved compared with previous releases. This topic provides recommendations for optimizing your wildcard search performance.

To optimize performance of wildcard search, use the following recommendations:

- **Account for increased time needed for indexing.** In general, if wildcard search is enabled in the MDEX Engine (even if it is not used by the users), it increases the time and disk space required for indexing. Therefore, consider first the business requirements for your Endeca application to decide whether you need to use wildcard search.



Note: To optimize performance, the MDEX Engine performs wildcard indexing for words that are shorter than 1024 characters. Words that are longer than 1024 characters are not indexed for wildcard search.

- **Do not use "low information" queries.** For optimal performance, Oracle recommends using wildcard search queries with at least 2-3 non-wildcarded characters in them, such as `abc*` and `ab*de`. Avoid wildcard searches with one non-wildcarded character, such as `a*`, since they are more expensive to process. Also be aware that the MDEX Engine ignores queries that contain only wildcards, such as `*`. Similarly, wildcard queries that contain only punctuation symbols, spaces and wildcards, such as `*.`, `*'`, or `* *`, are ignored.
- **Analyze the format of your typical wildcard query cases.** This lets you be aware of performance implications associated with one specific wildcard search pattern. Examine your queries to identify whether you have queries that contain punctuation syntax in between strings of text, such as

ab*c.def*. For strings with punctuation, the MDEX Engine generates lists of words that match each of the punctuation-separated wildcard expressions. In this case, the MDEX Engine uses the `--wildcard_max <count>` setting to optimize its performance. This setting does not affect wildcard searches for strings which do not contain punctuation.

You enable wildcard search in Developer Studio.

Wildcard search with punctuation and performance

The number of terms to which the MDEX Engine matches the wildcard search strings is limited by the `--wildcard_max <count>` number (the default is 100). This flag lets you specify to the MDEX Engine the maximum number of terms that can match a wildcard term in a wildcard search query that contains punctuation.

When a search reaches the `--wildcard_max` limit, the verbose Dgraph error log records a message similar to the following: `Wildcard term 1*0*.234* is too general: returns 1618 words, which is greater than max of 100. Using the most frequent 100 terms, which took 46.2 ms. to compute.`

Increasing the `--wildcard_max <count>` improves the completeness of results returned by wildcard search for strings with punctuation, but negatively affects performance. Thus you may want to find the number that provides a reasonable trade-off.

If your wildcard search queries contain punctuation, such as `1*0*.234*`, the MDEX Engine generates lists of words that match each of the punctuation-separated wildcard expressions, and uses these non-wildcard terms to locate related results in the documents (records).

This means that if the corpus of data contains other possible matches beyond the `--wildcard_max <count>` (and beyond the results that are already found), the MDEX Engine may not return them as results. Thus, the list of results returned by the Engine in a wildcard search with punctuation may not be exhaustive. This creates a trade-off situation in which you need to optimize performance cost versus business value of maximum completeness of returned results.

To summarize, if the business requirements of your application require a nearly 100% complete list of results even on very "low-information" wildcard queries with punctuation, such as `1*0*.234*`, increase the value of `wildcard_max`. Next, pay attention to the information returned in the search report. From it, you can estimate whether it may make sense to increase the `wildcard_max` value further.

Gradually increase the `--wildcard_max` value, while watching the performance of the MDEX Engine.



Note: If search queries contain only wildcards and punctuation, such as `*.*`, the MDEX Engine rejects them for performance reasons and returns no results.

Preventing expensive wildcard searches

Certain types of wildcard queries may cause the MDEX Engine to grow in memory footprint and take a long time to complete. Even though these types of queries are legitimate searches that would eventually return, they can cause the appearance of a timeout and potentially cause a site outage. As a best practice, Oracle recommends preventing these types of wildcard queries in your front-end application code.

The behavior of such wildcard queries does not typically indicate an actual timeout of the MDEX Engine; instead, it may indicate, for example, that the query search term is so broad that it takes a very long time to compute results. For example, to process a search for `"a*"`, the MDEX Engine must

return every record containing any word beginning with a; this is a more time-intensive query for the Dgraph to compute.

The following types of wildcard queries are potentially very expensive to compute for the MDEX Engine:

- Wildcard queries with short search terms, such as `*a*`, `*/*`, or `* *`.
- Wildcard queries with search terms that contain non-searchable characters, such as punctuation or dashes.
- Wildcard queries with search terms that have quoted phrases in them, such as `"pizza pie"`.

To prevent users from issuing such types of wildcard queries, utilize front-end application code to circumvent these scenarios for all queries that contain a wildcard character (`*`).



Note: If search queries contain only wildcards and punctuation, such as `*.*`, the MDEX Engine rejects them for performance reasons and returns no results.

Use the following recommendations in the front-end application, by utilizing application code at query time:

1. Remove all non-searchable characters from each wildcard query before issuing it to the MDEX Engine.

Stripping non-searchable characters should make little difference in your search results because the MDEX Engine treats non-searchable characters as white space both when indexing and when retrieving word matches.

2. Parse the queries to calculate their search term length to avoid very low information queries, such as `"a*"`. For example, you may want to prevent issuing to the MDEX Engine wildcarding queries that contain fewer than 3 non-wildcarded characters.

Filtering out such queries should make no difference in your search results because wildcard search for two characters or less would bring back an unusable results set in almost all instances.

3. Exclude wildcard queries with quoted phrase searches. This will not affect your search results because when users issue quoted phrase search, most likely they expect exact matches and do not require wildcards in this case.

You can accomplish these recommendations in the front-end application tier by programmatically analyzing search terms entered by the users before issuing them to the MDEX Engine, determining whether a query will be issued, and prompting the user to submit a better query (or using logic of your choice to handle this situation).



Note: In the majority of cases, none of these changes should impact the user experience.

Dimension search

The runtime performance of dimension search directly corresponds to the number of dimension values and the size of the resulting set of matching dimension values. In general, this feature performs at a much higher number of operations per second than navigation requests.

The most common performance problem occurs when the resulting set of dimension values is exceptionally large (greater than 1,000), thus creating a large results page. Always use the advanced dimension search and query parameters to limit the number of results per request. For details, see "Using Dimension Search" in the *MDEX Engine Basic Development Guide*.

Compound dimension search requests are generally more expensive than non-compound requests, and are comparable in performance to record search requests.

To summarize, if you submit a default dimension search query, the query is generally very fast. If you submit a compound dimension search query, performance is not as fast as for the default dimension search. In both cases, the query will be faster if you limit the results by using any of the advanced dimension search parameters. For example, you can use the `Di` parameter to specify the specific dimension (in the case of the default dimension search), or a list of dimension value IDs (in the case of compound dimension search) for which you expect matches returned by the MDEX Engine.



Note: Do not confuse the Dgraph configuration for dimension search with the Dgraph configuration to enable record search.

Precedence rules

This section discusses precedence rules and explains their performance impact.

About precedence rules

Precedence rules let you limit the presentation of certain Guided Navigation dimensions only to specified navigation states.

You configure precedence rules in Developer Studio.

Each precedence rule lets you identify a trigger dimension value and a target dimension, and presents the target dimension for Guided Navigation only in those query contexts in which:

- Users explicitly select the trigger dimension value as a refinement, or
- The trigger dimension value is assigned to all records in the current result set.

Example of a precedence rule

For example, suppose that an application includes a precedence rule linking the trigger dimension value “Part Category > Passives > Resistors” to a target dimension “Resistance”, which might contain refinements such as “10 ohms” and “22 ohms”.

In a navigation query where, for example, the user performs a search matching records tagged with a variety of values from “Part Category” including “Resistors” and other values, and where the user does not explicitly or implicitly select the dimension value “Part Category > Passives > Resistors”, the “Resistance” dimension is not returned for Guided Navigation.

This prevents the presentation of a contextually irrelevant navigation dimension to the user. Before the user has indicated some interest in resistors, presenting “Resistance” navigation choices may be unexpected, clutter the presentation of more relevant navigation choices, and detract from the overall experience.

If the user subsequently selects the “Part Category > Passives > Resistors” dimension value as a refinement, the “Resistance” dimension is presented for Guided Navigation (assuming that there are valid, available navigation refinements available for “Resistance”). Similarly, if the user performs a search that triggered “Part Category > Passives > Resistors” as an implicit refinement, for example if the user performed a text search for a manufacturer who only makes resistors, the “Resistance” dimension is returned for navigation.

This unique behavior provided by the MDEX Engine allows the contextual presentation of appropriate navigation dimensions to be more automatic and adaptive, as the front-end application need not be aware that the user's search has implied "Part Category > Passives > Resistors" for the "Resistance" dimension to be presented automatically as a navigation dimension.

Relevance ranking

Relevance ranking can impose a significant computational cost in the context of affected search operations (that is, operations where relevance ranking is enabled).



Important: The relevance ranking chapter in the *MDEX Engine Advanced Development Guide* outlines recommended strategies for both retail catalogs and document repositories. If you are developing an Endeca application on your own, you should start with the recommended relevance ranking strategy. Later, if the recommended strategy is not sufficient, you can experiment carefully with strategy tuning.

The set of modules that will provide acceptable performance depends heavily on the size and characteristics of the application data set.

In general, Oracle recommends testing the set of modules used for relevance ranking in a staging environment before using it in production. This is because the qualities of the data set may affect relevance ranking performance in unexpected ways. The following characteristics of the data set may negatively affect performance:

- The data set is too large to fit into RAM
- It contains large file content used in search
- It uses stemming or thesaurus heavily
- It has many dimensions or properties per record
- It frequently produces large result set sizes

Minimizing the performance impact of relevance ranking

You can minimize the performance impact of relevance ranking in your implementation by making module substitutions when appropriate, and ordering the modules you do select sensibly within your relevance ranking strategy.

Making module substitutions

Because of the linear cost of relevance ranking in the size of the result set, the actual cost of relevance ranking depends heavily on the set of ranking modules used. In general, modules that do not perform text evaluation introduce significantly lower computational costs than text-matching-oriented modules.

Although the relative cost of the various ranking modules is dependent on the nature of your data and the number of records, the modules can be roughly grouped into four tiers:

- Exact is very computationally expensive.
- Proximity, Phrase with Subphrase or Query Expansion options specified, and First are all high-cost modules, presented in the order of decreasing cost.
- WFreq can also be costly in some situations.
- The remaining modules (Static, Phrase with no options specified, Freq, Spell, Glom, Nterms, Interp, Numfields, Maxfields and Field) are generally relatively cheap.

In order to maximize the performance of your relevance ranking strategy, consider a less expensive way to get similar results. For example, replacing Exact with Phrase may improve performance with relatively little impact on results.



Note: Choose the set of modules used for relevance ranking most carefully when the data set is large or contains large file content that is used for search operations.

Ordering modules sensibly

Relevance ranking modules are only evaluated as needed. When higher-priority modules determine the order of records, lower-priority modules do not need to be calculated. This can have a dramatic impact on performance when higher-cost modules have a lower priority than a lower-cost module.

To optimize performance, make sure that the cheaper modules are placed before the more expensive ones in your strategy.

Dynamic business rules

Dynamic business rules (used in merchandising and content spotlighting) require very little data processing or indexing, so they do not impact the Dgraph memory footprint.

However, because the MDEX Engine evaluates dynamic business rules at query time, the larger the number of rules, the longer the evaluation and response time.

To improve query response-time performance of the Dgraph with dynamic business rules:

- Monitor and limit the number of rules that are evaluated for each request. Each rule that is evaluated for a request impacts the response time for that request.

To do this, specify the number of records returned in the Maximum Records text box of the Styles editor in Developer Studio. Setting the Maximum Records value prevents business rules from returning an entire set of matching records, potentially overloading the network, memory, and page size limits for a request. If the Maximum Records value is set to a large number, such as 1,000, then as many as 1,000 promoted records could be returned with each navigation request, causing significant performance degradation.

- Use `Nmr f` to specify the syntax for the rule filter. Rule filters restricts which rules can promote records for a navigation query. The `Nmr f` query parameter controls the use of a rule filter. `Nmr f` has a corresponding `ENEQuery` method and parameter.
- Set a rule limit for each rule zone.
- Configure triggers for all business rules. Business rules without triggers are evaluated for every navigation query and negatively affect performance.
- Review how rule sorting is used. Rule sorting allows you to sort the rule's promoted records by a specified property or dimension value. Per-rule sorts can increase the performance cost of dynamic business rules.

Analytics performance considerations

This section explores issues related to optimizing performance of Analytics queries.

For more information about how to use Analytics functions, and for examples and best practices, see the *MDEX Engine Analytics Guide*.

Each of the following considerations has an impact on the Analytics query performance:

- Review existing Analytics queries to understand their processing order and Analytics statement dependencies. For example, you may improve query performance if you narrow down the working record set which Analytics statements must process.

When a query contains an Analytics query, the Analytics processing is one of the last steps in the overall query processing order. The Analytics statements are calculated on the resulting record set (NavStateRecords) after any search, navigation, or filtering has been applied by the Endeca query. This has performance benefits, since the fewer records the Analytics statements need to process, the better.

- Test Analytics queries that contain a GROUP BY operation to measure RAM footprint and query response time. This will help identify the size of a result set that does not negatively affect performance.

GROUP BY operations result in a large number of aggregated records that are stored within the Dgraph RAM. This may cause an increase in the RAM footprint and the Dgraph processing time. It may be necessary to tune GROUP BY operations within Analytics statements in your queries.

- Build Analytics queries in a way that lets them utilize the caching of Analytics statements used in more than one query.

The Dgraph dynamic cache stores Analytics statements. If statement dependencies exist in your queries, you can utilize previously computed data within other Analytics statements. If one Analytics query includes multiple Analytics statements, each statement is cached separately, which results in a significant performance gain in cases when specific Analytics statements are shared across multiple queries.

Appendix A

The MDEX Engine Request Log

This section describes the MDEX Engine (Dgraph) request log, which you can use to analyze Endeca application performance.

About the MDEX Engine request log

The MDEX Engine request log (also called the Dgraph request log) is the file that captures Web application query information.

The MDEX Engine always generates a request log with a default name `dgraph.reqlog`. You use the `--log` option when running the MDEX Engine to specify a different path to store the request log.

You can extract queries from this log file and use them with the Endeca Eneper tool to analyze Web application performance. You can also use Perl to extract useful information from Dgraph request logs.

In addition, depending upon the size of your log files, you can import them into a tool that allows you to manipulate column-based data, such as Microsoft Excel.

Related Links

[Extracting information from request logs](#) on page 81

MDEX Engine request logs can be very large and difficult to read. You might find it useful to sort them on fields you are interested in, such as Processing Time or Total Request Duration. You can then look for a pattern or feature in the most time-consuming queries that might be the origin of the performance issue.

Request log file format

The content of the request log file varies slightly, depending upon whether it is treating Presentation API queries or Web services invocations.



Note: If a field is not relevant to the query in question, the request log entry for that query contains a dash (-) in that location.

Each entry has the following 14 columns:

```
[Timestamp] [Client IP Address]
[HTTP Exchange ID] [Response Size] [Total Request Time]
[Total Processing Time] [HTTP Return Code] [Number of Results]
```

[Queue Status] [Thread ID] [Query String] [Query Body]
[HTTP Headers]

These entries are listed in the order of the timestamp. Because of this, the entries are listed in the response order, not in the request order. The following table describes the log entries in more detail:

Column	Presentation API Queries or Web Services Invocations	Description
Timestamp	Both	Time stamp indicating the time the request was completed, in milliseconds, since the epoch (January 1, 1970, 00:00:00 UTC). For example: 1208947882000=2008-04-23 10:51:22 AM GMT The time is recorded in GMT (not the localized time of the server). You can convert it using a UTC epoch converter utility, such as UTC.
Client IP	Both	IP address of the requesting client.
HTTP Exchange ID	Both	Unique query identifier. This identifier allows you to correlate Dgraph request log items with error messages in the Dgraph log. In addition, it is used by the MDEX Server Statistics page to compose most expensive query statistics.  Note: The identifier is only unique within a single Dgraph instance, and is not persistent across Dgraph shutdown.
Response Size	Both	Number of bytes written to the client. May be less than or equal to the intended result size, for example, due to a premature session end.
Total Request Time	Both	The request lifetime, in milliseconds. Equal to the total amount of time between when the Dgraph reads the request from the network and finishes sending the result. May include queuing time, such as time spent waiting for earlier requests to be completed.  Note: In previous releases, the request lifetime ended when the connection was closed. If connection close did not time out, this lifetime would include the time to transport the response to the client, and the time for the client to read the response. Starting with 6.1.0, the request lifetime ends when the response has been successfully delivered to the socket layer.
Total Processing Time	Both	Processing time, in milliseconds. Equal to the total computation time required for the Dgraph to handle the request, excluding network and wait time. This value gives an accurate measure of how expensive the request was to compute, given current system state. (That is, if the machine in question was

Column	Presentation API Queries or Web Services Invocations	Description
		busy with other threads or processes, the time may be longer than on an otherwise unused machine.) For any given query, Processing Time is always smaller than Total Request Time.
HTTP Status Code	Both	The HTTP return code. A status code of 200 (OK) is returned if the request was successful. For details on other codes that can appear in this field, see the table below.
Number of Results	Presentation API Queries	Number of results from your query (or "-" if the HTTP request was not a query).  Note: This number reflects the number of results, not necessarily the number of results returned. That is, this is the number of results from your query, not accounting for your nbins and offset settings. nbins and offset are used to specify how many of the results are actually returned.
Queue Status	Both	The number of queries in the queue that have not started processing yet. The number is calculated before the current query is enqueued, and therefore the current query is not included.  Note: Starting with the MDEX Engine version 6.1.2, this column does not report the number of query threads that are idle because there is no longer a one-to-one relationship between threads and queries. Specifically, when you specify the <code>--threads</code> flag, the number you specify determines the total number of threads available to the MDEX Engine, which includes query processing threads and other threads that support query processing. This means there is a greater chance that a non-saturated Dgraph could experience minor queuing, even in the case when the number of query requests in the queue is less than the number of threads specified. For more information, see the chapter in this guide about using the multithreaded mode.
Thread ID	Both	The thread ID of the thread that was assigned the request (or "-" in single-threaded mode).
Query String	Both	The URL of the Presentation API query or of the Web service.
Query Body	Web Services Invocations	The URL-encoded POST body of the query. The actual entry in the request log is a single token, even though POST body can contain multiple lines of text.
HTTP Headers	Both	The URL-encoded HTTP headers that were sent with the query.

Column	Presentation API Queries or Web Services Invocations	Description
		The actual entry in the request log is a single token, even though HTTP headers can contain multiple lines of text.

Non-OK HTTP Status Codes

This table details the non-OK HTTP Status Codes that might appear in the Request Log.

Status Code	Name	Condition
100	Continue	In response to HTTP request header Expect: 100-continue (not an error)
400	Bad Request	Admin or config request with unsupported op
400	Bad Request	HTTP request line parse error, or HTTP request header parse error, or HTTP request Transfer-Encoding other than chunked
400	Bad Request	HTTP request with invalid chunk size or missing chunk terminator
400	Bad Request	HTTP request with invalid trailing header format
400	Bad Request	HTTP request with wildcard URL ("*") not valid for METHOD
400	Bad Request	HTTP request URL includes protocol other than "http", or protocol but no host, or neither protocol nor host and path does not start with "/"
400	Bad Request	HTTP request with version 1.1 but no Host
400	Bad Request	HTTP request with more data than expected
400	Bad Request	Conversion of POST body to string failed for web service request
403	Forbidden	Admin ops are disabled for the Dgraph, and admin?op=exit or admin?op=restart is requested
404	Not Found	Presentation API request with URI parse error or processing error
404	Not Found	Request has empty path, or admin or config request has additional path steps
404	Not Found	File server request for non-existent file or for a directory, or file outside of allowed root directory
404	Not Found	Web service request for unknown Web service
408	Request Time-out	Queue timeout exceeded for the request, or I/O timeout reading HTTP request
410	Gone	Presentation API request for unsupported feature
411	Length Required	HTTP POST request with Content-Length missing or empty or not a non-negative integer
412	Precondition Failed	HTTP request with "If-None-Match" header

Status Code	Name	Condition
415	Unsupported Media Type	Content-Type parse error in Web service request
500	Internal Server Error	Attempt to return informational status code to HTTP 1.0 client
500	Internal Server Error	Exception from XQuery evaluation in Web service request
500	Internal Server Error	Unhandled exception during request processing
500	Internal Server Error	admin?op=update is requested and no update directory was specified for the Dgraph
501	Not Implemented	HTTP request for unsupported Method (such as PUT)
501	Not Implemented	HTTP request includes an unsupported header that must not be ignored: ("Authorization", "Content-Encoding", "Content-Transfer-Encoding", "Range", "Content-Range", "If-Range")
501	Not Implemented	Presentation API request for disabled feature
503	Service Unavailable	HTTP request to server that is closed (in the process of shutting down)
505	HTTP Version Not Supported	HTTP request with version not "1.0" and not "1.1"

Related Links

[List of request log parameters](#) on page 88

This section lists request log parameters.

