

Lecture #21: Database Logging

15-445/645 Database Systems (Fall 2025)

<https://15445.courses.cs.cmu.edu/fall2025/>

Carnegie Mellon University

Andy Pavlo

1 Crash Recovery

Recovery algorithms are techniques to ensure database consistency, transaction atomicity, and durability despite failures. When a crash occurs, all the data in volatile memory that has been committed but not yet flushed to disk is at risk of being lost. This is not ideal since the user expects their committed data changes to be persisted. Recovery algorithms act to prevent loss of information after a crash.

Why this matters: Volatile memory is significantly faster than non-volatile storage (like SSDs or HDDs), so DBMSs use volatile memory extensively for performance. However, this creates the recovery challenge - we need smart algorithms that allow us to benefit from fast memory while ensuring we don't lose data.

Every recovery algorithm has two parts:

- Lec#21** • Actions during normal transaction processing to ensure that the DBMS can recover from a failure.
- Lec#22** • Actions after a failure to recover the database to a state that ensures atomicity, consistency, and durability.

The key primitives used in recovery algorithms are UNDO and REDO. Not all algorithms use both primitives.

- ✓ **UNDO:** The process of removing the effects of an incomplete or aborted transaction. *Ctrl+Z*
- ✓ **REDO:** The process of re-applying the effects of a committed transaction for durability. *Ctrl+Y*

2 Buffer Pool Management Policies

The DBMS needs to ensure the following guarantees:

- The changes for any transaction are durable once the DBMS has told somebody that it committed.
- No partial changes are durable if the transaction aborted.

A steal policy dictates whether the DBMS allows an uncommitted transaction to overwrite the most recent committed value of an object in non-volatile storage.

- **STEAL:** The DBMS can write uncommitted changes to disk before the transaction completes. (*Mnemonic: The system "steals" dirty pages from active transactions and writes them to disk.*)
- **NO-STEAL:** The DBMS cannot write uncommitted changes to disk before the transaction completes. (*Mnemonic: The system cannot "steal" dirty pages - they must stay in memory until commit time.*)

A force policy dictates whether the DBMS requires that all updates made by a transaction are reflected on non-volatile storage before the transaction is allowed to commit.

- **FORCE:** All changes must be written to disk at commit time. (*Mnemonic: The system "forces" all changes to disk at commit time.*)

- **NO-FORCE:** Changes are not required to be written to disk at commit time. (*Mnemonic: The system does not "force" changes to disk - they can be written later.*)

These policies create four possible combinations, but only two are commonly used in practice:

- A. **NO-STEAL + FORCE:** Simple recovery but poor runtime performance
- B. **STEAL + NO-FORCE:** Excellent runtime performance but more complex recovery

The other combinations (STEAL + FORCE, NO-STEAL + NO-FORCE) offer few practical advantages and are rarely implemented.

3 Naive NO-STEAL + FORCE A

The simplest buffer pool management policy to implement is called *NO-STEAL + FORCE*. In this policy, the DBMS never has to undo changes of an aborted transaction because the changes were not written to disk. It also never has to redo changes of a committed transaction because all the changes are guaranteed to be written to disk at commit time. An example of NO-STEAL + FORCE is shown in Figure 1.

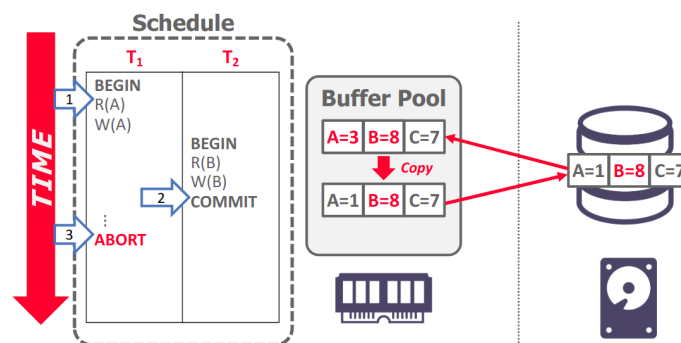


Figure 1: DBMS using NO-STEAL + FORCE Example – All changes from a transaction are only written to disk when the transaction is committed. Once the schedule begins at Step #1, changes from T_1 and T_2 are written to the buffer pool. Because of the FORCE policy, when T_2 commits at Step #2, all of its changes must be written to disk. To do this, the DBMS makes a copy of the page in memory, applies only the changes from T_2 , and writes it back to disk. This is because NO-STEAL forbids the system from writing uncommitted changes from T_1 to disk. At Step #3, it is trivial for the DBMS to rollback T_1 since no dirty changes from T_1 are on disk.

