



Getting started with Hibernate Search in Hibernate ORM

2025-07-31

Table of Contents

1. Assumptions	2
2. Dependencies	3
3. Configuration	5
4. Mapping	6
5. Initialization	10
5.1. Schema management	10
5.2. Initial indexing	10
6. Indexing	12
7. Searching	13
8. Analysis	15
9. What's next	20

This guide will walk you through the initial steps required to index and query your [Hibernate ORM](#) entities using [Hibernate Search](#).

If your entities are **not** defined in Hibernate ORM, see [this other guide](#) instead.

Chapter 1. Assumptions

This getting-started guide aims to be generic and does not tie itself to any framework in particular. If you use [Quarkus](#), [WildFly](#) or [Spring Boot](#), make sure to first have a look at [framework support](#) for tips and to learn about quirks.

For the sake of simplicity, this guide assumes you are building an application deployed as a single instance on a single node. For more advanced setups, you are encouraged to have a look at the [Examples of architectures](#).

Chapter 2. Dependencies

The Hibernate Search artifacts can be found in Maven's [Central Repository](#). If you get Hibernate Search from Maven, it is recommended to import Hibernate Search BOM as part of your dependency management to keep all its artifact versions aligned:

```
<dependencyManagement>
  <dependencies>
    <!-- Import Hibernate Search BOM to get all of its artifact versions aligned: -->
    <dependency>
      <groupId>org.hibernate.search</groupId>
      <artifactId>hibernate-search-bom</artifactId>
      <version>8.1.0.CR1</version>
      <type>pom</type>
      <scope>import</scope>
    </dependency>
    <!-- Any other dependency management entries -->
  </dependencies>
</dependencyManagement>
```

If you do not want to, or cannot, fetch the JARs from a Maven repository, you can get them from the [distribution bundle hosted at Sourceforge](#).

In order to use Hibernate Search, you will need at least two direct dependencies:

- a dependency to the **"mapper"**, which extracts data from your domain model and maps it to indexable documents;
- and a dependency to the **"backend"**, which allows indexing and searching these documents.

Below are the most common setups and matching dependencies for a quick start; read [Architecture](#) for more information.

Hibernate ORM + [Lucene](#)

Allows indexing of ORM entities in a single application node, storing the index on the local filesystem.

If you get Hibernate Search from Maven, use these dependencies:

```
<dependencies>
  <!--
    Add dependencies without specifying versions
    as that is already taken care by dependency management:
  -->
  <dependency>
    <groupId>org.hibernate.search</groupId>
    <artifactId>hibernate-search-mapper-orm</artifactId>
  </dependency>
  <dependency>
    <groupId>org.hibernate.search</groupId>
    <artifactId>hibernate-search-backend-lucene</artifactId>
  </dependency>
  <!-- Any other dependency entries -->
</dependencies>
```

If you get Hibernate Search from the distribution bundle, copy the JARs from `dist/engine`, `dist/mapper/orm`, `dist/backend/lucene`, and their respective `lib` subdirectories.

Hibernate ORM + [Elasticsearch](#) (or [OpenSearch](#))

Allows indexing of ORM entities on multiple application nodes, storing the index on a remote Elasticsearch or OpenSearch cluster (to be configured separately).

If you get Hibernate Search from Maven, use these dependencies:

```
<dependencies>
  <!--
    Add dependencies without specifying versions
    as that is already taken care by dependency management:
  -->
  <dependency>
    <groupId>org.hibernate.search</groupId>
    <artifactId>hibernate-search-mapper-orm</artifactId>
  </dependency>
  <dependency>
    <groupId>org.hibernate.search</groupId>
    <artifactId>hibernate-search-backend-elasticsearch</artifactId>
  </dependency>
  <!-- Any other dependency entries -->
</dependencies>
```

If you get Hibernate Search from the distribution bundle, copy the JARs from `dist/engine`, `dist/mapper/orm`, `dist/backend/elasticsearch`, and their respective `lib` subdirectories.

Chapter 3. Configuration

Once you have added all required dependencies to your application, it's time to have a look at the configuration file.



