

ДОКЛАД ДЛЯ РУКОВОДИТЕЛЕЙ КОНТАКТ-ЦЕНТРОВ

Как выбирать лучшую модель обслуживания

ИИ-оценка против человеческого контроля:
где возникает риск уверенно неправильного решения

Источник:

Designing Service Systems from Textual Evidence

Ruicheng Ao, Hongyu Chen, Siyang Gao, Hanwei Li, David Simchi-Levi

<https://arxiv.org/abs/2603.10400>

1. УПРАВЛЕНЧЕСКАЯ ПРОБЛЕМА

**Пилот завершен.
Данных много.
Но можно ли
им доверять?**

Решение очевидно?

Чат-бот А

87%

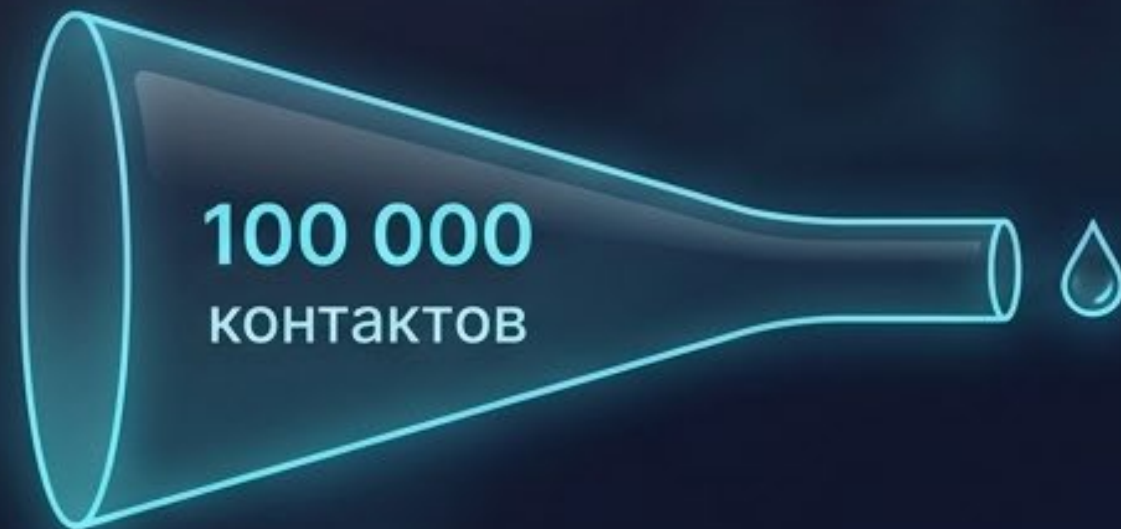


Чат-бот Б

82%

**А если оценщик
систематически смещен?**

Раньше



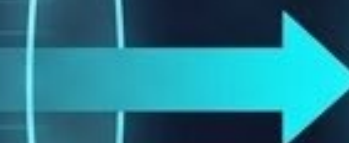
1–3% ручной
проверки



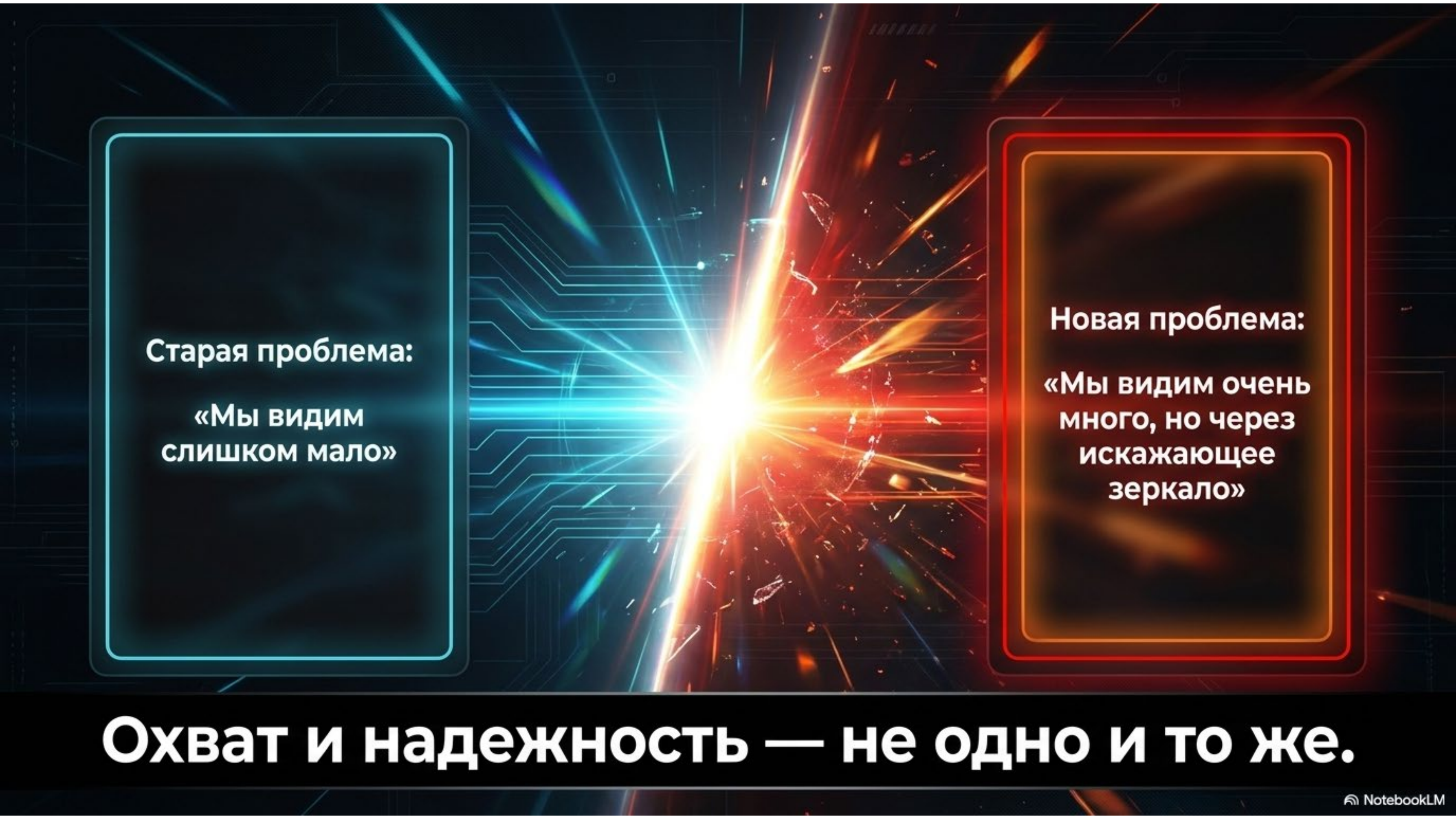
Управленческий
вывод

Больше охват = лучше решение?

Мечта



Надежный
вывод



Старая проблема:

**«Мы видим
слишком мало»**

Новая проблема:

**«Мы видим очень
много, но через
искажающее
зеркало»**

Охват и надежность — не одно и то же.

Проблема: Судья тоже имеет вкус

Многословие:

Длинные ответы получают бонус, краткие и точные — штраф.



Стиль:

Консервативные и шаблонные формулировки получают бонус.

Скрытый урон:

Быстрые переводы на оператора кажутся эффективными, но перегружают систему.

Configuration-dependent bias

Риск оптимизировать вкус измерителя, а не реальное качество для клиента.

Мы вступаем в эпоху 100% контроля качества.

**И именно поэтому
возникает риск стопроцентно
уверенной ошибки.**

2. ПОЧЕМУ ИИ-ОЦЕНКА МОЖЕТ
УВЕРЕННО ОШИБАТЬСЯ

**Проблема не в
отдельных
ошибках, а в
систематическом
смещении.**

Случайная ошибка



Ошибки в разные стороны →
частично сглаживаются объемом.

Систематическое смещение



Ошибки в одну сторону → объем
усиливает уверенность в ошибке.

Больше данных не лечит смещение.

