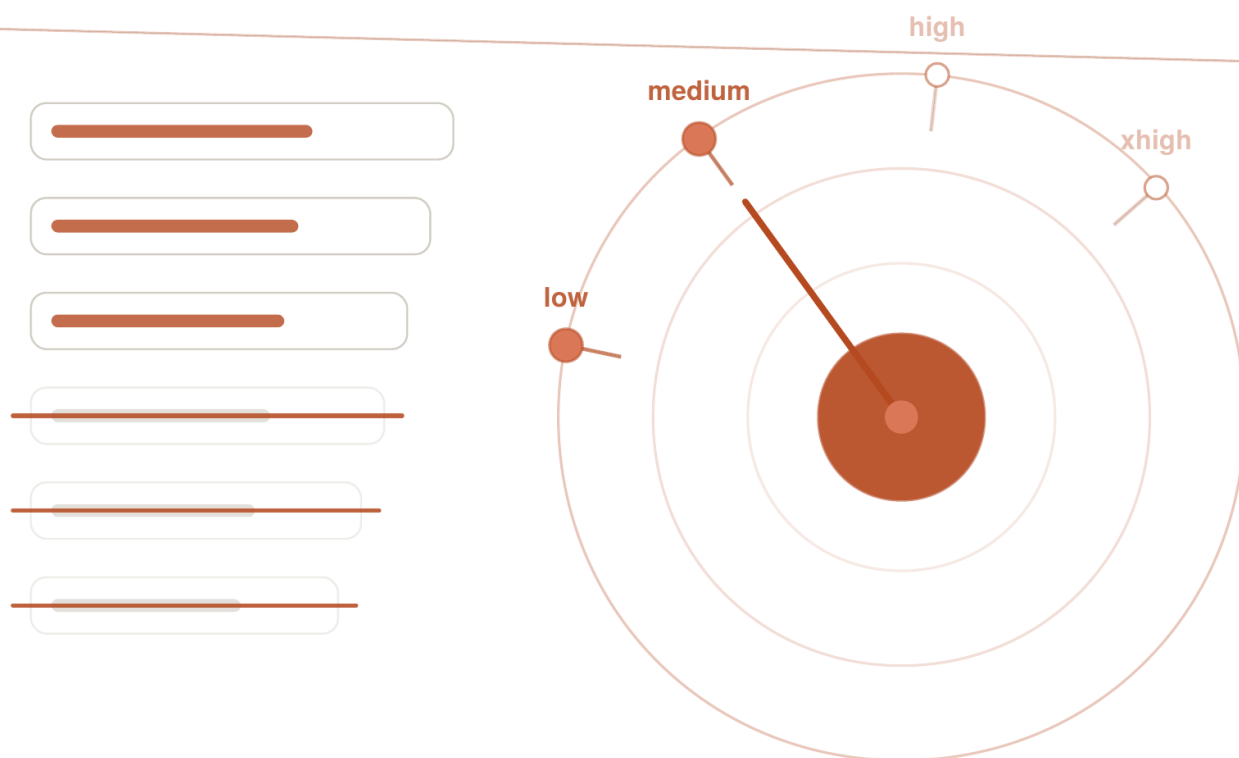


ПОД КАПОТОМ

Промптинг Claude Opus 5

Гайд Anthropic на русском. Что изменилось в поведении модели и какие строчки пора выкинуть из системного промпта.



Under the Hood · Илья Утов

@gorilla_under_hood

Содержание

Как читать этот перевод	3
-------------------------	---

ЧАСТЬ I · ЧТО ИЗМЕНИЛОСЬ

1. О чём этот гайд	5
2. Что улучшилось	6

ЧАСТЬ II · МЕНЬШЕ СЛОВ

3. Длина ответа и многословность	9
4. Отчёты о ходе работы	10
5. Объём письменных артефактов	11

ЧАСТЬ III · ГРАНИЦЫ И ДЕЛЕГИРОВАНИЕ

6. Границы задачи и лишняя перепроверка	13
7. Контроль запуска субагентов	14
8. Самокоррекция	14

ЧАСТЬ IV · ТОНКИЕ МЕСТА

9. Работа с выключенным мышлением	17
10. Разбор от канала	18

Источники	20
-----------	----

Как читать этот перевод

Anthropic выпустила отдельную страницу документации про то, как разговаривать с Claude Opus 5. Это редкий жанр: не маркетинг про бенчмарки, а список конкретных поведений модели, которые чаще всего приходится подкручивать промптом.

