

Оценка планирования ИИ-агентов в аналитике контакт-центров

не клиентское обслуживание,
а аналитика контакт-центра

Олег Зельдин. Апекс Берг

Gartner Predicts Over 40% of Agentic AI Projects Will Be Canceled by End of 2027

SYDNEY, Australia, June 25, 2025

Более 40% проектов с ИИ-агентами будут отменены к концу 2027 года из-за:

- “растущих затрат;”
- “неясной бизнес-ценности;”
- “недостаточного контроля рисков.  gartner.com”

Reuters передал это как рыночный сигнал: крупные поставщики активно инвестируют в ИИ-агентов, но Gartner предупреждает, что многие проекты пока остаются ранними экспериментами или пилотами, часто движимыми ажиотажем, а не ясной экономикой.  Reuters

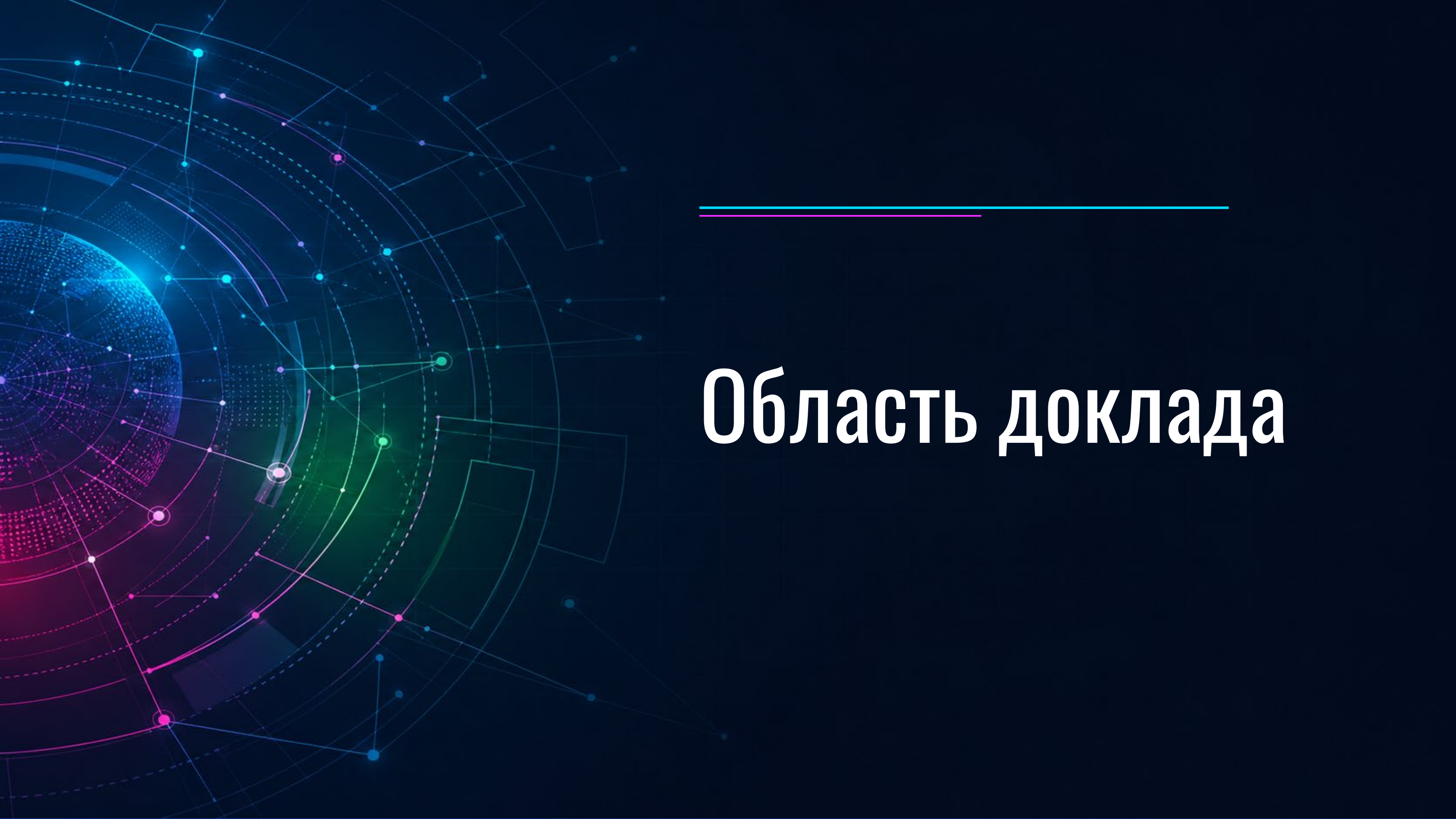
Еще важная деталь: Gartner использовал термин **agent washing** — когда существующие чат-боты, помощники или RPA-решения переупаковывают как «агентный ИИ», хотя у них нет существенных агентных возможностей. Gartner оценивал, что из тысяч поставщиков агентного ИИ реально агентными являются около 130.  gartner.com

Источник:

Tool-Aware Planning in Contact Center AI: Evaluating LLMs through Lineage-Guided Query Decomposition

[Varun Nathan](#), [Shreyas Guha](#), [Ayush Kumar](#)

<https://arxiv.org/abs/2602.14955>



Область доклада

Агент обслуживания

ведет диалог, отвечает,
передает оператору

Аналитический агент

работает в контуре
руководителя/аналитика,
понимает источники,
фильтры, строит цепочки
анализа



«Сравнить оценки контроля качества по критериям “профессионализм” и “процедуры решения проблемы” для нерешенных звонков, в которых тональность клиента изменилась с негативной на позитивную»

Оценка качества плана

покомпонентная
и по сравнению
с эталоном

Улучшение планирования

цикл исправления
ошибок и сохранения
истории

Сравнение моделей

тестирование 14
больших языковых
моделей на
планировании



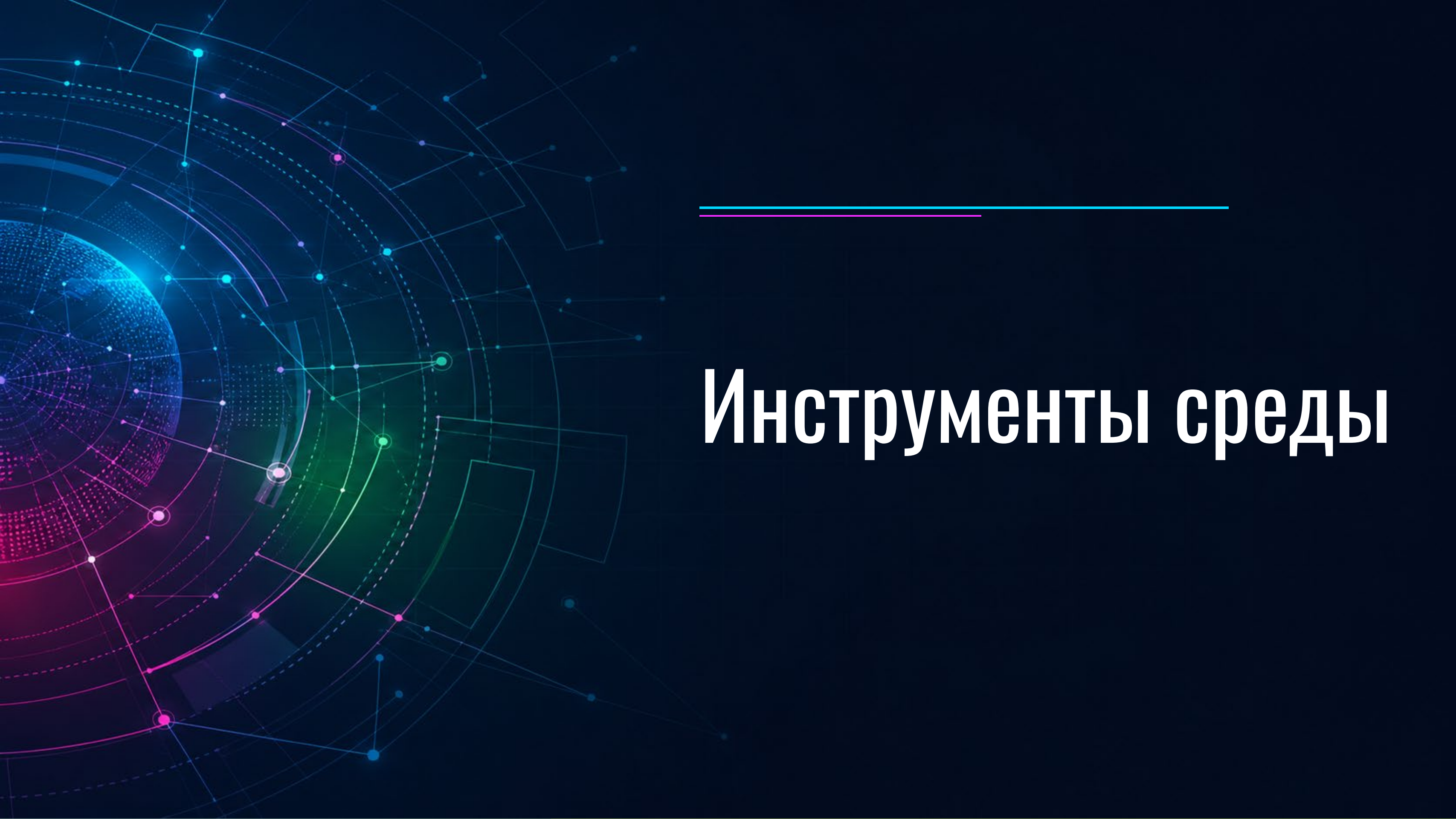
«Ответ»

проверка результата



«План»

проверка метода
получения результата

The background features a dark blue gradient with a complex network of glowing lines and nodes. On the left, a stylized globe is composed of a grid of points, with a bright blue and green glow emanating from its center. Concentric circular lines and various colored dots (blue, green, pink) are scattered across the frame, creating a sense of digital connectivity and data flow.

