

The Gold Masterclass

A trading course built from measured results — including the tests that failed, the numbers we had to correct, and the instruments this method does not work on.

What a liquidity sweep actually is

Why price reaches for a level, takes it, and turns — explained without needing to believe anyone is hunting you personally.

A liquidity sweep is what happens when price moves to a level where a lot of stop-loss orders are sitting, triggers them, and then reverses. It looks like a failed breakout. It is usually the opposite — a successful one, just not for the people who got stopped.

Why stops cluster

Traders are taught to put stops beyond obvious levels: below a swing low, above a swing high, under a round number. That advice is fine in isolation, but everyone gets the same advice. The result is that stop orders pile up in predictable places.

A stop-loss to sell is, mechanically, a market sell order waiting to happen. A cluster of them under a swing low is a pool of guaranteed selling. If you need to buy a large position without pushing the price up against yourself, that pool is exactly where you want to be buying — because someone else's forced selling is your liquidity.

THE CORE IDEA

A swing high or low is not a wall. It is a **shopping list**. Price goes there because that is where the orders are, not because the level is meaningful in itself.

What the pattern looks like

On a chart it is unremarkable, which is part of why it works:

1. Price ranges for a while, leaving a clear high and low.
2. It pushes *through* one of them — the breakout everyone was waiting for.
3. It fails to hold, and closes back inside the range on the same candle or the next one.
4. It moves away in the opposite direction, often quickly.

Step 3 is the whole signal. A break that closes beyond the level is a break. A break that pokes through and closes back inside is a sweep. The difference is one number — where the candle closed — and it changes the meaning entirely.

Why "smart money" language is mostly unnecessary

You will see this taught with a lot of vocabulary about institutions and manipulation. Most of it is unfalsifiable, and you do not need any of it. You do not have to believe anyone is hunting you personally. You only have to accept two boring facts:

- Stops cluster in predictable places, because trading education is homogeneous.
- Large orders need someone to trade against, and clustered stops provide that.

That is enough to explain the behaviour, and it has the advantage of being testable.

What we could actually measure. Over ~2.5 years of gold 15-minute data, fading sweeps of a 6-hour swing level produced a positive expectancy that survived a holdout period. It also produced a *lot* of losing trades — roughly two in three. The pattern is real. It is not reliable in the way beginners hope "real" means, and module 5 is about exactly that.

What to take from this module

- A sweep is a break that **closes back inside**. The close is the signal.
- Levels attract price because of the orders resting there, not because of the level.
- You do not need a conspiracy to explain it, and you should be suspicious of teaching that requires one.

Most sweeps are noise

The depth threshold that separates a real liquidity grab from an ordinary wick, and the measurement showing that more trades means less edge.

Once you can see sweeps, you see them everywhere — and that is the problem. Most of what looks like a sweep is an ordinary wick. The difference between a strategy and a pattern-matching habit is a threshold, and where you set it is a trade between how often you trade and how good each trade is.

The depth threshold

Our engine only counts a sweep if price pushes through the level by a minimum distance before closing back inside. On gold that distance is currently **120 points (\$1.20)**, which is roughly 0.3 of a typical 15-minute candle's range.

Below that, you are not looking at a liquidity grab. You are looking at noise that happens to touch a line you drew.

The measurement

The obvious temptation is to lower the threshold, because fewer restrictions means more trades and more trades *feels* like more opportunity. We tested exactly that. Same entry logic, same session, same stops and targets — only the required depth changed:

Sweep depth	Trades	Profit factor	Holdout PF
20 points	250	1.53	1.71
40 points	224	1.48	1.73
80 points	162	1.73	2.11
120 points (live)	123	2.16	2.52
160 points	98	2.26	2.48
200 points	79	2.43	2.35

Quality rises steadily with depth: doubling the threshold from 80 to 160 points cuts the number of trades by 40% and lifts profit factor from 1.73 to 2.26. Each trade is worth more, and there are fewer of them.

An honest correction, and a lesson in itself. An earlier version of this module showed shallow sweeps *losing* money — profit factor 0.94 at 20 points. That was measured before we widened the stop, and it is no longer true: at the current stop distance, shallow sweeps are profitable, just less so. We caught it re-checking our own figures before publishing.

The lesson is not that we made a mistake. It is that **a number is only true for the exact configuration it was measured on**. Change one setting and every other number you quoted may quietly stop being true. This is the single most common way trading education misleads people without anyone lying.

So what is the threshold actually for?

Not survival — quality. At the wide stop the strategy tolerates shallow sweeps, but you pay for the extra trades in edge per trade. Where you sit on that curve is a real decision:

- **Deeper (160–200):** fewer, better trades. Higher profit factor, longer gaps between signals, more patience required.
- **Shallower (40–80):** roughly double the activity at a materially lower profit factor — and more trades also means more exposure to costs and to execution slippage that a backtest models imperfectly.

