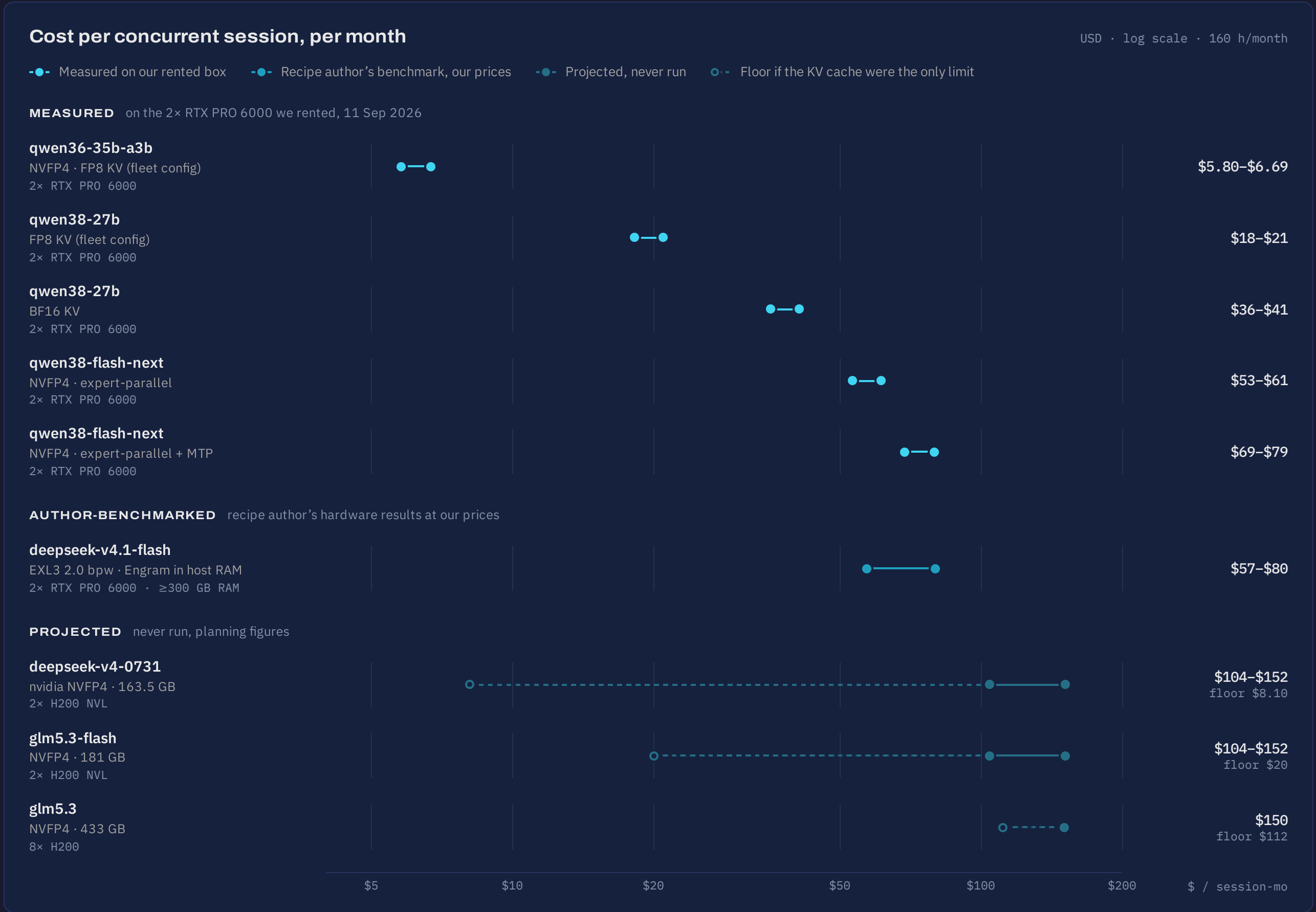


# Flotilla Session Rate Card

$$\text{box } \$/\text{h} \times 160 \text{ h} \div \text{concurrent 262K sessions per box} = \$/\text{session-month}$$



## Inputs behind every row

| Evidence             | Model and configuration                                         | Box                              | \$/h          | Box, 160 h    | Concurrent sessions               | \$/session-month            | Engine                  | Basis                                                                                            |
|----------------------|-----------------------------------------------------------------|----------------------------------|---------------|---------------|-----------------------------------|-----------------------------|-------------------------|--------------------------------------------------------------------------------------------------|
| ● Author-benchmarked | <b>deepseek-v4.1-flash</b><br>EXL3 2.0 bpw · Engram in host RAM | 2× RTX PRO 6000<br>· ≥300 GB RAM | \$2.86–\$4.00 | \$457–\$640   | 8<br>cache alone ≈ 14.5           | \$57–\$80                   | vLLM day-0 image + EXL3 | The recipe runs at most 8 sequences; its author benchmarked up to 4 streams                      |
| ● Projected          | <b>deepseek-v4-0731</b><br>nvidia NVFP4 · 163.5 GB              | 2× H200 NVL                      | \$5.21–\$7.58 | \$834–\$1,212 | 8 planned<br>cache alone ≈ 41–103 | \$104–\$152<br>floor \$8.10 | vLLM or SGLang          | Assumes 8 per 2 GPUs, borrowed from V4.1-Flash; cache size per token inferred from the V4.1 card |
| ● Projected          | <b>glm5.3-flash</b><br>NVFP4 · 181 GB                           | 2× H200 NVL                      | \$5.21–\$7.58 | \$834–\$1,212 | 8 planned<br>cache alone ≈ 42     | \$104–\$152<br>floor \$20   | vLLM per-model image    | Assumes 8 per 2 GPUs; only 11 of its 45 layers hold KV cache                                     |
| ● Projected          | <b>glm5.3</b><br>NVFP4 · 433 GB                                 | 8× H200                          | \$30.04       | \$4,806       | 32 planned<br>cache alone ≈ 43    | \$150<br>floor \$112        | vLLM or SGLang          | Assumes 8 per 2 GPUs; cache size per token derived from its config                               |

