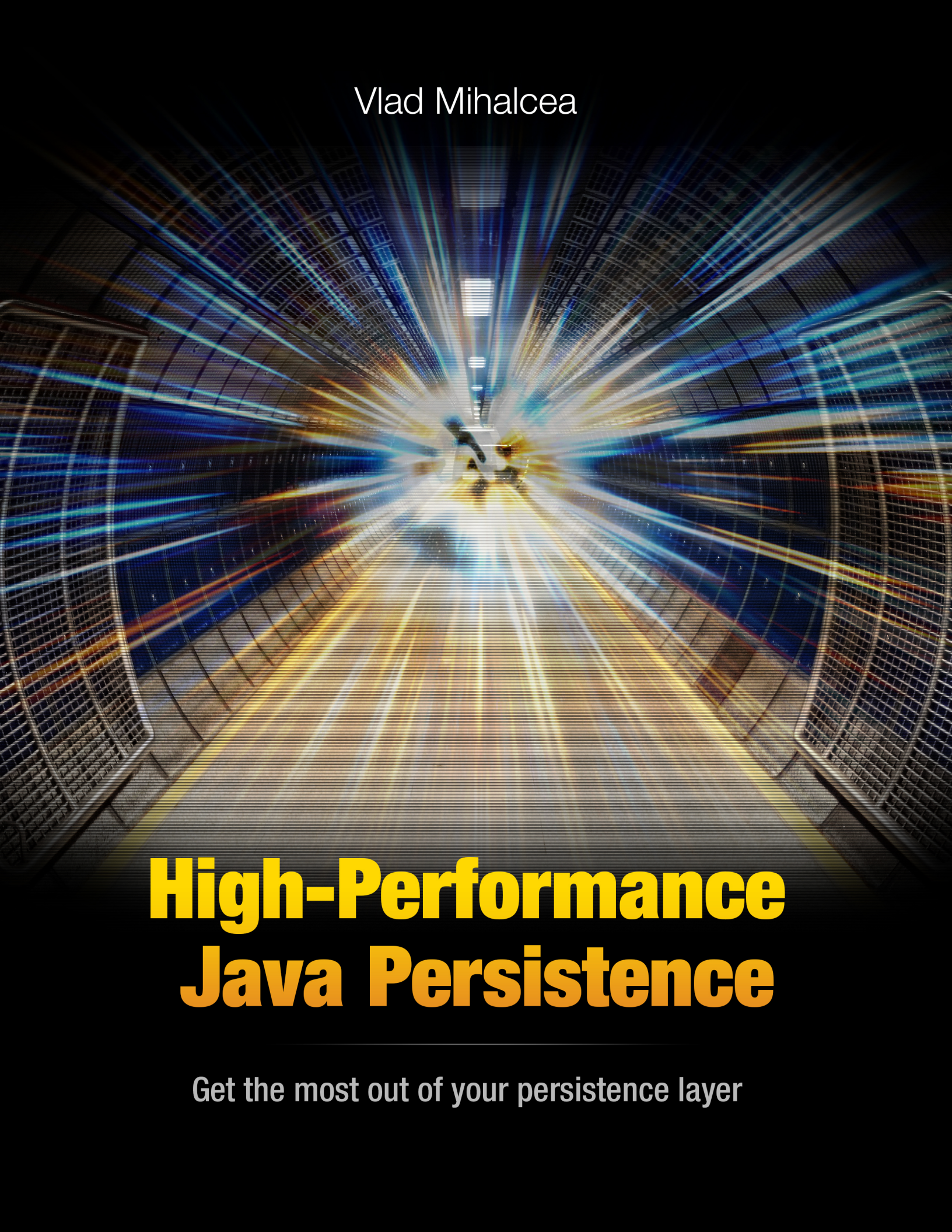




Vlad Mihalcea

High-Performance Java Persistence

Get the most out of your persistence layer

High-Performance Java Persistence

Get the most out of your persistence layer

Vlad Mihalcea

This book is for sale at <http://leanpub.com/high-performance-java-persistence>

This version was published on 2019-06-20



This is a [Leanpub](#) book. Leanpub empowers authors and publishers with the Lean Publishing process. [Lean Publishing](#) is the act of publishing an in-progress ebook using lightweight tools and many iterations to get reader feedback, pivot until you have the right book and build traction once you do.

© 2015 - 2019 Vlad Mihalcea

Tweet This Book!

Please help Vlad Mihalcea by spreading the word about this book on [Twitter](#)!

The suggested tweet for this book is:

Just bought "High-Performance Java Persistence"

<https://leanpub.com/high-performance-java-persistence/> by @vlad_mihalcea #Java #JPA
#Hibernate

To my wife and kids

Contents

1. Batch Updates	1
1.1 Batching Statements	1
1.2 Batching PreparedStatements	4
1.2.1 Choosing the right batch size	6
1.2.2 Bulk processing	7
1.3 Retrieving auto-generated keys	8
1.3.1 Sequences to the rescue	11
2. Why JPA and Hibernate matter	14
2.1 The impedance mismatch	15
2.2 JPA vs. Hibernate	16
2.3 Schema ownership	18
2.4 Entity state transitions	20
2.5 Write-based optimizations	22
2.6 Read-based optimizations	25
2.7 Wrap-up	28
3. Why jOOQ matters	29
3.1 How jOOQ works	29
3.2 DML statements	29
3.3 Java-based schema	31
3.4 Upsert	33
3.4.1 Oracle	34
3.4.2 SQL Server	35
3.4.3 PostgreSQL	35
3.4.4 MySQL	36
3.5 Batch updates	36
3.6 Inlining bind parameters	37
3.7 Complex queries	38
3.8 Stored procedures and functions	41
3.9 Streaming	43
3.10 Keyset pagination	47

1. Batch Updates

JDBC 2.0 introduced *batch updates* so that multiple DML statements can be grouped into a single database request. Sending multiple statements in a single request reduces the number of database roundtrips, therefore decreasing transaction response time. Even if the reference specification uses the term *updates*, any *insert*, *update* or *delete* statement can be batched, and JDBC supports batching for `java.sql.Statement`, `java.sql.PreparedStatement` and `java.sql.CallableStatement` too.

Not only each database driver is distinct, but even different versions of the same driver might require implementation-specific configurations.

1.1 Batching Statements

For executing static SQL statements, JDBC defines the `Statement` interface, which comes with a batching API as well. Other than for test sake, using a `Statement` for CRUD (Create, Read, Update, Delete), as in the example below, should be avoided for it's prone to [SQL injection attacks](#).

```
statement.addBatch(
    "INSERT INTO post (title, version, id) " +
    "VALUES ('Post no. 1', 0, default)");

statement.addBatch(
    "INSERT INTO post_comment (post_id, review, version, id) " +
    "VALUES (1, 'Post comment 1.1', 0, default)");

int[] updateCounts = statement.executeBatch();
```

The numbers of database rows affected by each statement is included in the return value of the `executeBatch()` method.

Oracle

For `Statement` and `CallableStatement`, the [Oracle JDBC Driver](#)^a does not actually support batching. For anything but `PreparedStatement`, the driver ignores batching, and each statement is executed separately.

^ahttp://docs.oracle.com/cd/E11882_01/java.112/e16548/oraperf.htm#JJDBC28752

The following graph depicts how different JDBC drivers behave when varying batch size, the test measuring the time it takes to insert 1000 *post* rows with 4 *comments* each:

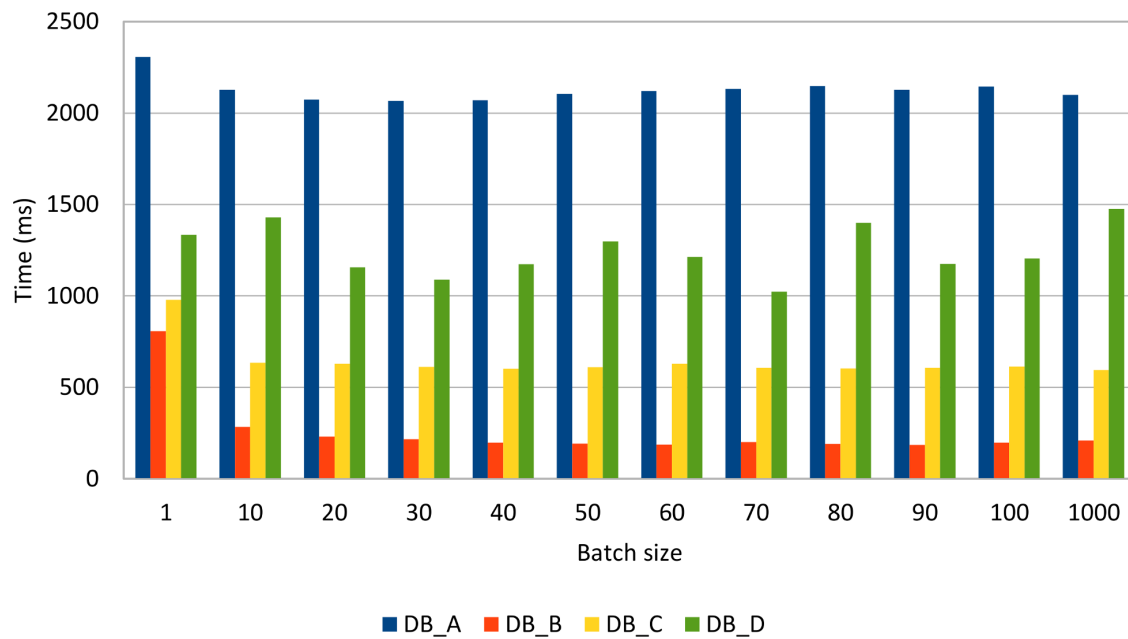


Figure 3.1: Statement batching

Reordering inserts, so that all *posts* are inserted before the *comment* rows, gives the following results:

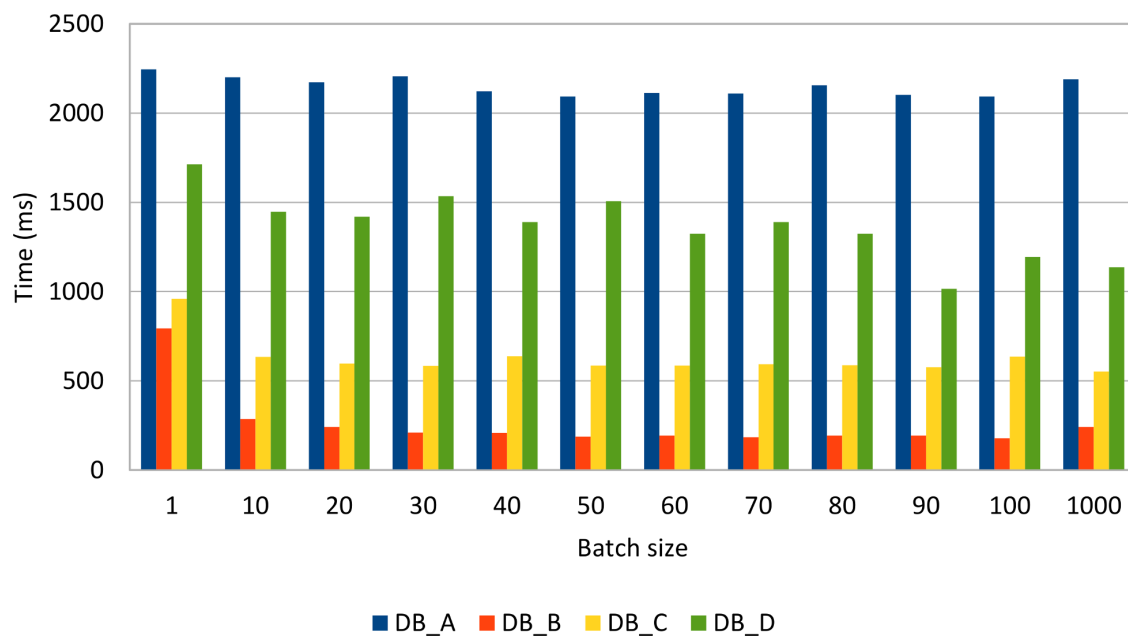


Figure 3.2: Reordered statement batching

Reordering statements does not seem to improve performance noticeably, although some drivers (e.g. MySQL) might take advantage of this optimization.

MySQL

Although it implements the JDBC specification, by default, the MySQL JDBC driver does not send the batched statements in a single request.

For this purpose, the JDBC driver defines the `rewriteBatchedStatements`^a connection property, so that statements get rewritten into a single `String` buffer. In order to fetch the auto-generated row keys, the batch must contain insert statements only.

For `PreparedStatement`, this property rewrites the batched insert statements into a multi-value insert. Unfortunately, the driver is not able to use server-side prepared statements when enabling rewriting.

Without setting this property, the MySQL driver simply executes each DML statement separately, therefore defeating the purpose of batching.