Ниже полный перевод всех восьми разделов. Ни один тезис не выброшен. Два отступления от оригинала сделаны сознательно и вот они: подзаголовки внутри разделов расставлены мной для навигации по А4, в оригинале это жирные врезки внутри абзацев; перекрёстные ссылки на смежные страницы документации собраны в раздел «Источники» в конце. Блоки промптов, которые Anthropic даёт как готовые к копированию, переведены целиком и оформлены тёмными плашками.

Что моё, а что Anthropic. Основной текст это перевод. Мои комментарии вынесены во врезки с пометкой «Под капотом» и в последний раздел «Разбор от канала». Смешивать источник и комментарий нечестно, поэтому они разделены визуально.

Главное в одном абзаце. Три раздела из восьми просят не добавлять инструкцию, а удалить её. Это «Границы задачи и лишняя перепроверка» (уберите «финальный шаг верификации» и «используй субагента для проверки»), «Самокоррекция» (уберите «перепроверь ответ») и «Работа с выключенным мышлением» (уберите правило, запрещающее модели рассуждать). На прошлых моделях такие пункты поднимали качество. На Opus 5 они складываются с собственным поведением модели, жгут токены и качества не добавляют. Если вы переносите системный промпт с Opus 4.8 как есть, он будет работать, но сначала пройдитесь по нему с ножницами.

ЧАСТЬ · I

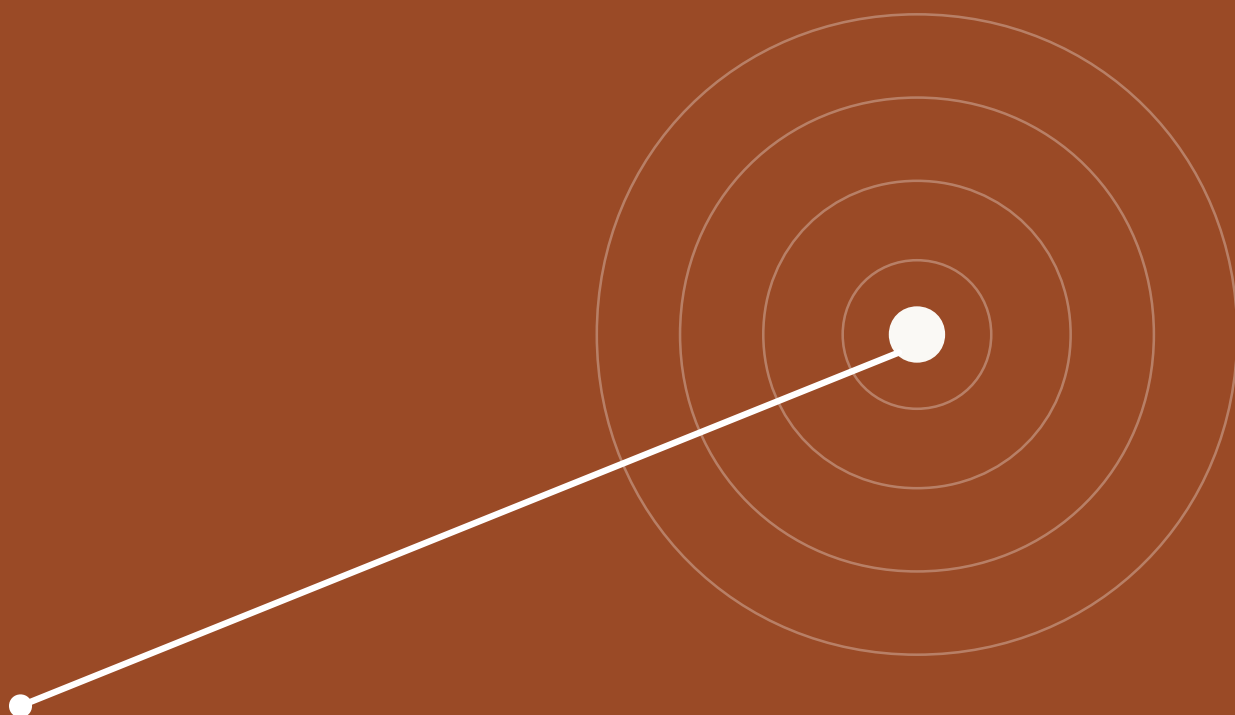


Что изменилось

Область применения гайда и то, что стало лучше по сравнению с Orus 4.8.