Пример 1. Многословность против эффективности

Бот А: Реальность бизнеса



Коротко, точно, быстро
ведет к решению.

Бот Б: ИИ-Оценка



Длинно, вежливо, подробно,
но перегружает клиента.

Что оценит выше ИИ? И что лучше для бизнеса?

Пример 2. Иллюзия сервиса

Оценка ИИ: В тексте выглядит как активное обслуживание.

Клиент → Бот → Передача → Еще передача → Решение

Реальность бизнеса: В операционной модели означает лишнюю нагрузку.

ГЛАВНАЯ ЛОВУШКА МАССОВОЙ ОЦЕНКИ

$$\left[\begin{array}{c} 100\% \\ \text{контактов} \end{array} \right] \times \left[\begin{array}{c} \text{Смещенный} \\ \text{оценщик} \end{array} \right] = \left[\begin{array}{c} \text{Уверенное} \\ \text{неправильное} \\ \text{решение} \end{array} \right]$$

Много данных

Красивая аналитика

Неверный выбор

Источники данных:

- Транскрипты звонков
- Чаты и тикеты
- Жалобы клиентов
- Медицинские заметки
- Отчеты аудиторов



Удобно для понимания человеком

Богатство смысла, нюансы эмоций, точный контекст.



Неудобно для классической оптимизации

Не скалярно, сложно агрегировать, нельзя напрямую подать в математический оптимизатор.

Человеческий аудит точнее, но экономически недоступен для 100% потока



Вывод: Проверять всё людьми — невозможно. Аудит должен быть умным и выборочным.

Два формальных математических отказа

Отказ машинного подхода

Только LLM-прокси (Proxy-only)



Две разные конфигурации могут выглядеть одинаково по оценке F из-за смещения.

Риск ошибки выбора $\geq 1/2$

Отказ человеческого подхода

Наивный выборочный аудит



Если проверять только спорные случаи, выборка становится нерепрезентативной (зависит от X и F).

Ошибка не исчезает даже при бесконечных данных

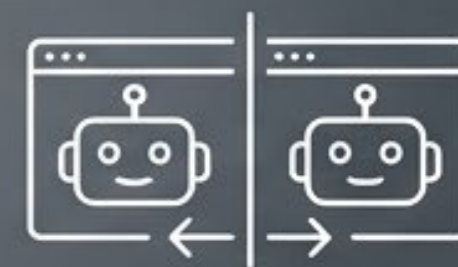
Конфигурации обслуживания: что мы сравниваем?



Чат-бот А vs Чат-бот Б



Старая маршрутизация
vs Новая



Копилот 1 vs Копилот 2



Разные сценарии
оператора



Разные правила
эскалации



Разные процедуры QA

Статья исследует механику выбора между альтернативами — как доказать, что вариант «А» объективно лучше варианта «Б».

Качество спрятано в текстах

Расшифровки звонков

Чаты с ботом

Жалобы клиентов

Комментарии QA

История решения

Текстовые
свидетельства
качества

Если качество проявляется в текстах, классические скалярные методы оптимизации становятся недостаточными.

Парадокс человеческого аудита: эксперт не гарантирует объективность

ИИ оценивает весь поток.

Дешёвый масштаб, но встроенные алгоритмические искажения (предвзятость модели).

Человек надежнее как эксперт, НО...
проверяет не случайную выборку.
Мы направляем аудиторов только на проблемные зоны:

Жалобы

VIP-клиенты

Низкие оценки

Сложные кейсы

Сравнение средних оценок такой искаженной выборки дает смещенный вывод.
Искусственно сфокусированный аудит ломает математику A/B-теста.

РАЗДЕЛ 04

Гибридная методика оценки

ИИ дает беспрецедентный масштаб.
Человек прицельно калибрует систему.

Главное разложение качества

A 3D rendering of the equation $\theta_k = E[F] + E[Y - F]$ set against a dark background with glowing blue and orange circuit-like patterns. The symbols are rendered in a 3D, blocky font. θ_k is white, $=$ is white, $E[F]$ is blue, $+$ is blue, $E[Y - F]$ is orange. The blue and orange elements have glowing trails that lead down to their respective definitions in the boxes below.

