

EDUSHARK TRAINING

Self-paced programmes in data and AI

Machine Learning

A 16-Week Practitioner Workbook

Build models, measure them honestly, and ship one.



Machine Learning: A 16-Week Practitioner Workbook

Publisher: EDUSHARK TRAINING

Web: <https://edusharktraining.com/machine-learning>

Contact: info@edusharktraining.com

This workbook is the offline form of a course published in full and free to read at edusharktraining.com. It may be passed on to anyone, including a whole team. All trademarks are the property of their respective owners.

Typeset with $\LaTeX$.

How to Use This Workbook

This book covers 16 weeks, 112 lessons and 128 self-assessment questions. Assumes some Python. No prior machine learning.

What is different about it

Every result in this book came out of code that was run, and the printed output is what the machine actually returned. Where a technique lost to something simpler, the number that says so is on the page next to it.

How each chapter is laid out

One chapter per week, seven days per chapter. Each day states what it is for, works through the material, and ends with a takeaway worth keeping. Code appears as it was run, and program output is set apart in grey so you can tell at a glance which is which:

```
1 | print(sum(x * x for x in range(4)))
```

14

Every chapter ends with a self-assessment quiz. The answers are on the same page, set small, because a question you cannot check is not worth attempting. Anything you get wrong points at the day to re-read.

Working alongside the site

Every chapter here matches a week published at <https://edusharktraining.com/machine-learning>. The website version carries the same material with the code laid out for copying. Use whichever suits where you are reading.

Contents

How to Use This Workbook	i
1 Python for Machine Learning	1
2 Data Wrangling and Exploratory Analysis	25
3 The Maths That Actually Matters	49
4 Regression, Regularisation and Leakage	73
5 Classification Fundamentals	101
6 Trees, Forests and Boosting	129
7 Model Evaluation and Validation	161
8 Feature Engineering	191
9 Unsupervised Learning and Clustering	219
10 Dimensionality Reduction and Anomaly Detection	245
11 Neural Networks From Scratch, Then Keras	269
12 Convolutional Networks and Representation Learning	301
13 Text, Embeddings and Sequence Models	345
14 Hyperparameter Optimisation and Interpretability	379
15 Deployment and MLOps	411
16 Capstone and Career	441

Python for Machine Learning

Week 1

Set up the environment, learn the array and the DataFrame, use a linear working dataset model before the week is out.

OBJECTIVES

Week 1 does two jobs. It gives you the Python you need, arrays, DataFrames, plots, and the three-method contract every scikit-learn object obeys, and it gets you to a real, measured model by day 7 rather than making you wait a month for one. You will also generate the customer dataset the rest of the course is taught on, deliberate defects and all.

Day 1 What Machine Learning Actually Is

Supervised, unsupervised, and the one idea the whole field rests on

Most programming you have done so far is *instructional*. You knew the rule, add tax at 20 percent, reject an order under £5, and you wrote it down. Machine learning inverts that. You have the answers and you want the rule, and the rule is one nobody can state precisely.

Nobody can write the condition for “this customer is about to leave”. But a company has thousands of customers who did leave and thousands who stayed, and the difference is in the data somewhere. A learning algorithm searches for it.

The three families you will meet

Almost everything in this course sits in one of three boxes. Knowing which box a problem belongs to decides which algorithms are even candidates.

Family	You give it	It learns to	Example
Supervised	Inputs <i>and</i> known answers	Predict the answer for new inputs	Will this customer churn?
Unsupervised	Inputs only	Find structure, groups, axes, oddities	Which customers behave alike?
Reinforcement	An environment and a reward	Choose actions that pay off over time	Which offer to show, learning as you go

Weeks 4 to 7 are supervised learning, weeks 9 and 10 are unsupervised. Reinforcement learning is a course of its own and we only sketch it.

Supervised splits again, by what you are predicting

Inside supervised learning the type of the answer decides the metric, the loss function and often the algorithm:

- **Regression:** the answer is a number on a continuous scale. How much will this customer spend next quarter?
- **Classification:** the answer is one of a fixed set of labels. Churn or stay. Fraud or legitimate. Which of nine defect types.

The same dataset can pose both. `monthly_charges` is a regression target; `churned` is a classification target. You will use both this course.

The one idea that separates ML from statistics

A statistician asks whether a relationship is real. A machine learning practitioner asks whether a model will work on data it has not seen. That second question has a name, and it is the discipline the whole field is built on.

Watch Out

The mistake that ruins most first projects

You fit a model, score it on the same rows you fitted it on, and get 99 percent. It feels like success. It is the machine learning equivalent of marking your own exam with the answer sheet open. Any model complex enough can memorise its training set. Week 7 is devoted to measuring this honestly; until then, treat every score computed on training data as meaningless.

What you need installed

Python 3.10 or newer, and a handful of libraries. Install [Python from python.org](https://python.org) and tick *Add Python to PATH* on Windows, or use your package manager on Linux and macOS. Check it took:

Work in a virtual environment, always

A virtual environment is a private copy of Python's package folder for one project. Without one, every project on your machine shares a single set of library versions, and the day a project needs an older scikit-learn you break the other five. Create one per project:

Your prompt gains a `(.venv)` prefix. That prefix is the whole point: it tells you which Python you are about to run. Now install the stack:

Pin what you installed

Run `pip freeze > requirements.txt` once the environment works. That file records exact versions, so `pip install -r requirements.txt` reproduces it on another machine, or on yours in six months, when the newest scikit-learn has renamed the argument your code passes.

How you will actually work

Two tools, two jobs. Use a notebook while you are exploring, because you want to see each intermediate result. Move to a `.py` file the moment the code matters, because notebooks hide

execution order and a notebook that only runs if you press the cells in the right sequence is not reproducible.

Day 1 takeaway

Day 2 NumPy and Vectorised Thinking

Arrays, shapes, broadcasting and why loops are the wrong habit

Every machine learning library in Python stands on NumPy. A scikit-learn model does not really take a DataFrame; it takes an array of numbers, and pandas hands one over on request. TensorFlow tensors follow NumPy's rules deliberately. Learn the array and the rest of the stack stops feeling arbitrary.

Why not a plain Python list

A Python list holds pointers to objects scattered across memory, each one carrying its own type information. A NumPy array holds raw numbers of one type packed end to end. That single difference buys both compactness and speed, because the loop that adds them runs in C instead of in the interpreter.

```
1 import numpy as np
2
3 py_list = list(range(1_000_000))
4 arr = np.arange(1_000_000)
5
6 print(type(arr), arr.dtype, arr.shape)
7 print('list bytes:', py_list.__sizeof__())
8 print('array bytes:', arr.nbytes)
```

```
<class 'numpy.ndarray'> int64 (1000000,)
list bytes: 8000040
array bytes: 8000000
```

The sizes look similar, but the list figure counts only the pointers. Each of the million integers is a separate object of roughly 28 bytes on top of that. The array figure is the whole story.

Vectorised thinking

The habit to build this week: stop writing loops over rows. An operation written against a whole array is shorter, faster and less likely to be wrong.

```
1 import numpy as np
2 import time
3
4 charges = np.random.default_rng(0).uniform(20, 120, 1_000_000)
5
6 t = time.perf_counter()
7 loop = [c * 1.2 for c in charges]
8 loop_time = time.perf_counter() - t
9
```

```

10 t = time.perf_counter()
11 vec = charges * 1.2
12 vec_time = time.perf_counter() - t
13
14 print(f'loop           {loop_time:4f}s')
15 print(f'vectorised    {vec_time:4f}s')
16 print(f'speedup       {loop_time / vec_time:0f}x')

```

```

loop           0.083917s
vectorised    0.002325s
speedup       36.093464x

```

Forty-something times, for arithmetic. On the operations that matter in model fitting the gap is wider still, because those call into optimised linear algebra libraries.

Shape is the thing you will get wrong

Nearly every confusing error in scikit-learn and TensorFlow is a shape error. Build the habit now of printing `.shape` whenever something surprises you.

```

1 import numpy as np
2
3 a = np.array([1, 2, 3])           # 1-D, shape (3,)
4 b = np.array([[1, 2, 3]])        # 2-D, one row
5 c = np.array([[1], [2], [3]])    # 2-D, one column
6
7 for name, x in [('a', a), ('b', b), ('c', c)]:
8     print(name, x.shape, x.ndim)

```

```

a (3,) 1
b (1, 3) 2
c (3, 1) 2

```

Watch Out

(3,) is not (3, 1)

A one-dimensional array of three numbers and a column of three numbers are different objects, and scikit-learn cares. Fitting on a 1-D feature array raises *Expected 2D array, got 1D array instead*. The fix is `x.reshape(-1, 1)`, where `-1` means *work out this dimension from the length*.

Broadcasting

When you combine arrays of different shapes, NumPy stretches the smaller one to fit, if the shapes are compatible. Compatible means: compare dimensions from the right, and each pair must be equal or one of them must be 1.

```

1 import numpy as np
2
3 X = np.array([[10.0, 200.0],
4               [12.0, 240.0],

```

```

5         [11.0, 100.0]])
6
7 col_mean = X.mean(axis=0)      # shape (2,) -- one mean per column
8 centred = X - col_mean        # (3, 2) - (2,) broadcasts down the rows
9
10 print('means ', col_mean)
11 print('centred\n', centred)
12 print('new means', centred.mean(axis=0).round(10))

```

```

means    [ 11. 180.]
centred
[[ -1.  20.]
 [  1.  60.]
 [  0. -80.]]
new means [0. 0.]

```

That is standardisation, in one line, and it is exactly what `StandardScaler` does in week 8. `axis=0` means “collapse the rows, keep the columns”, which is the direction you almost always want, because columns are features.

Boolean masks

Selecting rows by condition is the operation you will reach for constantly.

```

1 import numpy as np
2
3 charges = np.array([25.0, 88.5, 104.0, 19.0, 71.2])
4 high = charges > 70
5
6 print('mask ', high)
7 print('values ', charges[high])
8 print('count ', high.sum())      # True counts as 1
9 print('share ', high.mean())     # so the mean is the proportion

```

```

mask    [False  True  True False  True]
values  [ 88.5 104.   71.2]
count    3
share    0.6

```

`mask.mean()` is a proportion

Because `True` is 1 and `False` is 0, the mean of a boolean array is the fraction of elements that satisfy the condition. `(y == 1).mean()` is your churn rate. You will use this constantly and it saves defining a counter.

Day 2 takeaway

Day 3 pandas and the Course Dataset

Series, DataFrames, group-by, and generating the data you will use

NumPy gives you a grid of one type. Real data is a table with a name and a type per column, missing values, and dates. That is pandas.

Series and DataFrame

A **Series** is one column: values plus an index. A **DataFrame** is a dictionary of Series sharing that index.

```
1 import pandas as pd
2
3 s = pd.Series([25.0, 88.5, 104.0], name='monthly_charges')
4 print(s)
5 print('dtype:', s.dtype)
6
7 df = pd.DataFrame({
8     'customer_id': ['C1', 'C2', 'C3'],
9     'monthly_charges': [25.0, 88.5, 104.0],
10    'contract': ['Month-to-month', 'Two year', 'Month-to-month'],
11 })
12 print(df)
13 print(df.dtypes)
```

```
0    25.0
1    88.5
2   104.0
Name: monthly_charges, dtype: float64
dtype: float64
   customer_id  monthly_charges  contract
0          C1             25.0  Month-to-month
1          C2             88.5    Two year
2          C3            104.0  Month-to-month
customer_id      object
monthly_charges  float64
contract         object
dtype: object
```

Watch Out

object means pandas gave up

A dtype of **object** means the column holds Python objects, usually strings. That is right for **contract**. It is a bug for anything you intend to do arithmetic on. When a column you expected to be numeric shows as **object**, one value in it is not a number, and you will meet exactly that tomorrow.

Selecting: loc and iloc

Two accessors, and mixing them up is the most common pandas error. **.loc** selects by *label*, **.iloc** by *integer position*.

```
1 import pandas as pd
2
3 df = pd.DataFrame({
```

```

4     'tenure': [3, 40, 12, 65],
5     'charges': [70.0, 55.5, 99.9, 21.0],
6     'churned': [1, 0, 1, 0],
7 }, index=['a', 'b', 'c', 'd'])
8
9 print(df.loc['b', 'charges'])          # label row, label column
10 print(df.iloc[1, 1])                  # second row, second column
11 print(df.loc[df['churned'] == 1, ['tenure', 'charges']])

```

```

55.5
55.5
   tenure  charges
a        3    70.0
c       12    99.9

```

That last line is the pattern you will use every day: a boolean condition picks the rows, a list of names picks the columns.

Grouping answers most questions

“Does churn differ by contract type?” is a group-by.

```

1 import pandas as pd
2
3 df = pd.DataFrame({
4     'contract': ['Month-to-month'] * 4 + ['Two year'] * 4,
5     'churned': [1, 1, 0, 1, 0, 0, 0, 1],
6     'charges': [70, 82, 65, 90, 55, 40, 61, 58],
7 })
8
9 summary = df.groupby('contract').agg(
10     customers=('churned', 'size'),
11     churn_rate=('churned', 'mean'),
12     avg_charge=('charges', 'mean'),
13 )
14 print(summary.round(3))

```

```

      customers  churn_rate  avg_charge
contract
Month-to-month      4      0.75      76.75
Two year            4      0.25      53.50

```

Named aggregation, `new_name=('column', 'function')`, is worth learning properly. It produces flat, readable column names instead of the nested mess the older syntax gives you.

Make pandas print the whole frame

By default pandas folds wide tables and replaces the middle columns with an ellipsis, so the output you see depends on how wide your terminal happens to be. Set these two options once at the top of every notebook. Every printed result in this course was produced with them set, so your output will match what these pages show.

```

1 import pandas as pd
2

```

```

3 pd.set_option('display.width', 100)
4 pd.set_option('display.max_columns', 25)
5 print('pandas', pd.__version__)

```

pandas 2.2.3

Build the dataset for the rest of the course

Public churn datasets move, vanish behind logins, or change shape. So you will generate yours. Save this as `make_dataset.py` next to your notebook and run it once. It is deterministic: the same seed gives everyone the same file, so the numbers in these pages are the numbers you will see.

```

1 import numpy as np
2 import pandas as pd
3
4 rng = np.random.default_rng(42)
5 n = 3000
6
7 contract = rng.choice(['Month-to-month', 'One year', 'Two year'],
8                       size=n, p=[0.55, 0.25, 0.20])
9 internet = rng.choice(['Fibre optic', 'DSL', 'No internet'], size=n,
10                      p=[0.44, 0.34, 0.22])
11 payment = rng.choice(['Electronic check', 'Mailed check',
12                      'Bank transfer', 'Credit card'], size=n,
13                      p=[0.34, 0.23, 0.22, 0.21])
14 dependents = rng.choice(['Yes', 'No'], size=n, p=[0.3, 0.7])
15
16 # Two-year customers have by definition been around longer.
17 base = np.where(contract == 'Two year', 42,
18                np.where(contract == 'One year', 28, 16))
19 tenure = np.clip(rng.gamma(2.2, base / 2.2), 1, 72).round().astype(int)
20
21 price = np.where(internet == 'Fibre optic', 79,
22                np.where(internet == 'DSL', 55, 21))
23 monthly = np.clip(rng.normal(price, 9), 15, 125).round(2)
24 support = rng.poisson(np.where(internet == 'Fibre optic', 1.6, 0.9))
25
26 # The rule a model is meant to recover.
27 logit = (-2.8
28         + 1.55 * (contract == 'Month-to-month')
29         - 0.85 * (contract == 'Two year')
30         - 0.055 * tenure
31         + 0.021 * monthly
32         + 0.30 * support
33         + 0.42 * (payment == 'Electronic check')
34         - 0.25 * (dependents == 'Yes')
35         + rng.normal(0, 0.45, n))
36 churned = (1 / (1 + np.exp(-logit))) > rng.uniform(size=n).astype(int)
37
38 total = (monthly * tenure * rng.uniform(0.94, 1.06, n)).round(2)
39 start = np.datetime64('2019-01', 'M')
40 signup = (start + (72 - tenure).astype('timedelta64[M]')).astype('datetime64[D]')

```

```

41 df = pd.DataFrame({
42     'customer_id': ['C%05d' % i for i in range(1, n + 1)],
43     'signup_date': pd.to_datetime(signup).astype(str),
44     'tenure_months': tenure,
45     'contract': contract,
46     'internet_service': internet,
47     'payment_method': payment,
48     'has_dependents': dependents,
49     'support_calls': support,
50     'monthly_charges': monthly,
51     'total_charges': total,
52     'churned': churned,
53 })
54
55
56 # --- and now spoil it, the way collection genuinely spoils data ---
57 fresh = df.index[df['tenure_months'] <= 1]
58 df.loc[fresh, 'total_charges'] = np.nan
59 # A space, not an empty string: pandas reads '' as missing, which would
60 # hide the very defect this is meant to reproduce.
61 df['total_charges'] = df['total_charges'].map(
62     lambda v: ' ' if pd.isna(v) else ('%.2f' % v))
63
64 drift = rng.choice(df.index, size=int(n * 0.07), replace=False)
65 df.loc[drift, 'contract'] = df.loc[drift, 'contract'].map(
66     lambda s: rng.choice([s.upper(), s.lower(), ' ' + s, s + ' ']))
67
68 for col, frac in (('has_dependents', 0.06), ('monthly_charges', 0.055)):
69     gaps = rng.choice(df.index, size=int(n * frac), replace=False)
70     df.loc[gaps, col] = np.nan
71
72 df.loc[df.index[7], 'support_calls'] = 97
73 df = pd.concat([df, df.iloc[500:638]], ignore_index=True)
74 df = df.sample(frac=1, random_state=42).reset_index(drop=True)
75
76 df.to_csv('customers.csv', index=False)
77 print('customers.csv:', df.shape, 'churn rate %.1f%%' % (df['churned'].mean() *
    ↪ 100))

```

```
customers.csv: (3138, 11) churn rate 26.8%
```

The mess is on purpose

Blank totals, four spellings of one contract type, 138 duplicated rows, one customer with 97 support calls and gaps in two columns. Every one of those is a defect you will meet in real extracts, and every one teaches something specific in the weeks ahead. Cleaning it is tomorrow.

