

---

# Informatica Interview Questions And Answers For 3 Years Experience

---

*Informatica Interview Questions And  
Answers For 3 Years Experience*

*Downloaded from [old.setefer.com](http://old.setefer.com) by  
Cross Castaneda*

---

## **NAVIGATING THE MID-CAREER CHALLENGE: INFORMATICA INTERVIEW QUESTIONS AND ANSWERS FOR 3 YEARS EXPERIENCE**

---

Reaching the three-year mark in your data integration career is a significant milestone. You have moved past the initial learning curve, mastered the fundamentals of ETL development, and are now expected to contribute with greater autonomy and a deeper understanding of system architecture. This is the phase where interviewers stop testing your basic syntax knowledge and start probing your problem-solving abilities, your understanding of performance optimization, and your experience with complex real-world data pipelines.

If you are preparing for a new role, focusing on **Informatica Interview Questions And Answers For 3 Years Experience** is the smartest strategy you can adopt. At this level, you are not just a coder; you are a developer who understands the 'why' behind the 'how'. This article is designed to help you bridge that gap. We will cover the advanced concepts, the tricky scenarios,

and the best-practice approaches that hiring managers are looking for. We will move beyond the textbook definitions and dive into the practical, hands-on knowledge that sets a seasoned Informatica professional apart from a junior developer.

The questions we explore here are drawn from actual interview experiences. They are designed to assess your ability to design efficient mappings, debug performance bottlenecks, handle slowly changing dimensions, and manage the complexities of a production-grade Informatica environment. Whether you are aiming for a senior developer role or a team lead position, mastering these questions will give you a distinct competitive edge.

---

## **UNDERSTANDING THE CORE EXPECTATIONS FOR A THREE-YEAR PROFESSIONAL**

---

Before we jump into the questions, let us set the stage. What exactly does an interviewer expect from someone with three years of experience in Informatica? They expect you to have built and deployed multiple mappings from scratch. They assume you have encountered and resolved common performance issues.

They anticipate that you can work independently, design reusable components, and communicate effectively with business analysts and data architects.

At this level, your answers should demonstrate maturity. You should be able to explain not just what a particular transformation does, but also when to use it and when to avoid it. You should be comfortable discussing alternative approaches and justifying your design choices. The questions below are categorized to reflect these expectations.

# 1. Mapping Design and Advanced Transformations

**Q: You have a requirement to join three large source tables with millions of records each. One approach is to use a Joiner transformation, and another is to use a Source Qualifier override with a SQL join. Which would you choose and why?**

**A:** This is a classic question that tests your understanding of the Informatica processing engine. For three years of experience, you should be able to articulate the trade-offs clearly.

My preference would depend on the specific scenario. If the source is a relational database and the join logic is straightforward (an equi-join on indexed columns), I would typically use a **Source Qualifier override** with a native SQL join. The primary reason is performance. The database is highly optimized for join operations. It can leverage indexes, choose the

most efficient join algorithm (hash join, merge join, nested loop), and reduce the data volume before it ever reaches the Informatica server. This minimizes network traffic and the memory footprint on the Integration Service.

However, I would choose the **Joiner transformation** in specific situations. First, if the sources are from heterogeneous systems (e.g., one from Oracle and one from a flat file), the Joiner is the only option. Second, if the join condition is complex or involves non-equi joins that are difficult to express in SQL, the Joiner provides more flexibility. Third, if I need to perform incremental processing where one source is much smaller than the other, the Joiner's cached master mode can be very efficient. The key is to always make the smaller input the master partition to optimize the cache.

**Q: Describe a scenario where you used a Dynamic Lookup and explain its advantages over a standard Lookup.**

**A:** A Dynamic Lookup is a powerful feature when you are dealing with slowly changing dimensions (SCD), specifically Type 1 updates where you need to insert new records and update existing ones in the same mapping run.

I used a Dynamic Lookup in a project where we were loading customer data from a transactional source into a dimension table. The source had both new customers and updates to existing customer addresses. With a standard Lookup, you would need separate mappings or complex post-processing to handle inserts and updates. With a Dynamic Lookup, the lookup cache is updated as the mapping processes rows. When a new row comes in, it is inserted into the target and simultaneously added to the

cache. When an update row comes in, the lookup finds the existing record in the cache, and you can route the row to an Update strategy transformation for a Type 1 update.

The advantage is significant: it reduces the number of passes over the data and ensures consistency within a single mapping. The key to using it effectively is to understand that the cache is updated in the order the rows are processed. If you have a source sorted by a timestamp, the latest version will be in the cache. For Type 2 SCDs, you would typically use a combination of a standard Lookup and a Router, but for Type 1, the Dynamic Lookup is elegant and efficient.

## 2. Performance Tuning and Optimization

**Q: A mapping that used to run in 30 minutes is now taking over three hours. What is your systematic approach to diagnosing and fixing this performance degradation?**

**A:** Performance tuning is a core competency for anyone with three years of experience. I would approach this in a structured, step-by-step manner.

