

К статье «Форматы артефактов моделей и решение open vs closed»

Шаг 1 — Фильтр privacy и compliance Шаг 4 — Налог на инженерную ёмкость

Есть ли в обработке регулируемые данные (152-ФЗ, GDPR, HIPAA, PCI, etc.)? DevOps

- ☐ Подтверждён ли регион API, где вы планируете запускать SRE? ☐ Подтверждён ли регион API, где вы планируете запускать SRE?
- ☐ Просмотрел ли юрист обязательства Article 50 EU AI Act с момента подписания 2 августа 2022? ☐ Просмотрел ли юрист обязательства Article 50 EU AI Act с момента подписания 2 августа 2022?
- ☐ При self-host: есть ли документированная политика хранения модели и сканирования? ☐ При self-host: есть ли документированная политика хранения модели и сканирования?

Шаг 2 — Фильтр carability gap

Задача требует frontier reasoning (уровень GPQA) или это production (суммаризация, ASB, CV)?

- ☐ Прогнан ли хотя бы один open-weight на образце ВАШИХ данных в бенчмарке?
☐ При closed: задокументирован ли план fallback, если вендор удалит модель из API?
☐ Зафиксирована ли лицензия (Apache 2.0 vs custom) для каждого open-weight в стеке?

Шаг 3 — Crossover по стоимости

Прогноз tokens/day на старте ; через 12 месяцев

- ☐ Closed API: стоимость за миллион токенов (input + output, blended) \$ _____, per-request телеметрия по стоимости (closed API)
☐ GPU-час в вашем регионе (L4 / A100 / H100) \$ _____, tokens, open GPU-seconds) включена?
☐ Ожидаемый batch size и throughput на GPU _____, Квартальный пересчёт стоимости в _____ календаре?
☐ Точка пересечения (closed = open total) рассчитана и зафиксирована? _____, если open-weights качество _____, план отката, если open-weights качество _____?

Шаг 5 — Гибридный routing

Описано правило routing для 'лёгких 80%'

- ☐ (open) и 'сложных 20%' (closed)?
- ☐ request телеметрия по стоимости (closed output, blended tokens, open GPU-seconds) включена?
- ☐ 'Квартальный пересчёт стоимости в календаре?'
- ☐ на и зафиксирована?
- ☐ план отката, если open-weights качество просядет после bump-a модели?

Чек-лист по формату артефакта

Мобильное приложение: модель

- ☐ сконвертирована в Core ML (iOS) или LiteRT / ONNX (Android)?
Браузер: ONNX (WebGPU) или llama.cpp WASM (GGUF)? Размер бандла замерен?
- ☐ Edge-сервер: ONNX Runtime + TensorRT EP или vLLM Safetensors? Throughput замерен?
- ☐ Cloud LLM serving: Safetensors через vLLM или скомпилированный TensorRT-LLM engine?
- ☐ Если только pickle: отказались или просканировали picklescan и положили в карантин?
- ☐ Каждый выбор квантизации (Q4_K_M, Q5_K_M, Q8_0) подкреплён бенчмарком качества?