Day 3 takeaway

Day 4 First Contact With Messy Data

Dtypes that lie, categories that multiply, duplicates and gaps

You have a CSV. Before any modelling, you need to know what is actually in it, and the answer is never what the column names promise.

First contact

```
1 import pandas as pd
2
3 df = pd.read_csv('customers.csv')
4 print(df.shape)
5 print(df.dtypes)
```

```
(3138, 11)
customer_id      object
signup_date      object
tenure_months    int64
contract         object
internet_service object
payment_method   object
has_dependents   object
support_calls    int64
monthly_charges  float64
total_charges    object
churned          int64
dtype: object
```

Two things are already wrong. `signup_date` is text, not a date, expected, since CSV has no date type. But `total_charges` is text too, and that is a defect: it should be a number.

Why a numeric column arrives as text

```
1 import pandas as pd
2 df = pd.read_csv('customers.csv')
3
4 numeric = pd.to_numeric(df['total_charges'], errors='coerce')
5 bad = df.loc[numeric.isna(), ['customer_id', 'tenure_months', 'total_charges']]
6 print('unconvertible rows:', len(bad))
7 print(bad.head())
```

```
unconvertible rows: 15
   customer_id  tenure_months  total_charges
220      C00175              1             
531      C02127              1             
624      C02779              1             
749      C02411              1             
870      C00436              1              
```

Fifteen rows, every one of them a customer whose tenure is one month. They have not been billed yet, so the field holds a single space, and one non-numeric value among three thousand numbers forces the whole column to text. `errors='coerce'` turns what cannot be parsed into NaN instead of raising, which is what you want here.

⚠ Watch Out

Do not fill those with zero

Zero means “billed nothing”. These customers mean “not billed yet”. Writing zero invents a fact and drags the column mean down. Leave them missing and decide deliberately in week 8: missingness that carries meaning is a feature, not a nuisance.

Categories that are not as tidy as they look

```
1 import pandas as pd
2 df = pd.read_csv('customers.csv')
3
4 print('distinct contract values:', df['contract'].nunique())
5 print(df['contract'].value_counts())
```

```
distinct contract values: 15
contract
Month-to-month      1611
One year            710
Two year            592
  Month-to-month      43
MONTH-TO-MONTH       31
Month-to-month       25
month-to-month       22
TWO YEAR             18
  One year            18
One year             17
ONE YEAR             14
Two year             12
  Two year            10
two year              8
one year              7
Name: count, dtype: int64
```

Fifteen levels of a three-level variable. Group by this column and you get fifteen groups; one-hot encode it and you get fifteen columns, twelve of which are near-empty noise. Normalising is two calls:

```
1 import pandas as pd
2 df = pd.read_csv('customers.csv')
3
4 df['contract'] = df['contract'].str.strip().str.title()
5 print(df['contract'].value_counts())
```

```
contract
Month-To-Month      1732
One Year             766
Two Year             640
Name: count, dtype: int64
```

`.str.strip()` removes the padding, `.str.title()` puts casing back in one form. Three levels, as intended.

Duplicates

```

1 import pandas as pd
2 df = pd.read_csv('customers.csv')
3
4 print('fully duplicated rows:', df.duplicated().sum())
5 print('repeated customer ids:', df['customer_id'].duplicated().sum())
6
7 clean = df.drop_duplicates()
8 print('rows before', len(df), '→ after', len(clean))

```

```

fully duplicated rows: 138
repeated customer ids: 138
rows before 3138 → after 3000

```

⚠ Watch Out

Duplicates leak across a train/test split

138 rows appear twice. Split the data afterwards and some customers land in *both* halves. The model then gets tested on rows it trained on, and the test score reports memorisation as skill. De-duplicate before you split, never after.

Missing values

```

1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates()
3 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
4
5 missing = df.isna().sum()
6 print(missing[missing > 0])
7 print('\nrows with at least one gap:', df.isna().any(axis=1).sum())

```

```

has_dependents      180
monthly_charges     165
total_charges        15
dtype: int64

rows with at least one gap: 346

```

346 rows of 3000. Dropping them costs you eleven percent of the data and, worse, assumes the gaps are random. If customers who decline to state dependents churn differently, dropping them biases the model. Week 8 handles this properly with imputation inside a pipeline.

The target

```

1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates()
3
4 print(df['churned'].value_counts())
5 print('churn rate: %.1f%%' % (df['churned'].mean() * 100))

```

```
churned
0    2196
1     804
Name: count, dtype: int64
churn rate: 26.8%
```

Write that number down

26.8 percent churn means a model that predicts “nobody churns” and never gets anything right is **73.2 percent accurate**. That is your floor. Any accuracy you report from now on has to be compared against it, and week 7 will show you why accuracy is the wrong headline metric here at all.

Day 4 takeaway

Day 5 Seeing the Data Before Modelling It

matplotlib, seaborn, and the charts that actually decide things

Summary statistics hide shape. Two columns with the same mean and standard deviation can look nothing alike, and the difference decides which model will work. Plot before you model.

matplotlib, and the one thing to learn about it

matplotlib has two interfaces. The `plt.plot()` style keeps hidden state about which figure is current and breaks the moment you draw two things. Use the explicit style: ask for a figure and axes, then draw on the axes.

```
1 import matplotlib
2 matplotlib.use('Agg')          # no window needed; write straight to a file
3 import matplotlib.pyplot as plt
4 import pandas as pd
5
6 df = pd.read_csv('customers.csv').drop_duplicates()
7
8 fig, ax = plt.subplots(figsize=(7, 4))
9 ax.hist(df['tenure_months'], bins=36, color='#2563eb', edgecolor='white')
10 ax.set_xlabel('Tenure (months)')
11 ax.set_ylabel('Customers')
12 ax.set_title('How long customers have been with us')
13 fig.tight_layout()
14 fig.savefig('tenure.png', dpi=120)
15 plt.close(fig)
16 print('written')
```

```
written
```

Agg means headless

`matplotlib.use('Agg')` selects a backend that renders to a file with no display attached. In a notebook you do not need it. In a script, on a server, or in a scheduled job you do. Without it, matplotlib tries to open a window and the job hangs or crashes.

Comparing a distribution across groups

The question that matters is never “what does tenure look like”, it is “does tenure look different for the people who left”. Overlay the two.

```

1 import matplotlib
2 matplotlib.use('Agg')
3 import matplotlib.pyplot as plt
4 import pandas as pd
5
6 df = pd.read_csv('customers.csv').drop_duplicates()
7 stayed = df.loc[df['churned'] == 0, 'tenure_months']
8 left = df.loc[df['churned'] == 1, 'tenure_months']
9
10 fig, ax = plt.subplots(figsize=(7, 4))
11 ax.hist([stayed, left], bins=24, density=True,
12         label=['Stayed', 'Churned'], color=['#2563eb', '#dc2626'])
13 ax.set_xlabel('Tenure (months)')
14 ax.set_ylabel('Share of group')
15 ax.legend()
16 fig.tight_layout()
17 fig.savefig('tenure_by_churn.png', dpi=120)
18 plt.close(fig)
19
20 print(df.groupby('churned')['tenure_months'].describe().round(1))

```

	count	mean	std	min	25%	50%	75%	max
churned								
0	2196.0	26.9	18.7	1.0	13.0	22.0	37.0	72.0
1	804.0	14.1	9.1	1.0	8.0	12.0	18.0	72.0

density=True, not counts

2196 customers stayed and 804 left. Plot raw counts and the churned bars are dwarfed regardless of shape. `density=True` scales each group to its own area, so you compare shapes rather than group sizes. Getting this wrong is the most common misleading chart in churn analysis.

The medians are 22 months against 12. That gap is the single strongest signal in the dataset, and any model that fails to find it is broken.

seaborn for the statistical plots

seaborn sits on matplotlib and knows about DataFrames, so the plots you want during exploration are one line each.

```

1 import matplotlib
2 matplotlib.use('Agg')
3 import matplotlib.pyplot as plt

```

```

4 import pandas as pd
5 import seaborn as sns
6
7 df = pd.read_csv('customers.csv').drop_duplicates()
8 df['contract'] = df['contract'].str.strip().str.title()
9
10 fig, ax = plt.subplots(figsize=(7, 4))
11 sns.barplot(data=df, x='contract', y='churned', errorbar=('ci', 95), ax=ax)
12 ax.set_ylabel('Churn rate')
13 ax.set_xlabel('')
14 fig.tight_layout()
15 fig.savefig('churn_by_contract.png', dpi=120)
16 plt.close(fig)
17
18 print(df.groupby('contract')['churned'].agg(['size', 'mean']).round(3))

```

	size	mean
contract		
Month-To-Month	1660	0.414
One Year	731	0.137
Two Year	609	0.028

Because `churned` is 0 or 1, its mean *is* the churn rate, so a bar chart of the mean is a bar chart of the rate. 41 percent against under 3 percent: contract type matters enormously.

Correlation, and what it will not tell you

```

1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates()
3 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
4
5 num = ['tenure_months', 'monthly_charges', 'total_charges',
6        'support_calls', 'churned']
7 print(df[num].corr().round(2))

```

	tenure_months	monthly_charges	total_charges	support_calls	churned
tenure_months	1.00	0.00	0.83	0.03	-0.32
monthly_charges	0.00	1.00	0.45	0.15	0.21
total_charges	0.83	0.45	1.00	0.10	-0.21
support_calls	0.03	0.15	0.10	1.00	0.08
churned	-0.32	0.21	-0.21	0.08	1.00

⚠ Watch Out

0.83 between tenure and total charges

Of course: total charges is roughly monthly charges multiplied by tenure. The column is nearly a restatement of two others. Correlations that high between features cause unstable coefficients in linear models, a problem called multicollinearity that week 4 addresses. Note it now; do not delete anything yet.

Also note that correlation only sees straight-line relationships between numbers. It says nothing

about `contract`, which the bar chart just showed is the strongest predictor you have. A correlation matrix is a starting point, never a feature-selection method.

Day 5 takeaway

Day 6 The scikit-learn Contract

fit, predict, transform, and the rule about fitting on test data

scikit-learn is large, but it is not complicated, because nearly every object in it obeys the same three-method contract. Learn the contract once and you can use a model you have never heard of.

The contract

Method	Who has it	What it does
<code>fit(X, y)</code>	Everything	Learn from data. Returns the object itself.
<code>predict(X)</code>	Models	Produce an answer for new rows.
<code>transform(X)</code>	Preprocessors	Produce a modified version of the input.
<code>fit_transform(X)</code>	Preprocessors	Both, in one call. Training data only.
<code>predict_proba(X)</code>	Most classifiers	Class probabilities rather than a hard label.

The contract in action

```

1 from sklearn.linear_model import LinearRegression
2 from sklearn.ensemble import RandomForestRegressor
3 from sklearn.datasets import make_regression
4
5 X, y = make_regression(n_samples=200, n_features=3, noise=12, random_state=0)
6
7 # Identical calls, entirely different algorithms.
8 for model in [LinearRegression(), RandomForestRegressor(random_state=0)]:
9     model.fit(X, y)
10    pred = model.predict(X[:3])
11    print(f'{type(model).__name__:24s} {pred.round(1)}')
```

```

LinearRegression      [ 21.3 -94.4 204.5]
RandomForestRegressor [ 17.4 -87.2 195.7]
```

Learned state lives on the object, with a trailing underscore

A convention worth knowing: attributes that exist only after `fit` end in an underscore. If you see `coef_` or `n_features_in_`, that value was learned from data. Anything without the underscore is a setting you chose.

```

1 from sklearn.linear_model import LinearRegression
2 from sklearn.datasets import make_regression
3
4 X, y = make_regression(n_samples=200, n_features=3, noise=12, random_state=0)
5 model = LinearRegression()
6
7 print('before fit:', [a for a in dir(model) if a.endswith('_')
8                       and not a.startswith('_')])
9 model.fit(X, y)
10 print('after fit: coef_ =', model.coef_.round(2))
11 print('            intercept_ = %.2f' % model.intercept_)
12 print('            n_features_in_ =', model.n_features_in_)

```

```

before fit: []
after fit: coef_ = [25.31 60.45 57.45]
            intercept_ = -0.33
            n_features_in_ = 3

```

Transformers, and the rule about fitting them

A scaler learns something from the data, the mean and standard deviation of each column, and that makes it dangerous. It must learn from training data only.

```

1 import numpy as np
2 from sklearn.preprocessing import StandardScaler
3
4 train = np.array([[10.0], [20.0], [30.0], [40.0]])
5 test = np.array([[50.0], [60.0]])
6
7 scaler = StandardScaler()
8 train_scaled = scaler.fit_transform(train) # learn AND apply
9 test_scaled = scaler.transform(test)      # apply only
10
11 print('learned mean %.1f, scale %.2f' % (scaler.mean_[0], scaler.scale_[0]))
12 print('train:', train_scaled.ravel().round(2))
13 print('test: ', test_scaled.ravel().round(2))

```

```

learned mean 25.0, scale 11.18
train: [-1.34 -0.45  0.45  1.34]
test:  [2.24 3.13]

```

Watch Out

fit_transform on test data is a leak

Call `fit_transform` on the test set and the scaler learns the test set's mean. Information about data the model is not supposed to have seen flows into the features, the test score improves, and it is a lie. `fit_transform` on train, `transform` on test, every time. Week 8's pipelines make it impossible to get backwards, which is exactly why they exist.

Reading the documentation

Every estimator's docstring lists its hyperparameters with defaults and its learned attributes. You do not need to memorise them:

```
1 from sklearn.ensemble import RandomForestClassifier
2
3 params = RandomForestClassifier().get_params()
4 for k in ['n_estimators', 'max_depth', 'min_samples_leaf', 'random_state']:
5     print(f'{k:20s} {params[k]}')
```

```
n_estimators      100
max_depth         None
min_samples_leaf   1
random_state      None
```

Set random_state on everything

`random_state=None` means results change between runs. For anything you intend to compare, report or debug, pass an explicit integer. Two models whose scores differ by 0.4 percent are indistinguishable if you never fixed the seed.

Day 6 takeaway

Day 7 Your First End-to-End Model

Clean, split, pipeline, fit, and measure against an honest baseline

Everything so far, assembled. Today you build a working churn model, measure it honestly, and see why the obvious metric misleads.

Step 1: load and clean

```
1 import pandas as pd
2
3 df = pd.read_csv('customers.csv').drop_duplicates().copy()
4 df['contract'] = df['contract'].str.strip().str.title()
5 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
6
7 print(df.shape, '| churn rate %.3f' % df['churned'].mean())
```

```
(3000, 11) | churn rate 0.268
```

Step 2: split before anything else

The split comes first. Any decision you make after looking at the test set, which columns to keep, how to fill gaps, which model to use, contaminates it.

```
1 import pandas as pd
```

```

2 from sklearn.model_selection import train_test_split
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 num_cols = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
9 cat_cols = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10
11 X = df[num_cols + cat_cols]
12 y = df['churned']
13
14 X_train, X_test, y_train, y_test = train_test_split(
15     X, y, test_size=0.25, stratify=y, random_state=42)
16
17 print('train', X_train.shape, 'churn %.3f' % y_train.mean())
18 print('test ', X_test.shape, 'churn %.3f' % y_test.mean())

```

```

train (2250, 8) churn 0.268
test  (750, 8) churn 0.268

```

stratify=y

Without it, a random split can hand you a test set with a noticeably different churn rate than the training set, and the score you get reflects that accident rather than the model. `stratify` keeps the class proportions identical in both halves. Always pass it for classification.

Step 3: one object that preprocesses and predicts

Numeric columns need their gaps filled and their scales evened out. Categorical columns need turning into numbers. A `ColumnTransformer` applies different treatment to different columns; a `Pipeline` chains that to the model so the whole thing behaves like a single estimator.

```

1 import pandas as pd
2 from sklearn.model_selection import train_test_split
3 from sklearn.pipeline import Pipeline
4 from sklearn.compose import ColumnTransformer
5 from sklearn.impute import SimpleImputer
6 from sklearn.preprocessing import StandardScaler, OneHotEncoder
7 from sklearn.linear_model import LogisticRegression
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12 num_cols = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
13 cat_cols = ['contract', 'internet_service', 'payment_method', 'has_dependents']
14 X, y = df[num_cols + cat_cols], df['churned']
15 X_train, X_test, y_train, y_test = train_test_split(
16     X, y, test_size=0.25, stratify=y, random_state=42)
17
18 numeric = Pipeline([
19     ('impute', SimpleImputer(strategy='median')),
20     ('scale', StandardScaler()),

```

```

21 ])
22 categorical = Pipeline([
23     ('impute', SimpleImputer(strategy='most_frequent')),
24     ('encode', OneHotEncoder(handle_unknown='ignore')),
25 ])
26
27 preprocess = ColumnTransformer([
28     ('num', numeric, num_cols),
29     ('cat', categorical, cat_cols),
30 ])
31
32 model = Pipeline([
33     ('prep', preprocess),
34     ('clf', LogisticRegression(max_iter=1000, random_state=42)),
35 ])
36
37 model.fit(X_train, y_train)
38 print('fitted: %.3f' % model.score(X_test, y_test))

```

```
fitted: 0.780
```

Twelve lines of setup and one fit. Everything inside learns from the training data only, because the pipeline passes the split through in the right order for you.