^a<http://dev.mysql.com/doc/connector-j/en/connector-j-reference-configuration-properties.html>

The following graph demonstrates how statement rewriting performs against the default behavior of the MySQL JDBC driver:

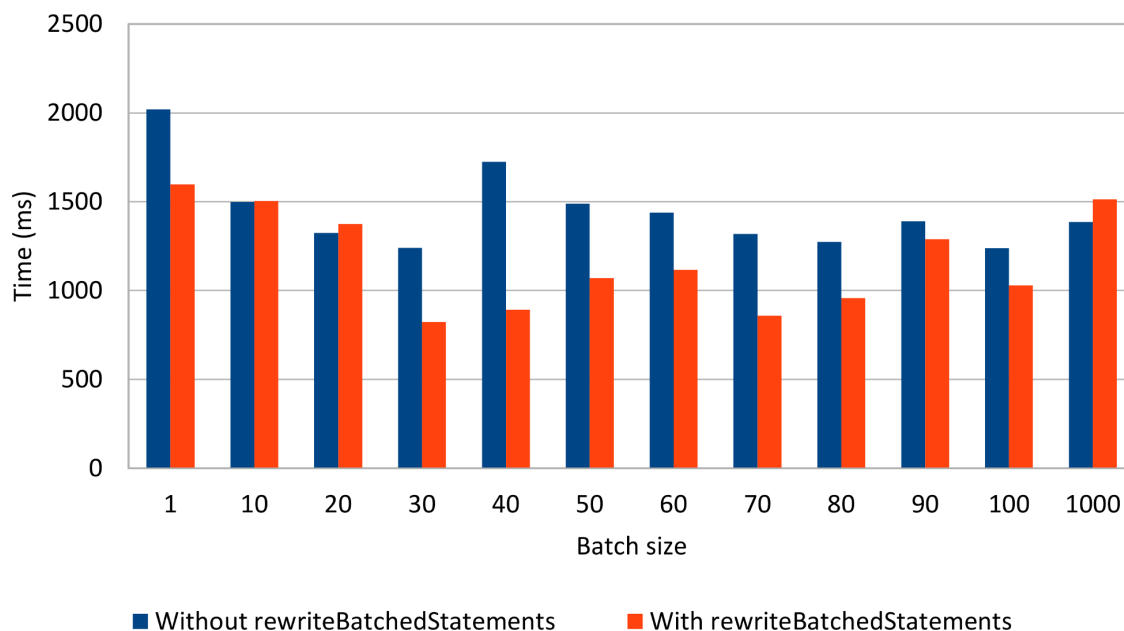


Figure 3.3: MySQL Statement batching

Rewriting non-parameterized statements seems to make a difference, as long as the batch size is not too large. In practice, it is common to use a relatively small batch size, to reduce both the client-side memory footprint and to avoid congesting the server from suddenly processing a huge batch load.

1.2 Batching PreparedStatement

For parameterized statements (a very common enterprise application requirement), the JDBC `Statement` is a poor fit because the only option for varying the executing SQL statement is through `String` manipulation. Using a `String` template or concatenating `String` tokens is risky as it makes the data access logic vulnerable to SQL injection attacks.

To address this shortcoming, JDBC offers the `PreparedStatement` interface for binding parameters in a safe manner. The driver must validate the provided parameter at runtime, therefore discarding unexpected input values.

Because a `PreparedStatement` is associated with a single DML statement, the batch update can group multiple parameter values belonging to the same prepared statement.

```
PreparedStatement postStatement = connection.prepareStatement(
    "INSERT INTO post (title, version, id) " +
    "VALUES (?, ?, ?)");

postStatement.setString(1, String.format("Post no. %1$d", 1));
postStatement.setInt(2, 0);
postStatement.setLong(3, 1);
postStatement.addBatch();

postStatement.setString(1, String.format("Post no. %1$d", 2));
postStatement.setInt(2, 0);
postStatement.setLong(3, 2);
postStatement.addBatch();

int[] updateCounts = postStatement.executeBatch();
```

SQL injection

For an enterprise application, security is a very important technical requirement. The *SQL Injection* attack exploits data access layers that do not use bind parameters. When the SQL statement is the result of `String` concatenation, an attacker could inject a malicious SQL routine that is sent to the database along the current executing statement.

SQL injection is usually done by ending the current statement with the `;` character and continuing it with a rogue SQL command, like modifying the database structure (deleting a table or modifying authorization rights) or even extracting sensitive information.

All DML statements can benefit from batching as the following tests demonstrate. Just like for the JDBC `Statement` test case, the same amount of data (1000 *post* and 4000 *comments*) is inserted, updated, and deleted while varying the batch size.

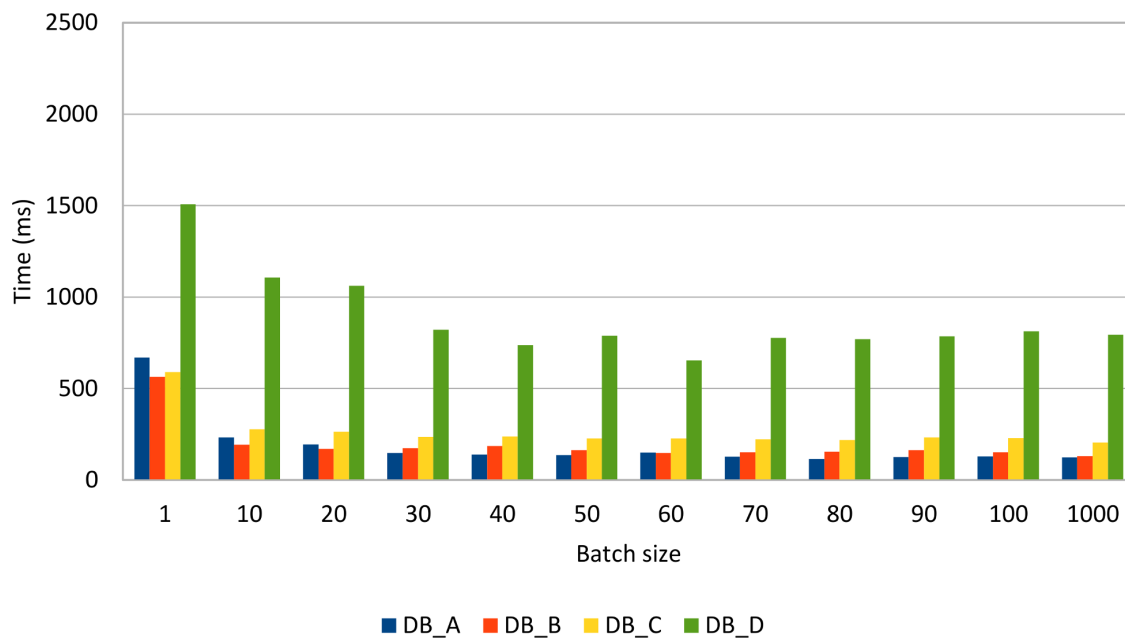


Figure 3.4: Insert PreparedStatement batch size

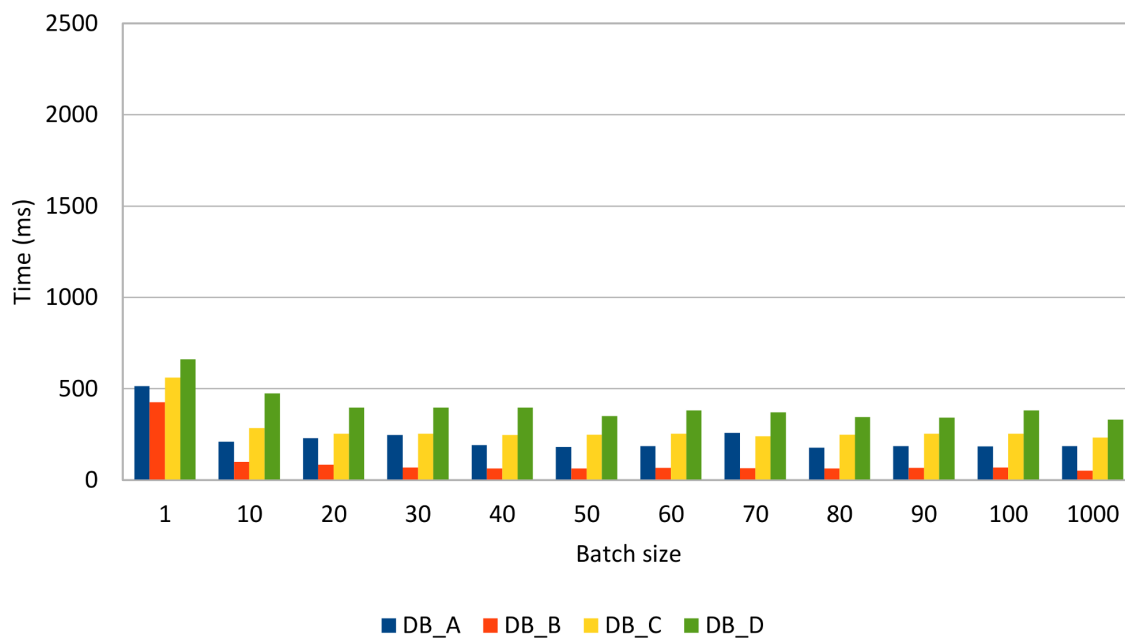


Figure 3.5: Update PreparedStatement batch size

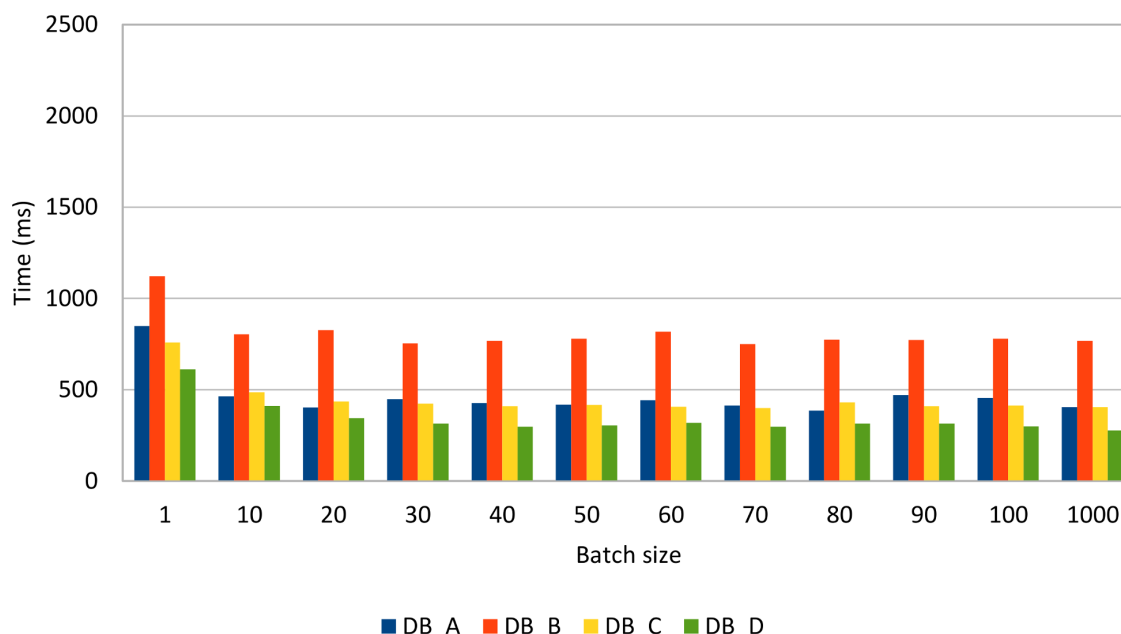


Figure 3.6: Delete PreparedStatement batch size

All database systems show a significant performance improvement when batching prepared statements. Some database systems are very fast when inserting or updating rows while others perform very well when deleting data.

Compared to the previous `Statement` batch insert results, it is clear that, for the same data load, the `PreparedStatement` use case performs just better. In fact, `Statement(s)` should not be used for batching CRUD operations, being more suitable for [bulk processing](#):

```
DELETE from post
WHERE spam = true AND created_on < current_timestamp - INTERVAL '30' day;
```

1.2.1 Choosing the right batch size

Finding the right batch size is not a trivial thing to do as there is no mathematical equation to solve the appropriate batch size for any enterprise application.

Like any other performance optimization technique, measuring the application performance gain in response to a certain batch size value remains the most reliable tuning option.

The astute reader has already figured out that even a low batch size can reduce the transaction response time, and the performance gain does not grow linearly with batch size. Although a larger batch value can save more database roundtrips, the overall performance gain does not necessarily increase linearly. In fact, a very large batch size can hurt application performance if the transaction takes too long to be executed.

As a rule of thumb, you should always measure the performance improvement for various batch sizes. In practice, a relatively low value (between 10 and 30) is usually a good choice.

1.2.2 Bulk processing

Apart from batching, SQL offers bulk operations to modify all rows that satisfy a given filtering criteria. *Bulk update* or *delete* statements can also benefit from indexing, just like select statements.

To update all records from the previous example, one would have to execute the following statements:

```
UPDATE post SET version = version + 1;
UPDATE post_comment SET version = version + 1;
```

Table 3.1: Bulk update time

DB_A time (ms)	DB_B time (ms)	DB_C time (ms)	DB_D time (ms)
26	13	58	9

The bulk alternative is one order of magnitude faster than batch updates. However, batch updates can benefit from application-level optimistic locking mechanisms, which are suitable for preventing *lost updates* when data is loaded in a read-only database transaction and written back in a successive transaction.

Like with updates, bulk deleting is also much faster than deleting in batches.

```
DELETE FROM post_comment WHERE version > 0;
DELETE FROM post WHERE version > 0;
```

Table 3.2: Bulk delete time

DB_A time (ms)	DB_B time (ms)	DB_C time (ms)	DB_D time (ms)
3	12	1	2

Long-running transaction caveats

Processing too much data in a single transaction can degrade application performance, especially in a highly concurrent environment. Whether using 2PL (Two-Phase Locking) or MVCC (Multiversion Concurrency Control), writers always block other conflicting writers.

Long running transactions can affect both batch updates and bulk operations if the current transaction modifies a very large number of records. For this reason, it is more practical to break a large batch processing task into smaller manageable ones that can release locks in a timely fashion.

1.3 Retrieving auto-generated keys

It is common practice to delegate the row identifier generation to the database system. This way, the developer does not have to provide a monotonically incrementing primary key since the database takes care of this upon inserting a new record.

As convenient as this practice may be, it is important to know that auto-generated database identifiers might conflict with the batch insert process.

Like many other database features, setting the auto incremented identifier strategy is database-specific so the choice goes between an *identity* column or a database *sequence* generator.

