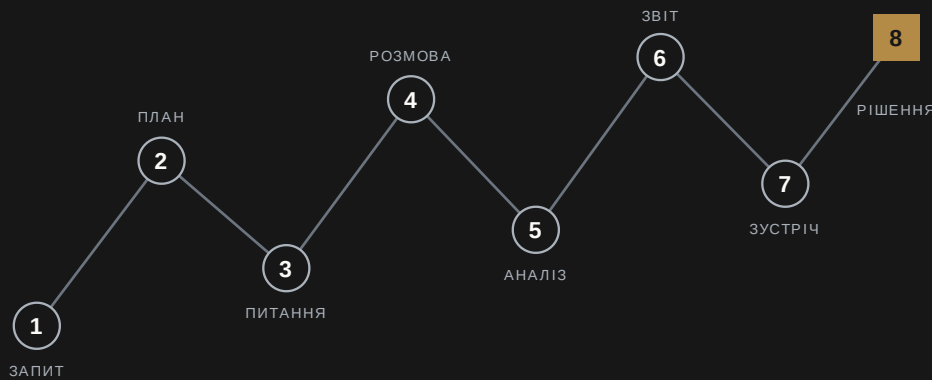




ТАК ТОБІ Й ЗБРЕХАЛИ

Путівник одним дослідженням — від запиту стейкхолдера до зміненого рішення. Вісім етапів, і на кожному на вас чекає своя брехня: від користувачів, від даних, від AI, від бізнесу і від власної голови. Я покажу, де саме, — на своїх дослідженнях і своїх помилках.



АВТОРКА

Ніна Сьомкіна (Nina Somkina) · Founder, Institute for Human Reasoning · 10+ років у дослідженнях · mixed-methods

ФОРМАТ

Research playbook · UA · перше видання, 2026
Don't outsource the thinking that matters.



Ніна Сьомкіна (Nina Somkina)

Founder at Institute for Human Reasoning
10+ years in Research · Mixed-methods ·
Consultant · Mentor · Pre-MBA

Записатися на консультацію

Дослідниця, консультантка і викладачка. Десять років у дослідженнях — і десять років колекціонування брехні, на яку купувалась сама.

Працюю в mixed-methods — там, де одне питання потребує і розмови, і цифр. Спеціалізуюсь на складних дослідженнях: відповідь не на поверхні, джерела суперечать одне одному, рішення дороге. Стратегічні сегментації, churn, запуски на нових ринках, робота майбутнього — для стартапів і великих бізнесів у підписковому B2C, банкінгу та B2B.

Консультую організації, як досліджувати якісно і швидко — і як ухвалювати рішення на даних, які вже є. Бо даних у компаніях зазвичай більше, ніж рішень на них.

Навчаю фаундерів, дизайнерів, продактів, райтерів, маркетологів і менеджерів критичному мисленню і роботі з новими інструментами — зокрема AI у дослідженнях. Викладаю у Projector Institute та Apollo IT School, менторю дослідників-початківців і команди без дослідника в штаті.

Для мене якісний рісерч — не самоціль. Важливо, щоб на даних ухвалили рішення, яке змінює тактику чи стратегію.

Дослідження, після якого нічого не змінилось, — витрачений час і гроші, хай яким гарним був звіт.

Засновниця Institute for Human Reasoning — дослідницької й освітньої організації про людське мислення та рішення у технологічному світі. Цей плейбук — її перший продукт.

Все, що ви прочитаете далі, сталося зі мною насправді. Включно з помилками. Особливо помилки.

Маршрут цього плейбука

Одне дослідження, пройдене з початку до кінця. Якщо у вас зараз живе дослідження — відкривайте одразу етап, на якому стоїте: номер сторінки праворуч. Якщо ще ні — читайте підряд: маршрут і є планом вашого першого пісерчу.