Step 4: measure it against the floor

```

1 import pandas as pd
2 from sklearn.model_selection import train_test_split
3 from sklearn.dummy import DummyClassifier
4 from sklearn.pipeline import Pipeline
5 from sklearn.compose import ColumnTransformer
6 from sklearn.impute import SimpleImputer
7 from sklearn.preprocessing import StandardScaler, OneHotEncoder
8 from sklearn.linear_model import LogisticRegression
9 from sklearn.metrics import accuracy_score, roc_auc_score, confusion_matrix
10
11 df = pd.read_csv('customers.csv').drop_duplicates().copy()
12 df['contract'] = df['contract'].str.strip().str.title()
13 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
14 num_cols = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
15 cat_cols = ['contract', 'internet_service', 'payment_method', 'has_dependents']
16 X, y = df[num_cols + cat_cols], df['churned']
17 X_train, X_test, y_train, y_test = train_test_split(
18     X, y, test_size=0.25, stratify=y, random_state=42)
19
20 preprocess = ColumnTransformer([
21     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
22                     ('s', StandardScaler())]), num_cols),
23     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
24                     ('e', OneHotEncoder(handle_unknown='ignore'))]), cat_cols),
25 ])
26
27 model = Pipeline([('prep', preprocess),
28                  ('clf', LogisticRegression(max_iter=1000, random_state=42))])

```

```

28 model.fit(X_train, y_train)
29
30 baseline = DummyClassifier(strategy='most_frequent').fit(X_train, y_train)
31
32 pred = model.predict(X_test)
33 proba = model.predict_proba(X_test)[: , 1]
34
35 print('baseline accuracy %.3f' % accuracy_score(y_test, baseline.predict(X_test)))
36 print('model accuracy    %.3f' % accuracy_score(y_test, pred))
37 print('model ROC AUC     %.3f' % roc_auc_score(y_test, proba))
38 print('\nconfusion matrix (rows = truth, cols = prediction)')
39 print(confusion_matrix(y_test, pred))

```

```

baseline accuracy 0.732
model accuracy    0.780
model ROC AUC     0.818

confusion matrix (rows = truth, cols = prediction)
[[499  50]
 [115  86]]

```

Watch Out

78 percent sounds good until you read the matrix

Predicting “nobody churns” scores 73.2 percent. Our model scores 78.0, just under five points of real improvement, not the triumph the raw number suggests. And the matrix shows where it goes wrong: of 201 customers who actually churned, it caught 86 and missed 115. It finds fewer than half of the people you most wanted to find.

This is not a broken model. An AUC of 0.818 means it ranks customers by risk genuinely well. It is a badly chosen *threshold*. The model outputs a probability, and something has to decide where to cut. The default cut is 0.5, which is almost never the right choice when one class is rarer than the other. Week 7 is about fixing exactly this.

Step 5: what did it learn

```

1 import pandas as pd, numpy as np
2 from sklearn.model_selection import train_test_split
3 from sklearn.pipeline import Pipeline
4 from sklearn.compose import ColumnTransformer
5 from sklearn.impute import SimpleImputer
6 from sklearn.preprocessing import StandardScaler, OneHotEncoder
7 from sklearn.linear_model import LogisticRegression
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12 num_cols = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
13 cat_cols = ['contract', 'internet_service', 'payment_method', 'has_dependents']
14 X, y = df[num_cols + cat_cols], df['churned']
15 X_train, X_test, y_train, y_test = train_test_split(

```

```

16 X, y, test_size=0.25, stratify=y, random_state=42)
17 preprocess = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                       ('s', StandardScaler())]), num_cols),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                       ('e', OneHotEncoder(handle_unknown='ignore'))]), cat_cols),
22 ])
23 model = Pipeline([('prep', preprocess),
24                   ('clf', LogisticRegression(max_iter=1000, random_state=42))])
25 model.fit(X_train, y_train)
26
27 names = model.named_steps['prep'].get_feature_names_out()
28 coefs = model.named_steps['clf'].coef_[0]
29 effect = pd.Series(coefs, index=names).sort_values()
30
31 print('pushes towards staying:')
32 print(effect.head(3).round(2))
33 print('\npushes towards churning:')
34 print(effect.tail(3).round(2))

```

```

pushes towards staying:
cat__contract_Two Year      -1.44
num__tenure_months         -0.81
cat__payment_method_Credit card -0.20
dtype: float64

pushes towards churning:
cat__payment_method_Electronic check    0.46
num__monthly_charges                   0.48
cat__contract_Month-To-Month           1.31
dtype: float64

```

Compare that against the generator you wrote on day 3: month-to-month contracts push churn up, two-year contracts and long tenure push it down, electronic-check payment pushes it up. The model recovered the rule that created the data. That is what a good fit looks like.

Your assignment

Swap `LogisticRegression` for `RandomForestClassifier(n_estimators=300, random_state=42)`, a one-line change, because both obey the same contract. Record accuracy and AUC. You will find the forest scores *worse* on AUC. Write down why you think that is; week 6 gives the answer, and it is not that you did something wrong.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. What makes a problem *supervised*?
 - a) The data is clean
 - b) Every training row carries the answer you want to predict
 - c) A human checks the output
 - d) The model is trained on a server
2. Why is a NumPy operation on a whole array preferred to a Python loop over it?
 - a) Loops cannot handle floats
 - b) It uses less memory
 - c) The loop is written in C
 - d) The operation runs as compiled code over contiguous memory instead of one interpreted step per element
3. `total_charges` arrives as an `object` column rather than a number. The most likely cause is:
 - a) pandas cannot read decimals
 - b) The column has too many rows
 - c) The file is corrupt
 - d) Some rows hold a non-numeric value such as a blank
4. In the scikit-learn contract, what does a trailing underscore on an attribute such as `mean_` signify?
 - a) It was learned from data during `fit`
 - b) It is deprecated
 - c) It is a hyperparameter you set
 - d) It is private
5. Why must `fit` never be called on the test set?
 - a) It is slower
 - b) The test set is smaller
 - c) The statistics it learns would carry information from data you are pretending not to have seen
 - d) scikit-learn raises an error
6. A churn model reports 74% accuracy. Before celebrating you should first check:
 - a) How long it took to fit
 - b) The learning rate
 - c) What accuracy you get by always predicting the majority class
 - d) The number of features
7. What does a `Pipeline` actually buy you?

- a) Automatic hyperparameter tuning
- b) Prettier code
- c) Faster fitting
- d) Every preprocessing step is refitted inside each cross-validation fold, on that fold's training rows only

8. `df.groupby('contract')['churned'].mean()` returns:

- a) The churn rate within each contract type
- b) The mean contract length
- c) A count of contract types
- d) The number of churners per contract type

Answers: 1b, 2d, 3d, 4a, 5c, 6c, 7d, 8a. Anything you got wrong points at the day to re-read, not at a gap in ability.

Data Wrangling and Exploratory Analysis

Week 2

Load, reshape, clean and profile data properly, the work that decides whether a model has anything to learn from.

OBJECTIVES

Most of the difference between a model that works and one that does not is made before any model is fitted. This week covers the loading traps that corrupt data silently, the reshaping and joining that gets it into modelling form without duplicating your target, the three mechanisms behind missing values and what each demands, and how to turn all of it into a report you run on every new extract.

Day 1 Loading Data That Fights Back

Encodings, dtypes, sentinel values and the missing-value list you did not know about

Week 1 read a CSV in one line and it worked. That is the exception. Real extracts arrive with the wrong encoding, dates as text, numbers with currency symbols, and a sentinel value someone chose in 1998 to mean “unknown”. Loading is where most silent errors enter.

pandas has opinions about what missing means

You met this in week 1 without knowing it. `read_csv` converts a list of strings to NaN automatically, and that list is longer than you would guess.

```
1 import pandas as pd
2 from pandas.io.parsers.readers import STR_NA_VALUES
3
4 print(sorted(STR_NA_VALUES))
```

```
['', '#N/A', '#N/A N/A', '#NA', '-1.#IND', '-1.#QNAN', '-NaN', '-nan', '1.#IND', '1.#QNAN',
  ↳ '<NA>', 'N/A', 'NA', 'NULL', 'NaN', 'None', 'n/a', 'nan', 'null']
```

⚠ Watch Out

'None' is on that list

A category legitimately called `None`: a customer with no internet service, a product with no discount, is read as a missing value. The column silently gains hundreds of gaps that were never in the file. This is a real defect that appeared while building this course's dataset,

and it is why the category is spelled No internet.

Two arguments control it. `keep_default_na=False` switches the whole list off; `na_values=` adds your own sentinels.

```

1 import pandas as pd
2 import io as _io
3
4 raw = 'service,code\nNone,1\nDSL,-999\nFibre,2\n'
5
6 default = pd.read_csv(_io.StringIO(raw))
7 print('default:')
8 print(default)
9
10 explicit = pd.read_csv(_io.StringIO(raw), keep_default_na=False,
11                        na_values=['-999'])
12 print('\nkeeping None as text, treating -999 as missing:')
13 print(explicit)

```

```

default:
  service  code
0     NaN    1
1     DSL  -999
2  Fibre    2

keeping None as text, treating -999 as missing:
  service  code
0    None  1.0
1     DSL  NaN
2  Fibre  2.0

```

Declare the types you expect

Letting pandas guess the dtype of every column is convenient and occasionally wrong. An identifier of digits becomes an integer and loses its leading zeros; a postcode becomes a float. Declare what matters.

```

1 import pandas as pd
2 import io as _io
3
4 raw = 'account,postcode\n007831,01234\n004120,00987\n'
5
6 print('guessed:')
7 print(pd.read_csv(_io.StringIO(raw)))
8
9 print('\ndeclared:')
10 print(pd.read_csv(_io.StringIO(raw), dtype={'account': str, 'postcode': str}))

```

```

guessed:
  account  postcode
0    7831    1234
1   4120     987

```

```

declared:
  account postcode
0  007831      01234
1  004120      00987

```

The leading zeros are gone in the first version and there is no way to get them back afterwards. Loss at load time is unrecoverable, which is why loading deserves more care than it usually gets.

Parse dates at load time

```

1 import pandas as pd
2
3 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
4 print(df['signup_date'].dtype)
5 print(df['signup_date'].min(), 'to', df['signup_date'].max())

```

```

datetime64[ns]
2019-01-01 00:00:00 to 2024-12-01 00:00:00

```

When the format is ambiguous

03/04/2024 is the third of April in Britain and the fourth of March in America, and pandas will pick one. If your source is not ISO format, pass the format explicitly: `pd.to_datetime(col, format='%d/%m/%Y')`. Guessing wrong produces dates that are valid, plausible and wrong for eleven months of the year.

Memory, and when it starts to matter

```

1 import pandas as pd
2
3 df = pd.read_csv('customers.csv')
4 before = df.memory_usage(deep=True).sum()
5
6 for col in ['contract', 'internet_service', 'payment_method']:
7     df[col] = df[col].astype('category')
8
9 after = df.memory_usage(deep=True).sum()
10 print('before %6.1f KB' % (before / 1024))
11 print('after  %6.1f KB' % (after / 1024))
12 print('saved  %5.1f%%' % (100 * (1 - after / before)))

```

```

before 1322.7 KB
after   782.4 KB
saved   40.8%

```

A category column stores each distinct value once and an integer code per row. With three distinct values across three thousand rows the saving is large, and it grows with the row count. On a dataset that fits comfortably in memory this is housekeeping; on one that does not, it is the difference between working and not.

Day 1 takeaway**Day 2 Reshaping and Joining**

Long against wide, and how a careless join duplicates your target

Data almost never arrives in the shape a model wants. Models want one row per observation and one column per feature. Sources give you one row per event, or one column per month, or the information split across three tables.

Long and wide

```

1 import pandas as pd
2
3 wide = pd.DataFrame({
4     'customer_id': ['C1', 'C2'],
5     'jan': [40.0, 55.0],
6     'feb': [42.0, 55.0],
7     'mar': [41.5, 60.0],
8 })
9 print('wide -- readable, not modellable')
10 print(wide)
11
12 long = wide.melt(id_vars='customer_id', var_name='month',
13                 value_name='charge')
14 print('\nlong -- one row per observation')
15 print(long)

```

```

wide -- readable, not modellable
  customer_id  jan  feb  mar
0          C1  40.0  42.0  41.5
1          C2  55.0  55.0  60.0

long -- one row per observation
  customer_id month  charge
0          C1   jan   40.0
1          C2   jan   55.0
2          C1   feb   42.0
3          C2   feb   55.0
4          C1   mar   41.5
5          C2   mar   60.0

```

`melt` goes wide to long, `pivot` goes back:

```

1 import pandas as pd
2
3 long = pd.DataFrame({
4     'customer_id': ['C1', 'C1', 'C1', 'C2', 'C2', 'C2'],
5     'month': ['jan', 'feb', 'mar'] * 2,
6     'charge': [40.0, 42.0, 41.5, 55.0, 55.0, 60.0],
7 })
8

```

```

9 back = long.pivot(index='customer_id', columns='month', values='charge')
10 print(back)
11
12 # Aggregating while pivoting, when there are duplicates per cell
13 summary = long.pivot_table(index='customer_id', values='charge',
14                             aggfunc=['mean', 'max'])
15 print('\n', summary)

```

```

month      feb  jan  mar
customer_id
C1         42.0  40.0  41.5
C2         55.0  55.0  60.0

              mean  max
              charge charge
customer_id
C1         41.166667  42.0
C2         56.666667  60.0

```

Which shape do you want

Long is right for storage and for plotting with seaborn. Wide is right for modelling, because a model needs a fixed set of columns. Converting long event logs into wide per-customer features is exactly what week 8 calls feature engineering.

Joining tables

The merge types are the same four you know from SQL, and the argument that matters most is how.

how	Keeps	Use when
inner	Rows matching in both	You need complete records only
left	All left rows, matched right ones	Enriching a base table, the usual choice
outer	Everything from both	Reconciling two sources
cross	Every combination	Building a grid; rarely what you want

```

1 import pandas as pd
2
3 customers = pd.DataFrame({
4     'customer_id': ['C1', 'C2', 'C3'],
5     'contract': ['Month-to-month', 'Two year', 'One year'],
6 })
7 tickets = pd.DataFrame({
8     'customer_id': ['C1', 'C1', 'C3', 'C9'],
9     'issue': ['billing', 'speed', 'billing', 'billing'],
10 })
11
12 left = customers.merge(tickets, on='customer_id', how='left')
13 print(left)
14 print('\nrows before %d, after %d' % (len(customers), len(left)))

```

	customer_id	contract	issue
0	C1	Month-to-month	billing
1	C1	Month-to-month	speed
2	C2	Two year	NaN
3	C3	One year	billing

rows before 3, after 4

⚠ Watch Out**A left join can multiply your rows**

C1 has two tickets, so C1 appears twice. Three customers became four rows. Join a one-per-customer table to a many-per-customer table and your target variable is now duplicated, which silently weights those customers more heavily in training. Aggregate the many side *first*, then join.

```

1 import pandas as pd
2
3 customers = pd.DataFrame({
4     'customer_id': ['C1', 'C2', 'C3'],
5     'contract': ['Month-to-month', 'Two year', 'One year'],
6 })
7 tickets = pd.DataFrame({
8     'customer_id': ['C1', 'C1', 'C3', 'C9'],
9     'issue': ['billing', 'speed', 'billing', 'billing'],
10 })
11
12 per_customer = tickets.groupby('customer_id').agg(
13     ticket_count=('issue', 'size'),
14     billing_tickets=('issue', lambda s: (s == 'billing').sum()),
15 )
16
17 safe = customers.merge(per_customer, on='customer_id', how='left')
18 safe[['ticket_count', 'billing_tickets']] = (
19     safe[['ticket_count', 'billing_tickets']].fillna(0).astype(int))
20 print(safe)
21 print('\nrows: %d -- unchanged' % len(safe))

```

	customer_id	contract	ticket_count	billing_tickets
0	C1	Month-to-month	2	1
1	C2	Two year	0	0
2	C3	One year	1	1

rows: 3 -- unchanged

Three rows in, three rows out, and two new features. That is the shape every join into a modelling table should have.

Validate the join instead of hoping

```

1 import pandas as pd

```

```

2 left_tbl = pd.DataFrame({'k': ['a', 'b', 'c'], 'x': [1, 2, 3]})
3 right_tbl = pd.DataFrame({'k': ['a', 'a', 'b'], 'y': [9, 8, 7]})
4
5
6 try:
7     left_tbl.merge(right_tbl, on='k', how='left', validate='one_to_one')
8 except Exception as e:
9     print(type(e).__name__ + ': ', e)

```

```
MergeError: Merge keys are not unique in right dataset; not a one-to-one merge
```

`validate` accepts `one_to_one`, `one_to_many`, `many_to_one` and `many_to_many`. Stating what you expect turns a silent row explosion into an immediate error.

Day 2 takeaway

Day 3 Missing Data, Properly

Why a value is missing decides how to fill it, and where to fit the imputer

Week 1 counted the gaps and moved on. Today you decide what to do about them, and the right answer depends on *why* they are missing.

Three mechanisms

Mechanism	Meaning	Example	Safe to impute?
MCAR	Missing completely at random	A sensor dropped packets	Yes
MAR	Missingness explained by other columns	Older customers skip the email field	Yes, using those columns
MNAR	Missingness depends on the missing value itself	High earners decline to state income	No, imputing biases it

You cannot test for MNAR from the data alone, which is the uncomfortable part. What you can do is check whether missingness relates to anything you *can* see, and record the fact that a value was missing.