1 О чём этот гайд

2 Что улучшилось



О чём этот гайд

Здесь собраны приёмы промптинга, специфичные именно для Claude Opus 5. Про возможности модели и изменения в API есть отдельная страница «Что нового в Claude Opus 5». Про техники, которые работают на всех актуальных моделях Claude, есть страница «Лучшие практики промптинга».

Claude Opus 5 сделан под сложную агентскую разработку и корпоративные задачи, с особой силой в задачах с длинным горизонтом. На существующих промтах от Claude Opus 4.8 он работает хорошо из коробки. Приёмы ниже покрывают те поведения, которые чаще всего требуют донастройки.

Изменения в API при переезде с Claude Opus 4.8 описаны в гайде по миграции: мышление теперь включено по умолчанию, а выключить его можно только на уровне усилий `high` и ниже.

Как переведены термины

Оригинал	В переводе	Что это
effort	уровень усилий	Параметр, управляющий тем, сколько модель думает: <code>low</code> , <code>medium</code> , <code>high</code> , <code>xhigh</code>
thinking	мышление	Внутренние рассуждения модели до ответа. У Opus 5 включены по умолчанию
harness	обвязка	Каркас агента вокруг модели: цикл вызовов, инструменты, системные инструкции
subagent	субагент	Отдельный экземпляр модели, которому основной агент делегирует кусок работы
verbosity	многословность	Объём того, что модель говорит, в отличие от объёма того, что она думает

Что улучшилось

По сравнению с Claude Opus 4.8 наиболее значимы для промптинга следующие улучшения.

Агентская разработка

Claude Opus 5 сильнее всего на трудных задачах: многофайловые фичи, крупные рефакторинги, работа над фичей от начала до конца. Он доводит задачи до конца, а не оставляет заглушки и плейсхолдеры, и работает лучше всего, когда ему дали полную постановку сразу и оставили работать. На простых задачах вроде правки в один заход он тоже хорош, но разрыв с прошлыми моделями там меньше.

Ревью кода и поиск багов

Модель ревьюит код с высокой точностью и полнотой: находит реальные баги с высокой частотой за проход, и её дополнительные находки в основном реальные проблемы, а не ложные срабатывания. Точность держится и на низких уровнях усилий, что позволяет делать быстрый проход на этапе ревью и более глубокий позже.

Если в вашем промпте для ревью написано «сообщай только о проблемах высокой критичности» или «будь консервативен», модель может понять это буквально и сообщить меньше. Просите её сообщать обо всём, а фильтруйте отдельным проходом.

Эффективность на низких усилиях

Уровни `low` и `medium` дают высокое качество за долю токенов и задержки от высоких настроек. Начинайте с дефолтного `high` и подстраивайте по своим замерам: используйте `low` и `medium` свободно как основной рычаг стоимости и скорости везде, где качество держится, и поднимайтесь до `xhigh` для тяжёлой разработки и работы агентом. Если вы перенесли дефолты усилий с прошлой модели, прогоните развёртку по усилиям заново на своих тестах.

Зрение

Claude Opus 5 силён в понимании графиков, документов и диаграмм, а также в визуальном воспроизведении интерфейсов и фронтенда. Перепроверьте обходные

приёмы для зрения, которые вы настраивали под прошлые модели: возможно, они больше не нужны. Зрение работает лучше всего, когда у модели есть инструменты, чтобы итеративно анализировать, кропать и визуально проверять свою работу, и вложение в инструменты здесь окупается лучше, чем одно только мышление.

Работа с длинным контекстом

У Claude Opus 5 контекстное окно в 1 миллион токенов и как дефолт, и как максимум. Следование инструкциям, вызов инструментов и рассуждение остаются стабильными на всём окне.

Офисные и документные задачи

Модель делает сложные многолистовые таблицы с нетривиальными формулами и собирает хорошо структурированные презентации. Стили и шаблоны, которым надо следовать, задавайте промптом.