Замість вступу

Який найкращий спосіб убити людину · що таке когнітивні викривлення, Система 1 і Система 2 · ≈180 способів збрехати собі · що я називаю «брехнею» · як читати позначки · три маршрути читання

05

Чесні цифри

4 000 → 28 · 10 000 → 0 → 0 · 260 і 30 · 126 → 0 opens → decisions · 76 → 58 → 7 · 35% → 7% · ⅓ «активних» · 2 тижні → 1 година · 30–40%

09

1 Запит

Переклад запиту · тест на замовника · три типи досліджень і чесні строки · roadmap і пріоритизація 42 запитів

11

2 План і люди

Два світи рекрутингу · скільки людей достатньо (без магії) · AI-агенти як масштаб · демократизація і continuous discovery

16

3 Питання

Опитувальник, глибинне інтерв'ю, юзабіліті-тест: три способи замовити брехню — і як їх ловити

23

4 Розмова

Три причини брехні без наміру · CPO і брехун · «учителька в класі» · кому легше зізнатись AI · команда інтерв'юерів

29

5 Аналіз

Ритуал перевірки AI · як бреше AI без контексту · японські прокляття, яких не існує · сегмент і персона · «дорого»

34

6 Звіт

0 opens ≠ 0 impact · полицки evidence · evidence ladder · формула інсайту · вирізки, скріншоти й емпатія · формати

42

7 Зустріч

«Ми і так знали» · коли шок C-level — це діагноз досліднику · пауза 2–4 дні · техніка Сократа · воркшоп

48

8 Рішення

Малі сигнали · decision check · пахуйте рішення, а не інтерв'ю · кейс про AI, який не додав цінності · фінал bank-case

52

AI на маршруті

57

Discover vs Scale · хто за що відповідає · 4 помилки AI і промпти проти них · ChatGPT vs Claude vs Gemini

Quality check

Rigor і impact · 6 вимірів якості · надійність ≠ валідність · триангуляція · insight hub vs репозиторій

61

Замість висновку

Чесні цифри з висновками · що почитати · чим користуватись · пак шаблонів · словничок · джерела

68

НАСКРІЗНИЙ КЕЙС

Через усі вісім етапів іде одне реальне дослідження — **bank-case**: преміальний сегмент банку, замовник — CEO, 4 000 контактів, 28 інтерв'ю. Шукайте блоки із золотою лінією: там видно, де дослідження перебуває на кожному етапі і який у нього decision state. Кейс анонімізовано — банк не називаємо; цифри і перебіг реальні.

Повний зміст

Усі підрозділи маршруту — для повернень: коли книга вже прочитана і потрібне конкретне місце.

Вступна частина	02
Хто веде цей маршрут	02
Маршрут цього плейбука	03
Повний зміст	04
Який найкращий спосіб убити людину?	05
Що я називаю «брехнею»	05
Система 1, Система 2 і викривлення	06
≈180 способів збрехати собі	07
Як читати позначки · три маршрути читання	07
Чесні цифри	09
Мала маршруту	10
Етап 1 · Запит	11
1.1 «Дізнайся, чому вони йдуть» — це не питання	12
1.2 Визнач тип — і не бреши собі про строки	13
Етап 2 · План і люди	16
2.1 Два світи: масовий B2C і преміальний банкінг	17
2.2 AI-агенти: коли розмов потрібно багато	18
2.3 Ваша база теж бреше	20
2.4 Один дослідник не тягне — і не мусить	20
Етап 3 · Питання	23
3.0 Три методи — три способи замовити брехню	24
3.1 Глибинне інтерв'ю: сім питань	25
3.2 Опитувальник: брехня у шкалі й порядку	26
3.3 Юзабіліті-тестування: репліки-підказки	27
Етап 4 · Розмова	29
4.1 Три причини брехні без наміру	30
4.2 Кому легше зізнатись AI, ніж людині	31
4.3 Якщо інтерв'юерів кілька — вирівняй руку	31
4.4 Етика розмови: згода, запис, дані	32
Етап 5 · Аналіз	34
5.1 Перша відповідь AI — це не аналіз	35
5.2 Три японські прокляття	36
5.3 Сегмент і персона: різні інструменти	38
5.4 Сегментація — судження. Але не будь-яке	39
«Дорого» — перекладайте, не читайте буквально	40