A critical limitation of this naive NO-STEAL + FORCE approach is that **all of the data (i.e., the write set) that a transaction needs to modify must fit into memory**. Otherwise, that transaction cannot execute because the DBMS is not allowed to write out dirty pages to disk before the transaction commits. More frequent writes can also lead to faster wearing of storage devices like SSDs.

4 Shadow Paging (Smarter NO-STEAL + FORCE) A

Shadow Paging is an improvement upon the naive NO-STEAL + FORCE scheme where the DBMS implements a **copy-on-write approach** to maintain two separate versions of the database:

- **master:** Contains only changes from committed transactions.
- **shadow:** Temporary database with changes made from uncommitted transactions.

Updates are only made in the shadow copy. When a transaction commits, the shadow copy is atomically switched to become the new master. The old master is eventually garbage collected. This is an example of a NO-STEAL + FORCE system. A high-level example of shadow paging is shown in Figure 2.

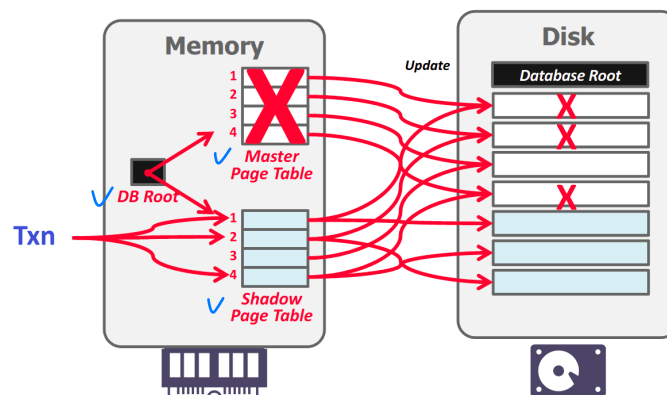


Figure 2: Shadow Paging – The database root points to a master page table which in turn points to the pages on disk (all of which contain committed data). When an updating transaction occurs, a shadow page table is created that points to the same pages as the master. Modifications are made to a temporary space on disk and the shadow table is updated. Read-only transactions will continue to read from the master page table. To complete the commit, the database root pointer is redirected to the shadow table, which becomes the new master. After that, all pages that were modified are now part of the committed database. (Note: For this example, we are assuming only one transaction is updating the database at a time. If multiple concurrent transactions are updating the database, multiple shadow page tables would be needed, and group fitting would be required to merge their changes when redirecting the root pointer on commit.)

Implementation

The DBMS organizes the database pages in a tree structure where the root is a single disk page. There are two copies of the tree, the *master* and *shadow*. The root always points to the current master copy. When a transaction executes, it only makes changes to the shadow copy.

When a transaction wants to commit, the DBMS must install its updates. To do this, it only has to overwrite the root to make it point to the shadow copy of the database, thereby swapping the master and shadow. Before overwriting the root, none of the transaction's updates are part of the disk-resident database. After

overwriting the root, all the transaction's updates are part of the disk resident database. This overwriting of the root can be done atomically.

Recovery

- ✗ **Undo:** Remove the shadow pages. Leave the master and DB root pointer alone.
- ✗ **Redo:** Not needed at all.

Disadvantages

Shadow paging has several key limitations:

- **High overhead:** Copying the entire page table is expensive. In practice, only paths in the tree that lead to updated leaf nodes need to be copied, not the entire tree.
- **Expensive commits:** Commits require the page table, root page, and every updated page to be flushed, causing lots of writes to random non-contiguous pages.
- **Data fragmentation:** Potentially related data gets divided between different pages as updates occur.
- **Concurrency limitations:** While technically possible to support multiple shadow pages for different transactions, this approach becomes increasingly inefficient with concurrent writers.