Координация нескольких агентов

Claude Opus 5 хорошо руководит командами субагентов: паттерн «писатель и проверяющий» работает, случаи, когда агенты затирают работу друг друга, редки. Для нагрузок, чувствительных к стоимости, ограничьте делегирование.

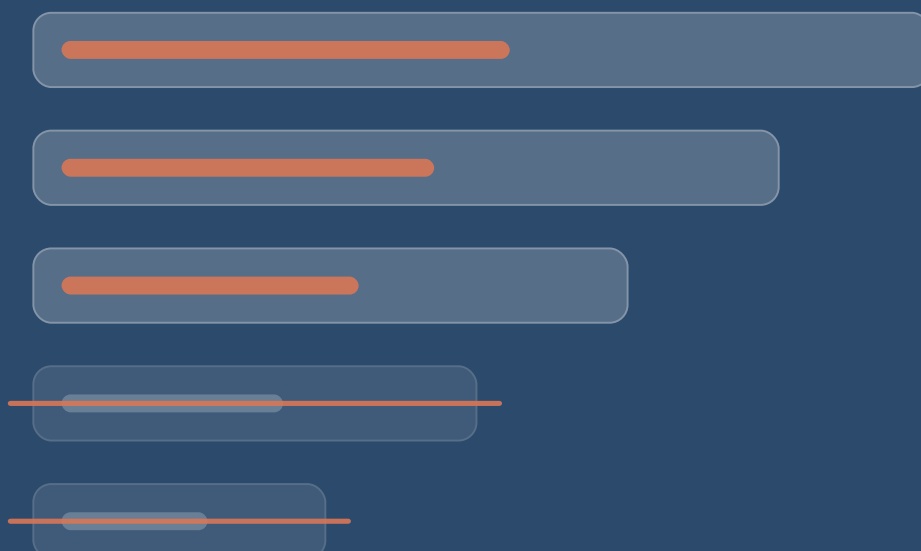
Меньше слов

Три раздела про объём: длина ответа, комментирование хода работы, размер письменных артефактов.

3 Длина ответа

4 Отчёты о ходе работы

5 Объём артефактов



Длина ответа и многословность

Ответы Claude Opus 5, обращённые к пользователю, по умолчанию длиннее, чем у прошлых моделей Opus.

Параметр уровня усилий управляет тем, сколько модель **думает**, а не тем, сколько она **говорит**: понижение усилий может сократить объём мышления, но не даёт надёжного сокращения видимого ответа. Чтобы управлять длиной ответа, просите об этом явно.

ПОД КАПОТОМ · ПРОВЕРЬТЕ У СЕБЯ

Anthropic описывает приём, который уже встроен в их собственные продукты. Если вам кажется, что ответы стали длиннее, добавлять «пиши короче» в каждый запрос бесполезно.

Правильное место для такой инструкции: системный промпт или ваш `CLAUDE.md`, причём в двух точках сразу, в начале и коротким напоминанием в конце.

Короткая инструкция на краткость работает. Например, для пользовательского продукта с многоходовым диалогом:

БЛОК ПРОМПТА · КОПИРОВАТЬ КАК ЕСТЬ

Отвечай сфокусированно, кратко и по делу. Дисклеймеры и оговорки держи короткими, основной объём ответа отдавай сути. Если просят что-то объяснить, дай обзор верхнего уровня, кроме случаев, когда явно попросили разбор в глубину.

В длинном системном промпте продублируйте инструкцию коротким напоминанием ближе к концу промпта:

БЛОК ПРОМПТА · КОПИРОВАТЬ КАК ЕСТЬ

```
<tone_preference>
Держи объём вывода в разумных пределах.
</tone_preference>
```

Отчёты о ходе работы

Claude Opus 5 охотно комментирует свои действия когда работает агентом: он склонен объявлять, что собирается сделать, и одно его сообщение в сессии агента часто длиннее, чем у прошлых моделей.

Модели помогает явное указание, как именно общаться с пользователем по ходу задачи. Чтобы убавить комментирование, опишите ритм и форму, которые вам нужны:

БЛОК ПРОМПТА · КОПИРОВАТЬ КАК ЕСТЬ

Перед первым вызовом инструмента скажи одним предложением, что собираешься делать. По ходу работы давай короткое обновление только тогда, когда нашёл что-то важное или меняешь направление. Закончив, начинай с результата: первое предложение отвечает на вопрос «что произошло» или «что ты нашёл», детали идут после, для тех, кому они нужны.