Етап 6 · Звіт	42
6.1 126 → 0 opens ≠ 0 impact	43
6.2 Полички, які бізнес розрізняє	44
6.3 Що таке інсайт — і формула, яку бізнес чує	45
6.4 Цитата, яку впізнав керівник	46
6.5 Формати, які працюють замість 120 слайдів	46
Етап 7 · Зустріч	48
7.1 «Це ж очевидно» — не атака	49
7.2 Одне дослідження — три презентації	50
Етап 8 · Рішення	52
8.1 Малі сигнали: дурять — або лебеді	53
8.2 Що чесно можна виміряти — і що ні	54
Шар · AI на маршруті	57
Робоча модель: human → AI → human → quant	58
Вісім етапів — хто що робить	58
Чотири помилки AI — і промпт проти кожної	59
ChatGPT, Claude, Gemini: хто з них хто	60
Шар · Quality check	61
Як якість оцінюють насправді	62
Надійність ≠ валідність	64
Триангуляція: найдешевший спосіб підняти валідність	65
Insight hub і репозиторій — не одне й те саме	66
Замість висновку	68
Чесні цифри — ще раз. Тепер із висновками	68
Вам брехатимуть далі	69
Дванадцять книжок	70
Статті, гайди і спільноти	71
Інструменти за циклом дослідження	72
Пак шаблонів	75
Додатки	76
А Словничок: терміни маршруту	76
Б Джерела	78

Який найкращий спосіб убити людину?

Я ставлю це питання на початку кожного потоку. Аудиторія — дорослі розумні люди: продакти, дизайнери, маркетологи. За дві хвилини у чаті — десятки впевнених варіантів. Із деталями. З аргументацією.

Потім я ставлю друге питання: «А хто з вас хоч раз убивав?»

Тиша.

Ніхто не має досвіду. Але всі щойно впевнено збрехали — собі й мені, без жодного злого наміру: мозок просто вміє генерувати переконливі відповіді про речі, яких ніколи не проживав. Це і є середовище, в якому працює дослідник.

Ваш стейкхолдер упевнено принесе вам зламане питання. Респонденти впевнено розкажуть про майбутнє, якого не буде. Панель підсуне людей не з тієї країни. AI впевнено напише висновок, якого немає в даних. А ваш мозок тихенько підбере докази під те, у що вам зручно вірити. І все це — на одному-єдиному дослідженні.

Дослідження — це не спосіб дістати істину. Це спосіб зловити брехню раніше, ніж вона коштуватиме грошей.

І ще одне, перш ніж рушимо. Цей плейбук — перший продукт Institute for Human Reasoning, і його головна думка вміщається в один рядок: **don't outsource the thinking that matters**. Відповіді сьогодні дешеві — їх дають аналітика, панелі, респонденти, AI. Дорогим лишається одне: судження людини, яка ці відповіді перевіряє і бере за них відповідальність. Вісім етапів попереду — тренування саме цієї здатності.

Що я називаю «брехнею»

Одразу домовимось про слово, бо воно в назві. «Брехня» тут — будь-яка переконлива відповідь, цифра чи висновок, яким не варто вірити буквально. Іноді це свідомий обман — як у респондента, що місяцями професійно дурив сервіс (етап 4). Частіше — ні: ввічливість, страх виглядати дурним, питання про уявне майбутнє, дашборд, що спізнюється на місяць, перекошена вибірка, галюцинація AI, мозок, що побудував відсутню відповідь. У назвах етапів лишаю «брехню» — вона тримає маршрут. У середині називатиму механізм точніше, бо лікується він по-різному.