Extracting information from request logs

MDEX Engine request logs can be very large and difficult to read. You might find it useful to sort them on fields you are interested in, such as Processing Time or Total Request Duration. You can then look for a pattern or feature in the most time-consuming queries that might be the origin of the performance issue.

Here are two approaches to extract information from request logs:

- Run the Request Log Analyzer.
- Write your own Perl code.

The Request Log Analyzer reads one or more MDEX Engine logs and reports on the nature and performance of the queries recorded in those logs. This report provides information on what actually happened in the past, instead of reporting on potential performance or capacity planning for the future. This script can be run manually in order to debug performance problems, and should also be run on a regular basis to continually monitor performance and call out trends in Dgraph traffic load, latency, throughput, and application behavior.

If you write Perl to extract, manipulate, and analyze the information in a request log, you may find the following setting useful in Perl scripts:

```
perl -nae
```

where:

- n indicates that it is a loop processing each line of the input file(s) in turn
- a turns on autosplit
- e indicates that it should execute the next argument, which should be Perl code

This script shows how many queries took more than five seconds. It splits the line on whitespace into an array called F. The sixth element in the array ([5]) corresponds to the Total Request Time and represents the amount of time the query took.

```
perl -nae 'print if $F[5] > 5000' logfile
```

If you are tracking system trends by time, you may find it useful to correlate the epochal time that the log displays with human-readable time. This script is used to convert the time stamps into a more readable form.

```
perl -nae 'print scalar localtime $F[0], " $_"'
```



Note: In this script, `Localtime` is set to the location where you are doing analysis, so if you are looking at a log from a different time zone, you may want to change the time zone. On UNIX systems the `TZ` environment variable can be set to effect this change. For example, `TZ=US/Pacific`.

Storing logs on a separate physical drive

There can be disk contention between MDEX logging and update processing that can cause sporadic increases in query processing latencies. Update processing includes both partial update processing and merging generations. One way to minimize disk contention is to store MDEX logs, such as the error log, request log, and update log on a separate physical drive from where the MDEX indices are stored.

To store logs on a separate physical drive:

- For the error log, specify the `--out <stdout/stderrfile>` flag to the Dgraph with a path to a different physical drive from the MDEX indices.
- For the request log, specify the `--log` flag to the Dgraph with a path to a different physical drive from the MDEX indices.
- For the update log, specify the `--update-log` flag to the Dgraph with a path to a different physical drive from the MDEX indices.

Request log rolling

The MDEX Engine request log is subject to log rotation when it goes over one gigabyte.

When this occurs, the existing logfile is renamed from, say, `dgraph.reqlog` to `dgraph.reqlog.PID.N`, where:

- PID is the Dgraph process ID

- N is the number of logs that this Dgraph has already rotated. N=0 the first time the Dgraph does log rotation, and then goes up by 1 each time.

Appendix B

The MDEX Engine Parameter Listing

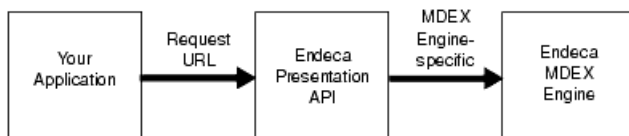
This section describes the parameters in the MDEX Engine request logs and provides mappings between the URL that is sent from the application to the Endeca Presentation API, and the URL that is sent from the API to the MDEX Engine.

Understanding the URL parameter mapping

Typically, when you analyze the MDEX Engine request query logs for troubleshooting purposes, you investigate a log entry for a query in question, and identify an MDEX Engine parameter in the query's log entry.

Next, you want to trace this log parameter to its corresponding settings in the user-visible URL that is sent from the application to the Endeca Presentation API and the URL that is sent from the API to the MDEX Engine. There is not a one-to-one correlation between the two URLs.

The Presentation API transforms the URL it receives from the application into an MDEX Engine-specific URL before sending it to the MDEX Engine.



Mappings between request log and UrlENEQuery URL parameters

This explains a mapping between the URL that is sent from the application to the Endeca Presentation API, and the URL that is sent from the API to the MDEX Engine.

It helps you translate the MDEX Engine request log file, which tells you exactly which URLs the MDEX Engine has processed. By extension, these are the URLs that the Presentation API has sent to the MDEX Engine. If the API has sent an incorrect URL to the MDEX Engine, it is a good indication that the API received an incorrect URL from the Web application in the first place.



Note: For a complete description of the ENE URL query parameters, see the *Endeca MDEX Engine Advanced Development Guide*.

Example mappings

Here are some sample mappings:

Web Application to API	API to MDEX Engine
/controller.jsp?N=0	/graph?node=0
/controller.jsp?N=0&Ntk=DESC&Ntt=merlot	/graph?node=0+attrs=DESC+merlot

Mapping parameters

The table in this section establishes a mapping between those MDEX Engine request log parameters that have corresponding `Ur1ENEQuery` URL parameters, such as **N** and **Ntt**.

Not all request log parameters have corresponding `Ur1ENEQuery` URL parameters. This table does not list those MDEX Engine request log parameters that do not have directly corresponding end-user parameters. It also does not indicate which methods or properties of the `ENEQuery` objects can be used to produce the specified request log parameters.

In this table, the ENE parameters in bold are the primary parameters, while those in non-bold are secondary parameters.

MDEX Engine parameter	Description	Maps to...
graph?	Navigation query	N
node	Navigation query parameter, navigation descriptors	N
offset	Navigation query parameter, record offset	No
offset	Navigation query parameter, aggregated record offset	Nao
group	Navigation query parameter, exposed refinements	Ne
allbins	Navigation query parameter, records per aggregated record	Np
analytics	Navigation query parameter, analytics expression to apply to a query	Na
sort	Navigation query parameter, sort	Ns
sort	Navigation query parameter, sort order	Nso
groupby	Navigation query parameter, rollup	Nu
attrs	Navigation query parameter, record search key, terms, and options	Ntk, Ntt, Ntx
relrank	Navigation query parameter, search interface, relevance	Nrk, Nrt, Nrr, Nrm

MDEX Engine parameter	Description	Maps to...
	ranking terms, relevance ranking strategy and match mode	
dym	Navigation query parameter, Did You Mean	Nty
autophrase	Navigation query parameter, compute phrasings	Ntpc
autophrasedwim	Navigation query parameter, rewrite query	Ntpr
merchpreviewtime	Navigation query parameter, merchandising preview time	Nmpt
merchrulefilter	Navigation query parameter, merchandising rule filter	Nmrf
pred	Navigation query parameter, range filters	Nf
filter	Navigation query parameter, record filters	Nr
structured	Navigation query parameter, Endeca Query Language	Nrs
refinement	Navigation query parameter, dynamic refinement ranking	Nrc
search?	Dimension search query	D
terms	Dimension search query parameter, search terms	D
options	Dimension search query parameter, options	Dx
node	Dimension search query parameter, dimension search scope	Dn
model	Dimension search query parameter, search dimension	Di
num	Dimension search query parameter, number of results	Dp
offset	Dimension search query parameter, offset	Do
rank	Dimension search query parameter, rank	Dk
pred	Dimension search query parameter, range filters	Df
filter	Dimension search query parameter, record filters	Dr

MDEX Engine parameter	Description	Maps to...
structured	Dimension search query parameter, Endeca Query Language	Drs
abin?	Aggregated record query	A
id	Aggregated record query parameter, record ID	A
node	Aggregated record query parameter, descriptors	An
groupby	Aggregated record query parameter, rollup	Au
pred	Aggregated record query parameter, range filters	Af
filter	Aggregated record query parameter, record filters	Ar
structured	Aggregated record query parameter, Endeca Query Language	Ars
bin?	Record query	R
id	Record query parameter, record ID	R

List of request log parameters

This section lists request log parameters.

It provides the following information:

- Lists the request log parameters and explains what they do.
- Identifies how the request log parameters correspond with the end user visible URL parameters. In other words, a mapping is established between the parameters that are visible in the end-user URL, known as the `Ur1ENEQuery` URL parameters, such as `N` and `Nt t`, and the parameters that are present in the request log, such as `node` and `attrs`.
- Lists those request log parameters that do not have directly corresponding end-user parameters, such as `allgroups` and `nbins`.
- Indicates which methods or properties of the `ENEQuery` objects can be used to produce the specified request log parameters.

In general, in your application, you use either the `Ur1ENEQuery` URL parameters, such as `N` and `Nt t`, or the methods or properties of the `ENEQuery` object class. In either case, both methods produce the MDEX request log parameters described in this section.

Example: interpreting error log messages

This example illustrates how to interpret the messages found in the MDEX Engine error log.

Suppose the following messages appear in your MDEX Engine error log:

```
ERROR 06/04/08 18:13:33.250 UTC DGRAPH {dgraph}: Bad dimension or property
name
[WineType] in select
```

To troubleshoot, look through the corresponding MDEX request log for entries that contain “select” and “WineType”. The results are as follows:

```
1212603213 127.0.0.1 - 3378 105.54 7.49 200 56300 -2 10
/graph?node=0&select=P_Name+P_Score+WineType&group=0&offset=0&nbins=10&pred=
P_Score%7CGTEQ+70&irversion=510
```

Check the documentation in this section for the select parameter that appears in the MDEX Engine URL, in the request log. You will find that it corresponds to the Java API `ENEQuery.setSelection()` method; there is no corresponding `UrlENEQuery` URL parameter. This means that the incorrect value is set through this method. You can now look through the application code and find the `setSelection()` call to try to determine why it is specifying an incorrect property or dimension name as part of the value for this method. In this example, it is because the code is specifying “WineType” rather than “Wine Type” with a space.

Description of query types

The parameters in the MDEX request log use the query type names that correspond to the types of user queries. This section and the table below list the query types and maps them to user queries.

Query type as indicated in the request log	Description of the corresponding user query type
admin, config	Administrative query
bin, abin	Record query
graph	Navigation and record search queries that return navigation data
search	Dimension search queries only

allbins

Description	Specifies the number of representative records returned with each aggregated record.
Valid in query types	graph
ENEQuery method or property	Java: <code>ENEQuery.setNavErecsPerAggrERec()</code> .NET: <code>ENEQuery.NavErecsPerAggrERec</code>
UrlENEQuery URL parameters	Np
Format	Numeric value

Values (order)	0 (no representative records), 1 (one representative record), 2 (all records associated with aggregated record). Value "0" equates to API constant ENEQuery.ZERO_ERECs_PER_AGGR, "1" to ENEQuery.ONE_ERECs_PER_AGGR "2" to ENEQuery.ALL_ERECs_PER_AGGR
Example	N/A

allgroups

Description	Specifies whether child refinements are exposed for all dimension values. Takes precedence over group if both are specified. The API includes one parameter or the other.  Note: allgroups=1 in the Dgraph URL can cause significant impact on performance of the MDEX Engine and indicates that all refinements are exposed for navigation. If you notice this setting in the queries check the validity of this setting for the application.
Valid in query types	graph
ENEQuery method or property	Java: ENEQuery.setNavAllRefinements() .NET: ENEQuery.NavAllRefinements
UrlENEQuery URL parameters	N/A
Format	Numeric Boolean value
Values (order)	0 (false), 1 (true)
Example	N/A

analytics

Description	Specifies an analytics expression to apply to a query.
Valid in query types	graph
ENEQuery method or property	Java: ENEQuery.setAnalyticsQuery() .NET: ENEQuery.AnalyticsQuery

UrlENEQuery URL parameters	Na
Format	String analytics expression
Values (order)	N/A
Example	analytics=Q%28A%28Test%28T%29SL%28S%28%28Vintage%29KEY%28Vin- tage%29%29%29%29%29

attrs

Description	Specifies search key, terms, and options for record searches
Valid in query types	graph
ENEQuery method or property	Java: ENEQuery.setNavERecSearches() .NET: ENEQuery.NavERecSearches
UrlENEQuery URL parameters	Ntk, Ntt, Ntx
Format	Space-separated string values for search key, literal plus character separator, space-separated string values for search terms, pipe character separator, space-separated string values for search options (mode, rel, and autoforce).
Values (order)	See above
Example	attrs=Inter- face+search+terms mode+matchall+rel+ex- act+autoforce+correction

autophrase

Description	Specifies whether the MDEX Engine computes autophrase matches for search terms.
Valid in query types	graph
ENEQuery method or property	Java: ENEQuery.setNavERecSearchComputeAlternativePhrasings() .NET: ENEQuery.NavERecSearchComputeAlternativePhrasings
UrlENEQuery URL parameters	Ntpc
Format	Numeric Boolean value
Values (order)	0 (false), 1 (true)
Example	N/A

autophrasedwim

Description	Specifies whether the MDEX Engine replaces phrases found in search terms with computed autophrase matches. Is functional only if the autophrase parameter is also set to 1 (true).
Valid in query types	graph
ENEQuery method or property	Java: <code>ENEQuery.setNavERecSearchRewriteQueryWithAnAlternativePhrasing()</code> .NET: <code>ENEQuery.NavERecSearchRewriteQueryWithAnAlternativePhrasing</code>
UrlENEQuery URL parameters	Ntpr
Format	Numeric Boolean value
Values (order)	0 (false), 1 (true)
Example	N/A

compound

Description	Specifies whether dimension search is performed as a compound dimension search.
Valid in query types	search
ENEQuery method or property	Java: <code>ENEQuery.setDimSearchCompound()</code> .NET: <code>ENEQuery.DimSearchCompound</code>
UrlENEQuery URL parameters	N/A
Format	Numeric Boolean value
Values (order)	0 (false), 1 (true)
Example	N/A

dym

Description	Specifies whether "did you mean" (DYM) spelling correction is enabled for a record search.
Valid in query types	graph
ENEQuery method or property	Java: <code>ENEQuery.setNavERecSearchDidYouMean()</code> .NET: <code>ENEQuery.NavERecSearchDidYouMean</code>
UrlENEQuery URL parameters	Nty

Format	Numeric Boolean value
Values (order)	(order) 0 (false), 1 (true)
Example	N/A

filter

Description	Specifies record filter to apply for navigation, dimension-search, or aggregated-record (abin) queries.
Valid in query types	graph, search, abin
ENEQuery method or property	(graph) Java: <code>ENEQuery.setNavRecordFilter()</code> .NET: <code>ENEQuery.NavRecordFilter</code> (search) Java: <code>ENEQuery.setDimSearchNavRecordFilter()</code> .NET: <code>ENEQuery.DimSearchNavRecordFilter</code> (abin) Java: <code>ENEQuery.setAggrERecNavRecordFilter()</code> .NET: <code>ENEQuery.AggrERecNavRecordFilter</code>
UrlENEQuery URL parameters	Nr (graph), or Dr (search), or Ar (abin)
Format	String values separated by plus signs
Values (order)	String values
Example	<code>filter=P_Region%3aPortugal, filter=8021</code>

format

Description	Description Specifies result object return format for a query.  Note: Format can only be set by hand. XML schema is unsupported and is subject to change.
--------------------	---

Valid in query types	graph, search, bin, abin
ENEQuery method or property	N/A
UrlENEQuery URL parameters	N/A
Format	String value
Values (order)	binary (default) or XML
Example	N/A

group

Description	Specifies dimension values for which child refinements should be exposed; Overridden by allgroups if both are specified. The API includes one parameter or the other. Only a single dimval from any given dimension can be specified (even if the dimension is configured for multiselect).
Valid in query types	graph
ENEQuery method or property	Java: <code>ENEQuery.setNavExposedRefinements()</code> .NET: <code>ENEQuery.NavExposedRefinements</code>
UrlENEQuery URL parameters	Ne
Format	Space-separated numeric dimval IDs
Values (order)	Numeric dimval IDs
Example	group=123+3893+1232123

groupby

Description	Specifies rollup (aggregation) key to apply for navigation or aggregated-record queries.
Valid in query types	graph, abin
ENEQuery method or property	(graph) Java: <code>ENEQuery.setNavRollupKey()</code> .NET: <code>ENEQuery.NavRollupKey</code> (abin) Java: <code>ENEQuery.setAggrERecRollupKey()</code> .NET: <code>ENEQuery.AggrERecRollupKey</code>
UrlENEQuery URL parameters	Nu (graph), or

	Au (abin)
Format	Space-separated string property or dimension names
Values (order)	String property or dimension names
Example	groupby=My+DimName, groupby=P_Winery

id

Description	Specifies a record to return (by record spec value or other identifier).  Note: Aggregated-record (abin) queries only support a single record identifier, not a space-separated list.
Valid in query types	bin, abin
ENEQuery method or property	(bin) Java: ENEQuery.setERecs() .NET: ENEQuery.ERecs (abin) Java: ENEQuery.setAggrERecSpec() .NET: ENEQuery.AggrERecSpec
UrlENEQuery URL parameters	R (bin), or A (abin)
Format	Space-separated string values
Values (order)	String values
Example	id=18114, id=Record+23, id=2+73

ignore

Description	Specifies whether the Dgraph ignores missing dimension value IDs in a query. When set to false, queries with missing dimval IDs fail with "Invalid category id... in query" errors; when set to true, such queries return successfully with "Detected missing category... (query will return zero results)" messages.
Valid in query types	graph

ENEQuery method or property	N/A
UrlENEQuery URL parameters	N/A
Format	Numeric Boolean value
Values (order)	0 (false), 1 (true, default)
Example	N/A

irversion

Description	Specifies a major version of API; set automatically by API and should not be changed.
Valid in query types	graph, search, bin, abin
ENEQuery method or property	N/A
UrlENEQuery URL parameters	N/A
Format	Three-digit numeric value
Values (order)	N/A
Example	irversion=500 (5.0.x), irversion=510 (5.1.x), irversion=601 (6.0.1)

keyprops

Description	Specifies whether to return key properties with the query results.
Valid in query types	graph
ENEQuery method or property	Java: ENEQuery.setNavKeyProperties() .NET: ENEQuery.NavKeyProperties
UrlENEQuery URL parameters	Nk
Format	String value
Values (order)	none (default), all "All" equates to API constant ENEQuery.KEY_PROPS_ALL "None" equates to ENEQuery.KEY_PROPS_NONE
Example	N/A

lang

Description	Specifies a language to use for a query.
Valid in query types	graph, search
ENEQuery method or property	Java: <code>ENEQuery.setLanguageId()</code> .NET: <code>ENEQuery.LanguageId</code>
UrlENEQuery URL parameters	LanguageId
Format	N/A
Values (order)	Standard language code string value
Example	lang=en for English, lang=zn_CH for simplified Chinese

log

Description	Specifies session and query ID values.
Valid in query types	graph, search, bin, abin
ENEQuery method or property	Java: <code>ENEQuery.setQueryInfo()</code> .NET: <code>ENEQuery.QueryInfo</code>
UrlENEQuery URL parameters	N/A
Format	String containing one or more URL-encoded key=value pairs, separated by ampersands.
Values (order)	key=value pairs
Example	log=sid%3d11586B%26rid%3d11586

merchdebug

Description	Specifies debugging output for business rule evaluation in the Dgraph error log. Configured by the --merch_debug flag.
Valid in query types	graph
ENEQuery method or property	Java: <code>ENEQuery.setMerchDebugOn()</code> .NET: <code>ENEQuery.MerchDebugOn</code>
UrlENEQuery URL parameters	N/A
Format	Numeric Boolean value
Values (order)	0 (false), 1 (true)
Example	N/A

merchpreviewtime

Description	Specifies preview time to use for business rules.
Valid in query types	graph
ENEQuery method or property	Java: ENEQuery.setNavMerchPreviewTime() .NET: ENEQuery.NavMerchPreviewTime
UrlENEQuery URL parameters	Nmpt
Format	String value
Values (order)	now (current time), or a date expressed in yyyy-mm-ddTmm:ss format (such as, 2007-07-12T08%3a15 for 8:15am, 12 July 2007).
Example	merchpreviewtime=now, merchpreviewtime=2007-08-28T12%3a51

merchrulefilter

Description	Specifies the filter for business rules.
Valid in query types	graph
ENEQuery method or property	Java: ENEQuery.setNavMerchRuleFilter() .NET: ENEQuery.NavMerchRuleFilter
UrlENEQuery URL parameters	Nmrf
Format	String value, formatted per record filters.
Values (order)	N/A
Example	merchrulefilter=endeca.internal.workflow.state%3aACTIVE

model

Description	Specifies dimension(s) to which dimension search will be restricted. Multiple values are only usable for compound dimension searches (such as, search for "ford tempo" against intersection of Make and Model dimensions). Simple dimension searches are restricted to a single dimension only, and return 0 results if multiple dimval IDs are specified.
Valid in query types	search