Oracle

Prior to Oracle 12c, an auto incremented generator had to be implemented on top of a database sequence.

```
CREATE SEQUENCE post_seq;

CREATE TABLE post (
  id NUMBER(19,0) NOT NULL,
  title VARCHAR2(255 CHAR),
  version NUMBER(10,0) NOT NULL,
  PRIMARY KEY (id));

CREATE OR REPLACE TRIGGER post_identity
BEFORE INSERT ON post
FOR EACH ROW
BEGIN
  SELECT post_seq.NEXTVAL
  INTO   :NEW.id
  FROM   dual;
end;
```

Oracle 12c adds support for identity columns as well, so the previous example can be simplified as follows.

```
CREATE TABLE post (
  id NUMBER(19,0) NOT NULL GENERATED ALWAYS AS IDENTITY,
  title VARCHAR2(255 CHAR),
  version NUMBER(10,0) NOT NULL,
  PRIMARY KEY (id));
```

SQL Server

Traditionally, SQL Server offered identity column generators, but, since SQL Server 2012, it now supports database sequences as well.

```
CREATE TABLE post (  
    id BIGINT IDENTITY NOT NULL,  
    title VARCHAR(255),  
    version INT NOT NULL,  
    PRIMARY KEY (id));
```

PostgreSQL

PostgreSQL 9.5 does not support identity columns natively, although it offers the `SERIAL` column type which can emulate an identity column.

```
CREATE TABLE post (  
    id SERIAL NOT NULL,  
    title VARCHAR(255),  
    version INT4 NOT NULL,  
    PRIMARY KEY (id));
```

The `SERIAL` (4 bytes) and `BIGSERIAL` (8 bytes) types are just a syntactic sugar expression as, behind the scenes, PostgreSQL relies on a database sequence anyway.

The previous definition is therefore equivalent to:

```
CREATE SEQUENCE post_id_seq;  
  
CREATE TABLE post (  
    id INTEGER DEFAULT NEXTVAL('post_id_seq') NOT NULL,  
    title VARCHAR(255),  
    version INT4 NOT NULL,  
    PRIMARY KEY (id));  
);
```

MySQL

MySQL 5.7 only supports identity columns through the `AUTO_INCREMENT` attribute.

```
CREATE TABLE post (  
    id BIGINT NOT NULL AUTO_INCREMENT,  
    title VARCHAR(255),  
    version INTEGER NOT NULL,  
    PRIMARY KEY (id));
```

Many database developers like this approach since the client does not have to care about supplying a database identifier upon inserting a new row.

```
INSERT INTO post (title, version) VALUES (?, ?);
```

To retrieve the newly created row identifier, the JDBC `PreparedStatement` must be instructed to return the auto-generated keys.

```
PreparedStatement postStatement = connection.prepareStatement(  
    "INSERT INTO post (title, version) VALUES (?, ?)",  
    Statement.RETURN_GENERATED_KEYS  
);
```

One alternative is to hint the driver about the column index holding the auto-generated key column.

```
PreparedStatement postStatement = connection.prepareStatement(  
    "INSERT INTO post (title, version) VALUES (?, ?)",  
    new int[] {1}  
);
```

The column name can also be used to instruct the driver about the auto-generated key column.

```
PreparedStatement postStatement = connection.prepareStatement(  
    "INSERT INTO post (title, version) VALUES (?, ?)",  
    new String[] {"id"}  
);
```

It is better to know all these three alternatives because they are not interchangeable on all database systems.

Oracle auto-generated key retrieval gotcha

When using `Statement.RETURN_GENERATED_KEYS`, Oracle returns a `ROWID` instead of the actually generated column value. A workaround is to supply the column index or the column name, and so the auto-generated value can be extracted after executing the statement.

According to the JDBC 4.2 specification, every driver must implement the `supportsGetGeneratedKeys()` method and specify whether it supports auto-generated key retrieval. Unfortunately, this only applies to single statement updates as the specification does not make it mandatory for drivers to support generated key retrieval for batch statements. That being said, not all database systems support fetching auto-generated keys from a batch of statements.

Table 3.3: Driver support for retrieving generated keys

Returns generated keys after calling	Oracle JDBC driver (11.2.0.4)	Oracle JDBC driver (12.1.0.1)	SQL Server JDBC driver (4.2)	PostgreSQL JDBC driver (9.4-1201-jdbc41)	MySQL JDBC driver (5.1.36)
<code>executeUpdate()</code>	Yes	Yes	Yes	Yes	Yes
<code>executeBatch()</code>	No	Yes	No	Yes	Yes

If the Oracle JDBC driver 11.2.0.4 cannot retrieve auto-generated batch keys, the 12.1.0.1 version works just fine. When trying to get the auto-generated batch keys, the SQL Server JDBC driver throws this exception: *The statement must be executed before any results can be obtained.*

1.3.1 Sequences to the rescue

As opposed to identity columns, database sequences offer the advantage of decoupling the identifier generation from the actual row insert. To make use of batch inserts, the identifier must be fetched prior to setting the insert statement parameter values.

```
private long getNextSequenceValue(Connection connection)
    throws SQLException {
    try(Statement statement = connection.createStatement()) {
        try(ResultSet resultSet = statement.executeQuery(
            callSequenceSyntax())) {
            resultSet.next();
            return resultSet.getLong(1);
        }
    }
}
```

For calling a sequence, every database offers a specific syntax:

Oracle

```
SELECT post_seq.NEXTVAL FROM dual;
```

SQL Server

```
SELECT NEXT VALUE FOR post_seq;
```

PostgreSQL

```
SELECT NEXTVAL('post_seq');
```

Because the primary key is generated up-front, there is no need to call the `getGeneratedKeys()` method, and so batch inserts are not driver dependent anymore.

```
try(PreparedStatement postStatement = connection.prepareStatement(
    "INSERT INTO post (id, title, version) VALUES (?, ?, ?)")) {
    for (int i = 0; i < postCount; i++) {
        if (i > 0 && i % batchSize == 0) {
            postStatement.executeBatch();
        }
        postStatement.setLong(1, getNextSequenceValue(connection));
        postStatement.setString(2, String.format("Post no. %1$d", i));
        postStatement.setInt(3, 0);
        postStatement.addBatch();
    }
    postStatement.executeBatch();
}
```

Many database engines use sequence number generation optimizations to lower the sequence call execution as much as possible. If the number of inserted records is relatively low, then the sequence call overhead (extra database roundtrips) will be insignificant. However, for batch processors inserting large amounts of data, the extra sequence calls can add up.

Optimizing sequence calls

The data access layer does not need to go to the database to fetch a unique identifier if the sequence incrementation step is greater than 1. For a step of N , the sequence numbers are 1, $N + 1$, $2N + 1$, $3N + 1$, etc. The data access logic can assign identifiers in-between the database sequence calls (e.g. 2, 3, 4, ..., $N - 1$, N), and so it can mitigate the extra network roundtrips penalty.

This strategy is going to be discussed in greater detail in the [Hibernate types and identifiers chapter](#).

2. Why JPA and Hibernate matter

Although JDBC does a very good job of exposing a common API that hides the database vendor-specific communication protocol, it suffers from the following shortcomings:

- The API is undoubtedly verbose, even for trivial tasks.
- Batching is not transparent from the data access layer perspective, requiring a specific API than its non-batched statement counterpart.
- Lack of built-in support for explicit locking and optimistic concurrency control.
- For local transactions, the data access is tangled with transaction management semantics.
- Fetching joined relations requires additional processing to transform the `ResultSet` into Domain Models or DTO (Data Transfer Object) graphs.

Although the primary goal of an ORM (Object-Relational Mapping) tool is to automatically translate object state transitions into SQL statements, this chapter aims to demonstrate that Hibernate can address all the aforementioned JDBC shortcomings.

Java Persistence history

The EJB 1.1 release offered a higher-level persistence abstraction through scalable enterprise components, known as Entity Beans. Although the design looked good on paper, in reality, the heavyweight RMI-based implementation proved to be disastrous from a performance perspective. Neither the EJB 2.0 support for *local interfaces* could revive the Entity Beans popularity, and, due to high-complexity and vendor-specific implementation details, most projects chose JDBC instead.

Hibernate was born out of all the frustration of using the Entity Bean developing model. As an open-source project, Hibernate managed to gain a lot of popularity, and so it soon became the *de facto* Java persistence framework.

In response to all the criticism associated with Entity Bean persistence, the Java Community Process advanced a lightweight POJO-based approach, and the JDO specification was born. Although JDO is data source agnostic, being capable of operating with both relation databases as well as NoSQL or even flat files, it never managed to hit mainstream popularity. For this reason, the Java Community Process decided that EJB3 would be based on a new specification, inspired by Hibernate and TopLink, and JPA (Java Persistence API) became the standard Java persistence technology.

The morale of this is that persistence is a very complex topic, and it demands a great deal of knowledge of both the database and the data access usage patterns.

2.1 The impedance mismatch

When a relational database is manipulated through an object-oriented program, the two different data representations start conflicting.

In a relational database, data is stored in tables, and the relational algebra defines how data associations are formed. On the other hand, an object-oriented programming (OOP) language allows objects to have both state and behavior, and bidirectional associations are permitted.

The burden of converging these two distinct approaches has generated much tension, and it has been haunting enterprise systems for a very long time.

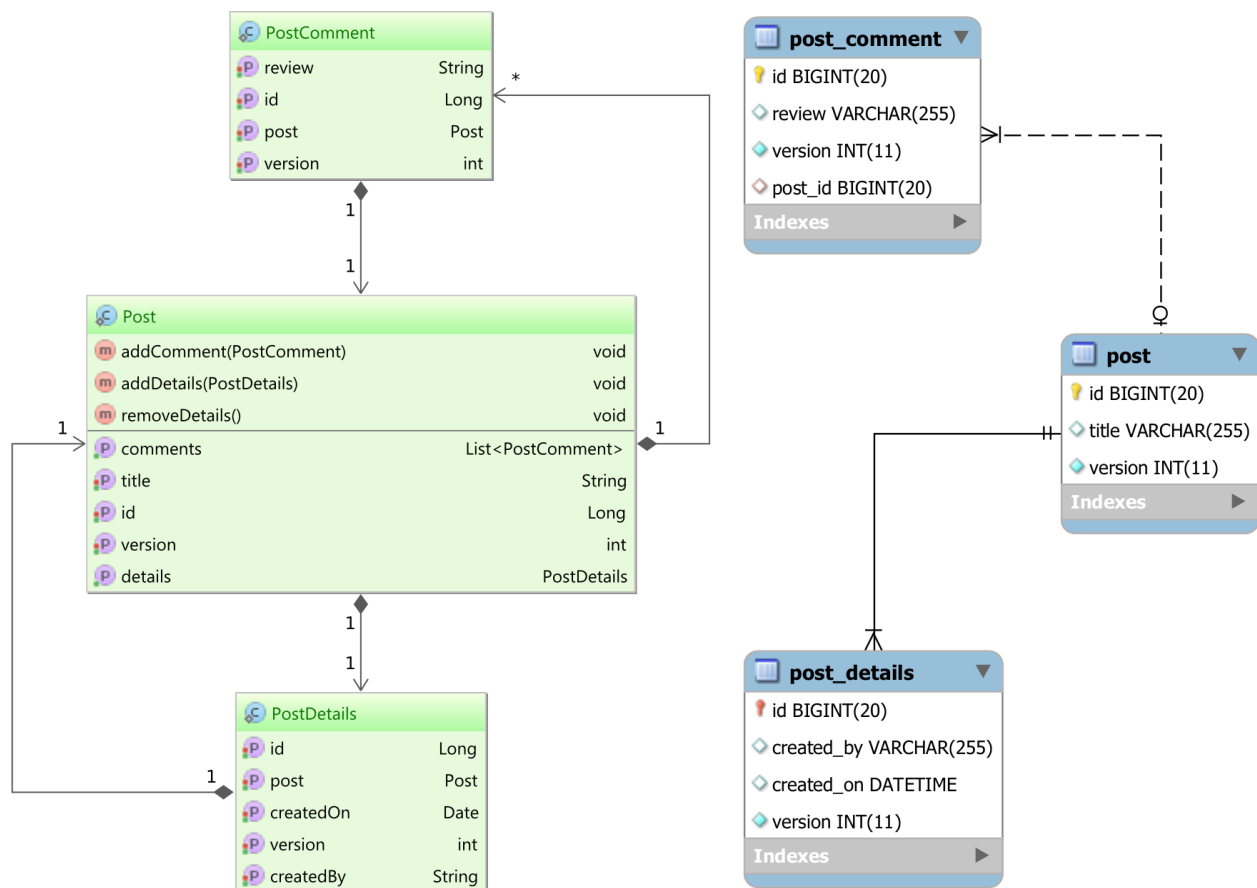


Figure 7.1: Object/Relational Mapping

The above diagram portrays the two different schemas that the data access layer needs to correlate. While the database schema is driven by the SQL standard specification, the Domain Model comes with an object-oriented schema representation as well.

The Domain Model encapsulates the business logic specifications and captures both data structures and the behavior that governs business requirements. OOP facilitates Domain Modeling, and many modern enterprise systems are implemented on top of an object-oriented programming language.

Because the underlying data resides in a relational database, the Domain Model must be adapted to the database schema and the SQL-driven communication protocol. The ORM design pattern helps to bridge these two different data representations and close the technological gap between them. Every database row is associated with a Domain Model object (Entity in JPA terminology), and so the ORM tool can translate the entity state transitions into DML statements.

From an application development point of view, this is very convenient since it is much easier to manipulate Domain Model relationships rather than visualizing the business logic through its underlying SQL statements.

2.2 JPA vs. Hibernate

JPA is only a specification. It describes the interfaces that the client operates with and the standard object-relational mapping metadata (Java annotations or XML descriptors). Beyond the API definition, JPA also explains (although not exhaustively) how these specifications are ought to be implemented by the JPA providers. JPA evolves with the Java EE platform itself (Java EE 6 featuring JPA 2.0 and Java EE 7 introducing JPA 2.1).