If you are new to Hibernate ORM, we recommend you start [there](#) to implement entity persistence in your application, and only then come back here to add Hibernate Search indexing.

The configuration properties of Hibernate Search are sourced from Hibernate ORM, so they can be added to any file from which Hibernate ORM takes its configuration:

- A `hibernate.properties` file in your classpath.
- The `hibernate.cfg.xml` file in your classpath, if using Hibernate ORM native bootstrapping.
- The `persistence.xml` file in your classpath, if using Hibernate ORM JPA bootstrapping.

Hibernate Search provides sensible defaults for all configuration properties, but depending on your setup you might want to set the following:

Example 3.1: Hibernate Search properties in `persistence.xml` for a "Hibernate ORM + Lucene" setup

```
<property name="hibernate.search.backend.directory.root"
  value="some/filesystem/path"/> ①
```

- ① Set the location of indexes in the filesystem. By default, the backend will store indexes in the current working directory.

Example 3.2: Hibernate Search properties in `persistence.xml` for a "Hibernate ORM + Elasticsearch/OpenSearch" setup

```
<property name="hibernate.search.backend.hosts"
  value="elasticsearch.mycompany.com"/> ①
<property name="hibernate.search.backend.protocol"
  value="https"/> ②
<property name="hibernate.search.backend.username"
  value="ironman"/> ③
<property name="hibernate.search.backend.password"
  value="j@rV1s"/>
```

- ① Set the Elasticsearch hosts to connect to. By default, the backend will attempt to connect to `localhost:9200`.
- ② Set the protocol. The default is `http`, but you may need to use `https`.
- ③ Set the username and password for basic HTTP authentication. You may also be interested in [AWS IAM authentication](#).

Chapter 4. Mapping

Let's assume that your application contains the Hibernate ORM managed classes **Book** and **Author** and you want to index them in order to search the books contained in your database.

Example 4.1: Book and Author entities BEFORE adding Hibernate Search specific annotations

```
import java.util.HashSet;
import java.util.Set;

import jakarta.persistence.Entity;
import jakarta.persistence.GeneratedValue;
import jakarta.persistence.Id;
import jakarta.persistence.ManyToMany;

@Entity
public class Book {

    @Id
    @GeneratedValue
    private Integer id;

    private String title;

    private String isbn;

    private int pageCount;

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

    public Book() {
    }

    // Getters and setters
    // ...

}
```

```
import java.util.HashSet;
import java.util.Set;

import jakarta.persistence.Entity;
import jakarta.persistence.GeneratedValue;
import jakarta.persistence.Id;
import jakarta.persistence.ManyToMany;

@Entity
public class Author {

    @Id
    @GeneratedValue
    private Integer id;

    private String name;

    @ManyToMany(mappedBy = "authors")
    private Set<Book> books = new HashSet<>();

    public Author() {
    }

    // Getters and setters
    // ...

}
```

```
}
```

To make these entities searchable, you will need to map them to an index structure. The mapping can be defined using [annotations](#), or using a [programmatic API](#); this getting started guide will show you a simple annotation mapping. For more details, refer to [Mapping entities to indexes](#).

Below is an example of how the model above can be mapped.