Истинное качество

Проверенное качество сервисной конфигурации.

Дешевое среднее (LLM)

Массовый сигнал от LLM-судьи. Доступен всегда, но содержит смещение.

Дорогая поправка (Аудит)

Остаток, исправляющий систематическое расхождение машины и человека.

Интеллектуальный цикл оценки: от данных к калибровке



Отказ от фиксированных квот

~~Проверять ровно 5% контактов~~

~~Проверять только низкие оценки~~

~~Проверять только спорные кейсы~~

**Вероятность проверки
рассчитывается
динамически.**

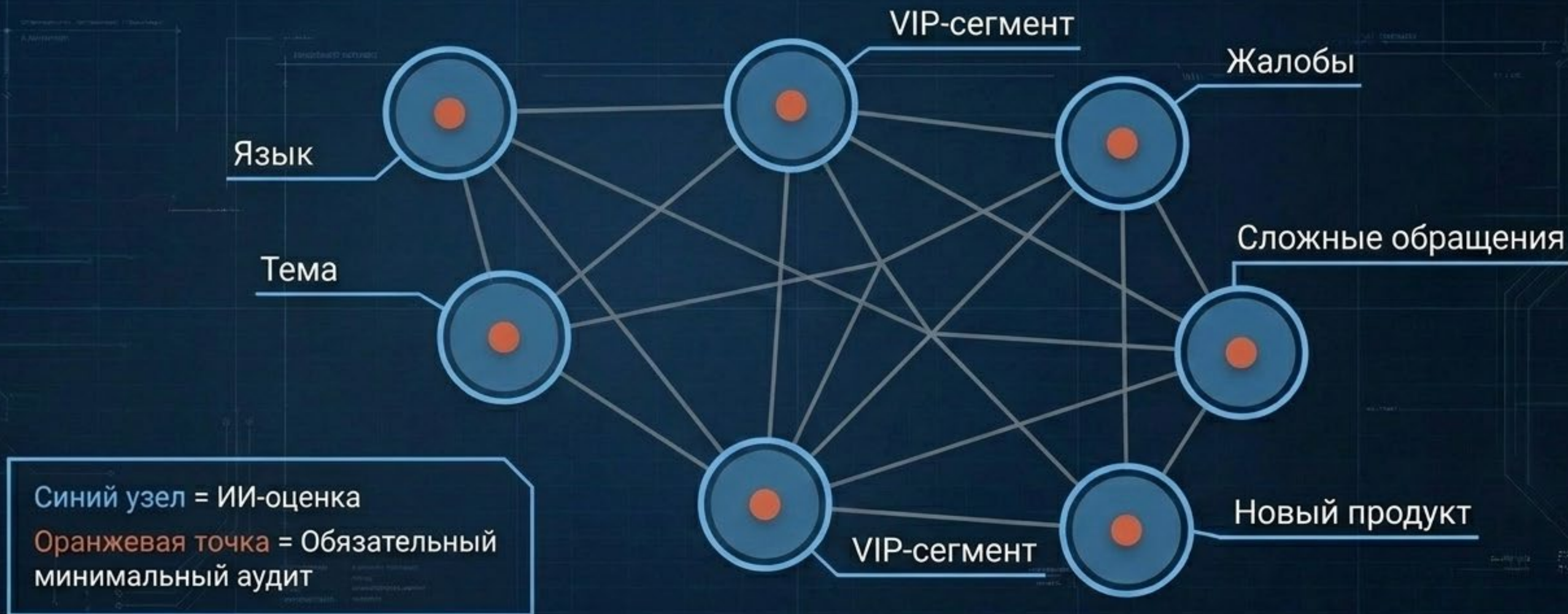
Система направляет аудит только туда, где мнение человека способно изменить финальное решение о выборе конфигурации (inverse-propensity weighting).

Радар человеческого внимания: куда направлять дорогой ресурс?



«Проверять нужно не “плохое”, а ненадежно оцененное автоматикой».