ENEQuery method or property	(search, simple) Java: <code>ENEQuery.setDimSearchDimension()</code> .NET: <code>ENEQuery.DimSearchDimension</code> (search, compound) Java: <code>ENEQuery.setDimSearchDimensions()</code> .NET: <code>ENEQuery.DimSearchDimensions</code>
UrlENEQuery URL parameters	Di
Format	Numeric dimval ID (simple dimension search), or space-separated list of numeric dimval IDs (compound dimension search).
Values (order)	N/A
Example	model=2344 (simple dimension search), model=1+18+25 (compound dimension search)

nbins

Description	Specifies maximum number of ERec objects to return for a navigation query, assuming that a query can be on non-aggregated records and on aggregated records. Does not map to any UrlENEQuery URL parameter.
Valid in query types	graph
ENEQuery method or property	- In non-aggregated navigation queries: Java: <code>ENEQuery.setNavNumERecs()</code> .NET: <code>ENEQuery.NavNumERecs</code> - In aggregated navigation queries: Java: <code>ENEQuery.setNavNumAggrERecs()</code> .NET: <code>ENEQuery.NavNumAggrERecs</code>
UrlENEQuery URL parameters	N/A
Format	Numeric value
Values (order)	10 (default)
Example	nbins=10 (default), nbins=500

nbulkbins

Description	Specifies maximum number of ERec objects to be returned via bulk export.
--------------------	--

	This parameter corresponds to different methods when querying aggregated records, that is, when a rollup key is applied.
Valid in query types	graph
ENEQuery method or property	(graph) Java: ENEQuery.setNavNumBulkERecs() .NET: ENEQuery.NavNumBulkERecs (graph, aggregated records) Java: ENEQuery.setNavNumBulkAggrERecs() .NET: ENEQuery.NavNumBulkAggrERecs
UrlENEQuery URL parameters	N/A
Format	Numeric value
Values (order)	Values (order) 0 (default), positive values, -1 (all records, or ENEQuery.MAX_BULK_ERECs_AVAILABLE)  Note: "-1" is equivalent to all records, or to setting ENEQuery.MAX_BULK_ERECs_AVAILABLE (that is, bulk-exporting all records matching the query) for the relevant methods.
Example	N/A

node

Description	Specifies selected (descriptor) dimension values.
Valid in query types	graph, search, abin
ENEQuery method or property	(graph) Java: ENEQuery.setNavDescriptors() .NET: ENEQuery.NavDescriptors (search) Java: ENEQuery.setDimSearchNavDescriptors() .NET: ENEQuery.DimSearchNavDescriptors (abin) Java: ENEQuery.setAggrERecNavDescriptors() .NET: ENEQuery.AggrERecNavDescriptors

	.NET: ENEQuery.AggrERecNavDescriptors
UrlENEQuery URL parameters	N (graph), Dn (search), An (abin)
Format	Space-separated numeric dimval IDs.
Values (order)	N/A
Example	node=0, node=125+234423+87

num

Description	Specifies the number of dimension value matches to return from each dimension as results of dimension search, per dimension.
Valid in query types	search
ENEQuery method or property	Java: ENEQuery.setDimSearchNumDimValues() .NET: ENEQuery.DimSearchNumDimValues
UrlENEQuery URL parameters	Dp
Format	Numeric value
Values (order)	N/A
Example	num=5

offset

Description	Specifies the number of values to skip before beginning to return record objects (for record search), or dimension value objects (for dimension search).
Valid in query types	graph, search
ENEQuery method or property	(graph) Java: ENEQuery.setNavERecsOffset() .NET: ENEQuery.NavERecsOffset (graph, aggregated records) Java: ENEQuery.setNavAggrERecsOffset() .NET: ENEQuery.NavAggrERecsOffset (search) Java: ENEQuery.setDimSearchResultsOffset() .NET: ENEQuery.DimSearchResultsOffset

UrlENEQuery URL parameters	No (graph) or Nao (graph, aggregated records), or Do (search)
Format	Numeric value
Values (order)	N/A
Example	offset=20 (begins returning objects from record or dimension value starting with 21 and onward).

op

Description	Specifies an operation to perform for command-type (non-query) URLs.
Valid in query types	admin, config
ENEQuery method or property	N/A
UrlENEQuery URL parameters	N/A
Format	String value
Values (order)	The following admin operations are supported: audit, auditreset, exit, flush, help, logroll, ping, restart, update, updatehistory, reload-services, stats, and statsreset. The following config operations are supported: help, log-disable, log-enable, log-status, and update.  Note: The config log-enable and log-disable operations can take several logging variables, which are documented in the MDEX Engine Logging Variables appendix to the <i>Endeca MDEX Engine Advanced Development Guide</i>.
Examples	admin?op=update, admin?op=stats, config?op=update

opts

Description	Specifies options, such as match mode, for dimension search. Also specifies a spell+nospell option for disabling spelling correction and DYM suggestions on individual queries.
Valid in query types	search

ENEQuery method or property	Java: ENEQuery.setDimSearchOpts() .NET: ENEQuery.DimSearchOpts
UrlENEQuery URL parameters	Dx
Format	Space-separated string values
Values (order)	N/A
Example	opts=mode+matchall+spell+nospell

pred

Description	Specifies a range filter expression for a query.
Valid in query types	graph, search, abin
ENEQuery method or property	(graph) Java: ENEQuery.setNavRangeFilters() .NET: ENEQuery.NavRangeFilters (search) Java: ENEQuery.setDimSearchNavRangeFilters() .NET: ENEQuery.DimSearchNavRangeFilters (abin) Java: ENEQuery.setAggrERecNavRangeFilters() .NET: ENEQuery.AggrERecNavRangeFilters
UrlENEQuery URL parameters	Nf (graph), or Df (search), or Af (abin)
Format	Space-separated string value
Values (order)	property or dimension name key, pipe character separator, operator (such as BTWN, GT), values.
Example	pred=P%5FPrice%7CBTWN+8+12 (restricts query to records where P_Price value is between 8 and 12).

pretendtime

Description	Specifies time value to use for time-triggered business rules.
--------------------	--

Valid in query types	graph
ENEQuery method or property	N/A
UrlENEQuery URL parameters	N/A
Format	String time value (m/ d/ yyyy hh:mm)
Values (order)	Value is the time of the Dgraph query.
Example	pretendtime=+2%2F+1%2F2007+11%3A49

profiles

Description	Specifies user profiles to apply to a query (used to restrict triggering of business rules).
Valid in query types	graph
ENEQuery method or property	Java: ENEQuery.setProfiles() .NET: ENEQuery.Profiles
UrlENEQuery URL parameters	N/A
Format	Space-separated list of string profile names
Values (order)	String profile names
Example	profiles=free_shipping+USA

rank

Description	Specifies whether to use relevance ranking to order dimension values returned by dimension search.
Valid in query types	search
ENEQuery method or property	Java: ENEQuery.setDimSearchRankResults() .NET: ENEQuery.DimSearchRankResults
UrlENEQuery URL parameters	Dk
Format	Numeric Boolean value
Values (order)	0 (default dimension value ranking), 1 (relevance ranking)
Example	N/A

refinement

Description	Specifies query-time dynamic refinement ranking settings.
Valid in query types	graph
ENEQuery method or property	Java: <code>ENEQuery.setNavRefinementConfigs()</code> .NET: <code>ENEQuery.NavRefinementConfigs</code>
UrlENEQuery URL parameters	Nrc
Format	Colon-separated list of space-separated values
Values (order)	string, number key, value pairs.
Example	<code>refinement=dimvalid:6300+dynrank:1+exposed:1+dynorder:0+dyncount:4</code>

rerank

Description	Specifies query-time relevance ranking settings.
Valid in query types	graph
ENEQuery method or property	- Through IAP 5.1.1: Java: <code>ENEQuery.setNavRelRankERecSearch()</code> .NET: <code>ENEQuery.NavRelRankERecSearch</code> - IAP 5.1.2 and later: Java: <code>ENEQuery.setNavRelRankERecRank()</code> .NET: <code>ENEQuery.NavRelRankERecRank</code>
UrlENEQuery URL parameters	Nrk Nrt Nrr Nrm
Format	Pipe-separated list of space-separated values
Values (order)	search key, search terms, relevance-ranking strategy, search mode
Example	<code>rerank=All napa+valley exact matchall</code>

select

Description	Specifies fields (properties and dimensions) to return on ERec objects from navigation query.
--------------------	---

Valid in query types	graph
ENEQuery method or property	Java: ENEQuery.setSelection() .NET: ENEQuery.Selection
UrlENEQuery URL parameters	N/A
Format	Space-separated list of string property/dimension name values
Values (order)	String property/dimension name values
Example	select=P_Name+Vintage

sort

Description	Description Specifies sort key(s) and order to use for records returned by a query.  Note: Current version only uses the Ns parameter (Nso is deprecated).
Valid in query types	graph
ENEQuery method or property	Java: ENEQuery.setNavActiveSortKeys() .NET: ENEQuery.NavActiveSortKeys
UrlENEQuery URL parameters	Ns Nso (deprecated)
Format	Pipe-separated list of string key order value pairs (two pipes between pairs)
Values (order)	asc (ascending), desc (descending)
Example	sort=P_Price asc Vintage desc

structured

Description	Specifies an Endeca Query Language (EQL) expression to apply to a query.
Valid in query types	graph, search, abins
ENEQuery method or property	(graph) Java: ENEQuery.setNavRecordStructureExpr() .NET: ENEQuery.NavRecordStructureExpr

	(search) Java: <code>ENEQuery.setDimSearchNavRecordStructureExpr()</code> .NET: <code>ENEQuery.DimSearchNavRecordStructureExpr</code> (abin) Java: <code>ENEQuery.setAggrERecStructureExpr()</code> .NET: <code>ENEQuery.AggrERecStructureExpr</code>
UrlENEQuery URL parameters	Nrs (graph), Drs (search) Ars (abin)
Format	String EQL expression
Values (order)	EQL expression
Example	<code>structured=collection%28%29%2frecored%5bP_Regi- on%3d%22Sonoma%22%5d</code>

terms

Description	Specifies search terms for dimension search.
Valid in query types	search
ENEQuery method or property	Java: <code>ENEQuery.setDimSearchTerms()</code> .NET: <code>ENEQuery.DimSearchTerms</code>
UrlENEQuery URL parameters	D
Format	Space-separated list
Values (order)	String values for terms
Example	<code>terms=my+search+terms</code>

Appendix C

Creating Eneperf input files with the Request Log Parser

The Request Log Parser is a lightweight tool that parses the MDEX Engine's request log and creates an input file, containing a list of query URLs, for use with the Eneperf load testing tool.

Installation location

The Request Log Parser is installed as `$ENDECA_MDEX_ROOT/bin/reqlogparser` (UNIX) and `%ENDECA_MDEX_ROOT%\bin\reqlogparser.exe` (Windows).

Log format requirements

The Request Log Parser supports request logs generated by IAP 4.x, IAP 5.x and MDEX Engine 6.x.

Invoking the Request Log Parser

The Request Log Parser is invoked as follows:

```
reqlogparser [--flags] request.log [request2.log...]
```

where `request.log` is a relative or absolute path to the MDEX Engine request log file to analyze. Multiple log files may be analyzed in a single run by specifying filenames in space-separated list form.

The Request Log Parser will write resulting parsed entries to standard output. To save results to a file, simply redirect standard output to a file.

If a filename is "-", then the request log is read from stdin. The "-" can be used in combination with other filenames but only one filename may be "-".

Per-file progress messages and a post-analysis summary message will be written to standard error unless the `--silent` flag is specified.

The Request Log Parser supports the following flags:

- `--help`: displays this message.

- `--filter <parameter_list>`: strip the specified query parameters and their values out of log entries. Parameters must be specified as a comma-separated list (e.g., "nbins,offset").
- `--input <number>`: specifies the number of entries to process. If not specified, the Request Log Parser will process all entries in the specified input files.
- `--keep-content-length`: if specified, will retain any "Content-Length" HTTP headers in MDEX Engine 6.x entries. These entries are removed by default.
- `--keep-ir`: if specified, will retain any irversion parameters found in entries. The irversion is an optional parameter sometimes specified by queries to indicate a particular version of the Endeca IAP. But this parameter often makes log entries less portable across versions, so by default the Request Log Parser removes it.
- `--noerrors`: removes entries with HTTP status codes 4xx (client errors) or 5xx (server errors). Equivalent to using both the `--no-client-errors` and `--no-server-errors` flags (both described below).
- `--no-client-errors`: removes entries with HTTP status codes 4xx (client errors).
- `--no-server-errors`: removes entries with HTTP status codes 5xx (server errors).
- `--output <number>`: specifies the number of valid entries to output. If not specified, the Request Log Parser will output all entries that have been processed.
- `--query-types <types>`: specifies which types of log entries to output. If not specified, the Request Log Parser will only output /graph, /search, /bin, /abin, and /ws queries (equivalent to `--query-types gsbaw`). Types include the following:

g: /graph	Navigation and search requests
s: /search	Dimension search request
b: /bin	Record request
a: /abin	Aggregate record request
w: /ws	Web Services queries
t: /admin?op=stats	Admin stats request
p: /admin?op=ping	Admin ping query
u: /admin?op=update	Admin update request
r: /admin?op=reload-services	Admin service reload request
d: /admin?op=updateaspell	Admin aspell-update request

- `--show-unrecognized`: causes the parser to write all unrecognized queries to standard error. Primarily useful for debugging.
- `--silent`: causes the parser to suppress display of per-file progress messages and post-parsing activity summary.
- `--version`: Print version information and exit.
- `--web-services`: causes the parser to display a breakdown of Web Services query subtypes (e.g., `rad_query` for `/ws/rad_query` entries) in the post-parsing activity summary report. This breakdown is in alphabetical order and shows the number of queries of each Web Services subtype parsed.

Example output from the Request Log Parser

By default, with no flags specified, the Request Log Parser generates the following output:

```
-----
-- reqlogparser: Parsing 'my/path/dgraph.log'...
-----
/search?terms=pinot+noir&opts=mode+matchall&rank=0&offset=0&compound=1
/graph?node=8024&group=0&offset=0&nbins=10 - -
/bin?id=37614 - -
/graph?node=8038&group=0&offset=0&nbins=10 - -
```

[...]

== SUMMARY OF PARSER ACTIVITY

```
=====
Total parsing time (seconds):                    5
Total log files read:                            1
Total entries parsed:                            100,000
Log file(s):                                     entries
-----
- my/path/dgraph.log                            100,000
-----
Total entries cleaned and retained:              99,994
-----
- Navigation/record-search queries (/graph)      99,200
- Dimension-search queries (/search)             790
- Record-detail queries (/bin)                   4
-----
Total entries discarded:                         6
-----
- Admin stats commands (/admin?op=stats)         4
- Startup messages ("DGRAPH STARTUP")            2
-----
```

Header Information

By default, the Request Log Parser outputs header information to standard error before each file is processed. For example:

```
-----
-- reqlogparser: Parsing 'my/path/dgraph1.log'...
-----
/search?terms=pinot+noir&opts=mode+matchall&rank=0&offset=0&compound=1
/graph?node=8024&group=0&offset=0&nbins=10
```

[...]

```
-----
-- reqlogparser: Parsing 'my/path/dgraph2.log'...
-----
/search?terms=pinot+noir&opts=mode+matchall&rank=0&offset=0&compound=1
/graph?node=8024&group=0&offset=0&nbins=10
```

[...]

Using the `--silent` flag will suppress this information.

Summary Information

By default, the Request Log Parser outputs a summary report to standard error after all files are processed. Even if multiple input files are specified, a single summary report is created.

The following metrics can be found in the summary report:

- **Total parsing time (seconds):** Amount of time it took for Request Log Parser script to complete.
- **Total log files parsed:** Number of raw log files specified as input.
- **Total entries parsed:** Number of raw input log entries across all input log files.
- **Log file(s)/entries:** Number of raw input log entries broken down by file.
- **Total entries cleaned and retained:** Total number of log entries that were parsed and output, as well as a breakdown by entry type (such as navigation/record search, dimension search, and record detail.)
- **Total entries discarded:** Total number of log entries that were discarded, as well as a breakdown by entry type (such as admin stats queries, admin exit commands, and startup/shutdown messages.)
- **Query parameters filtered/count:** List of query parameters specified for filtering, and a count of filtered occurrences for each. This entry is only displayed if the `--filter` flag is specified.
- **Web Services query subtype breakdown (alphabetical):** List of Web Services query subtypes encountered by the parser, with a count of occurrences for each. This entry is only displayed if the `-web-services` flag is specified.

Using the `--silent` flag will suppress display of this summary information.

Using the Request Log Parser with Eneperft

The Request Log Parser processes raw Dgraph logs into properly formatted input for Eneperft.

Eneperft is a performance debugging tool that can measure throughput to help you identify system bottlenecks. Eneperft drives a substantial load at the MDEX Engine and reveals how many operations per second the MDEX Engine responds with. You specify a log file as input and tell Eneperft how many times to run through it, as well as the number of client connections to simulate.

Eneperft is an executable that is included in the MDEX Engine package. It is located at `$ENDECA_MDEX_ROOT/bin/eneperft` (UNIX) and `%ENDECA_MDEX_ROOT%\bin\eneperft.exe` (Windows).

A Dgraph process (version 5.1) generates request logs in the following format:

```
1146617085 10.0.18.19 - - 15805 15.66 2.49 200 5000 0 - /graph?node=0
```

But as input, the Endeca IAP 5.1 version of Eneperft requires a single log file with URL-only entries in the following format:

```
/graph?node=0
```

MDEX Engine 6.x request log entries include two additional fields after this URL, for post body and HTTP headers. These do not appear in dgraph logs from earlier versions. The Request Log Parser will retain these additional fields when processing 6.x request log entries, as they are used by 6.x versions of Eneperft.

The following is an example of how the Request Log Parser can be used in conjunction with Eneperft to replay two separate dgraph.log files against an index running on 10.0.0.1:8000. In this example, the log file will be replayed ten times using a five simultaneous connection.

```
reqlogparser dgraph1.log dgraph2.log > dgraph_parsed.log
eneperft 10.0.0.1 8000 dgraph_parsed.log 5 10
```

Appendix D

Using the Eneperf Tool

Eneperf is a performance testing tool that is included in your Endeca installation. This section describes how to use Eneperf.

About Eneperf

Eneperf is a performance, analytics and debugging tool that can measure throughput to help you identify system bottlenecks. Eneperf makes HTTP queries against the MDEX Engine (Dgraph) based on your MDEX Engine request logs and gathers the resulting statistics, without processing the results in any way.

Because Eneperf is lightweight, it has a very slight impact on performance. In most cases, it can be run on the same machine as the Dgraph being tested. It can also be run on a remote machine.

Eneperf drives a substantial load at the MDEX Engine and reveals how many operations per second the MDEX Engine responds with. Eneperf lets you measure both query latency and throughput. You specify the log file and specify to Eneperf how many times to run through it, as well as the number of client connections to simulate.

Eneperf understands Endeca MDEX Engine URLs, which use the pipe symbol (|). Because the pipe symbol is not a legal character in the URL/URI standards, other programs, such as wget, may transform it inappropriately.

Using Eneperf

Eneperf is installed in the MDEX Engine bin directory. It has the following usage.

```
usage: eneperf [-v]
  [--header <header file path>]
  [--help] [--gzip]
  [--list] [--nreq <n>]
  [--nodnscache>] [--msec-between-updates]
  [--progress] [--pidcheck <pid>]
  [--prelude <log file path>] [--postlude <log file path>]
  [--quitenerror] [--rcvbuf <size bytes>]
  [--record <recording file prefix>] [--record_hdr]
  [--record_ord] [--record_roll <max KB per recording file>]
  [--reqstats] [--reqtimeout <secs>]
  [--runtime <max runtime (minutes)>]
```

```
[--seek <n>] [--seekrepeat] [--sleeponerror <secs>]
[--stats <num reqs>] [--throttle <max req/sec>]
[--updates-log] [--version]
[--warn <max req time warning threshold (msecs)>]
<host> <port> <log> <num connections> <num iterations>
```

Eneper has both required and optional settings.