Hibernate was already a full-featured Java ORM implementation by the time the JPA specification was released for the first time. Although it implements the JPA specification, Hibernate retains its native API for both backward compatibility and to accommodate non-standard features.

Even if it is best to adhere to the JPA standard, in reality, many JPA providers offer additional features targeting a high-performance data access layer requirements. For this purpose, Hibernate comes with the following non-JPA compliant features:

- extended identifier generators (hi/lo, pooled, pooled-lo)
- transparent prepared statement batching
- customizable CRUD (@SQLInsert, @SQLUpdate, @SQLDelete) statements
- static/dynamic entity/collection filters (e.g. @FilterDef, @Filter, @Where)
- mapping attributes to SQL fragments (e.g. @Formula)
- immutable entities (e.g. @Immutable)
- more flush modes (e.g. FlushMode.MANUAL, FlushMode.ALWAYS)
- querying the second-level cache by the natural key of a given entity
- entity-level cache concurrency strategies
(e.g. Cache(usage = CacheConcurrencyStrategy.READ_WRITE))
- versioned bulk updates through HQL
- exclude fields from optimistic locking check (e.g. @OptimisticLock(excluded = true))
- versionless optimistic locking (e.g. OptimisticLockType.ALL, OptimisticLockType.DIRTY)
- support for skipping (without waiting) pessimistic lock requests
- support for Java 8 Date and Time and stream()
- support for multitenancy



If JPA is the interface, Hibernate is one implementation and implementation details always matter from a performance perspective.

The JPA implementation details leak and ignoring them might hinder application performance or even lead to data inconsistency issues. As an example, the following JPA attributes have a peculiar behavior, which can surprise someone who is familiar with the JPA specification only:

- The `FlushModeType.AUTO`¹ does not trigger a flush for native SQL queries like it does for JPQL or Criteria API.
- The `FetchType.EAGER`² might choose a SQL join or a secondary select whether the entity is fetched directly from the `EntityManager` or through a JPQL (Java Persistence Query Language) or a Criteria API query.

That is why this book is focused on how Hibernate manages to implement both the JPA specification and its non-standard native features (that are relevant from an efficiency perspective).

Portability concerns

Like other non-functional requirements, portability is a feature, and there is still a widespread fear of embracing database-specific or framework-specific features. In reality, it is more common to encounter enterprise applications facing data access performance issues than having to migrate from one technology to the other (be it a relational database or a JPA provider).

The lowest common denominator of many RDBMS is a superset of the SQL-92 standard (although not entirely supported either). SQL-99 supports Common Table Expressions, but MySQL 5.7 does not. Although SQL-2003 defines the `MERGE` operator, PostgreSQL 9.5 favored the `UPSERT` operation instead. By adhering to a SQL-92 syntax, one could achieve a higher degree of database portability, but the price of giving up database-specific features can take a toll on application performance. Portability can be addressed either by subtracting non-common features or through specialization. By offering different implementations, for each supported database system (like the jOOQ framework does), portability can still be achieved.

The same argument is valid for JPA providers too. By layering the application, it is already much easier to swap JPA providers, if there is even a compelling reason for switching one mature JPA implementation with another.

¹<https://docs.oracle.com/javaee/7/api/javax/persistence/FlushModeType.html>

²<https://docs.oracle.com/javaee/7/api/javax/persistence/FetchType.html#EAGER>

2.3 Schema ownership

Because of data representation duality, there has been a rivalry between taking ownership of the underlying schema. Although theoretically, both the database and the Domain Model could drive the schema evolution, for practical reasons, the schema belongs to the database.

An enterprise system might be too large to fit into a single application, so it is not uncommon to split in into multiple subsystems, each one serving a specific goal. As an example, there can be front-end web applications, integration web services, email schedulers, full-text search engines and back-end batch processors that need to load data into the system. All these subsystems need to use the underlying database, whether it is for displaying content to the users or dumping data into the system.

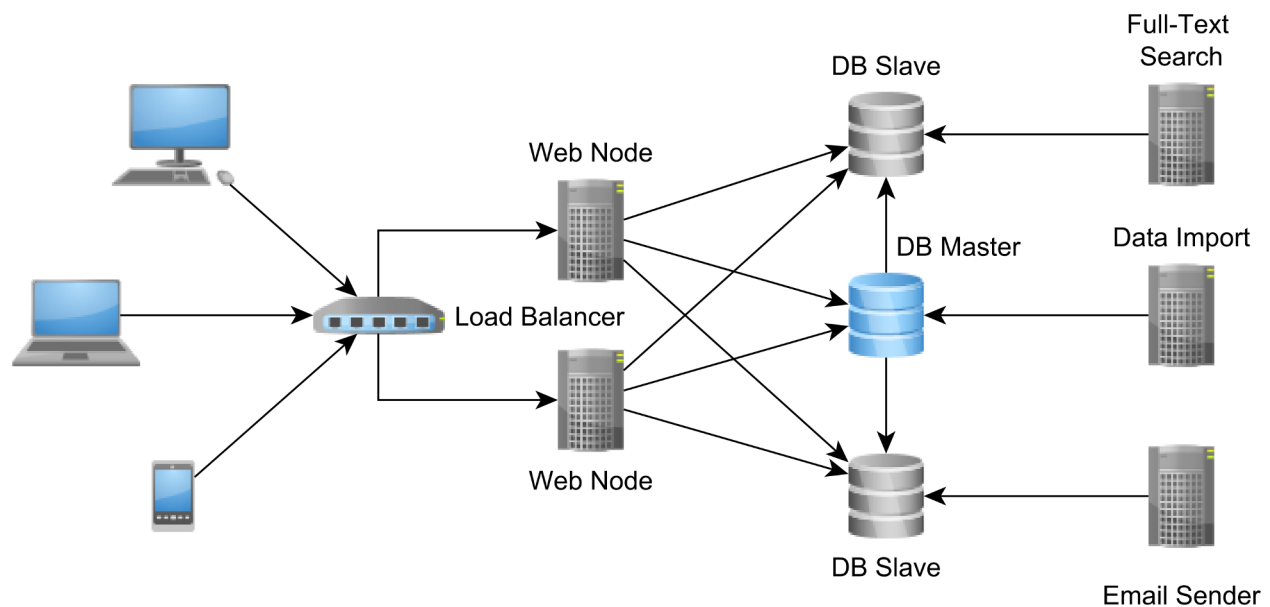


Figure 7.2: Database-centric integration

Although it might not fit any enterprise system, having the database as a central integration point can still be a choice for many reasonable size enterprise systems.

The relational database concurrency models offer strong consistency guarantees, therefore having a significant advantage to application development. If the integration point does not provide transactional semantics, it will be much more difficult to implement a distributed concurrency control mechanism.

Most database systems already offer support for various replication topologies, which can provide more capacity for accommodating an increase in the incoming request traffic. Even if the demand for more data continues to grow, the hardware is always getting better and better (and cheaper too), and database vendors keep on improving their engines to cope with more data.

For these reasons, having the database as an integration point is still a relevant enterprise system design consideration.

The distributed commit log

For very large enterprise systems, where data is split among different providers (relational database systems, caches, Hadoop, Spark), it is no longer possible to rely on the relational database to integrate all disparate subsystems.

In this case, [Apache Kafka](http://kafka.apache.org/)^a offers a fault-tolerant and scalable append-only log structure, which every participating subsystem can read and write concurrently.

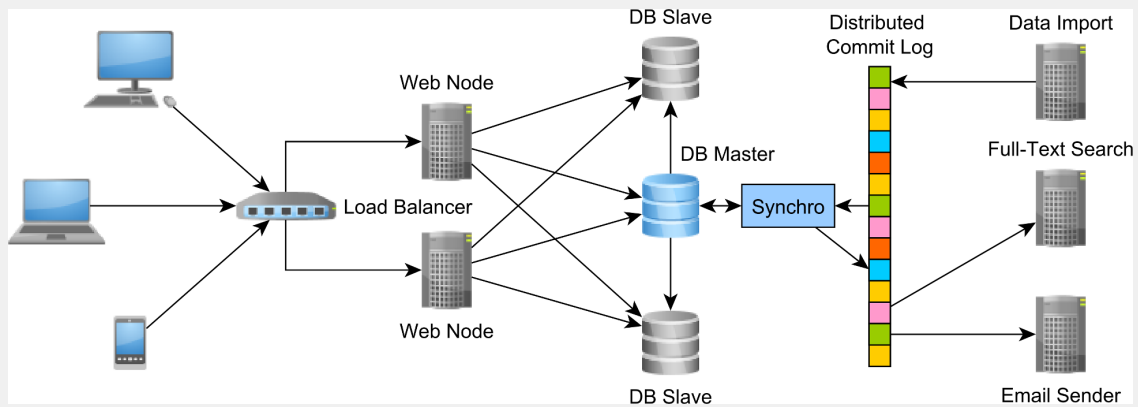


Figure 7.3: Distributed commit log integration

The commit log becomes the integration point, each distributed node individually traversing it and maintaining client-specific pointers in the sequential log structure. This design resembles a database replication mechanism, and so it offers durability (the log is persisted on disk), write performance (append-only logs do not require random access) and read performance (concurrent reads do not require blocking) as well.

^a<http://kafka.apache.org/>

No matter what architecture style is chosen, there is still need to correlate the transient Domain Model with the underlying persistent data.

The data schema evolves along the enterprise system itself, and so the two different schema representations must remain congruent at all times.

Even if the data access framework can auto-generate the database schema, the schema must be migrated incrementally, and all changes need to be traceable in the VCS (Version Control System) as well. Along with table structure, indexes and triggers, the database schema is, therefore, accompanying the Domain Model source code itself. A tool like [Flywaydb](http://flywaydb.org/)³ can automate the database schema migration, and the system can be deployed continuously, whether it is a test or a production environment.

³<http://flywaydb.org/>



The schema ownership goes to the database, and the data access layer must assist the Domain Model to communicate with the underlying data.

2.4 Entity state transitions

JPA shifts the developer mindset from SQL statements to entity state transitions. An entity can be in one of the following states:

Table 7.1: JPA entity states

State	Description
New (Transient)	A newly created entity which is not mapped to any database row is considered to be in the New or Transient state. Once it becomes managed, the Persistence Context issues an insert statement at flush time.
Managed (Persistent)	A Persistent entity is associated with a database row, and it is being managed by the currently running Persistence Context. State changes are detected by the <i>dirty checking</i> mechanism and propagated to the database as update statements at flush time.
Detached	Once the currently running Persistence Context is closed, all the previously managed entities become detached. Successive changes are no longer tracked, and no automatic database synchronization is going to happen.
Removed	A removed entity is only scheduled for deletion, and the actual database delete statement is executed during Persistence Context flushing.

The Persistence Context captures entity state changes, and, during flushing, it translates them into SQL statements. The JPA [EntityManager](http://docs.oracle.com/javaee/7/api/javax/persistence/EntityManager.html#persist-java.lang.Object-)⁴ and the Hibernate [Session](https://docs.jboss.org/hibernate/orm/current/javadocs/org/hibernate/Session.html)⁵ (which includes additional methods for moving an entity from one state to the other) interfaces are gateways towards the underlying Persistence Context, and they define all the entity state transition operations.

⁴<http://docs.oracle.com/javaee/7/api/javax/persistence/EntityManager.html#persist-java.lang.Object->

⁵<https://docs.jboss.org/hibernate/orm/current/javadocs/org/hibernate/Session.html>