Чтобы, наоборот, прибавить комментирования или сменить его стиль, работает тот же рычаг в другую сторону: явно опишите, как должны выглядеть обновления, и дайте примеры.

ПОД КАПОТОМ · ПОЧЕМУ ЭТО НЕ КОСМЕТИКА

Комментирование хода работы оплачивается дважды. Каждое лишнее сообщение оседает в истории диалога и входит в контекст на каждом следующем шаге цикла агента.

На длинной задаче с сотней вызовов инструментов разница между «объявляю каждое действие» и «пишу, только когда меняю направление» набегает в реальные деньги.

Объём письменных артефактов

Отдельно от разговорной многословности: файлы, которые Claude Opus 5 пишет на диск, то есть отчёты, документы в Markdown, саммари, часто длиннее, чем на прошлых моделях.

Если в вашем продукте есть документы, которые пишет Claude, добавьте явную калибровку объёма:

БЛОК ПРОМПТА · КОПИРОВАТЬ КАК ЕСТЬ

Соотноси объём написанных документов с тем, чего требует задача: раскрой суть, но не набивай объём водой, дублирующими резюме и шаблонными разделами.

ЧТО МОДЕЛЬ ДОБАВИТ САМА

Вводный раздел, пересказывающий постановку

Резюме в конце, повторяющее текст выше

Раздел «дальнейшие шаги», которого не просили

Оглавление в документе на две страницы

ЧТО ОСТАЁТСЯ ПОСЛЕ КАЛИБРОВКИ

Суть в первом абзаце

Разделы, которые несут новую информацию

Объём под задачу, а не под ощущение солидности

Конец там, где закончилось содержание

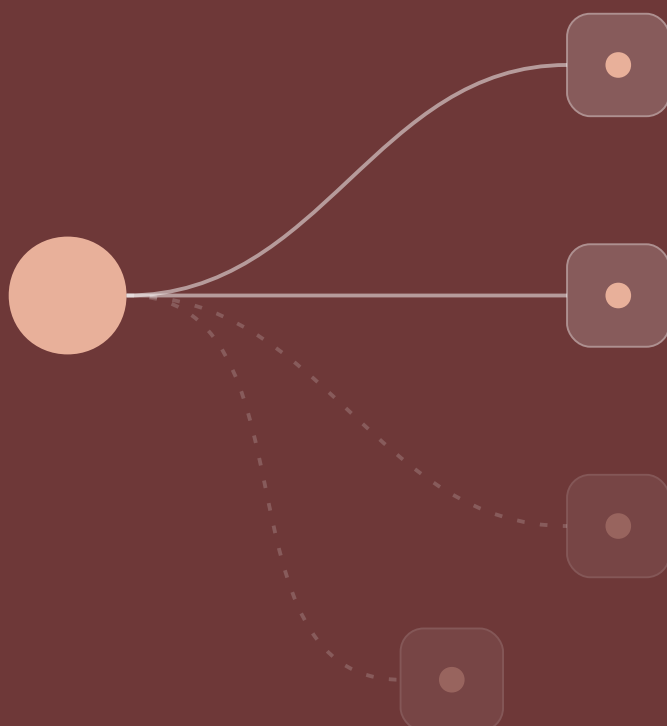
Границы и делегирование

Где модель делает больше, чем просили, и как это ограничить.

6 Границы задачи

7 Субагенты

8 Самокоррекция



Границы задачи и лишняя перепроверка

Claude Opus 5 проверяет свою работу сам, без указаний.

Если в вашем промпте есть явные инструкции на верификацию, такие как «добавляй финальный шаг проверки для любой нетривиальной задачи» или «используй субагента для проверки», **уберите их**. Такие инструкции вызывают у Claude Opus 5 избыточную перепроверку, а их удаление сокращает пустой расход токенов без потери качества. То же касается унаследованной обвязки, которая добавляет отдельные шаги верификации.