Required settings

The required settings (shown in order) are as follows.

```
<host> <port> <log> <num connections> <num iterations>
```

Their usage is as follows.

Setting	Description
<host>	Target host for requests.
<port>	Port on which the target host is listening for requests.
<log>	Log file of the query portion of the MDEX Engine URLs and optional associated information (that is, the portion that resides in the last three columns of the MDEX Engine request log). This log file is used for HTTP request generation. URLs and associated information from the <log> file are replayed in order. Each line of the <log> file contains three columns: • A URL (required) • A POST body (URL-encoded and optional) • HTTP headers (URL-encoded and optional). If a dash (-) is found in an optional column, the column is ignored.
<num connections>	Maximum number of outstanding requests to allow before waiting for replies. In other words, the number of simultaneous HTTP connection streams to keep open at all times. This number emulates multiple clients for the target server. For example, using <num connections> of 16 emulates 16 concurrent clients querying the target server at all times.
<num iterations>	Number of times to replay the URL query log. All outstanding requests are processed before a new iteration is started.

Host and port settings for running Eneper locally or remotely

You can run Eneper locally or from a remote machine.

- Running Eneper locally. Eneper is lightweight and has a very slight impact on performance. It can usually be run on the same machine as the Dgraph being tested with no impact on results.

To run Eneper on the same machine as the Dgraph, you point it to `localhost` and `<port>`. This configuration is useful for isolating MDEX Engine performance from any potential networking issues.

- Running Eneper on a remote host. Eneper can also be run from a remote host. Using Eneper to test the same MDEX Engine from the local machine and from across the network can expose networking problems if the throughputs are significantly different.



Note: Eneper can be run on a machine with a different architecture than one you are testing.

Log file settings suitable for Eneper input

MDEX Engine request logs can be used as Eneper input with some modifications.

URLs in the log should not include any machine connection parameters such as protocol, host, or port. These are added automatically. For example, a log entry of the following form is valid:

```
/graph?node=0
```

But a log entry of the following form is not valid:

```
http://myhost:5555/graph?node=0
```

You can achieve higher concurrent load by using a single large request log file (which might simply be repeated concatenations of a smaller log file) than by using multiple iterations of a small log file. The log file should preferably be at least 100 lines, even if it consists of the same query repeated over and over. Because Eneper drains all connections between each iteration, running a one-line log file through Eneper 100 times results in skewed throughput statistics.

If you are planning to measure performance of partial updates with Eneper, (as opposed to measuring performance of regular queries), create a separate updates log based on your existing request log.

That is, suppose your MDEX Engine request log contains both regular queries and updates operations. Then your updates log should contain only `config?op=update` operations. You can create this updates log manually, by extracting these operations from a regular log. You can then run Eneper against the updates log and the regular log, to measure the performance of your updates, by using the `--updates-log` and the `--log` settings together.



Note: This is only one way to measure performance of updates and should only be used in cases when you care about the time between the updates. (If you do not care about the timing between updates, you can use the regular log for your testing.)

About the number of connections and iterations

Eneper load is driven by the `num_connections` setting, which indicates the number of simultaneous connections Eneper tries to maintain at a time.

For example, if `num_connections` is set to 4, it sends four requests to the MDEX Engine. When one returns, another is sent out to replace it.

To adequately measure performance of the MDEX Engine, you need to identify the number of connections for Eneper that saturates the MDEX Engine thread pool.

The number of connections needed to saturate the MDEX Engine depends on the MDEX Engine threading configuration and the server characteristics, and generally correlates with the number of the MDEX Engine threads in use, (assuming the MDEX Engine is configured with enough threads). However, an MDEX Engine with four threads might be saturated by only three connections if the queries are complex and all CPUs are being fully utilized.

To identify an appropriate setting for `num_connections`, Oracle recommends running tests with the following settings:

- For debugging, run a test with `num_connections` set to one. This test sends only one request to the MDEX Engine at a time. Each query is processed alone; no other query computations are contending for the machine's resources. This test generates an MDEX Engine request log showing the canonical time for each query. You can examine the request log to identify slow queries without the concern that they happened to be slow because other queries were processed simultaneously. Note that using a log file with just one entry limits `num_connections` to one.
- For stress testing, run a test with `num_connections` set to the number of threads for the MDEX Engine. In this test, no requests are waiting in the queue. This lets you obtain an estimate of the maximum expected MDEX Engine performance. Because no queuing occurs, this test offers a conservative bias for throughput.

In addition, you can run a test with `num_connections` set to the "number of threads + one". In this test case, a minimal waiting in the queue for the MDEX Engine request may occur. This also lets you obtain an estimate of the maximum expected MDEX Engine performance. Because queuing does not occur, this test offers an aggressive bias for throughput.

- Do not use a small log with a large number of `num_connections`. Also, do not run a small log many times to simulate a large log.

Example: Selecting the number of connections

Commonly, you will wish to perform the load testing of the MDEX Engine to a level below saturation. Use the following examples to help you select an appropriate number of connections for Eneperft that will saturate MDEX Engine performance to the desired levels.

Typically, front-end applications have different requirements for response times and peak loads. Such as:

- An application that is used steadily across the year. For applications of this type, MDEX Engine performance must support average query response time under average loads. Occasional slowdowns under peak load are acceptable. Therefore, you need to measure average response time under average load.
- An application that is used during the peak seasons. For applications of this type, MDEX Engine performance must support peak response time under peak loads. It is acceptable for this application to have extra performance capacity during non-peak seasons.

To identify the projected throughput for the MDEX Engine, use the following formulas.

These formulas represent a highly simplified approach to calculating throughput. Although you can use more thorough methods, these formulas provide reasonable estimates, and can be used for initial assessment:

```
concurrent users / (expected page latency + think time) = page views/sec
page views / second x MDEX queries/page = ops/second for the MDEX Engine
```

Where:

- The number of concurrent users is the estimated number of users currently logged in to the application

- The number of simultaneous requests is the number of users currently making a request to the application. Typically, it is 20-30% of the number of concurrent users.
- Peak load is the expected maximum number of simultaneous requests, such as during a specific time period
- Think time is the time between requests issued by a single user. It is used to calculate simultaneous requests based on the estimated number of concurrent users.

For example, 100 concurrent users with a 5 second think time and a 1 second expected page latency will yield 17 pages/sec. 17 pages/second with 2 MDEX Engine queries per page will yield 34 ops/sec for the expected performance of the MDEX Engine. This means that to support 100 concurrent users in this application, the MDEX Engine must perform at 34 ops/sec.

In another example, if your implementation includes a load balancer serving four application servers, and two MDEX Engines with another load balancer, the following calculations provide you with the estimated performance for each of the MDEX Engines:

- 600 concurrent users are distributed across 4 application servers. This means 150 users per server.
- 150 users divided by 5 (4 sec think time and 1 sec expected page latency) yields 30 simultaneous page views per server.
- 30 page views with 2 MDEX Engine queries per page yield 60 MDEX Engine queries per server.
- 60 queries per server multiplied by 4 application servers yield 240 queries total.
- 240 queries are sent to the load balancer that distributes them across two MDEX Engines. Each MDEX Engine serves 120 queries.

This means that to support 100 concurrent users in this application, each MDEX Engine must perform at 120 ops/sec.

To summarize, you can use these recommendations to identify the number of connections (equal to the number of simultaneous requests in these examples) that you need to provide to Eneper to achieve the desired MDEX Engine performance.

Optional settings

Eneper contains the following optional settings.

Setting	Description
-v	Verbose mode. Print query URLs as they are requested.
--gzip	Add Accept-encoding: gzip to the HTTP request header.
--header <header_file_path>	Specify path of file containing HTTP header text, one header field per line. This setting, if used, overrides headers from the log file (which you can also specify).
--help	Print the help usage and exit.
--list	Treat the <log> parameter as the name of a file containing the names of a sequence of request logs, rather than directly naming a single request log. As a result, Eneper iterates over the sequence of logs. Each line in the <log> names a request log file to be replayed against Eneper in sequence, during each iteration.
--msec-between-updates	If you use this setting with --updates-log, it specifies the minimum time interval between sending partial update requests,

Setting	Description
	in milliseconds. Before sending a new update request, Eneper waits for a free connection (after the specified time interval expires). This setting must not be used together with <code>--list</code>, <code>--seek</code>, <code>--seekrepeat</code>, <code>--prelude</code>, <code>--postlude</code>, and <code>--throttle</code>.  Note: The <code>--msec-between-updates</code> setting is optional. If you use only the <code>--updates-log</code> setting, Eneper processes updates one after another. Eneper waits for the current update to finish and immediately sends another update. It does not wait for any period of time between sending individual updates to the Dgraph.
<code>--nreq <n></code>	Stop after n requests.
<code>--nodnscache</code>	Disable caching of DNS hostname lookups. By default, Eneper caches these lookups to improve performance.
<code>--pidcheck <pid></code>	On a connection error, check the specified Dgraph process to see if it is running. If the process is not running, terminate Eneper.
<code>--prelude <log_file_path></code>	Specify a <code><log_file_path></code> of the file with URLs to replay before those of the <code><log></code> parameter, for each iteration. Use this flag together with the <code>--list</code> flag to avoid repetition of requests in the several log files named in the <code><log></code> parameter.
<code>--postlude <log_file_path></code>	Specify a <code><log_file_path></code> of the file with URLs to replay after those of the <code><log></code> parameter, for each iteration. Use this flag together with the <code>--list</code> flag to avoid repetition of requests in the several log files named in the <code><log></code> parameter.
<code>--progress</code>	Display the percentage of the query log file processed.  Note: If you run Eneper in the two-stream mode for testing updates performance, it displays the progress only for the regular queries log, not for the updates log.
<code>--quionerror</code>	Terminate the Eneper process if it encounters a fatal HTTP connection error. By default, errors are ignored and do not stop the Eneper run.
<code>--rcvbuf <size_bytes></code>	Override the default TCP receive buffer size, set with the <code>SO_RCVBUF</code> socket option.
<code>--record <rec_file_prefix></code>	Record a log of all HTTP responses. Recorded data is placed in output files with the prefix <code><rec_file_prefix></code> . Data files are given the suffixes <code>.dat1</code> , <code>.dat2</code> , and so on. An index file with the suffix <code>.idx</code> is also produced.

Setting	Description
--record_hdr	In --record mode, record HTTP header information along with page content.
--record_ord	In --record mode, ensure that log entries are recorded in the same order that they are listed in the <log> file, even if they are processed out of order.
--record_roll <max_KB>	Set the maximum number of KB per recording file. Default is 1024 KB.
--reqstats	Maintain and report per-request timing statistics.  Note: This option produces accurate results only if you specify <num connections> as 1.
--reqtimeout	Places a limit on the time for any individual request. Default is 600 seconds.
--runtime <max_runtime>	Place a limit on the run time for Eneper. Eneper exits after <max_runtime> minutes. Minutes are the default unit.
--seek <n>	Skip a specified number of requests in the specified log file and start with log entry n. For example, in a log containing 100 requests, if you run Eneper with --seek 50, it issues 50 requests from 50 to 100.
--seekrepeat	Use in conjunction with --seek. Start each iteration with the log entry specified by --seek. --seekrepeat has an impact only if the number of iterations specified is greater than one. If it is so, when Eneper reaches the end of the log file, --seekrepeat indicates that it should start the next iteration from the log entry specified as a value to --seek (50 in the example above). The behavior without --seekrepeat and with --seek specified is to seek only on the first iteration and restart from the beginning of the file on subsequent iterations.
--sleeponerror <secs>	Sleep for a specified number of seconds before sending any new requests after a connection error occurs.
--stats <num_reqs>	Print statistics after the specified <num_reqs> are processed (sent and received).
--throttle <max_req/sec>	Place an approximate limit on the number of requests per second that Eneper generates.
--updates-log	Specifying the updates log allows running Eneper in a two-stream mode with two logs: regular query request logs and update request logs. In this mode, Eneper sends update requests from the updates log at regular intervals while sending queries from the query log. This setting can be used either together with the --msec-between-updates setting, or without it:

Setting	Description
	• If this setting is used together with <code>--msec-between-updates</code>, it specifies the updates log file that contains partial update requests. These requests are replayed at every interval in milliseconds specified with <code>--msec-between-updates</code>. • If this setting is used without <code>--msec-between-updates</code>, updates are sent to the Dgraph one after another, that is, Eneperft waits for the current update to finish and immediately sends another update. It does not wait for any period of time between sending individual updates to the Dgraph. This setting must not be used together with <code>--list</code>, <code>--seek</code>, <code>--seekrepeat</code>, <code>--prelude</code>, <code>--postlude</code>, and <code>--throttle</code>. Before running Eneperft in the two-stream mode, you need to create a separate log that contains only partial update requests. You should create such a log with several partial update requests pointing to a single update file using the <code>admin?op=update&updatefile=filename</code> command. For more information on running partial updates on a single file, see the <i>Partial Updates Guide</i>.
<code>--version</code>	Add the version of Eneperft that is used for this iteration. The version information is always displayed at the beginning of Eneperft output, as follows: <code>Endeca eneperft version <number></code>.
<code>--warn <max_req_threshold></code>	Print a warning message for any requests that take longer than the specified threshold time limit to return (useful for finding the “slow” requests in a log file). The threshold time limit is specified in milliseconds.

About generating incremental statistics

You use the `--stats` setting to specify how many queries you want to see statistics reported on. Typical values are 500 or 100. The `--reqstats` setting provides a finer level of detail.

Generating statistics on the fly

Eneperft can run for hours. If you neglected to set `--stats` yet want to obtain a statistics printout without stopping the process, you can send Eneperft a `usr1` signal.

For example, on UNIX, you could use the `kill` command to send a signal like this:

```
kill -usr1 pid
```

About setting the number of queries sent to the Dgraph

By default, Eneperft drives load as fast as the MDEX Engine can handle it. However, there is a setting, `--throttle`, that allows you to place an approximate limit on the number of queries per second sent to the MDEX Engine. That means you can drive load at a rate you select.

The `--throttle` setting is useful when you want to approximate a special case. For example, imagine you expect high-traffic load during the holiday season. You want to calculate maximum load, while maintaining a comfortable margin of error for the MDEX Engine by running it at 80% utilization.

You might prepare an estimate by multiplying the maximum load by 0.8. Alternatively, you could use `--throttle` to try different numbers of queries per second and to capture the CPU performance on the MDEX Engine machine, using a tool such as `vmstat` on Solaris. You could then calculate the average CPU utilization from these numbers, or plot a chart of utilization over time in Microsoft Excel.

The mapping of the `--throttle` setting to queries per second is not exact. Eneperf uses a simple method to calculate the waiting times to insert between queries. You get a real number of operations per second but it might be significantly lower than you want or expect. The `--throttle` setting to Eneperf can generate performance results that exceed the maximum throughput of the MDEX Engine and still result in throughput results for the MDEX Engine that are less than its maximum. Experiment with this setting to identify the best strategy for your situation.

Example of Eneperf output

This topic contains an example of Eneperf output and describes it briefly.

```
Running iteration 1...
Done:
58881 sent, 58881 received, 0 errors.
22 minutes, 42.63 seconds for iteration 1, 43.2112 req/sec.
22 minutes, 42.63 seconds elapsed (user: 6.20 seconds, system: 15.24 seconds).
Net: 1.18389e+06 KB (868.829 KB/sec).
Page Size: avg=91.34 KB, std dev=142.81 KB, max=1238.37 KB, min=0.16 KB.
Latency: avg=92.36 ms, std dev=238.27 ms, max=13441.11 ms, min=0.18 ms. 250 queries longer than 1s.
Eneperf completed:
58881 sent, 58881 received, 0 errors.
22 minutes, 42.63 seconds elapsed (user: 6.20 seconds, system: 15.24 seconds).
Net: 1.18389e+06 KB (868.829 KB/sec).
Page Size: avg=91.34 KB, std dev=142.81 KB, max=1238.37 KB, min=0.16 KB.
Latency: avg=92.36 ms, std dev=238.27 ms, max=13441.11 ms, min=0.18 ms. 250 queries longer than 1s.
    Best iteration time: 22 minutes, 42.63 seconds.
    Peak rate: 43.2112 req/sec.
    Avg iteration time: 22 minutes, 42.63 seconds.
    Avg rate: 43.2112 req/sec.
    Total rate: 43.2112 req/sec.
```

The entries from Eneperf output are described in the following table:

Sample Eneperf output entry	Description
Running iteration 1...	Is printed as each iteration begins. The numbers following this line, until "Eneperf completed:" occur for each iteration requested. The number of iterations requested is the last Eneperf parameter.
done:	Is printed once the iteration finishes.

Sample Eneper output entry	Description
58881 sent, 58881 received, 0 errors.	"Sent" is the number of queries sent. It is the sum of "Received" and "Errors" and the number of errors, where errors is the number of 404 or 400 HTTP codes that the Dgraph returns, (rather than errors in the Dgraph log). "Received" is the number of queries with a 200 HTTP status code that the Dgraph returns. "Errors" is the number of queries with 404 or 400 HTTP status code that the Dgraph returns, rather than errors in the Dgraph log.
22 minutes, 42.63 seconds for iteration 1, 43.2112 req/sec.	The time for the specific iteration, and the throughput for this iteration.
22 minutes, 42.63 seconds elapsed (user: 6.20 seconds, system: 15.24 seconds).	The total runtime up until this point. System time is the time spent in the operating system on behalf of the Dgraph. User time is the time spent in the Dgraph itself.
Net: 1.18389e+06 KB (868.829 KB/sec).	The total amount of data returned for the entire test (not just for one iteration).
Page Size: avg=91.34 KB, std dev=142.81 KB, max=1238.37 KB, min=0.16 KB.	Cumulative statistics on the amount of data returned for each query.
Latency: avg=92.36 ms, std dev=238.27 ms, max=13441.11 ms, min=0.18 ms. 250 queries longer than 1s.	Cumulative statistics on the latencies. The statistics include previous iterations. Latency information may be inaccurate when multiple connections are in use, particularly if the network is slow. If accuracy is critical, consider obtaining latency information from the Dgraph request log.
Peak rate: 43.2112 req/sec.	The processing rate of the iteration in the test with the best performance, but should not be confused with "peak" performance in the sense of a single second that showed the highest throughput. It is the total number of requests processed in that iteration divided by the time of the iteration in seconds. If the test includes only one iteration, peak rate is the processing rate for that iteration.
Avg iteration time: 22 minutes, 42.63 seconds.	The average time of the iterations in the test.

Sample Eneper output entry	Description
Avg rate: 43.2112 req/sec.	The average rate of the iterations in the test, in requests processed per second.
Total rate: 43.2112 req/sec.	The total of requests processed for all of the iterations in the test, divided by the total time of all of the iterations in the test.
Eneper completed:	All information after this statement is cumulative over the entire run. This line is printed once all iterations have completed.

About the format of logs for use with Eneper

In order to use Eneper, you need a log of URLs in the correct format. The lines in the log file you use with Eneper should not specify the run-time statistics, hostname and the port.

There are numerous ways that you can obtain such logs; this section provides you with guidelines and a few examples.

The Request Log Parser

In order to use Eneper, you need a log of URLs in the correct format. The Request Log Parser is an Endeca utility that converts the MDEX Engine log format into Eneper log format.

Alternatively, you can convert URLs yourself. For more information, see [Converting a MDEX Engine request log file for Eneper](#).

Recommendations for generating a representative log for Eneper

The test log that you will use with Eneper determines the contents and the results of your performance testing. Because the test log serves as input to Eneper, it should be representative of those aspects of the MDEX Engine performance you want to test.

Use these recommendations to create a representative log:

- Add queries of various types to your log to account for a variety of queries. Depending on the query type, some queries are processed much faster than others.

For example, dimension and record search queries are the fastest, queries on aggregate records, or navigation and search queries take longer, whereas navigation with Analytics, or navigation queries with RRN may take more time. Even within queries of the same type, individual queries can have large performance differences, depending on the query parameters.

- If you want to test a particular feature configuration for performance, ensure that your query log contains a fair percentage of queries of this type.
- If you are planning to test updates that run at regular intervals, create a separate updates log from your regular log that contains only `config?op=update` operations, and run Eneper against this updates log and the regular log at the same time. Use the `--updates-log` setting together with `log` and `--msec-between-updates` settings.