5 Journal File (SQLite Pre-2010)

When a transaction modifies a page, the DBMS copies the original page to a separate journal file before overwriting the master version. After restarting, if a journal file exists, the DBMS collects the original pages and blindly overwrites the pages with those original pages to restore data to the state before the uncommitted transaction.

The journal file approach allows for STEAL (unlike shadow paging), solving the memory limitation problem by permitting dirty pages to be written to disk before commit. However, it's typically limited to one writer at a time, which was acceptable for SQLite's design goals prior to 2010. After 2010, SQLite switched their implementation to use a write-ahead log instead for better performance.

6 Write-Ahead Logging

With *write-ahead logging*, the DBMS records all the changes made to the database in a log file (on stable storage) before the change is made to a disk page. The log contains sufficient information to perform the necessary undo and redo actions to restore the database after a crash. The DBMS must write to disk the log file records that correspond to changes made to a database object before it can flush that object to disk. An example of WAL is shown in Figure 3. 

WAL is an example of a STEAL + NO-FORCE system:

- **STEAL** allows dirty pages to be written to disk before commit, solving the memory limitation issue
- **NO-FORCE** means we don't need to write all changes to disk at commit time, improving performance
- However, this introduces the need to handle both **UNDO** (for rolled-back transactions that wrote to disk) and **REDO** (for committed transactions whose changes weren't forced to disk)

In shadow paging, the DBMS was required to perform writes to random non-contiguous pages on disk. Write-ahead logging instead works with sequential writes, which are much faster. Thus, almost every modern DBMS uses write-ahead logging (WAL) because it has the fastest runtime performance, even though the DBMS's recovery time with WAL is slower than shadow paging because it has to replay the log.

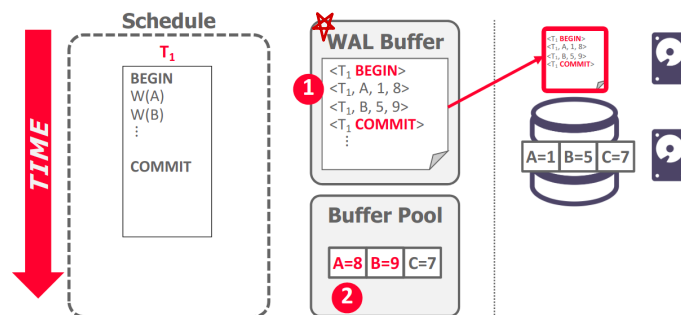


Figure 3: Write Ahead Logging – When the transaction begins, all changes are recorded in the WAL buffer in memory before being made to the buffer pool. At the time of commit, the WAL buffer is flushed out to disk. The transaction result can be written out to disk, once the WAL buffer is safely on disk.

Implementation

The DBMS first stages all of a transaction's log records in volatile storage. All log records pertaining to an updated page are then written to non-volatile storage before the page itself is allowed to be overwritten in non-volatile storage. A transaction is not considered committed until all its log records have been written to stable storage.

When the transaction starts, write a **<BEGIN>** record to the log for each transaction to mark its starting point.

When a transaction finishes, write a **<COMMIT>** record to the log and make sure all log records are flushed before it returns an acknowledgment to the application.

Each log entry contains information necessary to rewind or replay the changes to a single object:

- Transaction ID.
- Object ID.
- Before Value (used for UNDO).
- After Value (used for REDO).
- ...Some system dependent data (e.g. timestamp, checksum)

Note: When using Multi-Version Concurrency Control (MVCC), the "Before Value" may not be needed for UNDO operations since previous versions are already maintained by the system - an example of convenient synergy between concurrency control and recovery mechanisms.

The DBMS must flush all of a transaction's log entries to disk before it can tell the outside world that a transaction has successfully committed. The system can use the "group commit" optimization to batch multiple log flushes together to amortize overhead. (Flushes happen either when the log buffer is full, or if sufficient time has passed between successive flushes.) The DBMS can write dirty pages to disk whenever it wants to, as long as it is after flushing the corresponding log records.