Як побудований маршрут

Не за методами — за маршрутом. Одне дослідження від початку до кінця: запит → план → питання → розмова → аналіз → звіт → зустріч → рішення. На кожному етапі — брехня, яка там живе; історія, де я на неї купилась; викривлення, яке її обслуговує в голові; інструмент, який можна забрати завтра.

І одне наскрізне дослідження. Через усі вісім етапів пройде мій реальний кейс: банк, преміальний сегмент, просадка, яку бачив CEO і не могли пояснити менеджери. 4 000 контактів, 28 інтерв'ю, звіт, воркшоп — і чесний фінал. У кожному етапі буде блок «Bank-case: де ми зараз».

Але перш ніж рушити — три поняття на наступних двох сторінках. Без них незрозуміло, що таке «викривлення етапу».

Чому мозок бреше сам собі: Система 1, Система 2 і викривлення

У «Thinking, Fast and Slow» Деніел Канеман описав два режими мислення (Kahneman, 2011). Поділ спрощений, але для дослідника — робочий: він пояснює, звідки береться впевнена відповідь на питання, на яке немає даних.

СИСТЕМА 1 • ШВИДКА

Автоматична, миттєва, без зусиль. Впізнає обличчя, добудовує «2 + 2», відчуває, що співрозмовник роздратований, — і **генерує готову відповідь на будь-яке питання**, навіть на те, на яке у неї немає даних. Саме вона за дві хвилини написала в чат «найкращий спосіб убити». Вона не вміє сказати «не знаю».

Приклад. «Бита і м'яч коштують 1,10. Бита дорожча за м'яч на 1,00. Скільки коштує м'яч?» Система 1 миттєво каже «10 копійок» — і помиляється: правильна відповідь 5. Більшість людей, включно з дослідниками, не перевіряє.

СИСТЕМА 2 • ПОВІЛЬНА

Свідома, лінива, енергозатратна. Перемножує 17 × 24, порівнює два тарифи, перевіряє, чи є цитата в транскрипті. Вмикається тільки коли її змусити — і **охоче приймає те, що Система 1 уже вирішила**, дописуючи логічне пояснення заднім числом.

Приклад. Стейкхолдер «знає», чому користувачі йдуть (Система 1). Попросіть його пояснити по кроках, як саме людина доходить до відходу, — увімкнеться Система 2, і на третьому кроці пояснення зупиниться. Там і живе питання дослідження.

Що таке когнітивне викривлення

Когнітивне викривлення — не помилка і не дурість. Це **систематичний спосіб, у який Система 1 зрізає кут**: підставляє легше питання замість складного, бере найдоступніший приклад за типовий, шукає підтвердження замість спростування, тримається за те, у що вже вклалася. Головне слово — **систематичний**: та сама помилка повторюється у різних людей у схожих умовах. Отже, її можна передбачити — і закласти запобіжник у дизайн дослідження. Це і робить кожен блок «Викривлення етапу» далі.

У СТЕЙКХОЛДЕРА

«Знаю, чому вони йдуть, — онбординг задовгий». Система 1 підставила доступну картинку замість відповіді на складне питання. Це **ілюзія пояснювальної глибини** (етап 1).

Антидот: «розкажи по кроках».

У РЕСПОНДЕНТА

«Купили б за 9,99? — Так, звісно». Уявне майбутнє мозок не вміє прожити, тому вигадує правдоподібне. **Якір + гіпотетика** (етап 3).

Антидот: питати про останній реальний епізод.