Example 4.2: Book and Author entities AFTER adding Hibernate Search specific annotations

```
import java.util.HashSet;
import java.util.Set;

import jakarta.persistence.Entity;
import jakarta.persistence.GeneratedValue;
import jakarta.persistence.Id;
import jakarta.persistence.ManyToMany;

import org.hibernate.search.mapper.pojo.mapping.definition.annotation.FullTextField;
import org.hibernate.search.mapper.pojo.mapping.definition.annotation.GenericField;
import org.hibernate.search.mapper.pojo.mapping.definition.annotation.Indexed;
import org.hibernate.search.mapper.pojo.mapping.definition.annotation.IndexedEmbedded;
import org.hibernate.search.mapper.pojo.mapping.definition.annotation.KeywordField;

@Entity
@Indexed ①
public class Book {

    @Id ②
    @GeneratedValue
    private Integer id;

    @FullTextField ③
    private String title;

    @KeywordField ④
    private String isbn;

    @GenericField ⑤
    private int pageCount;

    @ManyToMany
    @IndexedEmbedded ⑥
    private Set<Author> authors = new HashSet<>();

    public Book() {
    }

    // Getters and setters
    // ...
}
```

```
import java.util.HashSet;
import java.util.Set;

import jakarta.persistence.Entity;
import jakarta.persistence.GeneratedValue;
import jakarta.persistence.Id;
import jakarta.persistence.ManyToMany;

import org.hibernate.search.mapper.pojo.mapping.definition.annotation.FullTextField;
```

```

@Entity ⑦
public class Author {

    @Id
    @GeneratedValue
    private Integer id;

    @FullTextField ③
    private String name;

    @ManyToMany(mappedBy = "authors")
    private Set<Book> books = new HashSet<>();

    public Author() {
    }

    // Getters and setters
    // ...

}

```

- ① **@Indexed** marks the **Book** entity type as **indexed**, i.e. an index will be created for that entity type, and that index will be kept up to date.
- ② By default, the JPA **@Id** is used to generate a document identifier. In some more complex scenarios it might be needed to map the ID explicitly with **@DocumentId**.
- ③ **@FullTextField** maps a property to a full-text index field with the same name and type. Full-text fields are broken down into tokens and normalized (lowercased, ...). Here we're relying on default analysis configuration, but most applications need to customize it; this will be addressed [further down](#).
- ④ **@KeywordField** maps a property to a non-analyzed index field. Useful for identifiers, for example.
- ⑤ Hibernate Search is not just for full-text search: you can index non-**String** types with the **@GenericField** annotation. A [broad range of property types](#) are supported out-of-the-box, such as primitive types (**int**, **double**, ...) and their boxed counterpart (**Integer**, **Double**, ...), enums, date/time types, **BigInteger**/**BigDecimal**, etc.
- ⑥ **@IndexedEmbedded** "embeds" the indexed form of associated objects (entities or embeddables) into the indexed form of the embedding entity.

Here, the **Author** class defines a single indexed field, **name**. Thus adding **@IndexedEmbedded** to the **authors** property of **Book** will add a single field named **authors.name** to the **Book** index. This field will be populated automatically based on the content of the **authors** property, and the books will be re-indexed whenever Hibernate Search [detects that the name property of their author changes](#). See [Mapping associated elements with @IndexedEmbedded](#) for more information.
- ⑦ **Author** is **@IndexedEmbedded** in other entities, but does not need to be searchable by itself, so it does not need an index and does not need to be annotated with **@Indexed**.

This is a very simple example, but is enough to get started. Just remember that Hibernate Search allows more complex mappings:

- Multiple **@*Field** annotations exist, some of them allowing full-text search, some of them allowing finer-grained configuration for field of a certain type. You can find out more about **@*Field**

annotations in [Mapping a property to an index field with @GenericField, @FullTextField,](#)

- Properties, or even types, can be mapped with finer-grained control using "bridges". This allows the mapping of types that are not supported out-of-the-box. See [Binding and bridges](#) for more information.

Chapter 5. Initialization

Before the application is started for the first time, some initialization may be required:

- The indexes and their schema need to be created.
- Data already present in the database (if any) needs to be indexed.

5.1. Schema management

Before indexing can take place, indexes and their schema need to be created, either on disk (Lucene) or through REST API calls (Elasticsearch).

Fortunately, by default, Hibernate Search will take care of creating indexes on the first startup: you don't have to do anything.

The next time the application is started, existing indexes will be re-used.



Any change to your mapping (adding new fields, changing the type of existing fields, ...) between two restarts of the application will require an update to the index schema.

This will require some special handling, though it can easily be solved by dropping and re-creating the index. See [Changing the mapping of an existing application](#) for more information.

5.2. Initial indexing

As we'll see [later](#), Hibernate Search takes care of triggering indexing every time an entity changes in the application.