Buffer Pool Policies Tradeoff

Most DBMSs employ the NO-FORCE + STEAL policy due to its superior runtime performance compared to FORCE + NO-STEAL. However, this tradeoff means:

Good → • **Runtime Performance:** STEAL + NO-FORCE is much faster during normal operation

Bad → • **Recovery Speed:** FORCE + NO-STEAL recovers faster after a crash

Change Data Capture

WAL can also serve purposes beyond recovery. Since the log contains a complete sequence of changes to the database, it can be used to propagate changes to external systems. This approach, known as Change Data Capture (CDC), allows other systems to subscribe to database changes and stay synchronized with the primary database. The same log that ensures durability for recovery can efficiently feed data to replicas, data warehouses, or microservices.

7 Logging Schemes

The contents of a log record can vary based on the implementation.

Physical Logging:

- Record the byte-level changes made to a specific location in the database.
- Example: git diff

Logical Logging:

- Record the high level operations executed by transactions.
- Not necessarily restricted to a single page.
- Requires less data written in each log record than physical logging because each record can update multiple tuples over multiple pages. (However, it is difficult to implement recovery with logical logging when there are concurrent transactions in a non-deterministic concurrency control scheme.) Additionally recovery takes longer because you must re-execute every transaction.
- Example: The UPDATE, DELETE, and INSERT queries invoked by a transaction.

Physiological Logging:

- Hybrid approach where log records target a single page but does not specify data organization of the page. That is, identify tuples based on a slot number in the page without specifying exactly where in the page the change is located. Therefore the DBMS can reorganize pages after a log record has been written to disk.
- Most common approach used in DBMSs.

physical to a page
logical within a page

8 Checkpoints

The main problem with a WAL-based DBMS is that the log file will grow forever. Some database systems have been running continuously for decades, potentially accumulating petabytes of logs if not managed properly. After a crash, the DBMS would have to replay the entire log, which can take an impractically long time. 🤯

Thus, the DBMS periodically takes a *checkpoint* where it flushes all buffers out to disk. Checkpoints serve two critical purposes:

- **Log Pruning:** Old log records before the checkpoint can be safely discarded

- **Recovery Efficiency:** The DBMS only needs to perform REDO/UNDO operations from the most recent checkpoint, not from the beginning of time

How often the DBMS should take a checkpoint depends on the application's performance and downtime requirements. (Taking a checkpoint too often causes the DBMS's runtime performance to degrade. But waiting a long time between checkpoints can potentially be just as bad, as the system's recovery time after a restart increases.) Therefore, it is common for DBMSs to provide tunable options for checkpoint frequency based on the application's recovery time requirements.

Blocking Checkpoint Implementation:

- ✓ The DBMS stops accepting new transactions and waits for all active transactions to complete.
- ✓ Flush all log records and dirty blocks currently residing in main memory to stable storage.
- ✓ Write a **<CHECKPOINT>** entry to the log and flush to stable storage.

In this implementation, the DBMS must halt everything when it takes a checkpoint to ensure a consistent snapshot. This is bad for runtime performance but makes recovery straightforward. Meanwhile, scanning the log to find uncommitted transactions is also time consuming.

Tradeoff!

Lecture #22: Database Crash Recovery

15-445/645 Database Systems (Fall 2025)

<https://15445.courses.cs.cmu.edu/fall2025/>

Carnegie Mellon University

Andy Pavlo

1 Crash Recovery

The DBMS relies on its recovery algorithms to ensure database consistency, transaction atomicity, and durability despite failures. Each recovery algorithm is comprised of two parts:

- Actions during normal transaction processing to ensure that the DBMS can recover from a failure
- Actions after a failure to recover the database to a state that ensures the atomicity, consistency, and durability of transactions.

The key to database resilience is the management of transaction integrity and durability, particularly in failure scenarios. This foundational concept sets the stage for the introduction of the ARIES recovery algorithm.

Algorithms for Recovery and Isolation Exploiting Semantics (ARIES) is a recovery algorithm developed at IBM research in early 1990s for the DB2 system.