Инструменты среды



языковая модель
для объединения
результатов (LLM)

расшифровки
разговоров (RAG)

структурированные
данные (T2S)



T2S / Snowflake

«counts»

«trends»

«QA metrics»

«flags»

«duration»

«call drivers»



RAG / Transcripts

«текст разговоров»

«темы»

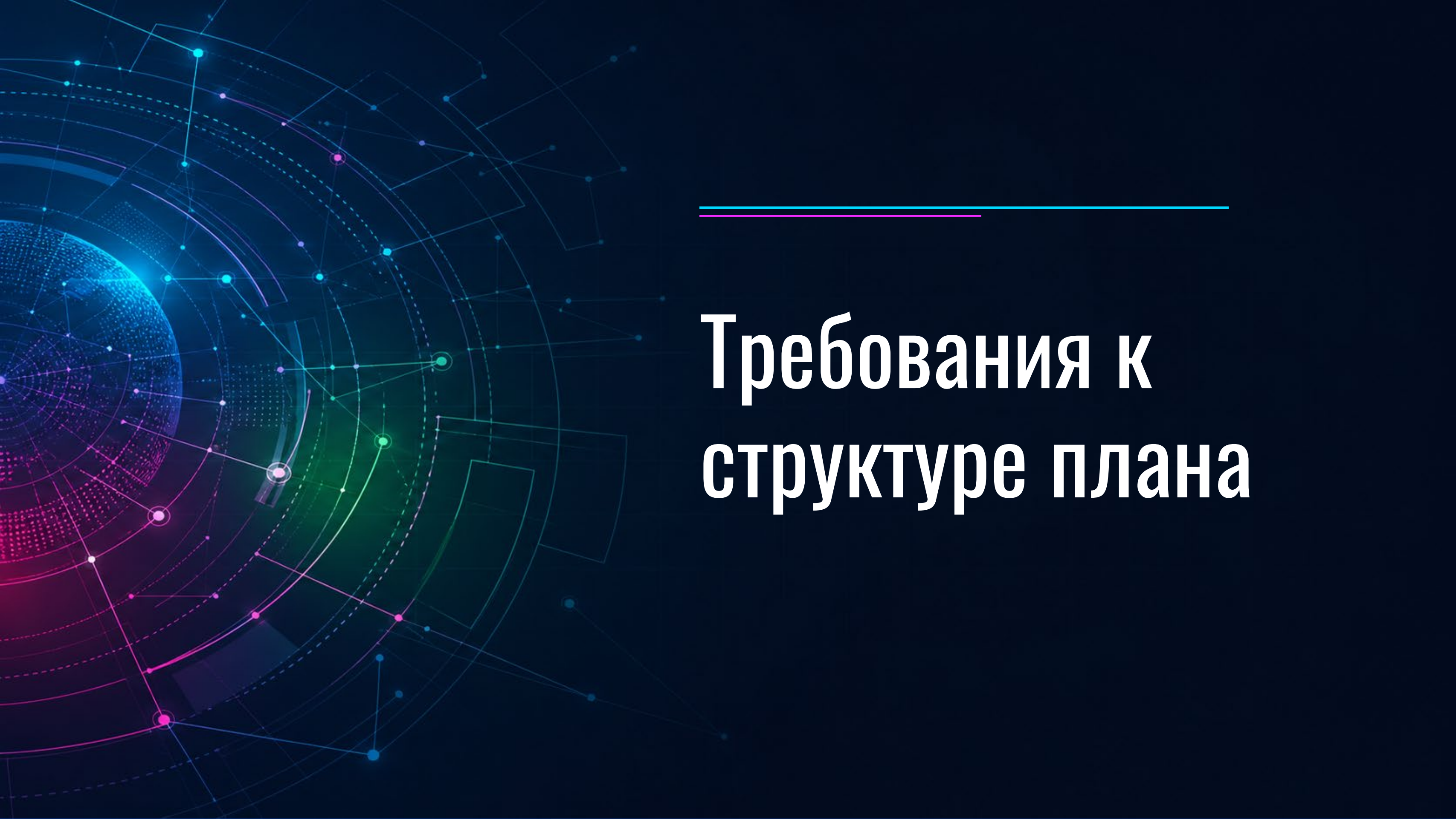
«формулировки»

«sentiment внутри звонка»

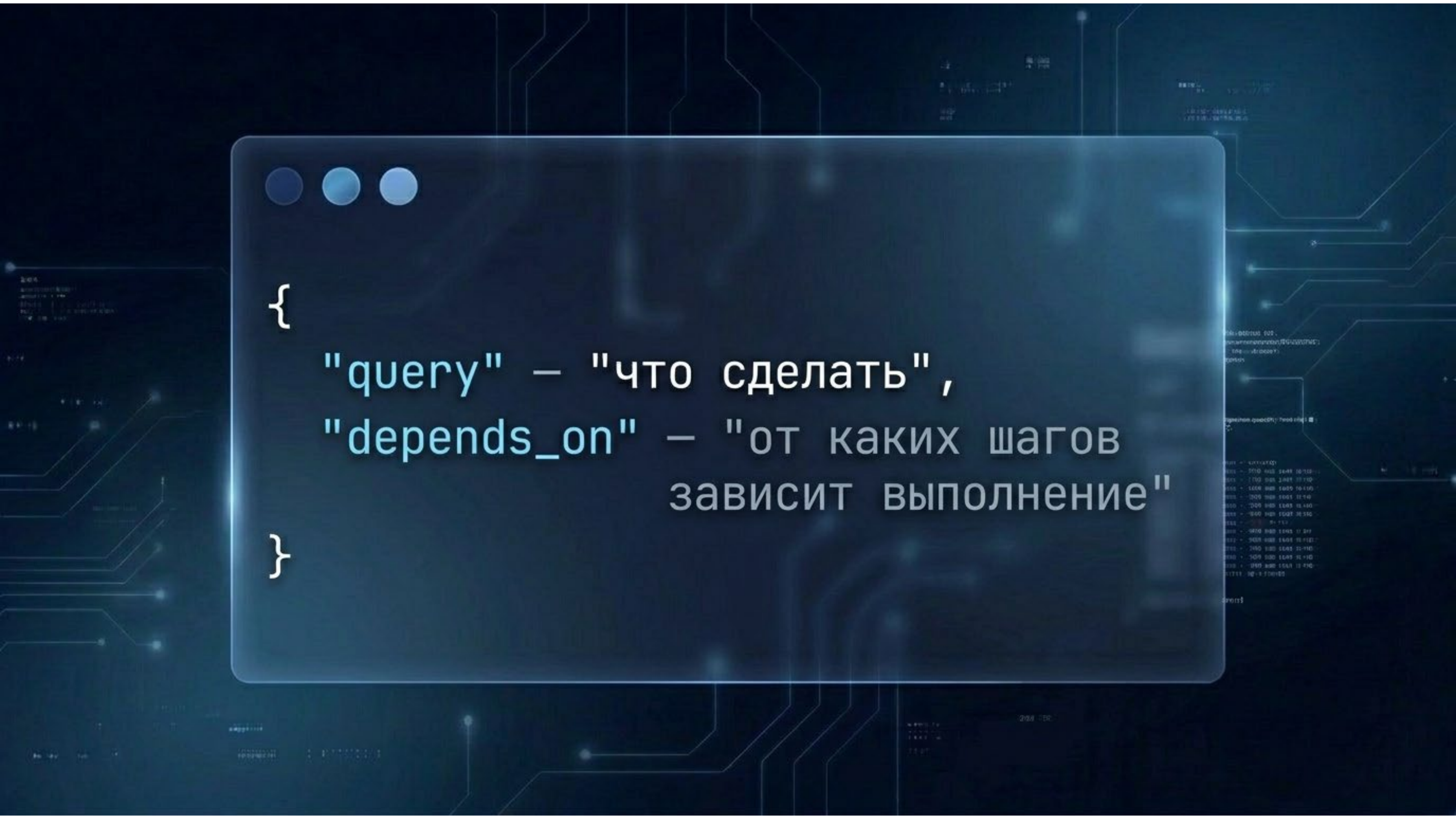


LLM

синтез, переформатирование, извлечение ID, объединение результатов

The background features a dark blue gradient. On the left side, there is a complex, glowing digital structure. It includes a sphere composed of many small, bright blue and purple dots, with lines radiating from it. Overlaid on this are several concentric, semi-transparent circular arcs in shades of blue and green. A network of thin, light blue lines connects various points across the scene, some of which are highlighted with larger, glowing dots in blue, green, and pink. The overall aesthetic is futuristic and technological.

Требования к структуре плана



```
{
```

```
  "query" — "что сделать",  
  "depends_on" — "от каких шагов  
                  зависит выполнение"
```

```
}
```

«2. Как устроен tool-aware plan»

«План — это не текстовое рассуждение, а структура: шаг, инструмент, инструкция, зависимости и исполнимый формат.»