- If queries are repeated in the log, or parts of them are repeated, this makes the log less useful for performance testing, since a large percentage of queries may be served entirely from the MDEX Engine cache. Therefore, do not replay a short query log multiple times.
- For a full-scale performance test, generate a log that runs for 30 minutes or more. In addition, you may want to create a smaller log that runs for 5-10 minutes to use it as a quick test.
- To create a representative log, use the existing MDEX Engine logs from the production system. Use the Request Log Parser to strip undesired columns and queries. For information, see “The Request Log Parser”.
- Translate existing Web application logs into the MDEX Engine format. For example:

```
/results.jsp?searchterm=ipod
```

turns into:

```
graph?node=0&group=0&offset=0&nbins=10&attrs=All+ipod|mode+matchall&dym=1
```

- Translate existing traffic reports, such as a list of top search terms, into the MDEX Engine format by programmatically generating URLs as produced by the MDEX Engine. For example, for the term “iron man”, generate:

```
graph?node=0&group=0&offset=0&nbins=10&attrs=All+iron+man|mode+matchall&dym=1
```

- Use the Request Log Parser to remove all admin queries from a request log (use the default or -q gb options for the parser). Typically, process health requests of type /admin?op=ping can run every few seconds, are typically very fast and not generated by end users. However, requests of type /admin?op=exit stop and restart the process and will impact your log.
- Remove dimension search queries from your Eneper log. This is because a single API request that includes a dimension search is turned into two MDEX Engine requests. For example, the following request:

```
?N=0&Ntk=All&Ntt=plum&Nty=1&D=plum
```

turns into:

```
/graph?node=0&group=0&offset=0&nbins=10&attrs=All+plum  
/search?terms=plum&rank=0&offset=0&compound=1
```

From the application perspective, this request constitutes one query, since the presentation API waits for both responses and recombines them into a single response object to the front-end application. However, the MDEX Engine and performance tools, such as Eneper and the Request Log Analyzer, treat such dimension search requests as two queries.

If you remove these dimension search queries, which are known to be fast, from the Eneper log and replace them with other queries, you can use Eneper to measure the MDEX Engine performance against this log. If the desired level of performance is achieved with such a log, you will achieve or exceed that performance when dimension searches are included again.

Running Eneper in two-stream mode: regular logs and logs with updates

You can run Eneper in a two-stream mode using two streams of request logs — regular query request logs and logs that contain partial update requests. This lets you test MDEX Engine performance with partial updates applied at regular intervals while running a regular query load.

To run Eneper in the two-stream mode, use the following Eneper settings together:

- --updates-log
- --msec-between-updates

- `--log`

When used in this mode, Eneper sends update requests from the updates log at regular intervals while sending queries from the query log.

In more detail, Eneper runs in the following way:

1. It uses the log file (specified with `--log`) and sends requests from this file for the duration that you specify by the `--msec-between-updates` setting.
2. At the specified time interval, it sends an update request from the updates log file (specified with `--updates-log`) and uses one of its connections for this request.
3. It continues to send query requests from the query log (`--log`), using the other connections.



Note: This behavior assumes that you are running Eneper with the number of connections set to more than one. If you use only one connection, Eneper will switch between update and regular query requests.

4. This process continues until either the regular query log or the updates log has been completely processed. For example:
 - If Eneper sends the last update request from the updates log, but the query log still contains queries, Eneper will send additional queries for the time interval specified with `--msec-between-updates` and then stop. (Since the two-stream mode is designed specifically to test updates performance, Eneper does not process regular queries after the last update in the updates log has been processed.)
 - If Eneper sends the last query from the regular log, but the updates log still contains additional update requests, it will not send these updates to the Dgraph. Therefore, ensure that the regular query log contains sufficient number of requests to last for the duration of your two-stream Eneper testing session.

The format of the updates request log is the same as the format of a regular query log for Eneper, except that the updates log should contain only `config?op=update` operations in order to provide meaningful performance results. (If your updates log contains regular queries, Eneper still processes this log successfully. However, the results are not meaningful for measuring updates performance.)

Using `--updates-log` and `--log` settings is useful to measure performance of those updates that run at regular intervals. To test updates that run at random times, you can continue using your regular log with Eneper.



Note: The actual time interval between sending update requests may be equal to or greater than the time specified with `--msec-between-updates`. This is because Eneper uses the same `num_connections` setting while processing the regular query log and updates log. This causes Eneper to wait for a preceding request to complete before it can process the next updates log request.

Before running Eneper in the two-stream mode, you need to create a separate log that contains only partial update requests. You should create such a log with several partial update requests pointing to a single update file using the `admin?op=update&updatefile=filename` command. For more information on running partial updates on a single file, see the *Partial Updates Guide*.



Note: The `--msec-between-updates` flag is optional. In other words, if you only specify the `--updates-log` flag, the updates are sent to the Dgraph one after another. Eneper waits for the current update to finish and immediately sends another update. It does not wait for any period of time between sending individual updates to the Dgraph.

Converting an MDEX Engine request log file for Eneperf

In order to use Eneperf, you need a log of URLs in the correct format. You can manually convert the log to the desired format, or use the Request Log Parser.

The lines in the log file you use with Eneperf should not specify the run-time statistics, hostname and the port. For example, raw URL requests could be formatted like these:

```
/search?terms=blackberry&rank=0&opts=mode+matchall&offset=0&compound=1
&irversion=510
/graph?node=0&group=10&offset=0&nbins=10&attrs=All+berry|mode+matchall
&dym=1&irversion=510
```

To convert a complete MDEX Engine request log file for Eneperf use:

Run the following command:

```
sed -e '/DGRAPH STARTUP/d' <logfile> |
sed -e '/\admin.*$/d' |
cut -d ' ' -f 12-
```

This does the following:

- It deletes DGRAPH STARTUP lines, because these lines contain no commands.
- It removes admin requests, such as admin?op=stats or admin?op=exit, that can cause problems in an Eneperf run.
- It obtains the last three columns in the log (the URL, POST body, and HTTP headers).

Performance testing .NET 2.0 applications that contain long or complex queries

In rare cases, if your .NET 2.0 (or later) application uses very complex record filters or Analytics statements, you may find that your Eneperf results differ from what is seen in production.

This discrepancy results from the way the .NET 2.0 API to the MDEX Engine handles very long or complex queries. Instead of the usual HTTP GET request to the MDEX Engine, it uses an HTTP POST request. However, the MDEX Engine logs the query as if it were a GET request. The different processing and validation that occurs for POST requests may result in performance differences.

To better simulate the performance of applications that contain such queries, you can use the Request Log Parser to pre-process the logs used to run the Eneperf test. For each request in the log that is longer than 65,000 characters, prepend '/graph' with a space after it to the request. Use the subsequent log as the input to Eneperf.



Note: This behavior only manifests itself in the case of very long or complex queries. Most applications never use queries of this sort.

Creating a log file by hand using substitute search terms

You can also approximate a log file to be used with Eneperf. This method is useful when you do not have a running MDEX Engine and archives of logs to work with.

For example, you may want to test the performance of search terms culled from some other system.

To create a log file by hand:

1. Create a list of search terms that you want to test.
2. Copy or create a URL and optional HTTP POST body in the appropriate format.
3. Compose a new log file by substituting your search terms into URL requests containing suitable options.

Debugging Eneper

Eneper generates error messages in various error conditions.

- If you make an error while typing the command line argument, Eneper returns its help message.
- if you accidentally mistype the MDEX Engine port, Eneper generates numerous failed connection error messages.
- If Eneper encounters socket connection errors, it reports error messages.

It is also possible for error messages to be displayed during normal operation. For example, if the log file contains a request to retrieve a record that is not present in the MDEX Engine data set, Eneper (as expected) presents a 404 (file not found) message.



Note: Queries that cause HTTP errors are not counted towards ops/sec performance results displayed by Eneper.

Appendix E

Using the Request Log Analyzer

The Request Log Analyzer is a performance testing tool that is included in your Endeca installation.

About the Request Log Analyzer

This tool simplifies and standardizes forensic analysis of Endeca MDEX Engine performance. The Request Log Analyzer reads one or more MDEX Engine logs and reports on the nature and performance of the queries recorded in those logs. This kind of analysis is called "forensic" because it reports on what actually happened in the past, instead of reporting on potential performance or capacity planning for the future.

There are two main applications for this script. First, it can and should be run manually in order to debug performance problems. Second, it can and should be run on a regular basis, either standalone or as part of a control script, in order to continually monitor performance and call out trends in Dgraph traffic load, latency, throughput, and application behavior. The default behavior of this script (without flags) is meant to be sufficient for daily or weekly reports, while the options available to the developer via flags are meant to give enough flexibility and power to perform serious debugging.

Installation location

The Request Log Analyzer is installed as `$ENDECA_MDEX_ROOT/bin/reqloganalyzer` (UNIX) and `%ENDECA_MDEX_ROOT%\bin\reqloganalyzer.exe` (Windows).

Log format requirements

The Request Log Analyzer supports request logs generated by IAP 4.x, IAP 5.x and MDEX Engine 6.x.

In order to efficiently process large volumes of log files, supply log files (and the entries they contain) to its command-line in date-time increasing order without overlap.

Invoking the Request Log Analyzer

The Request Log Analyzer is invoked as follows:

```
reqloganalyzer [--flags] dgraph.log|- [dgraph2.log ...]
```

where `dgraph.log` is a relative or absolute path to the `dgraph` log file to analyze. Multiple `Dgraph` logs may be analyzed by listing all logs, separated by spaces. To process log file data from the standard input, specify the file name as a single hyphen: `-`.

The Request Log Analyzer will write its results to standard output. To save results to a file, simply redirect standard output to a file.

The available flags for the Request Log Analyzer are detailed in the following topics. The flags are:

Show Flags:

- showHourly
- showProfiles
- showResponseCodes
- showRequestTypes
- showExtendedTypes
- showThreading
- showWorstEngines
- showWorstResponses
- showAll
- verbose
- numWorstEngines
- numWorstResponses

Threshold Flags:

- threshEngine
- threshResponse
- threshResponseDiff
- threshResponseSize
- threshQueueLength
- threshBinsRequested
- threshAggrBinsRequested
- threshBinsReturned
- threshAggrBinsReturned
- threshOffset
- threshNumNavDescriptors
- threshNumExposedRefinements
- threshNumSortKeys
- threshNumSearchTerms
- threshNumSearchKeys

Ignore Flags:

- ignoreAdmin
- ignoreInvalid
- ignoreEmpty
- ignoreErrors
- ignore

Timeframe Flags:

- timeLower
- timeUpper
- hourOffset

Miscellaneous Flags:

- precision
- help, -?

Show flags

By default, the Request Log Analyzer outputs a small number of widely applicable metrics, such as the average response time for all requests. There are additional metrics possible to display using the Request Log Analyzer; these flags toggle the calculation and display of these additional metrics. Note that enabling additional metrics can slow analysis time.

- `--showHourly`: calculates and outputs running statistics for each hour time span within the log as well as tracking the best-performing hour. This flag is useful for its statistics as well as for providing visual feedback that the script is actively running.
- `--showProfiles`: calculates and outputs detailed statistics about the nature of the request and response, such as the number of search terms or the number of sort keys. This flag is especially calculation-intensive and will notably slow analysis time.
- `--showResponseCodes`: calculates and outputs information about the performance of queries categorized by their HTTP response codes.
- `--showRequestTypes`: calculates and outputs statistics about the performance of queries categorized by their query type (search, navigation, record request, etc.)
- `--showExtendedTypes`: calculates and outputs statistics about the performance of queries categorized by their utilization of specific query features, such as wildcard searches, boolean matchmode, record filters, geocode filters, etc.
- `--showThreading`: calculates and outputs statistics about the behavior of a multithreaded MDEX Engine, such as average request queue length, average number of idle threads, and performance analysis of queued vs. unqueued queries.
- `--showWorstResponses`: calculates and outputs the N longest-running queries, based on round-trip response time. By default, N is 10 but is configurable using `--numWorstRequests`.
- `--showWorstEngines`: calculates and outputs the N longest-running queries, based on engine-only processing time. By default, N is 10 but is configurable using `--numWorstDgraphs`.
- `--verbose`: calculates and outputs all available statistics except request/response profiling. This flag is a shortcut and is equivalent to specifying `--showHourly --showResponseCodes --showRequestTypes --showSpecialSearches --showThreading --showWorstRequests --showWorstDgraphs`.
- `--showAll`: calculates and outputs all available statistics. This flag is a shortcut and is equivalent to specifying `--showHourly --showProfiles --showResponseCodes --showRequestTypes --showSpecialSearches --showThreading --showWorstRequests --showWorstDgraphs`. Since this flag includes `--showProfiles`, it is especially calculation-intensive and will notably slow analysis time.
- `--numWorstResponses`: specifies the number of longest-running queries to calculate and output, based on round-trip response time. This flag is only useful when `--showWorstRequests` is also enabled.
- `--numWorstEngines`: specifies the number of longest-running queries to calculate and output, based on engine-only processing time. This flag is only useful when `--showWorstDgraphs` is also enabled.

Threshold flags

The Request Log Analyzer includes functionality to report on the number of requests that exceed a threshold. For instance, by default the Request Log Analyzer reports the number of requests that took longer than 1.25 seconds total round-trip response time. The following threshold flags allow the user to specify the exact setting to use as a threshold for many metrics.

- `--threshEngine`: specifies the threshold for engine-only processing time, in milliseconds. The default is 500.

- `--threshResponse`: specifies the threshold for round-trip response time, in milliseconds. The default is 1250.
- `--threshResponseDiff`: specifies the threshold for response differential time (the difference between round-trip response time and engine-only processing time), in milliseconds. The default is 500.
- `--threshResponseSize`: specifies the threshold for response size, in bytes. The default is 393216, which is equivalent to 384K.
- `--threshQueueLength`: specifies the threshold for the number of queued requests. This metric is only calculated when `--showThreading` is enabled and is only valid for multithreaded MDEX Engines. The default is 5.
- `--threshBinsRequested`: specifies the threshold for the number of base records requested by the Endeca Presentation API. This metric is only calculated when `--showProfiles` is enabled. The default is 50.
- `--threshAggrBinsRequested`: specifies the threshold for the number of aggregate records requested by the Endeca PresentationAPI. This metric is only calculated when `--showProfiles` is enabled. The default is 50.
- `--threshBinsReturned`: specifies the threshold for the total number of base records found by the MDEX Engine (not the number returned in a single page). This metric is only calculated when `--showProfiles` is enabled. There is no default.
- `--threshAggrBinsReturned`: specifies the threshold for the total number of aggregate records found by the MDEX Engine (not the number returned in a single page). This metric is only calculated when `--showProfiles` is enabled. There is no default.
- `--threshOffset`: specifies the threshold for the pagination offset requested by the Endeca Presentation API. This metric is only calculated when `--showProfiles` is enabled. The default is 100.
- `--threshNumNavDescriptors`: specifies the threshold for the number of dimension values specified as descriptors by the Endeca API, not including the default root node (N=0). This metric is only calculated when `--showProfiles` is enabled. There is no default.
- `--threshNumExposedRefinements`: specifies the threshold for the number of open refinement dimensions specified by the Endeca Presentation API, not including requests for no open refinements or all open refinements. This metric is only calculated when `--showProfiles` is enabled. There is no default.
- `--threshNumSortKeys`: specifies the threshold for the number of explicit sort keys specified by the Endeca Presentation API. This metric is only calculated when `--showProfiles` is enabled. The default is 2.
- `--threshNumSearchTerms`: specifies the threshold for the total number of search terms specified the Endeca Presentation API, across all search keys. Note that this metric is approximate because of the variety of punctuation and search characters possible. This metric is only calculated when `--showProfiles` is enabled. The default is 6.
- `--threshNumSearchKeys`: specifies the threshold for the total number of search keys specified the Endeca Presentation API. This metric is only calculated when `--showProfiles` is enabled. The default is 3.

Ignore flags

By default, the Request Log Analyzer reports on all requests within a logfile. When performing analysis, it is sometimes useful to only report on certain types of requests. Excluding some requests provides a truer overall picture of the performance of the remaining queries, but can skew overall statistics. For instance, by excluding admin requests, the reports on average response size are more useful when

analyzing application-level query tuning. However, the reports on queue length can be misleading, since admin requests utilize a thread and contribute to queue length.

- `--ignoreAdmin`: excludes administrative and configuration requests (`/admin` and `/config`) from statistical analysis, though the Request Log Analyzer will still report on the number of admin requests found. Administrative requests signal the MDEX Engine to load partial updates, load new thesaurus entries or dynamic business rules, output the MDEX Engine's internal stats page, execute a health check against the ping page, and perform other administrative functions.
- `--ignoreInvalid`: excludes invalid requests from statistical analysis, though the Request Log Analyzer will still report on the number of invalid requests found. Invalid requests are those requests that cannot be handled by the MDEX Engine, such as a request for `/foo`, but do not include empty requests (see `--ignoreEmpty`).
- `--ignoreEmpty`: excludes empty requests from statistical analysis, though the Request Log Analyzer will still report on the number of empty requests found. Empty requests are those requests for the URL `""` (the empty string). Empty requests are sometimes generated by load balancer health-checks and can also be generated by telnet-ing directly to the MDEX Engine's port without issuing any further commands.
- `--ignoreErrors`: excludes error requests from statistical analysis, though the Request Log Analyzer will still report on the number of error requests found. Error requests are those that resulted in anything other than a 200 (OK) HTTP status code. These can be generated by a request for an unknown dimension value (HTTP status code 404) or by a request whose client was disconnected before the response could be written (HTTP status code 408).
- `--ignore`: excludes admin requests, invalid requests, empty requests, and error requests from statistical analysis. This flag is a shortcut and is equivalent to specifying `--ignoreAdmin --ignoreInvalid --ignoreEmpty --ignoreErrors`.

Timeframe flags

By default, the Request Log Analyzer reports on all requests within a logfile, and all time-based calculations such as operations per second are based on the time period between the first request in the log and the last request in the log. It is often useful to only analyze a certain time period within the log. This allows targeted analysis of specific events such as a load test or a traffic spike. Furthermore, it allows outlying requests such as developer-initiated requests before an MDEX Engine is released to production traffic to be excluded from calculation.

- `--timeLower`: the earliest datetime that should be analyzed. Requests that occurred before this datetime will be silently ignored. This datetime can be specified as in epoch time or in `YYYY-MM-DD-HH-MM-SS` format. Note that hours must be specified in 24-hour military time. To specify a lower time bound of December 14, 2005, 5:26:38 PM, use `--timeLower 1134599198` or `--timeLower 2005-12-14-17-26-38`.
- `--timeUpper`: the latest datetime that should be analyzed. Requests that occurred after this datetime will be silently ignored. This datetime can be specified as in epoch time or in `YYYY-MM-DD-HH-MM-SS` format. Note that hours must be specified in 24-hour military time. To specify an upper time bound of December 14, 2005, 5:26:38 PM, use `--timeUpper 1134599198` or `--timeUpper 2005-12-14-17-26-38`.
- `--hourOffset`: the difference, in hours, between the timezone of the server that created the MDEX Engine log and the server running the Request Log Analyzer. This timezone difference is important because the Request Log Analyzer will output time information in human-readable format even though the MDEX Engine logs times in epoch format. The translation of epoch to human-readable time honors the timezone of the server running the Request Log Analyzer. The `--hourOffset` flag allows human-readable times to honor the timezone of the server that wrote the log. If the server running the Request Log Analyzer is in EST (GMT-5) and the server that

wrote the log is in PST (GMT-8), specify `--hourOffset -3`. This flag also affects the translation of human-readable time to epoch time, when specifying `--timeLower` or `--timeUpper` in YYYY-MM-DD-HH-MM-SS format.

Miscellaneous Flags

- `--precision`: controls the number of significant decimal digits displayed in calculated statistics. The default is 3.
- `--help`: outputs help and usage information and then exits without performing any analysis.

Interpreting reports

Achieved vs. Potential Performance

The Request Log Analyzer measures the performance actually achieved by an MDEX Engine according to that engine's request logs. This "achieved performance" is completely dependent on the amount and nature of traffic sent to the MDEX Engine, and does not measure the capacity or upper bounds of performance that the MDEX Engine is actually capable of - the "potential performance".

For instance, consider an MDEX Engine running the reference sample wine dataset on a modern server and the latest release of Oracle Endeca Guided Search. This MDEX Engine is capable of handling well over 100 ops/second throughput at sub-second response times. Now consider that a single user leisurely clicks through the dataset, stopping to read about the descriptions and flavors of the featured wines. This user will generate a total of 45 requests over a 15 minute time span. When the Request Log Analyzer analyzes the logfile for this single user, it will report an achieved throughput of 0.05 operations per second (45 requests / (15 minutes * 60 seconds)), compared to the known potential throughput of 100+ operations per second.

Because of the possible large differences in achieved vs. potential performance, the Request Log Analyzer, as a standalone tool, is more suited to forensic analysis and behavior profiling than load testing and capacity planning. However, the Request Log Analyzer does work well in concert with other load testing tools such as Eneperfto analyze the performance characteristics of an MDEX Engine under load.

