



Hibernate Search 5.10.13.Final

Reference Guide

2022-10-27

Table of Contents

Preface	1
1. Getting started	2
1.1. System Requirements	2
1.2. Migration notes	2
1.3. Required libraries	2
1.4. Deploying on WildFly	3
1.5. Configuration	4
1.6. Indexing	7
1.7. Searching	8
1.8. Analyzer	9
1.9. What's next	11
2. Architecture	13
2.1. Overview	13
2.2. Backend	14
2.3. Reader strategy	18
3. Configuration	20
3.1. Enabling Hibernate Search and automatic indexing	20
3.2. Configuring the IndexManager	20
3.3. Directory configuration	21
3.4. Worker configuration	29
3.5. Reader strategy configuration	38
3.6. Serialization	39
3.7. Exception handling	40
3.8. Lucene configuration	41
3.9. Metadata API	48
3.10. Hibernate Search as a WildFly module	49
4. Mapping entities to the index structure	55
4.1. Mapping an entity	55
4.2. Boosting	73
4.3. Analysis	75
4.4. Bridges	85
4.5. Conditional indexing	99
4.6. Providing your own id	102
4.7. Programmatic API	102
5. Querying	119
5.1. Building queries	120
5.2. Retrieving the results	142
5.3. Filters	145
5.4. Faceting	150

5.5. Optimizing the query process	158
6. Manual index changes	159
6.1. Adding instances to the index	159
6.2. Deleting instances from the index	159
6.3. Rebuilding the whole index	160
6.4. Integration with JSR-352	166
7. Index Optimization	178
7.1. Automatic optimization	179
7.2. Manual optimization	179
7.3. Adjusting optimization	180
8. Monitoring	181
8.1. JMX	181
9. Spatial	182
9.1. Enable indexing of Spatial Coordinates	182
9.2. Performing Spatial Queries	185
9.3. Multiple Coordinate pairs	189
9.4. Insight: implementation details of spatial hashes indexing	190
10. Advanced features	194
10.1. Accessing the SearchFactory	194
10.2. Accessing the SearchIntegrator	194
10.3. Using an IndexReader	195
10.4. Accessing a Lucene Directory	196
10.5. Sharding indexes	196
10.6. Sharing indexes	199
10.7. Using external services	200
10.8. Customizing Lucene's scoring formula	202
10.9. Multi-tenancy	203
11. Integration with Elasticsearch	205
11.1. Status	205
11.2. Goal of the Elasticsearch integration	205
11.3. Getting started and configuration	206
11.4. Mapping and indexing	215
11.5. Queries	226
11.6. Index optimization	231
11.7. Direct access to the Elasticsearch client	231
11.8. Limitations	232
11.9. Specific versions of Elasticsearch to avoid	232
11.10. Acknowledgment	233
12. Further reading	234
13. Credits	235

Preface

Full text search engines like Apache Lucene are very powerful technologies to add efficient free text search capabilities to applications. However, Lucene suffers several mismatches when dealing with object domain models. Amongst other things indexes have to be kept up to date and mismatches between index structure and domain model as well as query mismatches have to be avoided.

Hibernate Search addresses these shortcomings - it indexes your domain model with the help of a few annotations, takes care of database/index synchronization and brings back regular managed objects from free text queries. To achieve this Hibernate Search is combining the power of [Hibernate](#) and [Apache Lucene](#).

Since version [5.6](#) Hibernate Search sports an experimental integration with [Elasticsearch](#). Elasticsearch is built on Apache Lucene so we can now expose very similar features, making most of this reference documentation a valid guide to both approaches. For details specific to Elasticsearch, jump to [Integration with Elasticsearch](#). The other sections of this guide will assume you're using Lucene without the Elasticsearch indirection.

Chapter 1. Getting started

Welcome to Hibernate Search. The following chapter will guide you through the initial steps required to integrate Hibernate Search into an existing Hibernate ORM enabled application. In case you are a Hibernate new timer we recommend you start [here](#).

1.1. System Requirements

Table 1. System requirements

Java Runtime	Requires Java version 8 or greater. You can download a Java Runtime for Windows/Linux/Solaris here .
Hibernate Search	<code>hibernate-search-5.10.13.Final.jar</code> and all runtime dependencies. You can get the jar artifacts either from the <code>dist/lib</code> directory of the Hibernate Search distribution or you can download them from the Maven Central repository.
Hibernate ORM	You will need <code>hibernate-core-5.3.28.Final.jar</code> and its dependencies (either from the distribution bundle or the maven repository).
JPA 2.2	Hibernate Search can be used without JPA but the following instructions will use JPA annotations for basic entity configuration (<code>@Entity</code> , <code>@Id</code> , <code>@OneToMany</code> ,...).

1.2. Migration notes

If you are upgrading an existing application from an earlier version of Hibernate Search to the latest release, make sure to check the out the [migration guide](#).

1.3. Required libraries

The Hibernate Search library is split in several modules to allow you to pick the minimal set of dependencies you need. It requires Apache Lucene, Hibernate ORM and some standard APIs such as the Java Persistence API and the Java Transactions API. Other dependencies are optional, providing additional integration points. To get the correct jar files on your classpath we highly recommend to use a dependency manager such as [Maven](#), or similar tools such as [Gradle](#) or [Ivy](#). These alternatives are also able to consume the artifacts from the [Using Maven](#) section.

1.3.1. Using Maven

The Hibernate Search artifacts can be found in Maven's [Central Repository](#) but are released first in the [JBoss Maven Repository](#). See also the [Maven Getting Started wiki page](#) to use the JBoss repository.

All you have to add to your pom.xml is:

Example 1. Maven artifact identifier for Hibernate Search

```
<dependency>
  <groupId>org.hibernate</groupId>
  <artifactId>hibernate-search-orm</artifactId>
  <version>5.10.13.Final</version>
</dependency>
```

Example 2. Optional Maven dependencies for Hibernate Search

```
<!-- Infinispan integration: -->
<dependency>
  <groupId>org.infinispan</groupId>
  <artifactId>infinispan-directory-provider</artifactId>
  <version>9.2.2.Final</version>
</dependency>
```

Only the *hibernate-search-orm* dependency is mandatory. *infinispan-directory-provider* is only required if you want to use Infinispan to store the Lucene indexes.

1.3.2. Manual library management

You can download zip bundles from Sourceforge containing all needed [Hibernate Search](#) dependencies. This includes - among others - the latest compatible version of Hibernate ORM. However, only the essential parts you need to start experimenting with are included. You will probably need to combine this with downloads from the other projects, for example the [Hibernate ORM distribution on Sourceforge](#) also provides the modules to enable caching or use a connection pool.

1.4. Deploying on WildFly

If you are creating an application to be deployed on WildFly you're lucky: Hibernate Search is included in the application server. This means that you don't need to package it along with your application, unless you want to use a different version than the one included. Since WildFly version 10 the embedded version of Hibernate Search is automatically activated when your application uses it. See [Hibernate Search as a WildFly module](#) for details.

If you want to use another version than the one included in WildFly, have a look at [Hibernate Search as a WildFly module](#).

1.5. Configuration

Once you have added all required dependencies to your application you have to add a couple of properties to your Hibernate configuration file. If you are using Hibernate directly this can be done in `hibernate.properties` or `hibernate.cfg.xml`. If you are using Hibernate via JPA you can also add the properties to `persistence.xml`. The good news is that for standard use most properties offer a sensible default. An example `persistence.xml` configuration could look like this:

Example 3. Basic configuration options to be added to `hibernate.properties`, `hibernate.cfg.xml` or `persistence.xml`

```
...
<property name="hibernate.search.default.directory_provider"
          value="filesystem"/>

<property name="hibernate.search.default.indexBase"
          value="/var/lucene/indexes"/>
...
```

First you have to tell Hibernate Search which `DirectoryProvider` to use. This can be achieved by setting the `hibernate.search.default.directory_provider` property. Apache Lucene has the notion of a `Directory` to store the index files. Hibernate Search handles the initialization and configuration of a Lucene `Directory` instance via a `DirectoryProvider`. In this tutorial we will use a directory provider which stores the index on the file system. This will give us the ability to inspect the Lucene indexes created by Hibernate Search (eg via [Luke](#)). Once you have a working configuration you can start experimenting with other directory providers (see [Directory configuration](#)). You also have to specify the default base directory for all indexes via `hibernate.search.default.indexBase`. This defines the path where indexes are stored.

Let's assume that your application contains the Hibernate managed classes `example.Book` and `example.Author` and you want to add free text search capabilities to your application in order to search the books contained in your database.

Example 4. Example entities *Book* and *Author* before adding Hibernate Search specific annotations

```
package example;
...
@Entity
public class Book {

    @Id
    @GeneratedValue
    private Integer id;

    private String title;

    private String subtitle;

    @ManyToMany
    private Set<Author> authors = new HashSet<Author>();

    private Date publicationDate;

    public Book() {}

    // standard getters/setters follow
    ...
}
```

```
package example;
...
@Entity
public class Author {

    @Id
    @GeneratedValue
    private Integer id;

    private String name;

    public Author() {}

    // standard getters/setters follow
    ...
}
```

To achieve this you have to add a few annotations to the *Book* and *Author* class. The first annotation *@Indexed* marks *Book* as indexable. By design Hibernate Search needs to store an *untokenized* id in the index to ensure index uniqueness for a given entity (for now don't worry if you don't know what *untokenized* means, it will soon be clear).

Next you have to mark the fields you want to make searchable. Let's start with *title* and *subtitle* and annotate both with *@Field*. The parameter *index=Index.YES* will ensure that the text will be indexed, while *analyze=Analyze.YES* ensures that the text will be analyzed using the default Lucene analyzer. Usually, analyzing or tokenizing means chunking a sentence into individual words and potentially excluding common words like "a" or "the". We will talk more about analyzers a little later on. The third parameter we specify is *store=Store.NO*, which ensures that the actual data will not be stored in the index. Whether data is stored in the index or not has nothing to do with the ability to search for it. It is not necessary to store fields in the index to allow Lucene to search for them: the

benefit of storing them is the ability to retrieve them via projections (see [Projection](#)).

Without projections, Hibernate Search will per default execute a Lucene query in order to find the database identifiers of the entities matching the query criteria and use these identifiers to retrieve managed objects from the database. The decision for or against projection has to be made on a case by case basis.

Note that `index=Index.YES`, `analyze=Analyze.YES` and `store=Store.NO` are the default values for these parameters and could be omitted.

After this short look under the hood let's go back to annotating the `Book` class. Another annotation we have not yet discussed is `@DateBridge`. This annotation is one of the built-in field bridges in Hibernate Search. The Lucene index is mostly string based, with special support for encoding numbers. Hibernate Search must convert the data types of the indexed fields to their respective Lucene encoding and vice versa. A range of predefined bridges is provided for this purpose, including the `DateBridge` which will convert a `java.util.Date` into a numeric value (a `long`) with the specified resolution. For more details see [Built-in bridges](#).

This leaves us with `@IndexedEmbedded`. This annotation is used to index associated entities (`@ManyToMany`, `.*ToOne`, `@Embedded` and `@ElementCollection`) as part of the owning entity. This is needed since a Lucene index document is a flat data structure which does not know anything about object relations. To ensure that the author names will be searchable you have to make sure that the names are indexed as part of the book itself. On top of `@IndexedEmbedded` you will also have to mark the fields of the associated entity you want to have included in the index with `@Field`. For more details see [Embedded and associated objects](#).

These settings should be sufficient for now. For more details on entity mapping refer to [Mapping an entity](#).

Example 5. Example entities after adding Hibernate Search annotations

```
package example;
...
@Entity
@Indexed
public class Book {

    @Id
    @GeneratedValue
    private Integer id;

    @Field(index=Index.YES, analyze=Analyze.YES, store=Store.NO)
    private String title;

    @Field(index=Index.YES, analyze=Analyze.YES, store=Store.NO)
    private String subtitle;

    @Field(index = Index.YES, analyze=Analyze.NO, store = Store.YES)
    @DateBridge(resolution = Resolution.DAY)
    private Date publicationDate;

    @IndexedEmbedded
    @ManyToMany
    private Set<Author> authors = new HashSet<Author>();
    public Book() {
    }

    // standard getters/setters follow here
    ...
}
```

```
@Entity
public class Author {

    @Id
    @GeneratedValue
    private Integer id;

    @Field
    private String name;

    public Author() {
    }

    // standard getters/setters follow here
    ...
}
```

1.6. Indexing

Hibernate Search will transparently index every entity persisted, updated or removed through Hibernate ORM. However, you have to create an initial Lucene index for the data already present in your database. Once you have added the above properties and annotations it is time to trigger an initial batch index of your books. You can achieve this by using one of the following code snippets (see also [Rebuilding the whole index](#)):

Example 6. Using Hibernate Session to index data

```
FullTextSession fullTextSession = org.hibernate.search.Search.getFullTextSession(session);
fullTextSession.createIndexer().startAndWait();
```

Example 7. Using JPA to index data

```
EntityManager em = entityManagerFactory.createEntityManager();
FullTextEntityManager fullTextEntityManager = org.hibernate.search.jpa.Search
    .getFullTextEntityManager(em);
fullTextEntityManager.createIndexer().startAndWait();
```

After executing the above code, you should be able to see a Lucene index under `/var/lucene/indexes/example.Book` (or based on a different path depending how you configured the property `hibernate.search.default.directory_provider`).

Go ahead and inspect this index with [Luke](#): it will help you to understand how Hibernate Search works.

1.7. Searching

Now it is time to execute a first search. The general approach is to create a Lucene query, either via the Lucene API ([Building a Lucene query using the Lucene API](#)) or via the Hibernate Search query DSL ([Building a Lucene query with the Hibernate Search query DSL](#)), and then wrap this query into a `org.hibernate.Query` in order to get all the functionality one is used to from the Hibernate API. The following code will prepare a query against the indexed fields, execute it and return a list of `Book` instances.

Example 8. Using Hibernate Session to create and execute a search

```
FullTextSession fullTextSession = Search.getFullTextSession(session);
Transaction tx = fullTextSession.beginTransaction();

// create native Lucene query using the query DSL
// alternatively you can write the Lucene query using the Lucene query parser
// or the Lucene programmatic API. The Hibernate Search DSL is recommended though
QueryBuilder qb = fullTextSession.getSearchFactory()
    .buildQueryBuilder().forEntity(Book.class).get();
org.apache.lucene.search.Query query = qb
    .keyword()
    .onFields("title", "subtitle", "authors.name")
    .matching("Java rocks!")
    .createQuery();

// wrap Lucene query in a org.hibernate.Query
org.hibernate.Query hibQuery =
    fullTextSession.createFullTextQuery(query, Book.class);

// execute search
List result = hibQuery.list();

tx.commit();
session.close();
```

Example 9. Using JPA to create and execute a search

```
EntityManager em = entityManagerFactory.createEntityManager();
FullTextEntityManager fullTextEntityManager =
    org.hibernate.search.jpa.Search.getFullTextEntityManager(em);
em.getTransaction().begin();

// create native Lucene query using the query DSL
// alternatively you can write the Lucene query using the Lucene query parser
// or the Lucene programmatic API. The Hibernate Search DSL is recommended though
QueryBuilder qb = fullTextEntityManager.getSearchFactory()
    .buildQueryBuilder().forEntity(Book.class).get();
org.apache.lucene.search.Query query = qb
    .keyword()
    .onFields("title", "subtitle", "authors.name")
    .matching("Java rocks!")
    .createQuery();

// wrap Lucene query in a javax.persistence.Query
javax.persistence.Query persistenceQuery =
    fullTextEntityManager.createFullTextQuery(query, Book.class);

// execute search
List result = persistenceQuery.getResultList();

em.getTransaction().commit();
em.close();
```

1.8. Analyzer

Let's make things a little more interesting now. Assume that one of your indexed book entities has the title "Refactoring: Improving the Design of Existing Code" and you want to get hits for all of the

following queries: "refactor", "refactors", "refactored" and "refactoring". In Lucene this can be achieved by choosing an analyzer class which applies word stemming during the indexing **as well as** the search process. Hibernate Search offers several ways to configure the analyzer to be used (see [Default analyzer and analyzer by class](#)):

- Setting the `hibernate.search.analyzer` property in the configuration file. The specified class will then be the default analyzer.
- Setting the `@Analyzer` annotation at the entity level.
- Setting the `@Analyzer` annotation at the field level.

When using the `@Analyzer` annotation one can either specify the fully qualified classname of the analyzer to use or one can refer to an analyzer definition defined by the `@AnalyzerDef` annotation. In the latter case the analyzer framework with its factories approach is utilized.

To find out more about the factory classes available you can either browse the Lucene JavaDoc or read the corresponding section on the [Solr Wiki](#).

You can use `@AnalyzerDef` on any:

- `@Indexed` entity regardless of where the analyzer is applied to;
- parent class of an `@Indexed` entity;
- `package-info.java` of a package containing an `@Indexed` entity.

This implies that analyzer definitions are global and their names must be unique. The `@AnalyzerDef` is a repeatable annotation, so you can declare multiple analyzer definitions on each of these types.

Why the reference to the Apache Solr wiki?



The analyzer factory framework was originally created in the Apache Solr project. Most of these implementations have been moved to Apache Lucene, but the documentation for these additional analyzers can still be found in the Solr Wiki. You might find other documentation referring to the "Solr Analyzer Framework"; just remember you don't need to depend on Apache Solr anymore: the required classes are part of the core Lucene distribution.

In the example below a `StandardTokenizerFactory` is used followed by two filter factories, `LowerCaseFilterFactory` and `SnowballPorterFilterFactory`. The standard tokenizer splits words at punctuation characters and hyphens. It is a good general purpose tokenizer. For indexing email addresses or internet hostnames it is not the best fit as it would split them up. You may either make use of Lucene's `ClassicTokenizerFactory` in such cases or implement a custom tokenizer and factory. The lowercase filter converts to lowercase the letters in each token whereas the snowball filter finally applies language specific stemming.

Generally, when using the Analyzer Framework you have to start with a tokenizer followed by an

arbitrary number of filters.



Support for defining analyzers or normalizers through annotations is deprecated, and will be removed in Hibernate Search 6, so the following example is out of date.

Use [a definition provider](#) instead.

Example 10. Using `@AnalyzerDef` and the Analyzer Framework to define and use an analyzer

```
@Entity
@Indexed
@AnalyzerDef(name = "customanalyzer",
    tokenizer = @TokenizerDef(factory = StandardTokenizerFactory.class),
    filters = {
        @TokenFilterDef(factory = LowerCaseFilterFactory.class),
        @TokenFilterDef(factory = SnowballPorterFilterFactory.class, params = {
            @Parameter(name = "language", value = "English")
        })
    })
public class Book {

    @Id
    @GeneratedValue
    @DocumentId
    private Integer id;

    @Field(analyzer = @Analyzer(definition = "customanalyzer"))
    private String title;

    @Field(analyzer = @Analyzer(definition = "customanalyzer"))
    private String subtitle;

    @IndexedEmbedded
    @ManyToMany
    private Set<Author> authors = new HashSet<Author>();

    @Field(index = Index.YES, analyze = Analyze.NO, store = Store.YES)
    @DateBridge(resolution = Resolution.DAY)
    private Date publicationDate;

    public Book() {
    }

    // standard getters/setters follow here
    ...
}
```

Using `@AnalyzerDef` only defines an Analyzer, you still have to apply it to entities and or properties using `@Analyzer`. Like in the above example the `customanalyzer` is defined but not applied on the entity: it's applied on the `title` and `subtitle` properties only. An analyzer definition is global, so you can define it on any entity and reuse the definition on other entities.

1.9. What's next

The above paragraphs helped you getting an overview of Hibernate Search. The next step after this

tutorial is to get more familiar with the overall architecture of Hibernate Search ([Architecture](#)) and explore the basic features in more detail. Two topics which were only briefly touched in this tutorial were analyzer configuration ([Default analyzer and analyzer by class](#)) and field bridges ([Bridges](#)). Both are important features required for more fine-grained indexing. More advanced topics cover clustering ([JMS Master/Slave back end](#), [Infinispan Directory configuration](#)) and large index handling ([Sharding indexes](#)).

Chapter 2. Architecture

2.1. Overview

Hibernate Search consists of an indexing component as well as an index search component. Both are backed by Apache Lucene by default, alternatively Elasticsearch can be used.

Each time an entity is inserted, updated or removed in/from the database, Hibernate Search keeps track of this event (through the Hibernate event system) and schedules an index update. All these updates are handled without you having to interact with the Apache Lucene APIs directly (see [Enabling Hibernate Search and automatic indexing](#)). Instead, the interaction with the underlying Lucene indexes is handled via so called IndexManagers.

Each Lucene index is managed by one index manager which is uniquely identified by name. In most cases there is also a one to one relationship between an indexed entity and a single IndexManager. The exceptions are the use cases of index sharding and index sharing. The former can be applied when the index for a single entity becomes too big and indexing operations are slowing down the application. In this case a single entity is indexed into multiple indexes each with its own index manager (see [Sharding indexes](#)). The latter, index sharing, is the ability to index multiple entities into the same Lucene index (see [Sharing indexes](#)).

The index manager abstracts from the specific index configuration. In the case of the default index manager this includes details about the selected backend, the configured reader strategy and the chosen DirectoryProvider. These components will be discussed in greater detail later on. It is recommended that you start with the default index manager which uses different Lucene Directory types to manage the indexes (see [Directory configuration](#)). You can, however, also provide your own IndexManager implementation (see [Configuring the IndexManager](#)).

Once the index is created, you can search for entities and return lists of managed entities saving you the tedious object to Lucene Document mapping. The same persistence context is shared between Hibernate and Hibernate Search. As a matter of fact, the `FullTextSession` is built on top of the Hibernate Session so that the application code can use the unified `org.hibernate.Query` or `javax.persistence.Query` APIs exactly the same way a HQL, JPA-QL or native query would do.

To be more efficient Hibernate Search batches the write interactions with the Lucene index. There are currently two types of batching.

Outside a transaction, the index update operation is executed right after the actual database operation. This is really a no batching setup.

In the case of an ongoing transaction, the index update operation is scheduled for the transaction commit phase and discarded in case of transaction rollback. The batching scope is the transaction. There are two immediate benefits:

- Performance: Lucene indexing works better when operation are executed in batch.
- ACIDity: The work executed has the same scoping as the one executed by the database transaction and is executed if and only if the transaction is committed. This is not ACID in the strict sense of it, but ACID behavior is rarely useful for full text search indexes since they can be rebuilt from the source at any time.

You can think of those two batch modes (no scope vs transactional) as the equivalent of the (infamous) autocommit vs transactional behavior. From a performance perspective, the *in transaction* mode is recommended. The scoping choice is made transparently. Hibernate Search detects the presence of a transaction and adjust the scoping (see [Worker configuration](#)).



It is recommended - for both your database and Hibernate Search - to execute your operations in a transaction, be it JDBC or JTA.



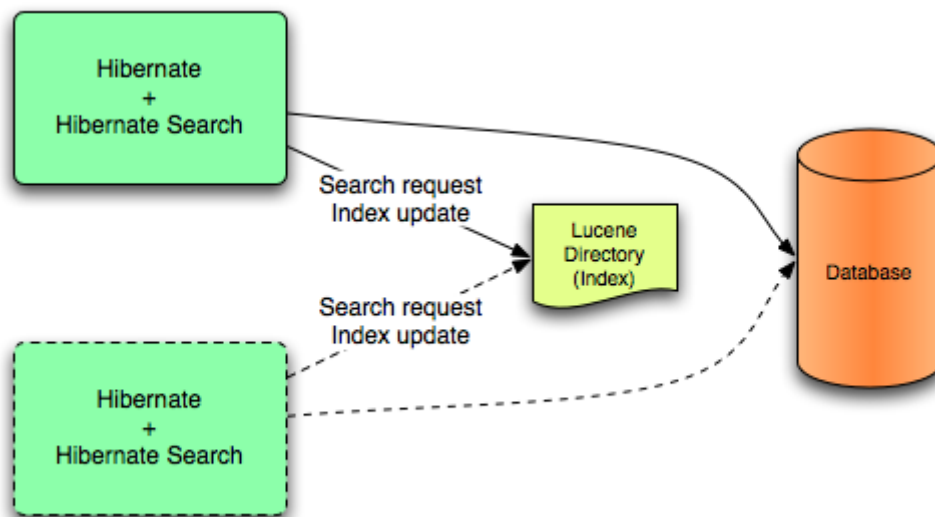
Hibernate Search works perfectly fine in the Hibernate / EntityManager long conversation pattern aka. atomic conversation.

2.2. Backend

Hibernate Search offers the ability to let the batched work be processed by different backends. Several backends are provided out of the box and you have the option to plug in your own. It is important to understand that in this context backend encompasses more than just the configuration option `hibernate.search.default.worker.backend`. This property just specifies an implementation of the `BackendQueueProcessor` interface (or the `Backend` interface, see [the configuration options](#)) which is a part of a backend configuration. In most cases, however, additional configuration settings are needed to successfully configure a specific backend setup, like for example the JMS backend.

2.2.1. Lucene

In this mode, all index update operations applied on a given node (JVM) will be executed to the Lucene directories (through the directory providers) by the same node. This mode is typically used in non clustered environment or in clustered environments where the directory store is shared.



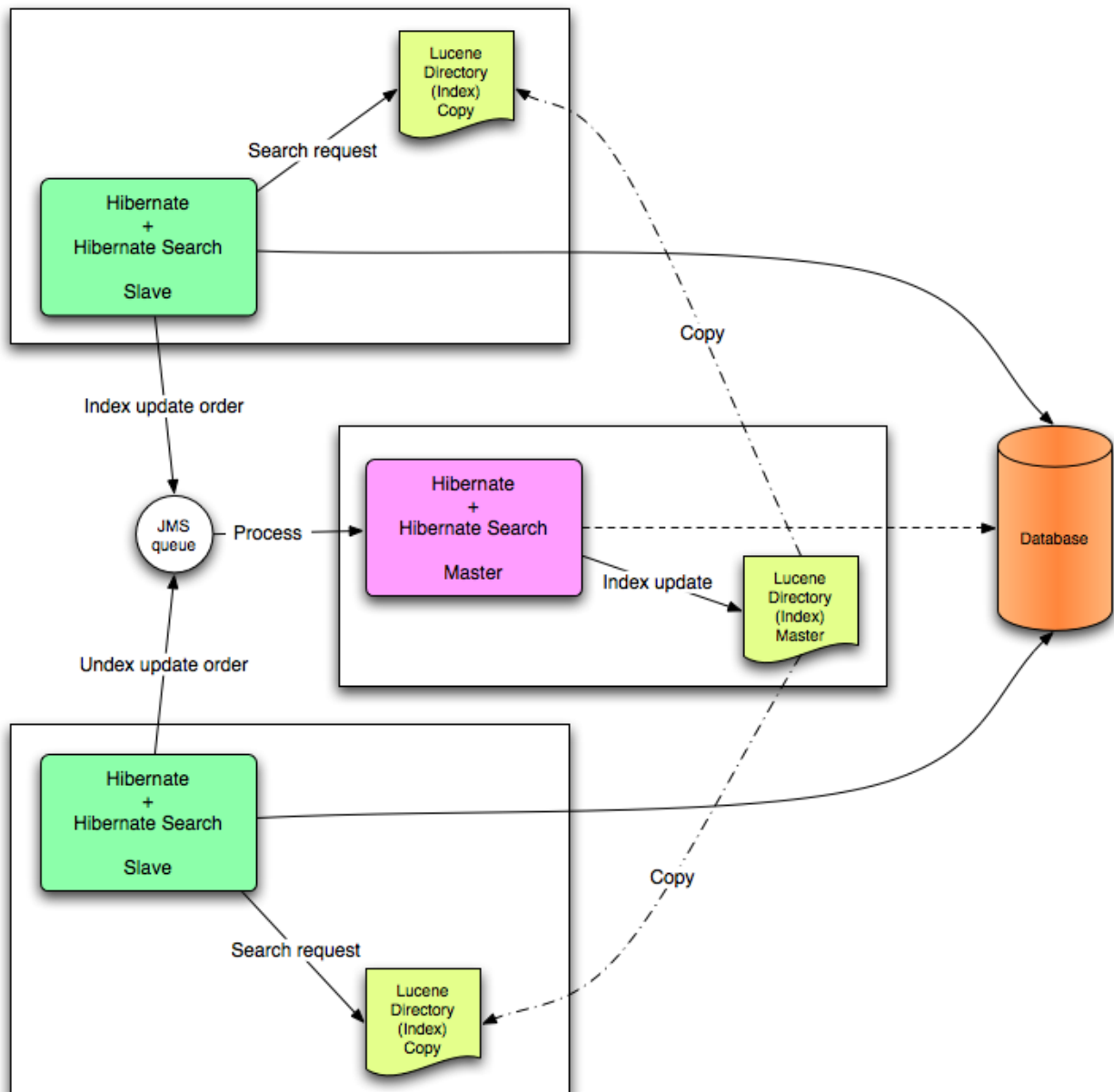
This mode targets non clustered applications, or clustered applications where the Directory is taking care of the locking strategy.

The main advantage is simplicity and immediate visibility of the changes in Lucene queries (a requirement in some applications).

An alternative backend viable for non-clustered and non-shared index configurations is the near- real-time backend.

2.2.2. JMS

All index update operations applied on a given node are sent to a JMS queue. A unique reader will then process the queue and update the master index. The master index is then replicated on a regular basis to the slave copies. This is known as the master/slaves pattern. The master is the sole responsible for updating the Lucene index. The slaves can accept read as well as write operations. However, while they process the read operations on their local index copy, they will delegate the update operations to the master.



This mode targets clustered environments where throughput is critical, and index update delays are affordable. Reliability is ensured by the JMS provider and by having the slaves working on a local copy of the index.

The JMS integration can be transactional. With this backend (and currently only this backend) you can have Hibernate Search send the indexing work into the queue within the same transaction applying changes to the relational database. This option requires you to use an XA transaction.

By default this backend's transactional capabilities are disabled: messages will be enqueued as a post-transaction event, consistently with other backends. To change this configuration see also [Worker configuration](#).

2.2.3. JGroups

The JGroups based backend works similar to the JMS one and is designed after the same

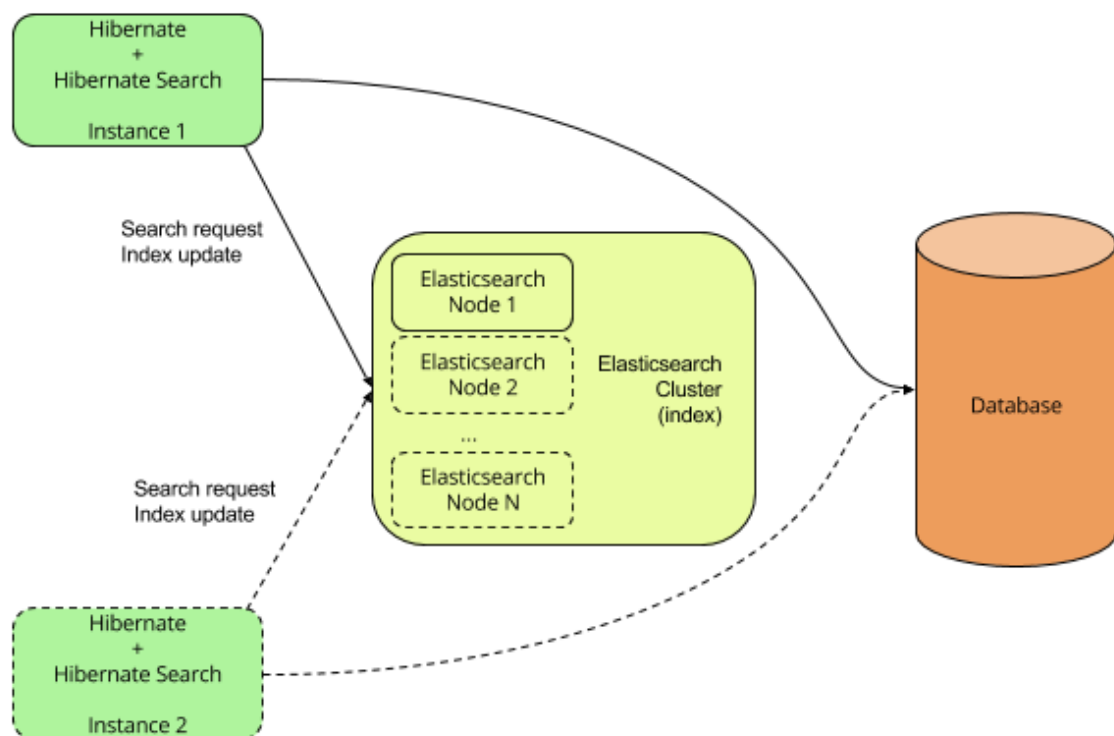
master/slaves pattern. However, instead of JMS, the JGroups toolkit is used as a replication mechanism. This backend can be used as an alternative to JMS when response time is critical, but i.e. JNDI service is not available.

Note that while JMS can usually be configured to use persistent queues, JGroups talks directly to other nodes over the network. Message delivery to other reachable nodes is guaranteed by the protocol, but if no master node is available, index operations are silently discarded. Also if the sending node crashes before the delivery is performed, the events might be lost. This backend can be configured to use asynchronous messages, or to wait for each indexing operation to be completed on the remote node before returning.

The JGroups backend can be configured with static master or slave roles, or can be setup to perform an auto-election of the master. This mode is particularly useful to have the system automatically pick a new master in case of failure, but during a reelection process some indexing operations might be lost. For this reason this mode is not suited for use cases requiring strong consistency guarantees. When configured to perform an automatic election, the master node is defined as an hash on the index name: the role is therefore possibly different for each index or shard.

2.2.4. Elasticsearch

In this mode, the index is not on the application server, but in an Elasticsearch cluster. Update operations are sent to the cluster and executed remotely, and so are search queries.



This mode allows to easily set up a clustered application, without the advanced configuration required by the JMS and JGroups modes.

More information can be found in the [Elasticsearch integration section](#).

2.3. Reader strategy

When executing a query, Hibernate Search interacts with the Apache Lucene indexes through a reader strategy. Choosing a reader strategy will depend on the profile of the application (frequent updates, read mostly, asynchronous index update etc). See also [Reader strategy configuration](#)

2.3.1. shared

With this strategy, Hibernate Search will share the same `IndexReader`, for a given Lucene index, across multiple queries and threads provided that the `IndexReader` is still up-to-date. If the `IndexReader` is not up-to-date, a new one is opened and provided. Each `IndexReader` is made of several `SegmentReaders`. This strategy only reopens segments that have been modified or created after last opening and shares the already loaded segments from the previous instance. This approach is quite efficient and guarantees that each query is run on the most recent index snapshot; the drawback is that for every query the strategy will have to verify if the `IndexReader` is still fresh, and if not perform a refresh; such a refresh is typically a cheap operation but if you have a significant amount of writes and queries happening concurrently then one of the other strategies might be preferred. This strategy is the default.

The name of this strategy is `shared`.

2.3.2. not-shared

Every time a query is executed, a Lucene `IndexReader` is opened. This strategy is not efficient since opening and warming up an `IndexReader` can be a relatively expensive operation, but is very simple code. Use it as an example implementation if you're interested to learn about Hibernate Search internals or want to extend it.

The name of this strategy is `not-shared`.

2.3.3. async

This implementation keeps an `IndexReader` open and ready to be used by all queries, while a background thread periodically verifies if there is need to open a fresh one, replaces the active one and disposes the outdated one. The frequency of checks - and refreshing - of this background thread is configurable, but defaults to 5000 milliseconds. The drawback of this design is that queries are effectively run on an index snapshot which might be approximately 5 seconds out of date (assuming the refresh period is not reconfigured); the benefit is that if your application writes frequently to the index, the query performance will be more consistent.

The name of this strategy is **async**.

2.3.4. Custom

You can write your own reader strategy that suits your application needs by implementing `org.hibernate.search.reader.ReaderProvider`. The implementation must be thread safe.

Chapter 3. Configuration

3.1. Enabling Hibernate Search and automatic indexing

Let's start with the most basic configuration question - how do I enable Hibernate Search?

3.1.1. Enabling Hibernate Search

The good news is that Hibernate Search is enabled out of the box when detected on the classpath by Hibernate ORM. If for some reason you need to disable it set `hibernate.search.autoregister_listeners` to false. Note that there is no performance penalty when the listeners are enabled but no entities are annotated as indexed.

3.1.2. Automatic indexing

By default, every time an object is inserted, updated or deleted through Hibernate, Hibernate Search updates the according Lucene index. It is sometimes desirable to disable that features if either your index is read-only or if index updates are done in a batch way (see [Rebuilding the whole index](#)).

To disable event based indexing, set

```
hibernate.search.indexing_strategy = manual
```



In most case, the JMS backend provides the best of both world, a lightweight event based system keeps track of all changes in the system, and the heavyweight indexing process is done by a separate process or machine.

3.2. Configuring the IndexManager

The role of the index manager component is described in [Architecture](#). Hibernate Search provides two possible implementations for this interface to choose from.

- **directory-based**: the default implementation which uses the Lucene Directory abstraction to manage index files.
- **near-real-time**: avoid flushing writes to disk at each commit. This index manager is also Directory based, but also makes uses of Lucene's NRT functionality.

To select an alternative you specify the property:

```
hibernate.search.[default|<indexname>].indexmanager = near-real-time
```

3.2.1. directory-based

The default IndexManager implementation. This is the one mostly referred to in this documentation. It is highly configurable and allows you to select different settings for the reader strategy, back ends and directory providers. Refer to [Directory configuration](#), [Worker configuration](#) and [Reader strategy configuration](#) for more details.

3.2.2. near-real-time

The NRTIndexManager is an extension of the default IndexManager, leveraging the Lucene NRT (Near Real Time) features for extreme low latency index writes. As a trade-off it requires a non-clustered and non-shared index. In other words, it will ignore configuration settings for alternative back ends other than **lucene** and will acquire exclusive write locks on the Directory.

To achieve these low latency writes, the IndexWriter will not flush every change to disk. Queries will be allowed to read updated state from the unflushed index writer buffers; the downside of this strategy is that if the application crashes or the IndexWriter is otherwise killed you'll have to rebuild the indexes as some updates might be lost.

Because of these downsides, and because a master node in cluster can be configured for good performance as well, the NRT configuration is only recommended for non clustered websites with a limited amount of data.

3.2.3. Custom

It is also possible to configure a custom IndexManager implementation by specifying the fully qualified class name of your custom implementation. This implementation must have a no-argument constructor:

```
hibernate.search.[default|<indexname>].indexmanager = my.corp.myapp.CustomIndexManager
```



Your custom index manager implementation doesn't need to use the same components as the default implementations. For example, you can delegate to a remote indexing service which doesn't expose a Directory interface.

3.3. Directory configuration

As we have seen in [Configuring the IndexManager](#) the default index manager uses Lucene's notion of a Directory to store the index files. The Directory implementation can be customized and Lucene comes

bundled with a file system and an in-memory implementation. `DirectoryProvider` is the Hibernate Search abstraction around a Lucene Directory and handles the configuration and the initialization of the underlying Lucene resources. [List of built-in DirectoryProvider](#) shows the list of the directory providers available in Hibernate Search together with their corresponding options.

To configure your `DirectoryProvider` you have to understand that each indexed entity is associated to a Lucene index (except of the case where multiple entities share the same index - [Sharing indexes](#)). The name of the index is given by the index property of the `@Indexed` annotation. If the index property is not specified the fully qualified name of the indexed class will be used as name (recommended).

Knowing the index name, you can configure the directory provider and any additional options by using the prefix `hibernate.search.<indexname>`. The name default (`hibernate.search.default`) is reserved and can be used to define properties which apply to all indexes. [Configuring directory providers](#) shows how `hibernate.search.default.directory_provider` is used to set the default directory provider to be the filesystem one. `hibernate.search.default.indexBase` sets then the default base directory for the indexes. As a result the index for the entity `Status` is created in `/usr/lucene/indexes/org.hibernate.example.Status`.

The index for the `Rule` entity, however, is using an in-memory directory, because the default directory provider for this entity is overridden by the property `hibernate.search.Rules.directory_provider`.

Finally the `Action` entity uses a custom directory provider `CustomDirectoryProvider` specified via `hibernate.search.Actions.directory_provider`.

Example 11. Specifying the index name

```
package org.hibernate.example;

@Indexed
public class Status { ... }

@Indexed(index="Rules")
public class Rule { ... }

@Indexed(index="Actions")
public class Action { ... }
```

Example 12. Configuring directory providers

```
hibernate.search.default.directory_provider = filesystem
hibernate.search.default.indexBase = /usr/lucene/indexes
hibernate.search.Rules.directory_provider = local-heap
hibernate.search.Actions.directory_provider = com.acme.hibernate.CustomDirectoryProvider
```



Using the described configuration scheme you can easily define common rules like the directory provider and base directory, and override those defaults later on on a per index basis.

Table 2. List of built-in *DirectoryProvider*

Name and description	Properties
local-heap: Directory using the local JVM heap. Local heap directories and all contained indexes are lost when the JVM shuts down. This option is only provided for use in testing configurations with small (trivial) indexes and low concurrency, where it could slightly improve performance. In setups requiring larger indexes and/or high concurrency, a file system based directory (see below) will achieve better performance. The directory will be uniquely identified (in the same deployment unit) by the <code>@Indexed.index</code> element	none
filesystem: File system based directory. The directory used will be <code><indexBase>/<indexName></code>	<code>indexBase</code> : base directory <code>indexName</code>: override <code>@Indexed.index</code> (useful for sharded indexes) <code>locking_strategy</code> : optional, see LockFactory configuration <code>filesystem_access_type</code>: allows to determine the exact type of <code>FSDirectory</code> implementation used by this <code>DirectoryProvider</code>. Allowed values are <code>auto</code> (the default value, selects <code>NIOFSDirectory</code> on non Windows systems, <code>SimpleFSDirectory</code> on Windows), <code>simple</code> (<code>SimpleFSDirectory</code>), <code>nio</code> (<code>NIOFSDirectory</code>), <code>mmap</code> (<code>MMapDirectory</code>). Make sure to refer to Javadocs of these Directory implementations before changing this setting. Even though <code>NIOFSDirectory</code> or <code>MMapDirectory</code> can bring substantial performance boosts they also have their issues.

Name and description	Properties
filesystem-master: File system based directory. Like <code>filesystem</code>. It also copies the index to a source directory (aka copy directory) on a regular basis. The recommended value for the refresh period is (at least) 50% higher than the time to copy the information (default 3600 seconds - 60 minutes). Note that the copy is based on an incremental copy mechanism reducing the average copy time. DirectoryProvider typically used on the master node in a JMS back end cluster.	<code>indexBase</code>: base directory <code>indexName</code>: override <code>@Indexed.index</code> (useful for sharded indexes) <code>sourceBase</code>: source (copy) base directory. <code>source</code>: source directory suffix (default to <code>@Indexed.index</code>). The actual source directory name being <code><sourceBase>/<source></code> <code>refresh</code>: refresh period in seconds (the copy will take place every refresh seconds). If a copy is still in progress when the following refresh period elapses, the second copy operation will be skipped. <code>locking_strategy</code>: optional, see LockFactory configuration <code>filesystem_access_type</code>: allows to determine the exact type of FSDirectory implementation used by this DirectoryProvider. Allowed values are <code>auto</code> (the default value, selects NIOFSDirectory on non Windows systems, SimpleFSDirectory on Windows), <code>simple</code> (SimpleFSDirectory), <code>nio</code> (NIOFSDirectory), <code>mmap</code> (MMapDirectory). Make sure to refer to Javadocs of these Directory implementations before changing this setting. Even though NIOFSDirectory or MMapDirectory can bring substantial performance boosts they also have their issues.

Name and description	Properties
filesystem-slave: File system based directory. Like <code>filesystem</code>, but retrieves a master version (source) on a regular basis. To avoid locking and inconsistent search results, 2 local copies are kept. The recommended value for the refresh period is (at least) 50% higher than the time to copy the information (default 3600 seconds - 60 minutes). Note that the copy is based on an incremental copy mechanism reducing the average copy time. If a copy is still in progress when refresh period elapses, the second copy operation will be skipped. DirectoryProvider typically used on slave nodes using a JMS back end.	<code>indexBase</code>: Base directory <code>indexName</code>: override <code>@Indexed.index</code> (useful for sharded indexes) <code>sourceBase</code>: Source (copy) base directory. <code>source</code>: Source directory suffix (default to <code>@Indexed.index</code>). The actual source directory name being <code><sourceBase>/<source></code> <code>refresh</code>: refresh period in second (the copy will take place every refresh seconds). <code>locking_strategy</code>: optional, see LockFactory configuration <code>retry_marker_lookup</code>: optional, default to 0. Defines how many times we look for the marker files in the source directory before failing. Waiting 5 seconds between each try. <code>retry_initialize_period</code>: optional, set an integer value in seconds to enable the retry initialize feature: if the slave can't find the master index it will try again until it's found in background, without preventing the application to start: full-text queries performed before the index is initialized are not blocked but will return empty results. When not enabling the option or explicitly setting it to zero it will fail with an exception instead of scheduling a retry timer. To prevent the application from starting without an invalid index but still control an initialization timeout, see <code>retry_marker_lookup</code> instead. <code>filesystem_access_type</code>: allows to determine the exact type of FSDirectory implementation used by this DirectoryProvider. Allowed values are <code>auto</code> (the default value, selects NIOFSDirectory on non Windows systems, SimpleFSDirectory on Windows), <code>simple</code> (SimpleFSDirectory), <code>nio</code> (NIOFSDirectory), <code>mmap</code> (MMapDirectory). Make sure to refer to Javadocs of these Directory implementations before changing this setting. Even though NIOFSDirectory or MMapDirectory can bring substantial performance boosts they also have their issues.