However, data already present in the database when you add the Hibernate Search integration is unknown to Hibernate Search, and thus has to be indexed through a batch process. To that end, you can use the mass indexer API, as shown in the following code:

Example 5.1: Using Hibernate Search MassIndexer API to manually (re)index the already persisted data

```
SearchSession searchSession = Search.session( entityManager ); ①

MassIndexer indexer = searchSession.massIndexer( Book.class ) ②
    .threadsToLoadObjects( 7 ); ③

indexer.startAndWait(); ④
```

① Get a Hibernate Search session, called `SearchSession`, from the `EntityManager`.

② Create an "indexer", passing the entity types you want to index. To index all entity types, call `massIndexer()` without any argument.

③ It is possible to set the number of threads to be used. For the complete list of options see [Reindexing large volumes of data with the MassIndexer](#).

④ Invoke the batch indexing process.



If no data is initially present in the database, mass indexing is not necessary.

Chapter 6. Indexing

Hibernate Search will transparently index every entity persisted, updated or removed through Hibernate ORM. Thus, this code would transparently populate your index:

Example 6.1: Using Hibernate ORM to persist data, and implicitly indexing it through Hibernate Search

```
// Not shown: get the entity manager and open a transaction
Author author = new Author();
author.setName( "John Doe" );

Book book = new Book();
book.setTitle( "Refactoring: Improving the Design of Existing Code" );
book.setIsbn( "978-0-58-600835-5" );
book.setPageCount( 200 );
book.getAuthors().add( author );
author.getBooks().add( book );

entityManager.persist( author );
entityManager.persist( book );
// Not shown: commit the transaction and close the entity manager
```

By default, in particular when using the Elasticsearch backend, changes will not be visible right after the transaction is committed. A slight delay (by default one second) will be necessary for Elasticsearch to process the changes.



For that reason, if you modify entities in a transaction, and then execute a search query right after that transaction, the search results may not be consistent with the changes you just performed.

See [Synchronization with the indexes](#) for more information about this behavior and how to tune it.

Chapter 7. Searching

Once the data is indexed, you can perform search queries.

The following code will prepare a search query targeting the index for the `Book` entity, filtering the results so that at least one field among `title` and `authors.name` contains the string `refactoring`. The matches are implicitly on words ("tokens") instead of the full string, and are case-insensitive: that's because the targeted fields are **full-text** fields with the `default analyzer` being applied.

Example 7.1: Using Hibernate Search to query the indexes

```
// Not shown: get the entity manager and open a transaction
SearchSession searchSession = Search.session( entityManager ); ①

SearchResult<Book> result = searchSession.search( Book.class ) ②
    .where( f -> f.match() ③
        .fields( "title", "authors.name" )
        .matching( "refactoring" ) )
    .fetch( 20 ); ④

long totalHitCount = result.total().hitCount(); ⑤
List<Book> hits = result.hits(); ⑥

List<Book> hits2 =
    /* ... same DSL calls as above... */
    .fetchHits( 20 ); ⑦
// Not shown: commit the transaction and close the entity manager
```

- ① Get a Hibernate Search session, called `SearchSession`, from the `EntityManager`.
- ② Initiate a search query on the index mapped to the `Book` entity.
- ③ Define that only documents matching the given predicate should be returned. The predicate is created using a factory `f` passed as an argument to the lambda expression.
- ④ Build the query and fetch the results, limiting to the top 20 hits.
- ⑤ Retrieve the total number of matching entities.
- ⑥ Retrieve matching entities.
- ⑦ In case you're not interested in the whole result, but only in the hits, you can also call `fetchHits()` directly.