Is the missingness informative?

```

1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates().copy()
3 df['contract'] = df['contract'].str.strip().str.title()
4 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
5
6 gap = df['has_dependents'].isna()
7 print('rows with the field blank: %d' % gap.sum())
8 print('churn rate when blank      : %.3f' % df.loc[gap, 'churned'].mean())
9 print('churn rate when present   : %.3f' % df.loc[~gap, 'churned'].mean())

```

```
rows with the field blank: 180
churn rate when blank      : 0.272
churn rate when present    : 0.268
```

Close enough to call it uninformative here, which is what you would expect, because the generator removed those values at random. Had the two rates differed sharply, the fact of the gap would itself be a feature worth keeping.

Keep the flag anyway

Adding an `was_missing` indicator column costs one bit per row and occasionally carries real signal: in medical data, a test that was not ordered tells you what the clinician was thinking. `SimpleImputer(add_indicator=True)` does it for you, inside the pipeline.

What each strategy actually does to the distribution

```
1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates().copy()
3 df['contract'] = df['contract'].str.strip().str.title()
4 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
5 from sklearn.impute import SimpleImputer
6 import numpy as np
7
8 col = df[['monthly_charges']]
9 truth = col.dropna()['monthly_charges']
10 print('observed : mean %.2f std %.2f n %d'
11       % (truth.mean(), truth.std(), len(truth)))
12
13 for strategy in ['mean', 'median', 'most_frequent']:
14     filled = SimpleImputer(strategy=strategy).fit_transform(col).ravel()
15     print('%-14s: mean %.2f std %.2f n %d'
16           % (strategy, filled.mean(), filled.std(), len(filled)))
```

```
observed : mean 59.41 std 23.47 n 2835
mean      : mean 59.41 std 22.81 n 3000
median    : mean 59.62 std 22.83 n 3000
most_frequent : mean 56.96 std 24.96 n 3000
```

Watch Out

Imputation always shrinks the variance

Every value filled with the mean or median sits exactly at the centre, so the spread of the column falls, 23.47 to 22.81 here. `most_frequent` is different: the mode is not the centre, so it can widen the spread instead, as it does above. With 5 percent missing either is a rounding error. With 40 percent it is a serious distortion, and any model that relies on the variance of that feature is being misled. Check how much you are filling before you decide it is fine.

Using the other columns

A customer's charge is not independent of their internet service. Filling from the group is better than filling from the whole column.

```
1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates().copy()
3 df['contract'] = df['contract'].str.strip().str.title()
4 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
5
6 overall = df['monthly_charges'].median()
7 by_service = df.groupby('internet_service')['monthly_charges'].median()
8 print('overall median   %.2f' % overall)
9 print(by_service.round(2))
10
11 smart = df['monthly_charges'].fillna(
12     df.groupby('internet_service')['monthly_charges'].transform('median'))
13 print('\nremaining gaps:', smart.isna().sum())
```

```
overall median   63.21
internet_service
DSL              54.97
Fibre optic     79.71
No internet     20.69
Name: monthly_charges, dtype: float64

remaining gaps: 0
```

Filling a fibre customer's charge with the overall median understates it by a wide margin. Grouping first respects a relationship you already know exists.

```
1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates().copy()
3 df['contract'] = df['contract'].str.strip().str.title()
4 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
5 from sklearn.impute import KNNImputer
6 from sklearn.preprocessing import StandardScaler
7 import numpy as np
8
9 num = df[['tenure_months', 'monthly_charges', 'support_calls']]
10 scaled = StandardScaler().fit_transform(num)
11 filled = KNNImputer(n_neighbors=5).fit_transform(scaled)
12
13 gaps = num['monthly_charges'].isna().values
14 back = StandardScaler().fit(num.dropna()).inverse_transform(filled)
15 print('imputed values, first five:')
16 print(np.round(back[gaps, 1][:5], 2))
```

```
imputed values, first five:
[55.45 54.23 54.45 81.7  55.19]
```

KNN imputation must be scaled first

It finds neighbours by distance, and distance is meaningless across columns measured in months and pounds. Scale, impute, and only then reverse the scaling. Getting this order wrong gives you neighbours chosen almost entirely by whichever column has the largest numbers.

The rule that matters more than the method

```

1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates().copy()
3 df['contract'] = df['contract'].str.strip().str.title()
4 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
5 from sklearn.model_selection import train_test_split
6 from sklearn.impute import SimpleImputer
7
8 X = df[['tenure_months', 'monthly_charges']]
9 y = df['churned']
10 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
11                                         random_state=42)
12
13 # Right: learn the median on train, apply it to both.
14 imp = SimpleImputer(strategy='median').fit(X_tr)
15 print('median learned from train only: %.2f' % imp.statistics_[1])
16
17 # Wrong: learn it from everything, including rows you will test on.
18 leaky = SimpleImputer(strategy='median').fit(X)
19 print('median learned from everything: %.2f' % leaky.statistics_[1])

```

```

median learned from train only: 63.05
median learned from everything: 63.21

```

The two numbers are close, and that is exactly what makes the mistake hard to catch. The leak is real regardless of its size, it grows with the proportion missing, and it always flatters the test score. Fit imputers on training data only, which pipelines enforce for you.

Day 3 takeaway**Day 4 Outliers, Skew and Robustness**

Finding extreme values, deciding what they mean, and which models care

An outlier is not a value that is large. It is a value that does not belong to the process you are modelling, and telling those apart requires knowing the domain, not just the numbers.

Find them three ways

```

1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates().copy()
3 df['contract'] = df['contract'].str.strip().str.title()
4 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')

```

```

5 import numpy as np
6
7 x = df['support_calls']
8
9 # 1. Standard deviations from the mean
10 z = (x - x.mean()) / x.std()
11 print('|z| > 3      : ', (z.abs() > 3).sum())
12
13 # 2. Interquartile range, the boxplot rule
14 q1, q3 = x.quantile([0.25, 0.75])
15 iqr = q3 - q1
16 print('outside 1.5 IQR:', ((x < q1 - 1.5 * iqr) | (x > q3 + 1.5 * iqr)).sum())
17
18 # 3. Just look at the extremes
19 print('largest values :', np.sort(x.unique())[-5:])

```

```

|z| > 3      : 1
outside 1.5 IQR: 8
largest values : [ 4  5  6  7 97]

```

⚠ Watch Out

The z-score method is broken by the thing it looks for

A z-score uses the mean and standard deviation, and one extreme value drags both. The outlier inflates the standard deviation, which shrinks its own z-score, which can push it under the threshold. The IQR rule uses quartiles, which a single extreme value cannot move. That is what *robust* means and it is why the boxplot rule is the safer default.

What to do about the one customer with 97 calls

```

1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates().copy()
3 df['contract'] = df['contract'].str.strip().str.title()
4 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
5
6 extreme = df.loc[df['support_calls'] > 20,
7                  ['customer_id', 'tenure_months', 'support_calls',
8                  'monthly_charges', 'churned']]
9 print(extreme)
10 print('\nsupport_calls: mean %.2f, median %.1f, max %d'
11       % (df['support_calls'].mean(), df['support_calls'].median(),
12          df['support_calls'].max()))

```

```

      customer_id  tenure_months  support_calls  monthly_charges  churned
710      C00008             48             97             76.42         0

```

```

support_calls: mean 1.27, median 1.0, max 97

```

One row. It is not a typo. A customer really can call ninety-seven times, and if anything they are the most interesting customer in the file. Deleting them removes a genuine, if rare, pattern. The options, in order of how much they assume:

1. **Leave it** and use a model that does not care. Trees split on order, not magnitude, so an extreme value changes nothing.
2. **Clip it** to a percentile, caps the influence without discarding the row.
3. **Transform** the column with a log, which compresses the long tail.
4. **Drop it**, only when you can show it is an error.

```

1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates().copy()
3 df['contract'] = df['contract'].str.strip().str.title()
4 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
5 import numpy as np
6
7 x = df['support_calls']
8 cap = x.quantile(0.99)
9
10 print('original : max %5.1f mean %.3f std %.3f' % (x.max(), x.mean(), x.std()))
11 print('clipped : max %5.1f mean %.3f std %.3f'
12       % (x.clip(upper=cap).max(), x.clip(upper=cap).mean(),
13          ↪ x.clip(upper=cap).std()))
14 print('log1p : max %5.1f mean %.3f std %.3f'
15       % (np.log1p(x).max(), np.log1p(x).mean(), np.log1p(x).std()))

```

```

original : max 97.0 mean 1.267 std 2.100
clipped  : max 5.0 mean 1.233 std 1.156
log1p    : max 4.6 mean 0.672 std 0.527

```

log1p, not log

`np.log1p(x)` computes $\log(1 + x)$, which is defined at zero. Plain `np.log` on a column containing zeros gives you `-inf`, and one infinity poisons every calculation downstream. Count columns almost always contain zeros.

Which models care

```

1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates().copy()
3 df['contract'] = df['contract'].str.strip().str.title()
4 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
5 from sklearn.model_selection import cross_val_score
6 from sklearn.pipeline import make_pipeline
7 from sklearn.preprocessing import StandardScaler
8 from sklearn.impute import SimpleImputer
9 from sklearn.linear_model import LogisticRegression
10 from sklearn.tree import DecisionTreeClassifier
11 import numpy as np
12
13 X = df[['tenure_months', 'support_calls', 'monthly_charges']]
14 y = df['churned']
15 X_capped = X.copy()
16 X_capped['support_calls'] = X_capped['support_calls'].clip(

```

```

17     upper=X_capped['support_calls'].quantile(0.99))
18
19 models = {
20     'logistic': make_pipeline(SimpleImputer(strategy='median'),
21                               StandardScaler(),
22                               LogisticRegression(max_iter=1000)),
23     'tree': make_pipeline(SimpleImputer(strategy='median'),
24                           DecisionTreeClassifier(max_depth=4,
25                                                  random_state=42)),
26 }
27 for name, m in models.items():
28     raw = cross_val_score(m, X, y, cv=5, scoring='roc_auc').mean()
29     cap = cross_val_score(m, X_capped, y, cv=5, scoring='roc_auc').mean()
30     print('%s raw %.4f capped %.4f change %.4f'
31           % (name, raw, cap, cap - raw))

```

```

logistic raw 0.7762 capped 0.7786 change +0.0024
tree      raw 0.7545 capped 0.7545 change +0.0000

```

The tree score is identical to four decimal places, because a tree only ever asks “is this value above the split point” and 97 answers that exactly as 20 does. The linear model moves, because after scaling that row sits many standard deviations out and pulls the fitted coefficient toward itself. The movement is small here, one contaminated row in three thousand, and that is the honest lesson: the direction is reliable, the size depends entirely on how much of your data is affected.

Skew, and why it matters for linear models

```

1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates().copy()
3 df['contract'] = df['contract'].str.strip().str.title()
4 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
5 import numpy as np
6
7 for col in ['tenure_months', 'monthly_charges', 'support_calls']:
8     s = df[col].dropna()
9     print('%-16s skew %6.2f log1p skew %6.2f'
10           % (col, s.skew(), np.log1p(s).skew()))

```

```

tenure_months    skew    1.22 log1p skew  -0.26
monthly_charges  skew   -0.47 log1p skew  -1.15
support_calls     skew   31.75 log1p skew   0.14

```

A skew near zero means symmetric. `support_calls` is extremely right-skewed at 31.75, and the log transform pulls it to 0.14. But look at `monthly_charges`: already near-symmetric at -0.47, and the log makes it distinctly *worse* at -1.15. A log is a tool for right-skewed columns, not a default. Linear models benefit from roughly symmetric inputs; tree models do not care either way.

Day 4 takeaway

Day 5 Dates and Time Features

Decomposition, elapsed time, cyclical encoding, and the split dates force

A date column is useless to a model as it stands. No algorithm can do anything with the integer 1704067200. What models can use is what the date *means*: how long ago, which month, which day of the week.

The accessor

```
1 import pandas as pd
2
3 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
4 d = df['signup_date']
5
6 parts = pd.DataFrame({
7     'date': d,
8     'year': d.dt.year,
9     'month': d.dt.month,
10    'quarter': d.dt.quarter,
11    'dayofweek': d.dt.dayofweek,
12    'is_month_start': d.dt.is_month_start,
13 })
14 print(parts.head())
```

	date	year	month	quarter	dayofweek	is_month_start
0	2024-03-01	2024	3	1	4	True
1	2024-05-01	2024	5	2	2	True
2	2024-10-01	2024	10	4	1	True
3	2024-10-01	2024	10	4	1	True
4	2023-04-01	2023	4	2	5	True

Elapsed time is usually the feature you want

```
1 import pandas as pd
2
3 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
4
5 # A fixed reference, not today's date -- see the warning below.
6 as_of = pd.Timestamp('2025-01-01')
7 df['days_since_signup'] = (as_of - df['signup_date']).dt.days
8 df['years_since_signup'] = df['days_since_signup'] / 365.25
9
10 print(df[['signup_date', 'days_since_signup', 'years_since_signup']]
11        .head().round(2))
```

	signup_date	days_since_signup	years_since_signup
0	2024-03-01	306	0.84
1	2024-05-01	245	0.67
2	2024-10-01	92	0.25
3	2024-10-01	92	0.25
4	2023-04-01	641	1.75

⚠ Watch Out

Never build a feature from today's date

`pd.Timestamp.now()` gives a different answer every day you run it. The model trained in March sees different numbers in June, and nothing in your code changed. Pass a fixed reference date, store it with the model, and use the same one at prediction time. This is one of the most common causes of a model that silently degrades in production.

Cyclical features

Month 12 and month 1 are adjacent, but as numbers they are eleven apart. A linear model reads December and January as maximally different. The fix is to place the value on a circle.

```
1 import numpy as np
2 import pandas as pd
3
4 months = pd.Series(range(1, 13))
5 angle = 2 * np.pi * (months - 1) / 12
6 cyc = pd.DataFrame({
7     'month': months,
8     'sin': np.sin(angle).round(3),
9     'cos': np.cos(angle).round(3),
10 })
11 print(cyc)
12
13 def gap(a, b):
14     ra = cyc.loc[cyc['month'] == a, ['sin', 'cos']].values[0]
15     rb = cyc.loc[cyc['month'] == b, ['sin', 'cos']].values[0]
16     return np.linalg.norm(ra - rb)
17
18 print('\nraw distance Dec to Jan: %d' % abs(12 - 1))
19 print('raw distance Jun to Jul: %d' % abs(6 - 7))
20 print('circular Dec to Jan: %.3f' % gap(12, 1))
21 print('circular Jun to Jul: %.3f' % gap(6, 7))
```

	month	sin	cos
0	1	0.000	1.000
1	2	0.500	0.866
2	3	0.866	0.500
3	4	1.000	0.000
4	5	0.866	-0.500
5	6	0.500	-0.866
6	7	0.000	-1.000
7	8	-0.500	-0.866
8	9	-0.866	-0.500
9	10	-1.000	-0.000
10	11	-0.866	0.500
11	12	-0.500	0.866


```
raw distance Dec to Jan: 11
raw distance Jun to Jul: 1
circular Dec to Jan: 0.518
circular Jun to Jul: 0.518
```

On the circle, December to January is the same distance as June to July, which is the truth. Use the same trick for hour of day and day of week.

The split that dates force on you

Watch Out

A random split leaks the future

If your rows are ordered in time and you split at random, the model trains on March and tests on February. It has seen the future. Every score is inflated, sometimes enormously, and the failure only shows up in production. For anything time-ordered, split by date, and scikit-learn gives you `TimeSeriesSplit` for cross-validation.

```
1 import numpy as np
2 from sklearn.model_selection import TimeSeriesSplit
3
4 x = np.arange(12)
5 for i, (train_idx, test_idx) in enumerate(TimeSeriesSplit(n_splits=4).split(x), 1):
6     print('fold %d train %-18s test %s'
7           % (i, str(x[train_idx]), str(x[test_idx])))
```

```
fold 1 train [0 1 2 3]          test [4 5]
fold 2 train [0 1 2 3 4 5]      test [6 7]
fold 3 train [0 1 2 3 4 5 6 7]  test [8 9]
fold 4 train [0 1 2 3 4 5 6 7 8 9] test [10 11]
```

Each fold trains only on rows that came before the ones it tests on, and the training window grows. No fold ever sees data from after its test period, which is the only honest way to score a time-ordered problem.

Day 5 takeaway

Day 6 Categories and Text

Cardinality, ordinal against nominal, unseen levels and rare levels

Categories and free text are where most of the information in business data lives, and where most of the leakage and most of the silent breakage comes from too.

Cardinality decides the treatment

```
1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates().copy()
3 df['contract'] = df['contract'].str.strip().str.title()
4 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
5
6 for col in df.select_dtypes('object').columns:
7     print('%-18s %5d distinct %5.1f%% of rows'
```

```
8 % (col, df[col].nunique(), 100 * df[col].nunique() / len(df))
```

```
customer_id      3000 distinct  100.0% of rows
signup_date      72 distinct   2.4% of rows
contract         3 distinct   0.1% of rows
internet_service 3 distinct   0.1% of rows
payment_method   4 distinct   0.1% of rows
has_dependents   2 distinct   0.1% of rows
```

Distinct values	Treatment	Why
2	Map to 0/1	One column is enough
3 to about 15	One-hot encode	Manageable width, no false ordering
15 to a few hundred	Group the rare ones, then one-hot	A level seen twice cannot be learned from
Thousands	Target or hashing encoding	One-hot would be wider than the dataset is long
Unique per row	Drop it	An identifier is not a feature

⚠ Watch Out

customer_id is not a feature

It is unique per row, so a sufficiently flexible model can memorise the target for every training row through it and score perfectly. On new customers it knows nothing. Drop identifiers explicitly rather than trusting yourself to remember.