There are three key concepts in the ARIES recovery protocol:

- **Write Ahead Logging:** Any change is recorded in log on stable storage before the database change is written to disk (STEAL + NO-FORCE). WAL Buffer
- **Repeating History During Redo:** On restart, retrace actions and restore database to exact state before crash.
- **Logging Changes During Undo:** Record undo actions to log to ensure action is not repeated in the event of repeated failures.

2 WAL Records A

Write-ahead log records extend the DBMS's log record format to include a globally unique **log sequence number (LSN)**. All log records have an LSN. A high level diagram of how log records with LSN's are written is shown in Figure 1.

Each data page contains a **pageLSN**, which is the LSN of the most recent update to that page. The pageLSN is updated every time a transaction modifies a record in the page.

The DBMS also keeps track of the max LSN flushed so far (**flushedLSN**). The flushedLSN in memory is updated every time the DBMS writes out the WAL buffer to disk.

Various components in the system keep track of LSNs that pertain to them. Section 2 shows a summary of these LSNs.

3 Normal Execution

Every transaction invokes a sequence of reads and writes, followed by a commit or abort. It is this sequence of events that recovery algorithms must have.

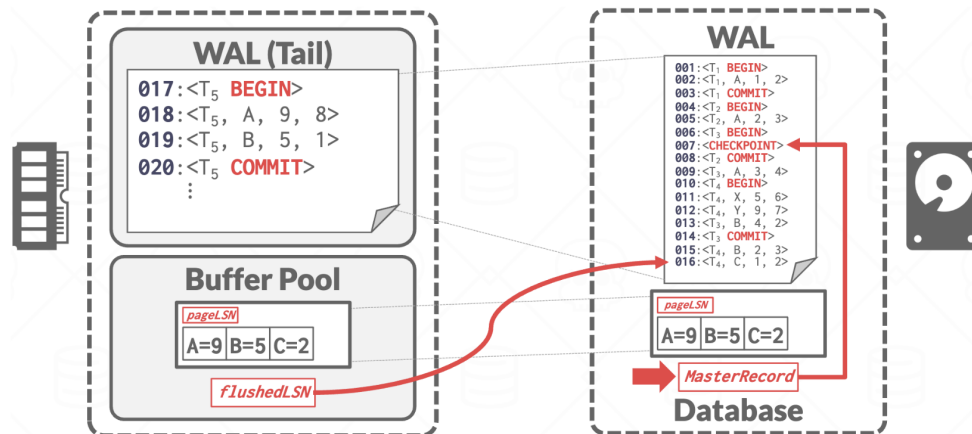


Figure 1: Writing Log Records – Each WAL has a counter of LSNs that is incremented at every step. The page also keeps a pageLSN, which stores the most recent log record that updated this page. The flushedLSN is a pointer to the last LSN that was written out to disk. Before the DBMS can write page i to disk, it must flush log at least to the point where $\text{pageLSN}_i \leq \text{flushedLSN}$. The MasterRecord points to the last successful checkpoint passed. *safe to evict*

Name	Location	Definition
✓ flushedLSN	Memory	Last LSN in log on disk
✓ pageLSN	page _x	Newest update to page _x
DPT ★ recLSN	Dirty Page Table	Oldest update to page _x since it was last flushed
ATT ★ lastLSN	Active Transaction Table	Latest record of txn T_i (managed by transaction)
MasterRecord	Disk	LSN of latest checkpoint

Table 1: Log Sequence Number Types – A list of the different LSNs that a DBMS maintains in its internal components and data structures.

Transaction Commit **B**

When a transaction goes to commit, the DBMS first writes COMMIT record to log buffer in memory. Then the DBMS flushes all log records up to and including the transaction's COMMIT record to disk. Note that these log flushes are sequential, synchronous writes to disk. There can be multiple log records per log page. A diagram of a transaction commit is shown in Figure 2.

Once the COMMIT record is safely stored on disk, the DBMS returns an acknowledgment back to the application that the transaction has committed. At some later point, the DBMS will write a special TXN-END record to log. This indicates that the transaction is completely finished in the system and there will not be anymore log records for it. These TXN-END records are used for internal bookkeeping and do not need to be flushed immediately. ★

Transaction Abort **C**

Aborting a transaction is a special case of the ARIES undo operation applied to only one transaction.