If for some reason you don't want to use lambdas, you can use an alternative, object-based syntax, but it will be a bit more verbose:

Example 7.2: Using Hibernate Search to query the indexes – object-based syntax

```
// Not shown: get the entity manager and open a transaction
SearchSession searchSession = Search.session( entityManager ); ①

SearchScope<Book> scope = searchSession.scope( Book.class ); ②

SearchResult<Book> result = searchSession.search( scope ) ③
    .where( scope.predicate().match() ④
        .fields( "title", "authors.name" )
        .matching( "refactoring" )
        .toPredicate() )
```

```

        .fetch( 20 ); ⑤

long totalHitCount = result.total().hitCount(); ⑥
List<Book> hits = result.hits(); ⑦

List<Book> hits2 =
    /* ... same DSL calls as above... */
    .fetchHits( 20 ); ⑧
// Not shown: commit the transaction and close the entity manager

```

- ① Get a Hibernate Search session, called `SearchSession`, from the `EntityManager`.
- ② Create a "search scope", representing the indexed types that will be queried.
- ③ Initiate a search query targeting the search scope.
- ④ Define that only documents matching the given predicate should be returned. The predicate is created using the same search scope as the query.
- ⑤ Build the query and fetch the results, limiting to the top 20 hits.
- ⑥ Retrieve the total number of matching entities.
- ⑦ Retrieve matching entities.
- ⑧ In case you're not interested in the whole result, but only in the hits, you can also call `fetchHits()` directly.

It is possible to get just the total hit count, using `fetchTotalHitCount()`.

Example 7.3: Using Hibernate Search to count the matches

```

// Not shown: get the entity manager and open a transaction
SearchSession searchSession = Search.session( entityManager );

long totalHitCount = searchSession.search( Book.class )
    .where( f -> f.match()
        .fields( "title", "authors.name" )
        .matching( "refactoring" ) )
    .fetchTotalHitCount(); ①
// Not shown: commit the transaction and close the entity manager

```

- ① Fetch the total hit count.



While the examples above retrieved hits as managed entities, it is just one of the possible hit types. See [Projection DSL](#) for more information.

Chapter 8. Analysis

Full-text search allows fast matches on words in a case-insensitive way, which is one step further than substring search in a relational database. But it can get much better: what if we want a search with the term "refactored" to match our book whose title contains "refactoring"? That's possible with custom analysis.

Analysis defines how both the indexed text and search terms are supposed to be processed. This involves *analyzers*, which are made up of three types of components, applied one after the other:

- zero or (rarely) more character filters, to clean up the input text: `A <strong>GREAT</strong> résumé` ⇒ `A GREAT résumé`.
- a tokenizer, to split the input text into words, called "tokens": `A GREAT résumé` ⇒ `[A, GREAT, résumé]`.
- zero or more token filters, to normalize the tokens and remove meaningless tokens. `[A, GREAT, résumé]` ⇒ `[great, resume]`.

There are built-in analyzers, in particular the default one, which will:

- tokenize (split) the input according to the Word Break rules of the [Unicode Text Segmentation algorithm](#);
- filter (normalize) tokens by turning uppercase letters to lowercase.

The default analyzer is a good fit for most language, but is not very advanced. To get the most out of analysis, you will need to define a custom analyzer by picking the tokenizer and filters most suited to your specific needs.

The following paragraphs will explain how to configure and use a simple yet reasonably useful analyzer. For more information about analysis and how to configure it, refer to the [Analysis](#) section.

Each custom analyzer needs to be given a name in Hibernate Search. This is done through analysis configurers, which are defined per backend:

1. First, you need to implement an analysis configurer, a Java class that implements a backend-specific interface: `LuceneAnalysisConfigurer` or `ElasticsearchAnalysisConfigurer`.
2. Second, you need to alter the configuration of your backend to actually use your analysis configurer.

As an example, let's 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 any of the following search terms: "Refactor", "refactors", "refactored" and "refactoring". One way to achieve this is to use an analyzer with the following components:

- A "standard" tokenizer, which splits words at whitespaces, punctuation characters and hyphens. It is a good general-purpose tokenizer.
- A "lowercase" filter, which converts every character to lowercase.
- A "snowball" filter, which applies language-specific [stemming](#).

- Finally, an "ascii-folding" filter, which replaces characters with diacritics ("é", "à", ...) with their ASCII equivalent ("e", "a", ...).

The examples below show how to define an analyzer with these components, depending on the backend you picked.