Ordinal is not the same as categorical

```
1 import pandas as pd
2 from sklearn.preprocessing import OrdinalEncoder, OneHotEncoder
3
4 sat = pd.DataFrame({'rating': ['low', 'high', 'medium', 'low']})
5
6 order = OrdinalEncoder(categories=[['low', 'medium', 'high']])
7 print('ordinal, order stated:')
8 print(order.fit_transform(sat).ravel())
9
10 print('\nordinal, order left to chance (alphabetical):')
11 print(OrdinalEncoder().fit_transform(sat).ravel())
```

```
ordinal, order stated:
[0. 2. 1. 0.]
```

```
ordinal, order left to chance (alphabetical):
[1. 0. 2. 1.]
```

Left to itself the encoder sorts alphabetically, making *high* 0, *low* 1 and *medium* 2. An ordering that says medium is more than low is more than high. If a variable has a genuine order, state it. If it does not, do not use an ordinal encoder at all.

One-hot, and the argument you must not omit

```

1 import pandas as pd
2 from sklearn.preprocessing import OneHotEncoder
3
4 train = pd.DataFrame({'contract': ['Month-to-month', 'One year', 'Two year']})
5 new = pd.DataFrame({'contract': ['Two year', 'Weekly trial']})
6
7 enc = OneHotEncoder(handle_unknown='ignore', sparse_output=False)
8 enc.fit(train)
9 print('columns:', list(enc.get_feature_names_out()))
10 print('unseen category encodes as all zeros:')
11 print(enc.transform(new))

```

```

columns: ['contract_Month-to-month', 'contract_One year', 'contract_Two year']
unseen category encodes as all zeros:
[[0. 0. 1.]
 [0. 0. 0.]]

```

⚠ Watch Out

Without `handle_unknown='ignore'` this raises

Marketing launches a Weekly trial contract, your scoring job meets a category the encoder has never seen, and it throws. At three in the morning. `handle_unknown='ignore'` encodes the unknown as all zeros, which is the sane default: the model treats it as none of the known categories rather than falling over.

Rare categories

```

1 import pandas as pd
2 import numpy as np
3
4 rng = np.random.default_rng(0)
5 cities = (['London'] * 400 + ['Manchester'] * 250 + ['Leeds'] * 120 +
6           ['Bristol'] * 60 + ['Hull'] * 8 + ['Truro'] * 5 + ['Kirkwall'] * 2)
7 s = pd.Series(rng.permutation(cities), name='city')
8
9 counts = s.value_counts()
10 print(counts)
11
12 keep = counts[counts >= 50].index
13 grouped = s.where(s.isin(keep), 'Other')
14 print('\nafter grouping rare levels:')
15 print(grouped.value_counts())

```

```

city
London      400
Manchester   250
Leeds        120
Bristol       60
Hull          8

```

```

Truro          5
Kirkwall       2
Name: count, dtype: int64

after grouping rare levels:
city
London        400
Manchester     250
Leeds          120
Bristol         60
Other          15
Name: count, dtype: int64

```

Kirkwall appears twice. A one-hot column that is 1 for two rows out of a thousand cannot support a reliable estimate, and in cross-validation it may be entirely absent from some folds. Grouping the tail into *Other* keeps the information that these are unusual without pretending you can learn each one.

min_frequency does this for you

`OneHotEncoder(min_frequency=50)` groups anything rarer into a single infrequent bucket, and it learns the threshold on training data only. Prefer it to a manual pass, because doing it by hand before the split means the decision was made using test rows.

Text, briefly

```

1 from sklearn.feature_extraction.text import TfidfVectorizer
2
3 notes = [
4     'billing error on the latest invoice',
5     'speed very slow in the evening',
6     'invoice wrong again, second billing error',
7     'router keeps dropping the connection',
8 ]
9 vec = TfidfVectorizer(stop_words='english', min_df=1)
10 X = vec.fit_transform(notes)
11 print('shape:', X.shape)
12 print('vocabulary:', sorted(vec.vocabulary_)[:10])
13 print('\nrow 0 non-zero terms:')
14 names = vec.get_feature_names_out()
15 row = X[0].toarray().ravel()
16 for i in row.argsort()[::-1][:4]:
17     if row[i] > 0:
18         print(' %-12s %.3f' % (names[i], row[i]))

```

```

shape: (4, 13)
vocabulary: ['billing', 'connection', 'dropping', 'error', 'evening', 'invoice', 'keeps',
  ↳ 'latest', 'router', 'second']

row 0 non-zero terms:

```

TF-IDF weights a term by how often it appears in this document against how rare it is across all of them, so common words score low and distinctive ones score high. It produces a sparse matrix

that goes straight into any scikit-learn model. Week 13 covers the sequence models that improve on it.

Day 6 takeaway

Day 7 A Reproducible Data Quality Report

Turning a week of manual checks into one function you run on every extract

Exploration produces understanding, and understanding evaporates. Today you turn a week of scattered checks into one function you can run against any new extract.

What a data quality report should answer

1. What shape is it, and did that change since last time?
2. Which columns have the wrong type?
3. Where are the gaps, and how big?
4. Which columns are constant, or nearly so?
5. Which are duplicated in content under different names?
6. Are there duplicate rows?
7. What is the target's base rate?

The report

```
1 import pandas as pd
2 import numpy as np
3
4 def profile(df, target=None):
5     """One row per column, with the checks worth running every time."""
6     rows = []
7     for col in df.columns:
8         s = df[col]
9         top = s.value_counts(dropna=True)
10        rows.append({
11            'column': col,
12            'dtype': str(s.dtype),
13            'missing_pct': round(100 * s.isna().mean(), 2),
14            'distinct': s.nunique(dropna=True),
15            'top_value': None if top.empty else str(top.index[0])[:18],
16            'top_share': 0.0 if top.empty else round(100 * top.iloc[0] / len(s), 1),
17        })
18    out = pd.DataFrame(rows)
19    out['constant'] = out['distinct'] <= 1
20    out['id_like'] = out['distinct'] == len(df)
21    return out
22
23 df = pd.read_csv('customers.csv')
```

```
24 print(profile(df).to_string(index=False))
```

	column	dtype	missing_pct	distinct	top_value	top_share	constant	id_like
	customer_id	object	0.00	3000	C00507	0.1	False	False
	signup_date	object	0.00	72	2019-01-01	4.2	False	False
	tenure_months	int64	0.00	72	72	4.2	False	False
	contract	object	0.00	15	Month-to-month	51.3	False	False
	internet_service	object	0.00	3	Fibre optic	45.1	False	False
	payment_method	object	0.00	4	Electronic check	33.6	False	False
	has_dependents	object	6.09	2	No	65.8	False	False
	support_calls	int64	0.00	9	1	34.1	False	False
	monthly_charges	float64	5.45	2232	15.0	5.3	False	False
	total_charges	object	0.00	2969		0.5	False	False
	churned	int64	0.00	2	0	73.2	False	False

The checks that belong beside it

```
1 import pandas as pd
2
3 def integrity(df, target=None):
4     print('rows %d, columns %d' % df.shape)
5     print('duplicate rows      %d' % df.duplicated().sum())
6     obj = df.select_dtypes('object').columns
7     coercible = [c for c in obj
8                  if pd.to_numeric(df[c], errors='coerce').notna().mean() > 0.9]
9     print('text but ~numeric   %s' % (list(coercible) or 'none'))
10    padded = [c for c in obj if (df[c].fillna('') !=
11    ↪ df[c].fillna('').str.strip()).any()]
12    print('has padded values   %s' % (list(padded) or 'none'))
13    if target:
14        rate = df[target].mean()
15        print('target base rate   %.3f (majority guess scores %.3f)'
16              % (rate, max(rate, 1 - rate)))
17
18 df = pd.read_csv('customers.csv')
19 integrity(df, target='churned')
```

```
rows 3138, columns 11
duplicate rows      138
text but ~numeric   ['total_charges']
has padded values   ['contract', 'total_charges']
target base rate    0.268 (majority guess scores 0.732)
```

Every defect week 1 found by hand, found automatically: the numeric column stored as text, the padded category values, the duplicated rows, and the base rate any model has to beat.

Correlated pairs worth knowing about

```
1 import pandas as pd
2 df = pd.read_csv('customers.csv').drop_duplicates().copy()
3 df['contract'] = df['contract'].str.strip().str.title()
4 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
```

```

5 import numpy as np
6
7 num = df.select_dtypes('number').drop(columns=['churned'])
8 corr = num.corr().abs()
9
10 # Upper triangle only, so each pair appears once and never against itself.
11 mask = np.triu(np.ones(corr.shape, dtype=bool), k=1)
12 pairs = (corr.where(mask).stack().sort_values(ascending=False))
13 print(pairs[pairs > 0.5].round(3))

```

```

tenure_months  total_charges    0.829
dtype: float64

```

`tenure_months` and `total_charges` at 0.83. Not a reason to delete a column yet. It is a reason to expect unstable coefficients from a linear model, which week 4 will demonstrate and fix with regularisation.

Your assignment

Run `profile` and `integrity` against a dataset you have not seen, anything from your own work, or a public CSV. Write down three things the report told you that you would not have checked by hand. Then add a check of your own for whatever it missed.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. pandas treats several strings as missing by default. Which of these is *not* in that list?
 - a) NULL
 - b) A single space
 - c) An empty string
 - d) NA
2. A customers table is joined to a tickets table and the row count rises. What has happened?
 - a) The join failed
 - b) Customers with several tickets now appear once per ticket, so any average over customers is now weighted by ticket count
 - c) Duplicates were created at random
 - d) The join key was numeric
3. Values are missing *because* the customer is new. That is:
 - a) Not really missing
 - b) Missing completely at random

- c) Missing at random
 - d) Missing not at random. The reason is related to the value itself
4. Which scaler is least disturbed by a single extreme outlier?
- a) `RobustScaler`
 - b) `Normalizer`
 - c) `StandardScaler`
 - d) `MinMaxScaler`
5. Why encode an hour of the day as a sine and cosine pair?
- a) So that 23:00 and 00:00 are close together, which they are not on a plain 0 to 23 scale
 - b) To reduce the number of columns
 - c) Because trees require it
 - d) It is faster
6. A category appears twice in a thousand rows. The safest treatment is:
- a) Drop those rows
 - b) Give it its own one-hot column
 - c) Group it into an *Other* bucket, learned on training data only
 - d) Replace it with the mean
7. `OneHotEncoder(handle_unknown='ignore')` meets a category it never saw in training. It:
- a) Drops the row
 - b) Raises an error
 - c) Encodes it as a row of zeros
 - d) Assigns it to the most common category
8. What does TF-IDF weight a term by?
- a) Its position in the sentence
 - b) How often it appears in the whole corpus
 - c) Its length
 - d) How often it appears in this document against how rare it is across all documents

Answers: 1b, 2b, 3d, 4a, 5a, 6c, 7c, 8d. Anything you got wrong points at the day to re-read, not at a gap in ability.

The Maths That Actually Matters

Week 3

Linear algebra, calculus, probability and statistics, only the parts that change what you do, each one implemented rather than proved.

OBJECTIVES

You do not need a maths degree for machine learning, but four ideas will keep tripping you up until you own them: a model is a matrix multiplication, training is stepping against a gradient, a classifier outputs probabilities that must be read carefully, and every score you measure carries uncertainty. This week implements each rather than proving it, and ends by solving the same regression three ways.

Day 1 Vectors, Matrices and Why Models Are Multiplications

Dot products, matrix shapes, distance, and the one line most models are

You do not need a maths degree to do machine learning well. You do need four ideas, and this week covers them by implementing each one rather than proving anything. Today: why almost every model is, underneath, a matrix multiplication.

A row of data is a vector

One customer, tenure 12, charge 79.50, support calls 2, is a point in three-dimensional space. Two thousand customers are two thousand points. Everything that follows is geometry on those points.

```
1 import numpy as np
2
3 customer = np.array([12.0, 79.50, 2.0])
4 weights = np.array([-0.05, 0.02, 0.30])
5
6 score = np.dot(customer, weights)
7 manual = sum(c * w for c, w in zip(customer, weights))
8 print('dot product : %.4f' % score)
9 print('by hand      : %.4f' % manual)
```

```
dot product : 1.5900
by hand      : 1.5900
```

The whole dataset at once

You never score one customer. You score all of them, and the matrix product does every dot product in one call.

```

1 import numpy as np
2
3 X = np.array([[12.0, 79.50, 2.0],
4               [48.0, 55.00, 0.0],
5               [ 3.0, 92.00, 5.0],
6               [60.0, 21.00, 1.0]])
7 w = np.array([-0.05, 0.02, 0.30])
8 b = -1.2
9
10 scores = X @ w + b
11 print('X shape', X.shape, '@ w shape', w.shape, '→ ', scores.shape)
12 print('scores', scores.round(3))

```

```

X shape (4, 3) @ w shape (3,) → (4,)
scores [ 0.39 -2.5   1.99 -3.48]

```

That line is a linear model

$X @ w + b$ is the entire forward pass of linear regression, and with a sigmoid wrapped round it, of logistic regression. Week 11's neural network is the same line repeated with a nonlinearity between the repeats. Understand this one expression and you understand the shape of most of the course.

Shapes have to line up

A matrix product $(n, k) @ (k, m)$ gives (n, m) . The inner dimensions must match, and this is the source of most errors you will hit in week 11.

```

1 import numpy as np
2
3 X = np.random.default_rng(0).normal(size=(100, 4))    # 100 rows, 4 features
4 W1 = np.random.default_rng(1).normal(size=(4, 8))    # layer: 4 in, 8 out
5 W2 = np.random.default_rng(2).normal(size=(8, 1))    # layer: 8 in, 1 out
6
7 h = X @ W1
8 out = h @ W2
9 print('X ', X.shape)
10 print('X@W1', h.shape)
11 print('h@W2', out.shape)
12
13 try:
14     X @ W2
15 except ValueError as e:
16     print('\nmismatch:', e)

```

```

X      (100, 4)
X@W1   (100, 8)
h@W2   (100, 1)

```

```
mismatch: matmul: Input operand 1 has a mismatch in its core dimension 0, with gufunc
  ↳ signature (n?,k),(k,m?)→(n?,m?) (size 8 is different from 4)
```

Distance, and what it means for a model

kNN, k-means and every clustering algorithm in week 9 rest on measuring how far apart two rows are.

```
1 import numpy as np
2
3 a = np.array([12.0, 79.5])
4 b = np.array([14.0, 81.0])
5 c = np.array([60.0, 21.0])
6
7 def euclid(p, q):
8     return np.sqrt((p - q) ** 2).sum())
9
10 print('a to b: %.2f' % euclid(a, b))
11 print('a to c: %.2f' % euclid(a, c))
12 print('numpy : %.2f' % np.linalg.norm(a - c))
```

```
a to b: 2.50
a to c: 75.67
numpy : 75.67
```

Watch Out

Distance is meaningless on unscaled columns

Tenure runs 1 to 72; total charges run to 8000. In the distance calculation, charges dominate entirely and tenure might as well not exist. Every distance-based algorithm requires scaling first, and this is the real reason **StandardScaler** appears in nearly every pipeline in this course.

```
1 import numpy as np
2 from sklearn.preprocessing import StandardScaler
3
4 raw = np.array([[12.0, 950.0],
5                 [14.0, 8000.0],
6                 [60.0, 1000.0]])
7
8 print('unscaled: row0-row1 %.1f, row0-row2 %.1f'
9       % (np.linalg.norm(raw[0] - raw[1]), np.linalg.norm(raw[0] - raw[2])))
10
11 scaled = StandardScaler().fit_transform(raw)
12 print('scaled : row0-row1 %.2f, row0-row2 %.2f'
13       % (np.linalg.norm(scaled[0] - scaled[1]),
14          np.linalg.norm(scaled[0] - scaled[2])))
```

```
unscaled: row0-row1 7050.0, row0-row2 69.3
scaled : row0-row1 2.13, row0-row2 2.17
```

Unscaled, row 0 looks far closer to row 2 than to row 1, because the charge difference swamps everything. Scaled, the answer flips, and the scaled answer is the one that respects both columns.

Day 1 takeaway

Day 2 Derivatives and Gradients

Slopes, the chain rule, and why the sigmoid derivative caps at 0.25

Training a model means adjusting numbers until a loss gets smaller. To know which direction to adjust, you need the derivative. That is the whole role calculus plays here.

A derivative is a slope

```

1 import numpy as np
2
3 def f(x):
4     return x ** 2 + 3 * x + 5
5
6 def numeric_derivative(f, x, h=1e-6):
7     return (f(x + h) - f(x - h)) / (2 * h)
8
9 # The exact derivative of  $x^2 + 3x + 5$  is  $2x + 3$ .
10 for x in [-4.0, -1.5, 0.0, 3.0]:
11     print('x=%5.1f  numeric %8.4f  exact %8.4f'
12           % (x, numeric_derivative(f, x), 2 * x + 3))

```

```

x= -4.0  numeric  -5.0000  exact  -5.0000
x= -1.5  numeric  -0.0000  exact   0.0000
x=  0.0  numeric   3.0000  exact   3.0000
x=  3.0  numeric   9.0000  exact   9.0000

```

At $x = -1.5$ the derivative is zero, and that is the bottom of the parabola. Finding where the derivative is zero is finding the minimum, which is exactly what training does.

