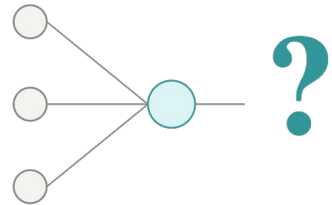


Testing Logistic Regression

From a trained model to a number that tells you how well it works — and what to do next.



You've trained your logistic regression model. So how do you know it actually works — not just on the data it learned from, but on data it's never seen before? That's the whole point of testing. In this video, we'll walk through how to evaluate your model and what to do with the result.

What you need to test

HELD-OUT VALIDATION SET - 200 USERS

X feature matrix						y labels
user	months_sub	logins_wk	tickets	plan_tier	age	churned
1	14	3	0	2	34	0
2	2	0	4	1	52	1
3	27	11	1	3	41	0
4	5	1	2	1	29	1
5	9	6	0	2	47	0
...	...	...	...	...	...	...

† 1 row = 1 user · 5 numbers — that's **one input** to the model.

$X \in \mathbb{R}^{m \times n}$ held-out features · $m = 200$ users (rows) · $n = 5$ features (columns)
 $y \in \{0, 1\}^m$ true labels · 200 zeros and ones (0 = stayed, 1 = churned)
 $\theta \in \mathbb{R}^n$ $\theta_1, \theta_2, \theta_3, \theta_4, \theta_5$ — 5 learned weights, one per feature (bigger = that feature matters more)

$\mathbb{R}$ = the real numbers (any decimal value). $\mathbb{R}^n$ = a list of n real numbers. $\mathbb{R}^{m \times n}$ = an m -by- n grid of real numbers.

Imagine a model that predicts whether a subscriber will cancel — or churn. To test it, you set aside data your model has never seen — what we'll call the held-out set.

Each row in this table is one user, described by five numbers like age and weekly logins. With 200 users, that's a matrix of 200 rows by 5 columns.

To test, you need three things: the held-out features X , the true labels y , and θ — the learned weights. We use new data to see how well our model generalizes. The process has three pieces: turn inputs into probabilities, threshold them into yes-or-no predictions, and check your accuracy.

Turn inputs into probabilities

FOR ONE USER - ROW 1

STEP A - WEIGHTED SUM

$$14 \cdot \theta_1 + 3 \cdot \theta_2 + 0 \cdot \theta_3 + 2 \cdot \theta_4 + 34 \cdot \theta_5 = z_1 = -0.3$$

raw score - any real number

↓ apply sigmoid σ

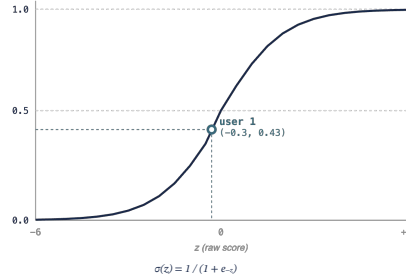
STEP B - PROBABILITY

$$\sigma(-0.3) = 0.43$$

probability - between 0 and 1

→ 43% chance user 1 will churn

SIGMOID - SQUASHES ANY NUMBER INTO [0, 1]

Repeat for all 200 users $\rightarrow h = \sigma(X\theta) \rightarrow$ 200 probabilities, one per user

First, we turn those inputs into probabilities. We multiply a user's five numbers by their weights and sum them to get a raw score.

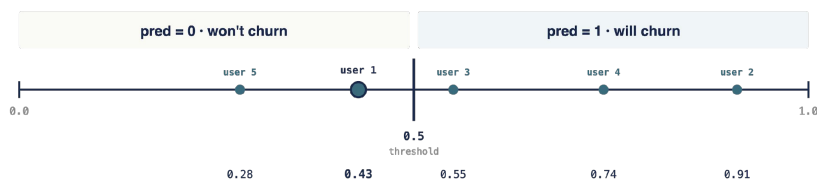
User 1 might get a negative zero point three. But raw scores can be any number, and probabilities must run from zero to one.

The sigmoid function fixes this by squashing any number into that zero-to-one range. Big positives end up close to one, big negatives close to zero. So sigmoid of negative zero point three is zero point four three — a forty-three percent chance of churning. Doing this for all 200 users gives us h , our list of probabilities.

Turn probabilities into predictions

From Step 1: h - 200 probabilities, one per user

↓ apply threshold $\text{pred} = (h \geq 0.5)$



user 1 $h_1 = 0.43 < 0.5$ + $\text{pred}_1 = 0 \cdot \text{won't churn}$

$$\text{pred} = (h \geq 0.5)$$

Result: $\text{pred} \in \{0, 1\}_{200}$ - 200 binary predictions, one per user

Now that we have probabilities, we face a new problem. Your business can't actually act on a probability. It needs a clear answer — will they churn, or not?

So we set a threshold, usually zero point five. If a probability is zero point five or higher, we predict one — they'll churn. Anything below, we predict zero — they won't.

