



CHECKLIST

VMware VM Performance Tuning for Latency Reduction Checklist

A practical, vendor-neutral checklist for reducing avoidable latency in VMware virtual machines by validating CPU, memory, storage, network, and guest settings before production use.

VMware VM Performance Tuning for Latency Reduction Checklist

Use this checklist to reduce avoidable VM latency without introducing unnecessary operational risk. It is designed for latency-sensitive workloads in VMware environments and assumes you will validate each change under representative load.

1) Define the problem before tuning

- ☐ Confirm the workload is actually experiencing latency, not just low throughput.
- ☐ Identify the symptom type: slow response, tail latency spikes, jitter, intermittent pauses, or burst-related delays.
- ☐ Note when the issue occurs:
 - ☐ Constantly
 - ☐ During peak load
 - ☐ During backup or maintenance windows
 - ☐ Only on specific hosts or datastores
- ☐ Record the business impact:



- ☐ Application response delays
- ☐ Transaction slowdown
- ☐ Batch window overruns
- ☐ Time-sensitive tool delays
- ☐ East-west service chatter issues

2) Capture a baseline

- ☐ Measure performance during representative workload conditions.
- ☐ Record baseline values for:
 - ☐ Application response time
 - ☐ Tail latency / p95 / p99 latency
 - ☐ Guest CPU ready or equivalent scheduling indicator
 - ☐ Host CPU saturation
 - ☐ Guest memory pressure
 - ☐ Datastore latency and queueing
 - ☐ Network retransmits, drops, or buffering signs
 - ☐ Document host, datastore, network, and VM placement details.
 - ☐ Confirm the baseline was captured before any tuning changes.

3) Identify the dominant wait class

CPU scheduling

- ☐ Check whether the VM is oversubscribed on vCPU.
- ☐ Look for high ready time or scheduler delay.
- ☐ Verify whether the workload can actually use all assigned vCPUs.



- ☐ Check for noisy neighbor activity on the same host.

Memory contention

- ☐ Check for ballooning, swapping, or memory pressure.
- ☐ Verify the VM is not crossing NUMA boundaries unnecessarily.
- ☐ Confirm the VM is not oversized for its real working set.
- ☐ Check host memory contention and placement.

Storage queueing

- ☐ Review datastore latency during peak I/O.
- ☐ Check for queue buildup or throttling.
- ☐ Confirm the virtual controller is appropriate for the workload.
- ☐ Verify backend storage is not the bottleneck.
- ☐ Review thin provisioning, multipath balance, and path congestion.

Network delay

- ☐ Check for packet drops, retransmits, or congestion.
- ☐ Verify the virtual NIC type is appropriate.
- ☐ Review offload settings and driver version.
- ☐ Confirm physical uplinks and fabric paths are healthy.

Guest-level factors

- ☐ Review guest power management settings.
- ☐ Verify virtual drivers are current.



- ☐ Confirm the application is not introducing avoidable synchronous I/O or poor threading behavior.

4) Right-size the VM

- ☐ Reduce vCPU count if the workload is not using them efficiently.
- ☐ Avoid adding vCPUs as a default latency fix.
- ☐ Match memory allocation to the actual workload working set.
- ☐ Avoid excess memory if it creates placement or NUMA inefficiency.
- ☐ Re-check whether VM sizing changes improve tail latency, not just averages.

5) Optimize CPU-related settings carefully

- ☐ Confirm host CPU contention is not the primary cause.
- ☐ Check if the VM is waiting because too many vCPUs must be scheduled together.
- ☐ Consider smaller, more schedulable VM shapes where appropriate.
- ☐ Avoid aggressive CPU reservations unless the latency requirement justifies them.
- ☐ Re-test after any change to vCPU count or placement.

6) Validate memory locality and pressure

- ☐ Keep the VM within a sensible NUMA boundary when possible.
- ☐ Check whether memory access patterns worsen during peak load.
- ☐ Look for signs of ballooning or swapping in the guest and on the host.
- ☐ Confirm that memory contention is not being caused by overcommit pressure.
- ☐ Avoid assuming that more RAM automatically means lower latency.