The gradient is the derivative in several directions

Models have thousands of parameters, not one. The gradient collects the derivative with respect to each of them into a vector that points in the direction of steepest increase.

```

1 import numpy as np
2
3 def loss(w):
4     # A bowl, steeper in the second direction than the first.
5     return (w[0] - 2) ** 2 + 4 * (w[1] + 1) ** 2
6
7 def gradient(f, w, h=1e-6):
8     g = np.zeros_like(w)
9     for i in range(len(w)):
10         step = np.zeros_like(w)
11         step[i] = h

```

```

12     g[i] = (f(w + step) - f(w - step)) / (2 * h)
13     return g
14
15 w = np.array([0.0, 0.0])
16 g = gradient(loss, w)
17 print('loss at [0,0]   %.3f' % loss(w))
18 print('gradient       %s' % g.round(4))
19 print('exact          %s' % np.array([2 * (w[0] - 2), 8 * (w[1] + 1)]))
20 print('\nstepping against the gradient:')
21 print('loss at w - 0.1*g   %.3f' % loss(w - 0.1 * g))

```

```

loss at [0,0]   8.000
gradient        [-4.  8.]
exact           [-4.  8.]

stepping against the gradient:
loss at w - 0.1*g   2.720

```

Against the gradient, not along it

The gradient points uphill. To reduce a loss you step in the opposite direction, which is why the update rule has a minus sign: $w = w - \text{learning_rate} * \text{gradient}$. That one line is gradient descent, and tomorrow you implement it.

The chain rule, which is what backpropagation is

Compose two functions and the derivative multiplies. This is the whole mechanism by which a neural network assigns credit to a weight buried four layers from the output.

```

1 import numpy as np
2
3 def g(x):
4     return 3 * x + 1
5
6 def f(u):
7     return u ** 2
8
9 def composed(x):
10    return f(g(x))
11
12 x = 2.0
13 # chain rule: d/dx f(g(x)) = f'(g(x)) * g'(x) = 2*g(x) * 3
14 chain = 2 * g(x) * 3
15 numeric = (composed(x + 1e-6) - composed(x - 1e-6)) / 2e-6
16 print('chain rule %.4f' % chain)
17 print('numeric    %.4f' % numeric)

```

```

chain rule 42.0000
numeric    42.0000

```

The derivatives you will actually meet

Function	Where it appears	Derivative
x^2	Squared error loss	$2x$
<code>sigmoid(x)</code>	Logistic output layer	<code>sigmoid(x)(1 - sigmoid(x))</code>
<code>relu(x)</code>	Hidden layers	1 if $x > 0$, else 0
<code>log(x)</code>	Log loss	$1/x$

```

1 import numpy as np
2
3 def sigmoid(x):
4     return 1 / (1 + np.exp(-x))
5
6 def relu(x):
7     return np.maximum(0, x)
8
9 xs = np.array([-2.0, -0.5, 0.0, 0.5, 2.0])
10 h = 1e-6
11 print('  x   sigmoid  d/dx numeric  d/dx formula  relu  d/dx')
12 for x in xs:
13     num = (sigmoid(x + h) - sigmoid(x - h)) / (2 * h)
14     formula = sigmoid(x) * (1 - sigmoid(x))
15     rnum = (relu(x + h) - relu(x - h)) / (2 * h)
16     print('%5.1f  %7.4f  %12.6f  %12.6f  %5.1f  %4.1f'
17           % (x, sigmoid(x), num, formula, relu(x), rnum))

```

x	sigmoid	d/dx numeric	d/dx formula	relu	d/dx
-2.0	0.1192	0.104994	0.104994	0.0	0.0
-0.5	0.3775	0.235004	0.235004	0.0	0.0
0.0	0.5000	0.250000	0.250000	0.0	0.5
0.5	0.6225	0.235004	0.235004	0.5	1.0
2.0	0.8808	0.104994	0.104994	2.0	1.0

⚠ Watch Out**The sigmoid derivative peaks at 0.25**

Look at the column: the largest value is 0.25, at $x = 0$, and it falls away fast. Multiply several of those together through a deep network and the gradient reaching the early layers is vanishingly small. That is the vanishing gradient problem, it is why sigmoid was abandoned for hidden layers, and it is why ReLU, whose derivative is exactly 1 for positive inputs, replaced it. Week 11 returns to this.

Day 2 takeaway**Day 3 Gradient Descent, Implemented**

The four-line loop, and the three ways learning rate and scale break it

Yesterday gave you the direction. Today you take the steps, and watch the three ways this goes wrong.

The algorithm, in full

1. Start with a guess for the parameters.
2. Compute the loss and its gradient.
3. Move the parameters a small step against the gradient.
4. Repeat until the loss stops falling.

```
1 import numpy as np
2
3 rng = np.random.default_rng(0)
4 n = 200
5 x = rng.uniform(0, 10, n)
6 y = 3.5 * x + 2.0 + rng.normal(0, 1.5, n)    # truth: slope 3.5, intercept 2
7
8 w, b = 0.0, 0.0
9 lr = 0.01
10
11 for step in range(1, 2001):
12     pred = w * x + b
13     error = pred - y
14     loss = (error ** 2).mean()
15     dw = 2 * (error * x).mean()
16     db = 2 * error.mean()
17     w -= lr * dw
18     b -= lr * db
19     if step in (1, 10, 100, 500, 2000):
20         print('step %4d  loss %8.4f  w %.4f  b %.4f' % (step, loss, w, b))
21
22 print('\ntruth      w 3.5000  b 2.0000')
23 print('recovered  w %.4f  b %.4f' % (w, b))
```

```
step    1  loss 541.8747  w 2.8686  b 0.4151
step   10  loss   2.8682  w 3.6719  b 0.5939
step  100  loss   2.5705  w 3.5999  b 1.1008
step  500  loss   2.3502  w 3.4833  b 1.9210
step 2000  loss   2.3448  w 3.4621  b 2.0705

truth      w 3.5000  b 2.0000
recovered  w 3.4621  b 2.0705
```

Two thousand steps of arithmetic recovered the rule that generated the data. No library, no solver. The same loop, at enormous scale, is what trains a neural network.

The learning rate is the whole game

```
1 import numpy as np
2
3 rng = np.random.default_rng(0)
```

```

4 x = rng.uniform(0, 10, 200)
5 y = 3.5 * x + 2.0 + rng.normal(0, 1.5, 200)
6
7 def descend(lr, steps=300):
8     w = b = 0.0
9     for _ in range(steps):
10         error = (w * x + b) - y
11         w -= lr * 2 * (error * x).mean()
12         b -= lr * 2 * error.mean()
13         # Overflow passes through astronomically large but still finite
14         # values on its way to inf, so test the magnitude, not isfinite.
15         if not np.isfinite(w) or abs(w) > 1e10:
16             return None, None, None
17     return w, b, (((w * x + b) - y) ** 2).mean()
18
19 for lr in [0.0001, 0.001, 0.01, 0.03, 0.05]:
20     w, b, loss = descend(lr)
21     if w is None:
22         print('lr %-7s diverged to infinity' % lr)
23     else:
24         print('lr %-7s w %7.4f b %7.4f loss %8.4f' % (lr, w, b, loss))

```

```

lr 0.0001 w 3.3268 b 0.4947 loss 7.8274
lr 0.001 w 3.6533 b 0.7253 loss 2.7752
lr 0.01 w 3.5162 b 1.6897 loss 2.3793
lr 0.03 diverged to infinity
lr 0.05 diverged to infinity

```

⚠ Watch Out

Too small wastes time, too large explodes

At 0.0001 the model has barely moved after 300 steps. At 0.05 each step overshoots the minimum by more than it started from, the error grows, and the parameters run to infinity within a few dozen iterations. If your loss becomes `nan` during training, the learning rate is the first thing to look at.

Scaling changes how hard the problem is

```

1 import numpy as np
2
3 rng = np.random.default_rng(1)
4 n = 500
5 tenure = rng.uniform(1, 72, n)
6 charges = rng.uniform(20, 120, n)
7 total = tenure * charges # ranges into the thousands
8 y = 0.4 * tenure + 0.05 * charges + 0.001 * total + rng.normal(0, 1, n)
9
10 def descend(X, y, lr, steps=400):
11     # A column of ones so the model can fit an intercept. Without it a
12     # scaled matrix has mean zero and cannot reach a non-zero target.
13     Xb = np.column_stack([np.ones(len(X)), X])

```

```

14     w = np.zeros(Xb.shape[1])
15     for _ in range(steps):
16         error = Xb @ w - y
17         w -= lr * 2 * (Xb.T @ error) / len(y)
18         if not np.all(np.isfinite(w)) or np.abs(w).max() > 1e10:
19             return None
20     return ((Xb @ w) - y) ** 2).mean()
21
22 raw = np.column_stack([tenure, charges, total])
23 scaled = (raw - raw.mean(axis=0)) / raw.std(axis=0)
24
25 for lr in [1e-7, 1e-5, 1e-3, 1e-2]:
26     r = descend(raw, y, lr)
27     s = descend(scaled, y, lr)
28     print('lr %-6s raw %-12s scaled %s'
29           % (lr, 'diverged' if r is None else '%.4f' % r,
30              'diverged' if s is None else '%.4f' % s))

```

```

lr 1e-07    raw 38.9770      scaled 518.7263
lr 1e-05    raw diverged    scaled 509.3600
lr 0.001    raw diverged    scaled 92.7342
lr 0.01     raw diverged    scaled 1.4854

```

On raw columns only a tiny learning rate survives, and it converges slowly. On scaled columns a much larger rate works and gets further. Scaling is not cosmetic. It changes the shape of the loss surface from a long thin valley into something closer to a bowl, and gradient descent handles bowls far better.

Batch, stochastic and mini-batch

Variant	Gradient computed on	Trade-off
Batch	Every row, every step	Smooth and accurate, slow on large data
Stochastic (SGD)	One row at a time	Fast and noisy; the noise can escape shallow minima
Mini-batch	32 to 512 rows	The compromise everyone actually uses

```

1 import numpy as np
2
3 rng = np.random.default_rng(0)
4 n = 2000
5 x = rng.uniform(0, 10, n)
6 y = 3.5 * x + 2.0 + rng.normal(0, 1.5, n)
7
8 def minibatch(batch_size, epochs=30, lr=0.01, seed=0):
9     r = np.random.default_rng(seed)
10    w = b = 0.0
11    for _ in range(epochs):
12        order = r.permutation(n)
13        for start in range(0, n, batch_size):
14            idx = order[start:start + batch_size]

```

```

15         xb, yb = x[idx], y[idx]
16         error = (w * xb + b) - yb
17         w -= lr * 2 * (error * xb).mean()
18         b -= lr * 2 * error.mean()
19     return w, b, (((w * x + b) - y) ** 2).mean()
20
21 for bs in [1, 32, 256, n]:
22     w, b, loss = minibatch(bs)
23     label = 'stochastic' if bs == 1 else ('batch' if bs == n else 'mini %d' % bs)
24     print('%-12s w %.4f b %.4f loss %.4f' % (label, w, b, loss))

```

```

stochastic  w 3.1202 b 2.0565 loss 6.8004
mini 32     w 3.4412 b 2.0094 loss 2.3413
mini 256    w 3.5696 b 1.5688 loss 2.2876
batch      w 3.6854 b 0.7612 loss 2.6329

```

The truth is $w = 3.5$, $b = 2.0$. The mini-batch runs get closest; full batch is furthest off, with the intercept still at 0.76 after 30 epochs. That is the point: batch made 30 parameter updates in total, mini-batch-32 made nearly two thousand, and stochastic made sixty thousand. Same data, same epochs, vastly different amounts of learning, which is why nobody trains on full batches. Stochastic is noisiest, visible in its slope of 3.12.

Day 3 takeaway

Day 4 Probability You Actually Need

Conditional probability, Bayes, base rates and the independence assumption

Classifiers do not output labels. They output probabilities, and something downstream turns those into a decision. Today is the probability you need to read that output correctly.

Conditional probability

```

1 import pandas as pd
2
3 df = pd.read_csv('customers.csv').drop_duplicates().copy()
4 df['contract'] = df['contract'].str.strip().str.title()
5
6 p_churn = df['churned'].mean()
7 p_m2m = (df['contract'] == 'Month-To-Month').mean()
8 p_churn_given_m2m = df.loc[df['contract'] == 'Month-To-Month', 'churned'].mean()
9 p_m2m_given_churn = (df.loc[df['churned'] == 1, 'contract'] ==
10     ↪ 'Month-To-Month').mean()
11
12 print('P(churn)                %.3f' % p_churn)
13 print('P(month-to-month)       %.3f' % p_m2m)
14 print('P(churn | month-to-month) %.3f' % p_churn_given_m2m)
15 print('P(month-to-month | churn) %.3f' % p_m2m_given_churn)

```

```

P(churn)                0.268

```

```
P(month-to-month)          0.553
P(churn | month-to-month) 0.414
P(month-to-month | churn) 0.854
```

⚠ Watch Out

These last two are not the same number

$P(\text{churn given month-to-month})$ is 0.41. $P(\text{month-to-month given churn})$ is 0.85. Confusing the two is the single most common probability error in analytics, and it has a name, the prosecutor's fallacy. "85 percent of churners were month-to-month" sounds like month-to-month customers mostly churn. They mostly do not.

Bayes' theorem, which converts between them

```
1 import pandas as pd
2
3 df = pd.read_csv('customers.csv').drop_duplicates().copy()
4 df['contract'] = df['contract'].str.strip().str.title()
5
6 p_churn = df['churned'].mean()
7 p_m2m = (df['contract'] == 'Month-To-Month').mean()
8 p_m2m_given_churn = (df.loc[df['churned'] == 1, 'contract'] ==
9     ↪ 'Month-To-Month').mean()
10
11 #  $P(A/B) = P(B/A) * P(A) / P(B)$ 
12 bayes = p_m2m_given_churn * p_churn / p_m2m
13 actual = df.loc[df['contract'] == 'Month-To-Month', 'churned'].mean()
14 print('Bayes %.4f' % bayes)
15 print('actual %.4f' % actual)
```

```
Bayes 0.4139
actual 0.4139
```

Why base rates dominate

A classic: a test for a rare condition is 99 percent accurate. You test positive. What is the chance you have it?

```
1 prevalence = 0.001          # 1 in 1000 people have it
2 sensitivity = 0.99          # correctly flags 99% of those who do
3 specificity = 0.99          # correctly clears 99% of those who do not
4
5 true_pos = prevalence * sensitivity
6 false_pos = (1 - prevalence) * (1 - specificity)
7 posterior = true_pos / (true_pos + false_pos)
8
9 print('per 100,000 people:')
10 print('  have it and test positive      %6.0f' % (100_000 * true_pos))
11 print('  healthy but test positive      %6.0f' % (100_000 * false_pos))
12 print('\nP(condition | positive test)  %.3f' % posterior)
```

```

per 100,000 people:
    have it and test positive      99
    healthy but test positive     999

P(condition | positive test)  0.090

```

Nine percent, from a 99 percent accurate test

Because the condition is rare, the far larger healthy group generates ten times more false positives than the sick group generates true ones. This is exactly the situation in fraud detection, disease screening and any rare-event problem, and it is why week 7 insists on precision and recall rather than accuracy.

Distributions you will meet

```

1 import numpy as np
2 from scipy import stats
3
4 rng = np.random.default_rng(0)
5 samples = {
6     'normal': rng.normal(50, 10, 5000),
7     'uniform': rng.uniform(0, 100, 5000),
8     'poisson': rng.poisson(1.5, 5000),
9     'lognormal': rng.lognormal(3, 0.6, 5000),
10 }
11 print('%-10s %8s %8s %8s %8s' % ('', 'mean', 'median', 'std', 'skew'))
12 for name, s in samples.items():
13     print('%-10s %8.2f %8.2f %8.2f %8.2f'
14           % (name, s.mean(), np.median(s), s.std(), stats.skew(s)))

```

	mean	median	std	skew
normal	49.95	49.72	9.95	0.02
uniform	50.03	50.20	28.81	-0.00
poisson	1.54	1.00	1.22	0.73
lognormal	24.00	20.19	15.68	2.24

Distribution	Models	Seen in this course
Normal	Measurement error, sums of many effects	Linear regression residuals
Bernoulli	A single yes/no outcome	The churn target
Poisson	Counts of rare events in a window	<code>support_calls</code>
Log-normal	Quantities that multiply	Charges, incomes, durations

Independence, and why naive Bayes is called naive

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()

```

```

6
7 a = (df['contract'] == 'Month-To-Month')
8 b = (df['payment_method'] == 'Electronic check')
9
10 print('P(A)          %.4f' % a.mean())
11 print('P(B)          %.4f' % b.mean())
12 print('P(A)*P(B) %.4f  ← what independence would predict' % (a.mean() * b.mean()))
13 print('P(A and B) %.4f ← what actually happens' % (a & b).mean())
14
15 c = (df['contract'] == 'Two Year')
16 d = (df['tenure_months'] > 40)
17 print('\nP(C)*P(D) %.4f' % (c.mean() * d.mean()))
18 print('P(C and D) %.4f ← strongly dependent' % (c & d).mean())

```

```

P(A)          0.5533
P(B)          0.3330
P(A)*P(B) 0.1843  ← what independence would predict
P(A and B) 0.1823 ← what actually happens

P(C)*P(D) 0.0327
P(C and D) 0.0850 ← strongly dependent

```

Contract type and payment method are close to independent, so multiplying their probabilities is nearly right. Contract type and tenure are not, and multiplying is badly wrong. Naive Bayes assumes every feature is independent given the class, usually false, and it often works anyway, which week 5 explains.

Day 4 takeaway

Day 5 Likelihood and Loss Functions

Why the loss you choose decides the estimator you get

A loss function is how a model knows it is wrong. Choosing the wrong one is not a tuning mistake. It means optimising for something you do not want.

Squared error, and what it assumes

```

1 import numpy as np
2
3 y = np.array([10.0, 12.0, 11.0, 13.0, 40.0]) # last one is an outlier
4
5 def mse(pred):
6     return ((y - pred) ** 2).mean()
7
8 def mae(pred):
9     return np.abs(y - pred).mean()
10
11 grid = np.linspace(9, 25, 1601)
12 print('minimises MSE at %.2f (the mean is %.2f)'
13       % (grid[np.argmin([mse(g) for g in grid])], y.mean()))

```

```

14 print('minimises MAE at %.2f    (the median is %.2f)'
15       % (grid[np.argmin([mae(g) for g in grid])], np.median(y)))

```

```

minimises MSE at 17.20    (the mean is 17.20)
minimises MAE at 12.00    (the median is 12.00)

```

Your loss chooses your estimator

Squared error is minimised by the mean; absolute error by the median. One outlier of 40 drags the MSE answer to 17.2 while the MAE answer stays at 12. If your target has extreme values and you do not want them to dominate, that is an argument about the loss function, not about the data.

Log loss, for probabilities

Accuracy treats a confident wrong answer and a hesitant wrong answer as identical. Log loss does not, and that is why it is what classifiers actually optimise.