Домен
контакт-центры

Дата иетовъзевит.	35
Релечант	х61нэ10
запросы	388
	9.0%

Задача
аналитические запросы

Инструменты

T2S/Snowflake, RAG/транскрипты, LLM-синтез

Фокус оценки
качество плана

Бенчмарк: 14 LLM, 600 запросов



$$P = \langle s_k \rangle_{k=1}^n, s_k = (t_k, p_k, D_k)$$

Инструмент

Инструкция / prompt

Зависимости от
предыдущих шагов

Выход:
JSON-план

T2S / Snowflake

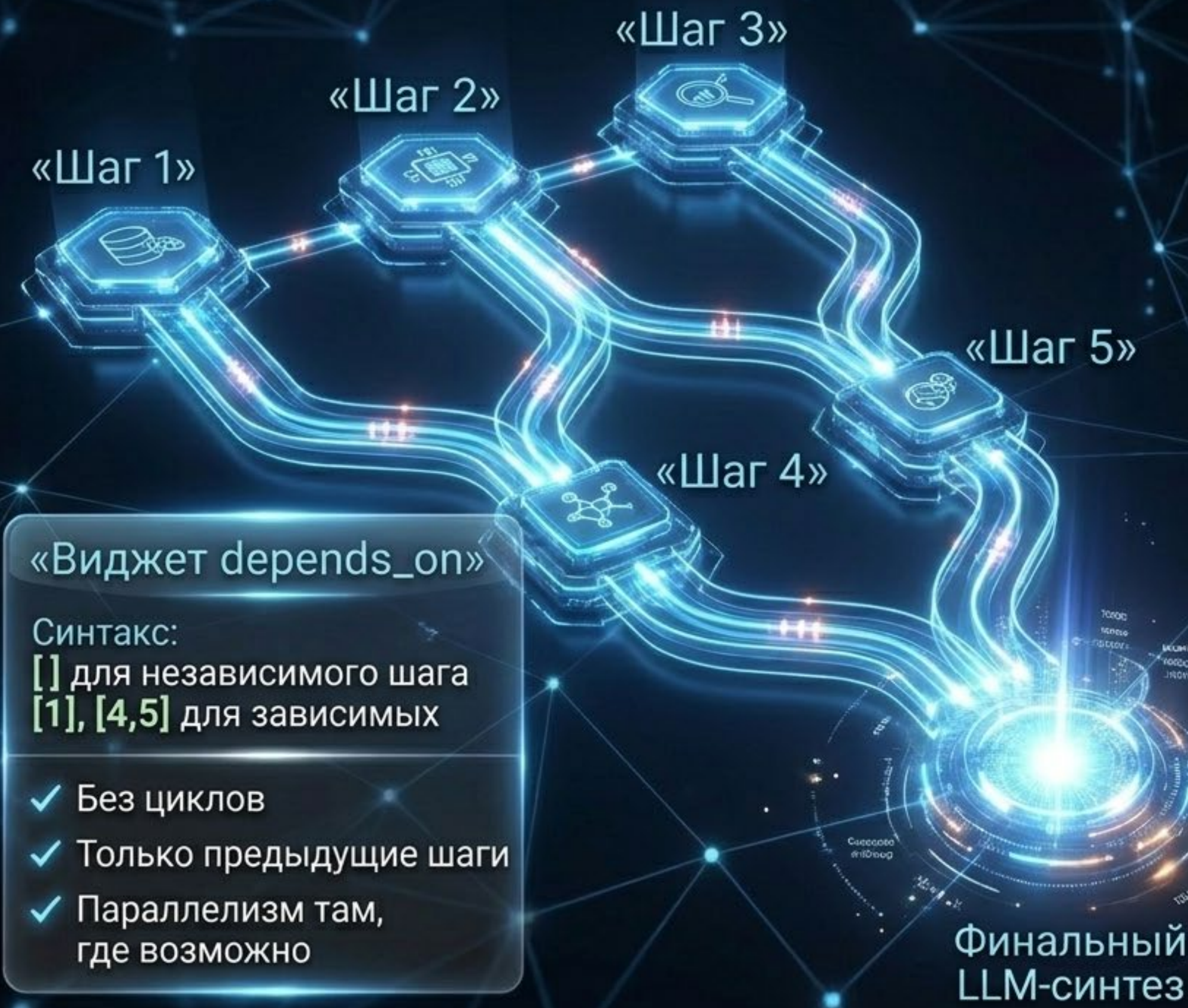
Структурированные
данные: counts,
counts, trends,
QA metrics, flags,
duration, call drivers

RAG / Transcripts

Текст разговоров,
темы,
формулировки,
sentiment внутри
звонка

LLM

Синтез,
переформатирование,
извлечение ID,
объединение
результатов



«Виджет depends_on»

Синтакс:

[] для независимого шага
[1], [4,5] для зависимых

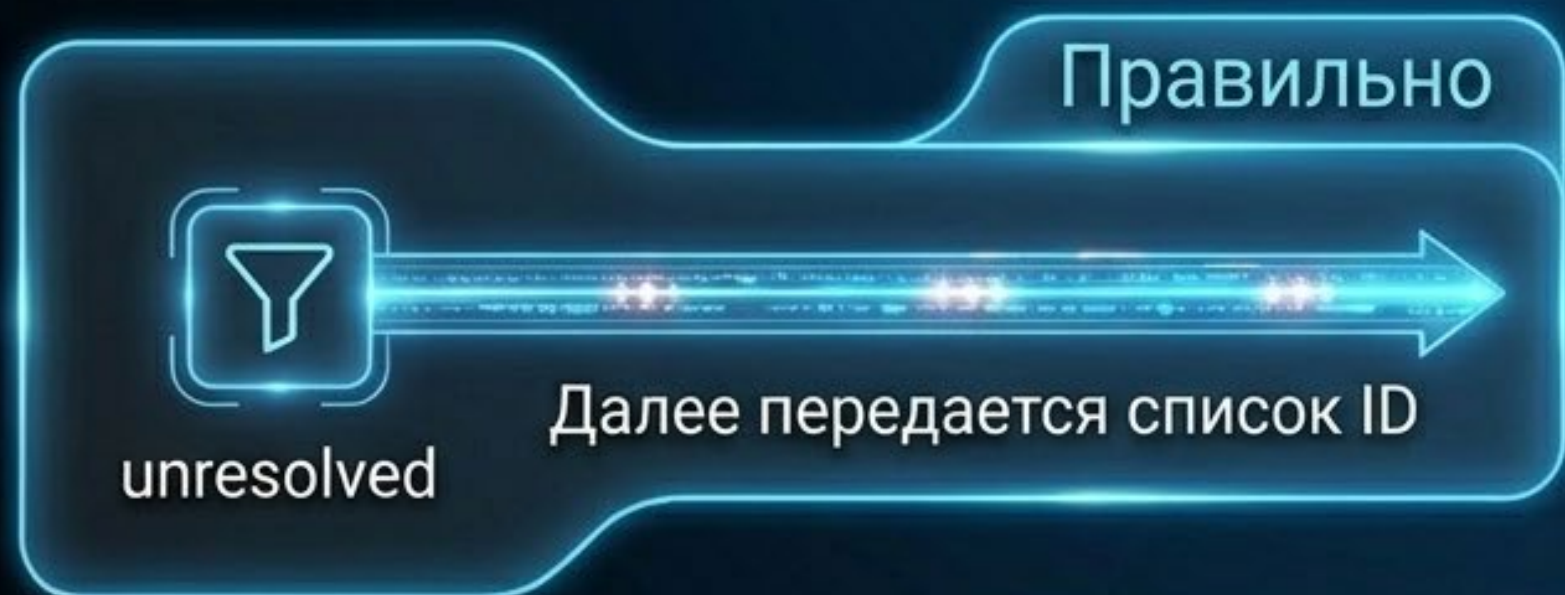
- ✓ Без циклов
- ✓ Только предыдущие шаги
- ✓ Параллелизм там, где возможно

Механика Placeholders



⚠ **Правило:** ссылаться только на завершённые предыдущие шаги

«Правило без повторной фильтрации»



Риски: лишние join, задержка, рассинхронизация условий

Пример правильной декомпозиции

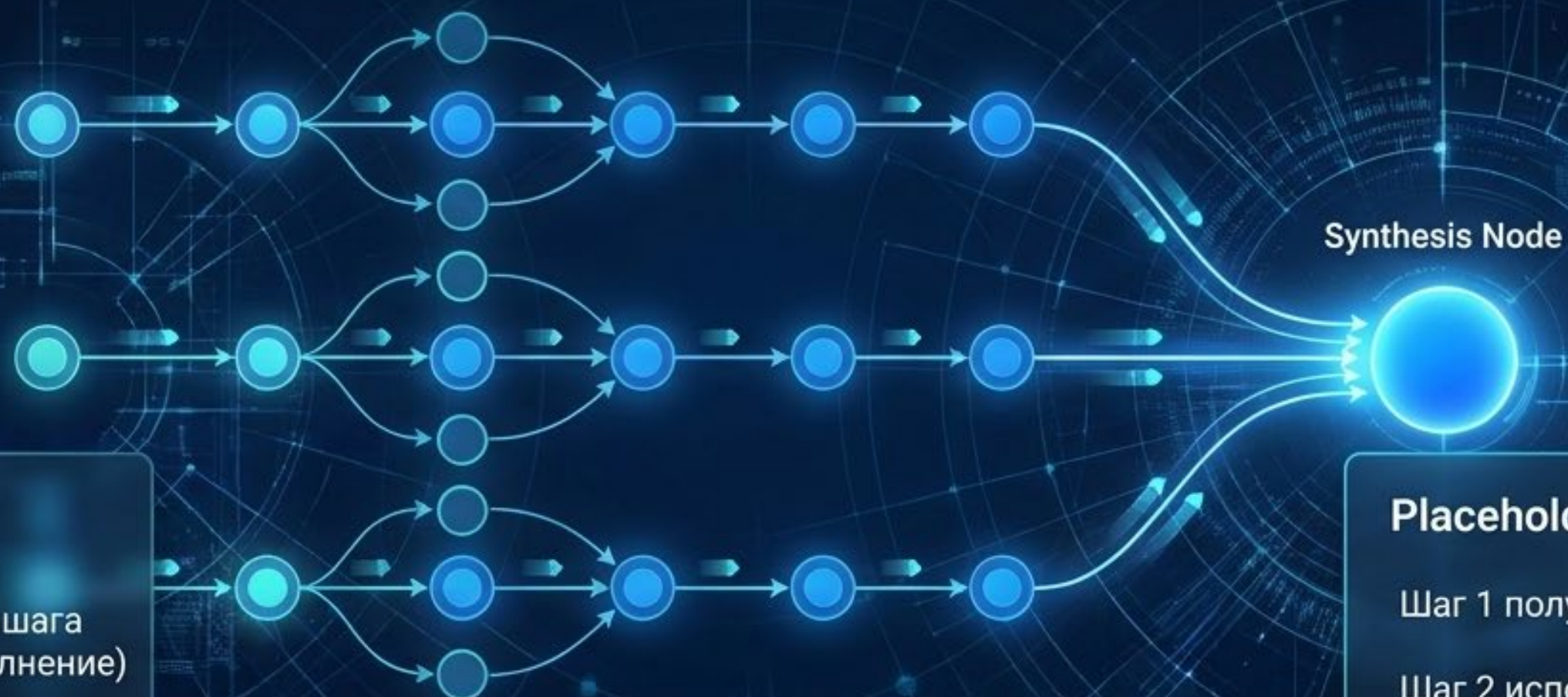


шаг 1 → шаг 2 → шаг 3



шаги 4 и 5 могут
выполняться
параллельно

Directed Acyclic Graph (DAG) & Dependencies



depends_on

[] — для независимого шага
(параллельное выполнение)

[1] или [4, 5] — для
зависимых шагов

⊗ Без циклов. Только
предыдущие шаги.