Name and description	Properties
infinispan: Infinispan based directory. Use it to store the index in a distributed grid, making index changes visible to all elements of the cluster very quickly. Also see Infinispan Directory configuration for additional requirements and configuration settings. Infinispan needs a global configuration and additional dependencies; the settings defined here apply to each different index.	locking_cachename: name of the Infinispan cache to use to store locks. data_cachename: name of the Infinispan cache to use to store the largest data chunks; this area will contain the largest objects, use replication if you have enough memory or switch to distribution. metadata_cachename: name of the Infinispan cache to use to store the metadata relating to the index; this data is rather small and read very often, it's recommended to have this cache setup using replication. chunk_size: large files of the index are split in smaller chunks, you might want to set the highest value efficiently handled by your network. Networking tuning might be useful.



If the built-in directory providers do not fit your needs, you can write your own directory provider by implementing the `org.hibernate.store.DirectoryProvider` interface. In this case, pass the fully qualified class name of your provider into the `directory_provider` property. You can pass any additional properties to configure your custom implementation using the prefix `hibernate.search.<indexname>`.

3.3.1. Infinispan Directory configuration

Infinispan is a distributed, scalable, cloud friendly data grid platform, which Hibernate Search can use to store the Lucene index. Your application can benefit from Infinispan's distribution capabilities making index updates available on all nodes with short latency.

This section describes how to configure Hibernate Search to use an Infinispan Lucene Directory.

When using an Infinispan Directory the index is stored in memory and shared across multiple nodes. It is considered a single directory distributed across all participating nodes: if a node updates the index, all other nodes are updated as well. Updates on one node can be immediately searched for in the whole cluster.

The default configuration replicates all data which defines the index across all nodes, thus consuming a significant amount of memory but providing the best query performance. For large indexes it's suggested to enable data distribution, so that each piece of information is replicated to a subset of all cluster members. The distribution option will reduce the amount of memory required for each node but is less efficient as it will cause high network usage among the nodes.

It is also possible to offload part or most information to a `CacheStore`, such as plain filesystem, Amazon S3, Cassandra, MongoDB or standard relational databases. You can configure it to have a `CacheStore` on each node or have a single centralized one shared by each node.

A popular choice is to use a replicated index aiming to keep the whole index in memory, combined with

a **CacheStore** as safety valve in case the index gets larger than expected.

See the [Infinispan documentation](#) for all Infinispan configuration options.

Requirements

To use the Infinispan directory via Maven, add the following dependencies:

Example 13. Maven dependencies for Hibernate Search using Infinispan

```
<dependency>
  <groupId>org.hibernate</groupId>
  <artifactId>hibernate-search-orm</artifactId>
  <version>5.10.13.Final</version>
</dependency>
<dependency>
  <groupId>org.infinispan</groupId>
  <artifactId>infinispan-directory-provider</artifactId>
  <version>9.2.2.Final</version>
</dependency>
```

This dependency changed in Hibernate Search version 5.2.



Previously the DirectoryProvider was provided by the Hibernate Search project and had Maven coordinates 'org.hibernate:hibernate-search-infinispan', but the Infinispan team is now maintaining this extension point so since this version please use the Maven definition as in the previous example.

The version printed above was the latest known compatible at the time of publishing this Hibernate Search version: it's possible that more recently improved versions of Infinispan have been published which are compatible with this same Hibernate Search version.

Architecture

Even when using an Infinispan directory it's still recommended to use the JMS Master/Slave or JGroups backend, because in Infinispan all nodes will share the same index and it is likely that **IndexWriter** instances being active on different nodes will try to acquire the lock on the same index. So instead of sending updates directly to the index, send it to a JMS queue or JGroups channel and have a single node apply all changes on behalf of all other nodes.

Configuring a non-default backend is not a requirement but a performance optimization as locks are enabled to have a single node writing.

To configure a JMS slave only the backend must be replaced, the directory provider must be set to **infinispan**; set the same directory provider on the master, they will connect without the need to setup the copy job across nodes. Using the JGroups backend is very similar - just combine the

backend configuration with the `infinispan` directory provider.

Infinispan Configuration

The most simple configuration only requires to enable the backend:

```
hibernate.search.[default|<indexname>].directory_provider = infinispan
```

That's all what is needed to get a cluster-replicated index, but the default configuration does not enable any form of permanent persistence for the index; to enable such a feature an Infinispan configuration file should be provided.

To use Infinispan, Hibernate Search requires a CacheManager; it can lookup and reuse an existing CacheManager, via JNDI, or start and manage a new one. In the latter case Hibernate Search will start and stop it (closing occurs when the Hibernate SessionFactory is closed).

To use an existing CacheManager via JNDI (optional parameter):

```
hibernate.search.infinispan.cachemanager_jndiname = [jndiname]
```

To start a new CacheManager from a configuration file (optional parameter):

```
hibernate.search.infinispan.configuration_resourcename = [infinispan configuration  
filename]
```

If both parameters are defined, JNDI will have priority. If none of these is defined, Hibernate Search will use the default Infinispan configuration included in `infinispan-directory-provider.jar`. This configuration should work fine in most cases but does not store the index in a persistent cache store.

As mentioned in [List of built-in DirectoryProvider](#), each index makes use of three caches, so three different caches should be configured as shown in the `default-hibernatesearch-infinispan.xml` provided in the `infinispan-directory-provider.jar`. Several indexes can share the same caches.

Infinispan relies on JGroups for its networking functionality, so unless you are using Infinispan on a single node, an Infinispan configuration file will refer to a JGroups configuration file. This coupling is not always practical and we provide a property to override the used JGroups configuration file:

```
hibernate.search.infinispan.configuration.transport_override_resourcename = jgroups-ec2.xml
```

This allows to just switch the JGroups configuration while keeping the rest of the Infinispan configuration.

The file `jgroups-ec2.xml` used in the example above is one of the several JGroups configurations included in Infinispan. It is a good starting point to run on Amazon EC2 networks. For more details and examples see [usage of pre-configured JGroups stacks](#) in the Infinispan configuration guide.

3.4. Worker configuration

It is possible to refine how Hibernate Search interacts with Lucene through the worker configuration. There exist several architectural components and possible extension points. Let's have a closer look.

First there is a Worker. An implementation of the Worker interface is responsible for receiving all entity changes, queuing them by context and applying them once a context ends. The most intuitive context, especially in connection with ORM, is the transaction. For this reason Hibernate Search will per default use the TransactionalWorker to scope all changes per transaction. One can, however, imagine a scenario where the context depends for example on the number of entity changes or some other application (lifecycle) events. For this reason the Worker implementation is configurable as shown in [Scope configuration](#).

Table 3. Scope configuration

Property	Description
hibernate.search.worker.scope	The fully qualified class name of the Worker implementation to use. If this property is not set, empty or <code>transaction</code> the default TransactionalWorker is used.
hibernate.search.default.worker.*	All configuration properties prefixed with <code>hibernate.search.default.worker</code> are passed to the Worker during initialization. This allows adding custom, worker specific parameters.
hibernate.search.worker.enlist_in_transaction	Defaults to <code>false</code> . Set it to <code>true</code> to have all indexing work sent to the queue within the same transaction as the Hibernate ORM Session. This options should only be enabled when all backends use JMS and the queues are configured to be transactional, XA enabled.

Once a context ends it is time to prepare and apply the index changes. This can be done synchronously or asynchronously from within a new thread. Synchronous updates have the advantage that the index is at all times in sync with the databases. Asynchronous updates, on the other hand, can help to minimize the user response time. The drawback is potential discrepancies between database and

index states. Lets look at the configuration options shown in [Execution configuration](#).



The following options can be different on each index; in fact they need the indexName prefix or use **default** to set the default value for all indexes.

Table 4. Execution configuration

Property	Description
hibernate.search.<indexName>.worker.execution	sync : synchronous execution (default) async : asynchronous execution

So far all work is done within the same Virtual Machine (VM), no matter which execution mode. The total amount of work has not changed for the single VM. Luckily there is a better approach, namely delegation. It is possible to send the indexing work to a different server by configuring `hibernate.search.default.worker.backend` - see [Backend configuration](#). Again this option can be configured differently for each index.

Table 5. Backend configuration

Property	Description
----------	-------------

hibernate.search.<indexName>.worker.backend	local: The default backend which runs index updates in the same JVM. Also used when the property is undefined or empty. jms: JMS backend. Index updates are send to a JMS queue to be processed by an indexing master. See JMS backend configuration for additional configuration options and JMS Master/Slave back end for a more detailed description of this setup. jgroupsMaster, jgroupsSlave or jgroups: Backend using JGroups as communication layer. See JGroups Master/Slave back end for a more detailed description of this setup. blackhole: Mainly a test/developer setting which ignores all indexing work You can also specify the fully qualified name of a class implementing BackendQueueProcessor. This way you can implement your own communication layer: the queue processors will receive all index works to be sent to index managers. Please note that instances of BackendQueueProcessor will only be created as soon as index managers are created, and those may be created lazily (for instance when the indexes are using dynamic sharding). If you want to initialize some resources eagerly, you can instead specify the fully qualified name of a class implementing Backend. This class which will be instantiated and initialized on startup, and will then be asked to produce BackendQueueProcessor instances as needed.
---	---

Table 6. JMS backend configuration

Property	Description
hibernate.search.<indexName>.worker.jms.connection_factory	Mandatory for the JMS back end. Defines the JNDI name to lookup the JMS connection factory from (/ConnectionFactory by default in WildFly)
hibernate.search.<indexName>.worker.jms.queue	Mandatory for the JMS back end. Defines the JNDI name to lookup the JMS queue from. The queue will be used to post work messages.
hibernate.search.<indexName>.worker.jms.login	Optional for the JMS slaves. Use it when your queue requires login credentials to define your login.

hibernate.search.<indexName>.worker.jms.password	Optional for the JMS slaves. Use it when your queue requires login credentials to define your password.
--	---

Since these components use JNDI, don't forget to configure the Hibernate ORM properties for the initial context lookup.

Table 7. Enable JNDI lookup capabilities of Hibernate ORM

Property	Description
hibernate.jndi.class	Name of the javax.naming.InitialContext implementation class to use
hibernate.jndi.url	Name of the JNDI InitialContext connection url

See also the [JNDI configuration in Hibernate ORM](#).



As you probably noticed, some of the shown properties are correlated which means that not all combinations of property values make sense. In fact you can end up with a non-functional configuration. This is especially true for the case that you provide your own implementations of some of the shown interfaces. Make sure to study the existing code before you write your own Worker or BackendQueueProcessor implementation.

3.4.1. JMS Master/Slave back end

This section describes in greater detail how to configure the Master/Slave Hibernate Search architecture.

JMS back end configuration.

Slave nodes

Every index update operation is sent to a JMS queue. Index querying operations are executed on a local index copy.

Example 14. JMS Slave configuration

```
### slave configuration

## DirectoryProvider
# (remote) master location
hibernate.search.default.sourceBase = /mnt/mastervolume/lucenedirs/mastercopy

# local copy location
hibernate.search.default.indexBase = /Users/prod/lucenedirs

# refresh every half hour
hibernate.search.default.refresh = 1800

# appropriate directory provider
hibernate.search.default.directory_provider = filesystem-slave

## Backend configuration
hibernate.search.default.worker.backend = jms
hibernate.search.default.worker.jms.connection_factory = /ConnectionFactory
hibernate.search.default.worker.jms.queue = queue/hibernatesearch
#optionally authentication credentials:
hibernate.search.default.worker.jms.login = myname
hibernate.search.default.worker.jms.password = wonttellyou
#optional jndi configuration (check your JMS provider for more information)

## Enqueue indexing tasks within an XA transaction with the database (optional)
hibernate.search.worker.enlist_in_transaction = true
```

The `enlist_in_transaction` option can be enabled if you need strict guarantees of indexing work to be stored in the queue within the same transaction of the database changes, however this will require both the RDBMs datasource and the JMS queue to be XA enabled.

Make sure to use a XA JMS queue and that your database supports XA as we are talking about coordinated transactional systems.

The default for `enlist_in_transaction` is `false` as often it is desirable to not have the database transaction fail in case there are issues with indexing.

It is possible to apply compensating operations to the index by implementing a custom `ErrorHandler` (see [Exception handling](#)), or simply re-synchronize the whole index state by starting the `MassIndexer` (see [Using a MassIndexer](#)).



A file system local copy is recommended for faster search results.

Master node

Every index update operation is taken from a JMS queue and executed. The master index is copied on a regular basis.

Example 15. JMS Master configuration

```
### master configuration

## DirectoryProvider
# (remote) master location where information is copied to
hibernate.search.default.sourceBase = /mnt/mastervolume/lucenedirs/mastercopy

# local master location
hibernate.search.default.indexBase = /Users/prod/lucenedirs

# refresh every half hour
hibernate.search.default.refresh = 1800

# appropriate directory provider
hibernate.search.default.directory_provider = filesystem-master

## Backend configuration
#The backend is not set: use the default one which is 'local'
```



It is recommended that the refresh period be higher than the expected copy time; if a copy operation is still being performed when the next refresh triggers, the second refresh is skipped: it's safe to set this value low even when the copy time is not known.

In addition to the Hibernate Search framework configuration, a Message Driven Bean has to be written and set up to process the index works queue through JMS.

Example 16. Message Driven Bean processing the indexing queue

```
import javax.ejb.ActivationConfigProperty;
import javax.ejb.MessageDriven;
import javax.jms.MessageListener;
import javax.persistence.EntityManagerFactory;
import javax.persistence.PersistenceUnit;

import org.hibernate.search.backend.jms.spi.AbstractJMSHibernateSearchController;
import org.hibernate.search.orm.spi.SearchIntegratorHelper;
import org.hibernate.search.spi.SearchIntegrator;

@MessageDriven(activationConfig = {
    @ActivationConfigProperty(propertyName="destinationType",
                             propertyValue="javax.jms.Queue"),
    @ActivationConfigProperty(propertyName="destination",
                             propertyValue="queue/hibernatesearch")
} )
public class MDBSearchController extends AbstractJMSHibernateSearchController
    implements MessageListener {

    @PersistenceUnit EntityManagerFactory emf;

    @Override
    protected SearchIntegrator getSearchIntegrator() {
        return SearchIntegratorHelper.extractFromEntityManagerFactory( emf );
    }
}
```

This example inherits from the abstract JMS controller class available in the Hibernate Search source code and implements a JavaEE MDB. This implementation is given as an example and can be adjusted to make use of non Java EE Message Driven Beans. Essentially what you need to do is to connect the specific JMS Queue with the `SearchFactory` instance of the `EntityManager`. As an advanced alternative, you can implement your own logic by not extending `AbstractJMSHibernateSearchController` but rather to use it as an implementation example.

3.4.2. JGroups Master/Slave back end

This section describes how to configure the JGroups Master/Slave back end. The master and slave roles are similar to what is illustrated in [JMS Master/Slave back end](#), only a different backend (`hibernate.search.default.worker.backend`) needs to be set.

A specific backend can be configured to act either as a slave using `jgroupsSlave`, as a master using `jgroupsMaster`, or can automatically switch between the roles as needed by using `jgroups`.



Either you specify a single `jgroupsMaster` and a set of `jgroupsSlave` instances, or you specify all instances as `jgroups`. Never mix the two approaches!

All backends configured to use JGroups share the same channel. The JGroups JChannel is the main communication link across all nodes participating in the same cluster group; since it is convenient to have just one channel shared across all backends, the Channel configuration properties are not defined on a per-worker section but are defined globally. See [JGroups channel configuration](#).

Table [JGroups backend configuration properties](#) contains all configuration options which can be set independently on each index backend. These apply to all three variants of the backend: `jgroupsSlave`, `jgroupsMaster`, `jgroups`. It is very unlikely that you need to change any of these from their defaults.

Table 8. JGroups backend configuration properties

Property	Description
<code>hibernate.search.<indexName>.jgroups.block_waiting_ack</code>	Set to either <code>true</code> or <code>false</code> . False is more efficient but will not wait for the operation to be delivered to the peers. Defaults to <code>true</code> when the backend is synchronous, to <code>false</code> when the backend is <code>async</code> .
<code>hibernate.search.<indexName>.jgroups.messages_timeout</code>	The timeout of waiting for a single command to be acknowledged and executed when <code>block_waiting_ack</code> is <code>true</code> , or just acknowledged otherwise. Value in milliseconds, defaults to <code>20000</code> .

hibernate.search.<indexName>.jgroups.delegate_backend	The master node receiving indexing operations forwards them to a standard backend to be performed. Defaults to lucene . See also Backend configuration for other options, but probably the only useful option is blackhole , or a custom implementation, to help isolating network latency problems.
---	--

Slave nodes

Every index update operation is sent through a JGroups channel to the master node. Index querying operations are executed on a local index copy. Enabling the JGroups worker only makes sure the index operations are sent to the master, you still have to synchronize configuring an appropriate directory (See **filesystem-master**, **filesystem-slave** or **infinispan** options in [Directory configuration](#)).

Example 17. JGroups Slave configuration

```
### slave configuration
hibernate.search.default.worker.backend = jgroupsSlave
```

Master node

Every index update operation is taken from a JGroups channel and executed. The master index is copied on a regular basis.

Example 18. JGroups Master configuration

```
### master configuration
hibernate.search.default.worker.backend = jgroupsMaster
```

Automatic master election



This feature is considered experimental. In particular during a re-election process there is a small window of time in which indexing requests could be lost.

In this mode the different nodes will autonomously elect a master node. When a master fails, a new node is elected automatically.

When setting this backend it is expected that all Hibernate Search instances in the same cluster use the same backend for each specific index: this configuration is an alternative to the static **jgroupsMaster** and **jgroupsSlave** approach so make sure to not mix them.

To synchronize the indexes in this configuration avoid **filesystem-master** and **filesystem-slave** directory providers as their behaviour can not be switched dynamically; use the Infinispan

Directory instead, which has no need for different configurations on each instance and allows dynamic switching of writers; see also [Infinispan Directory configuration](#).

Example 19. JGroups configuration for automatic master configuration

```
### automatic configuration
hibernate.search.default.worker.backend = jgroups
```



Should you use **jgroups** or the couple **jgroupsMaster**, **jgroupsSlave**?

The dynamic **jgroups** backend is better suited for environments in which your master is more likely to need to failover to a different machine, as in clouds. The static configuration has the benefit of keeping the master at a well known location: your architecture might take advantage of it by sending most write requests to the known master. Also optimisation and MassIndexer operations need to be triggered on the master node.

JGroups channel configuration

Configuring the JGroups channel essentially entails specifying the transport in terms of a network protocol stack. To configure the JGroups transport, point the configuration property `hibernate.search.services.jgroups.configurationFile` to a JGroups configuration file; this can be either a file path or a Java resource name.



If no property is explicitly specified it is assumed that the JGroups default configuration file **flush-udp.xml** is used. This example configuration is known to work in most scenarios, with the notable exception of Amazon AWS; refer to the [JGroups manual](#) for more examples and protocol configuration details.

The default cluster name is **Hibernate Search Cluster** which can be configured as seen in [JGroups cluster name configuration](#).

Example 20. JGroups cluster name configuration

```
hibernate.search.services.jgroups.clusterName = My-Custom-Cluster-Id
```

The cluster name is what identifies a group: by changing the name you can run different clusters in the same network in isolation.

JGroups channel instance injection

For programmatic configurations, one additional option is available to configure the JGroups channel: to pass an existing channel instance to Hibernate Search directly using the property

`hibernate.search.services.jgroups.providedChannel`, as shown in the following example.

```
import org.hibernate.search.backend.impl.jgroups.JGroupsChannelProvider;

org.jgroups.JChannel channel = ...
Map<String,String> properties = new HashMap<String,String>(1);
properties.put( JGroupsChannelProvider.CHANNEL_INJECT, channel );
EntityManagerFactory emf = Persistence.createEntityManagerFactory( "userPU", properties );
```

3.5. Reader strategy configuration

The different reader strategies are described in [Reader strategy](#). Out of the box strategies are:

- **shared**: share index readers across several queries. This strategy is very efficient.
- **not-shared**: create an index reader for each individual query. Very simple implementation.
- **async**: only opens a new index reader periodically. This is the most efficient implementation, but queries might return out of date values.

The default reader strategy is **shared**.

You can pick the reader strategy by changing the `.reader.strategy` configuration property, scoped to the "default" index or to a specific index.

For example:

```
hibernate.search.[default|<indexname>].reader.strategy = async
hibernate.search.[default|<indexname>].reader.async_refresh_period_ms = 8000
```

Adding the above properties switches to the **async** strategy, and configures it to refresh the index reader each 8 seconds.

Alternatively you can use a custom implementation of a `org.hibernate.search.indexes.spi.ReaderProvider`:

```
hibernate.search.[default|<indexname>].reader.strategy = my.corp.myapp.CustomReaderProvider
```

where `my.corp.myapp.CustomReaderProvider` is the custom strategy implementation.

3.6. Serialization

When using clustering features, Hibernate Search needs to find an implementation of the `SerializationProvider` service on the classpath.

An implementation of the service based on [Apache Avro](#) can be found using the following GAV coordinates:

```
org.hibernate:hibernate-search-serialization-avro:5.10.13.Final
```

You can add the coordinates to your pom file or download all the required dependencies and add them to your classpath. Hibernate Search will find the service implementation without any additional configuration.

Alternatively, you can create a custom service implementation:

Example 21. Serialization strategy definition

```
package example.provider.serializer

import org.hibernate.search.indexes.serialization.spi.Deserializer;
import org.hibernate.search.indexes.serialization.spi.SerializationProvider;
import org.hibernate.search.indexes.serialization.spi.Serializer;

public class ExampleOfSerializationProvider implements SerializationProvider {

    @Override
    public Serializer getSerializer() {
        Serializer serializer = ...
        return serializer;
    }

    @Override
    public Deserializer getDeserializer() {
        Deserializer deserializer = ...
        return deserializer;
    }
}
```

Hibernate Search uses the Java ServiceLoader mechanism to transparently discover services. In this case you will add the following file in your classpath:

Example 22. Service file for the `SerializationProvider` service

```
/META-INF/services/org.hibernate.search.indexes.serialization.spi.SerializationProvider
```

Example 23. Content of /META-

INF/services/org.hibernate.search.indexes.serialization.spi.SerializationProvider

```
example.provider.serializer.ExampleOfSerializationProvider
```

You will find more details about services in the section [Using external services](#).

3.7. Exception handling

Hibernate Search allows you to configure how exceptions are handled during the indexing process. If no configuration is provided then exceptions are logged to the log output by default. It is possible to explicitly declare the exception logging mechanism as seen below:

```
hibernate.search.error_handler = log
```

The default exception handling occurs for both synchronous and asynchronous indexing. Hibernate Search provides an easy mechanism to override the default error handling implementation.

In order to provide your own implementation you must implement the ErrorHandler interface, which provides the handle(ErrorContext context) method. ErrorContext provides a reference to the primary LuceneWork instance, the underlying exception and any subsequent LuceneWork instances that could not be processed due to the primary exception.

```
public interface ErrorContext {
    List<LuceneWork> getFailingOperations();
    LuceneWork getOperationAtFault();
    Throwable getThrowable();
    boolean hasErrors();
}
```

To register this error handler with Hibernate Search you must declare the fully qualified classname of your ErrorHandler implementation in the configuration properties:

```
hibernate.search.error_handler = CustomerErrorHandler
```

Alternatively, an ErrorHandler instance may be passed via the configuration value map used when bootstrapping Hibernate Search programmatically.

3.8. Lucene configuration

Even though Hibernate Search will try to shield you as much as possible from Lucene specifics, there are several Lucene specifics which can be directly configured, either for performance reasons or for satisfying a specific use case. The following sections discuss these configuration options.

3.8.1. Tuning indexing performance

Hibernate Search allows you to tune the Lucene indexing performance by specifying a set of parameters which are passed through to underlying Lucene `IndexWriter` such as `mergeFactor`, `maxMergeDocs` and `maxBufferedDocs`. You can specify these parameters either as default values applying for all indexes, on a per index basis, or even per shard.

There are several low level `IndexWriter` settings which can be tuned for different use cases. These parameters are grouped by the `indexwriter` keyword:

```
hibernate.search.[default|<indexname>].indexwriter.<parameter_name>
```

If no value is set for an `indexwriter` value in a specific shard configuration, Hibernate Search will look at the index section, then at the default section.

Example 24. Example performance option configuration

```
hibernate.search.Animals.2.indexwriter.max_merge_docs = 10
hibernate.search.Animals.2.indexwriter.merge_factor = 20
hibernate.search.Animals.2.indexwriter.max_buffered_docs = default
hibernate.search.default.indexwriter.max_merge_docs = 100
hibernate.search.default.indexwriter.ram_buffer_size = 64
```

The configuration in [Example performance option configuration](#) will result in these settings applied on the second shard of the Animal index:

- `max_merge_docs` = 10
- `merge_factor` = 20
- `ram_buffer_size` = 64MB
- `max_buffered_docs` = Lucene default

All other values will use the defaults defined in Lucene.

The default for all values is to leave them at Lucene's own default. The values listed in [List of indexing performance and behavior properties](#) depend for this reason on the version of Lucene you are using. The values shown are relative to version `2.4`. For more information about Lucene indexing

performance, please refer to the Lucene documentation.

Table 9. List of indexing performance and behavior properties

Property	Description	Default Value
hibernate.search.[default <indexname>].exclusive_index_use	Set to true when no other process will need to write to the same index. This will enable Hibernate Search to work in exclusive mode on the index and improve performance when writing changes to the index. [WARNING] ==== Support for non-exclusive index use (setting this property to false) is deprecated, and will be removed in Hibernate Search 6. ====	true (improved performance, releases locks only at shutdown)
hibernate.search.[default <indexname>].max_queue_length	Each index has a separate "pipeline" which contains the updates to be applied to the index. When this queue is full adding more operations to the queue becomes a blocking operation. Configuring this setting doesn't make much sense unless the worker.execution is configured as async .	1000
hibernate.search.[default <indexname>].index_flush_interval	The interval in milliseconds between flushes of write operations to the index storage. Ignored unless worker.execution is configured as async .	1000
hibernate.search.[default <indexname>].indexwriter.max_buffered_delete_terms	Determines the minimal number of delete terms required before the buffered in-memory delete terms are applied and flushed. If there are documents buffered in memory at the time, they are merged and a new segment is created.	Disabled (flushes by RAM usage)
hibernate.search.[default <indexname>].indexwriter.max_buffered_docs	Controls the amount of documents buffered in memory during indexing. The bigger the more RAM is consumed.	Disabled (flushes by RAM usage)
hibernate.search.[default <indexname>].indexwriter.max_merge_docs	Defines the largest number of documents allowed in a segment. Smaller values perform better on frequently changing indexes, larger values provide better search performance if the index does not change often.	Unlimited (Integer.MAX_VALUE)

Property	Description	Default Value
hibernate.search.[default <indexname>].indexwriter.merge_factor	Controls segment merge frequency and size. Determines how often segment indexes are merged when insertion occurs. With smaller values, less RAM is used while indexing, and searches on unoptimized indexes are faster, but indexing speed is slower. With larger values, more RAM is used during indexing, and while searches on unoptimized indexes are slower, indexing is faster. Thus larger values (> 10) are best for batch index creation, and smaller values (< 10) for indexes that are interactively maintained. The value must not be lower than 2.	10
hibernate.search.[default <indexname>].indexwriter.merge_min_size	Controls segment merge frequency and size. Segments smaller than this size (in MB) are always considered for the next segment merge operation. Setting this too large might result in expensive merge operations, even though they are less frequent. See also org.apache.lucene.index.LogDocMergePolicy.minMergeSize .	0 MB (actually ~1K)
hibernate.search.[default <indexname>].indexwriter.merge_max_size	Controls segment merge frequency and size. Segments larger than this size (in MB) are never merged in bigger segments. This helps reduce memory requirements and avoids some merging operations at the cost of optimal search speed. When optimizing an index this value is ignored. See also org.apache.lucene.index.LogDocMergePolicy.maxMergeSize .	Unlimited
hibernate.search.[default <indexname>].indexwriter.merge_max_optimize_size	Controls segment merge frequency and size. Segments larger than this size (in MB) are not merged in bigger segments even when optimizing the index (see merge_max_size setting as well). Applied to org.apache.lucene.index.LogDocMergePolicy.maxMergeSizeForOptimize .	Unlimited
hibernate.search.[default <indexname>].indexwriter.merge_calibrate_by_deletes	Controls segment merge frequency and size. Set to false to not consider deleted documents when estimating the merge policy. Applied to org.apache.lucene.index.LogMergePolicy.calibrateSizeByDeletes .	true

Property	Description	Default Value
hibernate.search.[default<indexname>].indexwriter.ram_buffer_size	Controls the amount of RAM in MB dedicated to document buffers. When used together max_buffered_docs a flush occurs for whichever event happens first. Generally for faster indexing performance it's best to flush by RAM usage instead of document count and use as large a RAM buffer as you can.	16 MB
hibernate.search.enable_dirty_check	Not all entity changes require an update of the Lucene index. If all of the updated entity properties (dirty properties) are not indexed Hibernate Search will skip the re-indexing work. Disable this option if you use a custom <code>FieldBridge</code> which need to be invoked at each update event (even though the property for which the field bridge is configured has not changed). This optimization will not be applied on classes using a <code>@ClassBridge</code> or a <code>@DynamicBoost</code> . Boolean parameter, use "true" or "false".	true
hibernate.search.[default<indexname>].indexwriter.infostream	Enable low level trace information about Lucene's internal components. Will cause significant performance degradation: should only be used for troubleshooting purposes.	false



When your architecture permits it, always keep `hibernate.search.default.exclusive_index_use=true` as it greatly improves efficiency in index writing. This is the default since Hibernate Search version 4.



To tune the indexing speed it might be useful to time the object loading from database in isolation from the writes to the index. To achieve this set the **blackhole** as worker backend and start your indexing routines. This backend does not disable Hibernate Search: it will still generate the needed changesets to the index, but will discard them instead of flushing them to the index. In contrast to setting the **hibernate.search.indexing_strategy** to **manual**, using **blackhole** will possibly load more data from the database because associated entities are re-indexed as well.

```
hibernate.search.[default|<indexname>].worker.backend blackhole
```

The recommended approach is to focus first on optimizing the object loading by enabling the **blackhole** backend, and then use the timings you achieve as a baseline to tune the indexing process.



The **blackhole** backend is not meant to be used in production, only as a tool to identify indexing bottlenecks.

Control segment size

The options **merge_max_size**, **merge_max_optimize_size**, **merge_calibrate_by_deletes** give you control on the maximum size of the segments being created, but you need to understand how they affect file sizes. If you need to hard limit the size, consider that merging a segment is about adding it together with another existing segment to form a larger one, so you might want to set the **max_size** for merge operations to less than half of your hard limit. Also segments might initially be generated larger than your expected size at first creation time: before they are ever merged. A segment is never created much larger than **ram_buffer_size**, but the threshold is checked as an estimate.

Example:

```
//to be fairly confident no files grow above 15MB, use:  
hibernate.search.default.indexwriter.ram_buffer_size = 10  
hibernate.search.default.indexwriter.merge_max_optimize_size = 7  
hibernate.search.default.indexwriter.merge_max_size = 7
```



When using the Infinispan Directory to cluster indexes make sure that your segments are smaller than the **chunk_size** so that you avoid fragmenting segments in the grid. Note that the **chunk_size** of the Infinispan Directory is expressed in bytes, while the index tuning options are in MB.

Troubleshooting: enable Lucene's Infostream

Apache Lucene allows to log a very detailed trace log from its internals using a feature called "infostream". To access these details, Hibernate Search can be configured to capture this internal trace from Apache Lucene and redirect it to your logger.

- Enable `TRACE` level logging for the category `org.hibernate.search.backend.lucene.infostream`
- Activate the feature on the index you want to inspect:
`hibernate.search.[default|<indexname>].indexwriter.infostream=true`

Keep in mind that this feature has a performance cost, and although most logger frameworks allow the `TRACE` level to be reconfigured at runtime, enabling the `infostream` property will slow you down even if the logger is disabled.

3.8.2. LockFactory configuration

Lucene Directorys have default locking strategies which work generally good enough for most cases, but it's possible to specify for each index managed by Hibernate Search a specific LockingFactory you want to use. This is generally not needed but could be useful.

Some of these locking strategies require a filesystem-level lock. They may be used with the `local-heap` directory provider, but in this case the `indexBase` configuration option (usually not needed when using a `local-heap` directory provider) must be specified to point to a filesystem location where the lock marker files will be stored.

To select a locking factory, set the `hibernate.search.<index>.locking_strategy` option to one of `simple`, `native`, `single` or `none`. Alternatively set it to the fully qualified name of an implementation of `org.hibernate.search.store.LockFactoryProvider`.

Table 10. List of available LockFactory implementations

name	Class	Description
simple	org.apache.lucene.store.SimpleFSLockFactory	Safe implementation based on Java's File API, it marks the usage of the index by creating a marker file. If for some reason you had to kill your application, you will need to remove this file before restarting it.

name	Class	Description
native	org.apache.lucene.store.NativeFSLockFactory	As does simple this also marks the usage of the index by creating a marker file, but this one is using native OS file locks so that even if the JVM is terminated the locks will be cleaned up. This implementation has known problems on NFS, avoid it on network shares. native is the default implementation for the filesystem, filesystem-master and filesystem-slave directory providers.
single	org.apache.lucene.store.SingleInstanceLockFactory	This LockFactory doesn't use a file marker but is a Java object lock held in memory; therefore it's possible to use it only when you are sure the index is not going to be shared by any other process. This is the default implementation for the local-heap directory provider.
none	org.apache.lucene.store.NoLockFactory	All changes to this index are not coordinated by any lock; test your application carefully and make sure you know what it means.

Configuration example:

```
hibernate.search.default.locking_strategy = simple
hibernate.search.Animals.locking_strategy = native
hibernate.search.Books.locking_strategy = org.custom.components.MyLockingFactory
```

The Infinispan Directory uses a custom implementation; it's still possible to override it but make sure you understand how that will work, especially with clustered indexes.

3.8.3. Index format compatibility

While Hibernate Search strives to offer a backwards compatible API making it easy to port your application to newer versions, it still delegates to Apache Lucene to handle the index writing and searching. This creates a dependency to the Lucene index format. The Lucene developers of course

attempt to keep a stable index format, but sometimes a change in the format can not be avoided. In those cases you either have to re-index all your data or use an index upgrade tool. Sometimes Lucene is also able to read the old format so you don't need to take specific actions (besides making backup of your index).

While an index format incompatibility is a rare event, it can happen more often that Lucene's Analyzer implementations might slightly change its behavior. This can lead to a poor recall score, possibly missing many hits from the results.

Hibernate Search exposes a configuration property `hibernate.search.lucene_version` which instructs the analyzers and other Lucene classes to conform to their behavior as defined in an (older) specific version of Lucene. See also `org.apache.lucene.util.Version` contained in the `lucene-core.jar`. Depending on the specific version of Lucene you're using you might have different options available. When this option is not specified, Hibernate Search will instruct Lucene to use the default version, which is usually the best option for new projects. Still it's recommended to define the version you're using explicitly in the configuration so that when you happen to upgrade Lucene the analyzers will not change behavior. You can then choose to update this value at a later time, when you for example have the chance to rebuild the index from scratch.

Example 25. Force Analyzers to be compatible with a Lucene 4.7 created index

```
hibernate.search.lucene_version = LUCENE_47
```

This option is global for the configured SearchFactory and affects all Lucene APIs having such a parameter, as this should be applied consistently. So if you are also making use of Lucene bypassing Hibernate Search, make sure to apply the same value too.

3.9. Metadata API

After looking at all these different configuration options, it is time to have a look at an API which allows you to programmatically access parts of the configuration. Via the metadata API you can determine the indexed types and also verify how they are mapped (see [Mapping entities to the index structure](#)) to the index structure.

The entry point into this API is the SearchFactory. It offers two methods, namely `getIndexedTypes()` and `getIndexedTypeDescriptor(Class<?>)`. The former returns a set of all indexed type, where as the latter allows to retrieve a so called IndexedTypeDescriptor for a given type. This descriptor allows you determine whether the type is indexed at all and, if so, whether the index is for example sharded or not (see [Sharding indexes](#)). It also allows you to determine the static boost of the type (see [Static index time boosting](#)) as well as its dynamic boost strategy (see [Dynamic index time boosting](#)). Most importantly, however, you get information about the indexed properties and generated Lucene Document fields. This is exposed via PropertyDescriptors respectively FieldDescriptors. The easiest way to get to know the API is to explore it via the IDE or its javadocs.



All descriptor instances of the metadata API are read only. They do not allow to change any runtime configuration.

3.10. Hibernate Search as a WildFly module

Hibernate Search is included in the WildFly application server, and since WildFly 10 the module is automatically activated (added to the classpath of your deployment) if you are using Hibernate ORM and have any indexed entities.

WildFly 17 includes an older version of Hibernate Search; if you want to use that then you won't need to download anything else and you can stop reading this chapter; however you might miss out on several new features and improvements. When deciding to use the older version make sure to read the version of this guide matching the version of Hibernate Search included in your application server: this guide applies to **5.10.13.Final** exclusively.

We strongly suggest using the latest version. Doing so is not much harder: provision a WildFly server using the dedicated tool, and set the `wildfly.jpa.hibernate.search.module` property in your `persistence.xml` to the chosen version. Further details are provided below.

The modules distributed by Hibernate Search **5.10.13.Final** are meant for WildFly 17.0.1.Final.



WildFly often will include older versions of Hibernate Search dependencies, so in order to update Hibernate Search, you will need to upgrade the dependencies as well, otherwise your application will not start. Fortunately, there are tools to automate this for you: see [Update and select the latest Hibernate Search version in WildFly](#) for more information.

3.10.1. Use the Hibernate Search version included in WildFly

The activation of the Hibernate Search modules in WildFly is automatic, provided you're having at least one entity annotated with `org.hibernate.search.annotations.Indexed`.

You can control this behaviour of the JPA subsystem explicitly; for example to make sure Hibernate Search and Apache Lucene classes are available to your application even though you haven't annotated any entity, set the following property in your `persistence.xml`:

```
wildfly.jpa.hibernate.search.module=org.hibernate.search.orm:main
```

3.10.2. Update and select the latest Hibernate Search version in WildFly

To use an updated version of Hibernate Search within WildFly, you will need to follow two steps:

1. Provision a WildFly server with an updated version of Hibernate Search, either [via Maven](#) or [via another tool](#).
2. [Select](#) the updated Hibernate Search version

Server provisioning via Maven

Maven users can use the `wildfly-server-provisioning-maven-plugin` to create a custom WildFly server including the Hibernate Search modules:

```
<plugins>
  <plugin>
    <groupId>org.wildfly.build</groupId>
    <artifactId>wildfly-server-provisioning-maven-plugin</artifactId>
    <version>1.2.6.Final</version>
    <executions>
      <execution>
        <id>server-provisioning</id>
        <goals>
          <goal>build</goal>
        </goals>
        <phase>compile</phase>
        <configuration>
          <config-file>server-provisioning.xml</config-file>
          <server-name>wildfly-with-updated-hibernate-search</server-name>
        </configuration>
      </execution>
    </executions>
  </plugin>
```

You will also need a `server-provisioning.xml` in the root of your project:

```
<server-provisioning xmlns="urn:wildfly:server-provisioning:1.1">
  <feature-packs>

    <feature-pack
      groupId="org.hibernate"
      artifactId="hibernate-search-jbossmodules-orm"
      version="5.10.13.Final"/>

    <feature-pack
      groupId="org.wildfly"
      artifactId="wildfly-feature-pack"
      version="17.0.1.Final" />

  </feature-packs>
</server-provisioning>
```

This will populate the server with the Hibernate Search ORM integration, along with an updated

version of every required module: the Hibernate Search engine, Lucene, Hibernate ORM, ...

You may want to also include other, optional modules. If so, check out the [list of available feature packs](#) and add a `<feature-pack>` markup with the coordinates of the feature pack you want.



Depending on the WildFly feature pack you chose, some transitive dependencies may not be available in Maven Central. In that case, you should [set up the JBoss Nexus repository](#).

Server provisioning via other build tools

The provisioning tools have been developed mainly for Maven users, but there are options for users of other build tools too.

Gradle users can use the [org.wildfly.build.provision](#) plugin, available on the official Gradle plugin portal.

Both the Maven and the Gradle plugins are just thin wrappers invoking a standalone provisioning library, so you could also integrate this library with any other JVM based tool.

See also [WildFly provisioning build tools](#).

Selecting the updated Hibernate ORM/Hibernate Search versions in WildFly

Next you will need to make sure the JPA subsystem of WildFly provides you with the version you have chosen, instead of the default version bundled with the application server.

First, you will have to select the updated version of Hibernate ORM. If you followed the instructions above, the server will include just the right version of Hibernate ORM: 5.3.28.Final. You just need to select it by setting the following property in your `persistence.xml`:

```
jboss.as.jpa.providerModule=org.hibernate.orm:5.3
```

Then select the updated version of Hibernate Search by setting this other property:

```
wildfly.jpa.hibernate.search.module=org.hibernate.search.orm:5.10.13.Final
```

Setting these properties will also prevent WildFly to pollute your classpath with the older copies of Hibernate ORM and Hibernate Search which are included in the server.



For more information about configuring Hibernate ORM and Hibernate Search in WildFly, see:

- [the WildFly JPA Reference Guide](#);
- [the specific section about Hibernate Search](#).

3.10.3. List of the Hibernate Search WildFly/JBoss feature packs

Engine feature pack

This feature pack includes all the main modules:

- *org.hibernate.search.engine*, for projects depending on the internal indexing engine that don't require other dependencies to Hibernate.
- *org.hibernate.search.backend-jms*, in case you want to use the JMS backend described in [JMS Architecture](#).

Several more modules are included but you should not need to refer to the other ones directly.

The feature pack is published on the JBoss Nexus repository and Maven Central as [org.hibernate:hibernate-search-jbossmodules-engine:5.10.13.Final:zip](#).

JGroups feature pack

This feature pack includes the modules *org.hibernate.search.backend-jgroups*, to use the JGroups backend described in [JGroups Architecture](#).

The feature pack is published on the JBoss Nexus repository and Maven Central as [org.hibernate:hibernate-search-jbossmodules-orm:5.10.13.Final:zip](#).

ORM feature pack

This feature pack includes the following modules:

- *org.hibernate.search.orm*, for users of Hibernate Search with Hibernate ORM.
- *org.hibernate.search.jsr352*, to use the JSR-352 integration.

The feature pack is published on the JBoss Nexus repository and Maven Central as [org.hibernate:hibernate-search-jbossmodules-orm:5.10.13.Final:zip](#).

Elasticsearch feature pack

This feature pack includes the module *org.hibernate.search.elasticsearch*, to use the Elasticsearch integration.

The feature pack is published on the JBoss Nexus repository and Maven Central as [org.hibernate:hibernate-search-jbossmodules-elasticsearch:5.10.13.Final:zip](#).

Elasticsearch-AWS feature pack

This package includes an internal module enabling the [AWS authentication mechanism](#) when connecting to an Elasticsearch cluster.

The feature pack is published on the JBoss Nexus repository and Maven Central as [org.hibernate:hibernate-search-jbossmodules-elasticsearch-aws:5.10.13.Final:zip](#).

3.10.4. More about modules

More information about the modules configuration in WildFly can be found in the [Class Loading in WildFly](#) wiki.



Modular classloading is a feature of JBoss EAP 7 as well, but if you are using JBoss EAP, you're reading the wrong version of the user guide! JBoss EAP subscriptions include official support for Hibernate Search and come with a different edition of this guide specifically tailored for EAP users.

See the [Chapter 13 of the Red Hat JBoss Enterprise Application Platform 7 Development Guide](#).

3.10.5. Using Apache Tika and other extensions with modules for WildFly

Apache Tika requires additional dependencies depending on the media formats you need. Since it is not practical to package all possible extensions you might possibly need into Hibernate Search modules, we leave it up to our users to define an "org.apache.tika" module where you can add any extensions you might need.

This version of Hibernate Search expects Apache Tika version **1.4** so the module should be defined as:

Example 26. Example definition of a custom Tika module

```
<module xmlns="urn:jboss:module:1.1" name="org.apache.tika" slot="1.4">
```

Include Apache Tika 1.4 and any extensions and dependencies from Tika that you might need in the module.

Next, you need to have your application depend on this same module as well:

Example 27. Manifest header to add in your application

```
Dependencies: org.apache.tika:1.4
```

3.10.6. Using Infinispan with Hibernate Search on WildFly

If you are updating the version of Hibernate Search in WildFly as described in the previous paragraph, you might need to update Infinispan as well. The process is very similar: download the modules from [Infinispan project downloads](#), picking a compatible version, and decompress the modules into the `modules` directory of your WildFly installation.

Chapter 4. Mapping entities to the index structure

4.1. Mapping an entity

In [Getting started](#) you have already seen that all the metadata information needed to index entities is described through annotations. There is no need for XML mapping files. You can still use Hibernate mapping files for the basic Hibernate configuration, but the Hibernate Search specific configuration has to be expressed via annotations.



There is no XML configuration available for Hibernate Search but we provide a programmatic mapping API that elegantly replaces this kind of deployment form (see [Programmatic API](#) for more information).