We run 120 because it sits where the curve is still steep — you gain a lot of quality for each step up, and give up a lot for each step down.

Why deeper sweeps are worth more

A plausible explanation, offered as reasoning rather than proof: a shallow poke through a level may not reach the bulk of the resting stops. If fewer orders were triggered, less forced flow occurred, and there is less reason for price to reverse hard. You have identified the location without the full event.

A deep push is evidence that something substantial was actually absorbed there.

Be careful with this reasoning. The story is consistent with the data, but the data does not prove it — other explanations fit the same numbers equally well. Treat the *measurement* as the finding and the explanation as a hypothesis. Module 8 is about why keeping those two things separate matters more than almost anything else in trading.

What to take from this module

- Depth is what separates a sweep from a wick. Ours is 120 points on gold.
- Quality rises with depth: 80 points gives PF 1.73, 160 points gives 2.26.
- Frequency is a trade against edge per trade, not a free dial.
- A measured number is only true for the exact configuration it was measured on — including ours, including this table.

Stop placement — the biggest finding

The same 98 trades went from profit factor 1.26 to 1.78 by moving one number. Why a tight stop gets hunted out of trades that would have worked.

This is the single biggest thing we found, and it is not about entries at all. We changed one number — where the stop goes — and profit factor went from **1.11 to 1.62** on essentially the same trades. Same entry rule, same session, same targets. Better stop.

The finding

Our stop sits beyond the swept level, plus a buffer. The buffer was 150 points. We tested it across a wide range, on the same trade list:

Stop buffer	Avg stop	Profit factor	Holdout PF
150 points	385 pts	1.11	1.29
250 points	459 pts	1.14	1.45
350 points	521 pts	1.37	1.81
400 points (live)	546 pts	1.49	1.69
500 points	584 pts	1.62	1.81

It is not one lucky setting. The whole region from about 350 to 600 points works, across every reward-to-risk ratio we tried, and it holds up on the holdout — data the choice was never fitted to. A single good cell is noise; a broad plateau is a finding.

Why a tight stop loses money here

Think about what you are doing when you fade a sweep. Price has just violently taken out a level. You are betting it comes back. But a level that got swept once is *a level that just proved it attracts price* — and it very often gets tested again before the reversal actually develops.

A tight stop sits right inside that re-test zone. You get taken out of a trade that then goes on to work, and you pay a full loss for the privilege. The trade thesis was correct. The stop was in the wrong place.

THE COUNTER-INTUITIVE PART

A wider stop is **not** more risk. Position size is calculated *from* the stop distance: a stop twice as far away means a position half as large, for the same dollars at risk. You are not risking more — you are risking the same amount, with more room to be right.

The trap this creates

Read that box again, because it is easy to hear "wider stops are better" and simply move your stop while keeping your position size. That is not what this says. That is doubling your risk, and it is how the finding turns into a blown account.

Widening the stop only works if the size comes down with it. If you trade a fixed lot size, this entire module does not apply to you until you fix your sizing first — which is module 6.

A real limit worth knowing. Wider stops mean smaller positions, and brokers have a minimum position size. On a standard gold contract the smallest allowed trade risks about \$6.70 at these stop distances — over 3% of a \$200 account. Below roughly \$500 the maths simply does not fit, and the honest answer is that the account is too small for this strategy rather than that the stop should be tightened.

What to take from this module

- Stop placement mattered more than entry selection in everything we tested.
- Tight stops get hunted out of trades that would have worked.
- Wider stop + proportionally smaller size = same risk, more room. **Both halves are required.**
- Trust plateaus, not single settings.

How to tell a real edge from a curve-fit one

Holdout testing, sample size and the multiple-comparisons trap — including the filter that improved 7 configurations out of 7 and still failed.

This is the most useful module here, and it is not about gold. It is about how to tell whether *any* strategy — yours, ours, or one you paid for — has a real edge or has simply been fitted to the past. Almost everything sold to retail traders fails this.

The problem in one sentence

If you test enough variations, some will look excellent purely by chance, and you will pick those — because picking the best-looking one is exactly what testing feels like it is for.

Test forty variations of a coin-flip strategy and about two will show a profit factor above 1.05 on noise alone. Nothing in the result itself tells you which kind you are holding.

The fix: a holdout

Split your data before you start. Use the first portion to make every decision — settings, filters, thresholds, all of it. Then run the finished thing **once** on the portion you never looked at.

If the edge is real, it shows up in both. If it was fitted, it evaporates in the second. The discipline is not the split — it is that you only get one look, and you accept the answer.

OUR STANDARD

A strategy must clear the bar **twice**: on the full sample, and on a 30% holdout it was never tuned against. In-sample performance alone rewards whoever searched hardest.