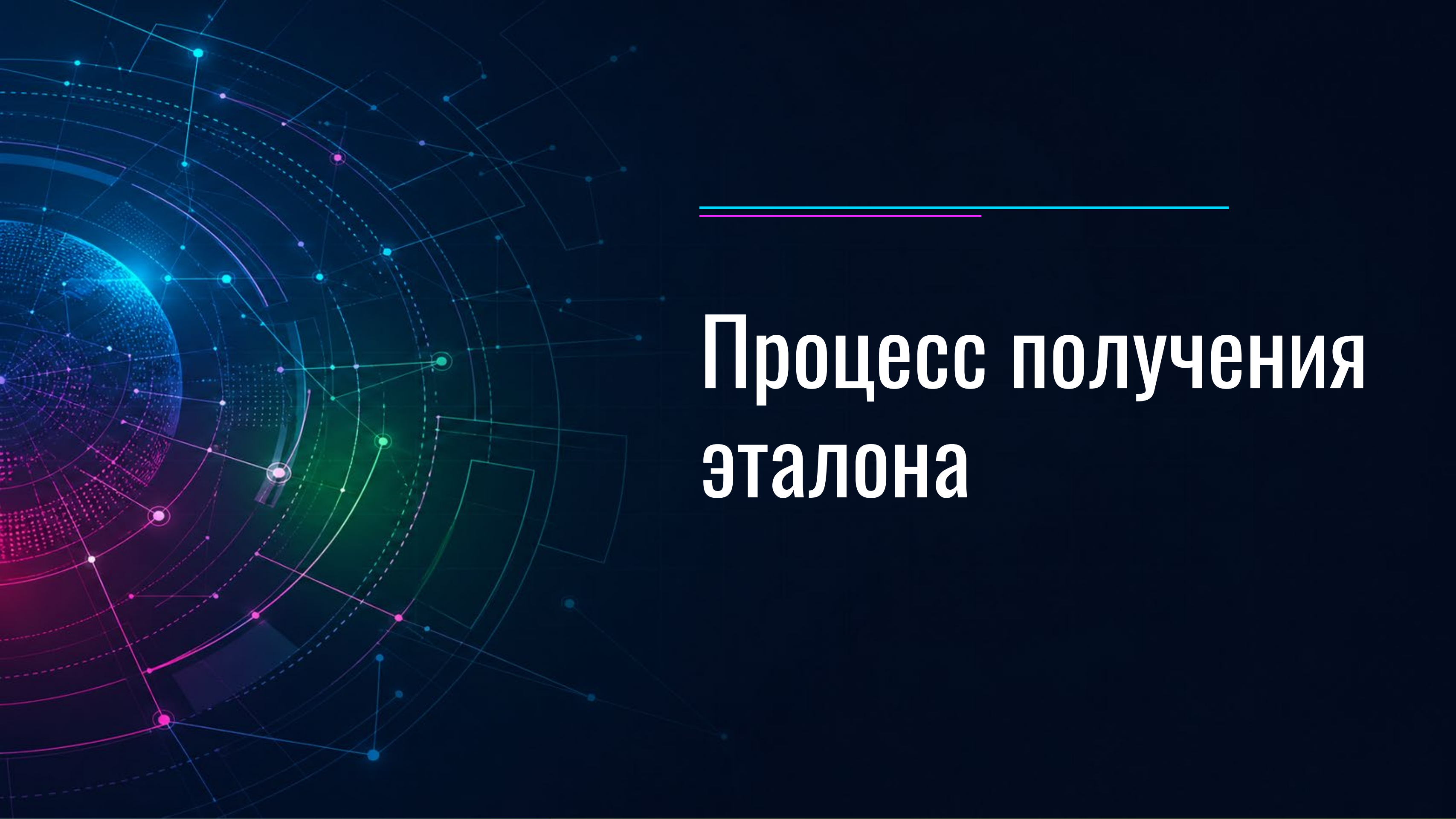
Placeholders

Шаг 1 получает **interaction_ids**

Шаг 2 использует **(1)**

Шаг 3 использует **(2)**

⚠ Ссылаться только
на завершённые
предыдущие шаги.

The background features a dark blue gradient with a complex network of glowing lines and nodes. On the left, a stylized globe is composed of numerous small, bright blue and purple dots. Concentric circular lines, some solid and some dashed, radiate from the center, creating a sense of depth and motion. The overall aesthetic is futuristic and technological.

Процесс получения эталона



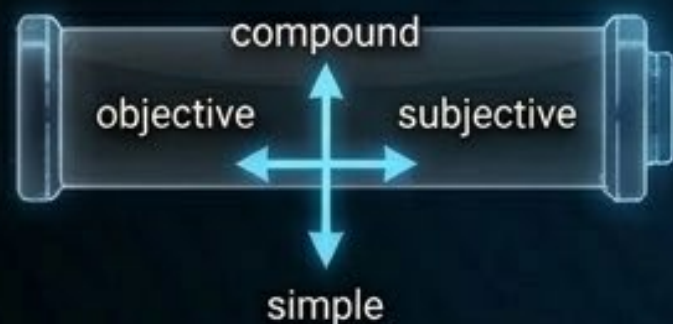
запрос

первый план
(черновой
вариант за один
проход)

проверка

улучшение

1. Контролируемая генерация запросов



2. One-shot генерация начального плана



3. Цикл Evaluator→Optimizer



4. Human verification финального плана



Схема evaluator→optimizer



Правила цикла



Плюсы offline-цикла

- ✓ Масштабируемая кураторская проверка
- ✓ Нет runtime-затрат
- ✓ Видимая история исправлений

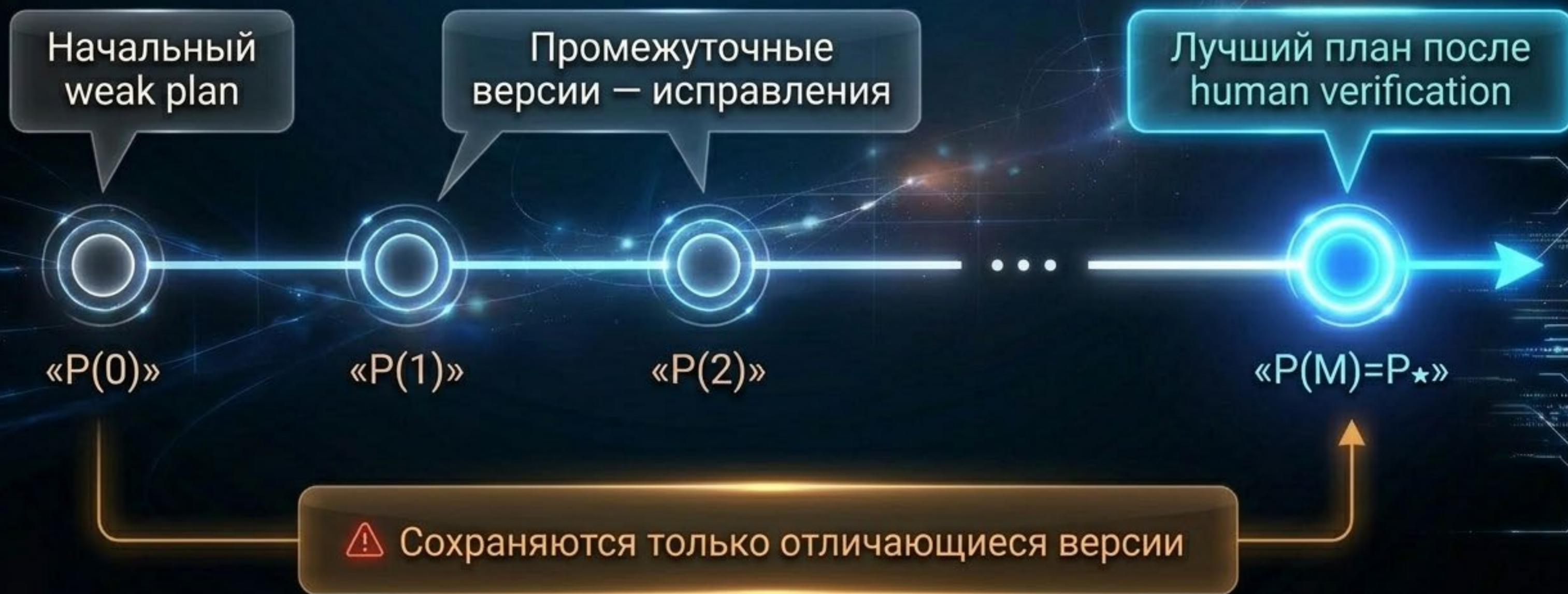


Ограничения метода

- ⚠ Инструменты не запускаются
- ⚠ SQL не проверяется фактически
- ⚠ Пустые join и race conditions не измеряются