An additional field is added to the log records called the prevLSN. This corresponds to the previous LSN for the transaction. The DBMS uses these prevLSN values to maintain a linked-list for each transaction that makes it easier to walk through the log to find its records. ★

A new type of record called the compensation log record (CLR) is also introduced. A CLR describes the

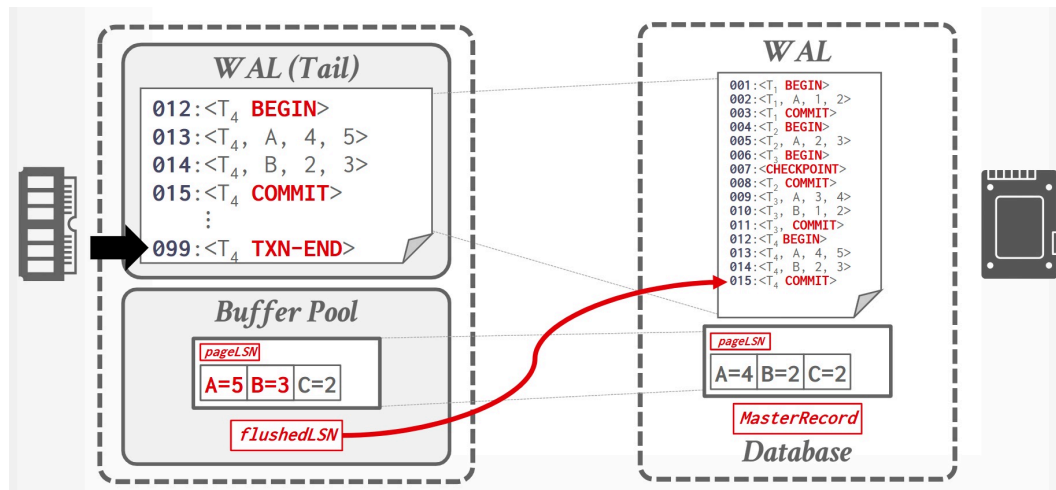


Figure 2: Transaction Commit – After the transaction commits (015), the log is flushed out and the flushedLSN is modified to point to the last log record generated. At some later point, a transaction end message is written in the log to signify that this transaction will not appear again. Then we can trim the in-memory log up to flushedLSN.

actions taken to undo the actions of a previous update record. It has all the fields of an update log record plus the undoNextLSN pointer (i.e., the next-to-be-undone LSN). The DBMS adds CLRs to the log like any other record but they never need to be undone. Moreover, the DBMS does not wait for CLRs to be flushed to disk before notifying the application that the transaction has been aborted. This approach ensures efficient transaction management, especially in scenarios involving transaction rollbacks.

To abort a transaction, the DBMS first appends a ABORT record to the log buffer in memory. It then undoes the transaction's updates in reverse order to remove their effects from the database. For each undone update, the DBMS creates CLR entry in the log and restore old value. After all of the aborted transaction's updates are reversed, the DBMS then writes a TXN-END log record. A diagram of this is shown in Figure 3.

4 Checkpointing

The DBMS periodically takes *checkpoints* where it writes the dirty pages in its buffer pool out to disk. This is used to minimize how much of the log it has to replay upon recovery.

The first two blocking checkpoint methods discussed below pause transactions during the checkpoint process. This pausing is necessary to ensure that the DBMS does not miss updates to pages during the checkpoint. Then, a better approach that allows transactions to continue to execute during the checkpoint but requires the DBMS to record additional information to determine what updates it may have missed is presented.

① Non-Fuzzy Checkpoints < CHECKPOINT >

The DBMS halts the execution of transactions and queries when it takes a checkpoint to ensure that it writes a consistent snapshot of the database to disk. The is the same approach discussed in previous lecture:

- Halt the start of any new transactions.
- Wait until all active transactions finish executing.
- Flush dirty pages to disk.