### Runtime behind every row

Measured on 2× RTX PRO 6000 Max-Q with vLLM nightly e7edf17c at TP=2 and a 262K window, 11 Sep 2026. Decode is aggregate throughput across concurrent streams. DeepSeek-V4.1-Flash's figures are its recipe author's, on the same card at 300 W with CUDA graphs and DSpark speculative decoding. Projected rows have never run: context is the model's configured maximum and the KV pool is our projection.

| Evidence             | Model and configuration                                         | Decode tok/s                                   | Cold prefill tok/s                         | Max context                      | KV pool                                      | Concurrency                                                                                                      |
|----------------------|-----------------------------------------------------------------|------------------------------------------------|--------------------------------------------|----------------------------------|----------------------------------------------|------------------------------------------------------------------------------------------------------------------|
| ● Measured           | <b>qwen36-35b-a3b</b><br>NVFP4 · FP8 KV (fleet config)          | 197 → 6,396<br>1 → 128 streams                 | 30.1K → 15.9K<br>8K → 132K prompt          | 262K<br>served window            | 14.2M tokens<br>measured ledger              | 256 max seqs; 64 concurrent 27K agent sessions at ≤0.37 s p95 TTFT                                               |
| ● Measured           | <b>qwen38-27b</b><br>FP8 KV (fleet config)                      | 92.6 → 2,908<br>1 → 64 streams                 | 10.0K → 5.9K<br>8K → 132K prompt           | 262K<br>served window            | 4.53M tokens<br>measured ledger              | 256 max seqs                                                                                                     |
| ● Measured           | <b>qwen38-27b</b><br>BF16 KV                                    | 96.7 → 2,873<br>1 → 64 streams                 | 10.1K → 6.5K<br>8K → 132K prompt           | 262K<br>served window            | 2.32M tokens<br>measured ledger              | 256 max seqs                                                                                                     |
| ● Measured           | <b>qwen38-flash-next</b><br>NVFP4 · expert-parallel             | 97.5 → 1,846<br>1 → 64 streams                 | 13.7K → 13.8K<br>8K → 132K, flat           | 262K<br>served window            | 1.55M tokens<br>measured ledger              | 256 max seqs; its attention forces a BF16 cache                                                                  |
| ● Measured           | <b>qwen38-flash-next</b><br>NVFP4 · expert-parallel + MTP       | 125.9 → 2,074<br>1 → 64 streams                | 12.6K → 13.3K<br>8K → 132K prompt          | 262K<br>served window            | 1.20M tokens<br>measured ledger              | 256 max seqs; MTP adds 12–29% decode at every level up to 64                                                     |
| ● Author-benchmarked | <b>deepseek-v4.1-flash</b><br>EXL3 2.0 bpw · Engram in host RAM | 113.3 → 277.2<br>1 → 4 streams @2K; 116.6 @46K | 4,044 → 8,085<br>46K prompt, 1 → 4 streams | 512K<br>configured; model max 1M | ≈3.8M tokens<br>8 GiB cache                  | 8 max seqs; benchmarked to 4 streams; agentic run held 240–254 tok/s over 3–4 streams with 87% prefix-cache hits |
| ● Projected          | <b>deepseek-v4-0731</b><br>nvidia NVFP4 · 163.5 GB              | not measured                                   | not measured                               | 1M<br>vendor recipe uses 393K    | ≈10.6–27.0M tokens<br>projected, 2× H200 NVL | 8 planned; vendor tested TP=8 on B200                                                                            |
| ● Projected          | <b>glm5.3-flash</b><br>NVFP4 · 181 GB                           | not measured                                   | not measured                               | 1M<br>configured maximum         | ≈10.9M tokens<br>projected, 2× H200 NVL      | 8 planned; only 11 of 45 layers hold KV cache                                                                    |
| ● Projected          | <b>glm5.3</b><br>NVFP4 · 433 GB                                 | not measured                                   | not measured                               | 1M<br>configured maximum         | ≈11.3M tokens<br>projected, 8× H200          | 32 planned                                                                                                       |

### Reading it

#### Two different limits fill a box.

The qwen rows run out of KV cache first, so their session counts come straight from the KV ledger we measured. DeepSeek-V4.1-Flash runs out of compute first: its recipe runs at most 8 sequences, though its cache could hold about 14.5. Its author benchmarked only up to 4 streams, and at 4 a session would cost \$114–\$160. For contrast, qwen36-35b-a3b held 64 concurrent 27K-token agent sessions at 0.37 s p95 time-to-first-token or better, which is \$4.92–\$5.68 each.

#### Projected rows are planning figures, not quotes.

deepseek-v4-0731, glm5.3-flash and glm5.3 have never run for us. They assume 8 concurrent sessions per 2 GPUs, borrowed from V4.1-Flash's recipe. The dashed line runs down to what they would cost if the KV cache were the only limit, which is where the lower figures quoted earlier came from. Quote the solid range and treat the dashed end as a best case.

#### Prices moved all day, and some costs are left out.

2× RTX PRO 6000 went for \$1.97–\$2.27/h. A V4.1-Flash host with at least 300 GB of RAM and CUDA 13 ran \$2.86–\$4.00/h, and the cheapest of those fails our 0.95 reliability floor. 2× H200 NVL ran \$5.21–\$7.58/h, and 8× H200 \$30.04/h. Bandwidth, disk and cold-start time are not included.