3. Lineage: как план исправляет сам себя

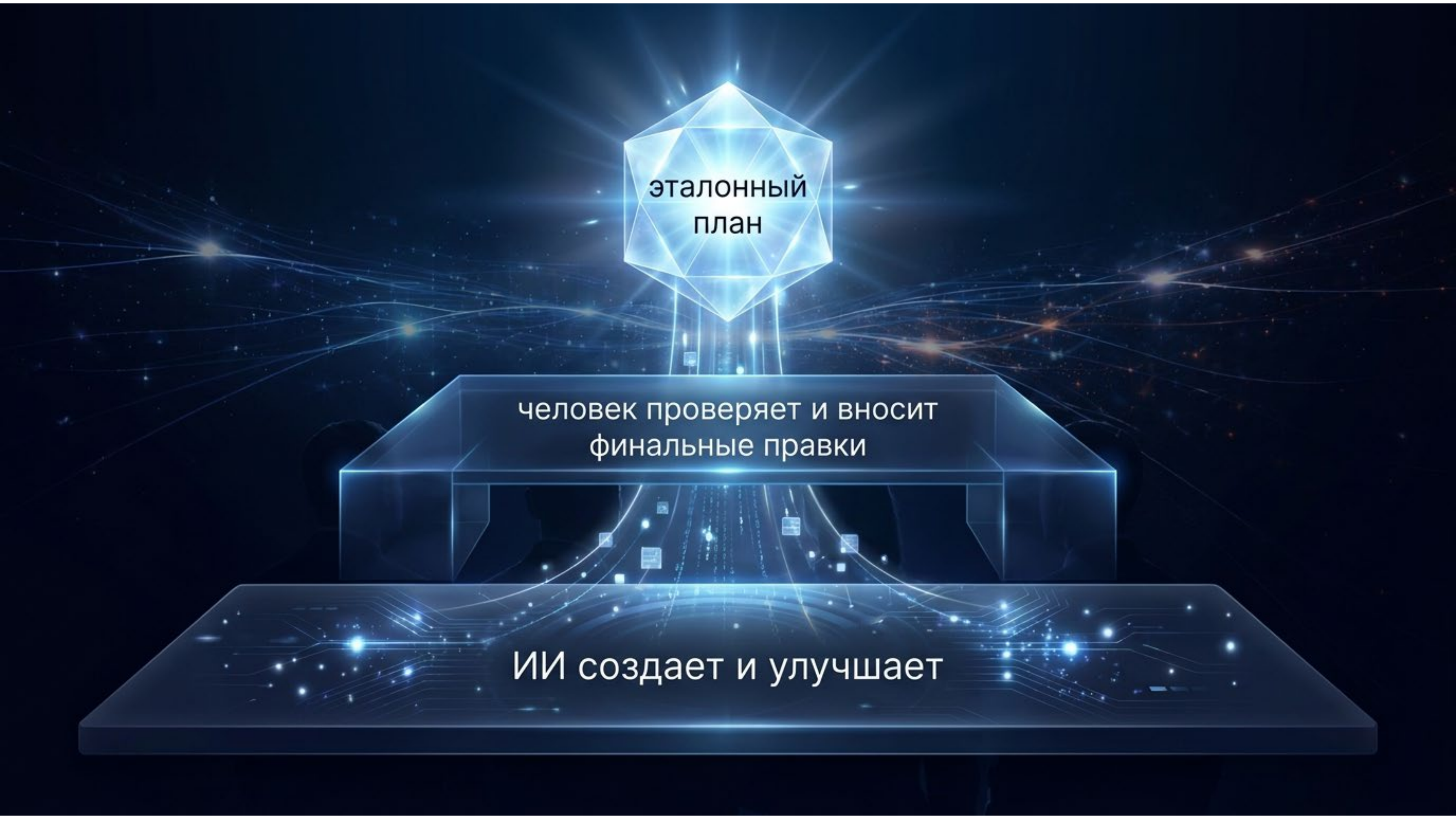
Тезис: Lineage хранит версии плана и показывает, где агент ошибался: инструмент, промпт, зависимость, формат, избыточность.



Lineage: Эволюция и Контроль Версий Плана



«Lineage хранит версии плана и показывает, где агент ошибался: инструмент, промпт, зависимость, формат, избыточность»



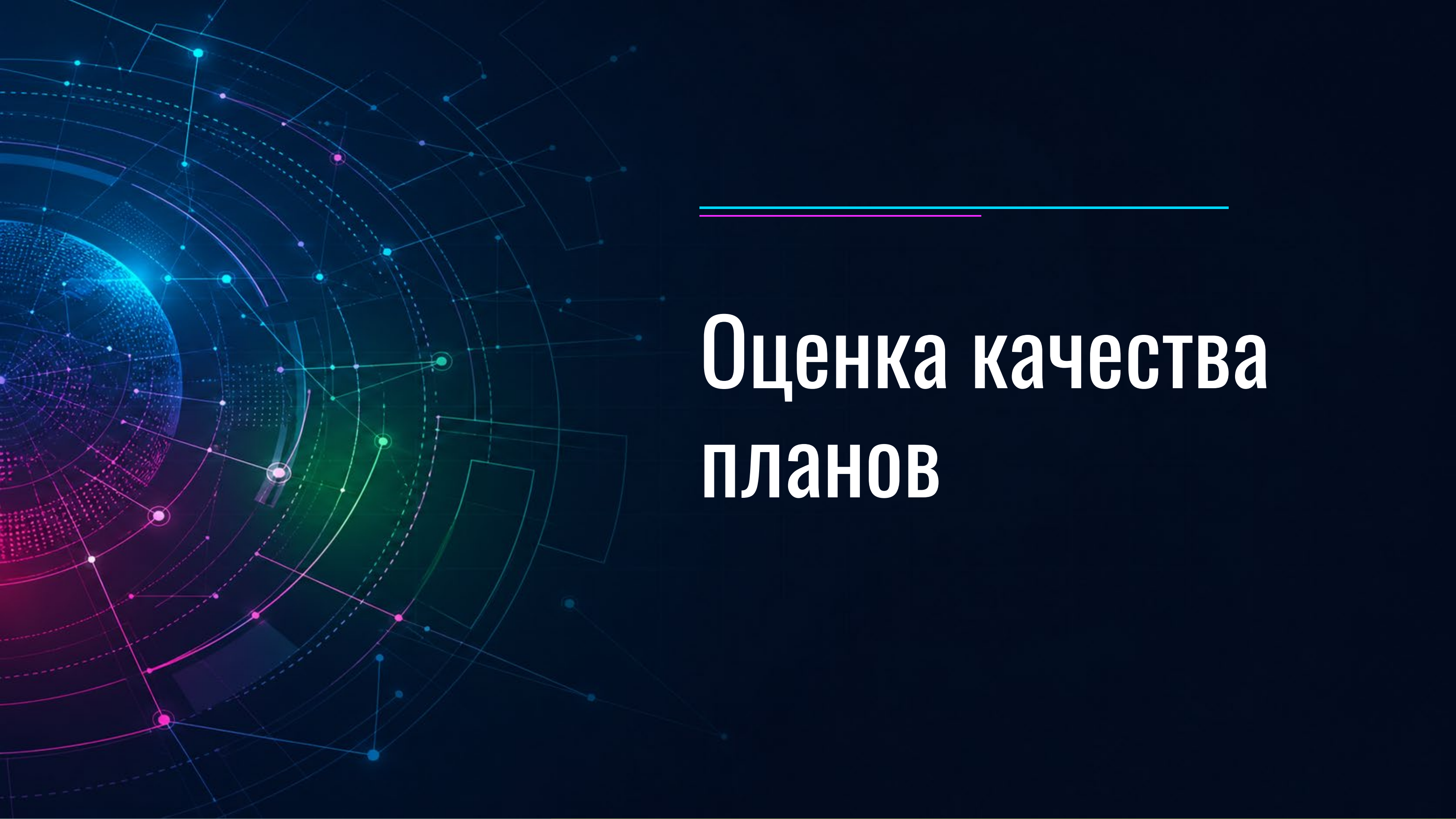
эталонный
план

человек проверяет и вносит
финальные правки

ИИ создает и улучшает

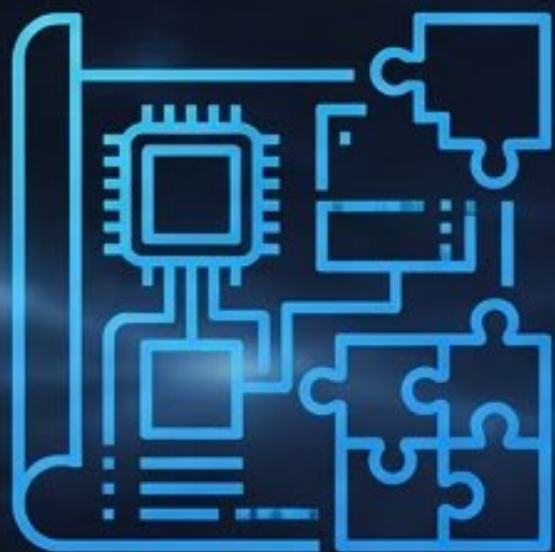
ПРОЦЕСС АНАЛИЗА ЗВОНКОВ



The background features a dark blue gradient with a complex network of glowing lines and nodes. On the left, a sphere is composed of numerous small, bright blue and purple dots, with lines radiating from its center. To the right of the sphere, a series of concentric, semi-transparent circular arcs are visible, some in shades of green and blue. The overall aesthetic is futuristic and technological.

