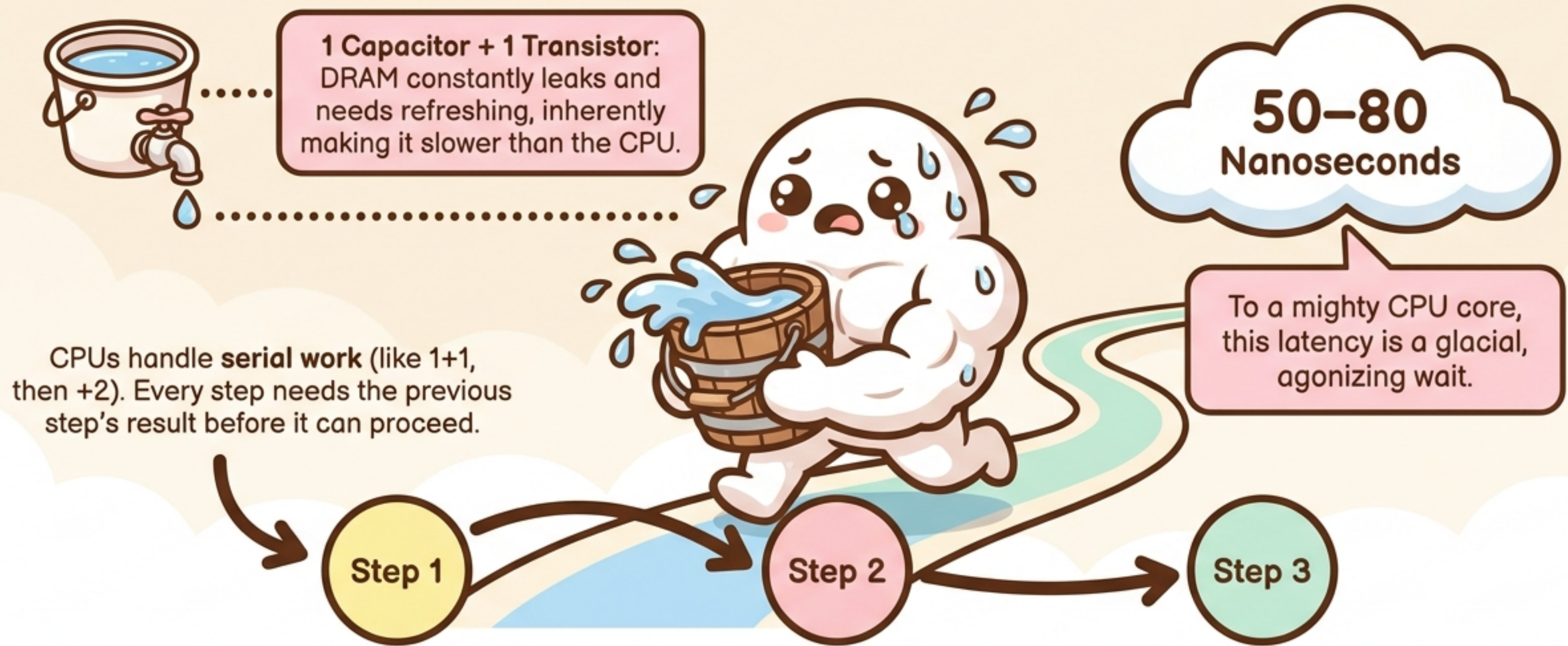


Main memory and video memory parted ways long ago—CPUs stayed behind to carry water by the bucket.



1 Mighty Core + 1 Narrow Lane (64-bit bus width) = The CPU demands its data right now, but can only fetch it one bucket at a time.

GPUs abandon individual speed to command a lockstep legion of thousands.

If a teacher wastes 1 minute of class, 45 students wait. But they wait together. Total time wasted? Still 1 minute.

CPU Style vs. GPU Style

CPU Style

Few Cores
Serial Tasks
Hates Waiting
(Low Latency required)

GPU Style

Thousands of Cores
Parallel Tasks
Doesn't Care About Waiting
(High Latency acceptable)

If 10,000 cores fetch data at once, a 200-nanosecond latency is perfectly fine.

GPUs don't need fast individuals; they need massive coordinated movement.



Feeding thousands of cores simultaneously requires building massive highway systems.