У ВАС

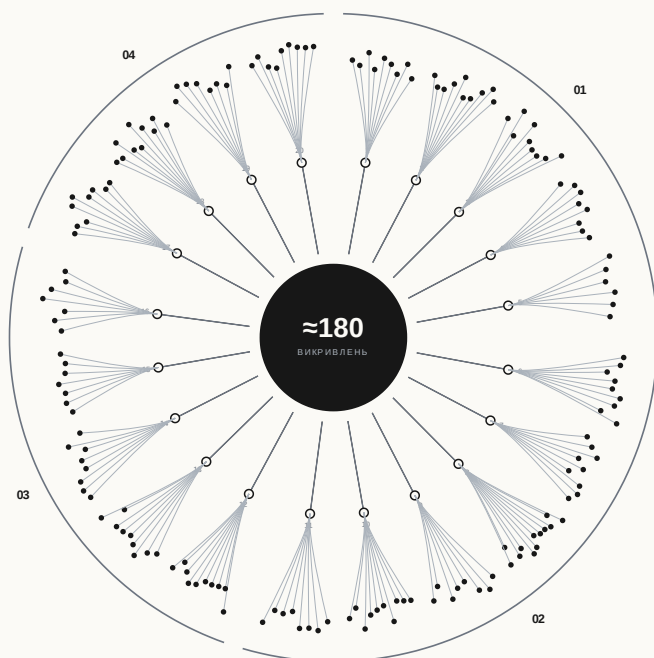
Із 28 транскриптів у звіт «самі» потрапляють цитати під вашу гіпотезу. Ви їх не обирали — вони здались «по темі».

Підтверджувальне викривлення (етап 5).

Антидот: «що в даних суперечить мені?»

Респондент, стейкхолдер, AI і ви — усі працюють на тій самій Системі 1 (AI — на статистичній копії мільйонів таких систем). Тому брехня в дослідженні майже ніколи не намір. Це побічний продукт мислення. І проти неї працює не «будьте уважніші», а процес: правило в гайді, чекліст, ритуал перевірки.

≈180 способів збрехати собі — і всього чотири причини



Кожна крапка — одне описане когнітивне викривлення; кожен вузол — механізм, який породжує кілька. Схема за структурою Cognitive Bias Codex (Benson, 2016; візуалізація — Manoogian, 2016): ~175–188 викривлень, згрупованих за чотирма проблемами, які мозок так «розв'язує».

01 · Забагато інформації

Мозок фільтрує — і фільтр упереджений.

- помічаємо те, що вже є в пам'яті або повторюється
- дивне і кумедне впадає в око сильніше
- помічаємо зміни, а не стани
- тягнемось до того, що підтверджує наші переконання
- краще бачимо чужі помилки, ніж свої

02 · Замало сенсу

Мозок добуває пропуски — правдоподібним.

- знаходимо історії і патерни навіть у рідких даних
- заповнюємо прогалини стереотипами
- спрощуємо ймовірності й числа
- «знаємо», що думають інші
- проєктуємо теперішнє на минуле і майбутнє
- цінуємо знайоме більше за нове

03 · Треба діяти швидко

Мозок обирає впевненість замість точності.

- переоцінюємо власну значущість і вплив
- тримаємось того, у що вже вклалися
- обираємо близьке і завершене замість відкладеного
- захищаємо автономію і статус
- віддаємо перевагу простим варіантам

04 · Що запам'ятати

Мозок редагує спогади — і не повідомляє.

- переписуємо спогади заднім числом
- відкидаємо деталі, лишаємо узагальнення
- зводимо події до ключових елементів
- запам'ятовуємо залежно від того, як переживали

Досить бачити чотири проблеми — і ставити перед кожним висновком одне питання: **«яку з чотирьох мій мозок щойно розв'язав за мене?»** Забагато даних — я щось відфільтрувала. Замало сенсу — щось добувала. Треба швидко — обрала впевнене замість точного. Пам'ять — переказ переказу. У плейбуку далі — два десятки конкретних викривлень: ті, які я ловлю в себе найчастіше.

Як читати позначки

MY CASE так сталося в моєму дослідженні — спостереження, не закон.

RULE OF THUMB практичний орієнтир, який у мене часто працює.