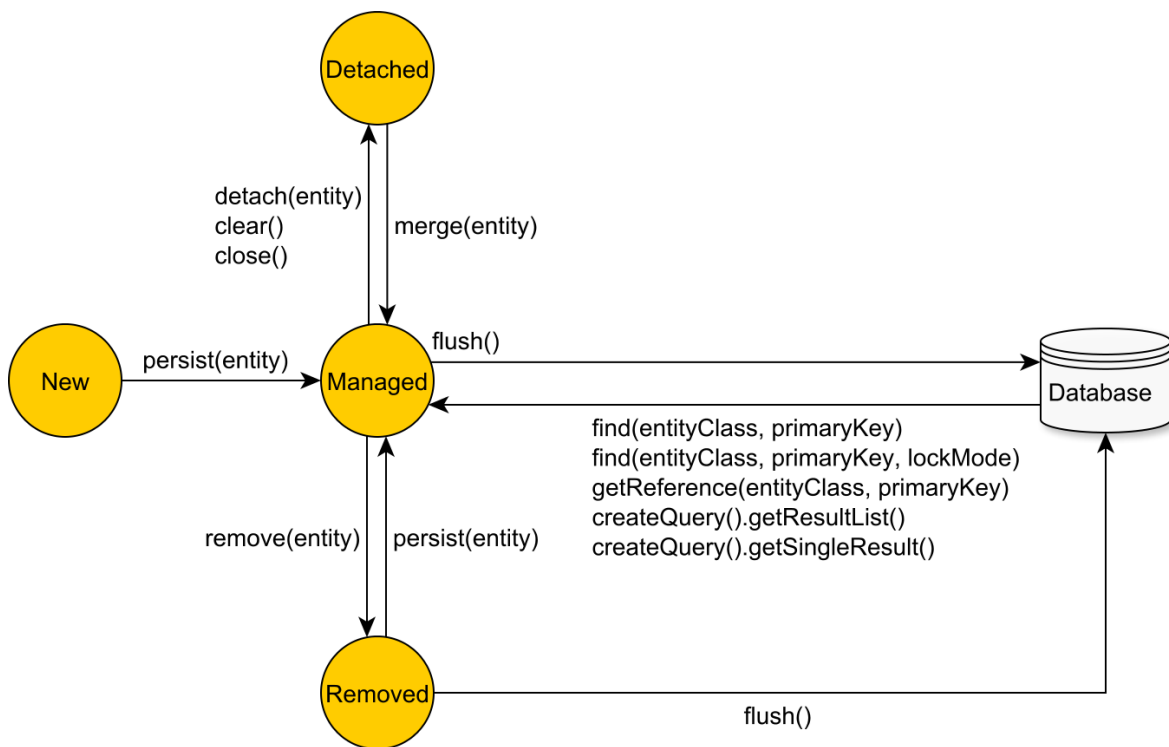


Figure 7.4: JPA entity state transitions

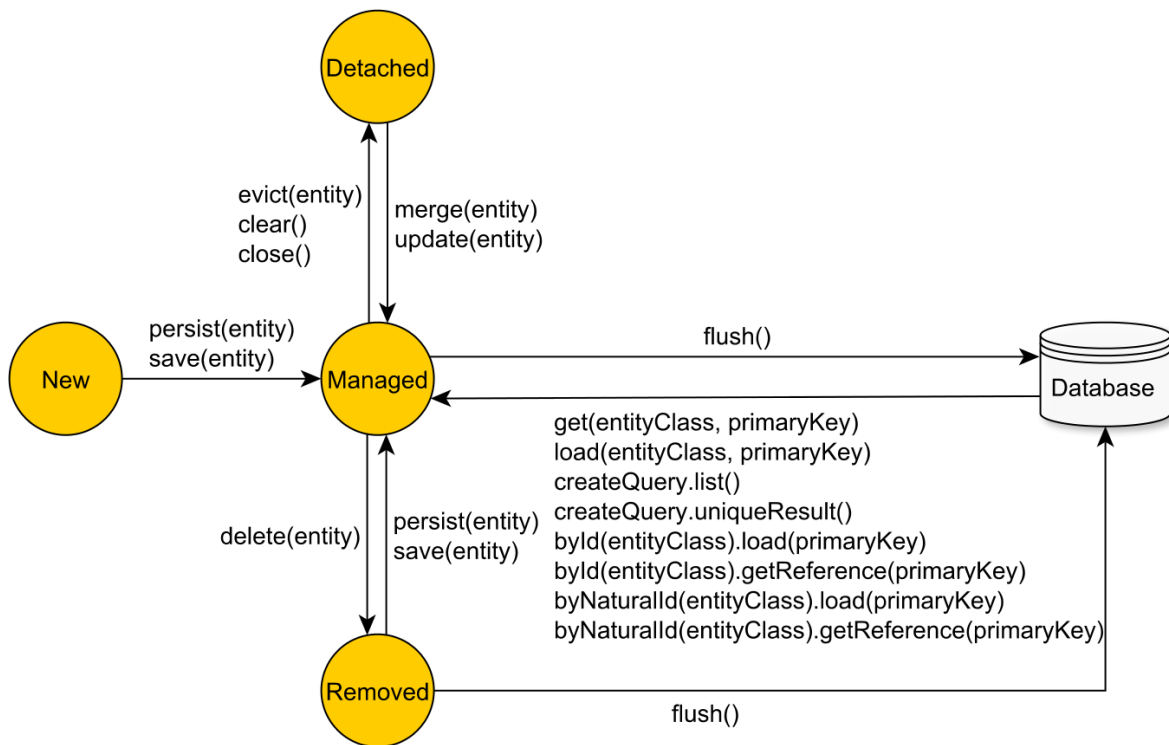


Figure 7.5: Hibernate entity state transitions

2.5 Write-based optimizations

SQL injection prevention

By managing the SQL statement generation, the JPA tool can assist in minimizing the risk of SQL injection attacks. The less the chance of manipulating SQL `String` statements, the safer the application can get. The risk is not completely eliminated because the application developer can still recur to concatenating SQL or JPQL fragments, so rigor is advised.



Hibernate uses `PreparedStatement(s)` exclusively, so not only it protects against SQL injection, but the data access layer can better take advantage of server-side and client-side statement caching as well.

Auto-generated DML statements

The enterprise system database schema evolves with time, and the data access layer must mirror all these modifications as well.

Because the JPA provider auto-generates insert and update statements, the data access layer can easily accommodate database table structure modifications. By updating the entity model schema, Hibernate can automatically adjust the modifying statements accordingly.

This applies to changing database column types as well. If the database schema needs to migrate a postal code from an `INT` database type to a `VARCHAR(6)`, the data access layer will need only to change the associated Domain Model attribute type from an `Integer` to a `String`, and all statements are going to be automatically updated. Hibernate defines a highly customizable JDBC-to-database type mapping system, and the application developer can override a default type association, or even add support for new database types (that are not currently supported by Hibernate).

The entity fetching process is automatically managed by the JPA implementation, which auto-generates the select statements of the associated database tables. This way, JPA can free the application developer from maintaining entity selection queries as well.

Hibernate allows customizing all the CRUD statements, in which case the application developer is responsible for maintaining the associated DML statements.

Although it takes care of the entity selection process, most enterprise systems need to take advantage of the underlying database querying capabilities. For this reason, whenever the database schema changes, all the native SQL queries need to be updated manually (according to their associated business logic requirements).

Write-behind cache

The Persistence Context acts as a transactional write-behind cache, deferring entity state flushing up until the last possible moment.

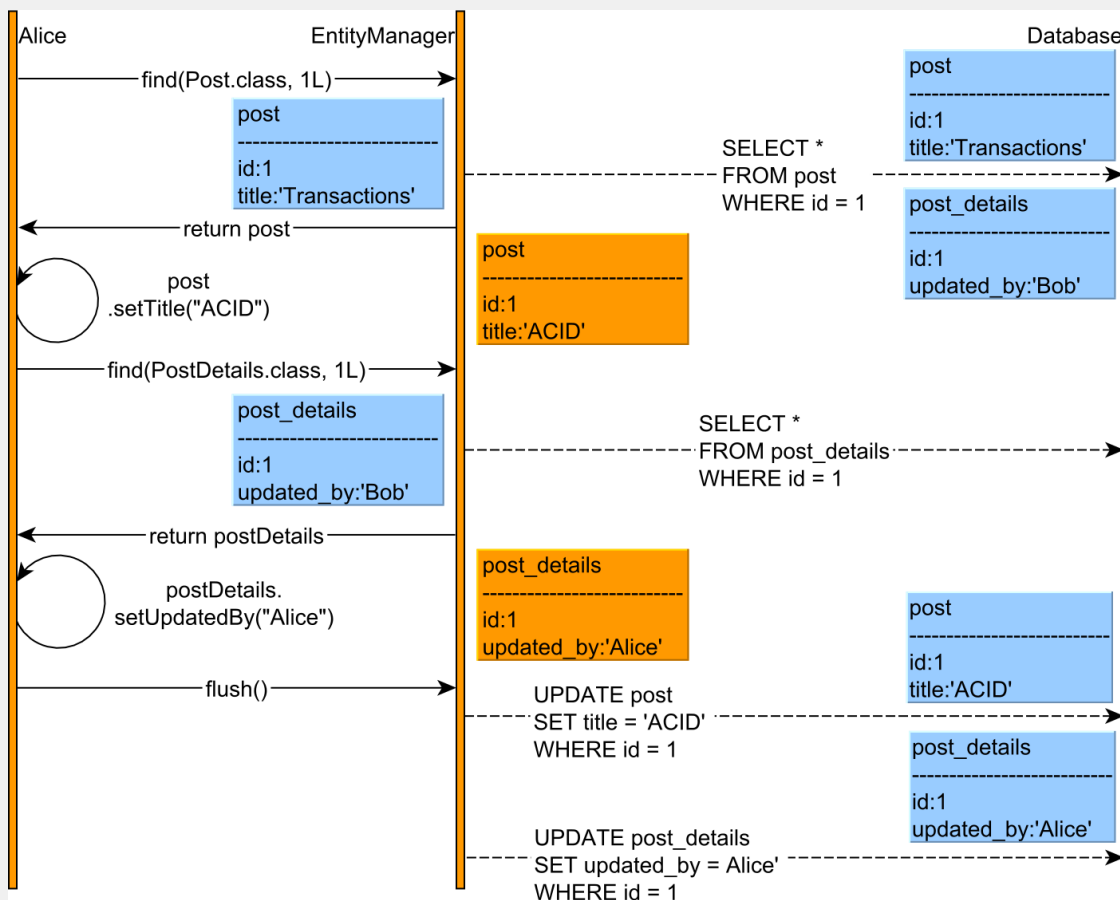


Figure 7.6: Persistence Context

Because every modifying DML statement requires locking (to prevent dirty writes), the write behind cache can reduce the database lock acquisition interval, therefore increasing concurrency.



However, caches introduce consistency challenges, and the Persistence Context requires a flush prior to executing any JPQL or native SQL query, as otherwise it might break the read-your-own-write consistency guarantee.

As detailed in the following chapters, Hibernate does not automatically flush pending changes when a native query is about to be executed, and the application developer must explicitly instruct what database tables are needed to be synchronized.

Transparent statement batching

Since all changes are being flushed at once, Hibernate may benefit from batching JDBC statements. Batch updates can be enabled transparently, even after the data access logic has been implemented. Most often, performance tuning is postponed until the system is already running in production, and switching to batching statements should not require a considerable development effort.



With just one configuration, Hibernate can execute all prepared statements in batches.

Application-level concurrency control

As previously explained, no database isolation level can protect against losing updates when executing a multi-request long conversation. JPA supports both optimistic and pessimistic locking.

The JPA optimistic locking mechanism allows preventing lost updates because it imposes a *happens-before* event ordering. However, in multi-request conversations, optimistic locking requires maintaining old entity snapshots, and JPA makes it possible through *Extended Persistence Contexts* or detached entities.



A Java EE application server can preserve a given Persistence Context across several web requests, therefore providing application-level repeatable reads. However, this strategy is not free since the application developer must make sure the Persistence Context is not bloated with too many entities, which, apart from consuming memory, it can also affect the performance of the Hibernate default dirty checking mechanism.

Even when not using Java EE, the same goal can be achieved using detached entities, which provide fine-grained control over the amount of data needed to be preserved from one web request to the other. JPA allows merging detached entities, which rebecome managed and automatically synchronized with the underlying database system.

JPA also supports a pessimistic locking query abstraction, which comes in handy when using lower-level transaction isolation modes.



Hibernate has a native pessimistic locking API, which brings support for timing out lock acquisition requests or skipping already acquired locks.

2.6 Read-based optimizations

Following the SQL standard, the JDBC `ResultSet` is a tabular representation of the underlying fetched data. The Domain Model being constructed as an entity graph, the data access layer must transform the flat `ResultSet` into a hierarchical structure.

Although the goal of the ORM tool is to reduce the gap between the object-oriented Domain Model and its relational counterpart, it is very important to remember that the source of data is not an in-memory repository, and the fetching behavior influences the overall data access efficiency.



The database cannot be abstracted out of this context, and pretending that entities can be manipulated just like any other plain objects is very detrimental to application performance. When it comes to reading data, the impedance mismatch becomes even more apparent, and, for performance reasons, it is mandatory to keep in mind the SQL statements associated with every fetching operation.

In the following example, the `posts` records are fetched along with all their associated `comments`. Using JDBC, this task can be accomplished using the following code snippet:

```
doInJDBC(connection -> {
    try (PreparedStatement statement = connection.prepareStatement(
        "SELECT * " +
        "FROM post AS p " +
        "JOIN post_comment AS pc ON p.id = pc.post_id " +
        "WHERE " +
        "    p.id BETWEEN ? AND ? + 1"
    )) {
        statement.setInt(1, id);
        statement.setInt(2, id);
        try (ResultSet resultSet = statement.executeQuery()) {
            List<Post> posts = toPosts(resultSet);
            assertEquals(expectedCount, posts.size());
        }
    } catch (SQLException e) {
        throw new DataAccessException(e);
    }
});
```

When joining *many-to-one* or *one-to-one* associations, each `ResultSet` record corresponds to a pair of entities, so both the parent and the child can be resolved in each iteration. For *one-to-many* or *many-to-many* relationships, because of how the SQL join works, the `ResultSet` contains a duplicated parent record for each associated child.

Constructing the hierarchical entity structure requires manual `ResultSet` transformation, and, to resolve duplicates, the parent entity references are stored in a `Map` structure.

```
List<Post> toPosts(ResultSet resultSet) throws SQLException {
    Map<Long, Post> postMap = new LinkedHashMap<>();
    while (resultSet.next()) {
        Long postId = resultSet.getLong(1);
        Post post = postMap.get(postId);
        if (post == null) {
            post = new Post(postId);
            postMap.put(postId, post);
            post.setTitle(resultSet.getString(2));
            post.setVersion(resultSet.getInt(3));
        }
        PostComment comment = new PostComment();
        comment.setId(resultSet.getLong(4));
        comment.setReview(resultSet.getString(5));
        comment.setVersion(resultSet.getInt(6));
        post.addComment(comment);
    }
    return new ArrayList<>(postMap.values());
}
```

The JDBC 4.2 `PreparedStatement` supports only positional parameters, and the first ordinal starts from 1. JPA allows named parameters as well, which are especially useful when a parameter needs to be referenced multiple times, so the previous example can be rewritten as follows:

```
doInJPA(entityManager -> {
    List<Post> posts = entityManager.createQuery(
        "select distinct p " +
        "from Post p " +
        "join fetch p.comments " +
        "where " +
        "    p.id BETWEEN :id AND :id + 1", Post.class)
    .setParameter("id", id)
    .getResultList();
    assertEquals(expectedCount, posts.size());
});
```

In both examples, the object-relation transformation takes place either implicitly or explicitly. In the JDBC use case, the associations must be manually resolved, while JPA does it automatically (based on the entity schema).