Expensive Features

Usage of expensive features such as wildcard search, Boolean search, exposure of all refinements, or large numbers of records per page is a common cause of performance problems.

The Request Log Analyzer can be used to measure the performance of queries utilizing expensive features against "standard" queries. Because the Request Log Analyzer can segment performance numbers based on these features (see Extended Query Types and Request Profiling), it is trivial to compare statistics. However, beware that statistics from a single MDEX Engine necessarily are inter-dependent. A standard query may be stuck in the request queue behind an expensive query, thus inflating its response time. Similarly, a standard query may be executing on a thread simultaneously with an expensive query on a parallel thread. If the expensive query causes resource (CPU, disk) contention, the standard query will see an inflated engine time.

Request Queuing and Multithreading

Excessive request queuing is another common cause of performance problems. The Request Log Analyzer can be used to detect the presence of excessive request queuing and can report on the performance of queued requests vs. requests that encountered no queue.

Request queuing is, in and of itself, not necessarily a bad thing. MDEX Engines are often more efficient at processing a small number of simultaneous requests quickly and then moving on to process requests that have been waiting in queue. In this model, because each individual request is processed quickly, requests are only in the queue for a very short time, and overall performance is good.

When the request queue gets very long, or when requests have to wait a long time in queue, the queue is a problem. In this situation, additional engine threads can help if the server has enough available resources (CPU, RAM, disk). Additional MDEX Engines in a load-balanced configuration will also help.

Note that the response differential metric - the difference between round-trip response time and engine-only processing time - includes time spent in both the request queue (waiting to be processed) and the response queue (waiting to be written back to the client). It is not possible to determine from the request logs exactly how long requests spend in the request queue alone.

Statistics

The Request Log Analyzer outputs a large amount of statistics for analysis. The types of statistics are explained in the following topics.

Common metrics

The following metrics are found in multiple sections of the Request Log Analyzer output:

- **Ops/Sec:** operations per second; the number of requests processed in a time period, divided by the number of seconds in that time period. This metric provides a sense of achieved throughput. In other words, this metric shows how well the MDEX Engine is servicing simultaneous clients. Note that this is achieved performance, not potential performance; see Result Interpretation.
- **Round-Trip Response Time or Response Time:** the total time required for the Endeca API to receive a response back from an MDEX Engine. This includes time spent in the request queue, time spent processing in the engine, and time spent writing the response back to the API. The round-trip response time is the time seen by end users of the system.
- **Response Times Over [threshold]:** the number of requests that took longer than [threshold] to complete, measured by round-trip response time. The [threshold] value is configurable with the `--threshResponse` flag. Typically, this metric is followed by the percentage of over-threshold requests compared to all requests.
- **Engine-Only Processing Time or Engine Time:** the total time required for the Endeca MDEX Engine to calculate the results of a query. This includes only the time spent processing in the engine, excluding time spent in the request queue and time spent writing the response back to the API. The engine-only processing time is a measure of the expensiveness of any particular query.
- **Engine Times Over [threshold]:** the number of requests that took longer than [threshold] to complete, measured by engine-only processing time. The [threshold] value is configurable with the `--threshEngine` flag. Typically, this metric is followed by the percentage of over-threshold requests compared to all requests.
- **Response Size:** the size of the data packet returned by the MDEX Engine to the API. As data packet sizes grow larger, more network resources are required to move the data. Additionally, larger data packets typically require more computing resources in the MDEX Engine and the Endeca API to pack and unpack the data.
- **Number of Requests:** when categorizing metrics by type, such as when using `--showResponseCodes` to display 404 Not Found responses, Number of Requests denotes the total quantity of requests in the category. This metric is typically followed by the percentage of these requests compared to all requests.

- **Requests Analyzed:** when calculating individual metrics, such as when using `--showProfiles` to display the number of base records requested, Requests Analyzed denotes the number of requests polled to calculate the metric.

Hourly results

These statistics are only available when the `--showHourly` flag is enabled.

The following sample may be line-wrapped:

DATE TIME	ROUND-TRIP OVER 750	NUM REQUESTS 1250	OPS/SEC AVG ENGINE TIME	AVG ROUND-TRIP ENGINE TIME OVER
----	-----	-----	-----	-----
2005-12-19 14:00-15:00	1 / 1	1.000 / 0.000	0.000 / 0.000	0 (0.00%) / 0.000
0 (0.00%) /	(0.00%)	0.000 / 0.000		0 (0.00%) / (0.00%)
2005-12-19 15:00-16:00	12890 / 12891	3.786 / 0.000	303.924 / 303.900	
325 (2.52%) / 325 (2.52%)	109.311 / 109.303	343 (2.66%) / 343		
(2.66%)				
2005-12-19 16:00-17:00	14169 / 27060	3.936 / 0.000	304.759 / 304.350	
375 (2.65%) / 700 (2.59%)	107.894 / 108.565	327 (2.31%) / 670		
(2.48%)				
2005-12-19 17:00-18:00	12182 / 39242	3.384 / 0.000	295.993 / 301.756	
292 (2.40%) / 992 (2.53%)	92.452 / 103.563	232 (1.90%) / 902		
(2.30%)				
2005-12-19 18:00-19:00	11189 / 50431	3.108 / 0.000	291.157 / 299.404	
286 (2.56%) / 1278 (2.53%)	86.383 / 99.751	164 (1.47%) / 1066		
(2.11%)				

Hourly statistics, when enabled, are output as they are calculated by the Request Log Analyzer. This is the only statistic to be output in this manner; all other statistics are only output when the Request Log Analyzer has finished processing all logs. For this reason, hourly statistics are useful to know that the Request Log Analyzer is continuing to run properly.

The metrics output by hourly statistics are:

1. **Date:** The timespan of the current hour. The Request Log Analyzers hourly statistics are measured by single hour boundaries; this is not configurable.
2. **Num Requests:** The number of requests found in the current hour, followed by the number of requests processed so far.
3. **Ops/Sec:** the achieved throughput for the current hour, followed by the achieved throughput for all requests processed so far.
4. **Avg Round-Trip Time:** the average round-trip response time for the current hour, followed by the round-trip response time for all requests processed so far.
5. **Round Trip Over [threshold]:** the number of requests in the current hour requiring longer than [threshold] to complete, based on round-trip response time, followed by the number of requests over [threshold] for all requests processed so far.
6. **Avg Engine Time:** the average engine-only processing time for the current hour, followed by the engine-only processing time for all requests processed so far.
7. **Engine Time Over [threshold]:** the number of requests in the current hour requiring longer than [threshold] to complete, based on engine-only processing time, followed by the number of requests over [threshold] for all requests processed so far.

Longest-running requests by round-trip response time

These statistics are only available when the `--showWorstResponses` flag is enabled. The `--numWorstResponses` flag controls the number of requests listed.

The following sample may be line-wrapped:

```
=====
== Longest-Running Requests
=====
1. 1733.90 ms:    1129914655 127.0.0.1 33868 1733.90 1711.04 200 29983 2 6
/graph?node=0&allgroups=1&groupby=RollupKey&offset=0&nbins=15&allbins=2&at-
trs=Keywords+COV+HD%2fBT+E%2fW+BLUE+2X+25%2fCS|mode%2bmatchany&dym=1&fil-
ter=4&irversion=460
2. 716.56 ms:    1129914664 127.0.0.1 132043 716.56 344.80 200 468 3 6
/graph?node=0&allgroups=1&groupby=RollupKey&offset=0&nbins=15&allbins=2&at-
trs=All+reinforced+silicon+tubing|mode%2bmatchpartialmax&dym=1&filter=4&irver-
sion=460
3. 698.01 ms:    1129914658 127.0.0.1 312 698.01 4.38 200 0 2 6
/search?terms=Black+&filter=4&rank=0&num=5&offset=0&model=1&irversion=460
```

This section lists the N longest-running requests, measured by the round-trip response time. The rank and round-trip response time are listed, followed by the request in the raw format found in the MDEX Engine logs. This raw format includes all request information request in context, which can be helpful to understand why the request may have been long-running.

Longest-running requests by engine-only processing time

These statistics are only available when the `--showWorstEngines` flag is enabled. The number of requests listed is controlled by the `--numWorstEngines` flag.

The following sample may be line-wrapped:

```
=====
== Longest-Running Dgraphs
=====
1. 525.09 ms:    1129914645 127.0.0.1 32426 545.12 525.09 200 2564 2 5
/graph?node=0&allgroups=1&groupby=RollupKey&offset=0&nbins=15&allbins=2&at-
trs=ProductDescription+Thermometers|mode%2bmatchpartialmax&dym=1&fil-
ter=4&irversion=460
2. 481.29 ms:    1129914646 127.0.0.1 61133 510.62 481.29 200 17903 2 6
/graph?node=0&allgroups=1&groupby=RollupKey&offset=0&nbins=15&allbins=2&at-
trs=All+glass|mode%2bmatchpartialmax&dym=1&filter=4&irversion=460
3. 457.31 ms:    1129914655 127.0.0.1 34035 519.63 457.31 200 1024 2 5
/graph?node=0&allgroups=1&groupby=RollupKey&offset=0&nbins=15&allbins=2&at-
trs=All+Windows|mode%2bmatchpartialmax&dym=1&filter=4&irversion=460
```

This section lists the N longest-running requests, measured by the engine-only processing time. The rank and engine-only processing time are listed, followed by the request in the raw format found in the MDEX Engine logs. This raw format includes all request information in context, which can be helpful to understand why the request may have been long-running.

Query types

These statistics are only available when the `--showQueryTypes` flag is enabled.

```
=====
== Query Types
```

```

=====
----- Navigation Requests -----
Number of Requests:          169 (12.793%)
Avg Response Time (ms):      54.633
Avg Engine Time (ms):        23.086
Avg Response Size (bytes):    14474.402
Request Times over 1250 ms:   0 (0.000%)
Engine Times over 750 ms:     0 (0.000%)

----- Search Requests -----
Number of Requests:          666 (50.416%)
Avg Response Time (ms):      108.178
Avg Engine Time (ms):        58.701
Avg Response Size (bytes):    28186.752
Request Times over 1250 ms:   1 (0.150%)
Engine Times over 750 ms:     2 (0.300%)

```

This section displays statistics categorized by the type of query processed by the MDEX Engine. The sum of the percentages of each query type should equal 100% (accounting for rounding), since any one query can only be of a single type.

The following query types are output within this section:

- **Web Service Requests:** requests that arrive via Web Services.
- **Navigation Requests:** requests that specify a navigation state, but do not include any text search terms, typically using the `N URL` parameter or the `setNavDescriptors()` API method.
- **Search Requests:** requests that specify text search terms, typically using the `Nt t URL` parameter or the `setNavERecSearches()` API method.
- **Record Requests:** requests that return one or more base records by specifying record specs, typically using the `R URL` parameter or the `setERecSpec()` or `setERecs()` API methods.
- **Aggregate Record Requests:** requests that return one or more aggregate records by specifying record specs, typically using the `A URL` parameter or the `setAggrERecSpec()` API method.
- **Dimension Search Requests:** requests that perform dimension search, typically using the `D` parameter or the `setDimSearchTerms()` API method. Note that even though the Endeca Presentation API treats dimension search as an additional feature of a navigation or search query, the MDEX Engine treats it as a separate query. Thus, the Request Log Analyzer reports dimension search requests as distinct queries.
- **Admin Ping Page Requests:** requests for the `/admin?op=ping` built-in MDEX Engine Ping page. This URL is used as a health check page by load balancers and other monitoring tools.
- **Admin Stats Page Requests:** requests for the `/admin?op=stats` built-in MDEX Engine Server Statistics page. This URL is used to monitor performance and characteristics of a running engine.
- **Admin Partial Update Requests:** requests for the `/admin?op=update` built-in partial update command. This URL is used to signal an MDEX Engine to look for and process any available partial update files.
- **Other Admin Requests:** requests for all other `/admin?` URLs. Note that since this is a catchall category, it would include requests for such nonexistent URLs as `/admin?nothing` that result in 404s.
- **Configuration Requests:** requests for all `/config?` URLs, such as the built-in command to signal a MDEX Engine to look for and process new thesaurus entries and dynamic business rules. Note that since this is a catchall category, it would include requests for such nonexistent URLs as `/config?nothing` that result in 404s.
- **Browser Requests:** when a web browser requests a MDEX Engine URL, it may request a favorites icon from `/favicon.ico`. Similarly, if a browser requests the Admin Stats Page, it may also request an accompanying XSLT from `/stats.xslt` for formatting. These two URLs are tracked together as Browser Requests. Note that the Endeca Navigation API never requests these URLs.

- **Empty Requests:** requests for the URL "" (the empty string). Empty requests are generated by telnet-ing directly to a MDEX Engine's port without issuing any further commands. Empty requests are sometimes generated by misconfigured load balancer health-checks. In MDEX Engine versions prior to 5.1, these requests are represented by the empty string in the request log; in versions 5.1 and later, these requests are represented in the request log as "-" (the dash character).
- **Invalid/Undecipherable Requests:** requests for any URL not included above, such as a request for "/index.html" or a request for "xyxyx".

Extended query types

These statistics are only available when the --showExtendedTypes flag is enabled.

```
=====
== Extended Query Types
=====
----- Did-You-Mean Enabled Searches -----
Number of Requests:          663 (50.189%)
Avg Response Time (ms):      107.990
Avg Engine Time (ms):        58.926
Avg Response Size (bytes):    28302.454
Request Times over 1250 ms:   1 (0.151%)
Engine Times over 750 ms:     2 (0.302%)

----- Searches Including Wildcards -----
Number of Requests:          0 (0.000%)
```

This section displays statistics categorized by the features enabled on each query processed by the MDEX Engine. (The output in this section does not include Web Service queries.) The sum of the percentages of each extended query type can be more or less than 100%, since any one query can enable none, one, or more than one of these features. The total may also be less than 100% because the percentages do not include Web Service queries. (For example, if most of the Dgraph queries are Web Service queries, then the totals in this section will likely add up to much less than 100%.)

The following extended query types are output within this section:

- **Did-You-Mean Enabled Searches:** search queries that enable the "Did you mean?" feature, typically using the Nty=1 URL parameter or the setNavERecSearchDidYouMean() API method.
- **Searches Including Wildcards:** search queries that include a wildcard character (*) within the search terms. Note that the inclusion of a wildcard character does not necessarily mean that the MDEX Engine performed a wildcard search, since it is not necessarily true that the target search interface was enabled for wildcarding. Wildcard searches are more expensive than standard searches.
- **Boolean Searches:** search queries that specified the matchboolean matchmode, typically using the Ntx URL parameter or by specifying the mode within the ERecSearch object's constructor in the Endeca Presentation API. Boolean searches can be more expensive than standard searches.
- **Requests Including a Record Filter:** navigation, search, dimension search, or aggregate record queries that specify a record filter, typically using the Nr URL parameter or the setNavRecordFilter() API method.
- **Requests Including a Range Filter:** navigation, search, dimension search, or aggregate record queries that specify a range filter, typically using the Nf URL parameter or the setNavRangeFilters() API method. This category does not include geocode filters, even though geocode filters use the same URL parameters and API methods.
- **Requests Including a Geocode Filter:** navigation, search, dimension search, or aggregate record queries that specify a geocode filter, typically using the Nf URL parameter or the setNavRangeFilters() API method.

Filters() API method. This category does not include other range filters, even though other range filters use the same URL parameters and API methods.

- **Aggregate requests:** navigation or search requests that specify an aggregate record rollup key, typically using the Nu URL parameter or the setNavRollupKey() API method.
- **Base requests:** navigation or search requests that do not specify an aggregate record rollup key.
- **Requests Exposing All Refinements (allgroups):** navigation or search requests that open all available dimension refinements by using the setNavAllRefinements(true) API method (there is no corresponding URL parameter). This can be an expensive configuration because it increases both MDEX Engine calculation requirements and response packet size.
- **Requests with an Explicit Sort Order:** navigation or search requests that specify one or more keys by which to sort results, typically using the Ns URL parameter or the setNavActiveSortKeys() API method. This does not include requests that use the default sort order specified in Dgidx, nor does it include search requests that are sorted by relevancy (the default for searches).
- **Multi-Key Search Requests:** search requests that specify more than one search key, even if the multiple searches utilize the same interface, typically by pipe-delimiting the Ntt and Ntk URL parameters or by specifying more than one ERecSearch in an ERecSearchList within the API. Multi-key searches are typically generated by search-within functionality or on advanced search pages that include parametric search functionality.

Response codes

These statistics are only available when the --showResponseCodes flag is enabled.

```
=====
== Response Codes
=====
----- Response Code: 200 -----
Number of Requests:          929698 (99.948%)
Avg Response Time (ms):      362.699
Avg Engine Time (ms):        44.742
Avg Response Size (bytes):    507214.896
Request Times over 1250 ms:   34070 (3.665%)
Engine Times over 750 ms:     1213 (0.130%)

----- Response Code: 404 -----
Number of Requests:          484 (0.052%)
Avg Response Time (ms):      902.004
Avg Engine Time (ms):        2.147
Avg Response Size (bytes):    10723.893
Request Times over 1250 ms:   150 (30.992%)
Engine Times over 750 ms:     0 (0.000%)
```

This section displays statistics categorized by the HTTP response code returned by the MDEX Engine for each query. The sum of the percentages of each response code should equal 100% (accounting for rounding), since any one query can only produce a single response code.

The following are the most common response codes produced by an MDEX Engine:

- **200:** OK. The MDEX Engine successfully processed the query.
- **404:** Not Found. The query sent to the MDEX Engine was in an incomprehensible format (e.g. "/favicon.ico"), or the query specified a navigation descriptor that does not exist. The most common causes of 404s are browser bookmarks/favorites that are no longer valid, direct browser queries to the MDEX Engine that produce favicon requests, and typos.
- **408:** Request Timeout. The client that produced the request (typically the Endeca API) went away before the MDEX Engine could write the response to it. Typically this means that the request took

a very long time to complete and the requesting application timed out before the MDEX Engine returned a response. Empty requests also often generate 408 codes.

For more information on HTTP response codes, see RFC2616 (Hypertext Transfer Protocol) and elsewhere.

Request profiling

These statistics are only available when the `--showProfiles` flag is enabled.

```
=====
== Request Profiling
=====
----- Number of Base Records Requested -----
Requests Analyzed:          169
Average:                    14.734
Standard Deviation:         1.119
Maximum:                     15
Minimum:                     10
Over Threshold (50):         0 (0.000%)

----- Number of Aggregate Records Requested -----
Requests Analyzed:          666
Average:                    14.932
Standard Deviation:         1.004
Maximum:                     15
Minimum:                     0
Over Threshold (50):         0 (0.000%)
```

This section displays metrics for the individual feature configuration specified for each request processed by the MDEX Engine. Because request profiling requires deeper analysis of request URLs, enabling profiling slows down the Request Log Analyzer's processing.

The output in this section does not include Web Service queries. Because of this, the percentages may add up to less than 100%. (For example, if 80% of all queries are Web Service queries, then the total of all percentages in this section will add up to 20%.)

The following metrics are output within this section:

- **Number of Base Records Requested:** for base navigation or search queries, the number of base records that the Endeca Presentation API specified should be returned in detail. The MDEX Engine response includes the full content of a certain number of records (the "current page" of records) as well as available navigation and other meta information about the remaining records. This metric measures the number of records included in the current page; the more records, the larger the resulting data packet.
- **Number of Aggregate Records Requested:** for aggregate navigation or search queries, the number of aggregate records that the Endeca Presentation API specified should be returned in detail. The MDEX Engine response includes the full content of a certain number of records (the "current page" of records) as well as available navigation and other meta information about the remaining records. This metric measures the number of records included in the current page; the more records, the larger the resulting data packet.
- **Pagination Offset:** for navigation or search queries, the pagination offset specified by the query. When paginating through result sets, the pagination offset sets which page of results should be the current page. Higher pagination offsets are more expensive for the MDEX Engine to compute.
- **Number of Navigation Descriptors (not including N=0):** for navigation or search queries, this metric measures the number of navigation nodes specified as descriptors by the query. This metric ignores queries that only specify the root node (N=0).