Example 8.1: Analysis configurer implementation and configuration in `persistence.xml` for a "Hibernate ORM + Lucene" setup

```
package
org.hibernate.search.documentation.mapper.orm.gettingstarted.withhsearch.customanalysis;

import org.hibernate.search.backend.lucene.analysis.LuceneAnalysisConfigurationContext;
import org.hibernate.search.backend.lucene.analysis.LuceneAnalysisConfigurer;

public class MyLuceneAnalysisConfigurer implements LuceneAnalysisConfigurer {
    @Override
    public void configure(LuceneAnalysisConfigurationContext context) {
        context.analyzer( "english" ).custom() ①
            .tokenizer( "standard" ) ②
            .tokenFilter( "lowercase" ) ③
            .tokenFilter( "snowballPorter" ) ③
            .param( "language", "English" ) ④
            .tokenFilter( "asciiFolding" );

        context.analyzer( "name" ).custom() ⑤
            .tokenizer( "standard" )
            .tokenFilter( "lowercase" )
            .tokenFilter( "asciiFolding" );
    }
}
```

```
<property name="hibernate.search.backend.analysis.configurer"
    value=
"class:org.hibernate.search.documentation.mapper.orm.gettingstarted.withhsearch.customanaly
sis.MyLuceneAnalysisConfigurer"/> ⑥
```

- ① Define a custom analyzer named "english", to analyze English text such as book titles.
- ② Set the tokenizer to a standard tokenizer. You need to pass Lucene-specific names to refer to tokenizers; see [Custom analyzers and normalizers](#) for information about available tokenizers, their name and their parameters.
- ③ Set the token filters. Token filters are applied in the order they are given. Here too, Lucene-specific names are expected; see [Custom analyzers and normalizers](#) for information about available token filters, their name and their parameters.
- ④ Set the value of a parameter for the last added char filter/tokenizer/token filter.
- ⑤ Define another custom analyzer, called "name", to analyze author names. On contrary to the first one, do not enable stemming (no `snowballPorter` token filter), as it is unlikely to lead to useful results on proper nouns.
- ⑥ Assign the configurer to the backend in the Hibernate Search configuration (here in `persistence.xml`). For more information about the format of bean references, see [Parsing of bean references](#).

Example 8.2: Analysis configurer implementation and configuration in `persistence.xml` for a "Hibernate ORM + Elasticsearch/OpenSearch" setup

```
package
org.hibernate.search.documentation.mapper.orm.gettingstarted.withhsearch.customanalysis;

import
org.hibernate.search.backend.elasticsearch.analysis.ElasticsearchAnalysisConfigurationContext;
import org.hibernate.search.backend.elasticsearch.analysis.ElasticsearchAnalysisConfigurer;

public class MyElasticsearchAnalysisConfigurer implements ElasticsearchAnalysisConfigurer {
    @Override
    public void configure(ElasticsearchAnalysisConfigurationContext context) {
        context.analyzer( "english" ).custom() ①
            .tokenizer( "standard" ) ②
            .tokenFilters( "lowercase", "snowball_english", "asciifolding" ); ③

        context.tokenFilter( "snowball_english" ) ④
            .type( "snowball" )
            .param( "language", "English" ); ⑤

        context.analyzer( "name" ).custom() ⑥
            .tokenizer( "standard" )
            .tokenFilters( "lowercase", "asciifolding" );
    }
}
```

```
<property name="hibernate.search.backend.analysis.configurer"
    value=
    "class:org.hibernate.search.documentation.mapper.orm.gettingstarted.withhsearch.customanalysis.MyElasticsearchAnalysisConfigurer"/> ⑦
```

- ① Define a custom analyzer named "english", to analyze English text such as book titles.
- ② Set the tokenizer to a standard tokenizer. You need to pass Elasticsearch-specific names to refer to tokenizers; see [Custom analyzers and normalizers](#) for information about available tokenizers, their name and their parameters.
- ③ Set the token filters. Token filters are applied in the order they are given. Here too, Elasticsearch-specific names are expected; see [Custom analyzers and normalizers](#) for information about available token filters, their name and their parameters.
- ④ Note that, for Elasticsearch, any parameterized char filter, tokenizer or token filter must be defined separately and assigned a new name.
- ⑤ Set the value of a parameter for the char filter/tokenizer/token filter being defined.
- ⑥ Define another custom analyzer, named "name", to analyze author names. On contrary to the first one, do not enable stemming (no `snowball_english` token filter), as it is unlikely to lead to useful results on proper nouns.
- ⑦ Assign the configurer to the backend in the Hibernate Search configuration (here in `persistence.xml`). For more information about the format of bean references, see [Parsing of bean references](#).