The fetching responsibility

Besides mapping database columns to entity attributes, the entity associations can also be represented in terms of object relationships. More, the fetching behavior can be hard-wired to the entity schema itself, which is most often a terrible thing to do.

Fetching multiple one-to-many or many-to-many associations is even more problematic because they might require a Cartesian Product, therefore affecting performance. Controlling the hard-wired schema fetching policy is cumbersome as it prevents overriding an eager retrieval with a lazy loading mechanism.



Each business use case has different data access requirements, and one policy cannot anticipate all possible use cases, so the **fetching strategy should always be set up on a query basis**.

Prefer projections for read-only views

Although it is very convenient to fetch entities along with all their associated relationships, it is better to take into consideration the performance impact as well. As previously explained, fetching too much data is not suitable because it increases the transaction response time.

In reality, not all use cases require loading entities, and not all read operations need to be served by the same fetching mechanism. Sometimes a custom projection (selecting only a few columns from an entity) is much more suitable, and the data access logic can even take advantage of database-specific SQL constructs that might not be supported by the JPA query abstraction.



As a rule of thumb, fetching entities is suitable when the logical transaction requires modifying them, even if that only happens in a successive web request. With this in mind, it is much easier to reason on which fetching mechanism to employ for a given business logic use case.

The second-level cache

If the Persistence Context acts as a transactional write-behind cache, its lifetime will be bound to that of a logical transaction. For this reason, the Persistence Context is also known as the first-level cache, and so it cannot be shared by multiple concurrent transactions.

On the other hand, the second-level cache is associated with an `EntityManagerFactory`, and all Persistence Contexts have access to it. The second-level cache can store entities as well as entity associations (one-to-many and many-to-many relationships) and even entity query results. Because JPA does not make it mandatory, each provider takes a different approach to caching (as opposed to EclipseLink, by default, Hibernate disables the second-level cache).

Most often, caching is a trade-off between consistency and performance. Because the cache becomes another source of truth, inconsistencies might occur, and they can be prevented only when all database modifications happen through a single `EntityManagerFactory` or a synchronized distributed caching solution. In reality, this is not practical since the application might be clustered on multiple nodes (each one with its own `EntityManagerFactory`) and the database might be accessed by multiple applications.



Although the second-level cache can mitigate the entity fetching performance issues, it requires a distributed caching implementation, which might not elude the network-ing penalties anyway.

2.7 Wrap-up

Bridging two highly-specific technologies is always a difficult problem to solve. When the enterprise system is built on top of an object-oriented language, the object-relational impedance mismatch becomes inevitable. The ORM pattern aims to close this gap although it cannot completely abstract it out.

In the end, all the communication flows through JDBC and every execution happens in the database engine itself. A high-performance enterprise application must resonate with the underlying database system, and the ORM tool must not disrupt this relationship.

Just like the problem it tries to solve, Hibernate is a very complex framework with many subtleties that require a thorough knowledge of both database systems, JDBC, and the framework itself. This chapter is only a summary, meant to present JPA and Hibernate into a different perspective that prepares the reader for high-performance object-relational mapping. There is no need to worry if some topics are not entirely clear because the upcoming chapters analyze all these concepts in greater detail.

3. Why jOOQ matters

When working with a relational database, it all boils down to SQL statements.

As previously explained, Hibernate entity queries are suitable for read-write logical transactions. For reporting, analytics or ETL (Extract, Transform, Load) native SQL queries are the best choice since they can take advantage of database-specific features like window functions or Common Table Expressions. Even for CRUD operations, there might be times when a database-specific syntax is more suitable like it's the case for the [upsert SQL operation](#)¹.

While Hibernate does a very good job to automate the vast majority of statements, it is unlikely that you can rely on Hibernate alone for every business use case. Therefore, native queries are a necessity for most enterprise applications.

As demonstrated in the [Native query DTO projection](#) section, both JPA and Hibernate provide a way to execute native SQL statements. Being able to execute any SQL statement is great, but, unfortunately, the JPA approach is limited to static statements only. To build native SQL statement dynamically, JPA and Hibernate are no longer enough.

3.1 How jOOQ works

jOOQ is a query builder framework that allows you generate a great variety of database-specific statements using a Java API. The `DSLContext` is the starting point to building any SQL statement, and it requires two things:

- a reference to a `JDBC Connection`
- a database dialect so that it can translate the Java API query representation into a database-specific SQL query

For instance, when using PostgreSQL 9.5, the `DSLContext` can be constructed as follows:

```
DSLContext sql = DSL.using(connection, SQLDialect.POSTGRES_9_5);
```

3.2 DML statements

With the `DSLContext` in place, it's time to show some simple DML statements like insert, update, delete, as well as a trivial select query. What's worth noticing is that the Java API syntax is almost identical to its SQL counterpart, so most jOOQ queries are self-describing.

¹[https://en.wikipedia.org/wiki/Merge_\(SQL\)](https://en.wikipedia.org/wiki/Merge_(SQL))

To delete all records for the `post` table, the following jOOQ statement must be used:

```
sql
.deleteFrom(table("post"))
.execute();
```

Which translates to the following SQL statement:

```
DELETE FROM post
```

To insert a new `post` table row, the following jOOQ statement can be used:

```
assertEquals(1, sql
    .insertInto(table("post")).columns(field("id"), field("title"))
    .values(1L, "High-Performance Java Persistence")
    .execute()
);
```

Just like in JDBC, the `execute` method return the affected row count for the current `insert`, `update`, or `delete` SQL statement.

When running the previous jOOQ query, the following SQL statement is being executed:

```
INSERT INTO post (id, title)
VALUES (1, 'High-Performance Java Persistence')
```

When updating the previously inserted record:

```
sql
.update(table("post"))
.set(field("title"), "High-Performance Java Persistence Book")
.where(field("id").eq(1))
.execute();
```

JOOQ generates the following SQL statement:

```
UPDATE post
SET title = 'High-Performance Java Persistence Book'
WHERE id = 1
```

Selecting the previously updated record is just as easy:

```
assertEquals("High-Performance Java Persistence Book", sql
    .select(field("title"))
    .from(table("post"))
    .where(field("id").eq(1))
    .fetch().getValue(0, "title")
);
```

To execute the statement and return the SQL query result set, the `fetch` method must be used. As expected, the previous jOOQ query generates the following SQL statement:

```
SELECT title FROM post WHERE id = 1
```

3.3 Java-based schema

All the previous queries were referencing the database schema explicitly, like the table name or the table columns. However, just like JPA defines a Metamodel API for Criteria queries, jOOQ allows generating a Java-based schema that mirrors the one in the database.

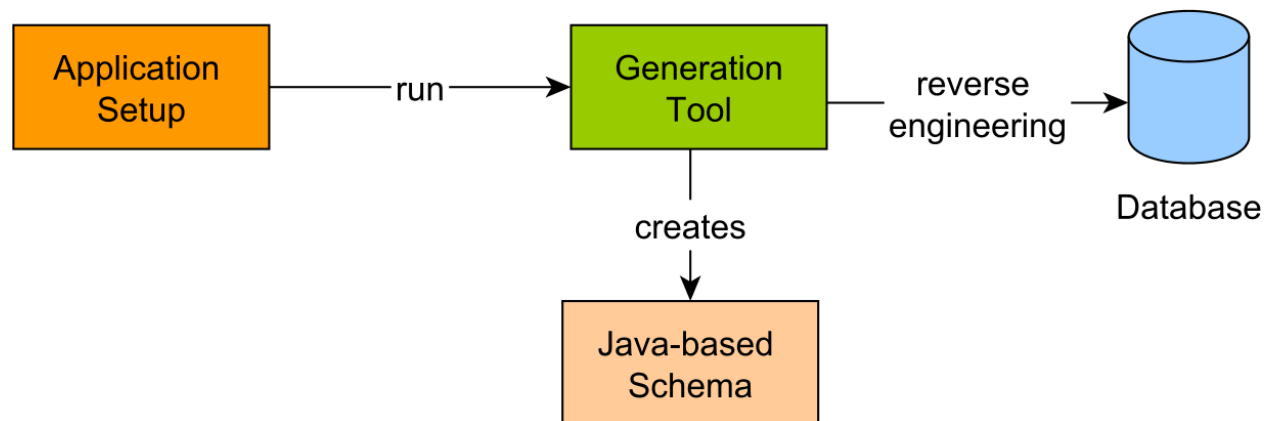


Figure 17.1: Java-based schema generation

There are many advantages to having access to the underlying database schema right from the Java data access layer. For instance, when executing a database stored procedure, the argument types can be bound at compile-time. The same argument holds for query parameters or the result set obtained from running a particular query.

When a column name needs to be modified, there is no risk of forgetting to update a given jOOQ statement because a Java-based schema violation will prevent the application from compiling properly. From a development perspective, the Java-based schema enables the IDE to autocomplete jOOQ queries, therefore increasing productivity and reducing the likelihood of typos.

After generating the Java-based schema, the application developer can use it to build any type-safe jOOQ query.

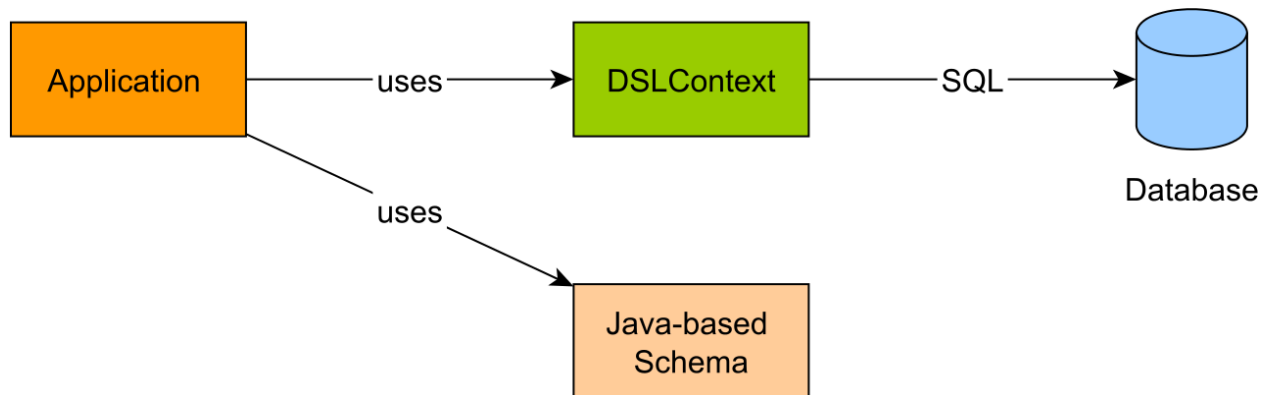


Figure 17.2: Typesafe schema usage

To rewrite the previous DML statements to use the Java-based schema, the generated schema classes need to be imported first:

```
import static com.vladmihalcea.book.hpjp.jooq.pgsql.schema.Tables.POST;
```

With the Java-based schema in place, the previous DML statements become even more descriptive:

```
sql.deleteFrom(POST).execute();

assertEquals(1, sql
    .insertInto(POST).columns(POST.ID, POST.TITLE)
    .values(1L, "High-Performance Java Persistence")
    .execute()
);

sql
    .update(POST)
    .set(POST.TITLE, "High-Performance Java Persistence Book")
    .where(POST.ID.eq(1L))
    .execute();

assertEquals("High-Performance Java Persistence Book", sql
    .select(POST.TITLE)
    .from(POST)
    .where(POST.ID.eq(1L))
    .fetch().getValue(0, POST.TITLE)
);
```



Although jOOQ can work just fine without a Java-based schema, it is much more practical to use typesafe queries whenever possible.

3.4 Upsert

In database terminology, an `upsert` statement is a mix between an insert and an update statement. First, the insert statement is executed and if it succeeds, the operation returns successfully. If the insert fails, it means that there is already a database row matching the same unique constraints with the insert statement. In this case, an update is issued against the database row that matches the given filtering criteria.

The SQL:2003 and SQL:2008 standards define the `MERGE` statement, which among other scenarios, it can be used to emulate the upsert operation. However, `MERGE` acts more like an if-then-else statement, therefore being possible to combine insert, update, and delete statements. While `upsert` implies the same database table, `MERGE` can also be used to synchronize the content of two different tables.

Oracle and SQL Server implement the `MERGE` operation according to the standard specification, whereas MySQL and PostgreSQL provide only an implementation for the upsert operation.

jOOQ implements the upsert operation, therefore, managing to translate the Java-based query to the underlying database-specific SQL syntax.

To visualize how upsert works, consider the following method which aims to insert a `post_details` record if there is none, or to update the existing record if there is already a row with the same primary key:

```
public void upsertPostDetails(
    DSLContext sql, BigInteger id, String owner, Timestamp timestamp) {
    sql
        .insertInto(POST_DETAILS)
        .columns(POST_DETAILS.ID, POST_DETAILS.CREATED_BY, POST_DETAILS.CREATED_ON)
        .values(id, owner, timestamp)
        .onDuplicateKeyUpdate()
        .set(POST_DETAILS.UPDATED_BY, owner)
        .set(POST_DETAILS.UPDATED_ON, timestamp)
        .execute();
}
```

Two users, *Alice* and *Bob*, are going to execute the `upsertPostDetails` method concomitantly, and, because of the upsert logic, the first user is going to insert the record while the second one is going to update it, without throwing any exception:

```
executeAsync(() -> {
    upsertPostDetails(sql, BigInteger.valueOf(1), "Alice",
        Timestamp.from(LocalDateTime.now().toInstant(ZoneOffset.UTC)));
});
executeAsync(() -> {
    upsertPostDetails(sql, BigInteger.valueOf(1), "Bob",
        Timestamp.from(LocalDateTime.now().toInstant(ZoneOffset.UTC)));
});
```

JOOQ is going to translate the upsert Java-based operation to the specific syntax employed by the underlying relational database.