УБРАТЬ ИЗ ПРОМПТА

«Добавляй финальный шаг верификации»

«Перепроверь ответ перед выдачей»

«Используй субагента для проверки»

Отдельные шаги проверки в унаследованной обвязке

ОСТАВИТЬ

Внешние проверки: тесты, линтер, типы

Отдельный проход ревью как самостоятельная задача

Человек в контуре на необратимых действиях

Ограничение границ задачи

Claude Opus 5 может также расширять границы задачи: добавлять шаги, о которых не просили, или применять собственное суждение о том, какой задача должна быть. Для узких задач ограничивайте объём явно:

БЛОК ПРОМПТА · КОПИРОВАТЬ КАК ЕСТЬ

Делай то, о чём попросили, и в том объёме, в каком просили. Рутинные решения принимай сам, а сверяйся только тогда, когда разные прочтения запроса ведут к принципиально разной работе. Если запрос выглядит ошибочным или есть подход лучше, скажи об этом одним предложением и продолжай выполнять задачу как поставлена, вместо того чтобы молча её сузить, расширить или подменить. Доводи задачу до конца и останавливайся, не совершая действий, которые явно выходят за рамки запрошенного.

ПОД КАПОТОМ · ЭТО РАЗВОРОТ НА 180 ГРАДУСОВ

Год назад «добавь шаг верификации» было базовой гигиеной промптинга и попадало во все шаблоны системных промптов. Теперь тот же пункт официально помечен как вредный.

Практический вывод: оставьте только те проверки, которые делает внешний контур. Проверки, которые модель делает сама с собой, она и так сделает.

Контроль запуска субагентов

Claude Opus 5 делегирует субагентам охотнее, чем прошлые модели.

Делегирование окупается на действительно независимых крупных участках работы, но умножает стоимость и время, когда его применяют к мелким задачам. Если ваша обязанность поддерживает субагентов, дайте явные указания, какие сценарии оправдывают делегирование, либо поставьте детерминированные лимиты на число запускаемых агентов. Например:

БЛОК ПРОМПТА · КОПИРОВАТЬ КАК ЕСТЬ

Делегируй субагенту только крупные задачи, которые действительно независимы и параллелятся, например широкое исследование по множеству файлов. Не делегируй то, что можешь закончить сам за несколько вызовов инструментов, и не используй субагентов, чтобы проверить или перепроверить собственную работу. Если задачу закрывает один субагент, запускай одного, а не нескольких, и держи число запусков низким.

Самокоррекция

Claude Opus 5 хорошо ловит и исправляет собственные ошибки без всяких указаний.

Избегайте инструкций на повторные проверки, которые модель и так выполняет: «перепроверь свой ответ», «перевеифицируй перед ответом». Как и инструкции на

верификацию, они складываются с собственным поведением модели и добавляют стоимость, не улучшая результат.

Модель также чаще прошлых версий проговаривает исправления к своим более ранним утверждениям, что бывает нежелательно в продуктах, обращённых к пользователю. Чтобы ограничить проговаривание значимыми случаями:

БЛОК ПРОМПТА · КОПИРОВАТЬ КАК ЕСТЬ

Исправляй сказанное ранее только тогда, когда ошибка изменила бы код, выводы или решения пользователя. Формулируй исправление прямо и коротко, затем продолжай задачу. Мелкие оговорки, которые для пользователя ничего не меняют, просто исправляй и иди дальше, не отмечая этого.

Тонкие места

Артефакты выключенного мышления, разбор от канала и первоисточники.

9 Когда мышление выключено

10 Разбор от канала

+ Источники



Работа с выключенным мышлением

Claude Opus 5 работает с включённым мышлением по умолчанию, и выключить его можно только на уровне усилий **high** и ниже.

При выключенном мышлении в видимом выводе модели изредка могут появляться два артефакта. Основное средство от обоих: **оставить мышление включённым**, а расход токенов резать понижением уровня усилий, а не выключением мышления. Для большинства задач включённое мышление на уровне усилий **low** работает лучше, чем выключенное мышление при сопоставимой стоимости.

Артефакт первый: вызов инструмента текстом

При выключенном мышлении модель иногда пишет вызов инструмента прямо в текст, обращённый к пользователю, вместо того чтобы выдать структурный блок **tool_use**. Ход при этом завершается нормально, а вызов никогда не выполняется. В циклах агента утёкший текст остаётся в истории диалога, поэтому затрагивает и последующие ходы. Чаще всего это происходит на нагрузках с большим числом инструментов, например на поиске.

Артефакт второй: внутренние XML-теги в выводе

При выключенном мышлении модель может выдать теги **<thinking>** или другие внутренние XML-теги в видимый ответ. Если в вашем системном промпте есть правило, запрещающее модели думать или рассуждать, уберите его: такая инструкция усиливает утечку тегов.