If you want to contribute the XML mapping implementation, see [HSEARCH-210](#).

4.1.1. Basic mapping

Lets start with the most commonly used annotations when mapping an entity.

@Indexed

Foremost you must declare a persistent class as indexable by annotating the class with `@Indexed`. All entities not annotated with `@Indexed` will be ignored by the indexing process.

Example 28. Making a class indexable with `@Indexed`

```
@Entity
@Indexed
public class Essay {
    ...
}
```

You can optionally specify the `Indexed.index` attribute to change the default name of the index. For more information regarding index naming see [Directory configuration](#).

You can also specify an optional indexing interceptor. For more information see [conditional indexing](#).

@Field

For each property of your entity, you have the ability to describe whether and how it will be indexed. Adding the `@Field` annotation declares a property as indexed and allows you to configure various aspects of the indexing process. Without `@Field` the property is ignored by the indexing process.



Hibernate Search tries to determine the best way to convert your property to an indexable value automatically. See [Built-in bridges](#) for more information about how property values are converted by default, and see the `bridge` attribute of the `@Field` annotation below for more information about how to customize conversion.

The following attributes of the `@Field` annotation help you control the indexing outcome:

- **name**: describes under which name the property should be stored in the Lucene Document. The default value is the property name following the JavaBeans convention (in particular it is de-capitalized).
- **store**: describes whether or not the property is stored in the Lucene index. You can store the value `Store.YES` (consuming more space in the index but allowing [projection](#)), store it in a compressed way `Store.COMPRESS` (this does consume more CPU), or avoid any storage `Store.NO` (this is the default value). When a property is stored, you can retrieve its original value from the Lucene Document. Storing the property has no impact on whether the value is searchable or not.
- **index**: describes whether the property is indexed or not. The different values are `Index.NO` (no indexing, meaning the value cannot be found by a query), `Index.YES` (the element gets indexed and is searchable). The default value is `Index.YES`. `Index.NO` can be useful for cases where a property is not required to be searchable, but needed for projection.



`Index.NO` in combination with `Analyze.YES` or `Norms.YES` is not useful, since analyze and norms require the property to be indexed

- **analyze**: determines whether the property is analyzed (`Analyze.YES`) or not (`Analyze.NO`). The default value is `Analyze.YES`.



Whether or not you want to analyze a property depends on whether you wish to search the element as is, or by the words it contains. It make sense to analyze a text field, but probably not a date field.



Fields used for faceting *must not* be analyzed.

Fields used for sorting *can* be analyzed, but *must not* be tokenized, so you should rather use [normalizers](#) on those fields.

- **norms**: describes whether index time scoring information should be stored (`Norms.YES`) or not (`Norms.NO`). Not storing the norms can save a considerable amount of memory, but index time boosting (in particular) will not be available in this case. The default value is `Norms.YES`.
- **termVector**: describes collections of term-frequency pairs. This attribute enables the storing of the term vectors within the documents during indexing. The default value is `TermVector.NO`.

The different values of this attribute are:

Value	Definition
<code>TermVector.YES</code>	Store the term vectors of each document. This produces two synchronized arrays, one contains document terms and the other contains the term's frequency.
<code>TermVector.NO</code>	Do not store term vectors.
<code>TermVector.WITH_OFFSETS</code>	Store the term vector and token offset information. This is the same as <code>TermVector.YES</code> plus it contains the starting and ending offset position information for the terms.
<code>TermVector.WITH_POSITIONS</code>	Store the term vector and token position information. This is the same as <code>TermVector.YES</code> plus it contains the ordinal positions of each occurrence of a term in a document.
<code>TermVector.WITH_POSITION_OFFSETS</code>	Store the term vector, token position and offset information. This is a combination of the <code>YES</code> , <code>WITH_OFFSETS</code> and <code>WITH_POSITIONS</code> .

- `indexNullAs`: Per default null values are ignored and not indexed. However, using `indexNullAs` you can specify a string which will be inserted as token for the null value. Per default this value is set to `org.hibernate.search.annotations.Field.DO_NOT_INDEX_NULL` indicating that null values should not be indexed. You can set this value to `DEFAULT_NULL_TOKEN` to indicate that a default null token should be used. This default null token can be specified in the configuration using `hibernate.search.default_null_token`. If this property is not set the string `_null_` will be used as default. When the field is of a Numeric Type (see [@NumericField](#)), the token will be encoded as the respective numeric type: the `indexNullAs` value needs to be set to a value which can be parsed into a number of the matching type, for example "-1".



When `indexNullAs` is used, it is important to use the chosen null token in search queries (see [Querying](#)) in order to find null values. It is also advisable to use this feature only with un-analyzed fields (`analyze=Analyze.NO`).



When implementing a custom `FieldBridge` or `TwoWayFieldBridge` it is up to the developer to handle the indexing of null values (see JavaDocs of `LuceneOptions.indexNullAs()`).

- `boost`: Refer to section about [boosting](#)
- `bridge`: Refer to section about [field bridges](#)

@NumericField

`@NumericField` is a companion annotation to `@Field`. It can be specified in the same scope as `@Field`, but only on properties of numeric type like byte, short, int, long, double and float (and their respective Java wrapper types). It allows to define a custom `precisionStep` for the numeric encoding of the property value.

`@NumericField` accepts the following parameters:

Value	Definition
<code>forField</code>	(Optional) Specify the name of of the related <code>@Field</code> that will be indexed numerically. It's only mandatory when the property contains more than a <code>@Field</code> declaration
<code>precisionStep</code>	(Optional) Change the way that the Trie structure is stored in the index. Smaller <code>precisionSteps</code> lead to more disk space usage and faster range and sort queries. Larger values lead to less space used and range query performance more close to the range query using string encoding. Default value is 4.

Lucene supports the numeric types: `Double`, `Long`, `Integer` and `Float`. For properties of types `Byte` and `Short`, an `Integer` field will be used in the index. Other numeric types should use the default string encoding (via `@Field`), unless the application can deal with a potential loss in precision, in which case a custom `TwoWayFieldBridge` can be used. See [Defining a custom FieldBridge for BigDecimal](#).

Example 29. Defining a custom `FieldBridge` for `BigDecimal`

```
public class BigDecimalNumericFieldBridge implements TwoWayFieldBridge {

    private static final BigDecimal storeFactor = BigDecimal.valueOf( 100 );

    @Override
    public void set(String name, Object value, Document document, LuceneOptions
luceneOptions) {
        if ( value != null ) {
            BigDecimal decimalValue = (BigDecimal) value;
            Long indexedValue = decimalValue.multiply( storeFactor ).longValue();
            luceneOptions.addNumericFieldToDocument( name, indexedValue, document );
        }
    }

    @Override
    public Object get(String name, Document document) {
        String fromLucene = document.get( name );
        BigDecimal storedBigDecimal = new BigDecimal( fromLucene );
        return storedBigDecimal.divide( storeFactor );
    }

    @Override
    public String objectToString(Object object) {
        return object.toString();
    }
}
```

You would use this custom bridge like seen in [Use of `BigDecimalNumericFieldBridge`](#). In this case three annotations are used - `@Field`, `@NumericField` and `@FieldBridge`. `@Field` is required to mark the property for being indexed (a standalone `@NumericField` is never allowed).

Example 30. Use of `BigDecimalNumericFieldBridge`

```
@Entity
@Indexed
public class Item {
    @Id
    @GeneratedValue
    private int id;

    @Field
    @NumericField
    @FieldBridge(impl = BigDecimalNumericFieldBridge.class)
    private BigDecimal price;

    public int getId() {
        return id;
    }

    public BigDecimal getPrice() {
        return price;
    }

    public void setPrice(BigDecimal price) {
        this.price = price;
    }
}
```

Alternatively, you can define the type of this field directly from the bridge by implementing `MetadataProvidingFieldBridge`, as shown in [Defining a custom FieldBridge for BigDecimal and implementing MetadataProvidingFieldBridge](#). Then the `@NumericField` annotation will not be needed anymore when using the bridge.

Example 31. Defining a custom FieldBridge for BigDecimal and implementing MetadataProvidingFieldBridge

```
public class BigDecimalNumericFieldBridge implements MetadataProvidingFieldBridge,
TwoWayFieldBridge {

    private static final BigDecimal storeFactor = BigDecimal.valueOf( 100 );

    @Override
    public void set(String name, Object value, Document document, LuceneOptions
luceneOptions) {
        if ( value != null ) {
            BigDecimal decimalValue = (BigDecimal) value;
            Long indexedValue = decimalValue.multiply( storeFactor ).longValue();
            luceneOptions.addNumericFieldToDocument( name, indexedValue, document );
        }
    }

    @Override
    public Object get(String name, Document document) {
        String fromLucene = document.get( name );
        BigDecimal storedBigDecimal = new BigDecimal( fromLucene );
        return storedBigDecimal.divide( storeFactor );
    }

    @Override
    public String objectToString(Object object) {
        return object.toString();
    }

    @Override
    public void configureFieldMetadata(String name, FieldMetadataBuilder builder) {
        builder.field( name, FieldType.LONG );
    }
}
```

@SortableField

As of Lucene 5 (and thus Hibernate Search 5.5) it is highly recommended to create a so-called "doc value field" for each field to sort on. Hibernate Search provides the `@SortableField` annotation for that purpose. This is an extension annotation to `@Field` and marks a field as sortable (internally, the required doc value field will be added to the index).

Example 32. Use of `@SortableField`

```
@Entity
@Indexed
public class Book {

    @Id
    @GeneratedValue
    private int id;

    @Field(name="Abstract", analyze=Analyze.NO)
    @SortableField
    private String summary;

    // ...
}
```

If there is a single `@Field` declared for a given property, `@SortableField` implicitly applies to this field. In case several fields exist for a single property, the `@Field` to be marked as sortable can be specified via `@SortableField#forField()`. Several sortable fields can be defined as the `@SortableField` annotation is repeatable.



The field to be marked as sortable must not be tokenized, so you should either disable analysis on this field, or use a [normalizer](#).

Note that sorting also works if a property is not explicitly marked with `@SortableField`. This has negative runtime performance and memory consumption implications, though. Therefore it is highly recommended to explicitly mark each field to be used for sorting.

Should you want to make a property sortable but not searchable, still an `@Field` needs to be declared (so its field bridge configuration can be inherited). It can be marked with `store = Store.NO` and `index = Index.NO`, causing only the doc value field required for sorting to be added, but not a regular index field.

Fields added through class-level bridges or custom field-level bridges (when not using the default field name) cannot be marked as sortable by means of the `@SortableField` annotation. Instead the field bridge itself has to add the required doc value fields, in addition to the document fields it adds. Furthermore such bridge needs to implement the `MetadataProvidingFieldBridge` interface which defines a method `configureFieldMetadata()` for marking the fields created by this bridge as sortable:

Example 33. Marking fields as sortable via a custom field bridge

```
/**
 * Custom field bridge for a Map property which creates sortable fields
 * with the values of two keys from the map.
 */
public class MyClassBridge implements MetadataProvidingFieldBridge {

    @Override
    public void set(String name, Object value,
        Document document, LuceneOptions luceneOps) {

        Map<String, String> map = (Map<String, String>) value;

        String firstName = map.get( "firstName" );
        String lastName = map.get( "lastName" );

        // add regular document fields
        luceneOps.addFieldToDocument( name + "_firstName", lastName, document );
        luceneOps.addFieldToDocument( name + "_lastName", lastName, document );

        // add doc value fields to allow for sorting
        document.addSortedDocValuesFieldToDocument( name + "_firstName", firstName );
        document.addSortedDocValuesFieldToDocument( name + "_lastName", lastName );
    }

    @Override
    public void configureFieldMetadata(String name, FieldMetadataBuilder builder) {
        builder
            .field( name + "_firstName", FieldType.STRING )
                .sortable( true )
            .field( name + "_lastName", FieldType.STRING )
                .sortable( true );
    }
}
```

The meta-data configured through `configureFieldMetadata()` will be used for sort validation upon query execution. The name passed to the method is the default field name also passed to `set()`. It needs to be used consistently with `set()`, e.g. as a prefix for all custom fields added.

Flagging uncovered sorts

By default Hibernate Search will raise an exception when running a query with sorts not covered by the sortable fields configured as described above.

You thus can optionally advise Hibernate Search to transparently create an uninverting index reader when detecting uncovered sorts. Note that while this allows to execute the query, relying on index uninverting negatively impacts performance.

To do so, specify the following option:

Example 34. Enabling automatic index uninverting for uncovered sorts

```
hibernate.search.index_uninverting_allowed = true
```

For instance you may leave this to the default (`false`) during testing to identify the sortable fields required for your queries and set it to `true` in production environments to fall back to index uninverting for uncovered sorts accidentally left over.

Document identifier: `@DocumentId` or `@Id`

Finally, the document identifier is a special field used by Hibernate Search to ensure there will only ever be one document per entity in the index, and to identify documents when they should be deleted. By design, a document identifier has to be stored and must not be tokenized. It is also always string encoded, even if the id is a number.



We strongly suggest avoiding to sort on the id field as it will lead to erratic behaviors. If you want to sort on your document id, we suggest declaring another field specifically for sorting using the `@Field` annotation.

If you use JPA and if you mapped a property with `@Id`, Hibernate Search will use this entity identifier as a document identifier by default. [Using the ORM entity identifier as document identifier](#) takes advantage of the `@Id` annotation: the document identifier will be the `id` property.

Example 35. Using the ORM entity identifier as document identifier

```
@Entity
@Indexed
public class Essay {
    ...

    @Id
    public Long getId() { return id; }

    ...
}
```

If you don't use JPA, or if you need to set your document ID to some other property than the entity ID, you can use the `@DocumentId` annotation. This annotation also allows to set the document identifier field name explicitly. [Specifying the document identifier explicitly](#) takes advantage of the `@DocumentId` annotation: the document identifier will be generated from the `orderNumber` property and stored into the `id_order_number` field.

Example 36. Specifying the document identifier explicitly

```
@Entity
@Indexed
public class Order {
    ...

    @Id
    @GeneratedValue
    public Long getId() { return id; }

    @DocumentId(name = "id_order_number")
    @NaturalId
    public String orderNumber;

    ...
}
```

Supported identifier types

By default, Hibernate Search can handle document identifiers of any type mentioned in the [Built-in bridges](#) section. For those types, you will not have anything more to do than selecting a property using `@Id` or `@DocumentId`.

If your identifier has a different type, for instance because you are using a JPA `@EmbeddedId`, you will have to write a custom two-way bridge and apply it using the `@FieldBridge` annotation on your property [as explained further down](#).

4.1.2. Mapping properties multiple times

Sometimes one has to map a property multiple times per index, with slightly different indexing strategies. For example, sorting a query by field requires the field to be un-analyzed. If one wants to search by words in this property and still sort it, one needs to index it twice - once analyzed and once un-analyzed. The `@Field` is repeatable, which allows to achieve this goal. Alternatively you can use the legacy, explicit plural form `@Fields`.

Example 37. Using `@Field` repeatedly to map a property multiple times

```
@Entity
@Indexed(index = "Book")
public class Book {

    @Field
    @Field(name = "summary_forSort", analyze = Analyze.NO, store = Store.YES)
    @SortableField(forField = "summary_forSort")
    public String getSummary() {
        return summary;
    }

    // ...
}
```

In [Using @Field repeatedly to map a property multiple times](#) the field `summary` is indexed twice, once as `summary` in a tokenized way, and once as `summary_forSort` in an un-tokenized way. `@Field` supports two attributes which are particularly useful when the annotation is repeated:

- **analyzer**: defines a `@Analyzer` annotation per field rather than per property
- **normalizer**: defines a `@Normalizer` annotation for the given field, instead of an analyzer. **analyzer** and **normalizer** are mutually exclusive: you can only use one of the two.
- **bridge**: defines a `@FieldBridge` annotation per field rather than per property

See below for more information about analyzers/normalizers and field bridges.

4.1.3. Embedded and associated objects

Associated objects as well as embedded objects can be indexed as part of the root entity index. This is useful if you expect to search a given entity based on properties of the associated objects.

In the example [Indexing associations](#) the aim is to return places where the associated city is Atlanta (in Lucene query parser language, it would translate into `address.city:Atlanta`). All place fields are added to the `Place` index, but also the address related fields `address.street`, and `address.city` will be added and made queryable. The embedded object id, `address.id`, is not added per default. To include it you need to also set `@IndexedEmbedded(includeEmbeddedObjectId=true, ...)`.



Only actual indexed fields (properties annotated with `@Field`) are added to the root entity index when embedded objects are indexed. The embedded object identifiers are treated differently and need to be included explicitly.

Example 38. Indexing associations

```
@Entity
@Indexed
public class Place {
    @Id
    @GeneratedValue
    private Long id;

    @Field
    private String name;

    @OneToOne(cascade = { CascadeType.PERSIST, CascadeType.REMOVE })
    @IndexedEmbedded
    private Address address;
    ....
}
```

```
@Entity
public class Address {
    @Id
    @GeneratedValue
    private Long id;

    @Field
    private String street;

    @Field
    private String city;

    @ContainedIn
    @OneToMany(mappedBy="address")
    private Set<Place> places;
    ...
}
```

Be careful. Because the data is de-normalized in the Lucene index when using the `@IndexedEmbedded` technique, Hibernate Search needs to be aware of any change in the `Place` object and any change in the `Address` object to keep the index up to date. To make sure the `Place` Lucene document is updated when its `Address` changes, you need to mark the other side of the bidirectional relationship with `@ContainedIn`.



`@ContainedIn` is useful on both associations pointing to entities and on embedded (collection of) objects.

Let's make [Indexing associations](#) a bit more complex by nesting `@IndexedEmbedded` as seen in [Nested usage of @IndexedEmbedded and @ContainedIn](#).

Example 39. Nested usage of `@IndexedEmbedded` and `@ContainedIn`

```
@Entity
@Indexed
public class Place {
    @Id
    @GeneratedValue
    private Long id;

    @Field
    private String name;

    @OneToOne(cascade = { CascadeType.PERSIST, CascadeType.REMOVE })
    @IndexedEmbedded
    private Address address;

    // ...
}
```

```
@Entity
public class Address {
    @Id
    @GeneratedValue
    private Long id;

    @Field
    private String street;

    @Field
    private String city;

    @IndexedEmbedded(depth = 1, prefix = "ownedBy_")
    private Owner ownedBy;

    @ContainedIn
    @OneToMany(mappedBy="address")
    private Set<Place> places;

    // ...
}
```

```
@Embeddable
public class Owner {
    @Field
    private String name;

    // ...
}
```

As you can see, any `@*ToMany`, `@*ToOne` or `@Embedded` attribute can be annotated with `@IndexedEmbedded`. The attributes of the associated class will then be added to the main entity index. In [Nested usage of `@IndexedEmbedded` and `@ContainedIn`](#) the index will contain the following fields

- `id`
- `name`

- `address.street`
- `address.city`
- `address.ownedBy_name`

The default prefix is `propertyName.`, following the traditional object navigation convention. You can override it using the `prefix` attribute as it is shown on the `ownedBy` property.



The prefix cannot be set to the empty string.

The `depth` property is necessary when the object graph contains a cyclic dependency of classes (not instances). For example, if `Owner` points to `Place`. Hibernate Search will stop including indexed embedded attributes after reaching the expected depth (or the object graph boundaries are reached). A class having a self reference is an example of cyclic dependency. In our example, because `depth` is set to 1, any `@IndexedEmbedded` attribute in `Owner` (if any) will be ignored.

Using `@IndexedEmbedded` for object associations allows you to express queries (using Lucene's query syntax) such as:

- Return places where name contains JBoss and where address city is Atlanta. In Lucene query this would be

```
+name:jboss +address.city:atlanta
```

- Return places where name contains JBoss and where owner's name contain Joe. In Lucene query this would be

```
+name:jboss +address.ownedBy_name:joe
```

In a way it mimics the relational join operation in a more efficient way (at the cost of data duplication). Remember that, out of the box, Lucene indexes have no notion of association, the join operation is simply non-existent. It might help to keep the relational model normalized while benefiting from the full text index speed and feature richness.



An associated object can itself (but does not have to) be `@Indexed`

When `@IndexedEmbedded` points to an entity, the association has to be bi-directional and the other side has to be annotated with `@ContainedIn`. If not, Hibernate Search has no way to update the root index when the associated entity is updated (in our example, a `Place` index document has to be updated when the associated `Address` instance is updated).

Sometimes, the object type annotated by `@IndexedEmbedded` is not the object type targeted by Hibernate and Hibernate Search. This is especially the case when interfaces are used in lieu of their implementation. For this reason you can override the object type targeted by Hibernate Search using

the `targetElement` parameter.

Example 40. Using the `targetElement` property of `@IndexedEmbedded`

```
@Entity
@Indexed
public class Address {
    @Id
    @GeneratedValue
    private Long id;

    @Field
    private String street;

    @IndexedEmbedded(depth = 1, prefix = "ownedBy_", targetElement = Owner.class)
    @Target(Owner.class)
    private Person ownedBy;

    // ...
}
```

```
@Embeddable
public class Owner implements Person { ... }
```

Limiting object embedding to specific paths

The `@IndexedEmbedded` annotation provides also an attribute `includePaths` which can be used as an alternative to `depth`, or in combination with it.

When using only `depth` all indexed fields of the embedded type will be added recursively at the same depth; this makes it harder to pick only a specific path without adding all other fields as well, which might not be needed.

To avoid unnecessarily loading and indexing entities you can specify exactly which paths are needed. A typical application might need different depths for different paths, or in other words it might need to specify paths explicitly, as shown in [Using the `includePaths` property of `@IndexedEmbedded`](#)

Example 41. Using the `includePaths` property of `@IndexedEmbedded`

```
@Entity
@Indexed
public class Person {

    @Id
    public int getId() {
        return id;
    }

    @Field
    public String getName() {
        return name;
    }

    @Field
    public String getSurname() {
        return surname;
    }

    @OneToMany
    @IndexedEmbedded(includePaths = { "name" })
    public Set<Person> getParents() {
        return parents;
    }

    @ContainedIn
    @ManyToOne
    public Human getChild() {
        return child;
    }

    // ... other fields omitted
}
```

Using a mapping as in [Using the `includePaths` property of `@IndexedEmbedded`](#), you would be able to search on a `Person` by `name` and/or `surname`, and/or the `name` of the parent. It will not index the `surname` of the parent, so searching on parent's surnames will not be possible but speeds up indexing, saves space and improve overall performance.

The `@IndexedEmbedded.includePaths` will include the specified paths *in addition to* what you would index normally specifying a limited value for `depth`. Using `includePaths` with a undefined (default) value for `depth` is equivalent to setting `depth=0`: only the included paths are indexed.

Example 42. Using the `includePaths` property of `@IndexedEmbedded`

```
@Entity
@Indexed
public class Human {

    @Id
    public int getId() {
        return id;
    }

    @Field
    public String getName() {
        return name;
    }

    @Field
    public String getSurname() {
        return surname;
    }

    @OneToMany
    @IndexedEmbedded(depth = 2, includePaths = { "parents.parents.name" })
    public Set<Human> getParents() {
        return parents;
    }

    @ContainedIn
    @ManyToOne
    public Human getChild() {
        return child;
    }

    // ... other fields omitted
}
```

In Using the `includePaths` property of `@IndexedEmbedded`, every human will have their name and surname attributes indexed. The name and surname of parents will be indexed too, recursively up to second line because of the `depth` attribute. It will be possible to search by name or surname, of the person directly, his parents or of his grand parents. Beyond the second level, we will in addition index one more level but only the name, not the surname.

This results in the following fields in the index:

- `id` - as primary key
- `_hibernate_class` - stores entity type
- `name` - as direct field
- `surname` - as direct field
- `parents.name` - as embedded field at depth 1
- `parents.surname` - as embedded field at depth 1
- `parents.parents.name` - as embedded field at depth 2
- `parents.parents.surname` - as embedded field at depth 2

- `parents.parents.parents.name` - as additional path as specified by `includePaths`. The first `parents.` is inferred from the field name, the remaining path is the attribute of `includePaths`



You can explicitly include the id of the embedded object using `includePath`, for example `@IndexedEmbedded(includePaths = { "parents.id" })`. This will work regardless of the `includeEmbeddedObjectId` attribute. However, it is recommended to just set `includeEmbeddedObjectId=true`.



Having explicit control of the indexed paths might be easier if you're designing your application by defining the needed queries first, as at that point you might know exactly which fields you need, and which other fields are unnecessary to implement your use case.

Indexing null embeddeds

Per default null values are ignored and not indexed. However, using `indexNullAs` you can specify that a field should be added when the embedded is null, with a value of your choice.

Per default `indexNullAs` is set to `org.hibernate.search.annotations.IndexedEmbedded.DO_NOT_INDEX_NULL`, indicating that null values should not be indexed. You can set this value to `IndexedEmbedded.DEFAULT_NULL_TOKEN` to indicate that a default null token should be used. This default null token can be specified in the configuration using `hibernate.search.default_null_token`. If this property is not set the string `_null_` will be used as default.

The field name used when indexing null values depend on the `prefix`:

- if the `prefix` is not set, the field name will be the Java property name
- if the `prefix` is set, the field name will be the prefix with the trailing dot (if any) removed. For instance with the prefix `my_embedded.`, the null field name will be `my_embedded` (without dot).



When `indexNullAs` is used, it is important to use the chosen null token in search queries (see [Querying](#)) in order to find null values.

4.1.4. Associated objects: building a dependency graph with `@ContainedIn`

While `@ContainedIn` is often seen as the counterpart of `@IndexedEmbedded`, it can also be used on its own to build an indexing dependency graph.

When an entity is reindexed, all the entities pointed by `@ContainedIn` are also going to be reindexed.

4.2. Boosting

Lucene has the notion of *boosting* which allows you to give certain documents or fields more or less importance than others. Lucene differentiates between index and search time boosting. The following sections show you how you can achieve index time boosting using Hibernate Search.

4.2.1. Static index time boosting



Index-time boosting is deprecated: support for this approach will be removed in Apache Lucene since version 7.

We recommend to use query-time boosting instead: see [Querying](#).

To define a static boost value for an indexed class or property you can use the `@Boost` annotation. You can use this annotation within `@Field` or specify it directly on method or class level.

Example 43. Different ways of using @Boost

```
@Entity
@Indexed
@Boost(1.7f)
public class Essay {
    ...

    @Id
    @DocumentId
    public Long getId() { return id; }

    @Field(name="Abstract", store=Store.YES, boost=@Boost(2f))
    @Boost(1.5f)
    public String getSummary() { return summary; }

    @Lob
    @Field(boost=@Boost(1.2f))
    public String getText() { return text; }

    @Field
    public String getISBN() { return isbn; }
}
```

In [Different ways of using @Boost](#), `Essay`'s probability to reach the top of the search list will be multiplied by 1.7. The summary field will be 3.0 ($2 * 1.5$, because `@Field.boost` and `@Boost` on a property are cumulative) more important than the `isbn` field. The text field will be 1.2 times more important than the `isbn` field. Note that this explanation is wrong in strictest terms, but it is simple and close enough to reality for all practical purposes. Please check the Lucene documentation or the excellent *Lucene In Action* from Otis Gospodnetic and Erik Hatcher.

4.2.2. Dynamic index time boosting



Index-time boosting is deprecated: support for this approach will be removed in Apache Lucene since version 7.

We recommend to use query-time boosting instead: see [Querying](#).

If query-time boosting is not enough for your use case, you can implement dynamic boosting using a [custom similarity](#) (advanced use).

The `@Boost` annotation used in [Static index time boosting](#) defines a static boost factor which is independent of the state of the indexed entity at runtime. However, there are use cases in which the boost factor may depend on the actual state of the entity. In this case you can use the `@DynamicBoost` annotation together with an accompanying custom `BoostStrategy`.

Example 44. Dynamic boost example

```
public enum PersonType {  
    NORMAL,  
    VIP  
}
```

```
@Entity  
@Indexed  
@DynamicBoost(impl = VIPBoostStrategy.class)  
public class Person {  
    private PersonType type;  
  
    // ...  
}
```

```
public class VIPBoostStrategy implements BoostStrategy {  
    public float defineBoost(Object value) {  
        Person person = ( Person ) value;  
        if ( person.getType().equals( PersonType.VIP ) ) {  
            return 2.0f;  
        }  
        else {  
            return 1.0f;  
        }  
    }  
}
```

In [Dynamic boost example](#) a dynamic boost is defined on class level specifying `VIPBoostStrategy` as implementation of the `BoostStrategy` interface to be used at indexing time. You can place the `@DynamicBoost` either at class or field level. Depending on the placement of the annotation either the whole entity is passed to the `defineBoost` method or just the annotated field/property value. It's up to you to cast the passed object to the correct type. In the example all indexed values of a VIP person would be double as important as the values of a normal person.



The specified `BoostStrategy` implementation must define a public no-arg constructor.

Of course you can mix and match `@Boost` and `@DynamicBoost` annotations in your entity. All defined boost factors are cumulative.

4.3. Analysis

Analysis is the process of converting text into single terms (words) and can be considered as one of the key features of a fulltext search engine. Lucene uses the concept of Analyzers to control this process. In the following section we cover the multiple ways Hibernate Search offers to configure the analyzers.

4.3.1. Default analyzer and analyzer by class

The default analyzer class used to index tokenized fields is configurable through the `hibernate.search.analyzer` property. The default value for this property is `org.apache.lucene.analysis.standard.StandardAnalyzer`.

You can also define the analyzer class per entity, property and even per `@Field` (useful when multiple fields are indexed from a single property).



Support for `@Analyzer` on entity types and properties is deprecated and is going to be removed in Hibernate Search 6.

Assign an analyzer to each index field separately using `@Field(analyzer = @Analyzer(...))` instead.

Example 45. Different ways of using @Analyzer

```
@Entity
@Indexed
@Analyzer(definition = "myEntityAnalyzer")
public class MyEntity {
    @Id
    @GeneratedValue
    @DocumentId
    private Integer id;

    @Field
    private String name;

    @Field
    @Analyzer(definition = "myPropertyAnalyzer")
    private String summary;

    @Field(analyzer = @Analyzer(definition = "myFieldAnalyzer"))
    private String body;

    ...
}
```

In this example, `myEntityAnalyzer` is used to index all tokenized properties (eg. `name`), except `summary` and `body` which are indexed with `myPropertyAnalyzer` and `myFieldAnalyzer` respectively.



`@Analyzer.definition` just points to an analyzer name; an actual analyzer must be assigned to that name through an [analyzer definition](#).



Mixing different analyzers in the same entity is most of the time a bad practice. It makes query building more complex and results less predictable (for the novice), especially if you are using a `QueryParser` (which uses the same analyzer for the whole query). As a rule of thumb, for any given field the same analyzer should be used for indexing and querying.

4.3.2. Named analyzers

Analyzers can become quite complex to deal with. For this reason Hibernate Search introduces the notion of analyzer definitions. An analyzer definition can be reused by many `@Analyzer` declarations and is composed of:

- a name: the unique string used to refer to the definition
- a list of char filters: each char filter is responsible to pre-process input characters before the tokenization. Char filters can add, change or remove characters; one common usage is for characters normalization
- a tokenizer: responsible for tokenizing the input stream into individual words

- a list of filters: each filter is responsible to remove, modify or sometimes even add words into the stream provided by the tokenizer

This separation of tasks - a list of char filters, and a tokenizer followed by a list of filters - allows for easy reuse of each individual component and let you build your customized analyzer in a very flexible way (just like Lego). Generally speaking the char filters do some pre-processing in the character input, then the Tokenizer starts the tokenizing process by turning the character input into tokens which are then further processed by the TokenFilters. Hibernate Search supports this infrastructure by utilizing the advanced analyzers provided by Lucene; this is often referred to as the Analyzer Framework.



Filters and char filters are applied in the order they are mentioned in the analyzer definition. Order matters!



Some of the analyzers and filters will require additional dependencies. For example to use the snowball stemmer you have to also include the `lucene-snowball` jar and for the `PhoneticFilterFactory` you need the `commons-codec` jar. Your distribution of Hibernate Search provides these dependencies in its `lib/optional` directory. Have a look at [Example of available tokenizers](#) and [Examples of available filters](#) to see which analyzers and filters have additional dependencies

Prior to Hibernate Search 5 it was required to add the Apache Solr dependency to your project as well; this is no longer required.

Programmatic analyzer definition

The first way to define analyzers is programmatically. You can of course use the [programmatic mapping API](#) to do so, but an easier way may be to use the `hibernate.search.lucene.analysis_definition_provider` configuration property.

This property can be set to the fully-qualified name of a class with a public, no-arg constructor in your application. This class must either implement `org.hibernate.search.analyzer.definition.LuceneAnalysisDefinitionProvider` directly or expose a `@Factory`-annotated method that returns such implementation.

Example 46. Analyzer definition using a `LuceneAnalysisDefinitionProvider`

```
public static class CustomAnalyzerProvider implements LuceneAnalysisDefinitionProvider {
    @Override
    public void register(LuceneAnalyzerDefinitionRegistryBuilder builder) {
        builder
            .analyzer( "myAnalyzer" )
                .tokenizer( StandardTokenizerFactory.class )
                .charFilter( MappingCharFilterFactory.class )
                    .param( "mapping",
                        "org/hibernate/search/test/analyzer/mapping-chars.properties" )
                .tokenFilter( ASCIIFoldingFilterFactory.class )
                .tokenFilter( LowerCaseFilterFactory.class )
                .tokenFilter( StopFilterFactory.class )
                    .param( "mapping",
                        "org/hibernate/search/test/analyzer/stoplist.properties" )
                    .param( "ignoreCase", "true" );
    }
}
```

Analyzer definition through annotations



Support for defining analyzers or normalizers through annotations is deprecated, and will be removed in Hibernate Search 6.

Use [a definition provider](#) instead.

The second way to define analyzers is by applying `@AnalyzerDef` annotations to your indexed classes. Which class you choose to put the annotation on does not matter: the definitions are globally scoped, so you will be able to reference them from any entity mapping.

Let's have a look at a concrete example now - [Analyzer definition using the `@AnalyzerDef` annotation](#). First a char filter is defined by its factory. In our example, a mapping char filter is used, and will replace characters in the input based on the rules specified in the mapping file. Next a tokenizer is defined. This example uses the standard tokenizer. Last but not least, a list of filters is defined by their factories. In our example, the StopFilter filter is built reading the dedicated words property file. The filter is also expected to ignore case.

Example 47. Analyzer definition using the @AnalyzerDef annotation

```
@AnalyzerDef(name="customanalyzer",
    charFilters = {
        @CharFilterDef(factory = MappingCharFilterFactory.class, params = {
            @Parameter(name = "mapping",
                value = "org/hibernate/search/test/analyzer/mapping-chars.properties")
        })
    },
    tokenizer = @TokenizerDef(factory = StandardTokenizerFactory.class),
    filters = {
        @TokenFilterDef(factory = ASCIIFoldingFilterFactory.class),
        @TokenFilterDef(factory = LowerCaseFilterFactory.class),
        @TokenFilterDef(factory = StopFilterFactory.class, params = {
            @Parameter(name="words",
                value= "org/hibernate/search/test/analyzer/stoplist.properties" ),
            @Parameter(name="ignoreCase", value="true")
        })
    })
})
public class Team {
    // ...
}
```

Referencing named analyzers

Once defined, an analyzer definition can be reused by an `@Analyzer` declaration as seen in [Referencing an analyzer by name](#).

Example 48. Referencing an analyzer by name

```
@Entity
@Indexed
@AnalyzerDef(name="customanalyzer", ... )
public class Team {
    @Id
    @DocumentId
    @GeneratedValue
    private Integer id;

    @Field
    private String name;

    @Field
    private String location;

    @Field(analyzer = @Analyzer(definition = "customanalyzer"))
    private String description;
}
```

Analyzer definitions are also available by their name in the `SearchFactory` which is quite useful when building queries.

```
Analyzer analyzer = fullTextSession.getSearchFactory().getAnalyzer("customanalyzer");
```

Fields in queries should be analyzed with the same analyzer used to index the field so that they speak a common "language": the same tokens are reused between the query and the indexing process. This rule has some exceptions but is true most of the time. Respect it unless you know what you are doing.

Normalizers

Analyzers are great when you need to search in text documents, but what if you want to sort the analyzed text? Then you're in for a bit of trouble, because analyzed text is multi-valued: when indexing a book with the title `"Refactoring: Improving the Design of Existing Code"`, the analyzed title is in fact the (unordered) set `{"refactoring", "improving", "design", "existing", "code"}`. If you tried to sort on the title after such an analysis, any of those words could be used, so your book could end up in the D's (because of "design"), or in the R's (because of "refactoring"), or in the E's, etc.

So in the end, you probably don't want your fields to be tokenized when you sort on those fields. Normalizers solve exactly this issue: they are analyzers, but without a tokenizer, and with some runtime checks that prevent the analysis to produce multiple tokens, thereby ensuring your sorts will always be consistent.

Hibernate Search provides normalizer equivalents for relevant analyzer annotations: `@Normalizer`, `@NormalizerDef`, `@NormalizerDefs`. As with analyzers, you can use named normalizers (for instance `@Normalizer(definition = "myNormalizer")` with `@NormalizerDef(name = "myNormalizer", filters = @TokenFilterDef(factory = LowerCaseFilterFactory.class))`). You can also refer to implementations directly (for instance `@Normalizer(impl = MyCollactionKeyAnalyzer.class)`), though that is deprecated and will no longer be possible in Hibernate Search 6.



You cannot define a "class-level", "property-level", or a "default" normalizer: each field requiring normalization must be assigned a normalizer explicitly using `@Field.normalizer`.

Available analysis components

Apache Lucene comes with a lot of useful default char filters, tokenizers and filters. You can find a complete list of char filter factories, tokenizer factories and filter factories at <http://wiki.apache.org/solr/AnalyzersTokenizersTokenFilters>. Let's check a few of them.

Table 11. Example of available char filters

Factory	Description	Parameters	Additional dependencies
<code>MappingCharFilterFactory</code>	Replaces one or more characters with one or more characters, based on mappings specified in the resource file	<code>mapping</code> : points to a resource file containing the mappings using the format: "á" ⇒ "a" "ñ" ⇒ "n" "ø" ⇒ "o"	<code>lucene-analyzers-common</code>
<code>HTMLStripCharFilterFactory</code>	Remove HTML standard tags, keeping the text	none	<code>lucene-analyzers-common</code>

Table 12. Example of available tokenizers

Factory	Description	Parameters	Additional dependencies
<code>StandardTokenizerFactory</code>	Use the Lucene StandardTokenizer	none	<code>lucene-analyzers-common</code>
<code>HTMLStripCharFilterFactory</code>	Remove HTML tags, keep the text and pass it to a <code>StandardTokenizer</code> .	none	<code>lucene-analyzers-common</code>
<code>PatternTokenizerFactory</code>	Breaks text at the specified regular expression pattern.	<code>pattern</code> : the regular expression to use for tokenizing <code>group</code> : says which pattern group to extract into tokens	<code>lucene-analyzers-common</code>

Table 13. Examples of available filters

Factory	Description	Parameters	Additional dependencies
<code>StandardFilterFactory</code>	Remove dots from acronyms and 's from words	none	<code>lucene-analyzers-common</code>
<code>LowerCaseFilterFactory</code>	Lowercases all words	none	<code>lucene-analyzers-common</code>
<code>StopFilterFactory</code>	Remove words (tokens) matching a list of stop words	<code>words</code> : points to a resource file containing the stop words <code>ignoreCase</code> : true if case should be ignore when comparing stop words, <code>false</code> otherwise	<code>lucene-analyzers-common</code>

Factory	Description	Parameters	Additional dependencies
<code>SnowballPorterFilterFactory</code>	Reduces a word to its root in a given language. (eg. protect, protects, protection share the same root). Using such a filter allows searches matching related words.	<code>language</code> : Danish, Dutch, English, Finnish, French, German, Italian, Norwegian, Portuguese, Russian, Spanish, Swedish and a few more	<code>lucene-analyzers-common</code>
<code>ASCIIFoldingFilterFactory</code>	Remove accents for languages like French	none	<code>lucene-analyzers-common</code>
<code>PhoneticFilterFactory</code>	Inserts phonetically similar tokens into the token stream	<code>encoder</code> : One of <code>DoubleMetaphone</code> , <code>Metaphone</code> , <code>Soundex</code> or <code>RefinedSoundex</code> <code>inject</code> : true will add tokens to the stream, false will replace the existing token <code>maxCodeLength</code> : sets the maximum length of the code to be generated. Supported only for Metaphone and DoubleMetaphone encodings	<code>lucene-analyzers-phonetic</code> and <code>commons-codec</code>
<code>CollationKeyFilterFactory</code>	Converts each token into its <code>java.text.CollationKey</code> , and then encodes the <code>CollationKey</code> with <code>IndexableBinaryStringTools</code> , to allow it to be stored as an index term.	<code>custom</code> , <code>language</code> , <code>country</code> , <code>variant</code> , <code>strength</code> , <code>decomposition</code> see Lucene's <code>CollationKeyFilter</code> javadocs for more info	<code>lucene-analyzers-common</code> and <code>commons-io</code>

We recommend to check out the implementations of `org.apache.lucene.analysis.util.TokenizerFactory` and `org.apache.lucene.analysis.util.TokenFilterFactory` in your IDE to see the implementations available.

4.3.3. Dynamic analyzer selection

So far all the introduced ways to specify an analyzer were static. However, there are use cases where it is useful to select an analyzer depending on the current state of the entity to be indexed, for example in a multilingual applications. For an `BlogEntry` class for example the analyzer could depend on the language property of the entry. Depending on this property the correct language specific stemmer should be chosen to index the actual text.

To enable this dynamic analyzer selection Hibernate Search introduces the `@AnalyzerDiscriminator` annotation. [Usage of @AnalyzerDiscriminator](#) demonstrates the usage of this annotation.

Example 49. Usage of @AnalyzerDiscriminator

```
@Entity
@Indexed
@AnalyzerDef(name = "en",
    tokenizer = @TokenizerDef(factory = StandardTokenizerFactory.class),
    filters = {
        @TokenFilterDef(factory = LowerCaseFilterFactory.class),
        @TokenFilterDef(factory = EnglishPorterFilterFactory.class)
    }
)
@AnalyzerDef(name = "de",
    tokenizer = @TokenizerDef(factory = StandardTokenizerFactory.class),
    filters = {
        @TokenFilterDef(factory = LowerCaseFilterFactory.class),
        @TokenFilterDef(factory = GermanStemFilterFactory.class)
    }
)
public class BlogEntry {

    @Id
    @GeneratedValue
    @DocumentId
    private Integer id;

    @Field
    @AnalyzerDiscriminator(impl = LanguageDiscriminator.class)
    private String language;

    @Field
    private String text;

    private Set<BlogEntry> references;

    // standard getter/setter
    // ...
}
```

```
public class LanguageDiscriminator implements Discriminator {

    public String getAnalyzerDefinitionName(Object value, Object entity, String field) {
        if ( value == null || !( entity instanceof Article ) ) {
            return null;
        }
        return (String) value;
    }

}
```

The prerequisite for using `@AnalyzerDiscriminator` is that all analyzers which are going to be used dynamically are predefined as [named analyzers](#). If this is the case, one can place the `@AnalyzerDiscriminator` annotation either on the class or on a specific property of the entity for which to dynamically select an analyzer. Via the `impl` parameter of the `@AnalyzerDiscriminator` you specify a concrete implementation of the `Discriminator` interface. It is up to you to provide an

implementation for this interface. The only method you have to implement is `getAnalyzerDefinitionName()` which gets called for each field added to the Lucene document. The entity which is getting indexed is also passed to the interface method. The `value` parameter is only set if the `AnalyzerDiscriminator` is placed on property level instead of class level. In this case the value represents the current value of this property.

An implementation of the `Discriminator` interface has to return the name of an existing analyzer definition or null if the default analyzer should not be overridden. [Usage of @AnalyzerDiscriminator](#) assumes that the language parameter is either 'de' or 'en' which matches the name of an analyzer.



There is no equivalent of `@AnalyzerDiscriminator` for [normalizers](#), at least for now. See [HSEARCH-2738](#).

4.3.4. Retrieving an analyzer

In some situations retrieving analyzers can be handy. For example, if your domain model makes use of multiple analyzers (maybe to benefit from stemming, use phonetic approximation and so on), you need to make sure to use the same analyzers when you build your query.



This rule can be broken but you need a good reason for it. If you are unsure, use the same analyzers. If you use the Hibernate Search query DSL (see [Building a Lucene query with the Hibernate Search query DSL](#)), you don't have to think about it. The query DSL does use the right analyzer transparently for you.

Whether you are using the Lucene programmatic API or the Lucene query parser, you can retrieve the scoped analyzer for a given entity. A scoped analyzer is an analyzer which applies the right analyzers depending on the field indexed. Remember, multiple analyzers can be defined on a given entity each one working on an individual field. A scoped analyzer unifies all these analyzers into a context-aware analyzer. While the theory seems a bit complex, using the right analyzer in a query is very easy.

Example 50. Using the scoped analyzer when building a full-text query

```
org.apache.lucene.queryparser.classic.QueryParser parser = new QueryParser(
    "title",
    fullTextSession.getSearchFactory().getAnalyzer( Song.class )
);

org.apache.lucene.search.Query luceneQuery =
    parser.parse( "title:sky Or title_stemmed:diamond" );

org.hibernate.Query fullTextQuery =
    fullTextSession.createFullTextQuery( luceneQuery, Song.class );

List result = fullTextQuery.list(); //return a list of managed objects
```

In the example above, the song title is indexed in two fields: the standard analyzer is used in the field

`title` and a stemming analyzer is used in the field `title_stemmed`. By using the analyzer provided by the search factory, the query uses the appropriate analyzer depending on the field targeted.



