

the WiseBot

Failed

WISEBOT • VERSION 2

Scalping

Three strategies, three losses.

Second in a series documenting every version of WiseBot. Like the first, this one didn't work. Here is what it tried, what the numbers were, and the more useful lesson hiding underneath the loss.

the WiseBot

the-wisebot.com
June 2026

What v1 taught me

The first WiseBot won 90% of its bets on Polymarket and still lost money, because the payoff was about 1:24 against it — a high win rate stapled to a terrible reward-to-risk ratio. (Full story in the previous paper.) I took three rules out of that wreck: a high win rate is not an edge; a losing system is usually a design problem, not a code problem; and no real capital goes in before an edge is actually proven.

So for the second version I gave myself one concrete constraint: **a healthier payoff geometry**. Whatever I traded next had to win and lose roughly comparable amounts — no more risking \$5 to make 21 cents.

There was also a personal reason I picked this particular game, and I'd rather be upfront about it.

The new bet

The year before, I'd been trading SOL perpetual futures by hand. For about two weeks it went beautifully — I was up roughly **+100%**. Then I gave it all back in a single day. Classic. If you've traded, you've either done this or watched a friend do it. I stopped after that.

That experience is exactly *why* I thought of automating it. My hunch was that I hadn't lacked an edge so much as discipline — that the blow-up came from a human holding on too long, sizing too big, trading on emotion. A bot, I reasoned, would take the same small edge without the human failure modes: no revenge trades, no hope, mechanical exits. So WiseBot v2 left Polymarket's binary markets for **continuous SOL perpetuals** — a market where positions can be sized sanely, exited at a defined level, and where wins and losses live on the same scale.

That was the theory. We'd see — and this time too, well, we saw.

Three strategies

I didn't want to bet the whole thing on one idea, so v2 ran three strategies side by side, each on its own slice of paper capital, so I could compare them honestly against the same market:

- **A random baseline.** It opened long or short essentially at random on a timer, with disciplined exits. Its only job was to measure what "pure exit discipline with no edge" actually earns — a control, like a placebo arm. If a real strategy can't beat random, it has no edge.
- **A mean-reversion strategy.** It bet that price, after stretching away from a short moving average, tends to snap back. Standard, well-known, the kind of thing that sometimes works in choppy markets.
- **A regime-filtered version** of the mean-reversion idea, which only allowed trades when volatility, borrow cost, and volume looked favorable — meant to sit out the conditions where the plain version gets chopped up.

All three shared the same exit rules: take profit at +0.5% net, stop loss at -0.4% net, a 30-minute time stop, modest leverage. That's a reward-to-risk near 1:1 — deliberately, the opposite of the v1 trap.

Paper-first discipline

This time the rule from v1 — *no capital before proof* — was enforced literally. v2 traded **only on paper** for a fixed 14-day window, with an explicit go/no-go decision at the end and pre-committed thresholds it had to clear: positive PnL, more than 50 trades, win rate above 50%, drawdown under 30%, a positive risk-adjusted return. No moving the goalposts after seeing the results.

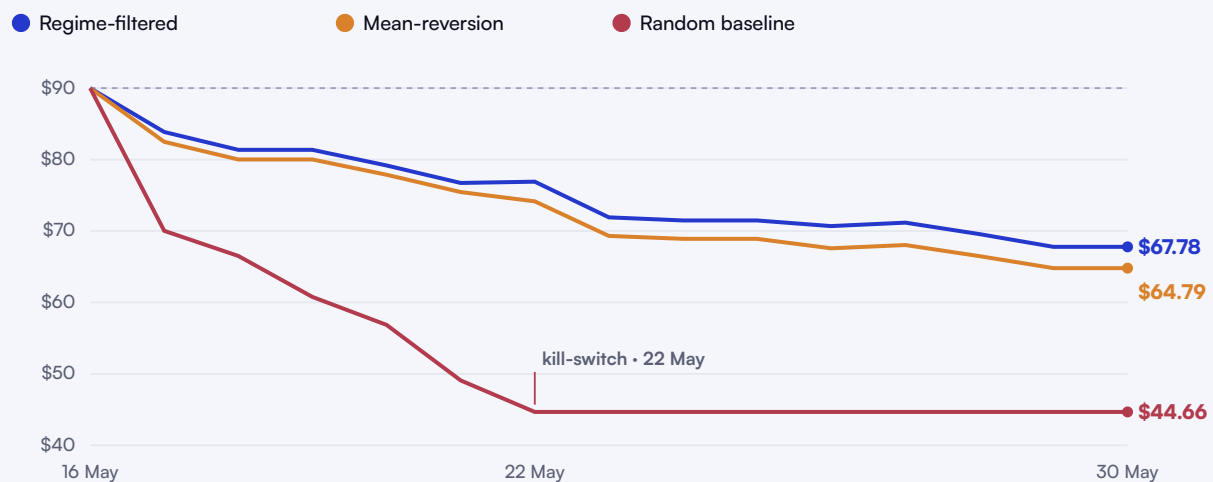
The results

Fourteen continuous days, 17–30 May. Here is what each strategy did, unedited:

Strategy	Ending equity (from \$90)	PnL	Trades	Win rate	Max drawdown
Regime-filtered	\$67.78	–24.7%	46	28%	24.7%
Mean-reversion	\$64.79	–28.0%	50	24%	28.0%
Random baseline	\$44.66	–50.4%	254	33%	50.4%

Equity over the 14-day paper run

Daily equity from the bot's own logs — \$90 start, all three drift down.



