

高雄醫學大學 114 學年度學士後醫學系招生考試試題參考答案疑義釋疑公告

科目	題號	釋疑答覆	釋疑結果
計算機概論與程式設計	9	Instruction-level parallelism is achieved through pipelining, which “overlaps the execution of multiple instructions” so that several are resident in different pipeline stages at the same time, thereby fulfilling the requirement that multiple instructions be executed simultaneously; in contrast, conventional hardware multithreading (both coarse- and fine-grained) issues only one instruction per cycle and simply switches among thread contexts to mask stalls, so no two instructions actually proceed concurrently in the pipeline, meaning it improves resource utilization rather than delivering true simultaneous execution. Hence, option (C) Pipelining is the correct choice because it achieves true instruction-level parallelism by overlapping successive stages of different instructions, allowing several instructions to reside in the processor concurrently. By contrast, option (A) Multithreading issues only one instruction per cycle and merely alternates among thread contexts to hide latency; it does not create simultaneous execution of multiple instructions within a single core.	本題為單選題，在只有一個正確答案的條件下，公告答案(C)比答案(A)更佳。故維持原答案。

	10	Based on standard operating system textbooks, the scheduling algorithm most commonly cited as potentially leading to starvation among the given options is Shortest Job First (SJF), particularly its preemptive variant Shortest Remaining Time First (SRTF). This is a direct consequence of the algorithm's core principle: if a continuous stream of new, short processes arrives, they can perpetually preempt longer processes, indefinitely postponing their execution and causing starvation. In contrast, Multilevel Queue (MLQ) is more accurately described as a scheduling framework or structure comprising multiple distinct queues. A key characteristic of MLQ is its configurability; the scheduling algorithms used within each queue (e.g., Round Robin, FCFS) and, critically, the policy for scheduling between these queues are determined by the system designer. Importantly, because MLQ is a flexible framework, it allows for the implementation of mechanisms specifically designed to prevent starvation, such as aging (gradually increasing the priority of processes waiting in lower queues) or allocating time slices to the queues themselves (time-slicing between queues), guaranteeing that each queue receives a certain percentage of CPU time. Hence, option (B) Shortest Job First (SJF) is the correct choice because its principle (particularly in the preemptive SRTF variant) of prioritizing shorter jobs means that longer jobs can be indefinitely postponed by the continuous arrival of new, short processes, leading to starvation. By contrast, option (D) Multilevel Queue (MLQ) is a configurable framework that can be explicitly mitigated through mechanisms like aging or time-slicing between queues, making it a less inherent or unavoidable cause of starvation compared to SJF.	本題為單選題，在只有一個正確答案的條件下，公告答案(B)比答案(D)更佳。故維持原答案。
--	----	---	---