Once analysis is configured, the mapping must be adapted to assign the relevant analyzer to each field:

Example 8.3: Book and Author entities after adding Hibernate Search specific annotations

```
import java.util.HashSet;
import java.util.Set;

import jakarta.persistence.Entity;
import jakarta.persistence.GeneratedValue;
import jakarta.persistence.Id;
import jakarta.persistence.ManyToMany;

import org.hibernate.search.mapper.pojo.mapping.definition.annotation.FullTextField;
import org.hibernate.search.mapper.pojo.mapping.definition.annotation.GenericField;
import org.hibernate.search.mapper.pojo.mapping.definition.annotation.Indexed;
import org.hibernate.search.mapper.pojo.mapping.definition.annotation.IndexedEmbedded;
import org.hibernate.search.mapper.pojo.mapping.definition.annotation.KeywordField;

@Entity
@Indexed
public class Book {

    @Id
    @GeneratedValue
    private Integer id;

    @FullTextField(analyzer = "english") ①
    private String title;

    @KeywordField
    private String isbn;

    @GenericField
    private int pageCount;

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

    public Book() {
    }

    // Getters and setters
    // ...
}
```

```
import java.util.HashSet;
import java.util.Set;

import jakarta.persistence.Entity;
import jakarta.persistence.GeneratedValue;
import jakarta.persistence.Id;
import jakarta.persistence.ManyToMany;

import org.hibernate.search.mapper.pojo.mapping.definition.annotation.FullTextField;

@Entity
public class Author {

    @Id
    @GeneratedValue
    private Integer id;

    @FullTextField(analyzer = "name") ①
    private String name;

    @ManyToMany(mappedBy = "authors")
```

```

private Set<Book> books = new HashSet<>();

public Author() {
}

// Getters and setters
// ...

}

```

- ① Replace the `@GenericField` annotation with `@FullTextField`, and set the `analyzer` parameter to the name of the custom analyzer configured earlier.



Mapping changes are not auto-magically applied to already-indexed data. Unless you know what you are doing, you should remember to [update your schema](#) and [reindex your data](#) after you changed the Hibernate Search mapping of your entities.

That's it! Now, once the entities will be re-indexed, you will be able to search for the terms "Refactor", "refactors", "refactored" or "refactoring", and the book entitled "Refactoring: Improving the Design of Existing Code" will show up in the results.

Example 8.4: Using Hibernate Search to query the indexes after analysis was configured

```

// Not shown: get the entity manager and open a transaction
SearchSession searchSession = Search.session( entityManager );

SearchResult<Book> result = searchSession.search( Book.class )
    .where( f -> f.match()
        .fields( "title", "authors.name" )
        .matching( "refactored" ) )
    .fetch( 20 );
// Not shown: commit the transaction and close the entity manager

```

Chapter 9. What's next

The above paragraphs gave you 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 review the [examples of architecture](#) to pick the most appropriate for your use case; distributed applications in particular require a specific setup involving a [coordination strategy](#).

You may also want to explore the basic features in more detail. Two topics which were only briefly touched in this tutorial were analysis configuration ([Analysis](#)) and bridges ([Binding and bridges](#)). Both are important features required for more fine-grained indexing.

When it comes to initializing your index, you will be interested in [schema management](#) and [mass indexing](#).

When querying, you will probably want to know more about [predicates](#), [sorts](#), [projections](#), [aggregations](#), [highlights](#).

And if the reference documentation is not enough, you can find pointers to external resources for Hibernate Search as well as related software in the [Further reading](#) section, including examples of applications using Hibernate Search.