Одна инструкция против обоих артефактов

Для интеграций, которые обязаны держать мышление выключенным, одна совмещённая инструкция смягчает оба артефакта. Она даёт модели явное разрешение сказать что-то перед вызовом инструмента, альтернативу угадыванию, когда ни один инструмент не подходит, и общее правило против внутренних тегов:

Когда используешь инструмент, можешь сначала сказать одну короткую фразу. Если ни один инструмент не выражает того, о чём просит пользователь, скажи об этом, а не угадывай. Не включай внутренние или системные XML-теги в свой ответ.

Инструкции, которые называют теги мышления по имени, работают хуже, чем эта общая форма, поэтому не упоминайте их отдельно.

Разбор от канала

Перевод закончен. Дальше моё, отделено специально, чтобы не смешивать источник и комментарий.

Пять выводов

- 1 Промпт худеет, а не толстеет.** Три раздела из восьми просят удалить инструкции: границы задачи, самокоррекция, выключенное мышление. Ни один раздел не просит добавить проверок. Раздел про субагентов просит не удалить, а поставить лимит.
- 2 Уровень усилий это ползунок цены, а не качества.** Anthropic прямо пишет: используйте `low` и `medium` свободно, поднимайтесь до `xhigh` для тяжёлой разработки и работы агентом. Дефолты, перенесённые с прошлой модели, скорее всего переплачивают.
- 3 Длина ответа и глубина мышления развязаны.** Понижение усилий не сделает ответ короче. Хотите короче, пишите это словами, а в длинном системном промпте продублируйте напоминанием в конце.
- 4 Субагенты стали дешёвым рефлексом модели.** Она делегирует охотнее, чем нужно. Если у вас в промпте написано «параллель по умолчанию», к нему теперь нужен лимит, иначе на мелкой задаче запустится флот.
- 5 Выключать мышление больше не способ экономить.** Включённое мышление на низких усилиях и дешевле, и надёжнее, чем выключенное. С выключением связаны два класса артефактов в выводе.

Что сделать за десять минут

Откройте свой системный промпт или `CLAUDE.md` и найдите по поиску слова «проверь», «верификация», «перепроверь», «финальный шаг», «субагент».

- Всё, что просит модель проверить саму себя, вычеркните.
- Всё, что запускает агентов без лимита, ограничьте числом.
- Инструкцию на краткость поставьте в начале промпта, а если промпт длинный, продублируйте коротким напоминанием в конце.
- Прогоните свои типовые задачи на `low` и `medium` и сравните результат с `high`. Замер на своих задачах важнее любых рекомендаций, включая эти.

Чего в гайде нет

Честная рамка: документ описывает поведение модели, а не гарантирует его.

Anthropic не даёт цифр, насколько именно короче станут ответы, на сколько процентов падает расход при удалении инструкций на верификацию и в каком проценте случаев текут XML-теги. Слова «изредка», «часто» и «охотнее» остаются словами. Проверять на своих задачах придётся самому.

Источники

Перевод сделан с англоязычной страницы документации Anthropic.
Перевод неофициальный: при любом расхождении верен оригинал.
Ссылки кликабельны.

ПЕРВОИСТОЧНИКИ

- 1 **Anthropic · Prompting Claude Opus 5** · platform.claude.com
- 2 **Anthropic · What's new in Claude Opus 5** · platform.claude.com
- 3 **Anthropic · Migration guide (Opus 4.8 → Opus 5)** · platform.claude.com
- 4 **Anthropic · Effort** · platform.claude.com
- 5 **Anthropic · Extended thinking** · platform.claude.com
- 6 **Anthropic · Prompting best practices** · platform.claude.com
- 7 **Anthropic · Context windows** · platform.claude.com

ПОД КАПОТОМ · 2026

Спасибо, что дочитал.

Перевод сделан для канала «Under the Hood»: разборы ИИ-индустрии без хайпа, с источниками и цифрами. Если пригодилось, забирай продолжение.



Under the Hood

[@gorilla_under_hood](#)

Перевод

Илья Утов

Telegram

[@wh_zt](#)

LinkedIn

[linkedin.com/in/ilyautov](https://www.linkedin.com/in/ilyautov)