3.4.1 Oracle

On Oracle, jOOQ uses the `MERGE` statement to implement the upsert logic:

```
MERGE INTO "POST_DETAILS" USING
    (SELECT 1 "one" FROM dual) ON ("POST_DETAILS"."ID" = 1)
WHEN MATCHED THEN
    UPDATE SET
        "POST_DETAILS"."UPDATED_BY" = 'Alice',
        "POST_DETAILS"."UPDATED_ON" = '2016-08-11 12:19:48.22'
WHEN NOT MATCHED THEN
    INSERT ("ID", "CREATED_BY", "CREATED_ON")
    VALUES (1, 'Alice', '2016-08-11 12:19:48.22')
```

```
MERGE INTO "POST_DETAILS" USING
    (SELECT 1 "one" FROM dual) ON ("POST_DETAILS"."ID" = 1)
WHEN MATCHED THEN
    UPDATE SET
        "POST_DETAILS"."UPDATED_BY" = 'Bob',
        "POST_DETAILS"."UPDATED_ON" = '2016-08-11 12:19:48.442'
WHEN NOT MATCHED THEN
    INSERT ("ID", "CREATED_BY", "CREATED_ON")
    VALUES (1, 'Bob', '2016-08-11 12:19:48.442')
```

3.4.2 SQL Server

Just like with Oracle, jOOQ uses `MERGE` to implement the upsert operation on SQL Server:

```
MERGE INTO [post_details] USING
  (SELECT 1 [one]) AS dummy_82901439([one]) ON [post_details].[id] = 1
WHEN MATCHED THEN
  UPDATE SET
    [post_details].[updated_by] = 'Alice',
    [post_details].[updated_on] = '2016-08-11 12:36:33.458'
WHEN NOT MATCHED THEN
  INSERT ([id], [created_by], [created_on])
  VALUES (1, 'Alice', '2016-08-11 12:36:33.458')
```

```
MERGE INTO [post_details] USING
  (SELECT 1 [one]) AS dummy_82901439([one]) ON [post_details].[id] = 1
WHEN MATCHED THEN
  UPDATE SET
    [post_details].[updated_by] = 'Bob',
    [post_details].[updated_on] = '2016-08-11 12:36:33.786'
WHEN NOT MATCHED THEN
  INSERT ([id], [created_by], [created_on])
  VALUES (1, 'Bob', '2016-08-11 12:36:33.786')
```

3.4.3 PostgreSQL

As opposed to Oracle and SQL Server, PostgreSQL offers the `ON CONFLICT` clause, which jOOQ uses for implementing upsert:

```
INSERT INTO "post_details" ("id", "created_by", "created_on")
VALUES (1, 'Alice', CAST('2016-08-11 12:56:01.831' AS timestamp))
ON CONFLICT ("id") DO
UPDATE SET
  "updated_by" = 'Alice',
  "updated_on" = CAST('2016-08-11 12:56:01.831' AS timestamp)
```

```
INSERT INTO "post_details" ("id", "created_by", "created_on")
VALUES (1, 'Bob', CAST('2016-08-11 12:56:01.865' AS timestamp))
ON CONFLICT ("id") DO
UPDATE SET
  "updated_by" = 'Bob',
  "updated_on" = CAST('2016-08-11 12:56:01.865' AS timestamp)
```

3.4.4 MySQL

Almost identical to PostgreSQL, MySQL uses the `ON DUPLICATE KEY` for upsert:

```
INSERT INTO `post_details` (`id`, `created_by`, `created_on`)
VALUES (1, 'Alice', '2016-08-11 13:27:53.898')
ON DUPLICATE KEY
UPDATE
  `post_details`.`updated_by` = 'Alice',
  `post_details`.`updated_on` = '2016-08-11 13:27:53.898'

INSERT INTO `post_details` (`id`, `created_by`, `created_on`)
VALUES (1, 'Bob', '2016-08-11 13:27:53.905')
ON DUPLICATE KEY
UPDATE
  `post_details`.`updated_by` = 'Bob',
  `post_details`.`updated_on` = '2016-08-11 13:27:53.905'
```

3.5 Batch updates

As [previously explained](#), JDBC batching plays a very important role in tuning the data access layer write operation performance. While Hibernate offers automated JDBC batching, for entities using identity columns, insert statements do not benefit from this feature. This is because Hibernate requires the entity identifier upon persisting the entity, and the only way to know the identity column value is to execute the insert statement.

Instead of implementing an automatic entity state management mechanism like Hibernate, jOOQ takes a WYSIWYG (what you see is what you get) approach to persistence. Even if nowadays many relational database systems offer sequences (Oracle, SQL Server 2012, PostgreSQL, MariaDB), the identity generator is still the only viable option for MySQL (e.g. `AUTO_INCREMENT`). However, since MySQL has a very significant market share, it is important to know that, with jOOQ, JDBC batching works just fine with insert statements.

To batch the insert statements associated to three `Post` entries, jOOQ offers the following API:

```
BatchBindStep batch = sql.batch(sql.insertInto(POST, POST.TITLE).values("?"));

for (int i = 0; i < 3; i++) {
    batch.bind(String.format("Post no. %d", i));
}
int[] insertCounts = batch.execute();
```

Running this test case on MySQL, jOOQ generates the following output:

```
INSERT INTO `post` (`title`) VALUES (Post no. 0), (Post no. 1), (Post no. 2)
```

As illustrated, jOOQ manages to batch all inserts in a single database roundtrip.



When using Hibernate with MySQL and need to perform lots of inserts, it is a good idea to execute the batch inserts with jOOQ.

3.6 Inlining bind parameters

By default, just like Hibernate, jOOQ uses `PreparedStatement(s)` and bind parameter values. This is a very good default strategy since prepared statements can benefit from statement caching, as [previously explained](#).

However, every rule has an exception. Because the bind parameter values might influence the execution plan, reusing a cached plan might be suboptimal in certain scenarios. Some database systems use statistics to monitor the efficiency of a cached execution plan, but the automatic adjustment process might take time.

For this reason, it is not uncommon to want to bypass the execution plan cache for certain queries that take skewed bind parameter values. Because the query string forms the cache key, by inlining the bind parameter values into the SQL statement, it is for sure that the database will either generate a new plan or pick the cached execution plan that was generated for the very same SQL statement.

This workaround can address the issue when bind parameter values are skewed, but it requires building the SQL statement dynamically. The worst thing to do would be to start concatenating string fragments and risk SQL injection attacks. Fortunately, jOOQ offers a way to inline the bind parameters right into the SQL statements without exposing the data access layer to SQL injection vulnerabilities. The jOOQ API ensures the bind parameter values match the expected bind parameter types.

Because by default jOOQ relies on `PreparedStatement(s)`, to switch to using an inlined `Statement`, it is required to provide the following setting upon creating the `DSLContext`:

```
DSLContext sql = DSL.using(connection, sqlDialect(),  
    new Settings().withStatementType(StatementType.STATIC_STATEMENT));
```

Afterward, when executing a parameterized query:

```
List<String> titles = sql
    .select(POST.TITLE)
    .from(POST)
    .where(POST.ID.eq(1L))
    .fetch(POST.TITLE);
```

jOOQ is going to inline all bind parameter values into the SQL statement `String`:

```
SELECT `post`.`title`
FROM   `post`
WHERE  `post`.`id` = 1
```

Without supplying the `StatementType.STATIC_STATEMENT` setting, when using *datasource-proxy* to intercept the executed SQL statement, the actual executed statement looks as follows:

```
Query: ["select `post`.`title` from `post` where `post`.`id` = ?"],
Params: [(1)]
```

Previous Hibernate and jOOQ SQL snippet format

Although most SQL snippets generated by Hibernate or jOOQ in this book give the impression that bind parameters are inlined, that was just for readability sake since the bind parameter values were manually inlined after extracting the SQL statement from the actual logs.

In reality, all Hibernate statements as well as all jOOQ statements using the default `StatementType.PREPARED_STATEMENT` type are using bind parameter placeholders as illustrated in the aforementioned *datasource-proxy* output.

3.7 Complex queries

In the [Native query DTO projection](#) section, there was an SQL query using Window Functions, Derived Tables, and Recursive CTE (Common Table Expressions). Not only that it's possible to rewrite the whole query in Java, but that can be done programmatically.

The `postCommentScores` method shows how Derived Tables and Window Functions work with jOOQ. In fact, the jOOQ API resembles almost identically the actual SQL statement.

```

public List<PostCommentScore> postCommentScores(Long postId, int rank) {
    return doInJOOQ(sql -> {
        return sql
            .select(field(name(TSG, "id"), Long.class),
                field(name(TSG, "parent_id"), Long.class),
                field(name(TSG, "review"), String.class),
                field(name(TSG, "created_on"), Timestamp.class),
                field(name(TSG, "score"), Integer.class)
            )
            .from(sql
                .select(field(name(ST, "id")), field(name(ST, "parent_id")),
                    field(name(ST, "review")), field(name(ST, "created_on")),
                    field(name(ST, "score")),
                    denseRank().over(orderBy(field(name(ST, "total_score")).desc()))
                    .as("rank"))
                .from(sql
                    .select(field(name(SBC, "id")),
                        field(name(SBC, "parent_id")), field(name(SBC, "review")),
                        field(name(SBC, "created_on")), field(name(SBC, "score")),
                        sum(field(name(SBC, "score"), Integer.class))
                            .over(partitionBy(field(name(SBC, "root_id"))))
                            .as("total_score")
                    )
                .from(sql
                    .withRecursive(withRecursiveExpression(sql, postId))
                    .select(field(name(PCS, "id")),
                        field(name(PCS, "parent_id")),
                        field(name(PCS, "root_id")), field(name(PCS, "review")),
                        field(name(PCS, "created_on")),
                        field(name(PCS, "score"))
                    )
                    .from(table(PCS)).asTable(SBC)
                ).asTable(ST)
            )
            .orderBy(
                field(name(ST, "total_score")).desc(),
                field(name(ST, "id")).asc()
            ).asTable(TSG)
        )
        .where(field(name(TSG, "rank"), Integer.class).le(rank))
        .fetchInto(PostCommentScore.class);
    });
}

```

Because following a very large query is sometimes difficult, with jOOQ, it's fairly easy to break a query into multiple building blocks. In this particular example, the `WITH RECURSIVE` query is encapsulated in its own method. Aside from readability, it is possible to reuse the `withRecursiveExpression` query method for other use cases, therefore reducing the likelihood of code duplication.

```
private CommonTableExpression<Record7<Long, Long, Long, Long, String, Timestamp,
Integer>> withRecursiveExpression(DSLContext sql, Long postId) {
    return name(POST_COMMENT_SCORE).fields("id", "root_id", "post_id",
        "parent_id", "review", "created_on", "score")
        .as(sql.select(
            POST_COMMENT.ID, POST_COMMENT.ID, POST_COMMENT.POST_ID,
            POST_COMMENT.PARENT_ID, POST_COMMENT.REVIEW,
            POST_COMMENT.CREATED_ON, POST_COMMENT.SCORE)
            .from(POST_COMMENT)
            .where(POST_COMMENT.POST_ID.eq(postId)
                .and(POST_COMMENT.PARENT_ID.isNull()))
            .unionAll(
                sql.select(
                    POST_COMMENT.ID,
                    field(name("post_comment_score", "root_id"), Long.class),
                    POST_COMMENT.POST_ID, POST_COMMENT.PARENT_ID,
                    POST_COMMENT.REVIEW, POST_COMMENT.CREATED_ON,
                    POST_COMMENT.SCORE)
                    .from(POST_COMMENT)
                    .innerJoin(table(name(POST_COMMENT_SCORE)))
                    .on(POST_COMMENT.PARENT_ID.eq(
                        field(name(POST_COMMENT_SCORE, "id"), Long.class)))
                    .where(POST_COMMENT.PARENT_ID.eq(
                        field(name(POST_COMMENT_SCORE, "id"), Long.class)))
                )
            );
}
```

To fetch the list of `PostCommentScore` entries, the application developer just has to call the `postCommentScores` method. However, the application requires the `PostCommentScore` entries to be arranged in a tree-like structure based on the `parentId` attribute. This was also the case with Hibernate, and that was the reason for providing a custom `ResultTransformer`. Therefore, a `PostCommentScoreRootTransformer` is added for the jOOQ query as well.

```
List<PostCommentScore> postCommentScores = PostCommentScoreRootTransformer.
    INSTANCE.transform(postCommentScores(postId, rank));
```

The `PostCommentScoreRootTransformer` class is almost identical to the `PostCommentScoreResultTransformer` used in the Hibernate Fetching chapter.

```
public class PostCommentScoreRootTransformer {

    public static final PostCommentScoreRootTransformer INSTANCE =
        new PostCommentScoreRootTransformer();

    public List<PostCommentScore> transform(
        List<PostCommentScore> postCommentScores) {
        Map<Long, PostCommentScore> postCommentScoreMap = new HashMap<>();
        List<PostCommentScore> roots = new ArrayList<>();

        for (PostCommentScore postCommentScore : postCommentScores) {
            Long parentId = postCommentScore.getParentId();
            if (parentId == null) {
                roots.add(postCommentScore);
            } else {
                PostCommentScore parent = postCommentScoreMap.get(parentId);
                if (parent != null) {
                    parent.addChild(postCommentScore);
                }
            }
            postCommentScoreMap.putIfAbsent(
                postCommentScore.getId(), postCommentScore);
        }
        return roots;
    }
}
```

3.8 Stored procedures and functions

When it comes to calling stored procedures or user-defined database functions, jOOQ is probably the best tool for this job. Just like it scans the database metadata and builds a Java-based schema, jOOQ is capable of generating Java-based stored procedures as well.