First, I would check the **session logs** for any obvious errors, warnings, or bottlenecks. I would look for high block counts, large cache sizes, or indications of data spillage to disk. The session log is your best friend for diagnostics.

Second, I would review recent changes. Was the source data

volume increased? Was a new transformation added? Was a database index dropped? I would check the source and target database statistics to ensure they are up to date.

Third, I would analyze the mapping performance using the **Performance Details** view in the Workflow Monitor. This tells me exactly how much time each transformation is taking. The bottleneck is usually one of the following:

- **Source bottleneck:** The Source Qualifier is taking too long. I would check if the SQL override is efficient and if indexes are being used.
- **Target bottleneck:** The target load is slow. I would consider increasing the commit interval, using array-based loading, or disabling constraints during the load.
- **Transformation bottleneck:** A Lookup or Joiner transformation might be spilling to disk. I would check the cache size and increase it if possible. I would also ensure the smaller table is used as the master for the Joiner.
- **Network bottleneck:** If the source and target are on different servers, network latency could be an issue. Pulling only necessary columns in the Source Qualifier helps.

Finally, I would consider partitioning the session if the source data is large and the server has multiple CPUs. Partitioning can parallelize the data flow and significantly reduce run time. My goal is always to identify the single point of failure and address it directly.

**Q: What is the impact of the commit interval on session performance, and how do you determine the optimal**

value?

**A:** The commit interval defines how many rows are processed before a commit is issued to the target database. A small commit interval (e.g., 1,000 rows) means more frequent commits, which increases database overhead and slows down the load. A large commit interval (e.g., 100,000 rows) reduces commit overhead but can lead to large rollback segments and potential session failures if the session crashes mid-load.

For a three-years-experience role, I would explain that there is no one-size-fits-all number. I typically start with a value between 10,000 and 50,000 rows for a large bulk load. I then monitor the database's redo log activity and the session's throughput. If the database is under heavy load, I might increase the interval. If the session is failing due to rollback segment issues, I decrease it. The key is to balance throughput with recoverability. In some data warehouses, loading with a commit interval of 0 (automatic commits at the end of the session) is preferred for full load scenarios, but this carries risk.

### 3. Slowly Changing Dimensions (SCD) and Data Warehousing Concepts

**Q: Explain how you would implement a Type 2 Slowly Changing Dimension in Informatica without using the**

**native SCD wizard. Walk me through the mapping logic.**

**A:** While the SCD wizard is useful for rapid development, understanding the manual implementation is crucial for complex scenarios. This question tests your foundational knowledge.

For a Type 2 SCD, where we track historical changes, the mapping logic is as follows:

1. **Source Qualifier:** Read the source data.
2. **Lookup:** Perform a lookup on the target dimension table using the business key (e.g., Customer ID) to get the existing record. We need to know if the record exists and if any of the tracked attributes have changed.
3. **Expression/Router:** Compare the incoming source attributes with the existing target attributes from the lookup. The Router transformation splits the data into two or more groups:
  - **New records:** The business key does not exist in the target. These are new insertions.
  - **Changed records:** The business key exists, but one or more tracked attributes have changed. These need to be handled for Type 2.
  - **Unchanged records:** The business key exists and the attributes are the same. No action is needed.
4. **Update Strategy (for Changed records):** For the changed records, we need to expire the current active record in the target. This involves updating the End\_Date of the existing record to the current date (or the source effective date). We set the Current\_Flag to 'N' or 0.
5. **Update Strategy (for New records):** For both new

records and the new version of changed records, we set the `Start_Date` to the current date, the `End_Date` to a far future date (e.g., '12/31/9999'), and the `Current_Flag` to 'Y' or 1.

6. **Target:** The target is configured to handle both insert and update operations. The session properties must allow for both row types.

The critical part is ensuring that the expiration and insertion happen in the correct order within the same session. Using a Dynamic Lookup can simplify this, but the manual approach gives you complete control over the logic.

## 4. Error Handling and Debugging

**Q: How do you manage error handling in a production Informatica environment? Give me an example of a robust error handling strategy.**

**A:** At the three-year experience level, you are expected to design

for resilience. A robust error handling strategy includes several layers.

First, within a mapping, I use the **Error Log transformation** to capture rows that fail validation or transformation. For example, if a date conversion fails, I route the bad row to an error table instead of failing the entire session. This allows the load to continue while capturing the problematic data for later analysis.

Second, I use **session-level error handling**. I configure the session to abort on error for critical loads, but for less critical ones, I set it to fail the row and continue. I also set the 'Stop on Errors' property carefully to avoid infinite loops.

Third, I implement **workflow-level error handling**. I use the 'On Error' event to trigger an email notification or a separate recovery workflow. For instance, if a session fails due to a database connection issue, I might have a workflow that waits for 5 minutes and then retries the session up to three times before sending a critical alert to the operations team.

Fourth, I maintain a **custom audit log table** that records every session run, including start time, end time, rows inserted, rows updated, rows rejected, and error messages. This is invaluable for