TRANSACTION ABORT – CLR EXAMPLE

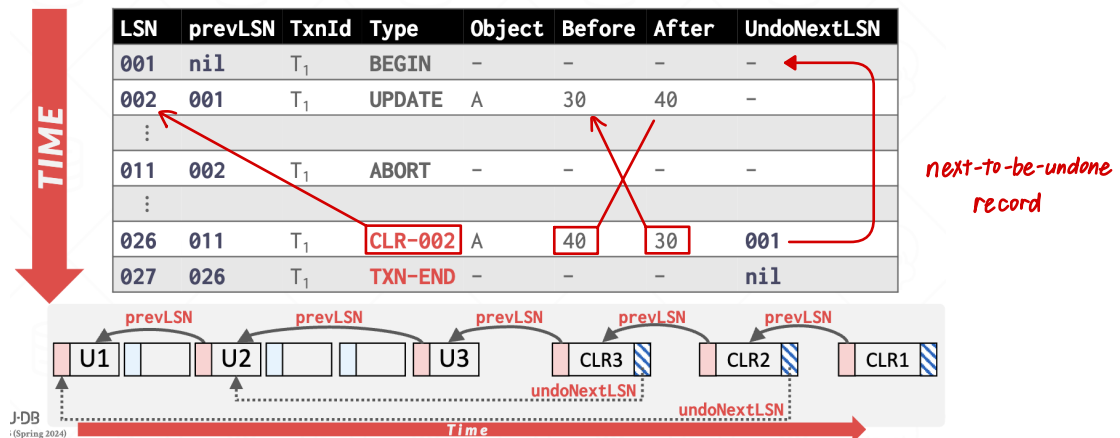


Figure 3: Transaction Abort – The DBMS maintains an LSN and prevLSN for each log record that the transaction creates. When the transaction aborts, all of the previous changes are reversed. After the log entries of the reversed changes make it to disk, the DBMS appends the TXN-END record to the log for the aborted transaction.

While this process impacts runtime performance, it significantly simplifies recovery.

② Slightly Better Blocking Checkpoints < CHECKPOINT ATT={...} DPT={...} >

Like previous checkpoint scheme except that you the DBMS does not have to wait for active transactions to finish executing. The DBMS now records the internal system state as of the beginning of the checkpoint.

- Halt the start of any new transactions.
- Pause transactions while the DBMS takes the checkpoint.

The checkpoint process requires recording the internal state at its commencement. This includes two key components: the Active Transaction Table (ATT), which tracks ongoing transactions, and the Dirty Page Table (DPT), which lists all modified pages not yet written to disk.

Active Transaction Table (ATT): The ATT represents the state of transactions that are actively running in the DBMS. A transaction's entry is removed after the DBMS completes the commit/abort process for that transaction. For each transaction entry, the ATT contains the following information:

- ✓ transactionId: Unique transaction identifier
- ✓ status: The current "mode" of the transaction (Running, Committing, Undo Candidate)
- ✓ lastLSN: Most recent LSN written by transaction

Note that the ATT contains every transaction without the TXN-END log record. This includes both transactions that are either committed or abort.

Dirty Page Table (DPT): The DPT contains information about the pages in the buffer pool that were modified by uncommitted transactions. There is one entry per dirty page containing the recLSN (i.e., the LSN of the log record that first caused the page to be dirty).

The DPT contains all pages that are dirty in the buffer pool. It doesn't matter if the changes were caused by a transaction that is running, committed, or aborted.

Overall, the ATT and DPT are vital in both checkpointing and recovery processes. During checkpointing, they capture the database's current state, with the ATT tracking active transactions and the DPT listing

unflushed modified pages. In recovery, such as with the ARIES protocol, these tables aid in restoring the database to its consistent pre-crash state.

③ Fuzzy Checkpoints <CHECKPOINT-END ATT={...} DPT={...}>

A *fuzzy checkpoint* is where the DBMS allows other transactions to continue to run. This is what ARIES uses in its protocol.

The DBMS uses additional log records to track checkpoint boundaries:

- **<CHECKPOINT-BEGIN>**: Indicates the start of the checkpoint. At this point, the DBMS takes a snapshot of the current ATT and DPT, which are referenced in the <CHECKPOINT-END> record. Transactions that start after the checkpoint initiation are not included in the ATT.
- **<CHECKPOINT-END>**: When the checkpoint has completed. It contains the ATT + DPT, captured just as the <CHECKPOINT-BEGIN> log record is written.