```

1 import numpy as np
2
3 def log_loss_one(y_true, p):
4     p = np.clip(p, 1e-15, 1 - 1e-15)
5     return -(y_true * np.log(p) + (1 - y_true) * np.log(1 - p))
6
7 print('truth=1, predicted probability → loss')
8 for p in [0.99, 0.9, 0.7, 0.5, 0.3, 0.1, 0.01]:
9     print('    p=%.2f    %7.4f' % (p, log_loss_one(1, p)))

```

```

truth=1, predicted probability → loss
p=0.99    0.0101
p=0.90    0.1054
p=0.70    0.3567
p=0.50    0.6931
p=0.30    1.2040
p=0.10    2.3026
p=0.01    4.6052

```

Watch Out

Confidently wrong is punished without limit

Predicting 0.01 for something that happens costs 4.6, against 0.69 for an honest 0.5. As the prediction approaches zero the loss goes to infinity, which is why the clip is there. A model trained on log loss learns to be uncertain when it should be, and that is what makes its probabilities usable for ranking and thresholding.

Where log loss comes from

It is not arbitrary. It is the negative log of the likelihood, the probability the model assigns to the data actually observed. Maximising that likelihood and minimising log loss are the same operation.

```

1 import numpy as np

```

```

2
3 y = np.array([1, 0, 1, 1, 0])
4 model_a = np.array([0.9, 0.2, 0.8, 0.7, 0.1]) # sensible
5 model_b = np.array([0.6, 0.5, 0.5, 0.6, 0.4]) # hedging
6
7 def likelihood(y, p):
8     return np.prod(np.where(y == 1, p, 1 - p))
9
10 def neg_log_lik(y, p):
11     return -np.sum(np.where(y == 1, np.log(p), np.log(1 - p)))
12
13 for name, p in [('A (sensible)', model_a), ('B (hedging)', model_b)]:
14     print('%-14s likelihood %.6f    neg log lik %.4f'
15           % (name, likelihood(y, p), neg_log_lik(y, p)))

```

```

A (sensible)    likelihood 0.362880    neg log lik 1.0137
B (hedging)     likelihood 0.054000    neg log lik 2.9188

```

Model A assigns the observed outcomes about eight times more probability than model B does, and correspondingly has the lower negative log likelihood. Products of many small probabilities underflow to zero, which is the practical reason everything is done in logs.

Matching the loss to the problem

Problem	Loss	sklearn scoring
Regression, outliers matter	Squared error	<code>neg_mean_squared_error</code>
Regression, outliers are noise	Absolute error	<code>neg_mean_absolute_error</code>
Regression, relative error matters	Squared log error	<code>neg_mean_squared_log_error</code>
Binary classification	Log loss	<code>neg_log_loss</code>
Multi-class	Cross-entropy	<code>neg_log_loss</code>
Ranking quality only	-	<code>roc_auc</code>

```

1 import numpy as np
2 from sklearn.metrics import (mean_squared_error, mean_absolute_error,
3                               mean_squared_log_error)
4
5 truth = np.array([100.0, 1000.0])
6 over = np.array([110.0, 1010.0]) # +10 on each
7
8 print('absolute error identical: %.1f' % mean_absolute_error(truth, over))
9 print('squared error identical : %.1f' % mean_squared_error(truth, over))
10 print('squared LOG error      : %.6f' % mean_squared_log_error(truth, over))
11 print('\n10 on 100 is a 10%% error; 10 on 1000 is 1%%.')
12 print('Only the log version treats those differently.')

```

```

absolute error identical: 10.0
squared error identical : 100.0
squared LOG error      : 0.004506

10 on 100 is a 10%% error; 10 on 1000 is 1%%.
Only the log version treats those differently.

```

Day 5 takeaway**Day 6 Uncertainty, Bias and Variance**

Confidence intervals on model scores, and the decomposition measured directly

Two models score 0.81 and 0.83 on your test set. Is the second better, or did it get luckier? Answering that is what this day is for, and it is the difference between a practitioner and someone who runs `fit`.

Any score computed on a sample is itself uncertain

```

1 import numpy as np
2
3 rng = np.random.default_rng(0)
4 true_rate = 0.268
5
6 for n in [50, 200, 750, 3000]:
7     estimates = [rng.binomial(n, true_rate) / n for _ in range(4000)]
8     lo, hi = np.percentile(estimates, [2.5, 97.5])
9     print('n=%5d    estimate %.3f +/- %.3f    95%% range %.3f to %.3f'
10           % (n, np.mean(estimates), np.std(estimates), lo, hi))

```

```

n=   50    estimate 0.267 +/- 0.063    95% range 0.140 to 0.400
n=  200    estimate 0.268 +/- 0.032    95% range 0.210 to 0.335
n=  750    estimate 0.268 +/- 0.016    95% range 0.236 to 0.300
n= 3000    estimate 0.268 +/- 0.008    95% range 0.251 to 0.284

```

With a 750-row test set, the size week 1 used, a churn rate measured at 0.268 could plausibly be anywhere from 0.237 to 0.300. Model scores carry the same kind of uncertainty, and two models within a percentage point of each other on 750 rows are not distinguishable.

Bootstrap a confidence interval for a model score

```

1 import numpy as np
2 import pandas as pd
3 from sklearn.model_selection import train_test_split
4 from sklearn.pipeline import make_pipeline
5 from sklearn.impute import SimpleImputer
6 from sklearn.preprocessing import StandardScaler
7 from sklearn.linear_model import LogisticRegression
8 from sklearn.metrics import roc_auc_score
9
10 df = pd.read_csv('customers.csv').drop_duplicates().copy()
11 X = df[['tenure_months', 'monthly_charges', 'support_calls']]
12 y = df['churned']
13 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
14                                           stratify=y, random_state=42)
15
16 model = make_pipeline(SimpleImputer(strategy='median'), StandardScaler(),

```

```

17         LogisticRegression(max_iter=1000)).fit(X_tr, y_tr)
18 proba = model.predict_proba(X_te)[: , 1]
19 point = roc_auc_score(y_te, proba)
20
21 rng = np.random.default_rng(0)
22 y_arr, p_arr = y_te.to_numpy(), proba
23 boot = []
24 for _ in range(2000):
25     idx = rng.integers(0, len(y_arr), len(y_arr))
26     if y_arr[idx].sum() in (0, len(idx)):
27         continue
28     boot.append(roc_auc_score(y_arr[idx], p_arr[idx]))
29
30 lo, hi = np.percentile(boot, [2.5, 97.5])
31 print('AUC %.4f   95% CI %.4f to %.4f   width %.4f'
32       % (point, lo, hi, hi - lo))

```

```
AUC 0.7764   95% CI 0.7392 to 0.8101   width 0.0708
```

Report the interval, not just the number

The interval here is about seven points wide. Any competing model whose score falls inside it has not been shown to be better. Quoting a single figure to three decimal places implies a precision the test set cannot support.

The bias-variance decomposition, measured

```

1 import numpy as np
2 from sklearn.tree import DecisionTreeRegressor
3 from sklearn.preprocessing import PolynomialFeatures
4 from sklearn.linear_model import LinearRegression
5 from sklearn.pipeline import make_pipeline
6
7 def truth(x):
8     return np.sin(1.5 * x) + 0.3 * x
9
10 rng = np.random.default_rng(0)
11 x_test = np.linspace(0, 6, 120).reshape(-1, 1)
12 f_test = truth(x_test.ravel())
13
14 def decompose(make_model, n_sets=200, n=40):
15     preds = np.zeros((n_sets, len(x_test)))
16     for i in range(n_sets):
17         x = rng.uniform(0, 6, n).reshape(-1, 1)
18         y = truth(x.ravel()) + rng.normal(0, 0.3, n)
19         preds[i] = make_model().fit(x, y).predict(x_test)
20     mean_pred = preds.mean(axis=0)
21     bias2 = ((mean_pred - f_test) ** 2).mean()
22     variance = preds.var(axis=0).mean()
23     return bias2, variance
24
25 candidates = {

```

```

26 'line (deg 1) ': lambda: make_pipeline(PolynomialFeatures(1),
    ↪ LinearRegression()),
27 'poly deg 5   ': lambda: make_pipeline(PolynomialFeatures(5),
    ↪ LinearRegression()),
28 'poly deg 15  ': lambda: make_pipeline(PolynomialFeatures(15),
    ↪ LinearRegression()),
29 'tree depth 2 ': lambda: DecisionTreeRegressor(max_depth=2),
30 'tree unlimited': lambda: DecisionTreeRegressor(),
31 }
32 print('%-16s %9s %9s %9s' % ('model', 'bias^2', 'variance', 'total'))
33 for name, maker in candidates.items():
34     b, v = decompose(maker)
35     print('%-16s %9.4f %9.4f %9.4f' % (name, b, v, b + v))

```

model	bias^2	variance	total
line (deg 1)	0.4725	0.0318	0.5043
poly deg 5	0.0167	0.0419	0.0586
poly deg 15	7.4814	2335.0026	2342.4839
tree depth 2	0.0842	0.0742	0.1584
tree unlimited	0.0016	0.1016	0.1032

The straight line has high bias and almost no variance: it is wrong in the same way every time. The degree-15 polynomial is the opposite failure. Its variance is over two thousand, because with only 40 noisy points each fit swings somewhere different. Note that its bias is bad too, at 7.48: when the individual fits are that unstable, even their average is nowhere near the truth. The best total error is the degree-5 polynomial, in between, and finding that point is what every regularisation technique in this course is for.

⚠ Watch Out

You cannot measure this on real data

The decomposition above needs the true function, which is why it uses a simulation. On real data you never know the truth, so you detect the same thing indirectly: a large gap between training and validation score means variance; both scores being poor means bias. Week 7 turns that into a diagnostic you can actually run.

Day 6 takeaway

Day 7 Linear Regression From Scratch, Three Ways

Normal equation, gradient descent and scikit-learn, agreeing

Everything this week, on one problem. You will solve the same regression three ways, exactly, iteratively, and with the library, and get the same answer three times.

The problem

```

1 import numpy as np
2 import pandas as pd

```

```

3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
6 d = df[['tenure_months', 'support_calls', 'monthly_charges']].dropna()
7
8 X = d[['tenure_months', 'support_calls']].to_numpy(dtype=float)
9 y = d['monthly_charges'].to_numpy(dtype=float)
10 print('X', X.shape, ' y', y.shape)
11 print('predicting monthly_charges from tenure and support calls')

```

```

X (2835, 2)  y (2835,)
predicting monthly_charges from tenure and support calls

```

Way 1: the exact solution

Squared error has a closed form. Setting the gradient to zero and solving gives the normal equation, and NumPy will do it in one call.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 d = df[['tenure_months', 'support_calls', 'monthly_charges']].dropna()
6 X = d[['tenure_months', 'support_calls']].to_numpy(dtype=float)
7 y = d['monthly_charges'].to_numpy(dtype=float)
8
9 Xb = np.column_stack([np.ones(len(X)), X])      # a column of 1s for the intercept
10
11 # w = (X'X)^-1 X'y, computed stably rather than by inverting
12 w_exact = np.linalg.lstsq(Xb, y, rcond=None)[0]
13 print('intercept      %8.4f' % w_exact[0])
14 print('tenure_months  %8.4f' % w_exact[1])
15 print('support_calls  %8.4f' % w_exact[2])

```

```

intercept      57.3815
tenure_months  -0.0037
support_calls   1.6707

```

Watch Out

Do not invert the matrix yourself

`np.linalg.inv(X.T @ X) @ X.T @ y` is the textbook formula and it is numerically fragile: when two features are highly correlated the matrix is near-singular and the inverse amplifies floating-point error enormously. `lstsq` solves the system directly and is stable. This matters on real data, where correlated features are the norm, recall tenure and total charges at 0.83.

Way 2: gradient descent

```

1 import numpy as np
2 import pandas as pd

```

```

3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 d = df[['tenure_months', 'support_calls', 'monthly_charges']].dropna()
6 X = d[['tenure_months', 'support_calls']].to_numpy(dtype=float)
7 y = d['monthly_charges'].to_numpy(dtype=float)
8
9 mu, sigma = X.mean(axis=0), X.std(axis=0)
10 Xs = (X - mu) / sigma
11 Xb = np.column_stack([np.ones(len(Xs)), Xs])
12
13 w = np.zeros(Xb.shape[1])
14 lr = 0.1
15 for step in range(1, 3001):
16     error = Xb @ w - y
17     w -= lr * 2 * (Xb.T @ error) / len(y)
18     if step in (1, 100, 1000, 3000):
19         print('step %5d  loss %10.4f' % (step, (error ** 2).mean()))
20
21 # Undo the scaling so the coefficients are comparable with way 1.
22 coef = w[1:] / sigma
23 intercept = w[0] - (w[1:] * mu / sigma).sum()
24 print('\nintercept      %8.4f' % intercept)
25 print('tenure_months  %8.4f' % coef[0])
26 print('support_calls  %8.4f' % coef[1])

```

```

step      1  loss  4080.0173
step    100  loss   537.9793
step   1000  loss   537.9793
step   3000  loss   537.9793

intercept      57.3815
tenure_months  -0.0037
support_calls   1.6707

```

Way 3: the library

```

1 import numpy as np
2 import pandas as pd
3 from sklearn.linear_model import LinearRegression
4
5 df = pd.read_csv('customers.csv').drop_duplicates().copy()
6 d = df[['tenure_months', 'support_calls', 'monthly_charges']].dropna()
7 X = d[['tenure_months', 'support_calls']].to_numpy(dtype=float)
8 y = d['monthly_charges'].to_numpy(dtype=float)
9
10 m = LinearRegression().fit(X, y)
11 print('intercept      %8.4f' % m.intercept_)
12 print('tenure_months  %8.4f' % m.coef_[0])
13 print('support_calls  %8.4f' % m.coef_[1])
14 print('\nR^2 %8.4f' % m.score(X, y))

```

```

intercept      57.3815
tenure_months  -0.0037

```

```
support_calls    1.6707
R^2  0.0233
```

Three routes, one answer. The library is not doing anything you have not now done by hand. It is doing it more carefully, and handling the cases where the naive version breaks.

Reading the coefficients

```
1 import numpy as np
2 import pandas as pd
3 from sklearn.linear_model import LinearRegression
4
5 df = pd.read_csv('customers.csv').drop_duplicates().copy()
6 d = df[['tenure_months', 'support_calls', 'monthly_charges']].dropna()
7 X = d[['tenure_months', 'support_calls']].to_numpy(dtype=float)
8 y = d['monthly_charges'].to_numpy(dtype=float)
9 m = LinearRegression().fit(X, y)
10
11 print('a customer with 12 months tenure and 2 support calls:')
12 print('  predicted charge %.2f' % m.predict([[12, 2]])[0])
13 print('one more support call, everything else held:')
14 print('  predicted charge %.2f' % m.predict([[12, 3]])[0])
15 print('  difference      %.4f ← exactly the coefficient' % m.coef_[1])
```

```
a customer with 12 months tenure and 2 support calls:
  predicted charge 60.68
one more support call, everything else held:
  predicted charge 62.35
  difference      1.6707 ← exactly the coefficient
```

"Holding everything else constant" is doing a lot of work

A coefficient is the change in the prediction for a one-unit change in that feature *with the others fixed*. When features are correlated you cannot actually hold the others fixed, increasing tenure increases total charges in the real world, so the coefficient describes the model, not the world. This is the single biggest source of over-claiming in applied work.

Your assignment

Add `total_charges` as a third feature and refit. It correlates with tenure at 0.83. Record what happens to the tenure coefficient, sign, size, or both. Then refit on a random 60 percent of the rows, twice, and compare. Week 4 names what you are seeing and fixes it.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. Most supervised models compute, at their core:
 - a) A random draw
 - b) A database join
 - c) A sort
 - d) A dot product between a feature vector and a weight vector
2. The derivative of the sigmoid reaches its maximum value of:
 - a) 0.5
 - b) 0.25
 - c) 0.0
 - d) 1.0
3. Features on very different scales slow gradient descent because:
 - a) The gradients overflow
 - b) The loss surface becomes an elongated valley, so a step size suitable for one direction is wrong for another
 - c) NumPy is slower on large numbers
 - d) The learning rate becomes negative
4. A test is 99% accurate and the disease affects 1 in 1,000. A positive result means the patient probably:
 - a) Has a 99% chance of having it
 - b) Has the disease
 - c) Does not have the disease, because the base rate swamps the test's accuracy
 - d) Needs a second opinion only if symptomatic
5. Choosing squared error as your loss is equivalent to assuming:
 - a) There are no outliers
 - b) The target is positive
 - c) Errors are uniformly distributed
 - d) Errors are normally distributed
6. Why does `softmax` subtract the maximum before exponentiating?
 - a) To speed it up
 - b) To prevent overflow. It changes nothing mathematically
 - c) To centre the logits
 - d) To make the output sum to one
7. High bias and low variance describes a model that:
 - a) Memorises the training data

- b) Is too simple to capture the pattern, and makes the same mistake regardless of the sample
- c) Has too many features
- d) Is well fitted

8. The normal equation, gradient descent and **LinearRegression** should:

- a) Agree only if scaled
- b) Give different answers
- c) Agree to within numerical precision on the same data
- d) Only agree for one feature

Answers: 1d, 2b, 3b, 4c, 5d, 6b, 7b, 8c. Anything you got wrong points at the day to re-read, not at a gap in ability.

