
Algorithm 1 KVCache-centric Scheduling Algorithm

Input: prefill instance pool P , decoding instance pool D , request R , cache block size B .
Output: the prefill and decoding instances (p, d) to process R .

- 1: $block_keys \leftarrow \text{PrefixHash}(R.prompt_tokens, B)$
- 2: $TTFT \leftarrow \text{inf}$
- 3: $p \leftarrow \emptyset$
- 4: $best_prefix_len, best_matched_instance \leftarrow \text{FindBestPrefixMatch}(P, block_keys)$
- 5: **for** $instance \in P$ **do**
- 6: $prefix_len \leftarrow instance.prefix_len$
- 7: $T_{queue} \leftarrow \text{EstimatePrefillQueueTime}(instance)$
- 8: **if** $\frac{best_prefix_len}{prefix_len} < kvcache_balancing_threshold$ **then** ▷ Cache-aware prefill scheduling
- 9: $T_{prefill} \leftarrow \text{EstimatePrefillExecutionTime}(\text{len}(R.prompt_tokens), prefix_len)$
- 10: **if** $TTFT > T_{queue} + T_{prefill}$ **then**
- 11: $TTFT \leftarrow T_{queue} + T_{prefill}$
- 12: $p \leftarrow instance$
- 13: **end if**
- 14: **else** ▷ Cache-aware and -balancing prefill scheduling
- 15: $transfer_len \leftarrow best_prefix_len - prefix_len$
- 16: $T_{transfer} \leftarrow \text{EstimateKVCacheTransferTime}(instance, best_matched_instance, transfer_len)$
- 17: $T_{prefill} \leftarrow \text{EstimatePrefillExecutionTime}(\text{len}(R.prompt_tokens), best_prefix_len)$
- 18: **if** $TTFT > T_{transfer} + T_{queue} + T_{prefill}$ **then**
- 19: $TTFT \leftarrow T_{transfer} + T_{queue} + T_{prefill}$
- 20: $p \leftarrow instance$
- 21: **end if**
- 22: **end if**
- 23: **end for**
- 24: $d, TBT \leftarrow \text{SelectDecodingInstance}(D)$ ▷ Load-balancing decoding scheduling
- 25: **if** $TTFT > TTFT_SLO$ **or** $TBT > TBT_SLO$ **then**
- 26: **reject** R ; **return**
- 27: **end if**
- 28: **if** $\frac{best_prefix_len}{p.prefix_len} > kvcache_balancing_threshold$ **then**
- 29: $\text{TransferKVCache}(best_matched_instance, p)$ ▷ KVCache hot-spot migration
- 30: **end if**
- 31: **return** (p, d)

6 KVCache-centric Scheduling

In this section, we mainly discuss how Conductor schedules the requests and KVCache blocks under normal conditions, leaving the discussion on overload scenarios for the next section.

6.1 Prefill Global Scheduling

Previous research on LLM serving typically uses a load-balancing strategy that evaluates the load on each instance based on the number of assigned requests. In Mooncake, however, the selection of prefill instances considers additional factors—not just load but also the prefix cache hit length and the distribution of reusable KVCache blocks. While there is a preference to route requests to prefill instances with longer prefix cache lengths to reduce computation costs, it may be beneficial to schedule them to other nodes to ensure overall system balance and meet TTFT SLOs. To address these complexities, we propose a cache-aware global scheduling algorithm that accounts for both the prefill time due to the prefix cache and the queuing time associated with the load on the instance.

Algorithm 1 details the mechanism for our cache-aware prefill scheduling. For every new request, its input tokens are divided into several blocks, and a hash key is computed for each block. This involves generating a hash key of tokens in a block concatenated with the hash key of the previous block (if available). The request’s block keys are then compared one by one against each prefill instance’s cache keys to identify the prefix match length ($prefix_len$). Similar reuse logic is already implemented in vLLM, but the open-source version of vLLM only supports local KVCache caching.

With this matching information, Conductor estimates the corresponding execution time based on the request length and $prefix_len$ (which varies by instance). It then adds the estimated waiting time for that request to get the TTFT on that instance. Finally, Conductor assigns the request to the instance