Upon the completion of the checkpoint, the LSN of the <CHECKPOINT-BEGIN> record is recorded in the MasterRecord.

Fuzzy checkpoints strike a balance between performance and recoverability by minimizing transaction disruption while still capturing sufficient recovery information.

5 ARIES Recovery

The ARIES protocol is comprised of three phases. Upon start-up after a crash, the DBMS will execute the following phases as shown in Figure 4:

1. **Analysis**: Read the WAL to identify dirty pages in the buffer pool and active transactions at the time of the crash. At the end of the analysis phase the ATT tells the DBMS which transactions were active at the time of the crash. The DPT tells the DBMS which dirty pages might not have made it to disk.

→ lastLSN (Analysis + Undo)
→ recLSN (Analysis + Redo)
2. **Redo**: Repeat all actions starting from an appropriate point in the log (even txns that will abort).
3. **Undo**: Reverse the actions of transactions that did not commit before the crash.

Analysis Phase

Start from last checkpoint found via the database's MasterRecord LSN.

- Range → 1. Scan log forward from the checkpoint.
- ATT → 2. If the DBMS finds a TXN-END record, remove its transaction from ATT.
3. All other records, add transaction to ATT with status **UNDO**, and on commit, change transaction status to **COMMIT**.
- DPT → 4. For UPDATE log records, if page *P* is not in the DPT, then add *P* to DPT and set *P*'s recLSN to the log record's LSN.

Redo Phase

The goal of this phase is for the DBMS to repeat history to reconstruct its state up to the moment of the crash. It will reapply all updates (even aborted transactions) and redo **CLRs**.

The DBMS scans forward from log record containing smallest recLSN in the DPT. For each update log record or CLR with a given LSN, the DBMS re-applies the update unless:

3 Cases don't need to REDO! X Affected page is not in the DPT, or

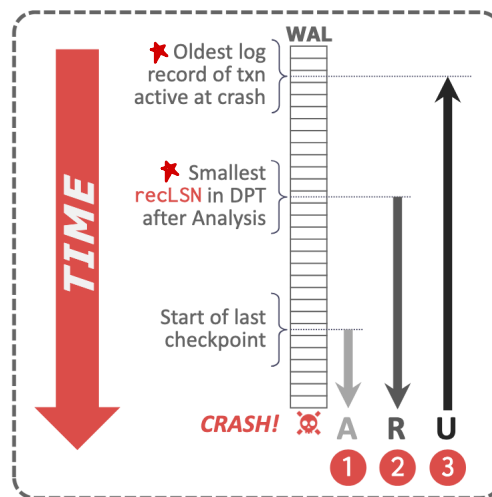


Figure 4: ARIES Recovery: The DBMS starts the recovery process by examining the log starting from the last BEGIN-CHECKPOINT found via MasterRecord. It then begins the Analysis phase by scanning forward through time to build out ATT and DPT. In the Redo phase, the algorithm jumps to the smallest recLSN, which is the oldest log record that may have modified a page not written to disk. The DBMS then applies all changes from the smallest recLSN. The Undo phase starts at the oldest log record of a transaction active at crash and reverses all changes up to that point.

- ✗ Affected page is in DPT but that log record's LSN is less than the page's recLSN, (the update was propagated to disk), or
- ✗ Affected pageLSN (on disk) $\geq$ LSN. (Note, we must fetch the page from the disk to read the page value.)

To redo an action, the DBMS re-applies the change in the log record and then sets the affected page's pageLSN to that log record's LSN. Also, there is no additional logging or forced flushes.

At the end of the redo phase, write TXN-END log records for all transactions with status COMMIT and remove them from the ATT.

Undo Phase

In the last phase, the DBMS reverses all transactions that were active at the time of crash. These are all transactions with UNDO status in the ATT after the Analysis phase.

The DBMS processes transactions in reverse LSN order using the lastLSN to speed up traversal. At each step, pick the largest lastLSN across all transactions in the ATT. As it reverses the updates of a transaction, the DBMS writes a CLR entry to the log for each modification.

Once the last transaction has been successfully aborted, the DBMS flushes out the log and then is ready to start processing new transactions.