You can also retrieve `named analyzers` by their definition name using `searchFactory.getAnalyzer(String)`.

4.4. Bridges

When discussing the basic mapping for an entity one important fact was so far disregarded. In Lucene all index fields have to be represented as strings. All entity properties annotated with `@Field` have to be converted to strings to be indexed. The reason we have not mentioned it so far is, that for most of your properties Hibernate Search does the translation job for you thanks to a set of built-in bridges. However, in some cases you need a more fine grained control over the translation process.

4.4.1. Built-in bridges

Hibernate Search comes bundled with a set of built-in bridges between a Java property type and its full text representation.

`null`

Per default `null` elements are not indexed. Lucene does not support `null` elements. However, in some situation it can be useful to insert a custom token representing the `null` value. See `@Field` for more information.

`java.lang.String`

Strings are indexed as they are.

`boolean, Boolean`

Are converted into their string representation and indexed.

`char, Character`

Are converted into their string representation and indexed.

`int, Integer, long, Long, float, Float, double, Double`

Are per default indexed numerically using a `Trie structure` which allows for efficient range queries and sorting, resulting in better query response times. You need to use a `NumericRangeQuery` to search for values. See also `@Field` and `@NumericField`.



Prior to Hibernate Search 5, numeric field encoding was only chosen if explicitly requested via `@NumericField`. As of Hibernate Search 5, this encoding is automatically chosen for the numeric types mentioned above.

To avoid numeric encoding you can explicitly specify a non numeric field bridge via `@Field.bridge` or `@FieldBridge`. The package `org.hibernate.search.bridge.builtin` contains a set of bridges which encode numbers as strings, for example `org.hibernate.search.bridge.builtin.IntegerBridge`.

`byte, Byte, short, Short`

Are converted into their string representation and indexed. They will only be encoded in numeric fields if explicitly marked with the `@NumericField` annotation (see `@NumericField`).

`BigInteger, BigDecimal`

Are converted into their string representation and indexed. Note that in this form the values cannot be compared by Lucene using for example a `TermRangeQuery`. For that the string representation would need to be padded. An alternative using numeric encoding with a potential loss in precision can be seen in [Defining a custom FieldBridge for BigDecimal](#).

`java.util.Date, java.util.Calendar`

Dates are indexed as `long` value representing the number of milliseconds since *January 1, 1970, 00:00:00 GMT*. You shouldn't really bother with the internal format. It is important, however, to query a numerically indexed date via a `NumericRangeQuery`.

Usually, storing the date up to the millisecond is not necessary. `@DateBridge` and `@CalendarBridge` allow to define the appropriate resolution you are willing to store in the index.

```
@Entity
@Indexed
public class Meeting {
    @Field(analyze=Analyze.NO)
    @DateBridge(resolution=Resolution.MINUTE)
    private Date date;
    // ...
}
```

You can also choose to encode the date as string using the `encoding=EncodingType.STRING` of `@DateBridge` and `@CalendarBridge`. In this case the dates are stored in the format `yyyyMMddHHmmssSSS` (using GMT time).



The default date bridge uses Lucene's `DateTools` to convert from `Date` or `Calendar` to its indexed value. This means that all dates are expressed in GMT time. If your requirements are to store dates in a fixed time zone you have to implement a custom date bridge.

`java.time.Year`

converts the year to the integer representation.

`java.time.Duration`

converts the duration to the total length in nanoseconds.

`java.time.Instant`

converts the instant to the number of milliseconds from Epoch. Note that these values are indexed with a precision to the millisecond.



Note that it must be possible to convert the `Instant` or the `Duration` to a `Long`. If these values are too big or too small an exception is thrown.

`LocalDate`, `LocalTime`, `LocalDateTime`, `MonthDay`, `OffsetDateTime`, `OffsetTime`, `Period`, `YearMonth`, `ZoneDateTime`, `ZoneId`, `ZoneOffset`

the bridges for these classes in the `java.time` package store the values as string padded with 0 when required to allow sorting.

`java.net.URI`, `java.net.URL`

`URI` and `URL` are converted to their string representation

`java.lang.Class`

Classes are converted to their fully qualified class name. The thread context classloader is used when the class is rehydrated

4.4.2. Tika bridge



Support for Tika bridges is deprecated, and will be removed in Hibernate Search 6.

If you need this feature, vote for it to be re-implemented using the new bridge API in a later version of Hibernate Search: [HSEARCH-3350](#). We'll also gladly help anyone interested in contributing a patch: feel free to [contact us](#).

Hibernate Search allows you to extract text from various document types using the built-in `TikaBridge` which utilizes [Apache Tika](#) to extract text and metadata from the provided documents.

The `@TikaBridge` annotation can be used with `String`, `URI`, `byte[]` or `java.sql.Blob` properties. In the case of `String` and `URI` the bridge interprets the values as file paths and tries to open a file to

parse the document. In the case of `byte[]` and `Blob` the values are directly passed to Tika for parsing.

When parsing a value, Tika uses:

- the document content as input
- metadata as input and output
- a `ParseContext` as input

You can find more information about the parsing process in [Parser interface](#). The Hibernate Search Tika bridge allows you to customize the whole process by passing implementations of three interfaces to the `TikaBridge`.

The first interface is the `TikaParserProvider`. It defines which parser will be used on the document content, i.e. how the document should be interpreted. The default behavior should cover most use cases, since it uses Tika's `AutoDetectParser` to detect the content type and use an appropriate parser, but using a custom provider can still be useful, for example to use a custom Tika configuration file.

The second interface is `TikaMetadataProcessor` which has two methods - `prepareMetadata()` and `set(String, Object, Document, LuceneOptions, Metadata metadata)`. The former allows to add additional metadata to the parsing process (for example the file name) and the latter allows you to index metadata discovered during the parsing process.



Sortable fields with Tika

If you want to add multiple fields in your `TikaMetadataProcessor`, and also want to make those fields sortable, you should make your processor implement the `MetadataProvidingTikaMetadataProcessor`. This is similar to implementing `MetadataProvidingFieldBridge` on a regular field bridge: see [@SortableField](#).

The `MetadataProvidingTikaMetadataProcessor` contract is under active development and considered experimental at this time. It may be altered in future revisions, e.g. by adding further methods, thus breaking existing implementations.

The third interface is the `TikaParseContextProvider`. It allows you to create a custom `ParseContext` for the document parsing.

Implementations of the interfaces mentioned above are passed to the `TikaBridge` by specifying their class as parameters on the `TikaBridge` annotation.

Example 51. Example mapping with Apache Tika

```
@Entity
@Indexed
public class Song {
    @Id
    @GeneratedValue
    long id;

    @Field
    @TikaBridge(metadataProcessor = Mp3TikaMetadataProcessor.class)
    String mp3FileName;

    // ...
}
```

```
QueryBuilder queryBuilder = fullTextSession.getSearchFactory()
    .buildQueryBuilder()
    .forEntity( Song.class )
    .get();
Query query = queryBuilder.keyword()
    .onField( "mp3FileName" )
    .ignoreFieldBridge() //mandatory
    .matching( "Apes" )
    .createQuery();
List result = fullTextSession.createFullTextQuery( query ).list();
```

In the [Example mapping with Apache Tika](#) the property `mp3FileName` represents a path to an MP3 file; the headers of this file will be indexed and so the performed query will be able to match the MP3 metadata.



`TikaBridge` does not implement `TwoWayFieldBridge`: queries built using the DSL (as in the [Example mapping with Apache Tika](#)) need to explicitly enable the option `ignoreFieldBridge()`.



To use the `@TikaBridge` in combination with WildFly modules, make sure to package your Tika extensions as described in [Using Apache Tika and other extensions with modules for WildFly](#).

4.4.3. Custom bridges

Sometimes, the built-in bridges of Hibernate Search do not cover some of your property types, or the `String` representation used by the bridge does not meet your requirements. The following paragraphs describe several solutions to this problem.

StringBridge

The simplest custom solution is to give Hibernate Search an implementation of your expected `Object` to `String` bridge. To do so you need to implement the

`org.hibernate.search.bridge.StringBridge` interface. All implementations have to be thread-safe as they are used concurrently.

Example 52. Custom `StringBridge` implementation

```
/**
 * Padding Integer bridge.
 * All numbers will be padded with 0 to match 5 digits
 *
 * @author Emmanuel Bernard
 */
public class PaddedIntegerBridge implements StringBridge {

    private int padding = 5;

    public String objectToString(Object object) {
        String rawInteger = ((Integer) object).toString();
        if (rawInteger.length() > padding)
            throw new IllegalArgumentException("Number too big to be padded");
        StringBuilder paddedInteger = new StringBuilder();
        for (int padIndex = rawInteger.length(); padIndex < padding; padIndex++) {
            paddedInteger.append('0');
        }
        return paddedInteger.append( rawInteger ).toString();
    }
}
```

Given the string bridge defined in [Custom StringBridge implementation](#), any property or field can use this bridge thanks to the `@FieldBridge` annotation:

```
@FieldBridge(impl = PaddedIntegerBridge.class)
private Integer length;
```

Parameterized bridge

Parameters can also be passed to the bridge implementation making it more flexible. [Passing parameters to your bridge implementation](#) implements a `ParameterizedBridge` interface and parameters are passed through the `@FieldBridge` annotation.

Example 53. Passing parameters to your bridge implementation

```
public class PaddedIntegerBridge implements StringBridge, ParameterizedBridge {

    public static String PADDING_PROPERTY = "padding";
    private int padding = 5; //default

    public void setParameterValues(Map<String,String> parameters) {
        String padding = parameters.get( PADDING_PROPERTY );
        if (padding != null) this.padding = Integer.parseInt( padding );
    }

    public String objectToString(Object object) {
        String rawInteger = ((Integer) object).toString();
        if (rawInteger.length() > padding)
            throw new IllegalArgumentException("Number too big to be padded");
        StringBuilder paddedInteger = new StringBuilder( );
        for (int padIndex = rawInteger.length(); padIndex < padding; padIndex++) {
            paddedInteger.append('0');
        }
        return paddedInteger.append(rawInteger).toString();
    }
}
```

```
//on the property:
@FieldBridge(impl = PaddedIntegerBridge.class,
            params = @Parameter(name="padding", value="10")
            )
private Integer length;
```

The `ParameterizedBridge` interface can be implemented by `StringBridge`, `TwoWayStringBridge`, `FieldBridge` implementations.

All implementations have to be thread-safe, but the parameters are set during initialization and no special care is required at this stage.

Type aware bridge

It is sometimes useful to get the type the bridge is applied on:

- the return type of the property for field/getter-level bridges
- the class type for class-level bridges

An example is a bridge that deals with enums in a custom fashion but needs to access the actual enum type. Any bridge implementing `AppliedOnTypeAwareBridge` will get the type the bridge is applied on injected. Like parameters, the type injected needs no particular care with regard to thread-safety.

FieldBridge

Some use cases require more than a simple object to string translation when mapping a property to a Lucene index. To give you the greatest possible flexibility you can also implement a bridge as a

FieldBridge. This interface gives you a property value and let you map it the way you want in your Lucene **Document**. You can for example store a property in two different document fields. The interface is very similar in its concept to the Hibernate ORM **UserTypes**.

Example 54. Implementing the FieldBridge interface

```
/**
 * Store the date in 3 different fields - year, month, day - to ease the creation of
 * RangeQuery per
 * year, month or day (eg get all the elements of December for the last 5 years).
 * @author Emmanuel Bernard
 */
public class DateSplitBridge implements FieldBridge {
    private final static TimeZone GMT = TimeZone.getTimeZone("GMT");

    public void set(String name, Object value, Document document,
                    LuceneOptions luceneOptions) {
        Date date = (Date) value;
        Calendar cal = GregorianCalendar.getInstance(GMT);
        cal.setTime(date);
        int year = cal.get(Calendar.YEAR);
        int month = cal.get(Calendar.MONTH) + 1;
        int day = cal.get(Calendar.DAY_OF_MONTH);

        // set year
        luceneOptions.addFieldToDocument(
            name + ".year",
            String.valueOf( year ),
            document );

        // set month and pad it if needed
        luceneOptions.addFieldToDocument(
            name + ".month",
            month < 10 ? "0" : "" + String.valueOf( month ),
            document );

        // set day and pad it if needed
        luceneOptions.addFieldToDocument(
            name + ".day",
            day < 10 ? "0" : "" + String.valueOf( day ),
            document );
    }
}
```

```
//property
@FieldBridge(impl = DateSplitBridge.class)
private Date date;
```

In **Implementing the FieldBridge interface** the fields are not added directly to **Document**. Instead the addition is delegated to the **LuceneOptions** helper; this helper will apply the options you have selected on **@Field**, like **Store** or **TermVector**, or apply the chosen **@Boost** value. It is especially useful to encapsulate the complexity of **COMPRESS** implementations. Even though it is recommended to delegate to **LuceneOptions** to add fields to the **Document**, nothing stops you from editing the **Document** directly and ignore the **LuceneOptions** in case you need to.



Classes like `LuceneOptions` are created to shield your application from changes in Lucene API and simplify your code. Use them if you can, but if you need more flexibility you're not required to.

ClassBridge

It is sometimes useful to combine more than one property of a given entity and index this combination in a specific way into the Lucene index. The `@ClassBridge` and `@ClassBridges` annotations can be defined at class level (as opposed to the property level). In this case the custom field bridge implementation receives the entity instance as the value parameter instead of a particular property. Though not shown in [Implementing a class bridge](#), `@ClassBridge` supports the `termVector` attribute discussed in section [Basic mapping](#).

Example 55. Implementing a class bridge

```
@Entity
@Indexed
@ClassBridge(name="branchnetwork",
    store=Store.YES,
    impl = CatFieldsClassBridge.class,
    params = @Parameter( name="sepChar", value=" " ) )
public class Department {
    private int id;
    private String network;
    private String branchHead;
    private String branch;
    private Integer maxEmployees
    // ...
}
```

```
public class CatFieldsClassBridge implements FieldBridge, ParameterizedBridge {
    private String sepChar;

    public void setParameterValues(Map parameters) {
        this.sepChar = (String) parameters.get( "sepChar" );
    }

    public void set(
        String name, Object value, Document document, LuceneOptions luceneOptions) {
        // In this particular class the name of the new field was passed
        // from the name field of the ClassBridge Annotation. This is not
        // a requirement. It just works that way in this instance. The
        // actual name could be supplied by hard coding it below.
        Department dep = (Department) value;
        String fieldValue1 = dep.getBranch();
        if ( fieldValue1 == null ) {
            fieldValue1 = "";
        }
        String fieldValue2 = dep.getNetwork();
        if ( fieldValue2 == null ) {
            fieldValue2 = "";
        }
        String fieldValue = fieldValue1 + sepChar + fieldValue2;
        luceneOptions.addFieldToDocument( name, fieldValue, document );
    }
}
```

In this example, the particular `CatFieldsClassBridge` is applied to the `department` instance, the field bridge then concatenate both branch and network and index the concatenation.

Two-way bridge

If you expect to retrieve values from the index (for example when using [projections](#)), you need to use what Hibernate Search calls "two-way bridges". Bridges applied to document identifiers in particular must be two-way bridges.

Those bridges are applied to your model using the `@FieldBridge` annotation just like any other bridge, but their implementation is slightly different.

The simpler two-way bridge interface, `TwoWayStringBridge`, is similar to `StringBridge` but exposes two methods:

- The `String objectToString(Object)` method transforms the property value into a `String`. This allows Hibernate Search to generate the value of the field when indexing. The method is also used to generate search queries when filtering on the field.
- The `Object stringToObject(String)` method transforms the value stored in the index back into the value of the associated entity property. This allows Hibernate Search to fetch a value when projecting. In the case of document identifiers, it also allows to load entities from index documents when fetching search results.

Example 56. Implementing a `TwoWayStringBridge` usable for id properties

```
public class PaddedIntegerBridge implements TwoWayStringBridge, ParameterizedBridge {

    public static String PADDING_PROPERTY = "padding";
    private int padding = 5; //default

    public void setParameterValues(Map parameters) {
        Object padding = parameters.get(PADDING_PROPERTY);
        if (padding != null) this.padding = (Integer) padding;
    }

    public String objectToString(Object object) {
        String rawInteger = ((Integer) object).toString();
        if (rawInteger.length() > padding)
            throw new IllegalArgumentException("Number too big to be padded");
        StringBuilder paddedInteger = new StringBuilder();
        for (int padIndex = rawInteger.length(); padIndex < padding ; padIndex++) {
            paddedInteger.append('0');
        }
        return paddedInteger.append(rawInteger).toString();
    }

    public Object stringToObject(String stringValue) {
        return new Integer(stringValue);
    }
}
```

```
//On an id property:
@DocumentId
@FieldBridge(impl = PaddedIntegerBridge.class,
    params = @Parameter(name="padding", value="10")
private Integer id;
```

The more powerful `TwoWayFieldBridge` interface allows you to store more than one field into the index, which can be useful for composite properties, but is more complex to implement:

- The `void set(String name, Object value, Document document, LuceneOptions luceneOptions)` method transforms the property value into one or several fields. This allows Hibernate Search to generate the document when indexing.
- The `Object get(String name, Document document)` method transforms the value stored

in the index back into the value of the associated entity property. This allows Hibernate Search to fetch a value when projecting. In the case of document identifiers, it also allows to load entities from index documents when fetching search results.

- The `String objectToString(Object)` method transforms the property value into a `String`. The method is used to generate search queries when filtering on the main field (the field whose name is passed to the `get/set` methods).

It is important for the two-way process to be idempotent, i.e.:

- for `TwoWayStringBridge`:
`object.equals(stringToObject(objectToString(object)))`.
- for `TwoWayFieldBridge`: for a given document, the object returned by `get()` must be the same as the one that was originally passed to `set()`.



Also, in order for `TwoWayFieldBridge` implementations to work correctly when used on a document identifier, the `void set(String name, Object value, Document document, LuceneOptions luceneOptions)` method **must** add a field to the document following these conventions:

- the field name must be the name provided in the `name` parameter
- the field value must be encoded as a `String`
- the field value must be unique to the given value of the `value` parameter
- the field value must match what the `objectToString` method would return for the given `value` parameter.

4.4.4. BridgeProvider: associate a bridge to a given return type

Custom field bridges are very flexible, but it can be tedious and error prone to apply the same custom `@FieldBridge` annotation every time a property of a given type is present in your domain model. That is what `BridgeProviders` are for.

Let's imagine that you have a type `Currency` in your application and that you want to apply your very own `CurrencyFieldBridge` every time an indexed property returns `Currency`. You can do it the hard way:

Example 57. Applying the same `@FieldBridge` for a type the hard way

```
@Entity @Indexed
public class User {
    @FieldBridge(impl=CurrencyFieldBridge.class)
    public Currency getDefaultCurrency();

    // ...
}

@Entity @Indexed
public class Account {
    @FieldBridge(impl=CurrencyFieldBridge.class)
    public Currency getCurrency();

    // ...
}

// continue to add @FieldBridge(impl=CurrencyFieldBridge.class) everywhere Currency is
```

Or you can write your own `BridgeProvider` implementation for `Currency`.

Example 58. Writing a `BridgeProvider`

```
public class CurrencyBridgeProvider implements BridgeProvider {

    //needs a default no-arg constructor

    @Override
    public FieldBridge provideFieldBridge(BridgeContext bridgeProviderContext) {
        if ( bridgeProviderContext.getReturnType().equals( Currency.class ) ) {
            return CurrencyFieldBridge.INSTANCE;
        }
        return null;
    }
}
```

```
# service file named META-INF/services/org.hibernate.search.bridge.spi.BridgeProvider
com.acme.myapps.hibernatesearch.CurrencyBridgeProvider
```

You need to implement `BridgeProvider` and create a service file named `META-INF/services/org.hibernate.search.bridge.spi.BridgeProvider`. This file must contain the fully qualified class name(s) of the `BridgeProvider` implementations. This is the classic Service Loader discovery mechanism.

Now, any indexed property of type `Currency` will use `CurrencyFieldBridge` automatically.

Example 59. An explicit `@FieldBridge` is no longer needed

```
@Entity @Indexed
public class User {

    @Field
    public Currency getDefaultCurrency();

    // ...
}

@Entity @Indexed
public class Account {

    @Field
    public Currency getCurrency();

    // ...
}

//CurrencyFieldBridge is applied automatically everywhere Currency is found on an indexed
property
```

A few more things you need to know:

- a `BridgeProvider` must have a no-arg constructor
- if a `BridgeProvider` only returns `FieldBridge` instances if it is meaningful for the calling context. Null otherwise. In our example, the return type must be `Currency` to be meaningful to our provider.
- if two or more bridge providers return a `FieldBridge` instance for a given return type, an exception will be raised.



What is a calling context

A calling context is represented by the `BridgeProviderContext` object and represents the environment for which we are looking for a bridge. `BridgeProviderContext` gives access to the return type of the indexed property as well as the `ServiceManager` which gives access to the `ClassLoaderService` for everything class loader related.

```
ClassLoaderService classLoaderService = serviceManager
    .getClassLoaderService();
CustomBridge cb = classLoaderService.classForName(
    "com.package.CustomBridge" );
```

4.4.5. Dependency injection in bridges

Hibernate Search offers SPIs that allow dependency injection frameworks to be used when instantiating custom bridges, enabling dependency injection in bridges.

When using the Hibernate Search ORM integration, Hibernate Search is wired to Hibernate ORM's dependency injection support, meaning that if your framework already provides dependency injection to Hibernate ORM (for `EntityListener`s` and `AttributeConverter`s` in particular), then dependency injection should automatically work for Hibernate Search bridges, too.

In order to get a bridge injected, you only need to implement a bridge using all the features provided by your dependency injection framework (`@Inject`, ...) and register it to your dependency injection context. Then you can apply the bridge as usual, using the implementation class or a uniquely identifying superinterface to refer to the bridge. Hibernate Search will try to use the dependency injection framework to retrieve the bridge, and if not found it will default to reflection-based instantiation (using the public, no-argument constructor).



Be careful to define the scope of your bridge as appropriate. For instance if the bridge implements `AppliedOnTypeAwareBridge`, a singleton bridge will not work, since the bridge may be used in multiple places on different properties with different types. In that case, a "prototype" (Spring terminology) or "dependent" (CDI terminology) scope is what you want, and it's generally the safest choice.

4.5. Conditional indexing

In some situations, you want to index an entity only when it is in a given state, for example:

- only index blog entries marked as published
- no longer index invoices when they are marked archived

This serves both functional and technical needs. You don't want your blog readers to find your draft entries and filtering them off the query is a bit annoying. Very few of your entities are actually required to be indexed and you want to limit indexing overhead and keep indexes small and fast.

Hibernate Search lets you intercept entity indexing operations and override them. It is quite simple:

- Write an `EntityIndexingInterceptor` class with your entity state based logic
- Mark the entity as intercepted by this implementation

Let's look at the blog example at [Index blog entries only when they are published and remove them when they are in a different state](#)

Example 60. Index blog entries only when they are published and remove them when they are in a different state

```
/**
 * Only index blog when it is in published state
 *
 * @author Emmanuel Bernard <emmanuel@hibernate.org>
 */
public class IndexWhenPublishedInterceptor implements EntityIndexingInterceptor<Blog> {
    @Override
    public IndexingOverride onAdd(Blog entity) {
        if (entity.getStatus() == BlogStatus.PUBLISHED) {
            return IndexingOverride.APPLY_DEFAULT;
        }
        return IndexingOverride.SKIP;
    }

    @Override
    public IndexingOverride onUpdate(Blog entity) {
        if (entity.getStatus() == BlogStatus.PUBLISHED) {
            return IndexingOverride.UPDATE;
        }
        return IndexingOverride.REMOVE;
    }

    @Override
    public IndexingOverride onDelete(Blog entity) {
        return IndexingOverride.APPLY_DEFAULT;
    }

    @Override
    public IndexingOverride onCollectionUpdate(Blog entity) {
        return onUpdate(entity);
    }
}
```

```
@Entity
@Indexed(interceptor=IndexWhenPublishedInterceptor.class)
public class Blog {
    @Id
    @GeneratedValue
    public Integer getId() { return id; }
    public void setId(Integer id) { this.id = id; }
    private Integer id;

    @Field
    public String getTitle() { return title; }
    public void setTitle(String title) { this.title = title; }
    private String title;

    public BlogStatus getStatus() { return status; }
    public void setStatus(BlogStatus status) { this.status = status; }
    private BlogStatus status;

    // ...
}
```

We mark the `Blog` entity with `@Indexed.interceptor`. As you can see, `IndexWhenPublishedInterceptor` implements `EntityIndexingInterceptor` and accepts `Blog` entities (it could have accepted super classes as well - for example `Object` if you create a generic

interceptor.

You can react to several planned indexing events:

- when an entity is added to your datastore
- when an entity is updated in your datastore
- when an entity is deleted from your datastore
- when a collection own by this entity is updated in your datastore

For each occurring event you can respond with one of the following actions:

- **APPLY_DEFAULT**: that's the basic operation that lets Hibernate Search update the index as expected - creating, updating or removing the document
- **SKIP**: ask Hibernate Search to not do anything to the index for this event - data will not be created, updated or removed from the index in any way
- **REMOVE**: ask Hibernate Search to remove indexing data about this entity - you can safely ask for **REMOVE** even if the entity has not yet been indexed
- **UPDATE**: ask Hibernate Search to either index or update the index for this entity - it is safe to ask for **UPDATE** even if the entity has never been indexed



Be careful, not every combination makes sense: for example, asking to **UPDATE** the index upon onDelete. Note that you could ask for **SKIP** in this situation if saving indexing time is critical for you. That's rarely the case though.

By default, no interceptor is applied on an entity. You have to explicitly define an interceptor via the **@Indexed** annotation (see [@Indexed](#)) or programmatically (see [Programmatic API](#)). This class and all its subclasses will then be intercepted. You can stop or change the interceptor used in a subclass by overriding **@Indexed.interceptor**. Hibernate Search provides **DontInterceptEntityInterceptor** which will explicitly not intercept any call. This is useful to reset interception within a class hierarchy.



Dirty checking optimization is disabled when interceptors are used. Dirty checking optimization does check what has changed in an entity and only triggers an index update if indexed properties are changed. The reason is simple, your interceptor might depend on a non indexed property which would be ignored by this optimization.



An **EntityIndexingInterceptor** can never override an explicit indexing operation such as **index(T)**, **purge(T, id)** or **purgeAll(class)**.

4.6. Providing your own id

You can provide your own id for Hibernate Search if you are extending the internals. You will have to generate a unique value so it can be given to Lucene to be indexed. This will have to be given to Hibernate Search when you create an `org.hibernate.search.Work` object - the document id is required in the constructor.

4.6.1. The ProvidedId annotation

Unlike `@DocumentId` which is applied on field level, `@ProvidedId` is used on the class level. Optionally you can specify your own bridge implementation using the bridge property. Also, if you annotate a class with `@ProvidedId`, your subclasses will also get the annotation - but it is not done by using the `java.lang.annotations.@Inherited`. Be sure however, to *not* use this annotation with `@DocumentId` as your system will break.

Example 61. Providing your own id

```
@ProvidedId(bridge = org.my.own.package.MyCustomBridge)
@Indexed
public class MyClass{
    @Field
    String MyString;
    ...
}
```

4.7. Programmatic API

Although the recommended approach for mapping indexed entities is to use annotations, it is sometimes more convenient to use a different approach:

- the same entity is mapped differently depending on deployment needs (customization for clients)
- some automation process requires the dynamic mapping of many entities sharing common traits

While it has been a popular demand in the past, the Hibernate team never found the idea of an XML alternative to annotations appealing due to its heavy duplication, lack of code refactoring safety, because it did not cover all the use case spectrum and because we are in the 21st century :)

The idea of a programmatic API was much more appealing and has now become a reality. You can programmatically define your mapping using a programmatic API: you define entities and fields as indexable by using mapping classes which effectively mirror the annotation concepts in Hibernate Search. Note that fan(s) of XML approach can design their own schema and use the programmatic API to create the mapping while parsing the XML stream.

In order to use the programmatic model you must first construct a `SearchMapping` object which you

can do in two ways:

- directly
- via a factory

You can pass the `SearchMapping` object directly via the property key `hibernate.search.model_mapping` or the constant `Environment.MODEL_MAPPING`. Use the Configuration API or the `Map` passed to the JPA Persistence bootstrap methods.

Example 62. Programmatic mapping

```
SearchMapping mapping = new SearchMapping();
// ... configure mapping
Configuration config = new Configuration();
config.getProperties().put( Environment.MODEL_MAPPING, mapping );
SessionFactory sf = config.buildSessionFactory();
```

Example 63. Programmatic mapping with JPA

```
SearchMapping mapping = new SearchMapping();
// ... configure mapping
Map props = new HashMap();
props.put( Environment.MODEL_MAPPING, mapping );
EntityManagerFactory emf = Persistence.createEntityManagerFactory( "userPU", props );
```

Alternatively, you can create a factory class (ie hosting a method annotated with `@Factory`) whose factory method returns the `SearchMapping` object. The factory class must have a no-arg constructor and its fully qualified class name is passed to the property key `hibernate.search.model_mapping` or its type-safe representation `Environment.MODEL_MAPPING`. This approach is useful when you do not necessarily control the bootstrap process like in a Java EE, CDI or Spring Framework container.

Example 64. Use a mapping factory

```
public class MyAppSearchMappingFactory {
    @Factory
    public SearchMapping getSearchMapping() {
        SearchMapping mapping = new SearchMapping();
        mapping
            .analyzerDef( "ngram", StandardTokenizerFactory.class )
            .filter( LowerCaseFilterFactory.class )
            .filter( NGramFilterFactory.class )
            .param( "minGramSize", "3" )
            .param( "maxGramSize", "3" );
        return mapping;
    }
}
```

```
<persistence ...>
  <persistence-unit name="users">
    ...
    <properties>
      <property name="hibernate.search.model_mapping"
        value="com.acme.MyAppSearchMappingFactory"/>
    </properties>
  </persistence-unit>
</persistence>
```

The `SearchMapping` is the root object which contains all the necessary indexable entities and fields. From there, the `SearchMapping` object exposes a fluent (and thus intuitive) API to express your mappings: it contextually exposes the relevant mapping options in a type-safe way. Just let your IDE auto-completion feature guide you through.

Today, the programmatic API cannot be used on a class annotated with Hibernate Search annotations, chose one approach or the other. Also note that the same default values apply in annotations and the programmatic API. For example, the `@Field.name` is defaulted to the property name and does not have to be set.

Each core concept of the programmatic API has a corresponding example to depict how the same definition would look using annotation. Therefore seeing an annotation example of the programmatic approach should give you a clear picture of what Hibernate Search will build with the marked entities and associated properties.

4.7.1. Mapping an entity as indexable

The first concept of the programmatic API is to define an entity as indexable. Using the annotation approach a user would mark the entity as `@Indexed`, the following example demonstrates how to programmatically achieve this.

Example 65. Marking an entity indexable

```
SearchMapping mapping = new SearchMapping();

mapping.entity(Address.class)
    .indexed()
        .indexName("Address_Index") //optional
        .interceptor(IndexWhenPublishedInterceptor.class); //optional

cfg.getProperties().put("hibernate.search.model_mapping", mapping);
```

As you can see you must first create a `SearchMapping` object which is the root object that is then passed to the Configuration object as property. You must declare an entity and if you wish to make that entity as indexable then you must call the `indexed()` method. The `indexed()` method has an optional `indexName(String indexName)` which can be used to change the default index name that is created by Hibernate Search. Likewise, an `interceptor(Class<? extends EntityIndexedInterceptor>)` is available. Using the annotation model the above can be achieved as:

Example 66. Annotation example of indexing entity

```
@Entity
@Indexed(index="Address_Index", interceptor=IndexWhenPublishedInterceptor.class)
public class Address {
    // ...
}
```

4.7.2. Adding DocumentId to indexed entity

To set a property as a document id:

Example 67. Enabling document id with programmatic model

```
SearchMapping mapping = new SearchMapping();

mapping.entity(Address.class).indexed()
    .property("addressId", ElementType.FIELD) //field access
    .documentId()
        .name("id");

cfg.getProperties().put("hibernate.search.model_mapping", mapping);
```

The above is equivalent to annotating a property in the entity as `@DocumentId` as seen in the following example:

Example 68. `@DocumentId` annotation definition

```
@Entity
@Indexed
public class Address {
    @Id
    @GeneratedValue
    @DocumentId(name="id")
    private Long addressId;

    // ...
}
```

4.7.3. Defining analyzers

Analyzers can be programmatically defined using the `analyzerDef(String analyzerDef, Class<? extends TokenizerFactory> tokenizerFactory)` method. This method also enables you to define filters for the analyzer definition. Each filter that you define can optionally take in parameters as seen in the following example :

Example 69. Defining analyzers using programmatic model

```
SearchMapping mapping = new SearchMapping();

mapping
    .analyzerDef( "ngram", StandardTokenizerFactory.class )
        .filter( LowerCaseFilterFactory.class )
        .filter( NGramFilterFactory.class )
            .param( "minGramSize", "3" )
            .param( "maxGramSize", "3" )
    .analyzerDef( "en", StandardTokenizerFactory.class )
        .filter( LowerCaseFilterFactory.class )
        .filter( EnglishPorterFilterFactory.class )
    .analyzerDef( "de", StandardTokenizerFactory.class )
        .filter( LowerCaseFilterFactory.class )
        .filter( GermanStemFilterFactory.class )
    .entity(Address.class).indexed()
        .property("addressId", ElementType.METHOD) //getter access
        .documentId()
            .name("id");

cfg.getProperties().put( "hibernate.search.model_mapping", mapping );
```

The analyzer mapping defined above is equivalent to the annotation model using `@AnalyzerDef`:

Example 70. Analyzer definition using annotation

```
@Indexed
@Entity
@AnalyzerDef(name = "ngram",
    tokenizer = @TokenizerDef(factory = StandardTokenizerFactory.class),
    filters = {
        @TokenFilterDef(factory = LowerCaseFilterFactory.class),
        @TokenFilterDef(factory = NGramFilterFactory.class,
            params = {
                @Parameter(name = "minGramSize", value = "3"),
                @Parameter(name = "maxGramSize", value = "3")
            })
    })
@AnalyzerDef(name = "en",
    tokenizer = @TokenizerDef(factory = StandardTokenizerFactory.class),
    filters = {
        @TokenFilterDef(factory = LowerCaseFilterFactory.class),
        @TokenFilterDef(factory = EnglishPorterFilterFactory.class)
    })
@AnalyzerDef(name = "de",
    tokenizer = @TokenizerDef(factory = StandardTokenizerFactory.class),
    filters = {
        @TokenFilterDef(factory = LowerCaseFilterFactory.class),
        @TokenFilterDef(factory = GermanStemFilterFactory.class)
    })
public class Address {
    // ...
}
```

4.7.4. Defining normalizers

Similarly to analyzers, normalizers can be programmatically defined using the `normalizerDef(String name)` method.

Example 71. Defining normalizers using programmatic model

```
SearchMapping mapping = new SearchMapping();

mapping
    .normalizerDef( "lowercase" )
    .filter( ASCIIIFoldingFilterFactory.class )
    .filter( LowerCaseFilterFactory.class )
    .normalizerDef( "charmapping" )
    .charFilter( MappingChar.class )
    .param( "mapping", "com/example/myapplication/lucene/mapping.properties" )
    .entity(Address.class).indexed()
    .property("addressId", ElementType.METHOD) //getter access
    .documentId()
    .name("id");

cfg.getProperties().put( "hibernate.search.model_mapping", mapping );
```

4.7.5. Defining full text filter definitions

The programmatic API provides easy mechanism for defining full text filter definitions which is available via `@FullTextFilterDef` (see [Filters](#)). The next example depicts the creation of full text filter definition using the `fullTextFilterDef` method.

Example 72. Defining full text definition programmatically

```
SearchMapping mapping = new SearchMapping();

mapping
    .analyzerDef( "en", StandardTokenizerFactory.class )
        .filter( LowerCaseFilterFactory.class )
        .filter( EnglishPorterFilterFactory.class )
    .fullTextFilterDef( "security", SecurityFilterFactory.class )
        .cache( FilterCacheModeType.INSTANCE_ONLY )
    .entity( Address.class )
        .indexed()
        .property( "addressId", ElementType.METHOD )
            .documentId()
            .name( "id" )
        .property( "street1", ElementType.METHOD )
            .field()
                .analyzer( "en" )
                .store( Store.YES )
            .field()
                .name( "address_data" )
                .analyzer( "en" )
                .store( Store.NO );

cfg.getProperties().put( "hibernate.search.model_mapping", mapping );
```

The previous example can effectively be seen as annotating your entity with `@FullTextFilterDef` like below:

Example 73. Using annotation to define full text filter definition

```
@Entity
@Indexed
@AnalyzerDef(name = "en",
    tokenizer = @TokenizerDef(factory = StandardTokenizerFactory.class),
    filters = {
        @TokenFilterDef(factory = LowerCaseFilterFactory.class),
        @TokenFilterDef(factory = EnglishPorterFilterFactory.class)
    })
@FullTextFilterDef(name = "security", impl = SecurityFilterFactory.class, cache =
FilterCacheModeType.INSTANCE_ONLY)
public class Address {

    @Id
    @GeneratedValue
    @DocumentId(name="id")
    public Long getAddressId() {...};

    @Field(store=Store.YES, analyzer=@Analyzer(definition="en"))
    @Field(name="address_data", analyzer=@Analyzer(definition="en"))
    public String getAddress1() {...};

    // ...
}
```

4.7.6. Defining fields for indexing

When defining fields for indexing using the programmatic API, call `field()` on the `property(String propertyName, ElementType elementType)` method. From `field()` you can specify the `name`, `index`, `store`, `bridge`, `analyzer` and `normalizer` parameters, as well as make the field sortable using `.sortableField()`.

Example 74. Indexing fields using programmatic API

```
SearchMapping mapping = new SearchMapping();

mapping
    .analyzerDef( "en", StandardTokenizerFactory.class )
        .filter( LowerCaseFilterFactory.class )
        .filter( EnglishPorterFilterFactory.class )
    .normalizerDef( "lowercase" )
        .filter( ASCIIFoldingFilterFactory.class )
        .filter( LowerCaseFilterFactory.class )
    .entity(Address.class).indexed()
        .property("addressId", ElementType.METHOD)
            .documentId()
                .name("id")
        .property("street1", ElementType.METHOD)
            .field()
                .analyzer("en")
                .store(Store.YES)
            .field()
                .name("address_data")
                .analyzer("en")
            .field()
                .name("address_sort")
                .sortableField()
                .normalizer("lowercase");

cfg.getProperties().put( "hibernate.search.model_mapping", mapping );
```

The above example of marking fields as indexable is equivalent to defining fields using `@Field` as seen below:

```

@Entity
@Indexed
@AnalyzerDef(name = "en",
    tokenizer = @TokenizerDef(factory = StandardTokenizerFactory.class),
    filters = {
        @TokenFilterDef(factory = LowerCaseFilterFactory.class),
        @TokenFilterDef(factory = EnglishPorterFilterFactory.class)
    })
@NormalizerDef(name = "lowercase",
    filters = {
        @TokenFilterDef(factory = ASCIIFoldingFilterFactory.class),
        @TokenFilterDef(factory = LowerCaseFilterFactory.class)
    })
)
public class Address {

    @Id
    @GeneratedValue
    @DocumentId(name="id")
    private Long getAddressId() {...};

    @Field(store=Store.YES, analyzer=@Analyzer(definition="en"))
    @Field(name="address_data", analyzer=@Analyzer(definition="en"))
    @Field(name="address_sort", normalizer=@Normalizer(definition="lowercase"))
    @SortableField(forField="address_sort")
    public String getAddress1() {...}

    // ...
}

```



When using a programmatic mapping for a given type **X**, you can only refer to fields defined on **X**. Fields or methods inherited from a super type are not configurable. In case you need to configure a super class property, you need to either override the property in **X** or create a programmatic mapping for the super class. This mimics the usage of annotations where you cannot annotate a field or method of a super class either, unless it is redefined in the given type.

4.7.7. Programmatically defining embedded entities

In this section you will see how to programmatically define entities to be embedded into the indexed entity similar to using the `@IndexedEmbedded` model. In order to define this you must mark the property as `indexEmbedded`. There is the option to add a prefix to the embedded entity definition which can be done by calling `prefix` as seen in the example below:

Example 76. Programmatically defining embedded entities

```
SearchMapping mapping = new SearchMapping();

mapping
    .entity(ProductCatalog.class)
    .indexed()
    .property("catalogId", ElementType.METHOD)
    .documentId()
    .name("id")
    .property("title", ElementType.METHOD)
    .field()
    .index(Index.YES)
    .store(Store.NO)
    .property("description", ElementType.METHOD)
    .field()
    .index(Index.YES)
    .store(Store.NO)
    .property("items", ElementType.METHOD)
    .indexEmbedded()
    .prefix("catalog.items"); //optional

cfg.getProperties().put( "hibernate.search.model_mapping", mapping )
```

The next example shows the same definition using annotation (`@IndexedEmbedded`):

Example 77. Using `@IndexedEmbedded`

```
@Entity
@Indexed
public class ProductCatalog {
    @Id
    @GeneratedValue
    @DocumentId(name="id")
    public Long getCatalogId() {...}

    @Field
    public String getTitle() {...}

    @Field
    public String getDescription();

    @OneToMany(fetch = FetchType.LAZY)
    @IndexColumn(name = "list_position")
    @Cascade(org.hibernate.annotations.CascadeType.ALL)
    @IndexedEmbedded(prefix="catalog.items")
    public List<Item> getItems() {...}

    // ...
}
```

4.7.8. Contained In definition

`@ContainedIn` can be defined as seen in the example below:

Example 78. Programmatically defining `containedIn`

```
SearchMapping mapping = new SearchMapping();

mapping
    .entity(ProductCatalog.class)
        .indexed()
        .property("catalogId", ElementType.METHOD)
            .documentId()
        .property("title", ElementType.METHOD)
            .field()
        .property("description", ElementType.METHOD)
            .field()
        .property("items", ElementType.METHOD)
            .indexEmbedded()

    .entity(Item.class)
        .property("description", ElementType.METHOD)
            .field()
        .property("productCatalog", ElementType.METHOD)
            .containedIn();

cfg.getProperties().put( "hibernate.search.model_mapping", mapping );
```

This is equivalent to defining `@ContainedIn` in your entity:

Example 79. Annotation approach for `@ContainedIn`

```
@Entity
@Indexed
public class ProductCatalog {

    @Id
    @GeneratedValue
    @DocumentId
    public Long getCatalogId() {...}

    @Field
    public String getTitle() {...}

    @Field
    public String getDescription() {...}

    @OneToMany(fetch = FetchType.LAZY)
    @IndexColumn(name = "list_position")
    @Cascade(org.hibernate.annotations.CascadeType.ALL)
    @IndexedEmbedded
    private List<Item> getItems() {...}

    // ...
}
```

```
@Entity
public class Item {

    @Id
    @GeneratedValue
    private Long itemId;

    @Field
    public String getDescription() {...}

    @ManyToOne( cascade = { CascadeType.PERSIST, CascadeType.REMOVE } )
    @ContainedIn
    public ProductCatalog getProductCatalog() {...}

    // ...
}
```

4.7.9. Date/Calendar Bridge

In order to define a calendar or date bridge mapping, call the `dateBridge(Resolution resolution)` or `calendarBridge(Resolution resolution)` methods after you have defined a `field()` in the `SearchMapping` hierarchy.

Example 80. Programmatic model for defining calendar/date bridge

```
SearchMapping mapping = new SearchMapping();

mapping
    .entity(Address.class)
        .indexed()
        .property("addressId", ElementType.FIELD)
            .documentId()
        .property("street1", ElementType.FIELD())
            .field()
        .property("createdOn", ElementType.FIELD)
            .field()
            .dateBridge(Resolution.DAY)
        .property("lastUpdated", ElementType.FIELD)
            .calendarBridge(Resolution.DAY);

cfg.getProperties().put( "hibernate.search.model_mapping", mapping );
```

See below for defining the above using `@CalendarBridge` and `@DateBridge`:

Example 81. `@CalendarBridge` and `@DateBridge` definition

```
@Entity
@Indexed
public class Address {

    @Id
    @GeneratedValue
    @DocumentId
    private Long addressId;

    @Field
    private String address1;

    @Field
    @DateBridge(resolution=Resolution.DAY)
    private Date createdOn;

    @CalendarBridge(resolution=Resolution.DAY)
    private Calendar lastUpdated;

    // ...
}
```

4.7.10. Declaring bridges

It is possible to associate bridges to programmatically defined fields. When you define a `field()` programmatically you can use the `bridge(Class<?> impl)` to associate a `FieldBridge` implementation class. The bridge method also provides optional methods to include any parameters required for the bridge class. The below shows an example of programmatically defining a bridge:

Example 82. Declaring field bridges programmatically

```
SearchMapping mapping = new SearchMapping();

mapping
    .entity(Address.class)
    .indexed()
    .property("addressId", ElementType.FIELD)
    .documentId()
    .property("street1", ElementType.FIELD)
    .field()
    .field()
        .name("street1_abridged")
        .bridge( ConcatStringBridge.class )
        .param( "size", "4" );

cfg.getProperties().put( "hibernate.search.model_mapping", mapping );
```

The above can equally be defined using annotations, as seen in the next example.