A worked example of it saving us

We tested adding a volatility filter — skip trades when the market is unusually quiet. On the training data it improved profit factor in **seven configurations out of seven**. Seven for seven. Every variation we tried got better.

That consistency is exactly what a real effect looks like, and we were ready to ship it.

Then we ran it on the holdout. It made things *worse*: profit factor fell from 1.35 to 1.29 on the data it had never seen.

Seven for seven was a mirage. We dropped the filter.

Sit with that for a second. If we had tested it the way most strategies are tested — try it, see improvement, adopt it — that filter would be in the live system right now, quietly making it worse, and we would be citing "improved in 7 of 7 tests" as evidence in our marketing. It would have sounded rigorous. It was wrong.

The three questions to ask any strategy

1. **How many trades?** Under 30 proves nothing in either direction. Under 100 is a small sample and should be described as one. Ours is 123, and we say so every time.
2. **Was it tested on data used to build it?** If there is no holdout, the number is a description of the past, not a prediction.
3. **How many variations were tried before this one?** "We found a setting that works" means something very different after two attempts than after two hundred.

If a course, signal service or EA cannot answer all three, that is your answer.

Costs are part of the test

One more way results lie. We tested our strategy on USDJPY and it looked tradeable — profit factor 1.12. Then we corrected the cost model to include commission properly. FX moves roughly \$0.61 per point per lot, so a \$7 round-trip commission is about 11 points, not the 2 that a spread-only model implied.

With the correct cost, that same strategy scored **0.92**. A loser. The difference between a product and a mistake was one arithmetic error in the cost model.

What to take from this module

- Split your data first. Decide on one half. Look at the other half once.
- Consistency on training data is not evidence — our 7-of-7 filter still failed.
- Count your attempts. Search width is part of the result.
- Model costs properly, including commission converted into points.
- Ask any seller the three questions. Ask us too.

Appendix — the backtests behind this course

Every figure quoted in the modules comes from the runs below. They were generated in a single pass on 2026-07-28 so that no two numbers in this document come from different configurations — which is the specific mistake an earlier draft of this course made, and which module 8 is about.

What was tested

Instrument / timeframe	XAUUSD and four US index CFDs, 15-minute
Sample	~2.5 years of real broker bars per instrument
Method	Train on the first 70%; confirm once on a 30% holdout never used for any decision
Entry	Liquidity sweep of the 120-point level, faded
Stop	400-point buffer beyond the swept level, clamped 150–1000 points
Target	5.0 : 1
Session	07:00–15:00 server, with higher-timeframe trend alignment
Costs	Gold 37 points (spread + commission). Indices: spread only, which is how this broker prices them.
Sizing	Fixed 0.02 lots on a \$1,836.08 account, so the money column is the real account

Results

Symbol	Trades	Win rate	Profit factor	Train	Holdout	Worst losing run	Verdict
XAUUSD	126	31.7%	2.11	2.0	2.19	10	PASS
US30	174	16.1%	0.86	0.63	1.3	19	FAIL
US500	162	13.6%	0.63	0.81	0.4	32	FAIL
USTEC	152	15.8%	0.8	1.04	0.56	34	FAIL
US2000	175	14.9%	0.65	0.67	0.6	37	FAIL

Read the train and holdout columns together. A strategy that only works on the half it was built on has not been shown to work at all. US30 is the clearest example: it scores 0.63 on the training half and 1.30 on the holdout. Those two numbers disagreeing that violently is not a promising signal — it is what noise looks like when you measure it twice.

Gold, in money

At 0.02 lots on a \$1,836.08 account over 2.54 years: **126 trades**, **31.7%** of them winners, net **\$1,366.28**, worst drawdown **\$129.78**, expectancy **\$10.84** per trade. The longest losing run was **10 trades**.

Read that last number again. Ten losses in a row is not a malfunction of this strategy — it is a normal feature of one that wins about a third of the time and pays 5:1. If a run like that would make you turn the system off, the position size is too big, and that is a sizing decision rather than a strategy decision.

Why the indices failed

The same rule, ATR-normalised so the thresholds mean the same thing on an index priced near 44,000 as on gold near 4,000, loses money on all four US indices. The win rates say why: around **14–16%** on the indices against **31.7%** on gold. Swept levels on an index tend to keep going rather than snap back.

This is worth more to you than another winning backtest would be. A method that produces a positive result on everything you point it at is not measuring anything.

The limits of all of this

- **The sample is small.** 126 trades over 2.54 years. It survives a holdout, which is real evidence, but it is not proof.
- **These are backtested and demo results.** Hypothetical performance has inherent limitations: it is produced with hindsight, risks no real money, and cannot fully model slippage, rejected orders, or the pressure of trading live.
- **A live track record is still being built.** When we have one, it will be published alongside these figures — including the losing months.
- **Nothing here is a forecast.** These are measurements of the past.