

CALL THE MAP

CTMM 2.2

How the Forecast Works

Call the Map forecasts every House, Senate and governor race in the 2026 midterms with one model, CTMM 2.2. This is a short account of how it thinks: where a race begins, what moves it, how sure it is, and what it cannot know.

September 2026 · about ten minutes

The companion technical methodology contains every equation, fitted parameter and validation table: [ctmm-2-2-technical-methodology.pdf](#).

What we forecast

On November 3, 2026, Americans elect all 435 members of the House, 35 senators and 36 governors. For each race the model publishes a forecast margin (how far ahead one side is, in points), a chance of winning, and a range showing where the result is likely to land. From the same simulated elections it publishes how many seats each party should expect and the chance each party controls each chamber.

A forecast is not a verdict. A race at 80% is one the favourite loses one time in five, and the model is built so that, over many races, its 80s lose about that often.

Where a race begins

Before any poll is taken, a race already has a shape. How the district or state voted last time, how it voted for president, whether an incumbent is running, whether it is a midterm under the president's party: these facts, and a few others like the age and education of the electorate, are turned into a starting margin by a regression fitted on every election since 2008. It is deliberately dull. It knows nothing about the candidates' fundraising, their ads, or their scandals, and it never looks at a poll. It is the ground the rest of the model stands on.

When district lines change, an estimated prior is labelled as an estimate. The current model retains that fitted estimate rather than artificially stopping it at 40 points. Senate state-history features likewise keep their actual prior-election scale, with coefficients refitted on matching historical inputs.

The national mood

A midterm is partly one election held 506 times. The generic ballot, the poll question asking which party's candidate a voter would choose, is averaged over the last sixty days and compared with what past elections looked like. The gap moves every race at once, by an amount fitted separately for each office from how much its results have historically swung with the nation. At the moment the generic ballot stands at +6.2, and the environment is adding about +1.3 points to a typical House race.

Polls

Where a race has polls, they pull the forecast toward what they say, but never all the way and never blindly. A poll counts for more when it is recent, when its sample is large, and when it is not sponsored by a campaign. Each polling firm counts once, however many times it releases: a firm that polls every week refreshes the picture, it does not multiply its vote. A race with no polls keeps its starting margin; the model does not invent a number.

Polls have also missed in the same direction across whole election years. The model measures that history and shifts its Senate and governor poll averages by about 1.6 and 2.5 points toward Republicans, carrying the uncertainty of that shift with it. House polls are left alone, because shifting them made the forecast worse.

Candidates

Who is on the ballot matters in three ways. Incumbents and their record are part of where a race begins. For governors, the model also counts how much elected experience each nominee brings, the one measure of candidate quality that improved forecasts when tested. And when the ballot is unusual, two Democrats in a California district, an independent in Nebraska, a ranked-choice count in Alaska, the simulation follows the actual rule, with the honest admission that it has little history to fit those cases on and so keeps a stated range of assumptions rather than one number.

Independent of expert ratings

The forecast does not use expert ratings. It learns from elections, polls and the other inputs described here, without pulling a race toward another publisher's judgment. Historical calibration and uncertainty have been refitted without those ratings too.

How sure the model is

A separate Senate correction addresses a pattern seen in historical tests: safe-seat margins were pulled too close. Its effect grows smoothly in lopsided races and can offset inward polling drag. Its strength is learned from past elections, not chosen to match another forecast. Selection explicitly accepts a small increase in close-race error in exchange for lower safe-seat error; a better central estimate is not a guarantee of better uncertainty ranges.

Most forecast misses are ordinary: a few points, the kind of thing polls and priors get wrong all the time. A small share, about one race in twenty, are much larger. The model separates the two. Every House and governor race carries a narrow 'core' of ordinary error, which shrinks with polling and grows in open seats and lopsided races, and a wide 'tail' for the rare large miss, both fitted to the record of 2018 to 2024. That is why a race at 19 points shows a chance near 99% rather than 100%: the tail keeps a real possibility open.

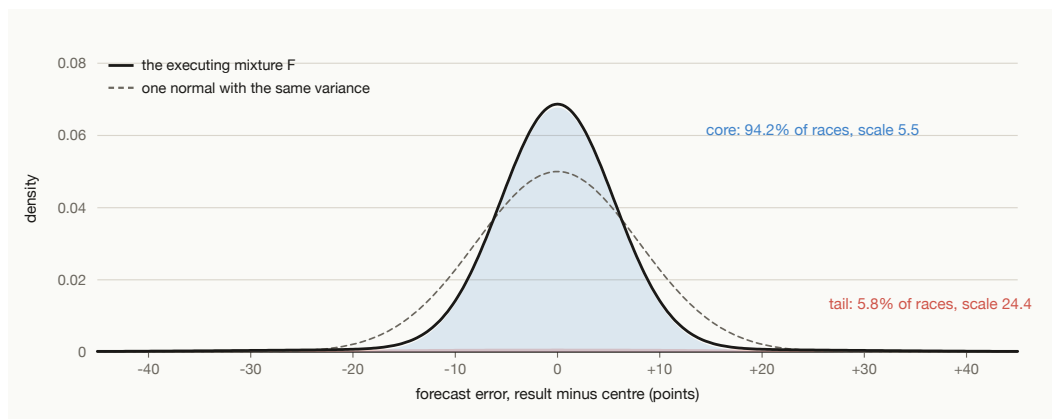


Figure 1: The core+tail error distribution F for an unpolled House race 53 days out with a centre of 5 points (core scale 5.55). The blue core carries ordinary error; the red tail, 5.8% of races, carries rare large misses. A single normal with the same variance (dashed) is too wide in the middle and too thin in the tails.