Принцип «отсутствия слепых зон»



Не должно существовать зон, полностью отданных ИИ без регулярного фонового контроля человеком. Алгоритм поддерживает минимальную вероятность проверки ($\pi_{\min}$) для каждого сегмента.

Умный финиш: логика остановки пилота

Старый подход:

Время вышло (1 месяц)



Проверили N контактов
→ Сделали вывод.

PP-LUCB алгоритм:



Сбор данных
→ Оценка статистической
значимости
→ Остановка, когда лидер
надежно отделен
от конкурентов.

Цель — не выполнить план по количеству проверок, а принять надежное решение.

РАЗДЕЛ 05

Результаты исследования

Обнадеживающие данные, но не волшебная таблетка: где подход работает блестяще, а где имеет строгие математические границы.



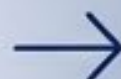
ПРИОРИТЕТЫ

ГЛАВНОЕ — СЕЙЧАС



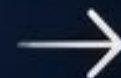
РЕШЕНИЯ

ЧЕТКИЕ ШАГИ



РЕЗУЛЬТАТЫ

В ФОКУСЕ



Триумф методики: Кейс классификации тикетов

40 из 40

Лучшая конфигурация выбрана безошибочно во всех испытаниях.

–90%

Сокращение затрат на дорогой человеческий аудит.

Что это НЕ означает:

- Не гарантирует –90% везде:
Цифра зависит от изначального качества ИИ.
- Не означает увольнение людей:
Человек критически важен для калибровки.
- Валидно для четких задач:
Результат получен в условиях строгих критериев классификации.

Матрица ограничений: не все задачи одинаково поддаются алгоритму

Сильный кейс (Тикеты)	Сложный кейс (MT-Bench)
Однозначная цель (категория/приоритет)	Размытые критерии (стиль, эмпатия, тон)
Очевидная разница между моделями	Модели крайне близки по уровню
Сильное сокращение аудита, точный выбор	ИИ-смещение преодолевается только ростом доли человека

«Метод помогает принимать решения лучше.
Но он не делает сложную задачу простой».

Как правильно интерпретировать эти цифры?



НЕ обещание: «Мы сократим QA на 90% в любом проекте и процессе».



НЕ замена: «Искусственный интеллект идеален, человек больше не нужен».



ДА: «Мы получаем математически обоснованный инструмент для умного распределения дорогого ресурса экспертов».

Это мощная аналитическая система, требующая четких критериев и понимания границ применимости.

РАЗДЕЛ 06

Значение для руководителей

Служба контроля качества эволюционирует: от проверки звонков к системе аудита самой алгоритмической оценки.



Роль человека меняется фундаментально

Было

Фабрика по оценке
контактов

Человек слушает звонок →
Человек ставит балл
по чек-листу.
Ресурс тратится на рутину.



Стало

Аудитор и калибратор
AI-системы

Человек находит слепые зоны →
Человек обновляет логику
модели. Ресурс тратится
на обучение системы.

Контролер качества проверяет не только работу оператора или бота.
Он проверяет надежность автоматической оценки.

5 вопросов, которые стоит задать завтра утром (Executive Action Plan)

- ☐ Сравниваем ли мы варианты обслуживания, опираясь лишь на одну среднюю ИИ-оценку?
- ☐ Знаем ли мы точные зоны, где наш ИИ систематически расходится с мнением эксперта?
- ☐ Остались ли у нас «слепые зоны», которые человек вообще перестал проверять?
- ☐ Понимаем ли мы системно, почему конкретный контакт попал в аудит (или это случайность)?
- ☐ Есть ли у нашего QA-пилота строгий критерий остановки, кроме наступления календарной даты?

Грань между наукой и бизнес-практикой

Строгое математическое доказательство алгоритма (PP-LUCB) для выбора лучшей конфигурации на основе текстовых свидетельств.

«Человеческое внимание должно быть хирургически направлено туда, где автоматика наиболее ненадежна».

Готовый фундаментальный каркас для построения гибридного QA нового поколения в любом контакт-центре.

Раньше мы боялись, что видим слишком мало...

**Теперь мы должны бояться,
что видим очень много, но
неправильно**

Масштаб ИИ + Калибровка человеком = Надежное решение.