

1 · Нужно ли вообще дообучать? (пробуйте по порядку)

- ☐ Сначала промпт: внятная инструкция + 3 хороших примера в запросе. Бесплатно, полдня. Если работает — СТОП.
- ☐ Затем поиск (RAG): разрыв — недостающие факты, что часто меняются (графики, политики, недавние события). Обновление за секунды, без обучения.
- ☐ Fine-tuning только ради ПОВЕДЕНИЯ: специализированный словарь, фиксированный формат вывода, доменное визуальное суждение или очень высокий объем вызовов.
- ☐ Сильнейшие системы делают все три. Поиск + дообучение вместе бьют каждый по отдельности (>11 пунктов в тестах 2026).

2 · Метод + шпаргалка по памяти GPU (модель 7B)

Полное (16-бит)	: ~120 ГБ	самое гибкое, самое дорогое, ВЫСОКИЙ риск забывания
Полное (8-бит)	: ~60 ГБ	всё ещё территория мульти-GPU
LoRA (16-бит)	: ~16 ГБ	ПО УМОЛЧАНИЮ — заморозка + малая надстройка, низкое забывание
QLoRA (4-бит)	: ~6 ГБ	дешевле всего; на 4-бит нельзя обучать vision tower

3 · Датасет — реальная стоимость, а не прогон обучения

- ☐ Несколько сотен примеров учат формату или скромному словарю; новое визуальное суждение — нижние тысячи.
- ☐ Качество бьёт количество: 1 000 чистых, разнообразных, верно размеченных клипов лучше 10 000 небрежных.
- ☐ Следите за лазейками: если каждый клип «нарушитель» снят ночью, модель выучит «ночь», а не «нарушитель».
- ☐ Закладывайте людей: ~1 000 клипов при ~20/час + ревью + тестовая выборка = ~75 человеко-часов. Прогон GPU — часы и \$10-30.

4 · Измерьте до выпуска (проверка на забывание)

- ☐ Разделите данные: отложите 10-20% как тестовую выборку, которую модель НИКОГДА не видит при обучении.
- ☐ Оцените ОРИГИНАЛЬНУЮ модель на нескольких общих задачах ДО дообучения; запишите цифры.
- ☐ После дообучения прогоните те же общие задачи. Домен вверх + общее рухнуло = вы переобучили.
- ☐ Чините переобучение: меньше эпох, LoRA вместо полного дообучения или подмешайте общие примеры обратно.