For example, the previous query can be encapsulated in a stored procedure which takes the `postId` and the `rankId` and returns a `REFCURSOR` which can be used to fetch the list of `PostCommentScore` entries.

```

CREATE OR REPLACE FUNCTION post_comment_scores(postId BIGINT, rankId INT)
RETURNS REFCURSOR AS
$BODY$
    DECLARE
        postComments REFCURSOR;
    BEGIN
        OPEN postComments FOR
            SELECT id, parent_id, review, created_on, score
            FROM (
                SELECT
                    id, parent_id, review, created_on, score,
                    dense_rank() OVER (ORDER BY total_score DESC) rank
                FROM (
                    SELECT
                        id, parent_id, review, created_on, score,
                        SUM(score) OVER (PARTITION BY root_id) total_score
                    FROM (
                        WITH RECURSIVE post_comment_score(id, root_id, post_id,
                            parent_id, review, created_on, score) AS (
                            SELECT
                                id, id, post_id, parent_id, review, created_on,
                                score
                            FROM post_comment
                            WHERE post_id = postId AND parent_id IS NULL
                            UNION ALL
                            SELECT pc.id, pcs.root_id, pc.post_id, pc.parent_id,
                                pc.review, pc.created_on, pc.score
                            FROM post_comment pc
                            INNER JOIN post_comment_score pcs
                            ON pc.parent_id = pcs.id
                            WHERE pc.parent_id = pcs.id
                        )
                        SELECT id, parent_id, root_id, review, created_on, score
                        FROM post_comment_score
                    ) score_by_comment
                ) score_total
            ORDER BY total_score DESC, id ASC
        ) total_score_group
        WHERE rank <= rankId;
    RETURN postComments;
END;
$BODY$ LANGUAGE plpgsql

```

When the Java-based schema was generated, jOOQ has created a `PostCommentScore` class for the `post_comment_scores` PostgreSQL function. The `PostCommentScore` jOOQ utility offers a very trivial API, so calling the `post_comment_scores` function is done like this:

```
public List<PostCommentScore> postCommentScores(Long postId, int rank) {  
    return doInJOOQ(sql -> {  
        PostCommentScores postCommentScores = new PostCommentScores();  
        postCommentScores.setPostid(postId);  
        postCommentScores.setRankid(rank);  
        postCommentScores.execute(sql.configuration());  
        return postCommentScores.getReturnValue().into(PostCommentScore.class);  
    });  
}
```



With jOOQ, calling database stored procedures and user-defined functions is as easy as calling a Java method.

3.9 Streaming

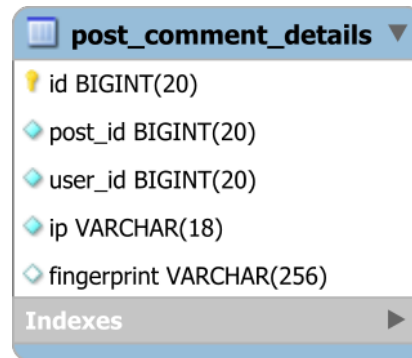
When processing large result sets, it's a good idea to split the whole data set into multiple subsets that can be processed in batches. This way, the memory is better allocated among multiple running threads of execution.

One way to accomplish this task is to split the data set at the SQL level, as explained in the [DTO projection pagination section](#). Streaming is another way of controlling the fetched result set size, and jOOQ makes it very easy to operate on database cursors.

To demonstrate how streaming works, let's consider a forum application which allows only one account for every given user. A fraud detection mechanism must be implemented to uncover users operating on multiple accounts.

To identify a user logins, the IP address must be stored in the database. However, the IP alone is not sufficient since multiple users belonging to the same private network might share the same public IP. For this reason, the application requires additional information to identify each particular user. Luckily, the browser sends all sorts of HTTP headers which can be combined and hashed into a user fingerprint. To make the fingerprint as effective as possible, the application must use the following HTTP headers: User Agent, Content Encoding, Platform, Timezone, Screen Resolution, Language, List of Fonts, etc.

The `user_id`, the `ip` and the `fingerprint` are going to be stored in a `post_comment_details` table. Every time a `post_comment` is being added, a new `post_comment_details` is going to be inserted as well. Because of the one-to-one relationship, the `post_comment` and the “`post_comment_details`” tables can share the same Primary Key.



post_comment_details ▼	
id	BIGINT(20)
post_id	BIGINT(20)
user_id	BIGINT(20)
ip	VARCHAR(18)
fingerprint	VARCHAR(256)
Indexes ▶	

Figure 17.3: The `post_comment_details` table

The fraud detection batch process runs periodically and validates the latest added `post_comment_details`. Because there can be many records to be scanned, a database cursor is used.

JOOQ offers a Java 8 stream API for navigating the underlying database cursor, therefore, the batch process job can be implemented as follows:

```
try (Stream<PostCommentDetailsRecord> stream = sql
    .selectFrom(POST_COMMENT_DETAILS)
    .where(POST_COMMENT_DETAILS.ID.gt(lastProcessedId))
    .stream()) {
    processStream(stream);
}
```



The `try-with-resources` statement is used to ensure that the underlying database stream always gets closed after being processed.

Because there can be thousands of posts added in a day, when processing the stream, a fixed-size `HashMap` is used to prevent the application from running out of memory.

To solve this issue, a custom-made `MaxSizeHashMap` can be used so that it provides a FIFO (first-in, first-out) data structure to hold the current processing data window. Implementing a `MaxSizeHashMap` is pretty straight forward since `java.util.LinkedHashMap` offers a `removeEldestEntry` extension callback which gets called whenever a new element is being added to the `Map`.

```

public class MaxSizeHashMap<K, V> extends LinkedHashMap<K, V> {
    private final int maxSize;

    public MaxSizeHashMap(int maxSize) {
        this.maxSize = maxSize;
    }

    @Override
    protected boolean removeEldestEntry(Map.Entry<K, V> eldest) {
        return size() > maxSize;
    }
}

```

The `IpFingerprint` class is used for associating multiple user ids to a specific IP and fingerprint. Because the `IpFingerprint` object is used as a Map key, the `equals` and `hashCode` methods must be implemented so that they use the associated IP and fingerprint.

```

public class IpFingerprint {
    private final String ip;
    private final String fingerprint;

    public IpFingerprint(String ip, String fingerprint) {
        this.ip = ip;
        this.fingerprint = fingerprint;
    }

    @Override
    public boolean equals(Object o) {
        if (this == o) return true;
        if (o == null || getClass() != o.getClass()) return false;
        IpFingerprint that = (IpFingerprint) o;
        return Objects.equals(ip, that.ip) &&
            Objects.equals(fingerprint, that.fingerprint);
    }

    @Override
    public int hashCode() {
        return Objects.hash(ip, fingerprint);
    }
}

```

With these utilities in place, the `processStream` must create a tree structure that can be navigated as follows: `post_id -> IpFingerprint -> list of user_id`.

The `processStream` method iterates the underlying database cursor and builds a `Map` where the key is the `post_id` and the value is a `Map` of fingerprints and user ids.

```
private void processStream(Stream<PostCommentDetailsRecord> stream) {
    Map<Long, Map<IpFingerprint, List<Long>>> registryMap =
        new MaxSizeHashMap<>(25);

    stream.forEach(postCommentDetails -> {
        Long postId = postCommentDetails
            .get(POST_COMMENT_DETAILS.POST_ID);
        String ip = postCommentDetails
            .get(POST_COMMENT_DETAILS.IP);
        String fingerprint = postCommentDetails
            .get(POST_COMMENT_DETAILS.FINGERPRINT);
        Long userId = postCommentDetails
            .get(POST_COMMENT_DETAILS.USER_ID);

        Map<IpFingerprint, List<Long>> fingerprintsToPostMap =
            registryMap.get(postId);
        if(fingerprintsToPostMap == null) {
            fingerprintsToPostMap = new HashMap<>();
            registryMap.put(postId, fingerprintsToPostMap);
        }

        IpFingerprint ipFingerprint = new IpFingerprint(ip, fingerprint);

        List<Long> userIds = fingerprintsToPostMap.get(ipFingerprint);
        if(userIds == null) {
            userIds = new ArrayList<>();
            fingerprintsToPostMap.put(ipFingerprint, userIds);
        }

        if(!userIds.contains(userId)) {
            userIds.add(userId);
            if(userIds.size() > 1) {
                notifyMultipleAccountFraud(postId, userIds);
            }
        }
    });
}
```

If the `user_id` list contains more than one entry, it means there have been multiple users identified by the same fingerprint, therefore, a notification must be sent to the system administrator.



Even if streaming is a very good fit for processing very large results sets, most often, it is much more appropriate to operate on smaller batches to avoid long-running transactions.

3.10 Keyset pagination

As explained in the [DTO projection pagination section](#), pagination can improve performance since the application only fetches as much data as it's required to be rendered by the current view. The default pagination technique supported by JPA and Hibernate is called the *offset* method, and it is efficient only for small result sets or when navigating the first pages of a large result set. The further the page, the more work is going to be done by the database to fetch the current subset of data. To overcome the offset pagination limitation, the application developer has two alternatives.

The first choice is to narrow down the result set as much as possible using multiple filtering criteria. From a user experience perspective, this is probably the best option as well since the user can select the exact subset of data that she is interested in operating. If the filtered subset is rather small, the offset pagination limitation is not going to be a big issue.

However, if the filtered subset is still large and there is no more filtered that can be further applied, then keyset pagination becomes a better alternative to using the SQL-level offset support. Keyset pagination uses the database table primary key to mark the position of the current fetching data subset.

If JPA 2.1 and Hibernate 5.2 do not offer support for keyset pagination, jOOQ provides a `seek()` method which translate to a database-specific keyset pagination query syntax.

Considering the front page of a forum application which displays all the posts in the descending order of their creation, the application requires a paginated view over the following `PostSummary` records:

 <code>PostSummary</code>	
 <code>PostSummary(Long, String, Timestamp)</code>	
 <code>getId()</code>	Long
 <code>getTitle()</code>	String
 <code>getCreatedOn()</code>	Timestamp

Figure 17.4: The `PostSummary` class

The keyset pagination query is rather trivial as illustrated by the following code snippet:

```
public List<PostSummary> nextPage(String user, int pageSize,
    PostSummary offsetPostSummary) {
    return doInJOOQ(sql -> {
        SelectSeekStep2<Record3<Long, String, Timestamp>, Timestamp, Long>
        selectStep = sql
            .select(POST.ID, POST.TITLE, POST_DETAILS.CREATED_ON)
            .from(POST)
            .join(POST_DETAILS).on(POST.ID.eq(POST_DETAILS.ID))
            .where(POST_DETAILS.CREATED_BY.eq(user))
            .orderBy(POST_DETAILS.CREATED_ON.desc(), POST.ID.desc());
        return (offsetPostSummary != null)
            ? selectStep
                .seek(offsetPostSummary.getCreatedOn(), offsetPostSummary.getId())
                .limit(pageSize)
                .fetchInto(PostSummary.class)
            : selectStep
                .limit(pageSize)
                .fetchInto(PostSummary.class);
    });
}
```

To fetch the first page, the offset `PostSummary` is `null`:

```
List<PostSummary> results = nextPage(pageSize, null);
```

When fetching the first page on PostgreSQL, jOOQ executes the following SQL query:

```
SELECT "post"."id", "post"."title", "post_details"."created_on"
FROM "post"
JOIN "post_details" on "post"."id" = "post_details"."id"
ORDER BY "post_details"."created_on" DESC, "post"."id" DESC
LIMIT 5
```

After fetching a page of results, the last entry becomes the offset `PostSummary` for the next page:

```
PostSummary offsetPostSummary = results.get(results.size() - 1);
results = nextPage(pageSize, offsetPostSummary);
```

When fetching the second page on PostgreSQL, jOOQ executes the following query:

```
SELECT "post"."id", "post"."title", "post_details"."created_on"
FROM "post"
JOIN "post_details" on "post"."id" = "post_details"."id"
WHERE (
  1 = 1 AND
  ("post_details"."created_on", "post"."id") <
  (CAST('2016-08-24 18:29:49.112' AS timestamp), 95)
)
ORDER BY "post_details"."created_on" desc, "post"."id" desc
LIMIT 5
```

On Oracle 11g, jOOQ uses the following SQL query:

```
SELECT "v0" "ID", "v1" "TITLE", "v2" "CREATED_ON"
FROM (
  SELECT "x"."v0", "x"."v1", "x"."v2", rownum "rn"
  FROM (
    SELECT
      "POST"."ID" "v0", "POST"."TITLE" "v1",
      "POST_DETAILS"."CREATED_ON" "v2"
    FROM "POST"
    JOIN "POST_DETAILS" on "POST"."ID" = "POST_DETAILS"."ID"
    WHERE (
      1 = 1 and (
        "POST_DETAILS"."CREATED_ON" <= '2016-08-25 03:04:57.72' and (
          "POST_DETAILS"."CREATED_ON" < '2016-08-25 03:04:57.72' or (
            "POST_DETAILS"."CREATED_ON" =
              '2016-08-25 03:04:57.72' and
            "POST"."ID" < 95
          )
        )
      )
    )
  )
  ORDER BY "v2" desc, "v0" desc
) "x"
WHERE rownum <= (0 + 5)
)
WHERE "rn" > 0
ORDER BY "rn"
```

Because Oracle 11g does not support comparison with row value expressions as well a dedicated SQL operator for limiting the result set, jOOQ must emulate the same behavior, hence the SQL query is more complex than the one executed on PostgreSQL.

Keyset pagination is a very handy feature when the size of the result set is rather large. To get a visualization of the performance gain obtain when switching to keyset pagination, Markus Winand has a [No-Offset article](http://use-the-index-luke.com/no-offset)^a explaining in great detail why offset is less efficient than keyset pagination.

^a<http://use-the-index-luke.com/no-offset>