Оценка качества планов

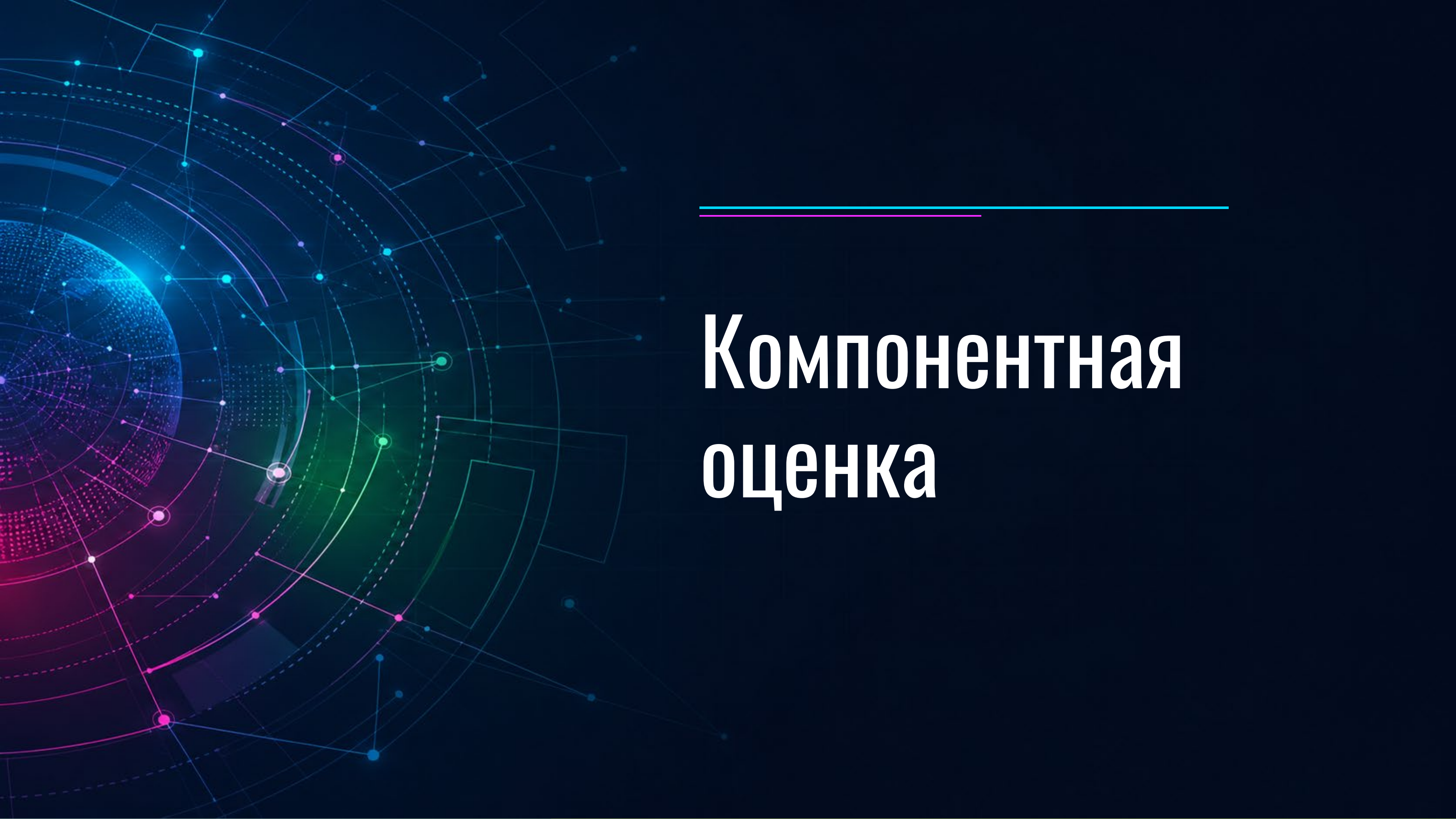
«Два способа оценки качества»



Покомпонентная
оценка: 7 критериев



Единая оценка:
сравнение с эталоном,
точность, полнота, F1

The background features a dark blue gradient with a complex network of glowing lines and nodes. On the left, a sphere is composed of numerous small, bright blue and purple dots, with lines radiating from its center. To the right of the sphere, a series of concentric, semi-transparent circular arcs are visible, some in shades of green and blue. The overall aesthetic is futuristic and technological.

Компонентная оценка

Effectiveness (70 баллов):
Корректность и исполнимость плана



Tool-Prompt Alignment (20)



Format Correctness (20)



Step Executability / Atomicity (15)



Query Adherence (15)



$$\text{SCORE}(P) = 100 \cdot \sum w_k m_k$$

Efficiency (30 баллов):
Оптимальность маршрута



Dependencies (10)



Redundancy (10)



Tool-Usage Completeness (10)

Effectiveness (70 баллов):
Корректность и исполнимость плана



Tool-Prompt Alignment (20)



Format Correctness (20)



Step Executability / Atomicity (15)



Query Adherence (15)



$$\text{SCORE}(P) = 100 \cdot \sum w_k m_k$$

ANALYTICAL FRAMEWORK | SYSTEM ARCHITECTURE v4.5



Таблица 1. Основные показатели	Показатель	2009	2009
Величина валового внутреннего продукта (в млрд. руб.)	Валовой внутренний продукт	44,5	1060
Величина валового внутреннего продукта (в млрд. руб.)	Валовой внутренний продукт	2009	308
Величина валового внутреннего продукта (в млрд. руб.)	Валовой внутренний продукт	1976	304
Величина валового внутреннего продукта (в млрд. руб.)	Валовой внутренний продукт	3055	53

КРИТЕРИЯ

Формат

Можно ли план прочитать и выполнить как структуру?

«Компонент 2: соответствие инструмента заданию»