Example 83. Declaring field bridges using annotation

```
@Entity
@Indexed
public class Address {

    @Id
    @GeneratedValue
    @DocumentId(name="id")
    private Long addressId;

    @Field
    @Field(name="street1_abridged",
        bridge = @FieldBridge( impl = ConcatStringBridge.class,
            params = @Parameter( name="size", value="4" ))
    private String address1;

    // ...
}
```

4.7.11. Mapping class bridge

You can define class bridges on entities programmatically. This is shown in the next example:

Example 84. Defining class bridges using API

```
SearchMapping mapping = new SearchMapping();

mapping
    .entity(Departments.class)
    .classBridge(CatDeptsFieldsClassBridge.class)
    .name("branchnetwork")
    .index(Index.YES)
    .store(Store.YES)
    .param("sepChar", " ")
    .classBridge(EquipmentType.class)
    .name("equiptype")
    .index(Index.YES)
    .store(Store.YES)
    .param("C", "Cisco")
    .param("D", "D-Link")
    .param("K", "Kingston")
    .param("3", "3Com")
    .indexed();

cfg.getProperties().put( "hibernate.search.model_mapping", mapping );
```

The above is similar to using `@ClassBridge` as seen in the next example:

Example 85. Using `@ClassBridge`

```
@Entity
@Indexed
@ClassBridge(name="branchnetwork",
    store= Store.YES,
    impl = CatDeptsFieldsClassBridge.class,
    params = @Parameter( name="sepChar", value=" " ) ),
@ClassBridge(name="equiptype",
    store= Store.YES,
    impl = EquipmentType.class,
    params = { @Parameter( name="C", value="Cisco" ),
        @Parameter( name="D", value="D-Link" ),
        @Parameter( name="K", value="Kingston" ),
        @Parameter( name="3", value="3Com" )
    })
public class Departments {
    // ...
}
```

4.7.12. Mapping dynamic boost

Index-time boosting is deprecated: support for this approach will be removed in Apache Lucene since version 7.



We recommend to use query-time boosting instead: see [Querying](#).

If query-time boosting is not enough for your use case, you can implement dynamic boosting using a [custom similarity](#) (advanced use).

You can apply a dynamic boost factor on either a field or a whole entity:

Example 86. `dynamicBoost` mapping using programmatic model

```
SearchMapping mapping = new SearchMapping();
mapping
    .entity(DynamicBoostedDescLibrary.class)
    .indexed()
    .dynamicBoost(CustomBoostStrategy.class)
    .property("libraryId", ElementType.FIELD)
    .documentId().name("id")
    .property("name", ElementType.FIELD)
    .dynamicBoost(CustomFieldBoostStrategy.class);
    .field()
    .store(Store.YES)

cfg.getProperties().put( "hibernate.search.model_mapping", mapping );
```

The next example shows the equivalent mapping using the `@DynamicBoost` annotation:

Example 87. Using the `@DynamicBoost`

```
@Entity
@Indexed
@DynamicBoost(impl = CustomBoostStrategy.class)
public class DynamicBoostedDescriptionLibrary {

    @Id
    @GeneratedValue
    @DocumentId
    private int id;

    private float dynScore;

    @Field(store = Store.YES)
    @DynamicBoost(impl = CustomFieldBoostStrategy.class)
    private String name;

    public DynamicBoostedDescriptionLibrary() {
        dynScore = 1.0f;
    }

    // ...
}
```

Chapter 5. Querying

The second most important capability of Hibernate Search is the ability to execute Lucene queries and retrieve entities managed by a Hibernate session. The search provides the power of Lucene without leaving the Hibernate paradigm, giving another dimension to the Hibernate classic search mechanisms (HQL, Criteria query, native SQL query).

Preparing and executing a query consists of four simple steps:

- Creating a `FullTextSession`
- Creating a Lucene query either via the Hibernate Search query DSL (recommended) or by utilizing the Lucene query API
- Wrapping the Lucene query using an `org.hibernate.Query`
- Executing the search by calling for example `list()` or `scroll()`

To access the querying facilities, you have to use a `FullTextSession`. This Search specific session wraps a regular `org.hibernate.Session` in order to provide query and indexing capabilities.

Example 88. Creating a FullTextSession

```
Session session = sessionFactory.openSession();
//...
FullTextSession fullTextSession = Search.getFullTextSession(session);
```

Once you have a `FullTextSession` you have two options to build the full-text query: the Hibernate Search query DSL or the native Lucene query.

If you use the Hibernate Search query DSL, it will look like this:

```
QueryBuilder b = fullTextSession.getSearchFactory()
    .buildQueryBuilder().forEntity(Myth.class).get();

org.apache.lucene.search.Query luceneQuery =
    b.keyword()
        .onField("history").boostedTo(3)
        .matching("storm")
        .createQuery();

org.hibernate.Query fullTextQuery = fullTextSession.createFullTextQuery(luceneQuery);
List result = fullTextQuery.list(); //return a list of managed objects
```

You can alternatively write your Lucene query either using the Lucene query parser or Lucene programmatic API.

Example 89. Creating a Lucene query via the QueryParser

```
SearchFactory searchFactory = fullTextSession.getSearchFactory();
org.apache.lucene.queryparser.classic.QueryParser parser =
    new QueryParser("title", searchFactory.getAnalyzer(Myth.class));
try {
    org.apache.lucene.search.Query luceneQuery = parser.parse("history:storm^3");
}
catch (ParseException e) {
    //handle parsing failure
}

org.hibernate.Query fullTextQuery = fullTextSession.createFullTextQuery(luceneQuery);
List result = fullTextQuery.list(); //return a list of managed objects
```



The Hibernate query built on top of the Lucene query is a regular `org.hibernate.Query`, which means you are in the same paradigm as the other Hibernate query facilities (HQL, Native or Criteria). The regular `list()` , `uniqueResult()`, `iterate()` and `scroll()` methods can be used.

In case you are using the Java Persistence APIs of Hibernate, the same extensions exist:

Example 90. Creating a Search query using the JPA API

```
EntityManager em = entityManagerFactory.createEntityManager();

FullTextEntityManager fullTextEntityManager =
    org.hibernate.search.jpa.Search.getFullTextEntityManager(em);

// ...
QueryBuilder b = fullTextEntityManager.getSearchFactory()
    .buildQueryBuilder().forEntity( Myth.class ).get();

org.apache.lucene.search.Query luceneQuery =
    b.keyword()
        .onField("history").boostedTo(3)
        .matching("storm")
        .createQuery();
javax.persistence.Query fullTextQuery =
    fullTextEntityManager.createFullTextQuery( luceneQuery );

List result = fullTextQuery.getResultList(); //return a list of managed objects
```



The following examples we will use the Hibernate APIs but the same example can be easily rewritten with the Java Persistence API by just adjusting the way the `FullTextQuery` is retrieved.

5.1. Building queries

Hibernate Search queries are built on top of Lucene queries which gives you total freedom on the type of Lucene query you want to execute. However, once built, Hibernate Search wraps further query

processing using `org.hibernate.Query` as your primary query manipulation API.

5.1.1. Building a Lucene query using the Lucene API

Using the Lucene API, you have several options. You can use the query parser (fine for simple queries) or the Lucene programmatic API (for more complex use cases). It is out of the scope of this documentation on how to exactly build a Lucene query. Please refer to the online Lucene documentation or get hold of a copy of *Lucene In Action* or *Hibernate Search in Action*.

5.1.2. Building a Lucene query with the Hibernate Search query DSL

Writing full-text queries with the Lucene programmatic API is quite complex. It's even more complex to understand the code once written. Besides the inherent API complexity, you have to remember to convert your parameters to their string equivalent as well as make sure to apply the correct analyzer to the right field (a ngram analyzer will for example use several ngrams as the tokens for a given word and should be searched as such).

The Hibernate Search query DSL makes use of a style of API called a fluent API. This API has a few key characteristics:

- it has meaningful method names making a succession of operations reads almost like English
- it limits the options offered to what makes sense in a given context (thanks to strong typing and IDE auto-completion).
- it often uses the chaining method pattern
- it's easy to use and even easier to read

Let's see how to use the API. You first need to create a query builder that is attached to a given indexed entity type. This `QueryBuilder` will know what analyzer to use and what field bridge to apply. You can create several `QueryBuilder` instances (one for each entity type involved in the root of your query). You get the `QueryBuilder` from the `SearchFactory`.

```
QueryBuilder mythQB = searchFactory.buildQueryBuilder().forEntity( Myth.class ).get();
```

You can also override the analyzer used for a given field or fields. This is rarely needed and should be avoided unless you know what you are doing.

```
QueryBuilder mythQB = searchFactory.buildQueryBuilder()  
    .forEntity( Myth.class )  
    .overridesForField("history", "stem_analyzer_definition")  
    .get();
```

Using the query builder, you can then build queries. It is important to realize that the end result of a `QueryBuilder` is a Lucene query. For this reason you can easily mix and match queries generated via Lucene's query parser or Query objects you have assembled with the Lucene programmatic API and use them with the Hibernate Search DSL. Just in case the DSL is missing some features.

Keyword queries

Let's start with the most basic use case - searching for a specific word:

```
Query luceneQuery = mythQB.keyword().onField("history").matching("storm").createQuery();
```

`keyword()` means that you are trying to find a specific word. `onField()` specifies in which Lucene field to look. `matching()` tells what to look for. And finally `createQuery()` creates the Lucene query object. A lot is going on with this line of code.

- The value `storm` is passed through the `history` `FieldBridge`: it does not matter here but you will see that it's quite handy when dealing with numbers or dates.
- The field bridge value is then passed to the analyzer used to index the field `history`. This ensures that the query uses the same term transformation than the indexing (lower case, n-gram, stemming and so on). If the analyzing process generates several terms for a given word, a boolean query is used with the `SHOULD` logic (roughly an `OR` logic).

We make the example a little more advanced now and have a look at how to search a field that uses ngram analyzers. ngram analyzers index succession of ngrams of your words which helps to recover from user typos. For example the 3-grams of the word `hibernate` are `hib`, `ibe`, `ber`, `rna`, `nat`, `ate`.

```

@AnalyzerDef(name = "ngram",
    tokenizer = @TokenizerDef(factory = StandardTokenizerFactory.class ),
    filters = {
        @TokenFilterDef(factory = StandardFilterFactory.class),
        @TokenFilterDef(factory = LowerCaseFilterFactory.class),
        @TokenFilterDef(factory = StopFilterFactory.class),
        @TokenFilterDef(factory = NGramFilterFactory.class,
            params = {
                @Parameter(name = "minGramSize", value = "3"),
                @Parameter(name = "maxGramSize", value = "3") } )
    }
)
@Entity
@Indexed
public class Myth {
    @Field(analyzer=@Analyzer(definition="ngram"))
    public String getName() { return name; }
    public String setName(String name) { this.name = name; }
    private String name;
    ...
}

Query luceneQuery = mythQb.keyword().onField("name").matching("Sisiphus")
    .createQuery();

```

The matching word "Sisiphus" will be lower-cased and then split into 3-grams: sis, isi, sip, phu, hus. Each of these n-gram will be part of the query. We will then be able to find the Sysiphus myth (with a y). All that is transparently done for you.



If for some reason you do not want a specific field to use the field bridge or the analyzer you can call the `ignoreAnalyzer()` or `ignoreFieldBridge()` functions.

To search for multiple possible words in the same field, simply add them all in the matching clause.

```

//search document with storm or lightning in their history
Query luceneQuery =
    mythQB.keyword().onField("history").matching("storm lightning").createQuery();

```

To search the same word on multiple fields, use the `onFields` method.

```

Query luceneQuery = mythQB
    .keyword()
    .onFields("history", "description", "name")
    .matching("storm")
    .createQuery();

```

Sometimes, one field should be treated differently from another field even if searching the same term,

you can use the `andField()` method for that.

```
Query luceneQuery = mythQB.keyword()
    .onField("history")
    .andField("name")
        .boostedTo(5f)
    .andField("description")
    .matching("storm")
    .createQuery();
```

In the previous example, only field name is boosted to 5.

Simple Query String queries

Simple Query String queries use the Lucene `SimpleQueryParser`, presented as a "parser for human-entered queries". The point of this parser is to be a very simple lenient state machine to parse queries entered by your end users.

The parser is capable to transform `keyword "some phrase" -keywordidontwant fuzzy~ prefix*` into a Lucene query, giving your users a little more power (phrase queries, fuzzy queries, boolean operators...).

The lenient part is important as it will try to build the best possible query without throwing a parsing exception, even if the query is not what we would consider syntactically correct.

Another nice feature is that it allows to search on multiple fields. You basically end up establishing the following contract with Lucene:

- users will enter a search query (more or less syntactically correct)
- it will search on the fields you have specified (and you can also specify a specific boost for each field)
- you can enable each of the features that you want to expose to the users (i.e. you can enable the phrase queries but not the boolean operators)
- building the query won't throw an exception



Even if we are talking about Lucene queries here, this feature is also supported by our Elasticsearch integration.

Here are a couple of examples introducing the features of our Simple Query String support one at a time.

The query below is the most simple form of Simple Query String query. It simply returns documents whose `history` field contains `storm`:

```
Query luceneQuery = mythQB
    .simpleQueryString()
    .onField("history")
    .matching("storm")
    .createQuery();
```

By default, a Simple Query String query uses the OR operator if the operator is not explicitly defined. If you prefer using the AND operator as default, you can add the `withAndAsDefaultOperator` flag. The following query returns results containing `storm` AND `tree`:

```
Query luceneQuery = mythQB
    .simpleQueryString()
    .onField("history")
    .withAndAsDefaultOperator()
    .matching("storm tree")
    .createQuery();
```

It is also possible to search on several fields at the same time, meaning you can search for the presence of keywords in e.g. `name`, `history` or `description`, potentially with different boosts:

```
Query luceneQuery = mythQB
    .simpleQueryString()
    .onField("name")
        .boostedTo(5f)
    .andFields("history", "description")
        .boostedTo(2f)
    .withAndAsDefaultOperator()
    .matching("storm tree")
    .createQuery();
```

The above query returns results containing `storm` and `tree` either in `name`, `history` or `description`.

The syntax supported by the `SimpleQueryParser` includes the following operators:

- boolean (AND using `+`, OR using `|`, NOT using `-`)
- prefix (`prefix*`)
- phrase ("`some phrase`")
- precedence (using parentheses)
- fuzzy (`fuzzy~2`)
- near operator for phrase queries ("`war peace`"~3).

The query below will also return results containing `storms` or `sturm`.

```
Query luceneQuery = mythQB
    .simpleQueryString()
    .onFields("name", "history", "description")
    .matching("storm~2")
    .createQuery();
```

The query below will return results containing **war** and either **peace** or **harmony**:

```
Query luceneQuery = mythQB
    .simpleQueryString()
    .onFields("name", "history", "description")
    .matching("war + (peace | harmony)")
    .createQuery();
```

And, finally, the following query will return results containing exactly the phrase **storm and tree** in either **name**, **history** or **description**:

```
Query luceneQuery = mythQB
    .simpleQueryString()
    .onFields("name", "history", "description")
    .matching("\"storm and tree\"")
    .createQuery();
```

Fuzzy queries

To execute a fuzzy query (based on the Levenshtein distance algorithm), start like a **keyword** query and add the fuzzy flag.

```
Query luceneQuery = mythQB
    .keyword()
    .fuzzy()
        .withThreshold(.8f)
        .withPrefixLength(1)
    .onField("history")
    .matching("starm")
    .createQuery();
```

threshold is the limit above which two terms are considering matching. It's a decimal between 0 and 1 and defaults to 0.5. **prefixLength** is the length of the prefix ignored by the "fuzzyness": while it defaults to 0, a non zero value is recommended for indexes containing a huge amount of distinct terms.

Wildcard queries

You can also execute wildcard queries (queries where some of parts of the word are unknown). The character `?` represents a single character and `*` represents any character sequence. Note that for performance purposes, it is recommended that the query does not start with either `?` or `*`.

```
Query luceneQuery = mythQB
    .keyword()
    .wildcard()
    .onField("history")
    .matching("sto*")
    .createQuery();
```



Wildcard queries do not apply the analyzer on the matching terms. Otherwise the risk of `*` or `?` being mangled is too high.

Phrase queries

So far we have been looking for words or sets of words, you can also search exact or approximate sentences. Use `phrase()` to do so.

```
Query luceneQuery = mythQB
    .phrase()
    .onField("history")
    .sentence("Thou shalt not kill")
    .createQuery();
```

You can search approximate sentences by adding a slop factor. The slop factor represents the number of other words permitted in the sentence: this works like a within or near operator

```
Query luceneQuery = mythQB
    .phrase()
    .withSlop(3)
    .onField("history")
    .sentence("Thou kill")
    .createQuery();
```

Range queries

After looking at all these query examples for searching for to a given word, it is time to introduce range queries (on numbers, dates, strings etc). A range query searches for a value in between given boundaries (included or not) or for a value below or above a given boundary (included or not).

```
//look for 0 <= starred < 3
Query luceneQuery = mythQB
    .range()
    .onField("starred")
    .from(0).to(3).excludeLimit()
    .createQuery();

//look for myths strictly BC
Date beforeChrist = ...;
Query luceneQuery = mythQB
    .range()
    .onField("creationDate")
    .below(beforeChrist).excludeLimit()
    .createQuery();
```

Spatial (or geolocation) queries

This set of queries has its own chapter, check out [Spatial](#).

More Like This queries



Support for more-like-this queries will be removed in Hibernate Search 6.

If you need more-like-this predicates in Hibernate Search 6, feel free to drop a comment to explain your use case on [HSEARCH-3272](#).

Have you ever looked at an article or document and thought: "I want to find more like this"? Have you ever appreciated an e-commerce website that gives you similar articles to the one you are exploring?

More Like This queries are achieving just that. You feed it an entity (or its identifier) and Hibernate Search returns the list of entities that are similar.



How does it work?

For each (selected) field of the targeted entity, we look at the most meaningful terms. Then we create a query matching the most meaningful terms per field. This is a slight variation compared to the original Lucene `MoreLikeThisQuery` implementation.

The query DSL API should be self explaining. Let's look at some usage examples.

```

QueryBuilder qb = fullTextSession.getSearchFactory()
    .buildQueryBuilder()
    .forEntity( Coffee.class )
    .get();

Query mltQuery = qb
    .moreLikeThis()
    .comparingAllFields()
    .toEntityWithId( coffeeId )
    .createQuery();
List<Object[]> results = (List<Object[]>) fullTextSession
    .createFullTextQuery( mltQuery, Coffee.class )
    .setProjection( ProjectionConstants.THIS, ProjectionConstants.SCORE )
    .list();

```

This first example takes the id of an `Coffee` entity and finds the matching coffees across all fields. To be fair, this is not across *all* fields. To be included in the More Like This query, fields need to store term vectors or the actual field value. Id fields (of the root entity as well as embedded entities) and numeric fields are excluded. The latter exclusion might change in future versions.

Looking at the `Coffee` class, the following fields are considered: `name` as it is stored, `description` as it stores the term vector. `id` and `internalDescription` are excluded.

```

@Entity @Indexed
public class Coffee {

    @Id @GeneratedValue
    public Integer getId() { return id; }

    @Field(termVector = TermVector.NO, store = Store.YES)
    public String getName() { return name; }

    @Field(termVector = TermVector.YES)
    public String getSummary() { return summary; }

    @Column(length = 2000)
    @Field(termVector = TermVector.YES)
    public String getDescription() { return description; }

    public int getIntensity() { return intensity; }

    // Not stored nor term vector, i.e. cannot be used for More Like This
    @Field
    public String getInternalDescription() { return internalDescription; }

    // ...
}

```

In the example above we used projection to retrieve the relative score of each element. We might use the score to only display the results for which the score is high enough.



For best performance and best results, store the term vectors for the fields you want to include in a More Like This query.

Often, you are only interested in a few key fields to find similar entities. Plus some fields are more important than others and should be boosted.

```
Query mltQuery = qb
    .moreLikeThis()
    .comparingField("summary").boostedTo(10f)
    .andField("description")
    .toEntityWithId( coffeeId )
    .createQuery();
```

In this example, we look for similar entities by summary and description. But similar summaries are more important than similar descriptions. This is a critical tool to make More Like This meaningful for your data set.

Instead of providing the entity id, you can pass the full entity object. If the entity contains the identifier, we will use it to find the term vectors or field values. This means that we will compare the entity state as stored in the Lucene index. If the identifier cannot be retrieved (for example if the entity has not been persisted yet), we will look at each of the entity properties to find the most meaningful terms. The latter is slower and won't give the best results - avoid it if possible.

Here is how you pass the entity instance you want to compare with:

```
Coffee coffee = ...; //managed entity from somewhere

Query mltQuery = qb
    .moreLikeThis()
    .comparingField("summary").boostedTo(10f)
    .andField("description")
    .toEntity( coffee )
    .createQuery();
```



By default, the results contain at the top the entity you are comparing with. This is particularly useful to compare relative scores. If you don't need it, you can exclude it.

```
Query mltQuery = qb
    .moreLikeThis()
    .excludeEntityUsedForComparison()
    .comparingField("summary").boostedTo(10f)
    .andField("description")
    .toEntity( coffee )
    .createQuery();
```

You can ask Hibernate Search to give a higher score to the very similar entities and downgrade the score of mildly similar entities. We do that by boosting each meaningful terms by their individual

overall score. Start with a boost factor of **1** and adjust from there.

```
Query mltQuery = qb
    .moreLikeThis()
    .favorSignificantTermsWithFactor(1f)
    .comparingField("summary").boostedTo(10f)
    .andField("description")
    .toEntity( coffee )
    .createQuery();
```

Remember, more like this is a very subjective meaning and will vary depending on your data and the rules of your domain. With the various options offered, Hibernate Search arms you with the tools to adjust this weapon. Make sure to continuously test the results against your data set.

Combining queries

You can combine queries to create more complex queries. The following aggregation operators are available:

- **SHOULD**: the query should contain the matching elements of the subquery
- **MUST**: the query must contain the matching elements of the subquery
- **MUST NOT**: the query must not contain the matching elements of the subquery

These aggregations have a similar effect as the classic boolean operators **AND**, **OR** and **NOT**, but have different names to emphasise that they will have an impact on scoring.

For example the **SHOULD** operator between two queries will have an effect similar to the boolean **OR**: if either of the two combined queries matches the entry, the entry will be included in the match; though the entries which match both queries will have an higher score than those which only match one of them.

The sub-queries can be any Lucene query including a boolean query itself.

*Example 91. Structure of a boolean **AND** query: the **must** method.*

```
Query combinedQuery = querybuilder
    .bool()
    .must( queryA )
    .must( queryB )
    .createQuery();
```

Example 92. Structure of boolean **OR** query: the **should** method.

```
Query combinedQuery = querybuilder
    .bool()
        .should( queryA )
        .should( queryB )
    .createQuery();
```

Example 93. Structure of a negation query: apply a **not** modifier to a **must**.

```
Query combinedQuery = querybuilder
    .bool()
        .must( queryA )
        .must( queryB ).not()
    .createQuery();
```

Let's look at a few more practical examples; note how the querybuilder usage can be nested and how 'should', 'must', and 'not' can be combined in many ways:

Example 94. Full example of combining fulltext queries

```
//look for popular modern myths that are not urban
Date twentiethCentury = ...;
Query luceneQuery = mythQB
    .bool()
        .must( mythQB.keyword().onField("description").matching("urban").createQuery() )
        .not()
        .must( mythQB.range().onField("starred").above(4).createQuery() )
        .must( mythQB
            .range()
            .onField("creationDate")
            .above(twentiethCentury)
            .createQuery() )
    .createQuery();

//look for popular myths that are preferably urban
Query luceneQuery = mythQB
    .bool()
        .should( mythQB.keyword().onField("description").matching("urban").createQuery() )
        .must( mythQB.range().onField("starred").above(4).createQuery() )
    .createQuery();

//look for all myths except religious ones
Query luceneQuery = mythQB
    .all()
        .except( mythQB
            .keyword()
            .onField( "description_stem" )
            .matching( "religion" )
            .createQuery()
        )
    .createQuery();
```

Query options

We already have seen several query options in the previous example, but let's summarize again the options for query types and fields:

- **boostedTo** (on query type and on field): boost the whole query or the specific field to a given factor
- **withConstantScore** (on query): all results matching the query have a constant score equals to the boost
- **filteredBy(Filter)** (on query): filter query results using the Filter instance
- **ignoreAnalyzer** (on field): ignore the [analyzer](#) (or [normalizer](#)) when processing this field
- **ignoreFieldBridge** (on field): ignore field bridge when processing this field

Let's check out an example using some of these options

```
Query luceneQuery = mythQB
    .bool()
    .should( mythQB.keyword().onField("description").matching("urban").createQuery() )
    .should( mythQB
        .keyword()
        .onField("name")
        .boostedTo(3)
        .ignoreAnalyzer()
        .matching("urban").createQuery() )
    .must( mythQB
        .range()
        .boostedTo(5).withConstantScore()
        .onField("starred").above(4).createQuery() )
    .createQuery();
```

As you can see, the Hibernate Search query DSL is an easy to use and easy to read query API and by accepting and producing Lucene queries, you can easily incorporate query types not (yet) supported by the DSL. Please give us feedback!

5.1.3. Building a Hibernate Search query

So far we only covered the process of how to create your Lucene query (see [Building queries](#)). However, this is only the first step in the chain of actions. Let's now see how to build the Hibernate Search query from the Lucene query.

Generality

Once the Lucene query is built, it needs to be wrapped into an Hibernate Query. If not specified otherwise, the query will be executed against all indexed entities, potentially returning all types of indexed classes.

Example 95. Wrapping a Lucene query into a Hibernate Query

```
FullTextSession fullTextSession = Search.getFullTextSession( session );
org.hibernate.Query fullTextQuery = fullTextSession.createFullTextQuery( luceneQuery );
```

It is advised, from a performance point of view, to restrict the returned types:

Example 96. Filtering the search result by entity type

```
fullTextQuery = fullTextSession
    .createFullTextQuery(luceneQuery, Customer.class);

// or

fullTextQuery = fullTextSession
    .createFullTextQuery(luceneQuery, Item.class, Actor.class);
```

In [Filtering the search result by entity type](#) the first example returns only matching **Customer** instances, the second returns matching **Actor** and **Item** instances. The type restriction is fully polymorphic which means that if there are two indexed subclasses **Salesman** and **Customer** of the baseclass **Person**, it is possible to just specify **Person.class** in order to filter on result types.

Pagination

Out of performance reasons it is recommended to restrict the number of returned objects per query. In fact is a very common use case anyway that the user navigates from one page to an other. The way to define pagination is exactly the way you would define pagination in a plain HQL or Criteria query.

Example 97. Defining pagination for a search query

```
org.hibernate.Query fullTextQuery =
    fullTextSession.createFullTextQuery(luceneQuery, Customer.class);
fullTextQuery.setFirstResult(15); //start from the 15th element
fullTextQuery.setMaxResults(10); //return 10 elements
```



It is still possible to get the total number of matching elements regardless of the pagination via `fullTextQuery.getResultSize()`

Sorting

Apache Lucene provides a very flexible and powerful way to sort results. While the default sorting (by relevance) is appropriate most of the time, it can be interesting to sort by one or several other properties.

In order to do so, you will have to build a Lucene **Sort** object, which can be done by using either the

Hibernate Search sort DSL or the native Lucene `SortField` class.

If you use the Hibernate Search sort DSL, it will look like this:

Example 98. Sorting the results with a DSL-built Lucene `Sort`

```
QueryBuilder builder = fullTextSession.getSearchFactory()
    .buildQueryBuilder().forEntity(Book.class).get();
Query luceneQuery = /* ... */;
FullTextQuery query = s.createFullTextQuery( luceneQuery, Book.class );
Sort sort = builder
    .sort()
    .byField("author").desc() // Descending order
    .andByField("title") // Default order (ascending)
    .createSort();
query.setSort(sort);
List results = query.list();
```



If you use the Sort DSL (like in the example above) to target fields indexed through **custom `FieldBridges`**, then those field bridges must implement **`MetadataProvidingFieldBridge`**, so as to define the type of those fields and to mention that they are sortable.

Alternatively, you may build your sort using the Lucene `SortField` class directly. Then you will have to always specify the sort field type manually.

Example 99. Sorting the results with a custom-built Lucene `Sort`

```
FullTextQuery query = s.createFullTextQuery( luceneQuery, Book.class );
org.apache.lucene.search.Sort sort = new Sort(
    new SortField("author", SortField.Type.STRING),
    new SortField("title", SortField.Type.STRING));
query.setSort(sort);
List results = query.list();
```



Be aware that fields used for sorting must not be tokenized (see `@Field`). Also they should be marked as sortable field using the `@SortableField` annotation (see `@SortableField`).

Handling missing values

By default, Hibernate Search will expect every document to contain exactly one value for the given sort fields. If it's not the case, you may decide whether documents with missing values will end up in first or last position in the result list:

Example 100. Returning documents with missing values in first position when sorting

```
QueryBuilder builder = fullTextSession.getSearchFactory()
    .buildQueryBuilder().forEntity(Book.class).get();
Query luceneQuery = /* ... */;
FullTextQuery query = s.createFullTextQuery( luceneQuery, Book.class );
Sort sort = builder
    .sort()
    .byField("author")
    .onMissingValue().sortFirst()
    .andByField("title")
    .createSort();
query.setSort(sort);
List results = query.list();
```

For numeric fields (and only numeric fields), you may even provide a default value to be used in case it's missing:

Example 101. Using a default value on a sort field

```
QueryBuilder builder = fullTextSession.getSearchFactory()
    .buildQueryBuilder().forEntity(Book.class).get();
Query luceneQuery = /* ... */;
FullTextQuery query = s.createFullTextQuery( luceneQuery, Book.class );
Sort sort = builder
    .sort()
    .byField("authorId")
    .onMissingValue().use(42L) // Assume the ID is 42 when it's missing
    .andByField("title")
    .createSort();
query.setSort(sort);
List results = query.list();
```



The `use(Object)` method expects a value matching the actual type of the field. Field bridges, if any, are **not** applied. Thus, in the example above, providing an `Author` instead of the identifier directly would not work.

Using native sorts within the sort DSL

If you happen to find a sorting feature that is supported by Lucene, but not yet implemented in Hibernate Search sort DSL. If so, you may still use the DSL to mix a native `SortField` with DSL-defined sorts:

Example 102. Mixing DSL-defined sorts with native Lucene `SortFields`

```
QueryBuilder builder = fullTextSession.getSearchFactory()
    .buildQueryBuilder().forEntity(Book.class).get();
Query luceneQuery = /* ... */;
FullTextQuery query = s.createFullTextQuery( luceneQuery, Book.class );
SortField complexSortField = /* ... */;
Sort sort = builder
    .sort()
    .byField("author")
    .andByNative(complexSortField)
    .createSort();
query.setSort(sort);
List results = query.list();
```

Fetching strategy

When you restrict the return types to one class, Hibernate Search loads the objects using a single query. It also respects the static fetching strategy defined in your domain model.

It is often useful, however, to refine the fetching strategy for a specific use case.

Example 103. Specifying `FetchMode` on a query

```
Criteria criteria =
    s.createCriteria(Book.class).setFetchMode("authors", FetchMode.JOIN);
s.createFullTextQuery(luceneQuery).setCriteriaQuery(criteria);
```

In this example, the query will return all Books matching the `luceneQuery`. The authors collection will be loaded from the same query using an SQL outer join.

When defining a criteria query, it is not necessary to restrict the returned entity types when creating the Hibernate Search query from the full text session: the type is guessed from the criteria query itself.



Only fetch mode can be adjusted, refrain from applying any other restriction. While it is known to work as of Hibernate Search 4, using restriction (ie a where clause) on your Criteria query should be avoided when possible. `getResultSize()` will throw a `SearchException` if used in conjunction with a Criteria with restriction.



You cannot use `setCriteriaQuery` if more than one entity type is expected to be returned.

Projection

For some use cases, returning the domain object (including its associations) is overkill. Only a small

subset of the properties is necessary. Hibernate Search allows you to return a subset of properties:

Example 104. Using projection instead of returning the full domain object

```
org.hibernate.search.FullTextQuery query =
    s.createFullTextQuery(luceneQuery, Book.class);
query.setProjection("id", "summary", "body", "mainAuthor.name");
List results = query.list();
Object[] firstResult = (Object[]) results.get(0);
Integer id = firstResult[0];
String summary = firstResult[1];
String body = firstResult[2];
String authorName = firstResult[3];
```

Hibernate Search extracts the properties from the Lucene index and convert them back to their object representation, returning a list of `Object[]`. Projections avoid a potential database round trip (useful if the query response time is critical). However, it also has several constraints:

- the properties projected must be stored in the index (`@Field(store=Store.YES)`), which increases the index size
- the properties projected must use a `FieldBridge` implementing `org.hibernate.search.bridge.TwoWayFieldBridge` or `org.hibernate.search.bridge.TwoWayStringBridge`, the latter being the simpler version.
- you can only project simple properties of the indexed entity or its embedded associations. This means you cannot project a whole embedded entity.
- projection does not work on collections or maps which are indexed via `@IndexedEmbedded`

Projection is also useful for another kind of use case. Lucene can provide metadata information about the results. By using some special projection constants, the projection mechanism can retrieve this metadata:

Example 105. Using projection in order to retrieve meta data

```
org.hibernate.search.FullTextQuery query =
    s.createFullTextQuery(luceneQuery, Book.class);
query.setProjection(
    FullTextQuery.SCORE,
    FullTextQuery.THIS,
    "mainAuthor.name" );
List results = query.list();
Object[] firstResult = (Object[]) results.get(0);
float score = firstResult[0];
Book book = firstResult[1];
String authorName = firstResult[2];
```

You can mix and match regular fields and projection constants. Here is the list of the available constants:

- **FullTextQuery.THIS**: returns the initialized and managed entity (as a non projected query would have done).
- **FullTextQuery.DOCUMENT**: returns the Lucene Document related to the object projected.
- **FullTextQuery.OBJECT_CLASS**: returns the class of the indexed entity.
- **FullTextQuery.SCORE**: returns the document score in the query. Scores are handy to compare one result against an other for a given query but are useless when comparing the result of different queries.
- **FullTextQuery.ID**: the id property value of the projected object.
- **FullTextQuery.DOCUMENT_ID**: the Lucene document id. Careful, Lucene document id can change overtime between two different IndexReader opening.
- **FullTextQuery.EXPLANATION**: returns the Lucene Explanation object for the matching object/document in the given query. Do not use if you retrieve a lot of data. Running explanation typically is as costly as running the whole Lucene query per matching element. Make sure you use projection!

Customizing object initialization strategies

By default, Hibernate Search uses the most appropriate strategy to initialize entities matching your full text query. It executes one (or several) queries to retrieve the required entities. This is the best approach to minimize database round trips in a scenario where none / few of the retrieved entities are present in the persistence context (ie the session) or the second level cache.

If most of your entities are present in the second level cache, you can force Hibernate Search to look into the cache before retrieving an object from the database.

Example 106. Check the second-level cache before using a query

```
FullTextQuery query = session.createFullTextQuery(luceneQuery, User.class);
query.initializeObjectWith(
    ObjectLookupMethod.SECOND_LEVEL_CACHE,
    DatabaseRetrievalMethod.QUERY
);
```

ObjectLookupMethod defines the strategy used to check if an object is easily accessible (without database round trip). Other options are:

- **ObjectLookupMethod.PERSISTENCE_CONTEXT**: useful if most of the matching entities are already in the persistence context (ie loaded in the Session or EntityManager)
- **ObjectLookupMethod.SECOND_LEVEL_CACHE**: check first the persistence context and then the second-level cache.

Note that to search in the second-level cache, several settings must be in place:



- the second level cache must be properly configured and active
- the entity must have enabled second-level cache (eg via `@Cacheable`)
- the `Session`, `EntityManager` or `Query` must allow access to the second-level cache for read access (ie `CacheMode.NORMAL` in Hibernate native APIs or `CacheRetrieveMode.USE` in JPA 2 APIs).



Avoid using `ObjectLookupMethod.SECOND_LEVEL_CACHE` unless your second level cache implementation is either EHCACHE or Infinispan; other second level cache providers don't currently implement this operation efficiently.

You can also customize how objects are loaded from the database (if not found before). Use `DatabaseRetrievalMethod` for that:

- **QUERY** (default): use a (set of) queries to load several objects in batch. This is usually the best approach.
- **FIND_BY_ID**: load objects one by one using the `Session.get` or `EntityManager.find` semantic. This might be useful if batch-size is set on the entity (in which case, entities will be loaded in batch by Hibernate Core). **QUERY** should be preferred almost all the time.

The defaults for both methods, the object lookup as well as the database retrieval can also be configured via configuration properties. This way you don't have to specify your preferred methods on each query creation. The property names are `hibernate.search.query.object_lookup_method` and `hibernate.search.query.database_retrieval_method` respectively. As value use the name of the method (upper- or lowercase). For example:

Example 107. Setting object lookup and database retrieval methods via configuration properties

```
hibernate.search.query.object_lookup_method = second_level_cache
hibernate.search.query.database_retrieval_method = query
```

Limiting the time of a query

You can limit the time a query takes in Hibernate Search in two ways:

- raise an exception when the limit is reached
- limit to the number of results retrieved when the time limit is raised

Raise an exception on time limit

You can decide to stop a query if when it takes more than a predefined amount of time. Note that this

is a best effort basis but if Hibernate Search still has significant work to do and if we are beyond the time limit, a `QueryTimeoutException` will be raised (`org.hibernate.QueryTimeoutException` or `javax.persistence.QueryTimeoutException` depending on your programmatic API).

To define the limit when using the native Hibernate APIs, use one of the following approaches

Example 108. Defining a timeout in query execution

```
Query luceneQuery = ...;
FullTextQuery query = fullTextSession.createFullTextQuery(luceneQuery, User.class);

//define the timeout in seconds
query.setTimeout(5);

//alternatively, define the timeout in any given time unit
query.setTimeout(450, TimeUnit.MILLISECONDS);

try {
    query.list();
}
catch (org.hibernate.QueryTimeoutException e) {
    //do something, too slow
}
```

Likewise `getResultSize()`, `iterate()` and `scroll()` honor the timeout but only until the end of the method call. That simply means that the methods of `Iterable` or the `ScrollableResults` ignore the timeout.



`explain()` does not honor the timeout: this method is used for debug purposes and in particular to find out why a query is slow

When using JPA, simply use the standard way of limiting query execution time.

Example 109. Defining a timeout in query execution

```
Query luceneQuery = ...;
FullTextQuery query = fullTextEM.createFullTextQuery(luceneQuery, User.class);

//define the timeout in milliseconds
query.setHint( "javax.persistence.query.timeout", 450 );

try {
    query.getResultList();
}
catch (javax.persistence.QueryTimeoutException e) {
    //do something, too slow
}
```



Remember, this is a best effort approach and does not guarantee to stop exactly on the specified timeout.

Limit the number of results when the time limit is reached

Alternatively, you can return the number of results which have already been fetched by the time the limit is reached. Note that only the Lucene part of the query is influenced by this limit. It is possible that, if you retrieve managed object, it takes longer to fetch these objects.



This approach is not compatible with the `setTimeout` approach.

To define this soft limit, use the following approach

Example 110. Defining a time limit in query execution

```
Query luceneQuery = ...;
FullTextQuery query = fullTextSession.createFullTextQuery(luceneQuery, User.class);

//define the timeout in seconds
query.limitExecutionTimeTo(500, TimeUnit.MILLISECONDS);
List results = query.list();
```

Likewise `getResultSize()`, `iterate()` and `scroll()` honor the time limit but only until the end of the method call. That simply means that the methods of `Iterable` or the `ScrollableResults` ignore the timeout.

You can determine if the results have been partially loaded by invoking the `hasPartialResults` method.

Example 111. Determines when a query returns partial results

```
Query luceneQuery = ...;
FullTextQuery query = fullTextSession.createFullTextQuery(luceneQuery, User.class);

//define the timeout in seconds
query.limitExecutionTimeTo(500, TimeUnit.MILLISECONDS);
List results = query.list();

if ( query.hasPartialResults() ) {
    displayWarningToUser();
}
```

If you use the JPA API, `limitExecutionTimeTo` and `hasPartialResults` are also available to you.

5.2. Retrieving the results

Once the Hibernate Search query is built, executing it is in no way different than executing a HQL or Criteria query. The same paradigm and object semantic applies. All the common operations are available: `list()`, `uniqueResult()`, `iterate()`, `scroll()`.

5.2.1. Performance considerations

If you expect a reasonable number of results (for example using pagination) and expect to work on all of them, `list()` or `uniqueResult()` are recommended. `list()` work best if the entity `batch-size` is set up properly. Note that Hibernate Search has to process all Lucene Hits elements (within the pagination) when using `list()`, `uniqueResult()` and `iterate()`.

If you wish to minimize Lucene document loading, `scroll()` is more appropriate. Don't forget to close the `ScrollableResults` object when you're done, since it keeps Lucene resources. If you expect to use scroll, but wish to load objects in batch, you can use `query.setFetchSize()`. When an object is accessed, and if not already loaded, Hibernate Search will load the next `fetchSize` objects in one pass.



Pagination is preferred over scrolling.

5.2.2. Result size

It is sometimes useful to know the total number of matching documents:

- for the Google-like feature "1-10 of about 888,000,000"
- to implement a fast pagination navigation
- to implement a multi step search engine (adding approximation if the restricted query return no or not enough results)

Of course it would be too costly to retrieve all the matching documents. Hibernate Search allows you to retrieve the total number of matching documents regardless of the pagination parameters. Even more interesting, you can retrieve the number of matching elements without triggering a single object load.

Example 112. Determining the result size of a query

```
org.hibernate.search.FullTextQuery query =
    s.createFullTextQuery(luceneQuery, Book.class);
//return the number of matching books without loading a single one
assert 3245 == query.getResultSize();

org.hibernate.search.FullTextQuery query =
    s.createFullTextQuery(luceneQuery, Book.class);
query.setMaxResult(10);
List results = query.list();
//return the total number of matching books regardless of pagination
assert 3245 == query.getResultSize();
```



Like Google, the number of results is an approximation if the index is not fully up-to-date with the database (asynchronous cluster for example).

5.2.3. ResultTransformer

As seen in [Projection](#) projection results are returned as Object arrays. This data structure is not always matching the application needs. In this case it is possible to apply a ResultTransformer which post query execution can build the needed data structure:

Example 113. Using ResultTransformer in conjunction with projections

```
org.hibernate.search.FullTextQuery query =
    s.createFullTextQuery(luceneQuery, Book.class);
query.setProjection("title", "mainAuthor.name");

query.setResultTransformer(
    new StaticAliasToBeanResultTransformer(
        BookView.class,
        "title",
        "author" )
);
List<BookView> results = (List<BookView>) query.list();
for (BookView view : results) {
    log.info("Book: " + view.getTitle() + ", " + view.getAuthor());
}
```

Examples of ResultTransformer implementations can be found in the Hibernate Core codebase.

5.2.4. Understanding results

You will find yourself sometimes puzzled by a result showing up in a query or a result not showing up in a query. Luke is a great tool to understand those mysteries. However, Hibernate Search also gives you access to the Lucene Explanation object for a given result (in a given query). This class is considered fairly advanced to Lucene users but can provide a good understanding of the scoring of an object. You have two ways to access the Explanation object for a given result:

- Use the `fullTextQuery.explain(int)` method
- Use projection

The first approach takes a document id as a parameter and return the Explanation object. The document id can be retrieved using projection and the `FullTextQuery.DOCUMENT_ID` constant.



The Document id has nothing to do with the entity id. Do not mess up these two notions.

In the second approach you project the Explanation object using the `FullTextQuery.EXPLANATION` constant.

Example 114. Retrieving the Lucene Explanation object using projection

```
FullTextQuery ftQuery = s.createFullTextQuery( luceneQuery, Dvd.class )
    .setProjection(
        FullTextQuery.DOCUMENT_ID,
        FullTextQuery.EXPLANATION,
        FullTextQuery.THIS );
@SuppressWarnings("unchecked") List<Object[]> results = ftQuery.list();
for (Object[] result : results) {
    Explanation e = (Explanation) result[1];
    display( e.toString() );
}
```

Be careful, building the explanation object is quite expensive, it is roughly as expensive as running the Lucene query again. Don't do it if you don't need the object

5.3. Filters

5.3.1. Full-text filters

Apache Lucene has a powerful feature that allows to filter query results according to a custom filtering process. This is a very powerful way to apply additional data restrictions, especially since filters can be cached and reused. Some interesting use cases are:

- security
- temporal data (eg. view only last month's data)
- population filter (eg. search limited to a given category)
- and many more

Hibernate Search pushes the concept further by introducing the notion of parameterizable named filters which are transparently cached. For people familiar with the notion of Hibernate Core filters, the API is very similar:

Example 115. Enabling fulltext filters for a given query

```
fullTextQuery = s.createFullTextQuery(query, Driver.class);
fullTextQuery.enableFullTextFilter("bestDriver");
fullTextQuery.enableFullTextFilter("security").setParameter("login", "andre");
fullTextQuery.list(); //returns only best drivers where andre has credentials
```

In this example we enabled two filters on top of the query. You can enable (or disable) as many filters as you like.

