

How the Scenario Lab numbers are produced

A plain-language method note for the CD2 campaign. Written 2026-07-24.

The short version: these are **not poll results**. They are the output of a simulation that takes real inputs (the poll, the voter file, the FEC money) and plays the election forward thousands of times under a set of stated assumptions. The value of the tool is the **ranking and the direction** of an effect, not the decimal. Anyone who quotes "+4 points" as a measured fact is misreading it.

1. What the model starts from (all real, all countable)

- a. **Preference** — the co/efficient poll (Jul 14-16 2026, n=624), broken out by media market. This is the only source of "who's ahead."
- b. **Turnout** — the CD2 voter file: 252,554 registered Republicans, per county, with who actually voted in the Aug 2024 primary. This is the only source of "how many ballots, and from where."
- c. **Money** — FEC Q2 filings. This drives how much media weight each side can put behind an attack.

Nothing in the model is invented by an AI. Where a number is a judgment call, it is labeled an assumption and exposed so you can argue with it.

2. How one attack turns into a number

Every attack is scored the same way:

damage to Gross = potency x medium reach x how definable Gross is

- **Potency** — how strong the charge is (0 to 1). The Georgia-Democrat hit is 0.90; out-of-state money is 0.52. Where we have Guidant's informed-ballot numbers, potency is anchored to them.
- **Medium reach** — how much of THIS electorate actually absorbs the message through that channel. Mail and cable score high (the electorate is 53% over 65); mobile digital scores low. This is why switching a mailer to a text collapses its effect.
- **How definable Gross is** — you cannot strip support that is already gone. He is already 27-fav / 35-unfav in Panama City, so there is little left to take there; his Tallahassee-market image is more intact, so the same attack does more there.

That damage is then **spread across the three markets** (Panama City, Leon, rural Tallahassee) in proportion to where Gross's support actually sits and where the attack plays best — a localism attack lands harder rural, a spending attack travels evenly.

3. The part that decides identity vs. hypocrisy: backfire

Every attack also carries a **backfire** — points the attacker (Rogers, or the PAC) loses when the attack reads as unfair. This is the whole ballgame for the values question, and it comes straight from the research:

- A **documented character / record / hypocrisy** attack carries low backfire (~0.4-0.6). Voters don't punish you for pointing out a lie.
- A **personal / identity** attack carries high backfire (~2.2). Voters see cruelty, feel sympathy for the target, and turn on the attacker.

So the identity frame ("he's gay") and the values-hypocrisy frame ("his website says one man and one woman while he's married to a man") differ almost entirely in this one number: **2.2 vs. 0.6**. Same underlying fact, opposite backfire, opposite result.

Two refinements added 2026-07-24:

- When several attacks run together, backfire is set by the **most personal** one, not the average. Running a clean hypocrisy ad alongside a cruel identity ad does not dilute the cruelty — you are still "the campaign that ran the cruel ad." That is why running both nets worse than the hypocrisy frame alone.
- Backfire on a sexuality-based frame is **scaled to the audience**. The backlash literature is general-population; this electorate is not. A low-turnout August GOP primary in the Panhandle / Big Bend is heavily White evangelical Protestant (PRRI: the least same-sex-marriage-supporting group, ~38%), so the sympathy backlash is much weaker here. The model scales the identity/values backfire down accordingly. It is **not zeroed** — a residual remains for earned-media and general-election spillover, which reach beyond the primary audience. This is why the pure-identity penalty is about -2 in this electorate rather than the -6 a general-population model would show.

4. From damage to "net win-odds"

Stripping points off Gross is not the same as gaining them yourself. When Gross loses support, it scatters — most goes to **undecided and the other six candidates**, and Rogers only catches a share. So a 2-point hit on Gross buys you much less than 2 points.

To turn that into the "+4" figure, the model runs the whole election **2,500 to 4,000 times**. Each run it jitters every input within its uncertainty (poll margin of error, how undecideds break, turnout, attack strength +/-40%), plays the eight-way plurality forward, and records who finishes first. "Net win-odds"

is the change in the **share of those runs Rogers finishes first**, versus the same simulation with the attack turned off.

That is why the output is a probability, not a vote count. "+4" means: across thousands of simulated elections, this move lifts Rogers's first-place finishes by about four points. It is a measure of how much the move improves your odds, under the model's assumptions.

4b. Two kinds of uncertainty (and why the race reads as a coin flip)

The simulation carries two different error bands, and the difference matters.

- **Sampling error** is the poll's own margin of error. Each run redraws every candidate's support within that MoE, independently for each market. Because it is independent, it partly averages out across the sixteen counties. This is the uncertainty the poll itself reports.
- **Systematic error** is the possibility that the July 16 poll is simply *biased* -- a bad likely-voter screen, a stale field, nonresponse. No amount of re-sampling the same poll can rule this out, because every run inherits the same starting numbers. So the model adds one **coherent** shock per run, applied identically across all markets (it does NOT average out), sized to a ~3.5-point standard deviation on the Rogers-minus-Gross margin -- mid-range for a congressional primary poll five weeks out.

The projection reports the win probability **both ways**. Sampling-only puts Rogers's lead near 53%. Adding the systematic band moves it to essentially a **coin flip** -- which is the honest read at this distance. It also adds a couple of points of run-to-run jitter to every Scenario Lab figure, which is why those carry ranges and an "≈".

One honest limit: this band does **not** shrink as VBM ballots return. Without ballot-matching, a returned ballot reveals who *turned out*, not who they *voted for*. Only a fresh poll or ballot-level preference data would tighten it.

5. What the numbers are NOT

- **Not measured effects.** They are calibrated to an *informed-ballot* poll question, which reads the charge to a respondent paying full attention. Real voters get it once, in passing. The published meta-analysis on negative advertising (Lau, Sigelman & Rovner 2007) finds real-world effects markedly smaller and less certain. **Treat every attack magnitude as an upper bound.**
- **Not a substitute for testing.** The only way to know a message's real number is to test it. The model tells you where to point the test and what to expect.

- **Not precise to the decimal.** A "+4" and a "+3" are the same finding. A "+4" and a "-6" are a real, directional difference you can act on. Read the sign and the rank; ignore the last digit.
-

6. The one line to remember

The model is a disciplined way to reason about direction and magnitude under stated assumptions. It earns trust by reproducing the pollster's own topline (validation check passes at 0.7pp) before it projects anything forward. Use it to choose between moves and to see the shape of a fight. Do not quote it as a poll.