Инструмент соответствует задаче



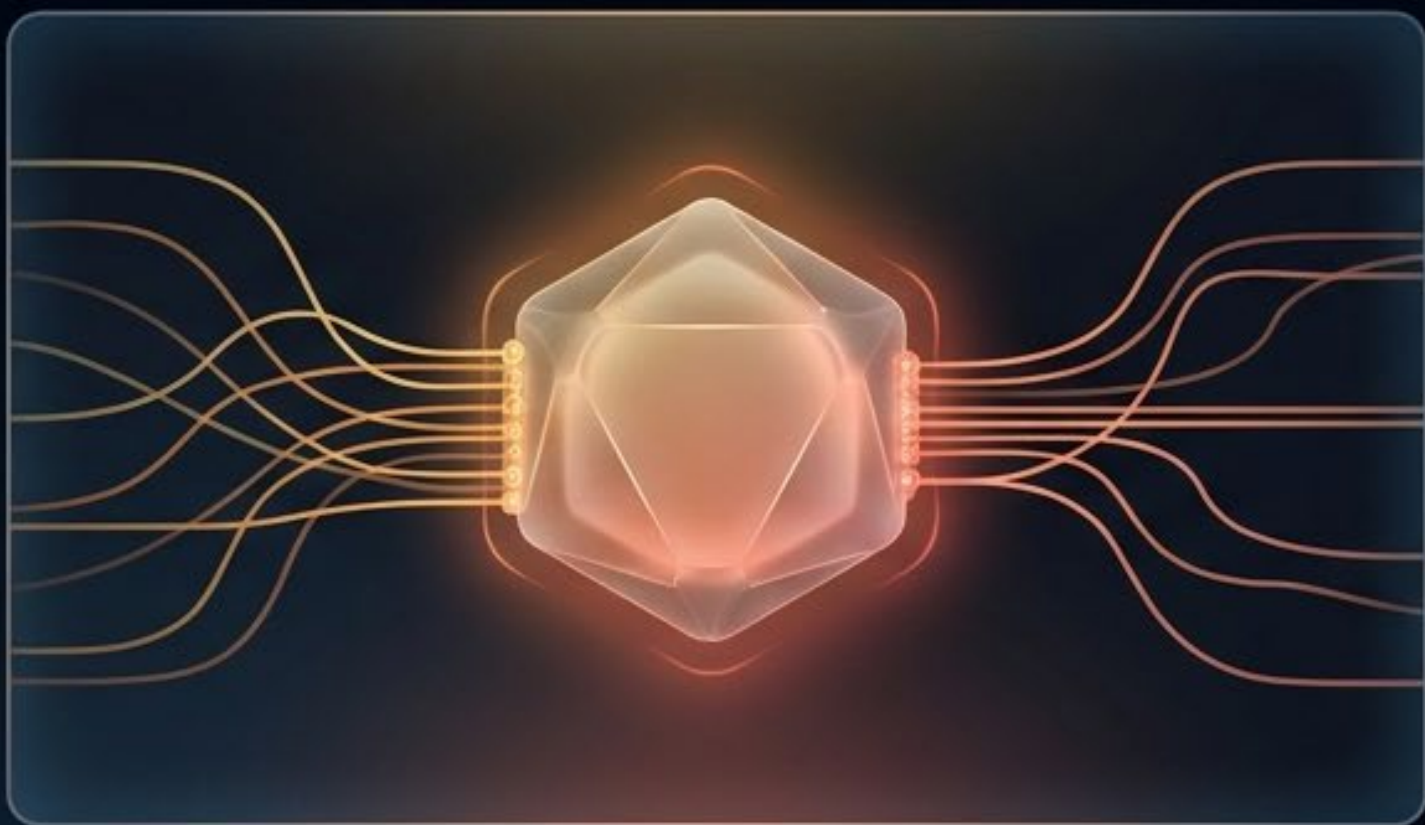
таблицы —
для метрик



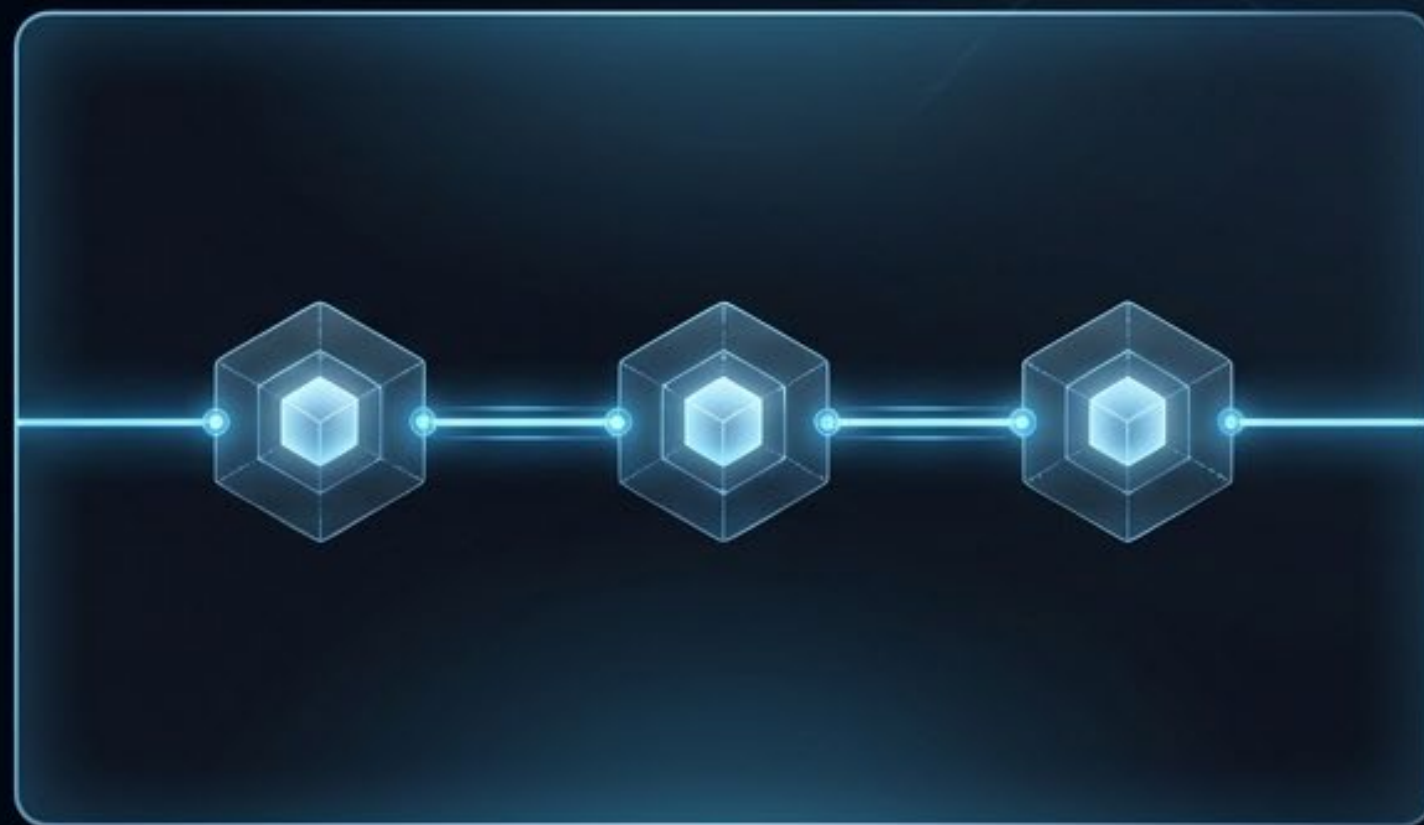
расшифровки —
для смысла
разговора

«Компонент 3: исполнимость шага»

Шаг можно реально выполнить



один шаг содержит сразу
несколько разных задач



Компонент 4: соответствие исходному вопросу

План отвечает именно на заданный вопрос



нерешенные
звонки



все звонки
с негативной
тональностью



$$\text{SCORE}(P) = 100 \cdot \sum w_k m_k$$

Efficiency (30 баллов): Оптимальность маршрута



Dependencies (10)



Redundancy (10)



Tool-Usage Completeness (10)

«Компонент 5: корректность зависимостей»

Зависимости между шагами заданы правильно



Компонент 6: отсутствие избыточности

Нет лишних и повторяющихся действий



повторное применение
уже учтенного фильтра

«Компонент 7: полнота использования инструментов»

Использованы все необходимые источники и инструменты



«Как из компонентов получается общая оценка»



«Единая оценка всего плана: точность, полнота и F1»



Оценка планов ИИ-агентов: Две ключевые метрики



ТОЧНОСТЬ (PRECISION)

«Какая доля из предложенного плана действительно правильная?»

Показывает, насколько шаги, сгенерированные ИИ, соответствуют эталонному (идеальному) плану.

Высокая точность может маскировать то, что ИИ предложил слишком мало шагов, боясь ошибиться.



ПОЛНОТА (RECALL)

«Какую часть обязательных шагов модель смогла покрыть?»

Показывает, сколько необходимых (критических) действий из эталонного плана ИИ **не пропустил**.

Высокая полнота может маскировать то, что ИИ добавил много лишнего "мусора" (галлюцинаций), лишь бы ничего не забыть.

F1-мера: Балансировка качества планирования

Защитный механизм, наказывающий ИИ за экстремальные перекосы в логике.

$$F1 = \frac{(2 \times \text{Точность} \times \text{Полнота})}{(\text{Точность} + \text{Полнота})}$$

Практический пример расчета агента:

- **Точность: 0.8** (80% предложенных шагов верны)
- **Полнота: 0.5** (Но ИИ нашел лишь 50% от всех нужных шагов)

$$\text{F1-мера: } \frac{2 \times 0.8 \times 0.5}{0.8 + 0.5} = \frac{0.8}{1.3} \approx \mathbf{0.62}$$

Вывод: Несмотря на высокую Точность (0.8), итоговая оценка снижена (0.62) из-за просадки в Полноте.

Прагматичный смысл метрики

F1-мера не позволяет ИИ-агенту выглядеть «хорошим» за счет чирерства по одному показателю. Она объективно наказывает за два сценария:

[Сценарий А]

ИИ предложил мало шагов, но все правильные (Точность ↑, Полнота ↓).

[Сценарий Б]

ИИ не упустил ничего, но добавил массу лишних, ненужных действий (Полнота ↑, Точность ↓).

Результат: F1 показывает реальное, сбалансированное качество мышления агента.

«One-shot judge»