Real daily equity, straight from the logs. The random baseline cliffs to –50% in five days and trips its kill-switch on 22 May, then flatlines; the two real strategies bleed down more slowly. None finished near the \$90 they started with.

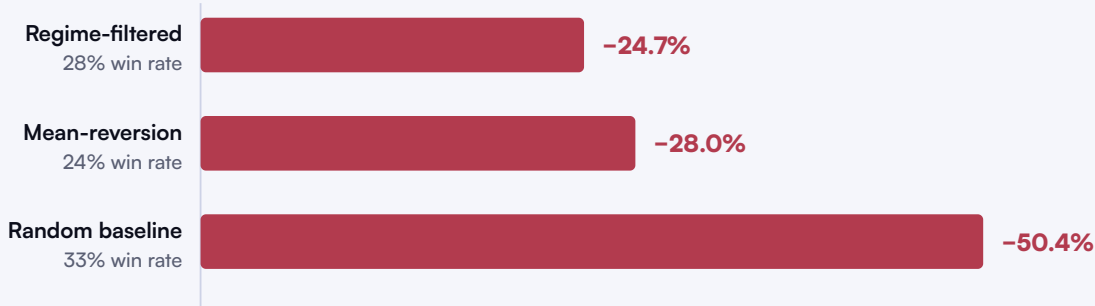
All three lost. The binding threshold — positive PnL — failed for every one of them. The "least bad" was the regime-filtered strategy at –24.7%, which is still a clear loss. The random baseline collapsed to –50% in five days and tripped its kill-switch on 22 May, exactly as a strategy with no edge should.

It's worth noting the one thing that *didn't* fail: the machine. Fourteen days, zero downtime, every daily log written, the kill-switch firing correctly on cue. The plumbing was sound. The trading thesis was not. Those are different things, and keeping them separate matters.

There's also a small, sharp detail in that table. The random baseline had the **highest** win rate (33%) and the **worst** result (-50%). That's v1's lesson echoing back: hit rate and profitability are not the same number, and confusing them is how you lose money twice.

Final result: all three strategies lost

Total ROI over the 14-day paper run — each started at \$90.



The random baseline had the highest win rate (33%) and the worst result. Hit rate $\neq$ profit.

Reward-to-risk was a healthy $\sim 1:1$ this time — but one low-volatility regime plus $\sim 0.12\%$ round-trip fees made the expectation negative. The machine ran flawlessly for 14 days; the trading thesis didn't hold.

The autopsy

Why did all three lose? Three reasons, in order of importance.

- 1. One single market regime.** SOL stayed in a low-volatility drift the entire two weeks — price went 86.6 $\rightarrow$ 82.6 (-4.6%) on tiny daily ranges. That is the worst possible environment for these strategies. Fourteen days sampled *one* mood of the market, not the range of moods a strategy has to survive.
- 2. The cost/payoff structure can't survive chop.** With $\sim 0.12\%$ in round-trip fees against a 0.5% profit target, fees eat roughly a quarter of a gross win before borrow costs. In a market that swings less than the take-profit distance, the wins barely cover costs and the 30-minute time stops close trades around breakeven-minus. With a win rate under a third and $\sim 1:1$ reward-to-risk, the expectation is structurally negative — the same *family* of mistake as v1, just less extreme.
- 3. The regime filter never found its regime.** The filtered strategy was built to wait for favorable conditions, but in 14 days the market never offered them. So it kept trading in the low-vol chop where it had no edge — and lost.

What I took away

Here's the meta-lesson, and it's the one that actually changed everything that came after.

Fourteen days of forward paper trading is too slow and too narrow to validate anything. It samples a single regime and calls it evidence. I had correctly refused to risk real money — but I'd been trying to validate an edge by *living through it in real time*, which means I could only ever see one market at a time, slowly.

The honest verdict on v2 was therefore narrow: it proved *absence of edge in low-volatility chop*, not *strategies broken in all conditions*. But it would have been waste to keep paying real fees to rediscover the same result, and dishonest to keep tuning parameters on two weeks of one regime and call the result a strategy. That's just overfitting with extra steps.

So v2 was a clean **no-go**. No real capital went in. And the next version started from a completely different discipline: before any live or forward trading, test the idea against *years* of history across bull, bear, and chop — let a backtest see hundreds of regimes in minutes instead of waiting for one to crawl past in real time.

That shift — from "watch it work" to "prove it across history first" — is what version three is built on. That's the next paper, and it's the first version where the honest answer isn't a flat "it failed."

Methodology note

The analysis and drafting of this document are AI-assisted (I work with Claude as my stated method, the same way the bot itself uses AI). Every figure here comes from the bot's own logged paper-trading data over 17–30 May 2026; every claim is mine and is meant to be verifiable. The AI helps me write clearly and pressure-test my reasoning — it does not invent the numbers, and it does not decide what's true. That stays with me.

the WiseBot

Verify me, don't trust me.

the-wisebot.com

Not financial advice. Past performance does not indicate future results.