RESEARCH-BACKED для цього є джерело, і я його називаю.

CHECK IN YOUR CONTEXT перевіряється на вашій аудиторії, перш ніж на це спиратись.

Bank-case — де зараз наскрізне дослідження.

Інструмент — вирвати й використати завтра як є.

Hint — коротка практична порада з поля.

Викривлення етапу — що у вашій голові робить брехню переконливою.

Заберіть із собою — підсумок етапу.

Три маршрути читання

Книга задумана як маршрут, але читати її можна трьома способами — залежно від того, скільки у вас часу і скільки досліджень за плечима.

Якщо у вас 30 хвилин. Прочитайте «Чесні цифри» на наступній сторінці, відкрийте етап, на якому зараз стоїть ваше живе дослідження, і заберіть його блок «Заберіть із собою». Останні пів години, що лишаться колись потім, — на Quality check: шість вимірів і рубрика 0·1·2.

Якщо це ваше перше дослідження. Читайте підряд, етап за етапом: маршрут і є планом вашого першого рісерчу. Блоки «Інструмент» виривайте і використовуйте як є, research-backed джерела лишіть на друге коло — вони нікуди не дінуться.

Якщо у вас за плечима десятки досліджень. Ідіть за позначками: my case і «Чесні цифри» — щоб звірити з власним досвідом; Quality check і AI-шар — цілком; словничок і джерела — як довідник для команди. І сперечайтесь зі мною на полях: половина цієї книги народилася саме так.

Те, що вам ніколи не покаже жоден вендор

Кожна цифра — з мого реального дослідження, і кожна розібрана всередині плейбука. Де категорію чи цифри змінено для анонімізації — це прямо написано в паспорті картки, зі збереженням пропорцій і висновків. Під кожною — мінімальний паспорт, щоб ви розуміли, що саме порівнюється, без NDA: категорія, метод, N, рішення, обмеження.

RECRUITMENT

4 000 → 28

Преміум-клієнти банку, що пішли за три місяці. Сапорт особисто зв'язався з кожним. 28 інтерв'ю — і подарунок із листом за підписом CEO, який важив більше за подарунок.

Наскрізний кейс · банк, преміальний сегмент · замовник CEO · глибинні інтерв'ю · N ≈ 4 000 контактів → 28 інтерв'ю · Decision impact: команда змінила framing проблеми і пріоритизацію причин відходу · Implementation: лишилась у клієнта · Business outcome: недоступний після handoff / NDA

SEGMENTATION

76 → 58 → 7

76 глибинних інтерв'ю, 58 в аналізі, 7 сегментів, підтверджених кількісно. Сегменти склалися на 52-му. Не на 12-му.

B2C · стратегічна сегментація · матриця пристрій × тип контенту × тариф · 3 дослідники · 3 тижні · 18 відсіяно за якістю: «все супер» без м'яса, нетверезий респондент, поганий звук

FAILURE

10 000 → 10 → 0 → 0

Churn-опитувальник на десять тисяч відписаних: десять анкет, нуль інтерв'ю. 3 розіграшем на \$300 — те саме. AI-агент через імейл із контрольною групою — знову нуль.

B2C підписка · churned users · замовник CEO, боард вимагав −5 п.п. відтоку квартал до кварталу · три спроби рекрутингу · Рішення: рекрутинг тих, хто пішов, — окрема проблема · Обмеження: один продукт

DECISION EVIDENCE

35% → 7%

Маркетинг приводив третину аудиторії категорії. Залишалась сьома частина. Мінус 28 пунктів — це дуже багато грошей.

B2C · acquisition-профіль проти JTBD тих, хто лишається · опитувальник + аналітика · Рішення: таргетинг · Обмеження: точний N і джерело denominator не збереглися — читай як порядок величини

SCALE / COST

260 і 30

260 інтерв'ю провів AI-агент за \$200 рекрутингового бюджету — категорійне discovery кавових напоїв, три категорії. Паралельно ~30 human-інтерв'ю з тією ж базою (\$150 рекрутингу) дали ті самі висновки.