Senate races use a related recipe with a heavier-tailed bell curve whose width was fitted the same way. In every office the published range is the middle 80% of the fitted distribution, and on the historical record results land inside that range about four times in five, which is what an 80% range is for.

Why races move together

If the polls are wrong about the national mood, they are wrong everywhere at once. So the model does not roll each race's dice separately. It simulates the whole election night 100,000 times, and in each simulated night it first draws a national swing, then a regional one, then a state one, and only then each race's own surprise. A wave that flips one close seat tends to flip the others; a good night for one party in the Senate tends to be a good night in the House. Chamber control comes out of those nights, not out of adding up independent races.

Chamber control

Each simulated night produces a seat count. The share of nights in which Democrats reach 218 House seats is the House control probability; the Senate adds the 34 Democratic-aligned seats not on the ballot and needs 51. Right now the model expects about 229.5 Democratic House seats with a 74% chance of a majority, and about 50.3 Democratic-aligned Senate seats with a 48% chance of control. Because every number comes from the same simulated nights, the seat expectation and the race-by-race chances always agree with each other.

How to read a probability

A 70% chance is not a prediction that the favourite wins; it is a statement that in ten such races, three would go the other way. A 95% chance still loses one in twenty. When the model says

‘Lean’ or ‘Likely’ it is describing the size of the margin and the width of the range together. The most useful habit is to read the range: a race at 4 points with a range from -8 to $+16$ is a toss-up with a lean, whatever the headline number says.

Probabilities are shown to the nearest whole percent. A race listed at 99% is not certain; it is a race the model gives roughly one chance in a hundred of surprising.

The history chart

The chart on each race page shows two kinds of point. Faded points, from 2026-08-01 to the launch, were computed later, by running the model on the polls and other admitted evidence available on each earlier date. They are the best account of what the model would have said, and they are not forecasts anyone could have read then. Solid points, from 2026-09-09 on, are forecasts issued live, each stored once when it was published and never rewritten. The launch marker separates the two, and movement across it should be read as a change of method, not as news.

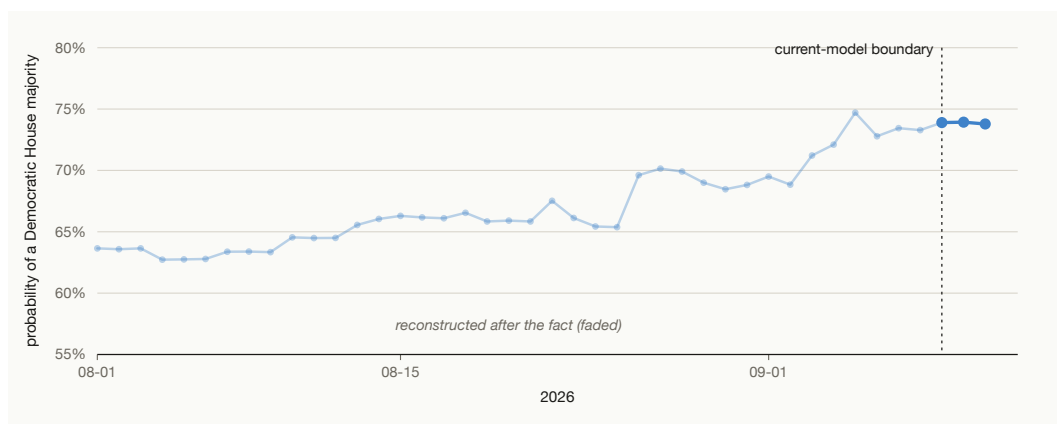


Figure 2: House control: faded reconstructions from 2026-08-01; solid current-model readings from 2026-09-09, including this guide's edition. A local preview does not establish publication.

What the model cannot know

Its modern validation covers four election cycles. It cannot know how a scandal, a debate or an October surprise will land; it learns from historical errors. Its Senate centre bias in that validation is $+2.7$ points, where a positive value means too Democratic. Better safe-seat margins do not mean every metric improves: the Senate repair accepts a small close-race error increase, and more realistic tail uncertainty slightly worsens overall distribution scores. It has no reliable history for two-Democrat races or ranked-choice counts and says which assumptions matter. A forecast may only use what was known on the day it was made.