Declaring filters is done through the `@FullTextFilterDef` annotation. You can use one or more `@FullTextFilterDef` on any: `*@Indexed` entity regardless of the query the filter is later applied to

* Parent class of an `@Indexed` entity * package-info.java of a package containing an `@Indexed` entity

This implies that filter definitions are global and their names must be unique. A `SearchException` is thrown in case two different `@FullTextFilterDef` annotations with the same name are defined. Each named filter has to specify a way to retrieve the actual filter implementation.

Example 116. Defining and implementing a Filter

```
@Entity
@Indexed
@FullTextFilterDef(name = "bestDriver", impl = BestDriversFilterFactory.class)
@FullTextFilterDef(name = "security", impl = SecurityFilterFactory.class)
public class Driver { ... }
```

```
public class BestDriversFilterFactory {

    @org.hibernate.search.annotations.Factory
    public Query create() {
        return new TermQuery( new Term( "score", "5" ) );
    }

}
```

`BestDriversFilterFactory` is an example of a simple Lucene filter which reduces the result set to drivers whose score is 5. In this example we use the factory pattern: the class assigned to `@FullTextFilterDef.impl` is a factory class, and the actual filter will be returned by a `@Factory` annotated, no-argument method on this class. Make sure the factory has a public constructor which does not require any parameter.

Alternatively, you can assign to `@FullTextFilterDef.impl` the exact type of your filter, i.e. a class extending `org.apache.lucene.search.Query`. The class will still have to provide a public, no-argument constructor.



In previous versions of Hibernate Search, filters had to extend the `org.apache.lucene.search.Filter` type instead of `org.apache.lucene.search.Query`.

This is still supported, but the `Filter` type has been deprecated and will be removed in a future version of Lucene.

Thus, it is advisable to convert your existing filters and filter factories so as not to rely on the `org.apache.lucene.search.Filter` type anymore, and to use simple queries instead. There should be a `Query` equivalent for every `Filter` subtype: `TermQuery` replaces `TermFilter`, `BooleanQuery` replaces `BooleanFilter`, etc.

Named filters come in handy where parameters have to be passed to the filter. For example a security

filter might want to know which security level you want to apply:

Example 117. Passing parameters to a defined filter

```
fullTextQuery = s.createFullTextQuery(query, Driver.class);
fullTextQuery.enableFullTextFilter("security").setParameter("level", 5);
```

Each parameter must have an associated setter on either the filter or filter factory of the targeted named filter definition.

Example 118. Using parameters in the actual filter implementation

```
public class SecurityFilterFactory {
    private Integer level;

    /**
     * injected parameter
     */
    public void setLevel(Integer level) {
        this.level = level;
    }

    @Factory
    public Query getFilter() {
        return new TermQuery( new Term( "level", level.toString() ) );
    }
}
```

Filters will be cached once created, based on all their parameter names and values. Caching happens using a combination of hard and soft references to allow disposal of memory when needed. The hard reference cache keeps track of the most recently used filters and transforms the ones least used to `SoftReferences` when needed. Once the limit of the hard reference cache is reached additional filters are cached as `SoftReferences`. To adjust the size of the hard reference cache, use `hibernate.search.filter.cache_strategy.size` (defaults to 128). For advanced use of filter caching, you can implement your own `FilterCachingStrategy`. The classname is defined by `hibernate.search.filter.cache_strategy`.

This filter caching mechanism should not be confused with caching the actual filter results. In Lucene it is common practice to wrap filters using the `IndexReader` around a `CachingWrapperQuery`. The wrapper will cache the set of matching documents to avoid expensive re-computation. It is important to mention that the computed set of matching documents is only cachable for the same `IndexReader` instance, because the reader effectively represents the state of the index at the moment it was opened. The document list cannot change within an opened `IndexReader`. A different/new `IndexReader` instance, however, works potentially on a different set of documents (either from a different index or simply because the index has changed), hence the filter result has to be recomputed.

Hibernate Search also helps with this aspect of caching. Per default the `cache` flag of `@FullTextFilterDef` is set to `FilterCacheModeType.INSTANCE_AND_DOCIDSETRESULTS` which

will automatically cache the filter instance as well as wrap the specified filter around a Hibernate specific implementation of `CachingWrapperQuery`. In contrast to Lucene's version of this class `SoftReferences` are used together with a hard reference count (see discussion about filter cache). The hard reference count can be adjusted using `hibernate.search.filter.cache_docidresults.size` (defaults to 5). The wrapping behavior can be controlled using the `@FullTextFilterDef.cache` parameter. There are three different values for this parameter:

Value	Definition
<code>FilterCacheModeType.NONE</code>	No filter instance and no result is cached by Hibernate Search. For every filter call, a new filter instance is created. This setting might be useful for rapidly changing data sets or heavily memory constrained environments.
<code>FilterCacheModeType.INSTANCE_ONLY</code>	The filter instance is cached and reused across concurrent filter calls. Filter results are not cached. This setting is useful when a filter uses its own specific caching mechanism or the filter results change dynamically due to application specific events making filter results caching in both cases unnecessary.
<code>FilterCacheModeType.INSTANCE_AND_DOCIDSETRESULTS</code>	Both the filter instance and the filter results are cached. This is the default value.

Last but not least - why should filters be cached? There are two areas where filter caching shines:

- the system does not update the targeted entity index often (in other words, the `IndexReader` is reused a lot)
- the filter's result is expensive to compute (compared to the time spent to execute the query)

5.3.2. Using filters in a sharded environment

It is possible, in a sharded environment to execute queries on a subset of the available shards. This can be done in two steps:

- create a sharding strategy that does select a subset of `IndexManagers` depending on some filter configuration
- activate the proper filter at query time

Let's first look at an example of sharding strategy that query on a specific customer shard if the customer filter is activated.

```

public class CustomerShardingStrategy implements IndexShardingStrategy {

    // stored IndexManagers in a array indexed by customerID
    private IndexManager[] indexManagers;

    public void initialize(Properties properties, IndexManager[] indexManagers) {
        this.indexManagers = indexManagers;
    }

    public IndexManager[] getIndexManagersForAllShards() {
        return indexManagers;
    }

    public IndexManager getIndexManagerForAddition(
        Class<?> entity, Serializable id, String idInString, Document document) {
        Integer customerID = Integer.parseInt(document.getFieldable("customerID").stringValue());
    };
    return indexManagers[customerID];
    }

    public IndexManager[] getIndexManagersForDeletion(
        Class<?> entity, Serializable id, String idInString) {
        return getIndexManagersForAllShards();
    }

    /**
     * Optimization; don't search ALL shards and union the results; in this case, we
     * can be certain that all the data for a particular customer Filter is in a single
     * shard; simply return that shard by customerID.
     */
    public IndexManager[] getIndexManagersForQuery(
        FullTextFilterImplementor[] filters) {
        FullTextFilter filter = getCustomerFilter(filters, "customer");
        if (filter == null) {
            return getIndexManagersForAllShards();
        }
        else {
            return new IndexManager[] { indexManagers[Integer.parseInt(
                filter.getParameter("customerID").toString())] };
        }
    }

    private FullTextFilter getCustomerFilter(FullTextFilterImplementor[] filters, String name)
    {
        for (FullTextFilterImplementor filter: filters) {
            if (filter.getName().equals(name)) return filter;
        }
        return null;
    }
}

```

In this example, if the filter named `customer` is present, we make sure to only use the shard dedicated to this customer. Otherwise, we return all shards. A given Sharding strategy can react to one or more filters and depends on their parameters.

The second step is simply to activate the filter at query time. While the filter can be a regular filter (as defined in [Full-text filters](#)) which also filters Lucene results after the query, you can make use of a special filter that will only be passed to the sharding strategy and otherwise ignored for the rest of the query. Simply use the `ShardSensitiveOnlyFilter` class when declaring your filter.

```
@Entity @Indexed
@FullTextFilterDef(name="customer", impl=ShardSensitiveOnlyFilter.class)
public class Customer {
    // ...
}
```

```
FullTextQuery query = ftEm.createFullTextQuery(luceneQuery, Customer.class);
query.enableFulltextFilter("customer").setParameter("CustomerID", 5);
@SuppressWarnings("unchecked")
List<Customer> results = query.getResultList();
```

Note that by using the `ShardSensitiveOnlyFilter`, you do not have to implement any Lucene filter. Using filters and sharding strategy reacting to these filters is recommended to speed up queries in a sharded environment.

5.3.3. Hibernate ORM filters

Hibernate ORM filters, enabled using `org.hibernate.Session.enableFilter(String)`, will only work partially when combined with a full-text query.

The filtered entities will be excluded from the results, but:

- the result size returned by `FullTextQuery.getResultSize()` will ignore the filters (it will return the size before applying the filters);
- if you use paging, the `setFirstResult` and `setMaxResults` modifiers will be applied before the filters, leading to some result pages being smaller than `maxResults` (or even empty) even though they are not the last.

These limitations are necessary to keep performance at a reasonable level. If they are not acceptable for your use case, you should use [full-text filters](#) in your full-text queries.

5.4. Faceting

[Faceted search](#) is a technique which allows to divide the results of a query into multiple categories. This categorization includes the calculation of hit counts for each category and the ability to further restrict search results based on these facets (categories). [Facets Example on Amazon](#) shows a faceting example. The search for 'Hibernate Search' results in fifteen hits which are displayed on the main part of the page. The navigation bar on the left, however, shows the category *Computers & Internet* with its subcategories *Programming*, *Computer Science*, *Databases*, *Software*, *Web Development*, *Networking* and *Home Computing*. For each of these subcategories the number of books is shown matching the main search criteria and belonging to the respective subcategory. This division of the category *Computers & Internet* is one facet of this search. Another one is for example the

average customer review rating.

The screenshot shows the Amazon search results page for "Hibernate Search". On the left, there are several facet categories: Department (Books), Format (Paperback), Author (Joe Vitale), Shipping Option (Free Super Saver Shipping), Avg. Customer Review (4 stars & up), and Condition (Used, New). The main search results area shows three books:

- Hibernate Search in Action** by Howard Lewis Ship. It has 3 customer reviews, a 4-star rating, and is available in Paperback format for \$49.99. The book is eligible for FREE Super Saver Shipping.
- Spring Persistence with Hibernate** by The Hibernate Team and Brian D. Mayfield. It has 5 customer reviews, a 4-star rating, and is available in Paperback format for \$44.99. It is also available in Kindle Edition.
- Lucene in Action, Second Edition** by Hatcher and Otis Gospodnetic. It has 1 customer review, a 4-star rating, and is available in Paperback format for \$49.99.

Figure 1. Facets Example on Amazon

In Hibernate Search the classes `QueryBuilder` and `FullTextQuery` are the entry point to the faceting API. The former allows to create faceting requests whereas the latter gives access to the so called `FacetManager`. With the help of the `FacetManager` faceting requests can be applied on a query and selected facets can be added to an existing query in order to refine search results. The following sections will describe the faceting process in more detail. The examples will use the entity `Cd` as shown in [Example entity for faceting](#):

Example 119. Example entity for faceting

```
@Entity
@Indexed
public class Cd {

    @Id
    @GeneratedValue
    private int id;

    @Field,
    private String name;

    @Field(analyze = Analyze.NO)
    @Facet
    private int price;

    @Field(analyze = Analyze.NO)
    @DateBridge(resolution = Resolution.YEAR)
    @Facet
    private Date releaseYear;

    @Field(analyze = Analyze.NO)
    @Facet
    private String label;

    // setter/getter
    // ...
}
```

In order to facet on a given indexed field, the field needs to be configured with the `@Facet` annotation. Also, the field itself cannot be analyzed.

`@Facet` contains a `name` and `forField` parameter. The `name` is arbitrary and used to identify the facet. Per default it matches the field name it belongs to. `forField` is relevant in case the property is mapped to multiple fields using `@Field` multiple times, as it is a repeatable annotation. In this case `forField` can be used to identify the index field to which it applies. Both `@Field` and `@Facet` are repeatable annotations.

Last but not least, `@Facet` contains a `encoding` parameter. Usually, Hibernate Search automatically selects the encoding:

- String fields are encoded as `FacetEncodingType.STRING`
- `byte`, `short`, `int`, `long` (including corresponding wrapper types) and `Date` as `FacetEncodingType.LONG`
- and `float` and `double` (including corresponding wrapper types) as `FacetEncodingType.DOUBLE`

In some cases it can make sense, however, to explicitly set the encoding. [Discrete faceting requests](#) for example only work for string encoded facets. In order to use a discrete facet for numbers the encoding must be explicitly set to `FacetEncodingType.STRING`.



Pre Hibernate Search 5.2 there was no need to explicitly use a `@Facet` annotation. In 5.2 it became necessary in order to use Lucene's native faceting API.

5.4.1. Creating a faceting request

The first step towards a faceted search is to create the `FacetingRequest`. Currently two types of faceting requests are supported. The first type is called *discrete faceting* and the second type *range faceting* request.

Discrete faceting request

In the case of a discrete faceting request, you start with giving the request a unique name. This name will later be used to retrieve the facet values (see [Interpreting a Facet result](#)). Then you need to specify on which index field you want to categorize on and which faceting options to apply. An example for a discrete faceting request can be seen in [Creating a discrete faceting request](#):

Example 120. Creating a discrete faceting request

```
QueryBuilder builder = fullTextSession.getSearchFactory()
    .buildQueryBuilder().forEntity(Cd.class).get();

FacetingRequest labelFacetingRequest = builder.facet()
    .name("labelFacetRequest")
    .onField("label")
    .discrete()
    .orderBy(FacetSortOrder.COUNT_DESC)
    .includeZeroCounts(false)
    .maxFacetCount(3)
    .createFacetingRequest();
```

When executing this faceting request a `Facet` instance will be created for each discrete value for the indexed field `label`. The `Facet` instance will record the actual field value including how often this particular field value occurs within the original query results. Parameters `orderBy`, `includeZeroCounts` and `maxFacetCount` are optional and can be applied on any faceting request. Parameter `orderBy` allows to specify in which order the created facets will be returned. The default is `FacetSortOrder.COUNT_DESC`, but you can also sort on the field value. Parameter `includeZeroCount` determines whether facets with a count of 0 will be included in the result (by default they are not) and `maxFacetCount` allows to limit the maximum amount of facets returned (by default 100 facets are returned).

There are several preconditions an indexed field has to meet in order to categorize (facet) on it:

- The indexed property must be of type `String`, `Date` or of the numeric type `byte`, `short`, `int`, `long`, `double` or `float` (or their respective Java wrapper types).
- The property has to be indexed with `Analyze.NO`.
- `null` values should be avoided.



When you need conflicting options, we suggest to index the property twice and use the appropriate field depending on the use case:

```
@Field(name="price")
@Field(name="price_facet",
      analyze=Analyze.NO,
      bridge=@FieldBridge(impl = IntegerBridge.class))
private int price;
```

Creating a range faceting request

The creation of a range faceting request is similar. We also start with a name for the request and the field to facet on. Then we have to specify ranges for the field values. A range faceting request can be seen in [Creating a range faceting request](#). There, three different price ranges are specified. `below` and `above` can only be specified once, but you can specify as many `from - to` ranges as you want. For each range boundary you can also specify via `excludeLimit` whether it is included into the range or not.

Example 121. Creating a range faceting request

```
QueryBuilder builder = fullTextSession.getSearchFactory()
    .buildQueryBuilder()
    .forEntity(Cd.class)
    .get();
FacetingRequest priceFacetingRequest = builder.facet()
    .name("priceFaceting")
    .onField("price_facet")
    .range()
    .below(1000)
    .from(1001).to(1500)
    .above(1500).excludeLimit()
    .createFacetingRequest();
```

5.4.2. Setting the facet sort order

The result of applying a faceting request is a list of `Facet` instances as seen in [Applying a faceting request](#). The order within the list is given by the `FacetSortOrder` parameter specified via `orderBy` when creating the faceting request. The default value is `FacetSortOrder.COUNT_DESC`, meaning facets are ordered by their count in descending order (highest count first). Other values are `COUNT_ASC`, `FIELD_VALUE` and `RANGE_DEFINITION_ORDER`. `COUNT_ASC` returns the facets in

ascending count order whereas `FIELD_VALUE` will return them in alphabetical order of the facet/category value (see [Interpreting a Facet result](#)). `RANGE_DEFINITION_ORDER` only applies for range faceting request and returns the facets in the same order in which the ranges are defined. For [Creating a range faceting request](#) this would mean the facet for the range of below 1000 would be returned first, followed by the facet for the range 1001 to 1500 and finally the facet for above 1500.

5.4.3. Applying a faceting request

In [Creating a faceting request](#) we have seen how to create a faceting request. Now it is time to apply it on a query. The key is the `FacetManager` which can be retrieved via the `FullTextQuery` (see [Applying a faceting request](#)).

Example 122. Applying a faceting request

```
// create a fulltext query
Query luceneQuery = builder.all().createQuery(); // match all query
FullTextQuery fullTextQuery = fullTextSession.createFullTextQuery(luceneQuery, Cd.class);

// retrieve facet manager and apply faceting request
FacetManager facetManager = fullTextQuery.getFacetManager();
facetManager.enableFaceting(priceFacetingRequest);

// get the list of Cds
List<Cd> cds = fullTextQuery.list();
...

// retrieve the faceting results
List<Facet> facets = facetManager.getFacets("priceFaceting");
...
```

You need to enable the faceting request before you execute the query. You do that via `facetManager.enableFaceting(<facetName>)`. You can enable as many faceting requests as you like. Then you execute the query and retrieve the facet results for a given request via `facetManager.getFacets(<facetname>)`. For each request you will get a list of `Facet` instances. Facet requests stay active and get applied to the fulltext query until they are either explicitly disabled via `disableFaceting(<facetName>)` or the query is `discarded`.

5.4.4. Interpreting a Facet result

Each facet request results in a list of `Facet` instances. Each instance represents one facet/category value. In the CD example ([Creating a discrete faceting request](#)) where we want to categorize on the CD labels, there would for example be a `Facet` for each of the record labels Universal, Sony and Warner. [Facet API](#) shows the API of `Facet`.

Example 123. Facet API

```
public interface Facet {
    /**
     * @return the faceting name this {@code Facet} belongs to.
     *
     * @see org.hibernate.search.query.facet.FacetingRequest#getFacetingName()
     */
    String getFacetingName();

    /**
     * Return the {@code Document} field name this facet is targeting.
     * The field needs to be indexed with {@code Analyze.NO}.
     *
     * @return the {@code Document} field name this facet is targeting.
     */
    String getFieldName();

    /**
     * @return the value of this facet. In case of a discrete facet it is the actual
     *         {@code Document} field value. In case of a range query the value is a
     *         string representation of the range.
     */
    String getValue();

    /**
     * @return the facet count.
     */
    int getCount();

    /**
     * @return a Lucene {@link Query} which can be executed to retrieve all
     *         documents matching the value of this facet.
     */
    Query getFacetQuery();
}
```

`getFacetingName()` and `getFieldName()` are returning the facet request name and the targeted document field name as specified by the underlying `FacetRequest`. For example "Creating a discrete faceting request" that would be `labelFacetRequest` and `label` respectively. The interesting information is provided by `getValue()` and `getCount()`. The former is the actual facet/category value, for example a concrete record label like Universal. The latter returns the count for this value. To stick with the example again, the count value tells you how many Cds are released under the Universal label. Last but not least, `getFacetQuery()` returns a Lucene query which can be used to retrieve the entities counted in this facet.

5.4.5. Restricting query results

Support for selecting facets ("drill-down") will be removed in Hibernate Search 6.



In order to drill-down and have your query return only specific facets, wrap it in a [boolean query](#) and add a filter clause ([filteredBy](#)) to constraint values of the facet field with either a [keyword query](#) or a [range query](#).

See also [this section of the migration guide](#).

A common use case for faceting is a "drill-down" functionality which allows you to narrow your original search by applying a given facet on it. To do this, you can apply any of the returned `Facet` instances as additional criteria on your original query via `FacetSelection`. `FacetSelection` is available via the `FacetManager` and allow you to select a facet as query criteria (`selectFacets`), remove a facet restriction (`deselectFacets`), remove all facet restrictions (`clearSelectedFacets`) and retrieve all currently selected facets (`getSelectedFacets`). [Restricting query results via the application of a `FacetSelection`](#) shows an example.

Example 124. Restricting query results via the application of a `FacetSelection`

```
// create a fulltext query
Query luceneQuery = builder.all().createQuery(); // match all query
FullTextQuery fullTextQuery = fullTextSession.createFullTextQuery( luceneQuery, clazz );

// retrieve facet manager and apply faceting request
FacetManager facetManager = fullTextQuery.getFacetManager();
facetManager.enableFaceting( priceFacetingRequest );

// get the list of Cd
List<Cd> cds = fullTextQuery.list();
assertTrue(cds.size() == 10);

// retrieve the faceting results
List<Facet> facets = facetManager.getFacets( "priceFaceting" );
assertTrue(facets.get(0).getCount() == 2)

// apply first facet as additional search criteria
FacetSelection facetSelection = facetManager.getFacetGroup( "priceFaceting" );
facetSelection.selectFacets( facets.get( 0 ) );

// re-execute the query
cds = fullTextQuery.list();
assertTrue(cds.size() == 2);
```

Per default selected facets are combined via disjunction (OR). In case a field has multiple values, like a potential `Cd.artists` association, you can also use conjunction (AND) for the facet selection.

Example 125. Using conjunction in `FacetSelection`

```
FacetSelection facetSelection = facetManager.getFacetGroup( "artistsFaceting" );
facetSelection.selectFacets( FacetCombine.AND, facets.get( 0 ), facets.get( 1 ) );
```

5.5. Optimizing the query process

Query performance depends on several criteria:

- the Lucene query itself: read the literature on this subject.
- the number of loaded objects: use pagination and / or index projection (if needed).
- the way Hibernate Search interacts with the Lucene readers: defines the appropriate [Reader strategy](#).

5.5.1. Logging executed search queries

Knowing the executed search queries is vital when working on performance optimizations. This is especially the case if your application accepts queries passed in by the user or e.g. dynamically builds queries using the Hibernate Search query DSL.

In order to log all search queries executed by Hibernate Search, enable **DEBUG** logging for the log category `org.hibernate.search.fulltext_query`.

Chapter 6. Manual index changes

As Hibernate core applies changes to the Database, Hibernate Search detects these changes and will update the index automatically (unless the EventListeners are disabled). Sometimes changes are made to the database without using Hibernate, as when backup is restored or your data is otherwise affected; for these cases Hibernate Search exposes the Manual Index APIs to explicitly update or remove a single entity from the index, or rebuild the index for the whole database, or remove all references to a specific type.

All these methods affect the Lucene Index only, no changes are applied to the Database.

6.1. Adding instances to the index

Using `FullTextSession.index(T entity)` you can directly add or update a specific object instance to the index. If this entity was already indexed, then the index will be updated. Changes to the index are only applied at transaction commit.

Example 126. Indexing an entity via `FullTextSession.index(T entity)`

```
FullTextSession fullTextSession = Search.getFullTextSession(session);
Transaction tx = fullTextSession.beginTransaction();
Object customer = fullTextSession.load( Customer.class, 8 );
fullTextSession.index(customer);
tx.commit(); //index only updated at commit time
```

In case you want to add all instances for a type, or for all indexed types, the recommended approach is to use a `MassIndexer`: see [Using a MassIndexer](#) for more details.

The method `FullTextSession.index(T entity)` is considered an explicit indexing operation, so any registered `EntityIndexingInterceptor` won't be applied in this case. For more information on `EntityIndexingInterceptor` see [Conditional indexing](#).

6.2. Deleting instances from the index

It is equally possible to remove an entity or all entities of a given type from a Lucene index without the need to physically remove them from the database. This operation is named purging and is also done through the `FullTextSession`.

Example 127. Purging a specific instance of an entity from the index

```
FullTextSession fullTextSession = Search.getFullTextSession(session);
Transaction tx = fullTextSession.beginTransaction();
for (Customer customer : customers) {
    fullTextSession.purge( Customer.class, customer.getId() );
}
tx.commit(); //index is updated at commit time
```

Purging will remove the entity with the given id from the Lucene index but will not touch the database.

If you need to remove all entities of a given type, you can use the `purgeAll` method. This operation removes all entities of the type passed as a parameter as well as all its subtypes.

Example 128. Purging all instances of an entity from the index

```
FullTextSession fullTextSession = Search.getFullTextSession(session);
Transaction tx = fullTextSession.beginTransaction();
fullTextSession.purgeAll( Customer.class );
//optionally optimize the index
//fullTextSession.getSearchFactory().optimize( Customer.class );
tx.commit(); //index changes are applied at commit time
```

As in the previous example, it is suggested to optimize the index after many purge operation to actually free the used space.

As is the case with method `FullTextSession.index(T entity)`, also `purge` and `purgeAll` are considered explicit indexing operations: any registered `EntityIndexingInterceptor` won't be applied. For more information on `EntityIndexingInterceptor` see [Conditional indexing](#).



Methods `index`, `purge` and `purgeAll` are available on `FullTextEntityManager` as well.



All manual indexing methods (`index`, `purge` and `purgeAll`) only affect the index, not the database, nevertheless they are transactional and as such they won't be applied until the transaction is successfully committed, or you make use of `flushToIndexes`.

6.3. Rebuilding the whole index

If you change the entity mapping to the index, chances are that the whole index needs to be updated; for example if you decide to index an existing field using a different analyzer you'll need to rebuild the index for affected types. Also if the Database is replaced (like restored from a backup, imported from a legacy system) you'll want to be able to rebuild the index from existing data. Hibernate Search provides two main strategies to choose from:

- Using `FullTextSession.flushToIndexes()` periodically, while using `FullTextSession.index()` on all entities.
- Use a `MassIndexer`.

When you change the entity mapping to the index, if you don't know whether the changes require a full reindexing or not, you can follow this rule of thumb:



- Rebuilding the whole index is **not** required if all you did was:
 - **adding** new indexed entities that will not have any persisted instance, e.g. adding an `@Indexed` annotation on an entity which has no rows in database.
 - **adding** new fields that will be empty for all currently persisted entities, e.g. adding a new property on an entity type and mapping it to a field, but with the guarantee that this property will initially be null for every instance of this entity;
 - and/or **removing** data from existing indexes/fields, e.g. removing an index field, or removing the need for a field to be stored.
- Rebuilding the whole index is required if you made any change to the mapping that is not mentioned above.

6.3.1. Using `flushToIndexes()`

This strategy consists in removing the existing index and then adding all entities back to the index using `FullTextSession.purgeAll()` and `FullTextSession.index()`, however there are some memory and efficiency constraints. For maximum efficiency Hibernate Search batches index operations and executes them at commit time. If you expect to index a lot of data you need to be careful about memory consumption since all documents are kept in a queue until the transaction commit. You can potentially face an `OutOfMemoryException` if you don't empty the queue periodically: to do this you can use `fullTextSession.flushToIndexes()`. Every time `fullTextSession.flushToIndexes()` is called (or if the transaction is committed), the batch queue is processed applying all index changes. Be aware that, once flushed, the changes cannot be rolled back.

Example 129. Index rebuilding using index() and flushToIndexes()

```
fullTextSession.setFlushMode(FlushMode.MANUAL);
fullTextSession.setCacheMode(CacheMode.IGNORE);
transaction = fullTextSession.beginTransaction();
//Scrollable results will avoid loading too many objects in memory
ScrollableResults results = fullTextSession.createCriteria( Email.class )
    .setFetchSize(BATCH_SIZE)
    .scroll(ScrollMode.FORWARD_ONLY);
int index = 0;
while(results.next()) {
    index++;
    fullTextSession.index(results.get(0)); //index each element
    if (index % BATCH_SIZE == 0) {
        fullTextSession.flushToIndexes(); //apply changes to indexes
        fullTextSession.clear(); //free memory since the queue is processed
    }
}
transaction.commit();
```

Try to use a batch size that guarantees that your application will not run out of memory: with a bigger batch size objects are fetched faster from database but more memory is needed.

6.3.2. Using a MassIndexer

Hibernate Search's MassIndexer uses several parallel threads to rebuild the index; you can optionally select which entities need to be reloaded or have it reindex all entities. This approach is optimized for best performance but requires to set the application in maintenance mode: making queries to the index is not recommended when a MassIndexer is busy.

Example 130. Index rebuilding using a MassIndexer

```
fullTextSession.createIndexer().startAndWait();
```

This will rebuild the index, deleting it and then reloading all entities from the database. Although it's simple to use, some tweaking is recommended to speed up the process: there are several parameters configurable.



During the progress of a MassIndexer the content of the index is undefined! If a query is performed while the MassIndexer is working most likely some results will be missing.

Example 131. Using a tuned MassIndexer

```
fullTextSession
    .createIndexer( User.class )
    .batchSizeToLoadObjects( 25 )
    .cacheMode( CacheMode.NORMAL )
    .threadsToLoadObjects( 12 )
    .idFetchSize( 150 )
    .transactionTimeout( 1800 )
    .progressMonitor( monitor ) //a MassIndexerProgressMonitor implementation
    .startAndWait();
```

This will rebuild the index of all `User` instances (and subtypes), and will create 12 parallel threads to load the `User` instances using batches of 25 objects per query; these same 12 threads will also need to process indexed embedded relations and custom `FieldBridges` or `ClassBridges`, to finally output a Lucene document. In this conversion process these threads are likely going to need to trigger lazy loading of additional attributes, so you will probably need a high number of threads working in parallel. When run in a JTA environment such as the WildFly application server, the mass indexer will use a timeout of 1800 seconds (= 30 minutes) for its transactions. Configure a timeout value which is long enough to load and index all entities of the type with the most instances, taking into account the configured batch size and number of threads to load objects. Note that these transactions are read-only, so choosing a substantially large value should pose no problem in general.

As of Hibernate Search 4.4.0, instead of indexing all the types in parallel, the `MassIndexer` is configured by default to index only one type in parallel. It prevents resource exhaustion especially database connections and usually does not slow down the indexing. You can however configure this behavior using `MassIndexer.typesToIndexInParallel(int threadsToIndexObjects)`:

Example 132. Configuring the MassIndexer to index several types in parallel

```
fullTextSession
    .createIndexer( User.class, Customer.class )
    .typesToIndexInParallel( 2 )
    .batchSizeToLoadObjects( 25 )
    .cacheMode( CacheMode.NORMAL )
    .threadsToLoadObjects( 5 )
    .idFetchSize( 150 )
    .progressMonitor( monitor ) //a MassIndexerProgressMonitor implementation
    .startAndWait();
```

Generally we suggest to leave `cacheMode` to `CacheMode.IGNORE` (the default), as in most reindexing situations the cache will be a useless additional overhead; it might be useful to enable some other `CacheMode` depending on your data: it could increase performance if the main entity is relating to enum-like data included in the index.



The MassIndexer was designed for speed and is unaware of transactions, so there is no need to begin one or committing. Also because it is not transactional it is not recommended to let users use the system during its processing, as it is unlikely people will be able to find results and the system load might be too high anyway.

MassIndexer using threads and JDBC connections

The MassIndexer was designed to finish the re-indexing task as quickly as possible, but this requires a bit of care in its configuration to behave fairly with your server resources.

There is a simple formula to understand how the different options applied to the MassIndexer affect the number of used worker threads and connections: each thread will require a JDBC connection.

```
threads = typesToIndexInParallel * (threadsToLoadObjects + 1);  
required JDBC connections = threads;
```

Let's see some suggestions for a roughly sane tuning starting point:

1. Option `typesToIndexInParallel` should probably be a low value, like 1 or 2, depending on how much of your CPUs have spare cycles and how slow a database round trip will be.
2. Before tuning a parallel run, experiment with options to tune your primary indexed entities in isolation.
3. Making `threadsToLoadObjects` higher increases the pre-loading rate for the picked entities from the database, but also increases memory usage and the pressure on the threads working on subsequent indexing.
4. Increasing parallelism usually helps as the bottleneck usually is the latency to the database connection: it's probably worth it to experiment with values significantly higher than the number of actual cores available, but make sure your database can handle all the multiple requests.
5. This advice might not apply to you: always measure the effects! We're providing this as a means to help you understand how these options are related.



Running the MassIndexer with many threads will require many connections to the database. If you don't have a sufficiently large connection pool, the MassIndexer itself and/or your other applications could starve being unable to serve other requests: make sure you size your connection pool accordingly to the options as explained in the above paragraph.



The "sweet spot" of number of threads to achieve best performance is highly dependent on your overall architecture, database design and even data values. All internal thread groups have meaningful names so they should be easily identified with most diagnostic tools, including simply thread dumps.

Using a custom MassIndexer implementation

The provided MassIndexer is quite general purpose, and while we believe it's a robust approach, you might be able to squeeze some better performance by writing a custom implementation. To run your own MassIndexer instead of using the one shipped with Hibernate Search you have to:

1. create an implementation of the `org.hibernate.search.spi.MassIndexerFactory` interface;
2. set the property `hibernate.search.massindexer.factoryclass` with the qualified class name of the factory implementation.

Example 133. Custom MassIndexerFactory example

```
package org.myproject
import org.hibernate.search.spi.MassIndexerFactory

// ...

public class CustomIndexerFactory implements MassIndexerFactory {

    public void initialize(Properties properties) {
    }

    public MassIndexer createMassIndexer(...) {
        return new CustomIndexer();
    }

}
```

```
hibernate.search.massindexer.factoryclass = org.myproject.CustomIndexerFactory
```

6.3.3. Useful parameters for batch indexing

Other parameters which affect indexing time and memory consumption are:

- `hibernate.search.[default|<indexname>].exclusive_index_use`
- `hibernate.search.[default|<indexname>].indexwriter.max_buffered_docs`
- `hibernate.search.[default|<indexname>].indexwriter.max_merge_docs`
- `hibernate.search.[default|<indexname>].indexwriter.merge_factor`
- `hibernate.search.[default|<indexname>].indexwriter.merge_min_size`

- `hibernate.search.[default|<indexname>].indexwriter.merge_max_size`
- `hibernate.search.[default|<indexname>].indexwriter.merge_max_optimize_size`
- `hibernate.search.[default|<indexname>].indexwriter.merge_calibrate_by_deletes`
- `hibernate.search.[default|<indexname>].indexwriter.ram_buffer_size`

Previous versions also had a `max_field_length` but this was removed from Lucene, it's possible to obtain a similar effect by using a `LimitTokenCountAnalyzer`.

All `.indexwriter` parameters are Lucene specific and Hibernate Search is just passing these parameters through - see [Tuning indexing performance](#) for more details.

The MassIndexer uses a forward only scrollable result to iterate on the primary keys to be loaded, but MySQL's JDBC driver will load all values in memory; to avoid this "optimization" set `idFetchSize` to `Integer.MIN_VALUE`.

6.4. Integration with JSR-352

The integration with JSR-352, Batch Applications for the Java Platform, is in development. We do think we have the basics covered and we are looking for feedback.

Hibernate Search provides a JSR-352 job to perform mass indexing. It covers not only the existing functionality of the mass indexer described above, but also benefits from some powerful standard features of the Java Batch Platform (JSR-352), such as failure recovery using checkpoints, chunk oriented processing, and parallel execution. This batch job accepts different entity type(s) as input, loads the relevant entities from the database, then rebuilds the full-text index from these.

However, it requires a batch runtime for the execution. Please notice that we don't provide any batch runtime, you are free to choose one that fits you needs, e.g. the default batch runtime embedded in your Java EE container. We provide full integration to the JBeret implementation (see [how to configure it here](#)). As for other implementations, they can also be used, but will require [a bit more configuration on your side](#).

Example 134. How to use the JSR-352 mass-indexing job?

```
Properties jobProps = MassIndexingJob
    .parameters()
    .forEntity( MyEntity.class )
    .build();
long executionId = BatchRuntime
    .getJobOperator()
    .start( MassIndexingJob.NAME, jobProps );
```

6.4.1. Job Parameters

Table 14. Job Parameters in JSR 352 Integration

Parameter Name	Builder Method	Requirement	Default value	Description
<code>entityTypes</code>	<code>forEntity(Class<?>), forEntities(Class<?>, Class<?>...)</code>	Required	-	The entity types to index in this job execution, comma-separated.
<code>purgeAllOnStart</code>	<code>purgeAllOnStart(boolean)</code>	Optional	True	Specify whether the existing index should be purged at the beginning of the job. This operation takes place before indexing.
<code>optimizeAfterPurge</code>	<code>optimizeAfterPurge(boolean)</code>	Optional	True	Specify whether the mass indexer should be optimized at the beginning of the job. This operation takes place after the purge operation and before indexing.
<code>optimizeOnFinish</code>	<code>optimizeOnFinish(boolean)</code>	Optional	True	Specify whether the mass indexer should be optimized at the end of the job. This operation takes place after indexing.

Parameter Name	Builder Method	Requirement	Default value	Description
<code>cacheMode</code>	<code>cacheMode(CacheMode)</code>	Optional	<code>IGNORE</code>	Specify the Hibernate <code>CacheMode</code> when loading entities. The default is <code>IGNORE</code> , and it will be the most efficient choice in most cases, but using another mode such as <code>GET</code> may be more efficient if many of the entities being indexed refer to a small set of other entities.
<code>idFetchSize</code>	<code>idFetchSize(int)</code>	Optional	1000	Specifies the fetch size to be used when loading primary keys. Some databases accept special values, for example MySQL might benefit from using <code>Integer#MIN_VALUE</code> , otherwise it will attempt to preload everything in memory.
<code>entityFetchSize</code>	<code>entityFetchSize(int)</code>	Optional	The value of <code>sessionClearInterval</code>	Specifies the fetch size to be used when loading entities from database. Some databases accept special values, for example MySQL might benefit from using <code>Integer#MIN_VALUE</code> , otherwise it will attempt to preload everything in memory.

Parameter Name	Builder Method	Requirement	Default value	Description
<code>customQueryHQL</code>	<code>restrictedBy(String)</code>	Optional	-	Use HQL / JPQL to index entities of a target entity type. Your query should contain only one entity type. Mixing this approach with the criteria restriction is not allowed. Please notice that there's no query validation for your input. See Indexing mode for more detail and limitations.
<code>customQueryCriteria</code> (Not available, query builder only)	<code>restrictedBy(Criterion)</code>	Optional	-	Add criterion to construct a customized selection of mass-indexing under the criteria approach. You can call this method multiple times to add multiple criteria: only entities matching every criterion will be indexed. However, mixing this approach with the HQL restriction is not allowed. See Indexing mode for more detail and limitations.

Parameter Name	Builder Method	Requirement	Default value	Description
<code>maxResultsPerEntity</code>	<code>maxResultsPerEntity(int)</code>	Optional	-	The maximum number of results to load per entity type. This parameter let you define a threshold value to avoid loading too many entities accidentally. The value defined must be greater than 0. The parameter is not used by default. It is equivalent to keyword <code>LIMIT</code> in SQL.
<code>rowsPerPartition</code>	<code>rowsPerPartition(int)</code>	Optional	20,000	The maximum number of rows to process per partition. The value defined must be greater than 0, and equal to or greater than the value of <code>checkpointInterval</code> .
<code>maxThreads</code>	<code>maxThreads(int)</code>	Optional	The number of partitions	The maximum number of threads to use for processing the job. Note the batch runtime cannot guarantee the request number of threads are available; it will use as many as it can up to the request maximum.
<code>checkpointInterval</code>	<code>checkpointInterval(int)</code>	Optional	2,000, or the value of <code>rowsPerPartition</code> if it is smaller	The number of entities to process before triggering a checkpoint. The value defined must be greater than 0, and equal to or less than the value of <code>rowsPerPartition</code> .

Parameter Name	Builder Method	Requirement	Default value	Description
<code>sessionClearInterval</code>	<code>sessionClearInterval(int)</code>	Optional	200, or the value of <code>checkpointInterval</code> if it is smaller	The number of entities to process before clearing the session. The value defined must be greater than 0, and equal to or less than the value of <code>checkpointInterval</code> .
<code>entityManagerFactoryReference</code>	<code>entityManagerFactoryReference(String)</code>	Required if there's more than one persistence unit	-	The string that will identify the <code>EntityManagerFactory</code> .
<code>entityManagerFactoryNamespace</code>	<code>entityManagerFactoryNamespace(String)</code>	-	-	See Selecting the persistence unit (EntityManagerFactory)

6.4.2. Indexing mode

The mass indexing job allows you to define your own entities to be indexed—you can start a full indexing or a partial indexing through 3 different methods: selecting the desired entity types, using HQL, or using Hibernate criteria.

Example 135. How to define the indexing mode?

```
// full indexing
Properties jobProps1 = MassIndexingJob
    .parameters()
    .forEntity( MyClass.class )
    .build();

// partial indexing using HQL
Properties jobProps2 = MassIndexingJob
    .parameters()
    .forEntity( MyClass.class )
    .restrictedBy( "select c from MyClass c where c.name in ( 'Foo', 'Bar' )" )
    .build();

// partial indexing using Criteria
Properties jobProps3 = MassIndexingJob
    .parameters()
    .forEntity( MyClass.class )
    .restrictedBy( Restrictions.in( "name", "Foo", "Bar" ) )
    .build();
```

While the full indexing is useful when you perform the very first indexing, or after extensive changes to your whole database, it may also be time consuming. If you want to reindex only part of your data, you need to add restrictions using HQL or criteria: they help you to define a customized selection, and only the entities inside that selection will be indexed. A typical use-case is to index the new entities

appeared since yesterday.

Note that, as detailed below, some features may not be supported depending on the indexing mode.

Table 15. Comparison of each indexing mode

Indexing mode	Scope	Parallel Indexing
Full Indexation	All entities	Supported
Criteria	Some entities	Supported
HQL	Some entities	Not supported



When using the HQL mode, there isn't any query validation before the job's start. If the query is invalid, the job will start and fail.

Also, parallel indexing is disabled in HQL mode, because our current parallelism implementations relies on selection order, which might not be provided by the HQL given by user.

Because of those limitations, we suggest you use this approach only for indexing small numbers of entities, and only if you know that no entities matching the query will be created during indexing.

6.4.3. Parallel indexing

For better performance, indexing is performed in parallel using multiple threads. The set of entities to index is split into multiple partitions. Each thread processes one partition at a time.

The following section will explain how to tune the parallel execution.



The "sweet spot" of number of threads, fetch size, partition size, etc. to achieve best performance is highly dependent on your overall architecture, database design and even data values.

You should experiment with these settings to find out what's best in your particular case.

Threads

The maximum number of threads used by the job execution is defined through method `maxThreads()`. Within the N threads given, there's 1 thread reserved for the core, so only N - 1 threads are available for different partitions. If N = 1, the program will work, and all batch elements will run in the same thread. The default number of threads used in Hibernate Search is 10. You can overwrite it with your preferred number.

```
MassIndexingJob.parameters()  
    .maxThreads( 5 )  
    ...
```



Note that the batch runtime cannot guarantee the requested number of threads are available, it will use as many as possible up to the requested maximum (JSR352 v1.0 Final Release, page 34). Note also that all batch jobs share the same thread pool, so it's not always a good idea to execute jobs concurrently.

Rows per partition

Each partition consists of a fixed number of elements to index. You may tune exactly how many elements a partition will hold with `rowsPerPartition`.

```
MassIndexingJob.parameters()  
    .rowsPerPartition( 5000 )  
    ...
```



This property has **nothing** to do with "chunk size", which is how many elements are processed together between each write. That aspect of processing is addressed by chunking.

Instead, `rowsPerPartition` is more about how parallel your mass indexing job will be.

Please see the [Chunking section](#) to see how to tune chunking.

When `rowsPerPartition` is low, there will be many small partitions, so processing threads will be less likely to starve (stay idle because there's no more partition to process), but on the other hand you will only be able to take advantage of a small fetch size, which will increase the number of database accesses. Also, due to the failure recovery mechanisms, there is some overhead in starting a new partition, so with an unnecessarily large number of partitions, this overhead will add up.

When `rowsPerPartition` is high, there will be a few big partitions, so you will be able to take advantage of a higher `chunk size`, and thus a higher fetch size, which will reduce the number of database accesses, and the overhead of starting a new partition will be less noticeable, but on the other hand you may not use all the threads available.



Each partition deals with one root entity type, so two different entity types will never run under the same partition.

6.4.4. Chunking and session clearing

The mass indexing job supports restart a suspended or failed job more or less from where it stopped.

This is made possible by splitting each partition in several consecutive *chunks* of entities, and saving process information in a *checkpoint* at the end of each chunk. When a job is restarted, it will resume from the last checkpoint.

The size of each chunk is determined by the `checkpointInterval` parameter.

```
MassIndexingJob.parameters()  
    .checkpointInterval( 1000 )  
    ...
```

But the size of a chunk is not only about saving progress, it is also about performance:

- a new Hibernate session is opened for each chunk;
- a new transaction is started for each chunk;
- inside a chunk, the session is cleared periodically according to the `sessionClearInterval` parameter, which must thereby be smaller than (or equal to) the chunk size;
- documents are flushed to the index at the end of each chunk.

In general the checkpoint interval should be small compared to the number of rows per partition.

Indeed, due to the failure recovery mechanism, the elements before the first checkpoint of each partition will take longer to process than the other, so in a 1000-element partition, having a 100-element checkpoint interval will be faster than having a 1000-element checkpoint interval.



On the other hand, **chunks shouldn't be too small** in absolute terms. Performing a checkpoint means your JSR-352 runtime will write information about the progress of the job execution to its persistent storage, which also has a cost. Also, a new transaction and session are created for each chunk which doesn't come for free, and implies that setting the fetch size to a value higher than the chunk size is pointless. Finally, the index flush performed at the end of each chunk is an expensive operation that involves a global lock, which essentially means that the less you do it, the faster indexing will be. Thus having a 1-element checkpoint interval is definitely not a good idea.