Discovery · кавові напої, 3 категорії · реклама → лендинг → AI-agent interview · N = 260, фільтр якості (<10 хв, не по темі) · триангуляція: ~30 human-інтерв'ю · бюджети — лише рекрутинг, без винагород · Обмеження: якісні дані, self-selected через рекламу · Категорії і цифри частково змінено зі збереженням сенсу; пропорції та висновки — реальні

STALE DATA

1/3 «активних»

Третина користувачів, яких аналітика показувала активними, відповіла: «я більше не користуюсь продуктом».

B2C підписка · опитувальник по базі «активних» · Обмеження: третина від тих, хто відповів; N не зберігся

ARTIFACT ADOPTION ≠ IMPACT

126 → 0 opens → decisions

126 інтерв'ю. Master table — нуль відкриттів. Consequential decisions — так. Деталі — NDA. Таблиця провалилась як інтерфейс. Дослідження — ні.

Мій перший великий рісєрч · N = 126 · Artifact adoption: 0 opens · Decision impact: consequential decisions were made · Business outcome / деталі: NDA

ANALYSIS COST

2 тижні → 1 година + валідація

12 інтерв'ю колись означали два тижні аналізу. Тепер — година мого капасіті плюс 4,5 дні команди на валідацію.

Економіка аналізу · N = 12 · чесна повна вартість = 1 год + 4,5 дні

AI CORRECTION · PERSONAL ESTIMATE

30–40%

Скільки я в середньому редагую у «готовому» AI-аналізі, перш ніж він стає правдою.

MY CASE оцінка з досвіду, не заміряна метрика

Вісім етапів — вісім видів брехні — вісім викривлень

Якщо у вас зараз живе дослідження — відкривайте етап, на якому ви стоїте. Якщо ще нема — читайте підряд: маршрут і є планом вашого першого рісерчу. Кожен етап має три постійні блоки: брехню етапу (що ззовні), викривлення етапу (що всередині вашої голови робить цю брехню переконливою) і bank-case (де в цей момент наскрізне дослідження).

ЕТАП	БРЕХНЯ ЕТАПУ — ЩО ЗЗОВНІ	ВИКРИВЛЕННЯ — ЩО ВСЕРЕДИНІ
1	Запит «Запит стейкхолдера — це і є питання дослідження»	Ілюзія пояснювальної глибини · фреймінг
2	План і люди «Знайдемо респондентів за три дні»	Помилка планування · той, хто вижив · самовідбір
3	Питання Половину брехні респондента замовляєте ви самі	Якір · порядок · соціальна бажаність
4	Розмова Людина навпроти бреше з ввічливості — а іноді професійно	Спостерігач · пост-фактум · раціоналізація · яскрава історія
5	Аналіз AI з готовим висновком і ваш мозок, що підбирає докази	Підтверджувальне · автоматизаційне · хибна причинність · плутання джерела
6	Звіт «Report і є результат дослідження»	Прокляття знання · IKEA-ефект
7	Зустріч «Це ж очевидно, ми і так знали»	Заднім числом · статус-кво
8	Рішення «Рісерч сам себе окупить» і «один інсайт — беремо у фічу»	Доступність · sunk cost · неприйняття втрат · дія заради дії

НАСКРІЗНИЙ ШАР · AI НА МАРШРУТІ

Хто за що відповідає на кожному з восьми етапів: що може взяти на себе агент, а що лишається за людиною — бо capability is not authority.

НАСКРІЗНИЙ ШАР · QUALITY CHECK

Як зрозуміти, що вашому дослідженню можна вірити: 6 вимірів якості, рубрика 0.1-2, надійність і валідність людською мовою, триангуляція, insight hub.