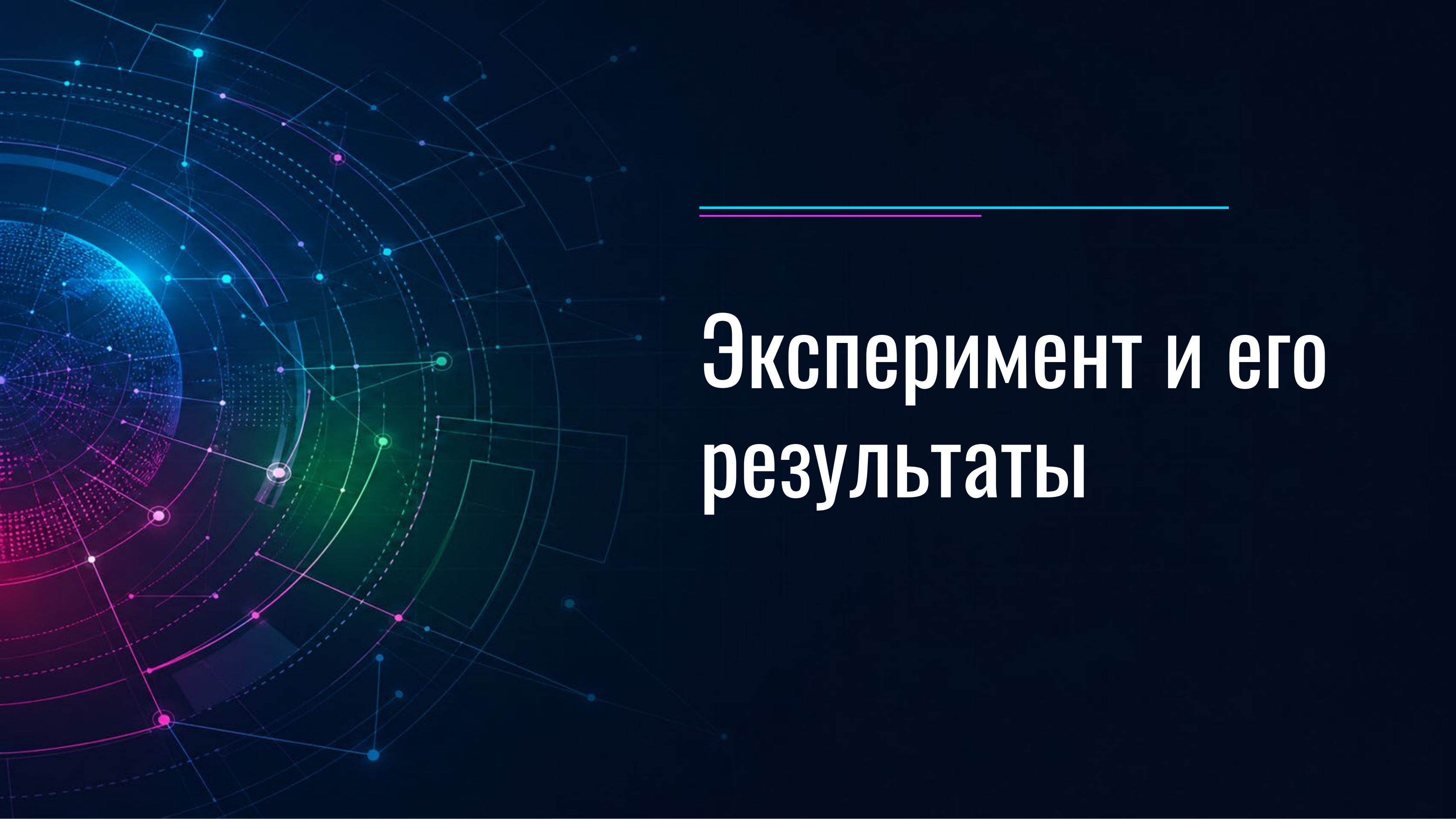


Метрики:
Precision,
Recall, F1
на уровне
шагов

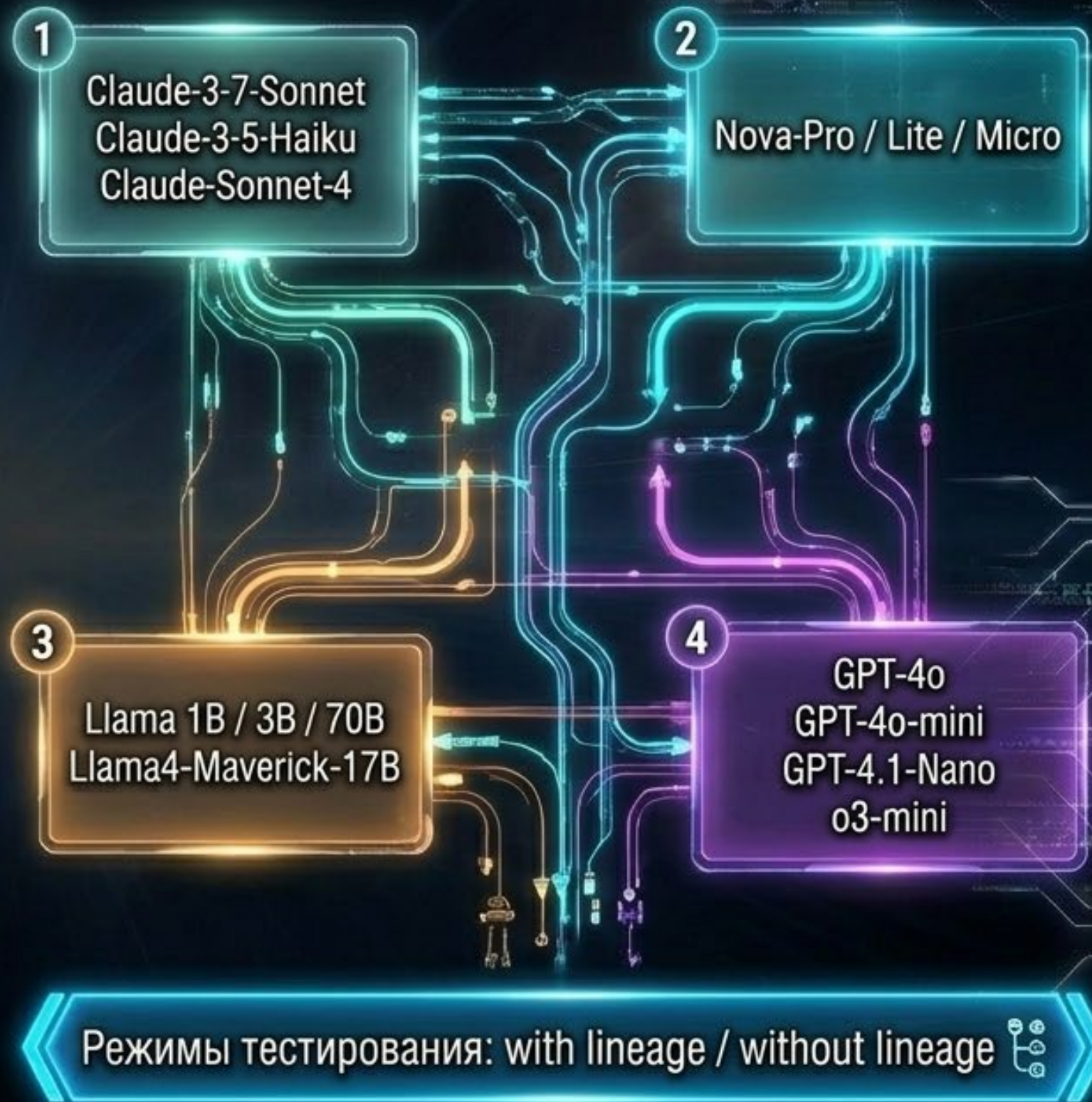
«Треугольник контроля»



Не парсится JSON → Extremely Bad

The background features a dark blue gradient. On the left side, there is a complex, glowing digital structure. It includes a sphere composed of many small, bright blue and purple dots, with lines radiating from it. Overlaid on this are several concentric, semi-transparent circular arcs in shades of blue and green. A network of thin, light blue lines connects various points across the scene, some of which are highlighted with larger, glowing dots in blue, green, and pink. The overall aesthetic is futuristic and technological.

Эксперимент и его результаты



«Как был устроен эксперимент»

600

запросов

14

моделей

2

режима подсказки

«Результат 1: даже сильные модели не дают стабильного эталона»



49,75% A +

лучший результат по единой оценке.

**«Результат 2: по компонентам лучше,
но ошибки остаются»**

84,8

лучший общий
результат по
покомпонентной
оценке.



Lineage prompting effect



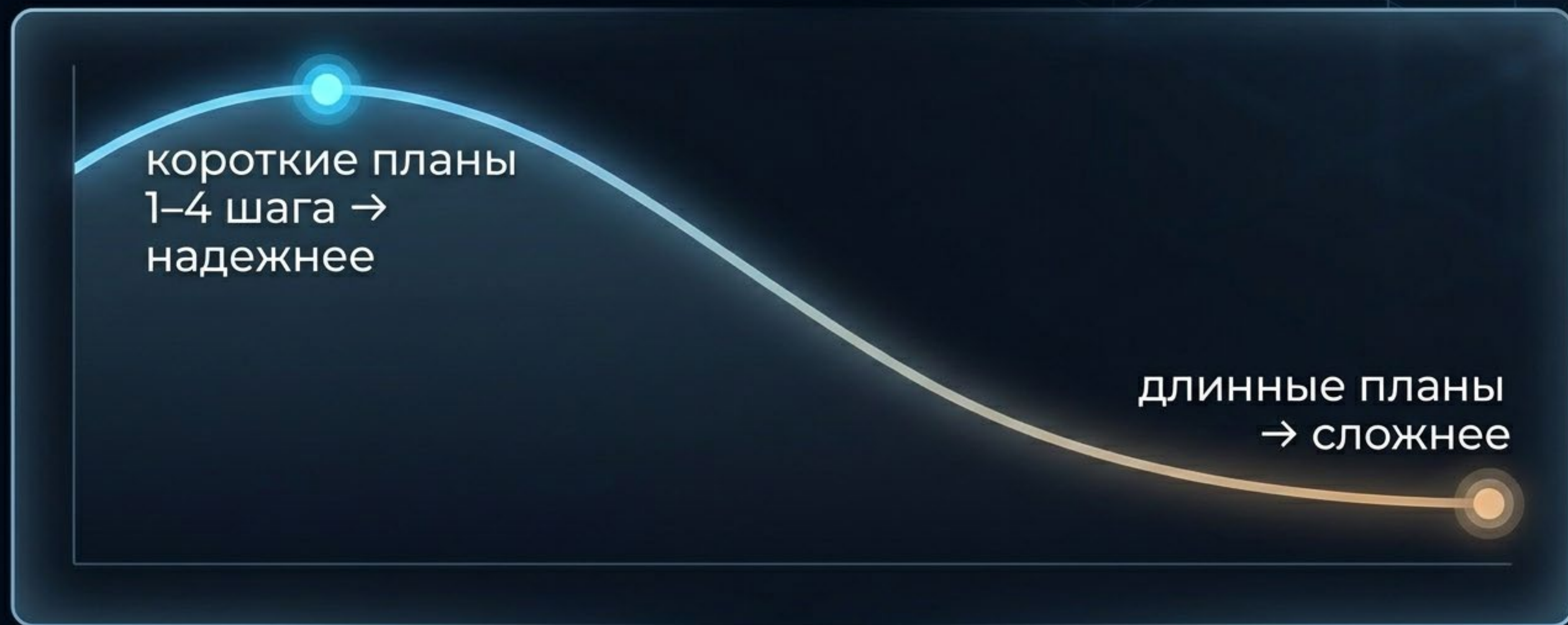
Step Executability улучшилась у 9 из 14,
но общий балл часто выше без lineage
(on 14 models)

Самые слабые метрики

Metric	with lineage vs without lineage	
Tool Usage Completeness ⚠	48,74	71,12
Tool-Prompt Alignment ⚠	64,36	69,10
Format ⚠	70,43	74,26

Note: Lower scores indicate weaker performance.
Colors denote critical failure zones.

«Результат 3: длинные и составные планы сложнее»



«Результат 4: история планов помогает не всегда»

6

6 моделей
улучшились

5

5 моделей
ухудшились

3

без
изменений

«Главный вывод эксперимента»

**«Качество планирования — отдельная задача.
Она еще не решена автоматически»**

Snowflake/T2S:
цифры, статусы, QA,
длительность,
флаги

Запросы,
требующие
обоих
источников

RAG/Transcripts:
смысл разговора,
формулировки,
причины, sentiment

Tool-Prompt Alignment



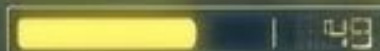
13.00

Format



13.00

Step Executability



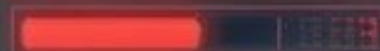
11.49

Query Adherence



13.40

Dependencies



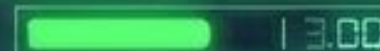
13.00

Redundancy



13.00

Tool-Usage Completeness



13.00

**AI-агенту нужен не только prompt,
а проверяемая операционная модель**

Исполнимый
план

Эталонные
планы

Lineage
исправлений

Метрики
допуска

Убедительный ответ VS Проверяемый маршрут

«Практический вывод: начинать с контроля планов»



карта инструментов



эталонные запросы



оценка планов



короткие задачи



**история
исправлений**

«Что можно внедрить у себя»

**30–50
вопросов**

Этап 1

**эталонные
планы**

Этап 2

**чек-лист
оценки**

Этап 3

**07
пилот**

Финальный этап

«Где нужен осторожный подход»

**сложные
запросы**

**длинные
планы**

**управленческие
решения
по людям**

Качественная ИИ-аналитика начинается с качественного плана