What about someone right at the edge — zero point five-one, or zero point four-nine? That's exactly why we need the threshold. It forces a clean decision. The result is pred — a list of 200 zeros and ones, one for each user.

Measure how often you got it right

FOR EACH USER - COMPARE PRED TO Y

user	pred	y	match?
1	0	0	✓
2	1	1	✓
3	1	0	✗
4	1	1	✓
5	0	0	✓
⋮	⋮	⋮	⋮

USER 1 - MATCH

 $y_1 = 0, \text{ pred}_1 = 0 \rightarrow 0 = 0 \rightarrow 1$ ✓

USER 3 - MISS

 $y_3 = 0, \text{ pred}_3 = 1 \rightarrow 1 \neq 0 \rightarrow 0$ ✗

ADD UP MATCHES - DIVIDE BY TOTAL

$$\text{accuracy} = \frac{\sum_{i=1}^m (y_i == \text{pred}_i)}{m}$$

$\sum_{i=1}^m$ ↓ add up across all 200	$(y_i == \text{pred}_i)$ ↓ 1 if match, else 0	m DENOMINATOR ↓ divide by 200 (total users)
---	--	--

$$\frac{168}{200} = 0.84$$

84%

ACCURACY

With predictions in hand, we can finally check how well our model did. For each user, compare your prediction to the true label. Count the matches. Divide by the total. That gives you accuracy — a single number that summarizes all 200 comparisons.

Accuracy runs from zero to one. One — or one-hundred percent — means you got every prediction right. Zero means you got none.

Let's say you got 168 out of 200 right. That's 168 divided by 200, or zero point eight-four — eighty-four percent accuracy.

So... is that good? **[pause for thought]**

Is 84% good?

It depends on the cost of being wrong.

SPAM FILTER

84%

TOO LOW

Letting through 16% of spam means a flooded inbox — unacceptable for users.

RARE DISEASE DETECTION

84%

EXCELLENT

Strong result for rare events where most random guesses fail. A genuine signal.

ANOTHER REASON ACCURACY MIGHT BE LOW - OVERFITTING

Training accuracy



98%

Validation accuracy



65%

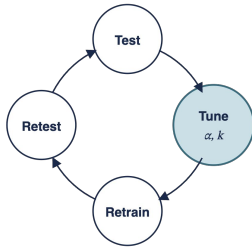
When training accuracy > validation accuracy, the model **memorized** the training data instead of **learning general patterns** — it aced the practice test but flunked the real one.

It depends. It depends on the cost of being wrong.

For predicting a rare disease, eighty-four percent might be excellent. For a spam filter, it's not nearly enough — sixteen percent of spam landing in your inbox every day is unacceptable.

There's another reason accuracy might be low: overfitting. That's when your model memorized the training data instead of learning the actual patterns. How do you spot it? Training accuracy is much higher than held-out accuracy. In other words, the model aced the practice test but flunked the real one. We'll explore how to handle overfitting another time.

Use accuracy to improve your model



HYPERPARAMETERS

α (alpha) · learning rate

How big a step the algorithm takes each weight update.

Typical range: 0.001 to 0.1

k · iterations

How many update steps total during training.

Typical range: 100 to 10,000

ACCURACY ACROSS ROUNDS · WHAT TUNING PRODUCES



Testing isn't the end of the process — it's a tool for continuous improvement.

Accuracy is your guide for what to change next. You have two main knobs to turn, or **hyperparameters**: alpha, the learning rate, and k, the number of iterations. Adjust these, and your accuracy changes.

But there's an important caveat. The moment you use this held-out data to tune your model, it becomes a **validation set**. To know your model's real-world accuracy, you need one final, completely untouched **test set** — checked only after all tuning is done.

Testing isn't the end of the process — it's a tool for continuous improvement.

Three things to remember

①

THE PROCESS - 3 INPUTS IN, 3 STEPS OUT

$$\left| X \cdot y \cdot \theta \right| \rightarrow \left| h = \sigma(X\theta) \right| \rightarrow \left| \text{pred} = (h \geq 0.5) \right| \rightarrow \left| \text{accuracy} = \sum(y_i == \text{pred}_i) / m \right|$$

inputs · probabilities · predictions · one number

②

THE INTERPRETATION

Whether your accuracy is good depends on the **cost of being wrong**.

84%

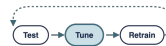
SPAM: TOO LOW

84%

DISEASE: EXCELLENT

③

THE RESET

Testing isn't the end — it's a tool for **continuous improvement**.

NEXT VIDEO

→

Tuning your model

So, the big picture. Testing tells you whether your model truly generalizes — not just memorizes. Whether the accuracy is good depends on the cost of being wrong. And testing isn't the finish line — it's where tuning begins.

In the next video, we'll start with the learning rate — how a small change in alpha can mean the difference between a model that quickly converges and one that wanders for thousands of iterations.

See you there.