6.4.5. Selecting the persistence unit (EntityManagerFactory)



Regardless of how the entity manager factory is retrieved, you must make sure that the entity manager factory used by the mass indexer will stay open during the whole mass indexing process.

JBeret

If your JSR-352 runtime is JBeret (used in WildFly in particular), you can use CDI to retrieve the `EntityManagerFactory`. Unless you use an already packaged `hibernate-search-jsr352` module for your application container, this will require you to add the `hibernate-search-jsr352-jberet` jar to your classpath.

If you use only one persistence unit, the mass indexer will be able to access your database automatically without any special configuration.

If you want to use multiple persistence units, you will have to register the `EntityManagerFactories` as beans in the CDI context. Note that entity manager factories will probably not be considered as beans by default, in which case you will have to register them yourself. You may use an application-scoped bean to do so:

```
@ApplicationScoped
public class EntityManagerFactoriesProducer {

    @PersistenceUnit(unitName = "db1")
    private EntityManagerFactory db1Factory;

    @PersistenceUnit(unitName = "db2")
    private EntityManagerFactory db2Factory;

    @Produces
    @Singleton
    @Named("db1") // The name to use when referencing the bean
    public EntityManagerFactory createEntityManagerFactoryForDb1() {
        return db1Factory;
    }

    @Produces
    @Singleton
    @Named("db2") // The name to use when referencing the bean
    public EntityManagerFactory createEntityManagerFactoryForDb2() {
        return db2Factory;
    }
}
```

Once the entity manager factories are registered in the CDI context, you can instruct the mass indexer to use one in particular by naming it using the `entityManagerReference` parameter.



Due to limitations of the CDI APIs, it is not currently possible to reference an entity manager factory by its persistence unit name when using the mass indexer with CDI.

Other DI-enabled JSR-352 implementations

If you want to use a different JSR-352 implementation that happens to allow dependency injection, you can use `hibernate-search-jsr352-core` under the following conditions:

1. You must map the following two scope annotations to the relevant scope in the dependency injection mechanism:
 - `org.hibernate.search.jsr352.inject.scope.spi.HibernateSearchJobScoped`
 - `org.hibernate.search.jsr352.inject.scope.spi.HibernateSearchPartitionScoped`
2. You must make sure that the dependency injection mechanism will register all injection-annotated classes (`@Named`, ...) from the `hibernate-search-jsr352-core` module in the dependency injection context. For instance this can be achieved in Spring DI using the `@ComponentScan` annotation.
3. You must register a single bean in the dependency injection context that will implement the `EntityManagerFactoryRegistry` interface.

Plain Java environment (no dependency injection at all)

The following will work only if your JSR-352 runtime does not support dependency injection at all, i.e. it ignores `@Inject` annotations in batch artifacts. This is the case for JBatch in Java SE mode, for instance.

If you use only one persistence unit, the mass indexer will be able to access your database automatically without any special configuration: you only have to make sure to create the `EntityManagerFactory` (or `SessionFactory`) in your application before launching the mass indexer.

If you want to use multiple persistence units, you will have to add two parameters when launching the mass indexer:

- `entityManagerFactoryReference`: this is the string that will identify the `EntityManagerFactory`.
- `entityManagerFactoryNamespace`: this allows to select how you want to reference the `EntityManagerFactory`. Possible values are:
 - `persistence-unit-name` (the default): use the persistence unit name defined in `persistence.xml`.

- `session-factory-name`: use the session factory name defined in the Hibernate configuration by the `hibernate.session_factory_name` configuration property.



If you set the `hibernate.session_factory_name` property in the Hibernate configuration and you don't use JNDI, you will also have to set `hibernate.session_factory_name_is_jndi` to `false`.

Chapter 7. Index Optimization

This section explains some low level tricks to keep your indexes at peak performance. We cover some Lucene details which in most cases you don't have to know about: Hibernate Search will handle these operations optimally and transparently in most cases without the need for further configuration. Still, it is good to know that there are ways to configure the behavior, if the need arises.

The index is physically stored in several smaller segments. Each segment is immutable and represents a generation of index writes. Index segments are periodically compacted, both to merge smaller segments and to remove stale entries; this merging process happens constantly in the background and can be tuned with the options specified in [Tuning indexing performance](#), but you can also define policies to fully run index optimizations when it is most suited for your specific workload.

With older versions of Lucene it was important to frequently optimize the index to maintain good performance, but with current Lucene versions this doesn't apply anymore. The benefit of explicit optimization is very low, and in certain cases even counter-productive. During an explicit optimization the whole index is processed and rewritten inflicting a significant performance cost. Optimization is for this reason a double-edged sword.

Another reason to avoid optimizing the index too often is that an optimization will, as a side effect, invalidate cached filters and field caches and internal buffers need to be refreshed.



Optimizing the index is often not needed, does not benefit write (update) performance at all, and is a slow operation: make sure you need it before activating it.

Of course optimizing the index does not only present drawbacks: after the optimization process is completed and new IndexReader instances have loaded their buffers, queries will perform at peak performance and you will have reclaimed all disk space potentially used by stale entries.

It is recommended to not schedule any optimization, but if you wish to perform it periodically you should run it:

- on an idle system or when the searches are less frequent
- after a lot of index modifications

When using a MassIndexer (see [Using a MassIndexer](#)) it will optimize involved indexes by default at the start and at the end of processing; you can change this behavior by using `MassIndexer.optimizeAfterPurge` and `MassIndexer.optimizeOnFinish` respectively. The initial optimization is actually very cheap as it is performed on an empty index: its purpose is to release the storage space occupied by the old index.

7.1. Automatic optimization

While in most cases this is not needed, Hibernate Search can automatically optimize an index after:

- a certain amount of write operations
- or after a certain amount of transactions

The configuration for automatic index optimization can be defined on a global level or per index:

Example 136. Defining automatic optimization parameters

```
hibernate.search.default.optimizer.operation_limit.max = 1000
hibernate.search.default.optimizer.transaction_limit.max = 100
hibernate.search.Animal.optimizer.transaction_limit.max = 50
```

With the above example an optimization will be triggered to the **Animal** index as soon as either:

- the number of additions and deletions reaches 1000
- the number of transactions reaches 50 (hibernate.search.Animal.optimizer.transaction_limit.max having priority over hibernate.search.default.optimizer.transaction_limit.max)

If none of these parameters are defined, no optimization is processed automatically.

The default implementation of `OptimizerStrategy` can be overridden by implementing `org.hibernate.search.store.optimization.OptimizerStrategy` and setting the `optimizer.implementation` property to the fully qualified name of your implementation. This implementation must implement the interface, be a public class and have a public constructor taking no arguments.

Example 137. Loading a custom OptimizerStrategy

```
hibernate.search.default.optimizer.implementation = com.acme.worlddomination.SmartOptimizer
hibernate.search.default.optimizer.SomeOption = CustomConfigurationValue
hibernate.search.humans.optimizer.implementation = default
```

The keyword `default` can be used to select the Hibernate Search default implementation; all properties after the `.optimizer` key separator will be passed to the implementation's `initialize` method at start.

7.2. Manual optimization

You can programmatically optimize (defragment) a Lucene index from Hibernate Search through the `SearchFactory`:

Example 138. Programmatic index optimization

```
FullTextSession fullTextSession = Search.getFullTextSession(regularSession);
SearchFactory searchFactory = fullTextSession.getSearchFactory();

searchFactory.optimize(Order.class);
// or
searchFactory.optimize();
```

The first example optimizes the Lucene index holding Orders; the second, optimizes all indexes.



`searchFactory.optimize()` has no effect on a JMS or JGroups backend: you must apply the optimize operation on the Master node.

7.3. Adjusting optimization

The Lucene index is constantly being merged in the background to keep a good balance between write and read performance; in a sense this is a form of background optimization which is always applied.

The following match attributes of Lucene's IndexWriter and are commonly used to tune how often merging occurs and how aggressive it is applied. They are exposed by Hibernate Search via:

- `hibernate.search.[default|<indexname>].indexwriter.max_buffered_docs`
- `hibernate.search.[default|<indexname>].indexwriter.max_merge_docs`
- `hibernate.search.[default|<indexname>].indexwriter.merge_factor`
- `hibernate.search.[default|<indexname>].indexwriter.ram_buffer_size`

See [Tuning indexing performance](#) for a description of these properties.

Chapter 8. Monitoring

Hibernate Search offers access to a `Statistics` object via `SearchFactory.getStatistics()`. It allows you for example to determine which classes are indexed and how many entities are in the index. This information is always available. However, by specifying the `hibernate.search.generate_statistics` property in your configuration you can also collect total and average Lucene query and object loading timings.

8.1. JMX



Support for statistics retrieved through JMX and indexing triggered through JMX is deprecated and will be removed in Hibernate Search 6.

You can also enable access to the statistics via JMX. Setting the property `hibernate.search.jmx_enabled` will automatically register the `StatisticsInfoMBean`. Depending on your the configuration the `IndexControlMBean` and `IndexingProgressMonitorMBean` will also be registered. In case you are having more than one JMX enabled Hibernate Search instance running within a single JVM, you should also set `hibernate.search.jmx_bean_suffix` to a different value for each of the instances. The specified suffix will be used to distinguish between the different MBean instances. Let's have a closer look at the mentioned MBeans.



If you want to access your JMX beans remotely via JConsole make sure to set the system property `com.sun.management.jmxremote` to true.

8.1.1. StatisticsInfoMBean

This MBean gives you access to Statistics object as described in the previous section.

8.1.2. IndexControlMBean

This MBean allows to build, optimize and purge the index for a given entity. Indexing occurs via the mass indexing API (see [Using a MassIndexer](#)). A requirement for this bean to be registered in JMX is, that the Hibernate SessionFactory is bound to JNDI via the `hibernate.session_factory_name` property. Refer to the Hibernate Core manual for more information on how to configure JNDI.

8.1.3. IndexingProgressMonitorMBean

This MBean is an implementation `MassIndexerProgressMonitor` interface. If `hibernate.search.jmx_enabled` is enabled and the mass indexer API is used the indexing progress can be followed via this bean. The bean will only be bound to JMX while indexing is in progress. Once indexing is completed the MBean is not longer available.

Chapter 9. Spatial

With the spatial extensions you can combine full-text queries with distance restrictions, filter results based on distances or sort results on such a distance criteria.

The spatial support of Hibernate Search has the following goals:

- Enable spatial search on entities: find entities within x km from a given location (latitude, longitude) on Earth
- Provide an easy way to enable spatial indexing via expressive annotations
- Provide a simple way for querying
- Hide geographical complexity

For example, you might search for restaurants somewhere in a 2 km radius around your office.

In order to use the spatial extensions for an indexed entity, you need to add the `@Spatial` annotation (`org.hibernate.search.annotations.Spatial`) and specify one or more sets of coordinates.

9.1. Enable indexing of Spatial Coordinates

There are different techniques to index point coordinates. Hibernate Search Spatial offers a choice between two strategies:

- index as numbers
- index as labeled spatial hashes

We will now describe both methods, so you can make a suitable choice. You can pick a different strategy for each set of coordinates. The strategy is selected by specifying the `spatialMode` attribute of the `@Spatial` annotation.

9.1.1. Indexing coordinates for range queries

When setting the `@Spatial.spatialMode` attribute to `SpatialMode.RANGE` (which is the default) coordinates are indexed as numeric fields, so that range queries can be performed to narrow down the initial area of interest.

Pros:

- Is quick on small data sets (< 100k entities)
- Is very simple: straightforward to debug/analyze
- Impact on index size is moderate

Cons:

- Poor performance on large data sets
- Poor performance if your data set is distributed across the whole world (for example when indexing points of interest in the United States, in Europe and in Asia, large areas collide because they share the same latitude. The latitude range query returns large amounts of data that need to be cross checked with those returned by the longitude range).

To index your entities for range querying you have to:

- add the `@Spatial` annotation on your entity
- add the `@Latitude` and `@Longitude` annotations on your properties representing the coordinates; these must be of type `Double`

Example 139. Sample Spatial indexing: Hotel class

```
import org.hibernate.search.annotations.*;

@Entity
@Indexed
@Spatial
public class Hotel {

    @Latitude
    Double latitude

    @Longitude
    Double longitude

    // ...
}
```

9.1.2. Indexing coordinates in a grid with spatial hashes

When setting `@Spatial.spatialMode` to `SpatialMode.HASH` the coordinates are encoded in several fields representing different zoom levels. Each box for each level is labeled so coordinates are assigned matching labels for each zoom level. This results in a grid encoding of labels called **spatial hashes**.

Pros :

- Good performance even with large data sets
- World wide data distribution independent

Cons :

- Index size is larger: need to encode multiple labels per pair of coordinates

To index your entities you have to:

- add the `@Spatial` annotation on the entity with the `SpatialMode` set to `GRID` :
`@Spatial(spatialMode = SpatialMode.HASH)`
- add the `@Latitude` and `@Longitude` annotations on the properties representing your coordinates; these must be of type `Double`

Example 140. Indexing coordinates in a grid using spatial hashes

```
@Spatial(spatialMode = SpatialMode.HASH)
@Indexed
@Entity
public class Hotel {

    @Latitude
    Double latitude;

    @Longitude
    Double longitude;

    // ...
}
```

9.1.3. Implementing the Coordinates interface

Instead of using the `@Latitude` and `@Longitude` annotations you can choose to implement the `org.hibernate.search.spatial.Coordinates` interface.

Example 141. Implementing the Coordinates interface

```
import org.hibernate.search.annotations.*;
import org.hibernate.search.spatial.Coordinates;

@Entity
@Indexed
@Spatial
public class Song implements Coordinates {

    @Id long id;
    double latitude;
    double longitude;
    // ...

    @Override
    Double getLatitude() {
        return latitude;
    }

    @Override
    Double getLongitude() {
        return longitude;
    }

    // ...
}
```

As we will see in the section [Multiple Coordinate pairs](#), an entity can have multiple `@Spatial` annotations; when having the entity implement `Coordinates`, the implemented methods refer to the

default `@Spatial` annotation with the default pair of coordinates.



The default (field) name in case `@Spatial` is placed on the entity level is `org.hibernate.search.annotations.Spatial.COORDINATES_DEFAULT_FIELD`.

An alternative is to use properties implementing the `Coordinates` interface; this way you can have multiple `Spatial` instances:

Example 142. Using attributes of type `Coordinates`

```
@Entity
@Indexed
public class Event {
    @Id
    Integer id;

    @Field(store = Store.YES)
    String name;

    double latitude;
    double longitude;

    @Spatial(spatialMode = SpatialMode.HASH)
    public Coordinates getLocation() {
        return new Coordinates() {
            @Override
            public Double getLatitude() {
                return latitude;
            }

            @Override
            public Double getLongitude() {
                return longitude;
            }
        };
    }
}

// ...
```

When using this form the `@Spatial.name` automatically defaults to the property name. In the above case to `location`.

9.2. Performing Spatial Queries

You can use the Hibernate Search query DSL to build a query to search around a pair of coordinates (latitude, longitude) or around a bean implementing the `Coordinates` interface.

As with any full-text query, the spatial query creation flow looks like:

1. retrieve a `QueryBuilder` from the `SearchFactory`
2. use the DSL to build a spatial query, defining search center and radius

3. optionally combine the resulting **Query** with other filters
4. call the **createFullTextQuery()** and use the resulting query like any standard Hibernate or JPA query

Example 143. Search for an Hotel by distance

```
QueryBuilder builder = fullTextSession.getSearchFactory()
    .buildQueryBuilder().forEntity( Hotel.class ).get();

org.apache.lucene.search.Query luceneQuery = builder
    .spatial()
    .within( radius, Unit.KM )
    .ofLatitude( centerLatitude )
    .andLongitude( centerLongitude )
    .createQuery();

org.hibernate.Query hibQuery = fullTextSession
    .createFullTextQuery( luceneQuery, Hotel.class );
List results = hibQuery.list();
```



In the above example we did not explicitly specify the field name to use. The default coordinates field name was used implicitly. To target an alternative pair of coordinates at query time, we need to specify the field name as well. See [Multiple Coordinate pairs](#).

A fully working example can be found in the test-suite of the [source code](#). Refer to **SpatialIndexingTest.testSpatialAnnotationOnClassLevel()** and its corresponding **Hotel** test class.

Alternatively to passing separate latitude and longitude values, you can also pass an instance implementing the **Coordinates** interface:

Example 144. DSL example with Coordinates

```
Coordinates coordinates = Point.fromDegrees(24d, 31.5d);
Query query = builder
    .spatial()
    .within( 51, Unit.KM )
    .ofCoordinates( coordinates )
    .createQuery();

List results = fullTextSession.createFullTextQuery( query, POI.class ).list();
```

9.2.1. Returning distance to query point in the search results

Returning distance to the center in the results

To retrieve the actual distance values (in kilometers) you need to use projection (see [Projection](#)):

Example 145. Distance projection example

```
double centerLatitude = 24.0d;
double centerLongitude= 32.0d;

QueryBuilder builder = fullTextSession.getSearchFactory()
    .buildQueryBuilder().forEntity(POI.class).get();
org.apache.lucene.search.Query luceneQuery = builder
    .spatial()
        .onField("location")
        .within(100, Unit.KM)
        .ofLatitude(centerLatitude)
        .andLongitude(centerLongitude)
    .createQuery();

FullTextQuery hibQuery = fullTextSession.createFullTextQuery(luceneQuery, POI.class);
hibQuery.setProjection(FullTextQuery.SPATIAL_DISTANCE, FullTextQuery.THIS);
hibQuery.setSpatialParameters(centerLatitude, centerLongitude, "location");
List results = hibQuery.list();
```

- Use `FullTextQuery.setProjection` with `FullTextQuery.SPATIAL_DISTANCE` as one of the projected fields.
- Call `FullTextQuery.setSpatialParameters` with the latitude, longitude and the name of the spatial field used to build the spatial query. Note that using coordinates different than the center used for the query will have unexpected results.



The default (field) name in case `@Spatial` is placed on the entity level is `org.hibernate.search.annotations.Spatial.COORDINATES_DEFAULT_FIELD`.



Distance projection and null values

When a spatial field on an entity has a `null` value for either its latitude or longitude (or both), the resulting projected distance will always be `null`.

Sorting by distance

To sort the results by distance to the center of the search you will have to build a `Sort` instance using Hibernate Search sort DSL:

Example 146. Distance sort example using the sort DSL

```
double centerLatitude = 24.0d;
double centerLongitude = 32.0d;

QueryBuilder builder = fullTextSession.getSearchFactory()
    .buildQueryBuilder().forEntity( POI.class ).get();
org.apache.lucene.search.Query luceneQuery = builder
    .spatial()
    .onField("location")
    .within(100, Unit.KM)
    .ofLatitude(centerLatitude)
    .andLongitude(centerLongitude)
    .createQuery();

FullTextQuery hibQuery = fullTextSession.createFullTextQuery(luceneQuery, POI.class);
List results = query.list();
Sort distanceSort = qb
    .sort()
    .byDistance()
    .onField("location")
    .fromLatitude(centerLatitude)
    .andLongitude(centerLongitude)
    .createSort();
hibQuery.setSort(distanceSort);
```

The sort must be constructed using the same coordinates on the same spatial field used to build the spatial query, otherwise the sorting will occur with another center than the query. This repetition is needed to allow you to define Queries with any tool.

Sorting and null values

When a spatial field on an entity has a **null** value for either its latitude or longitude (or both):



- if you are filtering the results using a distance query, the entity with missing coordinates will not appear in the query results and its rank in the sort is irrelevant.
- otherwise, the resulting distance will always be the greatest possible value, which means the entity will be ranked last if the sort is ascending, or first if the sort is descending.

Alternatively, you may also use a **DistanceSortField** directly, as it was done before the introduction of Hibernate Search sort DSL:

Example 147. Distance sort example without using the sort DSL

```
double centerLatitude = 24.0d;
double centerLongitude = 32.0d;

QueryBuilder builder = fullTextSession.getSearchFactory()
    .buildQueryBuilder().forEntity( POI.class ).get();
org.apache.lucene.search.Query luceneQuery = builder
    .spatial()
    .onField("location")
    .within(100, Unit.KM)
    .ofLatitude(centerLatitude)
    .andLongitude(centerLongitude)
    .createQuery();

FullTextQuery hibQuery = fullTextSession.createFullTextQuery(luceneQuery, POI.class);
Sort distanceSort = new Sort(
    new DistanceSortField(centerLatitude, centerLongitude, "location"));
hibQuery.setSort(distanceSort);
```

9.3. Multiple Coordinate pairs

You can associate multiple pairs of coordinates to the same entity, as long as each pair is uniquely identified by using a different name. This is achieved by stacking multiple `@Spatial` annotations and specifying the `name` attribute on the individual `@Spatial` annotation. At most one of them will be allowed to use the default name.

Example 148. Multiple sets of coordinates

```
import org.hibernate.search.annotations.*;

@Entity
@Indexed
@Spatial
@Spatial(name="work", spatialMode = SpatialMode.HASH)
public class UserEx {

    @Id
    Integer id;

    @Latitude
    Double homeLatitude;

    @Longitude
    Double homeLongitude;

    @Latitude(of="work")
    Double workLatitude;

    @Longitude(of="work")
    Double workLongitude;
```

To target an alternative pair of coordinates at query time, we need to specify the pair by name using `onField(String)`:

```
QueryBuilder builder = fullTextSession.getSearchFactory()
    .buildQueryBuilder().forEntity( UserEx.class ).get();

org.apache.lucene.search.Query luceneQuery = builder
    .spatial()
    .onField( "work" )
    .within( radius, Unit.KM )
    .ofLatitude( centerLatitude )
    .andLongitude( centerLongitude )
    .createQuery();

org.hibernate.Query hibQuery = fullTextSession.createFullTextQuery( luceneQuery,
    Hotel.class );
List results = hibQuery.list();
```

9.4. Insight: implementation details of spatial hashes indexing

The following chapter is meant to provide a technical insight in spatial hash (grid) indexing. It discusses how coordinates are mapped to the index and how queries are implemented.

9.4.1. At indexing level

When Hibernate Search indexes an entity annotated with `@Spatial`, it instantiates a `SpatialFieldBridge` to transform the latitude and longitude fields accessed via the `Coordinates` interface to the multiple index fields stored in the Lucene index.

Principle of the spatial index: the spatial index used in Hibernate Search is a [grid based spatial index](#) where grid ids are hashes derived from latitude and longitude.

To make computations easier the latitude and longitude field values will be projected into a flat coordinate system with the help of a [sinusoidal projection](#). Origin value space is :

$[-90 \rightarrow +90], [-180 \rightarrow +180]$

for latitude, longitude coordinates and projected space is:

$[-\pi \rightarrow +\pi], [-\pi/2 \rightarrow +\pi/2]$

for Cartesian x,y coordinates (beware of fields order inversion: x is longitude and y is latitude).

The index is divided into n levels labeled from 0 to n-1.

At the level 0 the projected space is the whole Earth. At the level 1 the projected space is divided into 4 rectangles (called boxes as in bounding box):

$[-\pi, -\pi/2] \rightarrow [0, 0]$, $[-\pi, 0] \rightarrow [0, +\pi/2]$, $[0, -\pi/2] \rightarrow [+ \pi, 0]$ and $[0, 0] \rightarrow [+ \pi, +\pi/2]$

At level $n+1$ each box of level n is divided into 4 new boxes and so on. The numbers of boxes at a given level is 4^n .

Each box is given an id, in this format: [Box index on the X axis][Box index on the Y axis]. To calculate the index of a box on an axis we divide the axis range in 2^n slots and find the slot the box belongs to. At the n level the indexes on an axis are from $-(2^n)/2$ to $(2^n)/2$. For instance, the 5th level has $4^5 = 1024$ boxes with 32 indexes on each axis (32×32 is 1024) and the box of Id "0|8" is covering the $[0, 8/32 \times \pi/2] \rightarrow [1/32 \times \pi, 9/32 \times \pi/2]$ rectangle in projected space.

Beware! The boxes are rectangles in projected space but the related area on Earth is not rectangular!

Now that we have all these boxes at all these levels, we index points "into" them.

For a point (lat, long) we calculate its projection (x, y) and then we calculate for each level of the spatial index, the ids of the boxes it belongs to.

At each level the point is in one and only one box. For points on the edges the boxes are considered exclusive on the left side and inclusive on the right i.e. $[start, end]$ (the points are normalized before projection to $[-90, +90], [-180, +180]$).

We store in the Lucene document corresponding to the entity to index one field for each level of the spatial hash grid. The field is named: *HSSI*[n]. [spatial index fields name] is given either by the parameter at class level annotation or derived from the name of the spatial annotated method of the entity, HSSI stands for Hibernate Search Spatial Index and n is the level of the spatial hashes grid.

We also store the latitude and longitude as a numeric field under [spatial index fields name]_HSSI_Latitude and [spatial index fields name]_HSSI_Longitude fields. They will be used to filter precisely results by distance in the second stage of the search.

9.4.2. At search level

Now that we have all these fields, what are they used for?

When you ask for a spatial search by providing a search discus (center+radius) we will calculate the box ids that do cover the search discus in the projected space, fetch all the documents that belong to these boxes (thus narrowing the number of documents for which we will have to calculate distance to the center) and then filter this subset with a real distance calculation. This is called two level spatial filtering.

Step 1: Compute the best spatial hashes grid level for the search discus

For a given search radius there is an optimal hash grid level where the number of boxes to retrieve shall be minimal without bringing back too many documents (level 0 has only 1 box but retrieve all documents). The optimal hash grid level is the maximum level where the width of each box is larger

than the search area. Near the equator line where projection deformation is minimal, this will lead to the retrieval of at most 4 boxes. Towards the poles where the deformation is more significant, it might need to examine more boxes but as the sinusoidal projection has a simple Tissot's indicatrix (see [Sinusoidal projection](#)) in populated areas, the overhead is minimal.

Step 2: Compute ids of the corresponding covering boxes at that level

Now that we have chosen the optimal level, we can compute the ids of the boxes covering the search discus (which is not a discus in projected space anymore).

This is done by `org.hibernate.search.spatial.impl.SpatialHelper.getSpatialHashCellsIds(Point center, double radius, int spatialHashLevel)`

It will calculate the bounding box of the search discus and then call `org.hibernate.search.spatial.impl.SpatialHelper.getSpatialHashCellsIds(Point lowerLeft, Point upperRight, int spatialHashLevel)` that will do the actual computation. If the bounding box crosses the meridian line it will cut the search in two and make two calls to `getSpatialHashCellsIds(Point lowerLeft, Point upperRight, int spatialHashLevel)` with left and right parts of the box.

There are some geo related hacks (search radius too large, search radius crossing the poles) that are handled in bounding box computations done by `Rectangle.fromBoundingCircle(Coordinates center, double radius)` (see <http://janmatuschek.de/LatitudeLongitudeBoundingCoordinates> for reference on those subjects).

The `SpatialHelper.getSpatialHashCellsIds(Point lowerLeft, Point upperRight, int spatialHashLevel)` project the defining points of the bounding box and compute the boxes they belong to. It returns all the box ids between the lower left to the upper right corners, thus covering the area.

Step 3: Lucene index lookup

The query is built with theses ids searching for documents having a `HSSI[n]` (n the level found at Step 1) field valued with one of the ids of Step 2.

See also the implementation of `org.hibernate.search.spatial.impl.SpatialHashFilter`.

This query will return all documents in the boxes covering the projected bounding box of the search discus. So it is too large and needs refining. But we have narrowed the distance calculation problems to a subset of our data.

Step 4: Refine

A distance calculation filter is set after the Lucene index lookup query of Step 3 to exclude false candidates from the result list.

See `SpatialQueryBuilderFromCoordinates.buildSpatialQuery(Coordinates center, double radius, String fieldName)`

Chapter 10. Advanced features

In this final chapter we are offering a smörgåsbord of tips and tricks which might become useful as you dive deeper and deeper into Hibernate Search.

10.1. Accessing the `SearchFactory`

The `SearchFactory` object keeps track of the underlying Lucene resources for Hibernate Search. It is a convenient way to access Lucene natively. The `SearchFactory` can be accessed from a `FullTextSession`:

Example 150. Accessing the `SearchFactory`

```
FullTextSession fullTextSession = Search.getFullTextSession(regularSession);
SearchFactory searchFactory = fullTextSession.getSearchFactory();
```

10.2. Accessing the `SearchIntegrator`

The interface `SearchIntegrator` gives access to advanced functionality and internals of Hibernate Search; these are typically only useful to integrate Hibernate Search with other libraries. These are called SPI - Service Provide Interface - to better differentiate them from normal APIs. You can recognize them by noticing an `spi` element in the package name, such as in `org.hibernate.search.spi.SearchIntegrator`.

You can access the `SearchIntegrator` SPI using the `SearchFactory` ([Accessing the `SearchFactory`](#)), or by extracting the instance from the Hibernate native `SessionFactory` or the JPA `EntityManagerFactory` as in the following examples:

Example 151. Accessing the `SearchIntegrator` from a `SearchFactory`

```
SearchIntegrator si = searchFactory.unwrap(SearchIntegrator.class);
```

Example 152. Accessing the `SearchIntegrator` from a `SessionFactory`

```
SearchIntegrator si = org.hibernate.search.orm.spi.SearchIntegratorHelper
    .extractFromSessionFactory( sf );
```

Example 153. Accessing the `SearchIntegrator` from an `EntityManagerFactory`

```
SearchIntegrator si = org.hibernate.search.orm.spi.SearchIntegratorHelper
    .extractFromEntityManagerFactory( emf );
```

10.3. Using an `IndexReader`

Queries in Lucene are executed on an `IndexReader`. Hibernate Search caches index readers to maximize performance and implements other strategies to retrieve updated `IndexReaders` in order to minimize IO operations. Your code can access these cached resources, but you have to follow some "good citizen" rules.

Example 154. Accessing an `IndexReader`

```
IndexReader reader = searchFactory.getIndexReaderAccessor().open(Order.class);
try {
    //perform read-only operations on the reader
}
finally {
    searchFactory.getIndexReaderAccessor().close(reader);
}
```

In this example the `SearchFactory` figures out which indexes are needed to query this entity. Using the configured `ReaderProvider` (described in [Reader strategy](#)) on each index, it returns a compound `IndexReader` on top of all involved indexes. Because this `IndexReader` is shared amongst several clients, you must adhere to the following rules:

- Never call `indexReader.close()`, but always call `readerProvider.closeReader(reader)`, using a finally block.
- Don't use this `IndexReader` for modification operations: it's a read-only instance, you would get an exception.

Aside from those rules, you can use the `IndexReader` freely, especially to do native Lucene queries. Using this shared `IndexReaders` will be more efficient than by opening one directly from - for example - the filesystem.

As an alternative to the method `open(Class... types)` you can use `open(String... indexNames)` in this case you pass in one or more index names; using this strategy you can also select a subset of the indexes for any indexed type if sharding is used.

Example 155. Accessing an `IndexReader` by index names

```
IndexReader reader = searchFactory
    .getIndexReaderAccessor()
    .open("Products.1", "Products.3");
```

10.4. Accessing a Lucene Directory

A Directory is the most common abstraction used by Lucene to represent the index storage; Hibernate Search doesn't interact directly with a Lucene Directory but abstracts these interactions via an `IndexManager`: an index does not necessarily need to be implemented by a Directory.

If you are certain that your index is represented as a Directory and need to access it, you can get a reference to the Directory via the `IndexManager`. You will have to cast the `IndexManager` instance to a `DirectoryBasedIndexManager` and then use `getDirectoryProvider().getDirectory()` to get a reference to the underlying Directory. This is not recommended, if you need low level access to the index using Lucene APIs we suggest to see [Using an IndexReader](#) instead.

10.5. Sharding indexes

In some cases it can be useful to split (shard) the data into several Lucene indexes. There are two main use cases:

- A single index is so big that index update times are slowing the application down. In this case static sharding can be used to split the data into a pre-defined number of shards.
- Data is naturally segmented by customer, region, language or other application parameter and the index should be split according to these segments. This is a use case for dynamic sharding.



By default sharding is not enabled.

10.5.1. Static sharding

To enable static sharding set the `hibernate.search.<indexName>.sharding_strategy.nbr_of_shards` property as seen in [Enabling index sharding](#).

Example 156. Enabling index sharding

```
hibernate.search.[default|<indexName>].sharding_strategy.nbr_of_shards = 5
```

The default sharding strategy which gets enabled by setting this property, splits the data according to the hash value of the document id (generated by the `FieldBridge`). This ensures a fairly balanced

sharding. You can replace the default strategy by implementing a custom `IndexShardingStrategy`. To use your custom strategy you have to set the `hibernate.search.[default|<indexName>].sharding_strategy` property to the fully qualified class name of your custom `IndexShardingStrategy`.

Example 157. Registering a custom `IndexShardingStrategy`

```
hibernate.search.[default|<indexName>].sharding_strategy = my.custom.RandomShardingStrategy
```

10.5.2. Dynamic sharding



Support for dynamic sharding is deprecated, and will be removed in Hibernate Search 6.

To implement a `ShardIdentifierProvider` in a future-proof way, make sure that all the shard identifiers are known on startup, and that no shard identifiers are added or removed at runtime. Only then will it be possible to migrate your implementation to Hibernate Search 6.

Dynamic sharding allows you to manage the shards yourself and even create new shards on the fly. To do so you need to implement the interface `ShardIdentifierProvider` and set the `hibernate.search.[default|<indexName>].sharding_strategy` property to the fully qualified name of this class. Note that instead of implementing the interface directly, you should rather derive your implementation from `org.hibernate.search.store.ShardIdentifierProviderTemplate` which provides a basic implementation. Let's look at [Custom `ShardIdentifierProvider`](#) for an example.

Example 158. Custom ShardIdentifierProvider

```
public static class AnimalShardIdentifierProvider extends ShardIdentifierProviderTemplate {

    @Override
    public String getShardIdentifier(Class<?> entityType, Serializable id,
        String idAsString, Document document) {
        if (entityType.equals(Animal.class)) {
            String typeValue = document.getField("type").stringValue();
            addShard(typeValue);
            return typeValue;
        }
        throw new RuntimeException("Animal expected but found " + entityType);
    }

    @Override
    protected Set<String> loadInitialShardNames(Properties properties, BuildContext
buildContext) {
        ServiceManager serviceManager = buildContext.getServiceManager();
        SessionFactory sessionFactory = serviceManager.requestService(
            HibernateSessionFactoryService.class).getSessionFactory();
        Session session = sessionFactory.openSession();
        try {
            Criteria initialShardsCriteria = session.createCriteria(Animal.class);
            initialShardsCriteria.setProjection(Projections.distinct(Property.forName("type")));
            List<String> initialTypes = initialShardsCriteria.list();
            return new HashSet<String>(initialTypes);
        }
        finally {
            session.close();
        }
    }
}
```

There are several things happening in `AnimalShardIdentifierProvider`. First off its purpose is to create one shard per animal type (e.g. mammal, insect, etc.). It does so by inspecting the class type and the Lucene document passed to the `getShardIdentifier()` method. It extracts the type field from the document and uses it as shard name. `getShardIdentifier()` is called for every addition to the index and a new shard will be created with every new animal type encountered. The base class `ShardIdentifierProviderTemplate` maintains a set with all known shards to which any identifier must be added by calling `addShard()`.

It is important to understand that Hibernate Search cannot know which shards already exist when the application starts. When using `ShardIdentifierProviderTemplate` as base class of a `ShardIdentifierProvider` implementation, the initial set of shard identifiers must be returned by the `loadInitialShardNames()` method. How this is done will depend on the use case. However, a common case in combination with Hibernate ORM is that the initial shard set is defined by the distinct values of a given database column. `CustomShardIdentifierProvider` shows how to handle such a case. `AnimalShardIdentifierProvider` makes in its `loadInitialShardNames()` implementation use of a service called `HibernateSessionFactoryService` (see also [Using external services](#)) which is available within an ORM environment. It allows to request a Hibernate `SessionFactory` instance which can be used to run a Criteria query in order to determine the initial set of shard identifiers.

Last but not least, the `ShardIdentifierProvider` also allows for optimizing searches by selecting

which shard to run a query against. By activating a filter (see [Using filters in a sharded environment](#)), a sharding strategy can select a subset of the shards used to answer a query (`getShardIdentifiersForQuery()`, not shown in the example) and thus speed up the query execution.



This `ShardIdentifierProvider` is considered experimental. We might need to apply some changes to the defined method signatures to accommodate for unforeseen use cases. Please provide feedback if you have ideas, or just to let us know how you're using this API.

10.6. Sharing indexes

It is technically possible to store the information of more than one entity into a single Lucene index. There are two ways to accomplish this:



Support for indexing multiple entity types in the same index is deprecated and is going to be removed in Hibernate Search 6.

Support for the `hibernate.search.[fully qualified entity name].indexName` configuration property is deprecated and is going to be removed in Hibernate Search 6.

- Configuring the underlying directory providers to point to the same physical index directory. In practice, you set the property `hibernate.search.[fully qualified entity name].indexName` to the same value. As an example, let's use the same index (directory) for the `Furniture` and `Animal` entities. We just set `indexName` for both entities to "Animal". Both entities will then be stored in the Animal directory:

```
hibernate.search.org.hibernate.search.test.shards.Furniture.indexName = Animal
hibernate.search.org.hibernate.search.test.shards.Animal.indexName = Animal
```

- Setting the `@Indexed` annotation's `index` attribute of the entities you want to merge to the same value. If we again wanted all `Furniture` instances to be indexed in the `Animal` index along with all instances of `Animal` we would specify `@Indexed(index="Animal")` on both `Animal` and `Furniture` classes.



This is only presented here so that you know the option is available. There is really not much benefit in sharing indexes.

10.7. Using external services

A `Service` in Hibernate Search is a class implementing the interface `org.hibernate.search.engine.service.spi.Service` and providing a default no-arg constructor. Theoretically that's all that is needed to request a given service type from the Hibernate Search `ServiceManager`. In practice you probably want to add some service life cycle methods (implement `Startable` and `Stoppable`) as well as actual methods providing some functionality.

Hibernate Search uses the service approach to decouple different components of the system. Let's have a closer look at services and how they are used.

10.7.1. Using a Service

Many of the pluggable contracts of Hibernate Search can use services. Services are accessible via the `BuildContext` interface as in the following example.

Example 159. Example of a custom `DirectoryProvider` using a `ClassLoaderService`

```
public CustomDirectoryProvider implements DirectoryProvider<RAMDirectory> {
    private ServiceManager serviceManager;
    private ClassLoaderService classLoaderService;

    public void initialize(
        String directoryProviderName,
        Properties properties,
        BuildContext context) {
        //get a reference to the ServiceManager
        this.serviceManager = context.getServiceManager();
    }

    public void start() {
        //get the current ClassLoaderService
        classLoaderService = serviceManager.requestService(ClassLoaderService.class);
    }

    public RAMDirectory getDirectory() {
        //use the ClassLoaderService
    }

    public stop() {
        //make sure to release all services
        serviceManager.releaseService(ClassLoaderService.class);
    }
}
```

When you request a service, an instance of the requested service type is returned to you. Make sure release the service via `ServiceManager.releaseService` once you don't need it anymore. Note that the service can be released in the `DirectoryProvider.stop` method if the `DirectoryProvider` uses the service during its lifetime or could be released right away if the service is only needed during initialization time.

10.7.2. Implementing a Service

To implement a service, you need to create an interface which identifies it and extends `org.hibernate.search.engine.service.spi.Service`. You can then add additional methods to your service interface as needed.

Naturally you will also need to provide an implementation of your service interface. This implementation must have a public no-arg constructor. Optionally your service can also implement the life cycle methods `org.hibernate.search.engine.service.spi.Startable` and/or `org.hibernate.search.engine.service.spi.Stoppable`. These methods will be called by the `ServiceManager` when the service is created respectively the last reference to a requested service is released.

Services are retrieved from the `ServiceManager.requestService` using the `Class` object of the interface you define as a key.

Managed services

To transparently discover services Hibernate Search uses the Java ServiceLoader mechanism. This means you need to add a service file to your jar under `/META-INF/services/` named after the fully qualified classname of your service interface. The content of the file contains the fully qualified classname of your service implementation.

Example 160. Service file for the Infinispan CacheManagerService service

```
/META-INF/services/org.infinispan.hibernate.search.spi.CacheManagerService
```

Example 161. Content of META-INF/services/org.infinispan.hibernate.search.spi.CacheManagerService

```
org.infinispan.hibernate.search.impl.DefaultCacheManagerService
```



Hibernate Search only supports a single service implementation of a given service. There is no mechanism to select between multiple versions of a service. It is an error to have multiple jars defining each a different implementation for the same service. If you want to override the implementation of a already existing service at runtime you will need to look at [Provided services](#).

Provided services



Provided services are usually used by frameworks integrating with Hibernate Search and not by library users themselves.

As an alternative to manages services, a service can be provided by the environment bootstrapping Hibernate Search. For example, Infinispan which uses Hibernate Search as its internal search engine, passes the `CacheContainer` to Hibernate Search. In this case, the `CacheContainer` instance is not managed by Hibernate Search and the start/stop methods defined by optional `Stoppable` and `Startable` interfaces will be ignored.

A Service implementation which is only used as a Provided Service doesn't need to have a public constructor taking no arguments.



Provided services have priority over managed services. If a provided service is registered with the same `ServiceManager` instance as a managed service, the provided service will be used.

The provided services are passed to Hibernate Search via the `SearchConfiguration` interface: as implementor of method `getProvidedServices` you can return a `Map` of all services you need to provide.



When implementing a custom `org.hibernate.search.cfg.spi.SearchConfiguration` we recommend you extend the base class `org.hibernate.search.cfg.spi.SearchConfigurationBase`: that will improve compatibility by not breaking your code when we need to add new methods to this interface.

10.8. Customizing Lucene's scoring formula

Lucene allows the user to customize its scoring formula by extending `org.apache.lucene.search.similarities.Similarity`. The abstract methods defined in this class match the factors of the following formula calculating the score of query q for document d :

$$\text{score}(q,d) = \text{coord}(q,d) \cdot \text{queryNorm}(q) \cdot \sum_{t \sim q} (\text{tf}(t \text{ in } d) \cdot \text{idf}(t)^2 \cdot t.\text{getBoost()} \cdot \text{norm}(t,d))$$

Factor	Description
<code>tf(t ind)</code>	Term frequency factor for the term (t) in the document (d).
<code>idf(t)</code>	Inverse document frequency of the term.
<code>coord(q,d)</code>	Score factor based on how many of the query terms are found in the specified document.
<code>queryNorm(q)</code>	Normalizing factor used to make scores between queries comparable.
<code>t.getBoost()</code>	Field boost.
<code>norm(t,d)</code>	Encapsulates a few (indexing time) boost and length factors.

It is beyond the scope of this manual to explain this formula in more detail. Please refer to Similarity's Javadocs for more information.

Hibernate Search provides two ways to modify Lucene's similarity calculation.

First you can set the default similarity by specifying the fully specified classname of your Similarity implementation using the property `hibernate.search.similarity`. The default value is `org.apache.lucene.search.similarities.ClassicSimilarity`.

Secondly, you can override the similarity used for a specific index by setting the `similarity` property for this index (see [Directory configuration](#) for more information about index configuration):

```
hibernate.search.[default|<indexname>].similarity = my.custom.Similarity
```

As an example, let's assume it is not important how often a term appears in a document. Documents with a single occurrence of the term should be scored the same as documents with multiple occurrences. In this case your custom implementation of the method `tf(float freq)` should return 1.0.



When two entities share the same index they must declare the same Similarity implementation.

10.9. Multi-tenancy

10.9.1. What is multi-tenancy?

The term multi-tenancy in general is applied to software development to indicate an architecture in which a single running instance of an application simultaneously serves multiple clients (tenants). Isolating information (data, customizations, etc) pertaining to the various tenants is a particular challenge in these systems. This includes the data owned by each tenant stored in the database. You will find more details on how to enable multi-tenancy with Hibernate in the [Hibernate ORM developer's guide](#).

10.9.2. Using a tenant-aware `FullTextSession`

Hibernate Search supports multi-tenancy on top of Hibernate ORM, it stores the tenant identifier in the document and automatically filters the query results.

The `FullTextSession` will be bound to the specific tenant ("client-A" in the example) and the mass indexer will only index the entities associated to that tenant identifier.

Example 162. Bind the session to a tenant

```
Session session = getSessionFactory()
    .withOptions()
    .tenantIdentifier( "client-A" )
    .openSession();

FullTextSession session = Search.getFullTextSession( session );
```

The use of a tenant identifier will have the following effects:

1. Every document saved or updated in the index will have an additional field `__HSearch_TenantId` containing the tenant identifier.
2. Every search will be filtered using the tenant identifier.
3. The MassIndexer (see [Using a MassIndexer](#)) will only affect the currently selected tenant.

Note that not using a tenant will return all the matching results for all the tenants in the index.

Chapter 11. Integration with Elasticsearch

11.1. Status



This feature is a work in progress. Make sure to read the [Limitations](#) section!

The integration with Elasticsearch is in development and should be considered experimental. We do think we have the basics covered and we are looking for feedback.

Patches can be sent as pull requests to the [Github repository](#), but also general feedback, suggestions and questions are very welcome. To get in touch or find other interesting links for contributors, see the [Community page](#) of the Hibernate website.

11.2. Goal of the Elasticsearch integration

The goal of integrating with Elasticsearch is to allow Hibernate Search users to benefit from the full-text capabilities integrated with Hibernate ORM but replacing the local Lucene based index with a remote Elasticsearch service.