- **Number of Exposed Dimensions (not including none or allgroups):** for navigation or search queries, this metric measures the number of refinement dimensions opened by the query. This metric ignores queries that open no dimensions and queries that open all dimensions.
- **Number of Explicit Sort Keys:** for navigation or search queries, this metric measures the number of record sort orders specified for the query. This metric ignores queries that do not specify a sort order, such as those queries that sort by search relevance or by the default sort order specified in `dgidx`.
- **Number of Search Terms:** for search queries, this metric measures the number of search terms specified in the query. The number of terms is calculated by splitting the entire search string on whitespace. This is an approximation; the MDEX Engine can also split terms on punctuation characters.
- **Number of Search Keys:** for search queries, this metric measures the number of search keys (interfaces) specified in the query. If the same key is specified twice, it counts as two keys.

Response profiling

These statistics are only available when the `--showProfiles` flag is enabled.

```
=====
== Response Profiling
=====
----- Requests that Returned Zero Results -----
Number of Requests:                0 (0.000%)

----- Number of Base Records in Result Set -----
Requests Analyzed:                  169
Average:                           19562.497
Standard Deviation:                 107058.210
Maximum:                           654243
Minimum:                           0

----- Number of Aggregate Records in Result Set -----
Requests Analyzed:                  666
Average:                           3874.508
Standard Deviation:                 36769.412
Maximum:                           546747
Minimum:                           0
```

This section displays information about the individual aspects of the responses returned by the MDEX Engine. Because response profiling requires deeper analysis of request URLs, enabling profiling slows down Request Log Analyzer processing.

The output in this section does not include Web Service queries. Because of this, the percentages may add up to less than 100%. (For example, if 80% of all queries are Web Service queries, then the total of all percentages in this section will add up to 20%.)

The following statistics are output within this section:

- **Requests that Returned Zero Results:** for base navigation or search queries, the number of queries that resulted in no records being found for the main result set of that query. Note that records may have been returned as supplements by dynamic business rules for the query.
- **Number of Base Records in Result Set:** for base navigation or search queries, the number of records found by the query. This is the total number of records found, not the number specified by the API to be returned in detail for the current page. The minimum and maximum metrics for this statistic will typically be 0 and the full number of records in the index and are therefore not particularly interesting.

- **Number of Aggregate Records in Result Set:** for aggregate navigation or search queries, the number of records found by the query. This is the total number of records found, not the number specified by the API to be returned in detail for the current page. The minimum and maximum metrics for this statistic will typically be 0 and the full number of records in the index and are therefore not particularly interesting.
- **No Children Per Aggregate:** for aggregate navigation or search queries, this statistic measures the performance of those queries specifying that no base child should be returned for each aggregate record.
- **One Child Per Aggregate:** for aggregate navigation or search queries, this statistic measures the performance of those queries specifying that one base child should be returned for each aggregate record (this is the default configuration).
- **All Children Per Aggregate:** for aggregate navigation or search queries, this statistic measures the performance of those queries specifying that all base children should be returned for each aggregate record. This can be an expensive configuration because it increases both MDEX Engine calculation requirements and response packet size.

Peak performance

These statistics are only available when the `--showHourly` flag is enabled.

```
=====
== Peak Performance
=====
----- testlogs/P_US_product_hot.log Fri Oct 21 12:11:05 2005:00-Fri Oct 21
13:11:05 2005 (1129914643-1129914665) -----
Number of Requests:                1321 (100.000%)
Avg Response Time (ms):            87.527
Avg Engine Time (ms):              34.816
Avg Response Size (bytes):         16274.311
Request Times over 1250 ms:        1 (0.076%)
Engine Times over 750 ms:          2 (0.151%)

ops/sec:                          57.435
```

This section displays the single hour time slice from all logfiles that contained the most requests. While other statistics report against the entire set of requests, viewing the peak performance gives a closer approximation of potential performance vs. achieved performance (see Result Interpretation).

Threading and queueing information

These statistics are only available when the `--showThreading` flag is enabled.

```
=====
== Threading/Queueing Information
=====
----- Queued Requests -----
Number of Requests:                3 (0.040%)
Avg Response Time (ms):            8598.690
Avg Engine Time (ms):              0.000
Avg Response Size (bytes):         0.000
Request Times over 1250 ms:        2 (66.667%)
Engine Times over 750 ms:          0 (0.000%)

----- Requests that encountered no queue -----
Number of Requests:                7513 (99.960%)
Avg Response Time (ms):            313.060
```

```

Avg Engine Time (ms):          144.244
Avg Response Size (bytes):     19739.288
Request Times over 1250 ms:    258 (3.434%)
Engine Times over 750 ms:      240 (3.194%)

----- Engine Queue Length -----
Requests Analyzed:             7516
Average:                       3.991
Standard Deviation:            1.734
Maximum:                       8
Minimum:                       0
Over Threshold (5):            3 (0.040%)

----- Idle Engine Threads -----
Requests Analyzed:             7516
Average:                       1.996
Standard Deviation:            0.075
Maximum:                       2
Minimum:                       0

```

This section displays information about the performance and behavior of a multithreaded MDEX Engine. A multithreaded MDEX Engine is able to process N requests simultaneously, where N is the number of threads specified using the `--threads` dgraph flag. As requests are received by the MDEX Engine, they are either handled immediately by an available engine thread or placed into the request queue and handled on a first-come-first-serve basis as MDEX Engine threads become available.

This section is only valid for multithreaded MDEX Engines. MDEX Engines that do not specify the `--threads` flag do not record the appropriate information to their logs and therefore will not produce usable statistics for this section.

The following statistics are output within this section:

- **Queued Requests:** this statistic measures performance of those requests that were placed into the request queue before being processed.
- **Requests that encountered no queue:** this statistic measures performance of those requests that were handled immediately by an available thread and were not placed into the request queue.
- **Engine Queue Length:** this metric measures the length of the request queue. When engine threads are available, the queue length is 0. When all threads are in use, the queue length is equivalent to the number of requests already in the queue ahead of the current request. *The minimum value for this metric should always be zero and is therefore not interesting.*
- **Idle Engine Threads:** this metric measures the number of available engine threads. When engine threads are available, this number is equivalent to the number of available threads. When all threads are in use, this number is 0. *The maximum value for this metric should always be equivalent to the total number of threads specified with the `--threads` flag and is therefore not interesting. If at any point a request encounters a queue, the minimum value for this metric will be zero.*

Note that requests to the MDEX Engine that generate a 404 error code arbitrarily write the value 10000 to the MDEX Engine log as their queue length. The Request Log Analyzer always completely ignores these requests when calculating queue lengths, though for other metrics these requests are controlled by the `--ignore*` flags. Because the Request Log Analyzer always ignores these requests, the statistics that track queued and unqueued requests will not add up to 100% in certain circumstances: if the appropriate `--ignore*` flags are not specified, the percentages will represent only the subset of those requests that did not contain a 10000 as their queue length.

Summary information

These statistics are always output.

```
=====
== Summary information
=====
First date analyzed:      Thu Nov  3 14:12:20 2005
Last date analyzed:      Thu Dec 15 16:04:52 2005
Time period analyzed:    42d 1h 52m 32s

Lines analyzed, total:    383545
Valid requests in time period: 383391
Avg ops/sec:             0.105

Bytes/sec:               2238.883
Mb/sec:                  0.017

Max ops in one second:   22
Max ops in five seconds: 75 (15.000 ops/sec)

Interpolated one-second ops avg: 1.796 ops/sec
Interpolated five-second ops avg: 1.110 ops/sec

----- Round-trip Response Time, ms -----
Requests Analyzed:      383391
Average:                345.183
Standard Deviation:     561.897
Maximum:                40077.56
Minimum:                1.55
Over Threshold (1250):  11269 (2.939%)

----- MDEX Engine-only Processing Time, ms -----
Requests Analyzed:      383391
Average:                48.861
Standard Deviation:     132.229
Maximum:                15475.75
Minimum:                0.00
Over Threshold (750):   1325 (0.346%)

----- Response Differential, ms (round-trip minus engine-only) -----
Requests Analyzed:      383391
Average:                296.321
Standard Deviation:     544.660
Maximum:                40003.33
Minimum:                1.3
Over Threshold (500):   42730 (11.145%)

----- Response Size, bytes -----
Requests Analyzed:      383391
Average:                21230.479
Standard Deviation:     49118.382
Maximum:                895065
Minimum:                0
Over Threshold (393216): 456 (0.119%)
```

This section outputs information about the performance and characteristics of the MDEX Engine as a whole.

The following statistics are output in this section:

- **First date analyzed:** the timestamp of the earliest request analyzed within the logs. Translation of the time recorded in the log to human-readable time is controlled with the `--hourOffset` flag.
- **Last date analyzed:** the timestamp of the latest request analyzed within the logs. Translation of the time recorded in the log to human-readable time is controlled with the `--hourOffset` flag.
- **Time period analyzed:** the timespan between the earliest analyzed request and the latest analyzed request.
- **Lines analyzed, total:** the number of requests inspected by the Request Log Analyzer within this set of logfiles. Not all of these requests will have been analyzed. MDEX Engine startup and shutdown status messages write a line to the logfile and thus contribute to this metric but are otherwise ignored by the Request Log Analyzer. Other lines are controlled by the `--ignore*` flags.
- **Valid requests in time period:** the number of requests actually analyzed by the Request Log Analyzer within this set of logfiles.
- **Avg ops/sec:** the achieved throughput, in operations (queries) per second. This metric is calculated by dividing the number of requests analyzed by the number of seconds in the time period analyzed. Note that this metric can be heavily skewed by large stretches of time that contain no requests.
- **Bytes/sec:** the achieved network bandwidth usage, in bytes. This metric is calculated by dividing the sum of all data packet sizes analyzed by the number of seconds in the time period analyzed. Note that this metric can be heavily skewed by large stretches of time that contain no requests.
- **Mb/sec:** the achieved network bandwidth usage, in megabytes.
- **Max ops in one second:** the highest achieved throughput metric for a one-second timespan. The Request Log Analyzer slices each logfile into one-second pieces, finds the one-second timespan that contains the most analyzed requests, and reports that number here.
- **Max ops in five seconds:** the highest achieved throughput metric for a five-second timespan. The Request Log Analyzer slices each logfile into five-second pieces, finds the five-second timespan that contains the most analyzed requests, and reports the throughput of that timespan here. *Note that the slicing algorithm is arbitrary; it is possible that the last two seconds of a slice plus the next three seconds of the following consecutive slice contain more requests than any of the arbitrary slices.*
- **Interpolated one-second ops avg:** a version of achieved throughput that attempts to address the weaknesses of the previous achieved throughput measurement. This metric is calculated by dividing the number of requests analyzed by the number of seconds in the logs that actually contain requests. For most logfiles this metric will be a more accurate achieved throughput. For logfiles containing a large number of requests that required more than one second to process, this metric will be inaccurate.
- **Interpolated five-second ops avg:** a version of achieved throughput that attempts to address the weaknesses of the previous achieved throughput measurement. This metric is calculated by dividing the number of requests analyzed by the number of five-second time slices in the logs that actually contain requests, then adjusting by 5. For most logfiles this metric will be a more accurate achieved throughput. For logfiles containing a large number of requests that required more than one second to process, this metric will be inaccurate.
- **Round-trip Response Time, ms:** The round-trip response time for all analyzed requests. See Common Metrics for an explanation of round-trip response time.
- **MDEX Engine-only Processing Time, ms:** The engine-only processing time for all analyzed requests. See Common Metrics for an explanation of engine-only processing time.
- **Response Differential, ms (round-trip minus engine-only):** The response differential for all analyzed requests. The response differential is measured as the difference between the round-trip response time and the engine-only processing time. This metric measures how long requests spend in the request queue and the response queue combined.
- **Response Size, bytes:** The response size for all analyzed requests. See Common Metrics for an explanation of round-trip response size.

Appendix F

MDEX Engine Statistics and Auditing

The MDEX Engine Statistics page displays MDEX Engine (Dgraph) performance statistics. The MDEX Engine Auditing page tracks usage for licensing and performance purposes. This section describes these pages.

About the MDEX Engine Statistics page

The MDEX Engine Statistics page provides a detailed breakdown of what the Dgraph is doing, and is a useful source of information about your Endeca implementation's configuration and performance.

The statistics page is also called the Dgraph Stats page or Admin Stats page.

It provides information such as startup time, last data indexing time, and indexing data path. This allows you to focus your tuning and load-balancing efforts. By examining this page, you can see where the Dgraph is spending its time. Begin your tuning efforts by identifying the features in the Hot Spot Analysis section with the highest totals.

Viewing the MDEX Engine Statistics page

You can request the MDEX Engine Statistics page for the Dgraph with the URL listed below.

```
http://DgraphServerNameOrIP:DgraphPort/admin?op=stats
```

For example, if your Dgraph is running on your local machine and listening on port 8000, specify this:

```
http://localhost:8000/admin?op=stats
```

You can determine the host and port on the EAC Admin Console of Endeca Workbench by opening the MDEX Engine component or by exploring your Deployment Template `AppConfig.xml` file.

To reset the statistics, make the following request:

```
http://DgraphServerNameOrIP:DgraphPort/admin?op=statsreset
```

To view the statistics information for a single request, clear statistics, issue a request and inspect statistics again.

The statistics page information is valid as long as the MDEX Engine keeps running; it is cleared upon the MDEX Engine restart.

The source data for the Dgraph statistics is stored in XML. By default, the MDEX Engine Statistics page is rendered into HTML through an Endeca XSLT stylesheet, `stats.xslt`, that is installed in the `ENDECA_MDEX_ROOT/conf/dtd/xform` directory.

If your browser supports XSLT transformations (for example, Internet Explorer 6 and later), you can view the statistics as transformed by `stats.xslt`, or you can modify the shipped `stats.xslt` stylesheet to provide a different transformation of the data.

If your browser does not support XSLT transformations, or if you want to see the raw XML, rename or remove `ENDECA_MDEX_ROOT/conf/dtd/xform/stats.xslt`.

To ensure that the statistics page displays properly in your browser, JavaScript must be enabled, as part of the settings for the browser. (If your browser's security setting is set to high, this may disable JavaScript.)

Sections of the MDEX Engine Statistics page

The MDEX Engine Statistics page for the Dgraph is divided into tabs. Information on all of the tabs is presented through the URL of the statistics page as described in the following sections.

The Performance Summary tab

The **Performance Summary** tab contains the highest level statistics. They reflect and help to monitor those characteristics that are external to the actual processing of queries, such as the queue of incoming queries, the thread pool, and the overall throughput of the process.

The **Performance Summary** tab contains the following sections:

Section	Description
Performance	Various statistics (average, standard deviation, minimum, maximum, and total) on: • The total number of requests received • Total CPU usage (in seconds of total user time and total system time). • The memory resource usage. • Resident Set Size (RSS) statistics.
Throughput (req/sec)	Five-minute, one-minute, and ten-second average throughput statistics (only for multithreaded mode). When thread becomes available, the throughput statistics is measured.

The General Information tab

The **General Information** tab contains the following sections.

Section	Description
Information	Basic connection and machine details, such as process ID, parent process ID, user ID, user name, effective user ID, group ID, effective group ID, current working

Section	Description
	directory, hostname, server port for the Dgraph, start time, information about the data (path, tag and date), and the number of index generations.
Arguments	A list of all arguments the Dgraph was started with.

The Index Preparation tab

The **Index** tab tracks index preparation and precomputed sorts statistics, including timing.

It contains the following sections:

Section	Description
Update Totals	The number of non-XQuery updates run against the Dgraph, and performance of updates (count, average, standard deviation, min, max and total), on the following items: • Record changes, including the number of adds, updates, deletes and replacements • Dimension changes • Record change errors • Dimension change errors • Update latency, including various finer-grained performance statistics of indexing processing.
XQuery Update Totals	The number of XQuery updates run against the Dgraph, and performance of updates (count, average, standard deviation, min, max and total), on the following items: • Record changes, including the number of adds, updates, deletes and replacements • Dimension changes • Record change errors • Dimension change errors • Update latency, including various finer-grained performance statistics of indexing processing.  Note: The XQuery update feature is Early Access in this release. For details, see the <i>Web Services and XQuery Developer's Guide</i>.
Precomputed Sorts	Displays how much time the Dgraph has spent computing sorts, including computing sorts and incremental sort updates.



Note: For some of the statistics on this page, it is possible to drill down for further information by clicking on the black arrow that appears outside the rightmost column.

The Cache tab

The **Cache** tab contains information about the MDEX Engine cache.

Section	Description
Main Cache	Provides details on totals, including number of entries in the cache, size of entries, number or lookups in the cache, number of rejections, percentage of hit rate and miss rate, number and size of evictions from the cache, number of reinsertions, total reinsertion time and average creation and eviction times. In particular, if you need to analyze the MDEX Engine cache, examine the results in the following columns. Analyzing these results may help you tune your cache and re-design your front-end application to improve performance of the MDEX Engine. • Number of rejections. Counts greater than zero in this column indicate that the cache is undersized and you may want to increase it. • Number of reinsertions. Large counts in this column indicate that simultaneous queries are computing the same values, and it may be possible to improve performance by sequencing queries, if the application design permits. • Total reinsertion time. Examining this column is useful for quantifying the overall performance impact of queries that contribute to the "Number of reinsertions" column. This column represents the aggregated time that has been spent calculating identical results in parallel with other queries. This is the amount of compute time that potentially can be saved by sequencing queries in a re-design of the front-end application.

The Details tab

The **Details** tab contains the following sections:

Section	Description
Most Expensive Queries	The URL and total time in milliseconds for the ten queries with the largest total computation time (that is, queue time plus Dgraph processing time plus write time) made in the session. The queries are ordered by processing time. Each time a new Dgraph transaction that yields results is completed, this tab may become updated with a new query, if it makes the list of current top ten most expensive queries. Each query is described with these characteristics: • Query rank • Computation processing time (in milliseconds) • URL Unlike in Presentation API mode, where the URL contains all of the information about the query, in Web services mode the URL only contains the service name. The bulk of the query is contained in the POST body. Therefore, if the Dgraph is running in Web services mode, a serial number is appended to the URL, as in the following example: <code>/ws/myservice:57</code>. This serial number corresponds to the HTTP Exchange ID in the MDEX Engine Request Log. You can use it to retrieve additional information about the contents of the query from the Request Log's Query Body field.

Section	Description
Hotspots	Details on the performance of specific features, such as clustering, record search, record filter, range filter, content spotlighting and snipping. This section also contains the following page render and record sorting statistics: • Page render total. After the MDEX Engine knows which records and values must be returned, this time represents the total time spent generating and returning those results to the Presentation API. This time includes retrieving records from memory or disk, ordering them based on the specified sort or relevance ranking strategies, as well as other information returned to the API, such as content spotlighting results. • Prefetching horizontal record. The cost to retrieve records from the data layer of the MDEX Engine. • Statistics related to various sorting strategies. The MDEX Engine examines information about the data being returned and selects the best sorting strategy.  Note: These statistics may change. They are used for internal debugging and tuning of the MDEX Engine sorting selection strategy and are not useful to the end user.
Results	The following items are listed in the Results section. The statistics includes count, average, standard deviation, min, max and total, where applicable: • Number of records in result set • Result page size in bytes • Result page format performance in milliseconds
Server	Statistical information for the MDEX Engine server: • HTTP: Total request time • HTTP: Time reading request • HTTP: Time in scheduler • HTTP: Time writing response • HTTP: Request bytes read (including HTTP overhead) • HTTP: Response body size (including HTTP overhead) • Scheduler: Queue time before processing • Scheduler: Processing time • Scheduler: Queue time after processing • Scheduler: Queries queued This metric describes the queue length. • Scheduler: Queries in process This metric describes the number of queries that are in process. • Scheduler: Update queue time • XQuery: Total time in XQuery engine • XQuery: Total time in XQuery external functions

Section	Description
	 Note: This statistic only includes the time spent in the following functions: <code>internal:query()</code>, <code>mdex:dimension-value-id-from-path()</code>, and <code>mdex:add-navigation-descriptors()</code>. • XQuery: Time retrieving documents with <code>fn:doc()</code> • XQuery: Time storing documents with <code>fn:put()</code> • XQuery: Result serialization time • Most expensive MAX invocations • Custom timing list This metric, which can list things like expensive queries, only appears when you implement custom metric gathering with the <code>ep:stats-timing</code> pragma. See the <i>Web Services and XQuery Developer's Guide</i> for more information.
Navigation	Information about the number of navigation pages, as well as navigation performance, query size, and result size by average, standard deviation, minimum, maximum, and total.
Record Sorting	The number and type of sorts performed (does not include timing), and the percentage of those sorts for each sort type.
Analytics	Information pertaining to the analytics features in Endeca Analytics, such as total processing time, query parsing, time checking and evaluation times.
Disk usage	Disk usage statistics for the indices: • current total disk usage value (MB) • disk usage high water mark value (MB)
Search	A finer-grained analysis of the performance of individual features. This information is used for the internal analysis by Endeca.
Data Layer Performance	Statistical information about the data layer performance. This information is used for the internal analysis by Endeca.