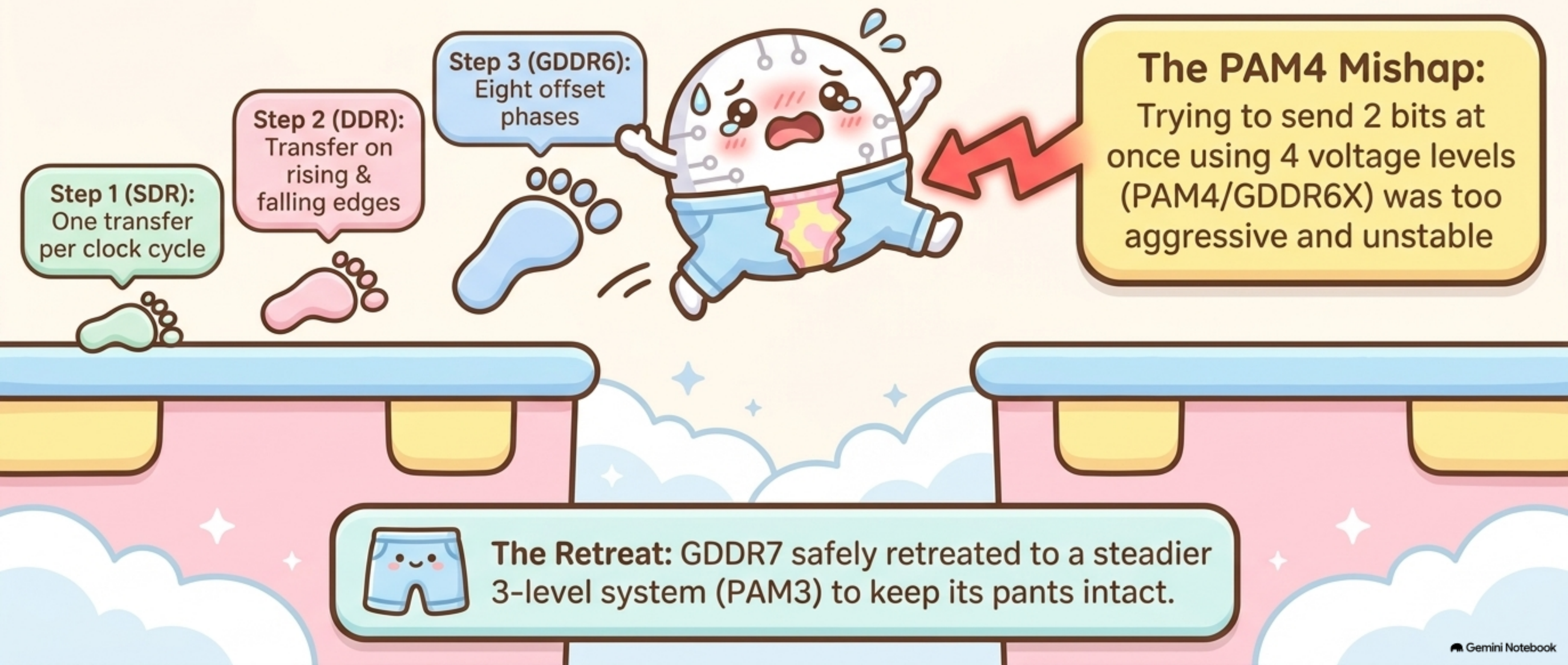
$$\text{Bandwidth} = \text{Bus Width} \times \text{Transfer Rate}$$

Bus Width (Lanes):
Standard main memory has been stuck at 64 lanes (bits) forever. High-end GPUs build 384 or even a flat 512 lanes.

Waa!

Transfer Rate (Trucks):
Since trucks can only go so fast, VRAM gets clever by cramming more trucks into more trucks into the tight gaps between clock cycles.

Pushing transfer rates too aggressively leads to embarrassing physical limits.



When highways run out of room, AI demands we bolt a skyscraper directly onto the GPU.

1

The Breaking Point: AI chips require a 1,024-bit bus width per stack. This means nearly 4,000 wires—too many for standard fiberglass motherboards to handle.

2

The Interposer: A mirror-smooth slab of pure silicon acting as the new foundation to handle the extreme wiring density.

3

TSV (Through-Silicon Vias): Microscopic copper pillars drilled directly through vertically stacked memory dies.



The 2 Million RMB Synthesis: AI servers cost millions because you aren't just buying the GPU—you are paying for this incredibly complex memory roof. As a side effect, this takes up standard for memory factory lines, quietly raising the price of everyone's regular RAM.