There could be various reasons to prefer this over an "embedded Lucene" approach:

- wish to separate the service running your application from the Search service (Microservices)
- benefit from Elasticsearch's out of the box horizontal scalability and high availability features
- integrate with an existing Elasticsearch cluster
- explore the data updated by an Hibernate powered application using the Elasticsearch dashboard integrations such as Kibana

There are a couple of drawbacks compared to the embedded Lucene approach though:

- incur a performance penalty of remote RPCs both for index updates and to run queries
- need to manage an additional service
- possibly need to buy and manage additional servers

Which solution is best will depend on the specific needs of your system and your organization.

Why not use Elasticsearch directly?

The #1 reason is that Hibernate Search integrates perfectly with Hibernate ORM. All changes done to your objects will trigger the necessary index changes transparently.



- it will honor the transaction boundary - i.e. not do the indexing work if the transaction ends up in rollback
- changes to cascaded objects are handled
- changes to nested object embedded in a root indexed entity are handled
- changes will be sent in bulk - i.e. optimized systematically for you
- it generates JSON encoded requests from your entities, parses responses
- manages the connection pools
- etc.

There is no more paradigm shift in your code. You are working on Hibernate ORM managed objects, doing your queries on object properties with a nice Java DSL.

11.3. Getting started and configuration

To experiment with the Elasticsearch integration you will have to download Elasticsearch and run it: Hibernate Search connects to an Elasticsearch node but does not provide one.

One option is to use the [Elasticsearch Docker image](#).

11.3.1. Elasticsearch version

Hibernate Search expects an Elasticsearch cluster running version 2.x or 5.x. The version running on your cluster will be automatically detected on startup, and Hibernate Search will adapt based on the detected version.

The targeted version is mostly transparent to Hibernate Search users, but there are a few differences in how Hibernate Search behaves depending on the Elasticsearch version that may affect you. The following table details those differences.

	2.x	5.0/5.1	5.2+
Configuration required for purges	Enable the delete-by-query plugin	None	

	2.x	5.0/5.1	5.2+
Datatype used for String fields in Elasticsearch	<code>string</code>	<code>text</code> (if analyzed) or <code>keyword</code> (if not). The <code>string</code> datatype has been deprecated in Elasticsearch 5.0.	<code>text</code> (if analyzed with an <code>analyzer</code>) or <code>keyword</code> (if not analyzed or analyzed with a <code>normalizer</code>).
Norms	Not implemented	Implemented	
Implementation of <code>@Field.indexNullAs</code> for analyzed text fields	<code>null_value</code> is added to the mapping, <code>null</code> values are indexed as such.	<code>null_value</code> is not available for <code>text</code> fields, null values are replaced with the <code>indexNullAs</code> value explicitly when indexing. This does not apply to <code>keyword</code> fields.	
Configuration required for spatial projections without a spatial sort	Enable and configure the groovy plugin	None	



The testsuite of Hibernate Search runs against Elasticsearch 5.6.16 by default.

Upgrading Elasticsearch

When upgrading your Elasticsearch cluster, some [administrative tasks](#) are still required on your cluster: Hibernate Search will not take care of those.

On top of that, there are some fundamental differences between some versions of Elasticsearch, and for that reason, some upgrades will require you to delete your indexes manually and [reindex your data](#).

The upgrades that require such changes are the following:

- upgrading from 2.x to 5.x, because Hibernate Search does not support the `string` datatype on Elasticsearch 5.x;
- upgrading from 5.0/5.1 to 5.2+ if you used [normalizers](#), because Hibernate Search normalizers are mapped to native Elasticsearch normalizers starting from Elasticsearch 5.2, thereby changing the datatype of fields using normalizers from `text` to `keyword`.

11.3.2. Dependencies in your Java application

In addition to the usual dependencies like Hibernate ORM and Hibernate Search, you will need the new `hibernate-search-elasticsearch` jar.

Example 163. Maven dependencies for Hibernate Search with Elasticsearch

```
<dependency>
  <groupId>org.hibernate</groupId>
  <artifactId>hibernate-search-elasticsearch</artifactId>
  <version>5.10.13.Final</version>
</dependency>
```

If you use Amazon's proprietary IAM authentication through request signing to access Elasticsearch, you will need an additional dependency to handle authentication: `hibernate-search-elasticsearch-aws`.

Example 164. Maven dependencies for Hibernate Search with AWS-hosted Elasticsearch

```
<dependency>
  <groupId>org.hibernate</groupId>
  <artifactId>hibernate-search-elasticsearch-aws</artifactId>
  <version>5.10.13.Final</version>
</dependency>
```

11.3.3. Elasticsearch configuration

Hibernate Search can work with an Elasticsearch server without altering its configuration.

However some features offered by Hibernate Search require specific configuration:

- on Elasticsearch 2.x only (not necessary on 5.x): if you want to be able to use the Hibernate Search `MassIndexer` with `purgeAllOnStart` enabled - it is enabled by default -, or to use `FullTextSession.purge()` or `FullTextSession.purgeAll()`, install the `delete-by-query` plugin
- on Elasticsearch 2.x only (not necessary on 5.x): if you want to retrieve the distance in a geolocation query, enable the `lang-groovy` plugin, see [Elasticsearch Spatial queries](#)
- if you want to use paging (as opposed to `scrolling`) on result sets larger than 10000 elements (for instance access the 10001st result), you may increase the value of the `index.max_result_window` property (default is 10000).

11.3.4. Hibernate Search configuration

Configuration is minimal. Add the configuration properties to your `persistence.xml` or where you put the rest of your Hibernate Search configuration.

Select Elasticsearch as the backend

```
hibernate.search.default.indexmanager elasticsearch
```

Hostname and port for Elasticsearch

```
hibernate.search.default.elasticsearch.host http://127.0.0.1:9200 (default)
```

You may also select multiple hosts (separated by whitespace characters), so that they are assigned requests in turns (load balancing):

```
hibernate.search.default.elasticsearch.host http://es1.mycompany.com:9200
http://es2.mycompany.com:9200
```

In the example above, the first request will go to `es1`, the second to `es2`, the third to `es1`, and so on.

Also note having multiple hosts will enable failover: if one node happens to fail to serve a request (timeout, server error, invalid HTTP response, ...), the same request will be sent to the next one; if the second request is served without error, the failure will be blamed on the node and no error will be reported to the application.

The failover feature will also be enabled when you only have one configured host but other hosts have been added thanks to automatic discovery (see below).

Optional path prefix of the Elasticsearch instance

```
hibernate.search.default.elasticsearch.path_prefix /my/path
```

In a typical Elasticsearch setup, it should not be set.

Use the path prefix if your Elasticsearch instance is located at a specific context path i.e. is, for instance, hosted at <http://localhost/my/elasticsearch/>.

The path prefix setting is global so your Elasticsearch cluster configuration must be consistent in this matter.

Username for Elasticsearch connection

```
hibernate.search.default.elasticsearch.username ironman
```

 (default is empty, meaning anonymous access)

Password for Elasticsearch connection

```
hibernate.search.default.elasticsearch.password j@rV1s
```

 (default is empty)

If you use HTTP instead of HTTPS in any of the Elasticsearch host URLs (see above), your password will be transmitted in clear text over the network.

Select the index creation strategy

```
hibernate.search.default.elasticsearch.index_schema_management_strategy  
CREATE
```

 (default)

Let's see the options for the `index_schema_management_strategy` property:

Value	Definition
<code>none</code>	The index, its mappings and the analyzer definitions will not be created, deleted nor altered. Hibernate Search will not even check that the index already exists.

Value	Definition
<code>validate</code>	The index, its existing mappings and analyzer definitions will be checked to be compatible with the metamodel of your application. The index, its mappings and analyzer definitions will not be created, deleted nor altered.
<code>update</code>	The index, its mappings and analyzer definitions will be created, existing mappings will be updated if there are no conflicts. Caution: if analyzer definitions have to be updated, the index will be closed automatically during the update.
<code>create</code>	The default: an existing index will not be altered, a missing index will be created along with their mappings and analyzer definitions.
<code>drop-and-create</code>	Indexes will be deleted if existing and then created along with their mappings and analyzer definitions to match the requirements of your application. This will delete all content from the indexes! Useful during development.
<code>drop-and-create-and-drop</code>	Similar to <code>drop-and-create</code> but will also delete the index at shutdown. Commonly used for tests.



Since Elasticsearch on Amazon Web Services (AWS) [does not support the `\_close/\_open` operations](#), the `update` strategy will fail when trying to update analyzer definitions on an AWS Elasticsearch cluster.

The only workaround is to avoid the `update` strategy on AWS.

Strategies in production environments

It is strongly recommended to use either `none` or `validate` in a production environment.

The alternatives `drop-and-create` and `drop-and-create-and-drop` are obviously unsuitable in this context unless you want to reindex everything upon every startup, and `update` may leave your mapping half-updated in case of conflict.



To be precise, if your mapping changed in an incompatible way, such as a field having its type changed, updating the mapping may be impossible without manual intervention. In this case, the `update` strategy will prevent Hibernate Search from starting, but it may already have successfully updated the mappings for another index, making a rollback difficult.

When updating analyzer definitions Hibernate Search will temporarily stop the affected indexes during the update. This means the `update` strategy should be used with caution when multiple clients use Elasticsearch indexes managed by Hibernate Search: those clients should be synchronized in such a way that while Hibernate Search is starting, no other client needs to access the index.

For these reasons, migrating your mapping on a live cluster should be properly planned within the deployment process.



Mapping validation is as permissive as possible. Fields or mappings that are unknown to Hibernate Search will be ignored, and settings that are more powerful than required (e.g. a field annotated with `@Field(index = Index.NO)` in Search but marked as `"index": analyzed` in Elasticsearch) will be deemed valid.

One exception: date formats must match exactly the formats specified by Hibernate Search, due to implementation constraints.

Maximum time to wait for the successful execution of a request to the Elasticsearch server before failing (in ms)

```
hibernate.search.default.elasticsearch.request_timeout 60000 (default)
```

The execution time of a request includes the time needed to establish a connection, to send the request, and to receive the whole response, optionally retrying in case of node failures.

Maximum time to wait for a connection to the Elasticsearch server before failing (in ms)

```
hibernate.search.default.elasticsearch.connection_timeout 3000 (default)
```

Maximum time to wait for a response from the Elasticsearch server before failing (in ms)

`hibernate.search.default.elasticsearch.read_timeout 60000` (default)

Maximum number of simultaneous connections to the Elasticsearch cluster

`hibernate.search.default.elasticsearch.max_total_connection 20` (default)

Maximum number of simultaneous connections to a single Elasticsearch server

`hibernate.search.default.elasticsearch.max_total_connection_per_route 2`
(default)

Whether to enable automatic discovery of servers in the Elasticsearch cluster (true or false)

`hibernate.search.default.elasticsearch.discovery.enabled false` (default)

When using automatic discovery, the Elasticsearch client will periodically probe for new nodes in the cluster, and will add those to the server list (see `host` above). Similarly, the client will periodically check whether registered servers still respond, and will remove them from the server list if they don't.

Time interval between two executions of the automatic discovery (in seconds)

`hibernate.search.default.elasticsearch.discovery.refresh_interval 10` (default)

This setting will only be taken into account if automatic discovery is enabled (see above).

Scheme to use when connecting to automatically discovered nodes (http or https)

`hibernate.search.default.elasticsearch.discovery.default_scheme http` (default)

This setting will only be taken into account if automatic discovery is enabled (see above).

Maximum time to wait for the indexes to become available before failing (in ms)

`hibernate.search.default.elasticsearch.index_management_wait_timeout 10000`
(default)

This setting is ignored when the `NONE` strategy is selected, since the index will not be checked on startup (see above).

This value must be lower than the read timeout (see above).

Status an index must at least have in order for Hibernate Search to work with it (one of "green", "yellow" or "red")

`hibernate.search.default.elasticsearch.required_index_status green` (default)

Only operate if the index is at this level or safer. In development, set this value to `yellow` if the number of nodes started is below the number of expected replicas.

Whether to perform an explicit refresh after a set of operations has been executed against a specific index (**true** or **false**)

```
hibernate.search.default.elasticsearch.refresh_after_write false (default)
```

This is useful in unit tests to ensure that a write is visible by a query immediately without delay. This keeps unit tests simpler. You should not rely on the synchronous behaviour for your production code except in rare cases as Elasticsearch is optimised for asynchronous writes: leave at **false** for optimal performance.

When [scrolling](#), the minimum number of previous results kept in memory at any time

```
hibernate.search.elasticsearch.scroll_backtracking_window_size 10000 (default)
```

When [scrolling](#), the number of results fetched by each Elasticsearch call

```
hibernate.search.elasticsearch.scroll_fetch_size 1000 (default)
```

When [scrolling](#), the maximum duration `ScrollableResults` will be usable if no other results are fetched from Elasticsearch, in seconds

```
hibernate.search.elasticsearch.scroll_timeout 60 (default)
```

The behavior with respect to dynamically added fields in the Elasticsearch mapping

```
hibernate.search.default.elasticsearch.dynamic_mapping strict (default)
```

Possible values are:

- **true**: Add unknown fields to the schema dynamically
- **false**: Ignore unknown fields
- **strict**: Throw an exception on unknown fields

This setting may be overridden in custom field bridges on a per-field basis: see [Dynamic mapping](#).

Whether JSON included in logs should be pretty-printed (indented, with line breaks) (**true** or **false**)

```
hibernate.search.elasticsearch.log.json_pretty_printing false (default)
```

Properties prefixed with `hibernate.search.default` can be given globally as shown above and/or be given for specific indexes:

```
hibernate.search.someindex.elasticsearch.index_schema_management_strategy update
```



This excludes properties related to the internal Elasticsearch client, which at the moment is common to every index manager (but this will change in a future version). Excluded properties are `host`, `username`, `password`, `read_timeout`, `connection_timeout`, `max_total_connection`, `max_total_connection_per_route`, `http_client_configurer`, `discovery.enabled`, `discovery.refresh_interval` and `discovery.scheme`.

Authentication on Amazon Web Services

Amazon's proprietary IAM authentication through request signing [requires additional dependencies](#). With these dependencies, you will be able to use the following configuration options.

Enable or disable AWS request signing

```
hibernate.search.default.elasticsearch.aws.signing.enabled false (default)
```

The value must be either `true` or `false`.

AWS access key

```
hibernate.search.default.elasticsearch.aws.access_key AKIDEXAMPLE
```

This property has no default and must be provided for the AWS authentication to work.

AWS secret key

```
hibernate.search.default.elasticsearch.aws.secret_key  
wJalrXUtnFEMI/K7MDENG+bPxRfiCYEXAMPLEKEY
```

This property has no default and must be provided for the AWS authentication to work.

AWS region

```
hibernate.search.default.elasticsearch.aws.region us-east-1
```

This property has no default and must be provided for the AWS authentication to work.

Should you need help with finding the correct values to use, please refer to the AWS documentation on [security credentials](#) and [regions](#).



Setting these AWS-specific configuration properties is not enough to enable the Elasticsearch integration: you should also have a look at the [main configuration properties](#), most notably `hibernate.search.default.indexmanager` and `hibernate.search.default.elasticsearch.host`.

11.4. Mapping and indexing

Like in Lucene embedded mode, indexes are transparently updated when you create or update entities mapped to Hibernate Search. Simply use familiar annotations from [Mapping entities to the index structure](#).

The name of the index will be the lowercased name provided to `@Indexed` (non qualified class name by default). Hibernate Search will map the fully qualified class name to the Elasticsearch type.

11.4.1. Annotation specificities

Field.indexNullAs

The `org.hibernate.search.annotations.Field` annotation allows you to provide a replacement value for null properties through the `indexNullAs` attribute (see [@Field](#)), but this value must be provided as a string.

In order for your value to be understood by Hibernate Search (and Elasticsearch), the provided string must follow one of those formats:

- For string values, no particular format is required.
- For numeric values, use formats accepted by `Double.parseDouble`, `Integer.parseInt`, etc., depending on the actual type of your field.
- For booleans, use either `true` or `false`.
- For dates (`java.util.Calendar`, `java.util.Date`, `java.time.*`), use the ISO-8601 format.

The full format is `yyyy-MM-dd'T'HH:mm:ss.nZ[ZZZ]` (for instance `2016-11-26T16:41:00.006+01:00[CET]`). Please keep in mind that part of this format must be left out depending on the type of your field, though. For a `java.time.LocalDateTime` field, for instance, the provided string must not include the zone offset (`+01:00`) or the zone ID (`[UTC]`), because those don't make sense.

Even when they make sense for the type of your field, the time and time zone may be omitted (if omitted, the time zone will be interpreted as the default JVM time zone).

Index-time boosting

The `org.hibernate.search.annotations.DynamicBoost` annotation is not (and cannot be) supported with Elasticsearch, because the platform lacks per-document, index-time boosting capabilities.

The `@Boost` annotation will work, but since it's deprecated we suggest to use query-time boosting instead: see [Querying](#).

11.4.2. Analyzers



Analyzers are treated differently than in Lucene embedded mode.

Built-in or server-defined analyzers

Using the `definition` attribute in the `@Analyzer` annotation, you can refer to the name of the built-in Elasticsearch analyzer, or custom analyzers already registered on your Elasticsearch instances.

More information on analyzers, in particular those already built in Elasticsearch, can be found in [the Elasticsearch documentation](#).

Example of custom analyzers defined in the `elasticsearch.yml`

```
# Custom analyzer
index.analysis:
  analyzer.custom-analyzer:
    type: custom
    tokenizer: standard
    filter: [custom-filter, lowercase]
  filter.custom-filter:
    type: stop
    stopwords: [test1, close]
```

From there, you can use the custom analyzers by name in your entity mappings.

Example of mapping that refers to custom and built-in analyzers on Elasticsearch

```
@Entity
@Indexed(index = "tweet")
public class Tweet {

    @Id
    @GeneratedValue
    private Integer id;

    @Field(analyzer = @Analyzer(definition = "english")) // Elasticsearch built-in analyzer
    private String englishTweet;

    @Field(analyzer = @Analyzer(definition = "whitespace")) // Elasticsearch built-in analyzer
    private String whitespaceTweet;

    @Field(name = "tweetNotAnalyzed", analyzer = Analyze.NO, store = Store.YES)
    // Custom analyzer:
    @Field(
        name = "tweetWithCustom",
        analyzer = @Analyzer(definition = "custom-analyzer")
    )
    private String multipleTweets;
}
```

You may also reference a built-in Lucene analyzer implementation using the `@Analyzer.impl` attribute: Hibernate Search will translate the implementation to an equivalent Elasticsearch built-in type, if possible.



Support for `@Analyzer.impl/@Normalizer.impl` is deprecated, and will be removed in Hibernate Search 6.

Use `@Analyzer.definition` instead, as explained above.



Using the `@Analyzer.impl` attribute is not recommended with Elasticsearch because it will never allow you to take full advantage of Elasticsearch analysis capabilities. You cannot, for instance, use custom analyzer implementations: only built-in Lucene implementations are supported.

This feature is designed to help migrating an application that already used Hibernate Search, moving from an embedded Lucene instance to an Elasticsearch cluster.

Example of mapping that refers to a built-in analyzer on Elasticsearch using a Lucene implementation class

```
@Entity
@Indexed(index = "tweet")
public class Tweet {

    @Id
    @GeneratedValue
    private Integer id;

    @Field(analyzer = @Analyzer(impl = EnglishAnalyzer.class)) // Elasticsearch built-in
    "english" analyzer
    private String englishTweet;

    @Field(analyzer = @Analyzer(impl = WhitespaceAnalyzer.class)) // Elasticsearch built-in
    "whitespace" analyzer
    private String whitespaceTweet;

}
```

Custom analyzers using a definition provider

As with [Lucene](#), you can define analyzers programmatically.

In the case of Elasticsearch, this requires to set the `hibernate.search.elasticsearch.analysis_definition_provider` configuration property.

The property can be set to the fully-qualified name of a class with a public, no-arg constructor in your application. This class must either implement `org.hibernate.search.elasticsearch.analyzer.definition.ElasticsearchAnalysisDefinitionProvider` directly or expose a `@Factory`-annotated method that returns such implementation.

Example of mapping that defines analyzers on Elasticsearch using the `Elasticsearch*Factory` types

```
public static class CustomAnalyzerProvider implements ElasticsearchAnalysisDefinitionProvider
{
    @Override
    public void register(ElasticsearchAnalysisDefinitionRegistryBuilder builder) {
        builder.analyzer( "tweet_analyzer" )
            .withTokenizer( "whitespace" )
            .withCharFilters( "custom_html_strip" )
            .withCharFilters( "p_br_as_space" );

        builder.charFilter( "custom_html_strip" )
            .type( "html_strip" )
            .param( "escaped_tags", "br", "p" );

        builder.charFilter( "p_br_as_space" )
            .type( "pattern_replace" )
            .param( "pattern", "<p/?>|<br/?>" )
            .param( "replacement", " " )
            .param( "tags", "CASE_INSENSITIVE" );
    }
}
```

Custom analyzers using the `@AnalyzerDef` annotation



Support for defining analyzers or normalizers through annotations is deprecated, and will be removed in Hibernate Search 6.

Use [a definition provider](#) instead.

You can also define analyzers within your Hibernate Search mapping using the `@AnalyzerDef` annotation, like you would [do with an embedded Lucene instance](#). When Hibernate Search creates the Elasticsearch indexes, the relevant definitions will then be automatically added as a [custom analyzer](#) in the index settings.

Two different approaches allow you to define your analyzers with Elasticsearch.

The first, recommended approach is to use the factories provided by the `hibernate-search-elasticsearch` module:

- `org.hibernate.search.elasticsearch.analyzer.ElasticsearchCharFilterFactory`
- `org.hibernate.search.elasticsearch.analyzer.ElasticsearchTokenFilterFactory`
- `org.hibernate.search.elasticsearch.analyzer.ElasticsearchTokenizerFactory`

Those classes can be passed to the `factory` attribute of the `@CharFilterDef`, `@TokenFilterDef` and `@TokenizerDef` annotations.

The `params` attribute may be used to define the `type` parameter and any other parameter accepted by Elasticsearch for this type.

The parameter values will be interpreted as JSON. The parser is not strict, though:

- quotes around strings may be left out in some cases, as when a string only contains letters.
- when quotes are required (e.g. your string may be interpreted as a number, and you don't want that), you may use single quotes instead of double quotes (which are painful to write in Java).



You may use the `name` attribute of the `@CharFilterDef`, `@TokenFilterDef` and `@TokenizerDef` annotations to define the exact name to give to that definition in the Elasticsearch settings.

Example of mapping that defines analyzers on Elasticsearch using the `Elasticsearch*Factory` types

```
@Entity
@Indexed(index = "tweet")
@AnalyzerDef(
    name = "tweet_analyzer",
    charFilters = {
        @CharFilterDef(
            name = "custom_html_strip",
            factory = ElasticsearchCharFilterFactory.class,
            params = {
                @Parameter(name = "type", value = "'html_strip'"),
                // One can use Json arrays
                @Parameter(name = "escaped_tags", value = "['br', 'p']")
            }
        ),
        @CharFilterDef(
            name = "p_br_as_space",
            factory = ElasticsearchCharFilterFactory.class,
            params = {
                @Parameter(name = "type", value = "'pattern_replace'"),
                @Parameter(name = "pattern", value = "'<p/?>|<br/?>'"),
                @Parameter(name = "replacement", value = "' '"),
                @Parameter(name = "tags", value = "'CASE_INSENSITIVE'")
            }
        )
    },
    tokenizer = @TokenizerDef(
        factory = ElasticsearchTokenizerFactory.class,
        params = {
            @Parameter(name = "type", value = "'whitespace'"),
        }
    )
)
public class Tweet {

    @Id
    @GeneratedValue
    private Integer id;

    @Field(analyzer = @Analyzer(definition = "tweet_analyzer"))
    private String content;
}
```

The second approach is to configure everything as if you were using Lucene: use the Lucene factories, their parameter names, and format the parameter values as required in Lucene. Hibernate Search will automatically convert these definitions to the Elasticsearch equivalent.

Referencing Lucene factories is not recommended with Elasticsearch because it will never allow you to take full advantage of Elasticsearch analysis capabilities.

Here are the known limitations of the automatic translation:



- a few factories have unsupported parameters, because those have no equivalent in Elasticsearch. An exception will be raised on startup if a parameter is not supported.
- the `hyphenator` parameter for `HyphenatedWordsFilterFactory` must refer to a file on the Elasticsearch servers, on the contrary to other factories where the files are accessed by Hibernate Search directly. This is due to an Elasticsearch limitation (there is no way to forward the content of a local hyphenation pattern file).
- some built-in Lucene factories are not (and cannot) be translated, because of incompatible parameters between the Lucene factory and the Elasticsearch equivalent. This is in particular the case for `HunspellStemFilterFactory`.

Therefore, Lucene factories should only be referenced within analyzer definitions when migrating an application that already used Hibernate Search, moving from an embedded Lucene instance to an Elasticsearch cluster.

```
@Entity
@Indexed(index = "tweet")
@AnalyzerDef(
    name = "tweet_analyzer",
    charFilters = {
        @CharFilterDef(
            name = "custom_html_strip",
            factory = HTMLStripCharFilterFactory.class,
            params = {
                @Parameter(name = "escapedTags", value = "br,p")
            }
        ),
        @CharFilterDef(
            name = "p_br_as_space",
            factory = PatternReplaceCharFilterFactory.class,
            params = {
                @Parameter(name = "pattern", value = "<p/?>|<br/?>"),
                @Parameter(name = "replacement", value = " ")
            }
        )
    },
    tokenizer = @TokenizerDef(
        factory = WhitespaceTokenizerFactory.class
    )
)
public class Tweet {

    @Id
    @GeneratedValue
    private Integer id;

    @Field(analyzer = @Analyzer(definition = "tweet_analyzer"))
    private String content;
}
```

Normalizers

Depending on the Elasticsearch version you're using, [normalizers](#) will be handled differently:

- On version 5.2 and above, Elasticsearch supports normalizers natively and thus Hibernate Search normalizers will be translated directly to Elasticsearch normalizers.
- On version 5.1 and below, Elasticsearch does not support normalizers natively, thus Hibernate Search normalizers will be translated to Elasticsearch analyzers with a keyword tokenizer.



On Elasticsearch 5.1 and below, the fact that normalizers are translated to Elasticsearch analyzers means in particular that no normalizer definition should have the same name as an analyzer definition.

All the features mentioned above for analyzers are also available when using normalizers:

- They can be defined using an `ElasticsearchAnalysisDefinitionProvider` (use `builder.normalizer(name)`).

- They can be defined using Lucene factories with `@NormalizerDef`, which will be translated to the equivalent types in Elasticsearch, though `@NormalizerDef` is deprecated (use the analysis definition provider instead).
- They can be defined using the special factories `ElasticsearchCharFilterFactory` and `ElasticsearchTokenFilterFactory` with `@NormalizerDef`, though `@NormalizerDef` is deprecated (use the analysis definition provider instead).
- The definitions can be referenced using `@Field(normalizer = @Normalizer(definition = "myDefinition"))`

11.4.3. Field bridges

Built-in field bridges

The default field bridges are the same as those mentioned in the [main documentation](#), with the following exceptions.



You can refer to the Elasticsearch documentation for a description of [date formats](#).

`boolean`, `Boolean`

are converted to native JSON booleans.

`Date`, `Calendar`

are converted to the Elasticsearch `date` type with the default format `(strict_date_optional_time|epoch_millis)`.

`Instant`

are converted to the Elasticsearch `date` type with the default format `(strict_date_optional_time|epoch_millis)`.

`LocalDate`

are converted to the Elasticsearch `date` type with format `strict_date|yyyyyyyyyy-MM-dd`.

`LocalDateTime`

are converted to the Elasticsearch `date` type with format `strict_date_hour_minute_second_fraction|yyyyyyyyyy-MM-dd'T'HH:mm:ss.SSSSSSSS`.

`LocalTime`

are converted to the Elasticsearch `date` type with format `strict_hour_minute_second_fraction`.

`OffsetDateTime`

are converted to the Elasticsearch `date` type with format `strict_date_time|yyyyyyyyyy-MM-`

`dd'T'HH:mm:ss.SSSSSSSSZ.`

OffsetTime

are converted to the Elasticsearch `date` type with format `strict_time`.

ZonedDateTime

are converted to the Elasticsearch `date` type with format `yyyy-MM-dd'T'HH:mm:ss.SSSZZ['ZZZ']'|'yyyyyyyy-MM-dd'T'HH:mm:ss.SSSSSSSSZ['ZZZ']'`.

Year

are converted to the Elasticsearch `date` type with format `strict_year'|'yyyyyyyyyy`.

YearMonth

are converted to the Elasticsearch `date` type with format `strict_year_month'|'yyyyyyyyyy-MM`.

MonthDay

are converted to the Elasticsearch `date` type with format `--MM-dd`.

Custom field bridges

You can write custom field bridges and class bridges. For class bridges and field bridges creating multiple fields, make sure to make your bridge implementation also implement the `MetadataProvidingFieldBridge` contract.

```
public static class FirstAndMiddleNamesFieldBridge implements MetadataProvidingFieldBridge {
    @Override
    public void set(String name, Object value, Document document, LuceneOptions luceneOptions)
    {
        Explorer explorer = (Explorer) value;

        String firstName = explorer.getFirstName();
        luceneOptions.addFieldToDocument( name + "_firstName", firstName, document );
        document.add( new SortedDocValuesField( name + "_firstName", new BytesRef( firstName )
    ) );

        String middleName = explorer.getMiddleName();
        luceneOptions.addFieldToDocument( name + "_middleName", middleName, document );
        document.add( new SortedDocValuesField( name + "_middleName", new BytesRef( middleName
    ) ) );
    }

    @Override
    public void configureFieldMetadata(String name, FieldMetadataBuilder builder) {
        builder
            .field( name + "_firstName", FieldType.STRING )
            .sortable( true )
            .field( name + "_middleName", FieldType.STRING )
            .sortable( true );
    }
}
```



This interface and `FieldBridge` in general are likely going to evolve in the next major version of Hibernate Search to remove its adherence to Lucene specific classes like `Document`.

Sub-fields

When your `MetadataProvidingFieldBridge` registers a field whose name is the name of an existing field, with a dot and another string appended, like `name + ".mySubField"`, Hibernate Search will translate it as an object with a property `mySubField` in the JSON document sent to Elasticsearch.

As a result, sub-fields can only have a parent field with the `OBJECT` type: obviously, Elasticsearch would reject a String or an Integer with `mySubField` property. So every time a field named `foo.bar` is registered, its parent field `foo` must be registered with the `OBJECT` type, as in the following example. Not doing so will result in errors when Hibernate Search generates the Elasticsearch schema.

```
@Override
public void configureFieldMetadata(String name, FieldMetadataBuilder builder) {
    builder
        .field( name, FieldType.OBJECT )
        .field( name + ".firstName", FieldType.STRING )
        .sortable( true )
        .field( name + ".middleName", FieldType.STRING )
        .sortable( true );
}
```

Dynamic mapping

By default Hibernate Search (or, more accurately, the Elasticsearch schema generated by Hibernate Search) will not allow custom bridges to create fields that weren't previously registered through the `configureFieldMetadata` method.

If you don't know in advance the name of the fields created by your custom bridge, you can enable "dynamic mapping" globally or per-index through [a configuration option](#), or you can do it on a per-field basis using the `FieldMetadataCreationContext.mappedOn` method and the `org.hibernate.search.elasticsearch.bridge.spi.Elasticsearch` extension:

```
@Override
public void configureFieldMetadata(String name, FieldMetadataBuilder builder) {
    builder
        .field( name, FieldType.OBJECT )
        .mappedOn( Elasticsearch.class )
        .dynamic( DynamicType.TRUE ); // Will allow to index fields named <name>.foo,
    <name>.bar, etc.
}
```



Dynamic mapping involves Elasticsearch guessing the type and attributes of fields automatically, and has its own limitations. See [the Elasticsearch documentation](#) for more information.

11.4.4. Tika bridges

If your metadata processors create fields with a different name from the one passed as a parameter, make sure to make your processor also implement the `MetadataProvidingTikaMetadataProcessor` contract.

11.5. Queries

You can write queries like you usually do in Hibernate Search: native Lucene queries and DSL queries (see [Querying](#)). We do automatically translate the most common types of Apache Lucene queries and all queries generated by the Hibernate Search DSL except more like this (see below).

Unsupported Query DSL features

Queries written via the DSL work. Open a JIRA otherwise.



The notable exception is more like this queries. Hibernate Search has a more advanced algorithm than Lucene (or Elasticsearch/Solr) which is not easily portable with what Elasticsearch exposes.

If you need this feature, contact us.

On top of translating Lucene queries, you can directly create Elasticsearch queries by using either its String format or a JSON format:

Example 165. Creating an Elasticsearch native query from a string

```
FullTextEntityManager fullTextEm = org.hibernate.search.jpa.Search.  
getFullTextEntityManager(entityManager);  
QueryDescriptor query = ElasticsearchQueries.fromQueryString("title:tales");  
List<?> result = fullTextEm.createFullTextQuery(query, ComicBook.class).getResultList();
```

Example 166. Creating an Elasticsearch native query from JSON

```
FullTextEntityManager fullTextEm = org.hibernate.search.jpa.Search.  
getFullTextEntityManager(entityManager);  
QueryDescriptor query = ElasticsearchQueries.fromJson(  
    "{ 'query': { 'match' : { 'lastName' : 'Brand' } } }");  
List<?> result = fullTextEm.createFullTextQuery(query, GolfPlayer.class).getResultList();
```

Date/time in native Elasticsearch queries

By default Elasticsearch interprets the date/time strings lacking the time zone as if they were represented using the UTC time zone. If overlooked, this can cause your native Elasticsearch queries to be completely off.



The simplest way to avoid issues is to always explicitly provide time zone IDs or offsets when building native Elasticsearch queries. This may be achieved either by directly adding the time zone ID or offset in date strings, or by using the `time_zone` parameter (range queries only). See the Elasticsearch documentation for more information.

Simple query string queries and analyzers

The simple query string DSL is supported by the Elasticsearch integration.



In the general case, your query will automatically use the analyzers defined in the schema but be aware that, in the advanced case when you want to override the analyzer used, Elasticsearch only supports overriding the analyzer globally.

Thus, when overriding the analyzer, it must be consistent across all the queried fields.

11.5.1. Spatial queries

The Elasticsearch integration supports spatial queries by using either the DSL or native Elasticsearch queries.

For regular usage, there are no particular requirements for spatial support.



On Elasticsearch 2.x only (not necessary on 5.x), if you want to calculate the distance from your entities to a point **without sorting** by the distance to this point, you need to enable the Groovy plugin by adding the following snippet to your Elasticsearch configuration:

Enabling Groovy support in your `elasticsearch.yml`

```
script.engine.groovy.inline.search: true
```

11.5.2. Paging and scrolling

You may handle large result sets in two different ways, with different limitations.

For (relatively) smaller result sets, you may use the traditional offset/limit querying provided by the `FullTextQuery` interfaces: `setFirstResult(int)` and `setMaxResults(int)`. Limitations:

- This will only get you as far as the 10000 first documents, i.e. when requesting a window that includes documents beyond the 10000th result, Elasticsearch will return an error. If you want to raise this limit, see the `index.max_result_window` property in [Elasticsearch's settings](#).

If your result set is bigger, you may take advantage of scrolling by using the `scroll` method on `org.hibernate.search.FullTextQuery`. Limitations:

- This method is not available in `org.hibernate.search.jpa.FullTextQuery`.
- The Elasticsearch implementation has poor performance when an offset has been defined (i.e. `setFirstResult(int)` has been called on the query before calling `scroll()`). This is because Elasticsearch does not provide such feature, thus Hibernate Search has to scroll through every previous result under the hood.
- The Elasticsearch implementation allows only limited backtracking. Calling `scrollableResults.setRowNumber(4)` when currently positioned at index `1006`, for example, may result in a `SearchException` being thrown, because only 1000 previous elements had been kept in memory. You may work this around by tweaking the property: `hibernate.search.elasticsearch.scroll_backtracking_window_size` (see [Elasticsearch integration configuration](#)).
- The `ScrollableResults` will become stale and unusable after a given period of time spent without fetching results from Elasticsearch. You may work this around by tweaking two properties: `hibernate.search.elasticsearch.scroll_timeout` and `hibernate.search.elasticsearch.scroll_fetch_size` (see [Elasticsearch integration configuration](#)). Typically, you will solve timeout issues by reducing the fetch size and/or increasing the timeout limit, but this will also increase the performance hit on Elasticsearch.

11.5.3. Sorting

Sorting is performed the same way as [with the Lucene backend](#).

If you happen to need an advanced Elasticsearch sorting feature that is not natively supported in `SortField` or in Hibernate Search sort DSL, you may still create a sort from JSON, and even mix it with DSL-defined sorts:

Example 167. Mixing DSL-defined sorts with native Elasticsearch JSON sorts

```
FullTextEntityManager fullTextEm = org.hibernate.search.jpa.Search.  
getFullTextEntityManager(entityManager);  
org.hibernate.search.query.dsl.QueryBuilder queryBuilder =  
    fullTextEm.getSearchFactory()  
        .buildQueryBuilder().forEntity(Book.class).get();  
Query luceneQuery = /* ... use the query builder as explained in the Querying section ...  
*/;  
FullTextQuery query = fullTextEm.createFullTextQuery( luceneQuery, Book.class );  
Sort sort = queryBuilder.sort()  
    .byNative("authors.name", "{ 'order': 'asc', 'mode': 'min' }")  
    .andByField("title")  
    .createSort();  
query.setSort(sort);  
List results = query.getResultList();
```

11.5.4. Projections

All fields are stored by Elasticsearch in the JSON document it indexes, there is no specific need to mark fields as stored when you want to project them. The downside is that to project a field, Elasticsearch needs to read the whole JSON document. If you want to avoid that, use the `Store.YES` marker.

You can also retrieve the full JSON document by using `org.hibernate.search.elasticsearch.ElasticsearchProjectionConstants.SOURCE`.

```
FullTextEntityManager fullTextEm = org.hibernate.search.jpa.Search.getFullTextEntityManager  
(entityManager);  
FullTextQuery query = fullTextEm.createFullTextQuery(  
    qb.keyword()  
        .onField( "tags" )  
        .matching( "round-based" )  
        .createQuery(),  
    VideoGame.class  
)  
    .setProjection( ElasticsearchProjectionConstants.SCORE,  
        ElasticsearchProjectionConstants.SOURCE );  
Object[] projection = (Object[]) query.getSingleResult();
```

If you're looking for information about execution time, you may also use `org.hibernate.search.elasticsearch.ElasticsearchProjectionConstants.TOOK` and `org.hibernate.search.elasticsearch.ElasticsearchProjectionConstants.TIMED_OUT`:

```

FullTextEntityManager fullTextEm = org.hibernate.search.jpa.Search.getFullTextEntityManager(
entityManager);
FullTextQuery query = fullTextEm.createFullTextQuery(
    qb.keyword()
      .onField( "tags" )
      .matching( "round-based" )
      .createQuery(),
    VideoGame.class
)
.setProjection(
    ElasticsearchProjectionConstants.SOURCE,
    ElasticsearchProjectionConstants.TOOK,
    ElasticsearchProjectionConstants.TIMED_OUT
);

Object[] projection = (Object[]) query.getSingleResult();
Integer took = (Integer) projection[1]; // Execution time (milliseconds)
Boolean timedOut = (Boolean) projection[2]; // Whether the query timed out

```

11.5.5. Filters

The Elasticsearch integration supports the definition of full text filters.

Your filters need to implement the `ElasticsearchFilter` interface.

```

public class DriversMatchingNameElasticsearchFilter implements ElasticsearchFilter {

    private String name;

    public DriversMatchingNameElasticsearchFilter() {
    }

    public void setName(String name) {
        this.name = name;
    }

    @Override
    public String getJsonFilter() {
        return "{ 'term': { 'name': '" + name + "' } }";
    }

}

```

You can then declare the filter in your entity.

```

@Entity
@Indexed
@FullTextFilterDef(name = "namedDriver",
    impl = DriversMatchingNameElasticsearchFilter.class)
public class Driver {

    @Id
    @DocumentId
    private int id;

    @Field(analyze = Analyze.YES)
    private String name;

    // ...

}

```

From then you can use it as usual.

```
FullTextQuery ftQuery = /* ... */;  
ftQuery.enableFullTextFilter( "namedDriver" ).setParameter( "name", "liz" );
```

For static filters, you can simply extend the `SimpleElasticsearchFilter` and provide an Elasticsearch filter in JSON form.

11.6. Index optimization

The optimization features documented in [Index Optimization](#) are only partially implemented. That kind of optimization is rarely needed with recent versions of Lucene (on which Elasticsearch is based), but some of it is still provided for the very specific case of indexes meant to stay read-only for a long period of time:

- The automatic optimization is not implemented and most probably never will be.
- The manual optimization (`searchFactory.optimize()`) is implemented.

11.7. Direct access to the Elasticsearch client

When writing complex applications with advanced requirements, it may be necessary from time to time to send requests to the Elasticsearch cluster directly, in particular if Hibernate Search does not support this kind of requests out of the box.

To that end, you can access the Elasticsearch client used by Hibernate Search internally with the following snippet of code:

```
SearchFactory searchFactory = fullTextSession.getSearchFactory();  
IndexFamily indexFamily = searchFactory.getIndexFamily( ElasticsearchIndexFamilyType.get() );  
ElasticsearchIndexFamily elasticsearchIndexFamily = indexFamily.unwrap(  
    ElasticsearchIndexFamily.class );  
RestClient restClient = elasticsearchIndexFamily.getClient( RestClient.class );
```



The client itself is not part of the Hibernate Search API, but of the official Elasticsearch REST client API.

If we have to switch the internal implementation to another client with a different Java type, which we may do without prior notice, the snippet of code above will throw an exception.

11.7.1. Logging executed requests

Search queries are logged to the `org.hibernate.search.fulltext_query` category at `DEBUG`

level, as when using an embedded Lucene instance (the query format is Elasticsearch's, though).

In addition, you can enable the logging of every single request sent to the Elasticsearch cluster by enabling **DEBUG** logging for the log category `org.hibernate.search.elasticsearch.request`. The logs will include, in particular, the HTTP method, path, query parameters, execution time and response status. At **DEBUG** level, the request body and response body will be omitted; use the **TRACE** level to also print out the request body and response body.



By default, JSON in logs will be printed inline, without indent. If you wish the JSON to be pretty-printed, use the `hibernate.search.elasticsearch.log.json_pretty_printing` configuration property.

11.8. Limitations

Not everything is implemented yet.

At the time of writing this, the following features are known to not work yet; feel free to check the JIRA links to get updates:

- Query timeouts: [HSEARCH-2399](#)
- MoreLikeThis queries: [HSEARCH-2395](#)
- `@IndexedEmbedded.indexNullAs`: [HSEARCH-2389](#)
- `Statistics`: [HSEARCH-2421](#)
- `@AnalyzerDiscriminator`: [HSEARCH-2428](#)
- Dynamic sharding: [HSEARCH-2725](#)
- Mixing Lucene based indexes and Elasticsearch based indexes (partial support is implemented though)
- Hibernate Search does not make use of nested objects nor parent child relationship mapping [HSEARCH-2263](#). This is largely mitigated by the fact that Hibernate Search does the denormalization itself and maintain data consistency when nested objects are updated.
- There is room for improvements in the performances of the MassIndexer implementation
- Our new Elasticsearch integration module does not work in OSGi environments. If you need this, please vote for: [HSEARCH-2524](#).

11.9. Specific versions of Elasticsearch to avoid

We might occasionally hit issues in specific versions of Elasticsearch. This is not the right place to track Elasticsearch defects but we'll do our best to track specific issues which affect compatibility with Hibernate Search.

- Mapping `java.time.ZonedDateTime` won't work with Elasticsearch 2.4.1 because of [a JodaTime bug affecting Elasticsearch: HSEARCH-2414](#).

Solution: Upgrade to Elasticsearch 2.4.2 or later.

11.10. Acknowledgment

More information about Elasticsearch can be found on the [Elasticsearch website](#) and its [reference documentation](#).

Chapter 12. Further reading

Last but not least, a few pointers to further information. We highly recommend you to get a copy of [Hibernate Search in Action](#). This excellent book covers Hibernate Search in much more depth than this online documentation can and has a great range of additional examples. If you want to increase your knowledge of Lucene we recommend [Lucene in Action \(Second Edition\)](#).

Because Hibernate Search's functionality is tightly coupled to Hibernate ORM it is a good idea to understand Hibernate. Start with the [online documentation](#) or get hold of a copy of [Java Persistence with Hibernate, Second Edition](#).

If you have any further questions regarding Hibernate Search or want to share some of your use cases have a look at the [Hibernate Search Wiki](#) and the [Hibernate Search Forum](#). We are looking forward hearing from you.

In case you would like to report a bug use the [Hibernate Search JIRA](#) instance. Feedback is always welcome!

Chapter 13. Credits

The full list of contributors to Hibernate Search can be found in the `copyright.txt` file in the Hibernate Search sources, available in particular in our [git repository](#).

The following contributors have been involved in this documentation:

- Emmanuel Bernard
- Hardy Ferentschik
- Gustavo Fernandes
- Sanne Grinovero
- Mincong Huang
- Nabeel Ali Memon
- Gunnar Morling
- Yoann Rodière
- Guillaume Smet