Note: For some of the statistics on this page, it is possible to drill down for further information by clicking the black arrow that appears outside the rightmost column.



Note: If you modified the `stats.xslt` style sheet that is included in the installation, the information might display differently.

About the Endeca MDEX Engine Auditing page

The MDEX Engine Auditing page lets you view the aggregate MDEX Engine metrics over time. It provides the output of XML reports that track ongoing usage statistics.

These statistics persist through process restarts.

This data can be used to verify compliance with licensing terms, and is also useful for tracking product usage.



Note: Each Dgraph in an implementation is audited separately.

Viewing the MDEX Engine Auditing page

You can request the MDEX Engine Auditing page with the URL below.

To view the MDEX Engine Auditing page:

Access the following URL:

```
http://DgraphServerNameOrIP:DgraphPort/admin?op=audit
```

For example, if your Dgraph is running on your local machine and listening on port 8000, specify this:

```
http://localhost:8000/admin?op=audit
```

The information on the MDEX Engine Auditing page is persistent and is valid across the MDEX Engine restarts.

The source data for the auditing reports is stored in XML. By default, the MDEX Engine Auditing page is rendered into HTML through an Endeca XSLT stylesheet, `audit.xslt`, that is installed in the `ENDECA_MDEX_ROOT/conf/dtd/xform` directory.

Audit persistence file details

The naming convention for the audit persistence file is:
`audit-<data_prefix>-<persistence_number>.xml`.

For example, an audit persistence file on the sample wine implementation might look like this:
`audit-wine-0.xml`.

This convention ensures that each Dgraph creates a unique file. It makes it possible to maintain the audit persistence files for numerous Dgraphs in an application in the same directory without contention.

By default, the audit persistence file is written to a directory called `persist` that is located in the application's working directory. To direct it elsewhere, use the Dgraph flag `--persistdir` when you first create the Dgraph. Do not move or rename this directory after it has been created.

You should not delete the audit persistence file or attempt to edit it manually. Upon startup, the Dgraph checks for the presence of this file, and if it cannot find it or read it, it issues a warning message and creates a new one. However, if you see such a warning message when you first create a Dgraph, you can safely disregard it.



Note: The auditing function adds information prefixed by the word `Endeca.*` to records. This namespace is reserved for administrative use and should not be used for other purposes.

Sections of the MDEX Engine Auditing page

The MDEX Engine Auditing page consists of two tabs: Audit Stats and General Information.

Auditing statistics are gathered in one of two ways:

- The Query Load statistic tracks the hour with the most queries in each calendar week, starting when you first run the Dgraph and persisting through process restarts.
- All other auditing statistics constantly monitor the peak value over the course of a calendar week, and report the exact time when a value greater than the current peak value appears, starting when you first run the Dgraph and persisting through process restarts.

Because these metrics are calculated over the course of a week, a change such as a deleted record is not reflected until the following week, when the peak value count is reset.

The Audit Stats tab

The Audit Stats tab contains the following information.

Section	Description
Query Load	The peak number, in the week beginning at the displayed time, of queries that the Dgraph has received in any single hour, plus the time at which that peak occurred. This field contains the sum of the next two fields, Net Query Load and WS Query Load. Depending on the modes in which you run your Dgraph, there may be values in both of these fields or only one of them.
Net Query Load	The peak number, in the week beginning at the displayed time, of queries that the Dgraph has received in any single hour while running in Presentation API mode, plus the time at which that peak occurred.
WS Query Load	The peak number, in the week beginning at the displayed time, of queries that the Dgraph has received in any single hour while running in Web services mode, plus the time at which that peak occurred.
Number of Records	The peak number, in the week beginning at the displayed time, of records, plus the time at which that peak was reached.
Number of Columns	The peak value, in the week beginning at the displayed time, for the total number of properties and dimensions across all records, plus the time at which that peak was reached.
Number of Words	The peak value, in the week beginning at the displayed time, for the total number of words (counting multiple occurrences of the same word) across all records, plus the time at which that peak was reached.
Number of Assignments	The peak value, in the week beginning at the displayed time, for the total number of populated dimension and property values across all records, plus the time at which that peak was reached.
Size of Data	The peak value, in the week beginning at the displayed time, for the total size occupied by all records, plus the time at which that peak was reached.  Note: This may vary, depending on operating system platform.

The General Information tab

The General Information tab contains the following sections.

Section	Description
Information	Basic connection and machine details.
Arguments	A list of all arguments the Dgraph was started with.



Note: This tab is identical to the one of the same name on the MDEX Engine Server Statistics page.

Appendix G

Useful Third-Party Tools

This section lists some third-party tools that you may find useful during the Endeca performance monitoring process. The tools listed here are not supported by Endeca and are subject to change. In addition, these suggestions are not meant to overrule your choice of other tools.

Cross-platform tools

The following tools are available in both UNIX and Windows versions.

Tool	Description
Wireshark	Wireshark is an open source network protocol analyzer for both UNIX and Windows. It allows you to examine data from a live network or a capture file on disk. For information and downloads, see http://www.wireshark.org/download.html.
Tcpdump/Windump	Tcpdump (and its Windows version, Windump) are network traffic analysis tools. These tools can be used to watch and diagnose network traffic according to various complex rules. You can download Tcpdump from http://www.tcpdump.org. You can download Windump from http://www.winpcap.org/windump.  Note: Tcpdump comes with most Linux distributions by default.

Solaris and Linux tools

The following tools are available for both Solaris and Linux.

Tool	Description
Netperf	Netperf is a network benchmarking tool that can be used to measure the throughput of many different types of TCP and UDP connections. Netperf provides tests for both unidirectional throughput, and end-to-end latency.  Note: Be sure to compile netperf with histogram support. To simulate the network traffic to a MDEX Engine with average result pages of 50,000 bytes, run netperf like this: <pre>netperf -l 600 -v 2 -H remotehost -p 8899 -t TCP_CRR -- -r 200, 50000</pre> where: • -l is the length of the test in seconds • -v specifies verbose output level • -H indicates the host where netserver is running • -p indicates the port that was given to the netserver process • -t indicates the test to run. TCP_CRR is the TCP test that opens a new TCP connection for each request/response • -r specifies the request/response characteristics, in this case a 200 byte request (approximately the size of a URL) and a 50K result For information and downloads, see http://www.netperf.org.
Top	Top is a UNIX utility you can use to quickly identify top CPU-using processes. It is a popular and common tool for monitoring system-wide process activity. For information and downloads, see http://www.groupsys.com/top.
Sar	Sar reports system activity on single processor systems. It reports the status of counters in the operating system that are incremented as the system performs various activities. These include counters for CPU utilization, buffer usage, disk I/O activity, TTY device activity, switching and system-call activity, file access, queue activity, inter-process communications, swapping and paging. On Solaris, sar is part of the system activity reporter package. On Linux, it is part of the downloadable sysstat package.
iostat	The iostat utility iteratively reports terminal, disk, and tape I/O activity, as well as CPU utilization. On Solaris, iostat is built in to the operating system. On Linux, it is part of the downloadable sysstat package.

Solaris-specific tools

The following utilities are built into Solaris.

Tool	Description
prstat	On Solaris the prstat command displays information about active processes on the system. By default, prstat displays information about all processes sorted by CPU usage.
cpusar and mpsar	On multiprocessor machines, cpusar reports per-CPU statistics, and mpsar reports system-wide statistics.
Kstat	Kstat reports many kernel parameters and statistics.
lockstat	The lockstat utility gathers and displays kernel locking and profiling statistics. It allows you to identify what are the processes and kernel really doing. Lockstat allows you to specify which events to watch, how much data to gather for each event, and how to display the data.
SE Toolkit	The SE Toolkit is a collection of scripts for performance analysis that gives advice on performance improvements.

Linux-specific tools

The following tools are available for Linux.

Tool	Description
sysstat	The sysstat utilities package is a download for Linux that contains performance monitoring tools such as iostat, sar, and mpstat. iostat and sar are described in “Solaris and Linux tools”. For information and downloads, see http://perso.wanadoo.fr/sebastien.godard.
Mpstat	Mpstat is the Linux multiprocessor load display utility. It displays system processor activity information on your screen for each of the processors serialized on your system.

Windows tools

The following tools are available for Windows.

Tool	Description
Task Manager	The Windows Task Manager provides information about programs and processes running on your computer. It also displays the most commonly used performance measures for processes. You can access the Task Manager by right-clicking an empty area on the task bar on your Windows machine.
Performance Monitor	The Performance Monitor provides details about the resources used by specific components of the operating system and by programs that have been designed to collect performance data.

Tool	Description
	You can access the Performance Monitor from the Control Panel by selecting Administrative Tools > Performance.
Other performance tools	Sysinternals (http://www.sysinternals.com) offers useful freeware tools, including the following: • Process Explorer, which shows you information about which handles and DLL processes have opened or loaded. • TCPView, which shows you detailed listings of all TCP and UDP endpoints on your system, including the local and remote addresses and state of TCP connections. On Windows NT, 2000, and XP TCPView also reports the name of the process that owns the endpoint.

Appendix H

Tuning the Network Performance

You only need to perform the procedures described in this appendix if you are installing in a production environment—they are not required for a typical developer installation. You will not see the benefits of this tuning until the Endeca server is placed under very heavy load.

Tuning network performance on Windows

Endeca provides two registry scripts that you can use, singly or in combination, to tune your server's network performance.

- The `tcp_time_wait_tune.reg` script tunes the server's network performance by changing the default time wait interval from 240,000 to 60,000 milliseconds. This change accelerates the rate at which the server re-uses ports when establishing TCP connections.

To determine if you need to run the tuning script, open the Registry Editor and look for the following key:

```
HKEY_LOCAL_MACHINE\SYSTEM\CurrentControlSet\
Services\Tcpip\Parameters\TcpTimedWaitDelay
```



Note: In the Registry Editor Explorer pane, expand the folders until you reach Parameters. Then click on the Parameters folder and look for the **TcpTimedWaitDelay** setting in the right pane.

If this key does not exist, that means that the system is using the default time-out of 240,000 milliseconds.

- The `tcp_max_ports_tune.reg` script increases the number of ports available for TCP connections from 5,000 to 65,534. The affected key appears in the Registry Editor as follows:

```
HKEY_LOCAL_MACHINE\SYSTEM\CurrentControlSet\
Services\Tcpip\Parameters\MaxUserPort
```

To tune network performance on Windows:

1. In the `%ENDECA_MDEX_ROOT%\bin` directory, double-click one of the following scripts:
 - `tcp_time_wait_tune.reg`
 - `tcp_max_ports_tune.reg`
2. When the information box reading “Are you sure you want to add the information in *script name* to the registry?” appears, click **Yes**.

3. The system displays a confirmation message that reads “Information in *script name* has been successfully added to the registry.” Click **OK**.
4. Optionally, repeat these steps for the other tuning script.
5. Reboot the server for the registry changes to take effect.

Tuning network performance on Solaris

This section applies only to Solaris installations, not to Linux installations.

The Endeca installation includes a script that tunes the server's network performance by changing the default time wait interval from 240,000 to 30,000 milliseconds. This change accelerates the rate at which the server re-uses ports when establishing TCP connections.

To determine if you need to run the tuning script, run the following command:

```
netstat -an | grep TIME_WAIT | wc -l
```

If the resulting number is consistently greater than 5,000, apply the tuning script and wait 4 minutes. The number of connections in a time wait state will drain off and you should find that the 5,000+ number drops by at least a factor of two.

To run the tuning script:

1. Change directories to the \$ENDECA_MDEX_ROOT/bin directory.
2. As root, type the following at the prompt:

```
./tcp_time_wait_tune.sh
```

3. Press **Enter**.

A message appears indicating that the `tcp_time_wait_interval` has been set to 30,000.

Configuring the FIN_WAIT_2 timeout interval

The `FIN_WAIT_2` timeout interval is the number of seconds that the HTTP server waits after sending the response for the client to close down its end of the socket. If this timeout expires, the server forcibly shuts down the connection.

This timeout interval is important for two reasons:

- Waiting for some time before shutting down the socket ensures that clients get complete responses.
- Timing out after certain period protects against buggy clients, which may never close their end of the socket. This can tie up resources on the server machine, leading to performance degradation and, in the extreme case, denial of service.

When the MDEX Engine finishes sending a response to a client, it does a "soft close" of the socket. This allows the client to finish reading data, and to close its end of the socket whenever it is ready. The state of the server-side socket during the interval between the server closing one end, and the client closing the other, is known as `FIN_WAIT_2`. All operating systems supported in this release automatically clean up sockets that stay in `FIN_WAIT_2` for too long.

In general, you should not need to change this from the default value. If you do need to change the setting, follow the instructions below for your operating system.

Configuring FIN_WAIT_2 timeout on Linux

On Linux systems, the `tcp_fin_wait` timeout is stored in `/proc/sys/net/ipv4/tcp_fin_timeout`.

You can change the value of this parameter using the `sysctl` command.

To get the value, issue the following command:

```
/sbin/sysctl net.ipv4.tcp_fin_timeout
```

To set the value, issue the following command:

```
/sbin/sysctl -w net.ipv4.tcp_fin_timeout=30
```



Note: Root permissions are typically required to set this value.

Configuring FIN_WAIT_2 timeout on Solaris

On Solaris systems, you can modify the `FIN_WAIT_2` timeout interval in `/dev/tcp`.

The default value is 675000ms.

To get the value, issue the following command:

```
ndd -get /dev/tcp tcp_fin_wait_2_flush_interval
```

To set the value, issue the following command:

```
ndd -set /dev/tcp tcp_fin_wait_2_flush_interval 30000
```



Note: Root permissions are typically required to set this value.

Configuring FIN_WAIT_2 timeout on Windows

On Windows systems, the variable to control the `FIN_WAIT_2` timeout interval can be modified in the Windows Registry.

The Registry entry that controls this setting is `HKEY_LOCAL_MACHINE\SYSTEM\CurrentControlSet\Services\Tcpip\Parameters`. You need to specify the `TcpFinWait2Delay` value for the above entry in the registry. The default value is 240s.



Note: Administrator privileges are required to set this value.

1. In the Windows Registry, go to `HKEY_LOCAL_MACHINE\SYSTEM\CurrentControlSet\Services\Tcpip\Parameters`
2. If the `TcpFinWait2Delay` value already appears in the details window, tune the value. The valid range is between 30s and 300s.
3. If the value does not exist, right click and select **Add a new DWORD value**. Add `TcpFinWait2Delay` and set its value.
4. Restart your system for the change to take effect.

Index

64-bit architecture 15

A

- allbins query type 89
- allgroups query type 90
- analytics query type 90
- application server issues 46
- attrs query type 91
- Auditing page 41
- autophrase query type 91
- autophrasedwim query type 92

B

- bandwidth
 - defined 12
- Boolean search 69

C

- cache
 - examining 150
- Cheetah utility 42
- compound query type 92
- concurrent users
 - defined 12
- connection setting for Eneper 115
- CPU considerations 46
- Cpusar performance analysis tool 159

D

- Dgidx
 - hardware performance recommendations 25
- Dgraph
 - introduced 11
- Dgraph cache
 - about 19
 - defined 12
 - impact on virtual process size 22
 - performance recommendations 20
- Dgraph memory usage
 - performance impact of 16
- Dgraph request log file 77
- Dgraph warming 44
- dictionary pruning 67
- dimension search 72
- dimension values
 - displaying refinements 54
 - dynamic statistics on 54

- disabled refinements
 - performance impact 55
- disk access considerations 45
- displaying
 - multi-select dimensions 53
 - refinement dimension values 54
- dym query type 92
- dynamic business rules 75
- dynamic refinement ranking
 - improving 55
- dynamic statistics on dimension values 54

E

- Eneper
 - introduced 113
 - logs for use with 123
 - optional settings 120
 - required settings 114
 - running locally 115
 - running remotely 115
 - setting the number of queries 121
 - usage 113
 - two-stream mode 124
- Eneper, about 42
- ephemeral port 47
- EQL filters
 - troubleshooting 65
 - typical errors 66
- Ethereal performance analysis tool 157
- expression evaluation of record filters 59

F

- feature performance
 - overview 51
- file system (FS) cache
 - about 20
- filter query type 93
- flat dimension hierarchy 53
- format query type 93
- front-end application
 - coding recommendations for 47
- full-duplex
 - recommended for NICs 25

G

- gathering performance information
 - guidelines for 39
- Gigabit Ethernet 25
- group query type 94
- groupby query type 94

H

hidden dimensions 53

I

id query type 95
 ignore query type 95
 iostat performance analysis tool 158
 inversion query type 96
 iteration setting for Eneper 115

K

keyprops query type 96

L

lang query type 97
 large OR filter performance 59
 large scale negation 59
 latency and maximum latency
 defined 11
 Linux
 performance tools 159
 sysstat package 159
 tuning 28
 load balancers
 recommendations 29
 load balancing 29
 Lockstat performance analysis tool 159
 log file (eneperf)
 about 123
 converting Dgraph request log 126
 settings 115
 with updates 124
 log query type 97

M

MDEX Engine architecture
 optimized for performance 13
 MDEX Engine request log
 converting for use with Eneperf 126
 MDEX Engine Statistics page
 about 147
 presentation transformed with XSLT 147
 viewing 147
 memory usage
 costs of record filters 58
 impact on performance 16
 recommendations 17
 merchdebug query type 97
 merchpreviewtime query type 98
 merchrulefilter query type 98
 model query type 98
 Mpsar performance analysis tool 159
 Mpstat performance analysis tool 159

multi-select dimensions
 displaying 53
 multithreaded mode
 associated costs 35
 introduced 33
 Solaris 37
 Linux 37
 VMware 37
 Windows 37

N

nbins query type 99
 nbulkbins query type 99
 Netperf performance analysis tool 158
 network recommendations 25
 node query type 100
 num query type 101

O

offline query type 101
 op query type 102
 operation
 defined 12
 opts query type 102
 out-of-memory situations
 Dgraph 18
 solutions 18

P

part list performance 59
 performance impact
 disabled refinements 55
 wildcard search 70
 performance in multithreaded mode 36
 Perl guidelines
 for MDEX Engine request log 82
 port setting for Eneperf 115
 precedence rules
 defined 73
 pred query type 103
 pretendtime query type 103
 Process Explorer performance analysis tool 160
 profiles query type 104
 Prstat performance analysis tool 159

R

RAID recommendations 15
 RAM estimates
 for the MDEX Engine 23
 rank query type 104
 read_ahead_kb parameter for tuning on Linux 28
 record counts for refinements
 aggregated records 55

- record filters
 - expression evaluation 59
 - large scale negation 59
 - memory costs 58
- Record Relationship Navigation
 - performance impact 63
- record search 69
- records
 - sorting by dimension or property 57
- refinement dimension values
 - displaying 54
- refinement query type 105
- relevance ranking
 - performance impact 74
- relink query type 105
- Request Log Analyzer, about 129
- Request Log Parser
 - modifying, for long or complex queries 126
- resident set size (RSS)
 - defined 12
- RSS
 - defined 12

S

- SAN-backed storage recommendations 16
- Sar performance analysis tool 158
- select query type 105
- snippeting
 - performance impact 66
- socket timeout interval
 - configuring 162
- Solaris performance tools 159
- sort query type 106
- spelling correction 66
- stat-abins
 - aggregated record counts per refinement 55
- statistics
 - viewing for MDEX Engine 147
- stats.xslt file 147
- stemming and thesaurus 68
- structured query type 106
- sustained throughput for Dgraph
 - defined 11

T

- Task Manager performance analysis tool 159
- Tcpdump performance analysis tool 157
- TCPView performance analysis tool 160
- terms query type 107

- thesaurus development
 - guidelines for 68
- threaded Dgraph 44
- throttle setting to Eneperf 121
- throughput
 - defined 11
- Top performance analysis tool 158
- tuning
 - cache 150
 - wildcard search 71

U

- uncovering network problems 47
- update logs and Eneperf 124
- updates performance, testing with Eneperf 124
- utilization
 - defined 12

V

- virtual process size
 - defined 12

W

- warming
 - effect on performance 44
- warming for the Dgraph process 12
- wildcard search
 - performance impact 70
 - preventing expensive queries 71
 - with punctuation 71
- wildcard_max setting 71
- Windows
 - Task Manager 159
- Windows 2008
 - support 26
- Windows performance tools 160
- Windump performance analysis tool 157
- working set size
 - defined 12
- WSS
 - defined 12

X

- XSLT
 - transforming MDEX Engine statistics 147