7) Check storage path efficiency



- ☐ Verify the datastore latency is acceptable under the actual workload.
- ☐ Confirm the virtual disk controller is suitable for the application.
- ☐ Review queue depth and path balance.
- ☐ Check whether synchronous writes are causing burst delays.
- ☐ Validate the storage backend during peak activity, not just at idle.
- ☐ Review multipath settings and ensure paths are balanced.
- ☐ Confirm thin provisioning behavior is not creating unexpected overhead.

8) Check network path and virtual NIC settings

- ☐ Confirm the VM uses an appropriate virtual NIC model.
- ☐ Review driver version and compatibility.
- ☐ Verify offload settings are aligned with the workload.
- ☐ Check uplink utilization and congestion.
- ☐ Look for intermittent latency spikes rather than only average delay.
- ☐ Validate east-west traffic patterns if the application is chatty or clustered.

9) Review guest OS tuning

- ☐ Set guest power management to favor consistent performance where appropriate.
- ☐ Update virtual hardware drivers.
- ☐ Verify interrupt handling and I/O behavior are not degraded by outdated drivers.
- ☐ Check for background services or agents that may add jitter.
- ☐ Confirm time synchronization and system behavior are stable.

10) Make one change at a time

- ☐ Change only one tuning variable at a time.



- ☐ Record the exact setting changed and why.
- ☐ Keep a rollback plan for each change.
- ☐ Avoid broad optimization passes that make root cause analysis harder.
- ☐ Re-test after every change under the same workload shape.

11) Re-test using the same workload conditions

- ☐ Run the same representative workload used in the baseline.
- ☐ Compare tail latency, not just mean values.
- ☐ Check whether the change reduced variance and spikes.
- ☐ Verify there was no new bottleneck created elsewhere.
- ☐ Confirm the result is repeatable across multiple test runs.

12) Decide whether to keep, refine, or roll back

- ☐ Keep the change if it clearly improves the targeted latency symptom.
- ☐ Refine the change if it helped but introduced a new constraint.
- ☐ Roll back the change if the effect is unclear, unstable, or harmful.
- ☐ Document the final decision and measured outcome.
- ☐ Update the production standard or runbook if the change is successful.

13) Production-safety checks before broader rollout

- ☐ Confirm the tuning change is safe for failover and recovery operations.
- ☐ Validate the change will not reduce operational flexibility beyond acceptable limits.
- ☐ Check whether the tuning affects consolidation efficiency.
- ☐ Ensure patching, maintenance, and host operations remain manageable.
- ☐ Review whether the improvement justifies the added complexity.



14) Common signs the tuning is likely to help

- ☐ The VM is oversized and harder to schedule than necessary.
- ☐ The workload is bursty and sensitive to queueing.
- ☐ Host contention is visible during the issue window.
- ☐ Storage latency rises during peak writes.
- ☐ Network delay appears as intermittent jitter or drops.
- ☐ Guest drivers or power settings are outdated or inconsistent.

15) Common signs tuning may not help much

- ☐ The application itself is the main source of delay.
- ☐ The issue happens only under abnormal load spikes.
- ☐ The bottleneck is outside the VM layer.
- ☐ The environment is already lightly loaded and stable.
- ☐ Measurements do not show a clear dominant wait class.

16) Final acceptance checklist

Before moving a tuned VM to production or expanding the change to other workloads, confirm:

- ☐ Baseline data was captured and saved.
- ☐ The root cause or dominant wait class was identified.
- ☐ The change was minimal and targeted.
- ☐ Tail latency improved under representative load.
- ☐ There was no unacceptable impact on throughput, resilience, or manageability.
- ☐ The rollback path remains available.
- ☐ The tuning decision is documented for future reference.



Quick rule of thumb

If a VM is latency-sensitive, start by reducing avoidable contention, aligning the VM to the real workload, and validating storage and network paths. Tune selectively, measure carefully, and keep the environment production-safe.