Regression, Regularisation and Leakage

Week 4

Predict a number, measure it honestly, and meet the two failures that make a good-looking model worthless.

OBJECTIVES

This week you predict a continuous number and learn to distrust the score. It covers the four regression metrics and what each one hides, how flexibility turns into overfitting and how regularisation buys it back, and finishes on the two failures every practitioner meets: correlated features that make coefficients meaningless, and leaked features that make a useless model look excellent.

Day 1 Linear Regression on a Real Target

Fitting, reading the score against the right baseline, and what residuals reveal

Week 3 fitted a line by hand. This week fits real regressions, and the first job is choosing a target worth predicting.

The problem

Predict `monthly_charges` from everything else. It is a genuine business question, what should this customer be paying, given who they are, and unlike week 3's toy version it has real signal in it.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 print('rows %d' % len(df))
13 print('target: mean %.2f std %.2f min %.2f max %.2f'
14       % (df[TARGET].mean(), df[TARGET].std(), df[TARGET].min(), df[TARGET].max()))
15 print('\nby internet service:')
16 print(df.groupby('internet_service')[TARGET].agg(['size', 'mean', 'std']).round(2))

```

```

rows 2835
target: mean 59.41  std 23.47  min 15.00  max 105.96

by internet service:
              size  mean  std
internet_service
DSL            965  55.19  8.99
Fibre optic    1282  79.56  9.18
No internet    588  22.38  7.34

```

Three tight clusters. Internet service almost determines the price, which is a promising sign: there is real structure for a model to find.

Fit it

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X = df[NUM + CAT]
27 y = df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25, random_state=42)
29 from sklearn.linear_model import LinearRegression
30 from sklearn.metrics import mean_absolute_error, root_mean_squared_error, r2_score
31
32 model = Pipeline([('prep', preprocessor()),
33                   ('reg', LinearRegression())]).fit(X_tr, y_tr)
34 pred = model.predict(X_te)
35
36 print('MAE %6.3f' % mean_absolute_error(y_te, pred))
37 print('RMSE %6.3f' % root_mean_squared_error(y_te, pred))
38 print('R2 %6.4f' % r2_score(y_te, pred))
39 print('\nstd of the target was %.2f, so predicting the mean'

```

```
40 ' would give RMSE about that.' % y_te.std()
```

```
MAE    7.200
RMSE   9.033
R2     0.8498
```

```
std of the target was 23.32, so predicting the mean would give RMSE about that.
```

Compare RMSE against the target's standard deviation

Predicting the mean for everybody gives an RMSE equal to the standard deviation, about 23.3 here. Our model gets just over 9. That ratio is the honest way to read a regression score, and it is what R^2 formalises: 0.85 means the model removed 85 percent of the variance the mean-only model left behind.

Residuals are where the truth is

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X = df[NUM + CAT]
27 y = df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25, random_state=42)
29 from sklearn.linear_model import LinearRegression
30
31 model = Pipeline([('prep', preprocessor()),
32                  ('reg', LinearRegression())]).fit(X_tr, y_tr)
33 resid = y_te - model.predict(X_te)
34
35 print('mean    %8.4f    ← should be near zero' % resid.mean())
```

```

36 print('median %8.4f' % resid.median())
37 print('std      %8.4f' % resid.std())
38 print('skew    %8.4f  ← should be near zero' % resid.skew())
39 print('\nlargest under-predictions:')
40 print(resid.nlargest(3).round(2).to_string())
41 print('largest over-predictions:')
42 print(resid.nsmallest(3).round(2).to_string())

```

```

mean    -0.4344  ← should be near zero
median  -0.6659
std      9.0292
skew     0.0188  ← should be near zero

largest under-predictions:
132     26.92
1074    25.72
2735    25.56
largest over-predictions:
3076    -30.66
767     -27.13
2637    -26.33

```

The four assumptions, and which ones matter

Assumption	What breaks if violated	How much you should care
Linearity	Systematic curved pattern in residuals	A lot. The model is simply wrong
Independent errors	Standard errors understated	A lot for time series, rarely otherwise
Constant variance	Predictions unreliable at the extremes	Moderate, affects intervals more than point estimates
Normal errors	Confidence intervals are off	Least of the four; irrelevant for pure prediction

⚠ Watch Out

These are assumptions about residuals, not about your columns

A common misreading: people check whether their *features* are normally distributed. Linear regression makes no such requirement. The normality assumption is about the errors, and it only matters when you want confidence intervals on the coefficients. For prediction you can ignore it.

Check linearity by plotting residuals against predictions

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')

```

```

7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20                             ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22                             ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X = df[NUM + CAT]
27 y = df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25, random_state=42)
29 import matplotlib
30 matplotlib.use('Agg')
31 import matplotlib.pyplot as plt
32 from sklearn.linear_model import LinearRegression
33
34 model = Pipeline([('prep', preprocessor()),
35                  ('reg', LinearRegression())]).fit(X_tr, y_tr)
36 pred = model.predict(X_te)
37 resid = y_te - pred
38
39 fig, ax = plt.subplots(figsize=(7, 4))
40 ax.scatter(pred, resid, s=8, alpha=0.4, color='#2563eb')
41 ax.axhline(0, color='#dc2626', linewidth=1)
42 ax.set_xlabel('Predicted')
43 ax.set_ylabel('Residual')
44 fig.tight_layout()
45 fig.savefig('residuals.png', dpi=120)
46 plt.close(fig)
47
48 # The same check, numerically: is the mean residual flat across the range?
49 bins = pd.qcut(pred, 5, labels=False)
50 print(pd.DataFrame({'pred_bin': bins, 'resid': resid.to_numpy()}
51                  .groupby('pred_bin')['resid'].agg(['size', 'mean', 'std']).round(3))

```

	size	mean	std
pred_bin			
0	142	0.577	7.503
1	142	0.261	9.132
2	141	-0.570	9.695
3	142	-0.045	9.648
4	142	-2.397	8.816

The bin means sit between +0.58 and -2.40 against a residual standard deviation of 9, so there is no strong trend, a slight tendency to over-predict the most expensive customers, and little else. That is roughly what a correctly specified linear model looks like. A clear U-shape here would mean the relationship is curved and you need the polynomial features from day 3.

Day 1 takeaway

Day 2 Regression Metrics and Their Traps

MAE, RMSE, R-squared and MAPE, what each hides

Four metrics, each answering a different question, and each with a way of misleading you.

What each one actually measures

Metric	Units	Reads as	Weakness
MAE	Same as target	Typical error size	Treats a huge miss as merely several small ones
RMSE	Same as target	Error, weighted toward big misses	One bad outlier dominates it
R ²	None	Share of variance explained	Depends on the spread of your test set
MAPE	Percent	Relative error	Explodes when the target is near zero

```

1 import numpy as np
2 from sklearn.metrics import (mean_absolute_error, root_mean_squared_error,
3                             r2_score)
4
5 truth = np.array([20.0, 40.0, 60.0, 80.0, 100.0])
6
7 spread = np.array([25.0, 35.0, 65.0, 75.0, 105.0]) # five misses of 5
8 single = np.array([20.0, 40.0, 60.0, 80.0, 75.0]) # one miss of 25
9
10 for name, p in [('five small misses', spread), ('one big miss', single)]:
11     print('%-18s MAE %5.2f RMSE %5.2f R2 %6.3f'
12           % (name, mean_absolute_error(truth, p),
13             root_mean_squared_error(truth, p), r2_score(truth, p)))

```

```

five small misses MAE 5.00 RMSE 5.00 R2 0.969
one big miss      MAE 5.00 RMSE 11.18 R2 0.844

```

Identical total absolute error, very different RMSE. If one large failure is much worse for your business than several small ones, a delivery estimate, a capacity forecast, use RMSE. If every unit of error costs the same, use MAE.

⚠ Watch Out

R^2 is not comparable across datasets

R^2 measures improvement over predicting the mean, so it depends on how variable the test set happens to be. The same model scores higher on a diverse test set than on a homogeneous one. Comparing your R^2 to a number from a paper on different data tells you nothing.

```
1 import numpy as np
2 from sklearn.metrics import r2_score, root_mean_squared_error
3
4 rng = np.random.default_rng(0)
5
6 for spread in [5, 20, 60]:
7     truth = rng.normal(50, spread, 500)
8     pred = truth + rng.normal(0, 4, 500)      # identical error every time
9     print('target std %2d    RMSE %.2f    R2 %.4f'
10           % (spread, root_mean_squared_error(truth, pred), r2_score(truth, pred)))
```

```
target std  5    RMSE 3.76    R2 0.4489
target std 20    RMSE 4.14    R2 0.9581
target std 60    RMSE 3.92    R2 0.9956
```

The model makes exactly the same size of error in all three cases. R^2 swings from 0.45 to 0.996. RMSE, which is in the units of the thing you care about, does not move.

MAPE and the zero problem

```
1 import numpy as np
2
3 def mape(truth, pred):
4     return np.mean(np.abs((truth - pred) / truth)) * 100
5
6 truth = np.array([100.0, 50.0, 10.0, 0.5])
7 pred = np.array([110.0, 55.0, 11.0, 0.55]) # every one is 10% high
8 print('all errors are 10%: MAPE %.1f%%' % mape(truth, pred))
9
10 truth2 = np.array([100.0, 50.0, 10.0, 0.01])
11 pred2 = np.array([110.0, 55.0, 11.0, 0.5])
12 print('one near-zero actual: MAPE %.1f%%' % mape(truth2, pred2))
```

```
all errors are 10%: MAPE 10.0%
one near-zero actual: MAPE 1232.5%
```

An error of 0.49 on a true value of 0.01 registers as 4,900 percent and swamps everything else. Never use MAPE on a target that can approach zero, and never on one that can be negative.

Cross-validation, not one split

```
1 import numpy as np
```

```

2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.linear_model import LinearRegression
25 from sklearn.model_selection import cross_val_score, KFold
26 import numpy as np
27
28 X, y = df[NUM + CAT], df[TARGET]
29 model = Pipeline([('prep', preprocessor()), ('reg', LinearRegression())])
30
31 cv = KFold(n_splits=5, shuffle=True, random_state=42)
32 for metric in ['r2', 'neg_mean_absolute_error', 'neg_root_mean_squared_error']:
33     scores = cross_val_score(model, X, y, cv=cv, scoring=metric)
34     if metric.startswith('neg_'):
35         scores = -scores
36     print('%-32s %.4f +/- %.4f' % (metric, scores.mean(), scores.std()))

```

```

r2                0.8593 +/- 0.0110
neg_mean_absolute_error    7.0785 +/- 0.2229
neg_root_mean_squared_error 8.7868 +/- 0.2839

```

sklearn negates losses so higher is always better

Every scorer in scikit-learn follows the rule “greater is better”, so errors come back negative. `neg_mean_absolute_error` of -3.6 means an MAE of 3.6. Forgetting to flip the sign is a common way to conclude your model got worse when it improved.

Day 2 takeaway

Day 3 Polynomials, Overfitting and Learning Curves

Adding flexibility, watching it backfire, and diagnosing which fix you need

A straight line cannot fit a curve. The obvious fix, give the model more flexibility, is also the fastest way to make it worse, and watching that happen is the most useful hour of this week.

Add curvature

```
1 import numpy as np
2 from sklearn.preprocessing import PolynomialFeatures
3
4 x = np.array([[2.0, 3.0]])
5 for degree in [1, 2, 3]:
6     pf = PolynomialFeatures(degree=degree, include_bias=False)
7     out = pf.fit_transform(x)
8     print('degree %d → %d features: %s'
9           % (degree, out.shape[1], list(pf.get_feature_names_out())))
```

degree 1 → 2 features: ['x0', 'x1']
degree 2 → 5 features: ['x0', 'x1', 'x0^2', 'x0 x1', 'x1^2']
degree 3 → 9 features: ['x0', 'x1', 'x0^2', 'x0 x1', 'x1^2', 'x0^3', 'x0^2 x1', 'x0 x1^2',
↪ 'x1^3']

Degree 2 adds the squares and the interaction $x_0 x_1$. That interaction term is often the valuable one: it lets the effect of one feature depend on another, which a plain linear model cannot express.

Watch Out

The feature count explodes

Ten features at degree 2 gives 65 columns; at degree 3, 285. At degree 5 you have 3,003 columns from ten. The model gains enough freedom to fit the noise exactly, and the number of rows you would need to constrain it grows just as fast.

Watch it overfit

```
1 import numpy as np
2 from sklearn.preprocessing import PolynomialFeatures
3 from sklearn.linear_model import LinearRegression
4 from sklearn.pipeline import make_pipeline
5 from sklearn.metrics import root_mean_squared_error
6
7 rng = np.random.default_rng(0)
8 n = 40
9 x = rng.uniform(0, 6, n).reshape(-1, 1)
10 y = np.sin(1.5 * x.ravel()) + 0.3 * x.ravel() + rng.normal(0, 0.3, n)
11
12 x_val = rng.uniform(0, 6, 400).reshape(-1, 1)
13 y_val = np.sin(1.5 * x_val.ravel()) + 0.3 * x_val.ravel() + rng.normal(0, 0.3, 400)
14
15 print('%6s %10s %10s' % ('degree', 'train', 'validation'))
16 for degree in [1, 2, 3, 5, 9, 15, 25]:
17     m = make_pipeline(PolynomialFeatures(degree), LinearRegression()).fit(x, y)
```

```

18 tr = root_mean_squared_error(y, m.predict(x))
19 va = root_mean_squared_error(y_val, m.predict(x_val))
20 print('%6d %10.4f %10.4f' % (degree, tr, va))

```

degree	train	validation
1	0.7065	0.7743
2	0.5418	0.6249
3	0.5140	0.6155
5	0.3146	0.3479
9	0.3019	0.3262
15	0.2953	0.3388
25	0.4085	0.4709

Training error falls as flexibility rises, in principle it must, since a richer model can always fit the given points better. Validation error bottoms out at degree 9 and then climbs. The best degree is where validation is lowest, and nothing in the training column would have told you where that was.

Look at degree 25, though: training error goes *up*, from 0.2953 to 0.4085. That is not overfitting, it is arithmetic failing. Raising a value near 6 to the 25th power produces columns spanning twenty orders of magnitude: the matrix becomes so ill-conditioned that the solver cannot work accurately, and the fit degrades. A training error that rises with complexity is a numerical warning, not a statistical one.

The learning curve: more data, or a better model?

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median'))], num)),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                          ('o', OneHotEncoder(handle_unknown='ignore'))], cat)),
22     ])
23
24 from sklearn.linear_model import LinearRegression
25 from sklearn.model_selection import learning_curve
26 import numpy as np
27
28 X, y = df[NUM + CAT], df[TARGET]

```

```

29 model = Pipeline([('prep', preprocessor()), ('reg', LinearRegression())])
30
31 sizes, train_scores, val_scores = learning_curve(
32     model, X, y, cv=5, scoring='neg_root_mean_squared_error',
33     train_sizes=np.linspace(0.1, 1.0, 6), random_state=42)
34
35 print('%8s %10s %10s %8s' % ('rows', 'train', 'val', 'gap'))
36 for n, tr, va in zip(sizes, -train_scores.mean(axis=1), -val_scores.mean(axis=1)):
37     print('%8d %10.4f %10.4f %8.4f' % (n, tr, va, va - tr))

```

rows	train	val	gap
226	8.3919	9.3860	0.9941
635	8.6457	8.8709	0.2252
1043	8.7090	8.8098	0.1008
1451	8.7100	8.8003	0.0903
1859	8.7191	8.7845	0.0654
2268	8.7448	8.7775	0.0327

How to read a learning curve

The two curves converging to a *low* error means you are done. Converging to a *high* error means high bias: more data will not help, you need a better model. A persistent gap means high variance: more data *will* help, or regularise. Here the gap closes to almost nothing at a low error, so this model is well matched to this problem.

Finding the right complexity

```

1 import numpy as np
2 from sklearn.preprocessing import PolynomialFeatures
3 from sklearn.linear_model import LinearRegression
4 from sklearn.pipeline import make_pipeline
5 from sklearn.model_selection import cross_val_score
6
7 rng = np.random.default_rng(0)
8 x = rng.uniform(0, 6, 120).reshape(-1, 1)
9 y = np.sin(1.5 * x.ravel()) + 0.3 * x.ravel() + rng.normal(0, 0.3, 120)
10
11 best = None
12 for degree in range(1, 16):
13     m = make_pipeline(PolynomialFeatures(degree), LinearRegression())
14     rmse = -cross_val_score(m, x, y, cv=5,
15                             scoring='neg_root_mean_squared_error').mean()
16     flag = ''
17     if best is None or rmse < best[1]:
18         best, flag = (degree, rmse), ' ← best so far'
19     if degree <= 8 or flag:
20         print('degree %2d  cv RMSE %8.4f%s' % (degree, rmse, flag))
21
22 print('\nchosen degree: %d' % best[0])

```

```

degree 1  cv RMSE  0.7715  ← best so far
degree 2  cv RMSE  0.6255  ← best so far

```

```

degree 3  cv RMSE  0.6250  ← best so far
degree 4  cv RMSE  0.3330  ← best so far
degree 5  cv RMSE  0.3355
degree 6  cv RMSE  0.3052  ← best so far
degree 7  cv RMSE  0.3086
degree 8  cv RMSE  0.3116

chosen degree: 6

```

Cross-validation picks the complexity for you, and it does so without ever looking at the test set. That is the pattern for every hyperparameter in this course.

Day 3 takeaway

Day 4 Ridge Regression

Keeping a flexible model honest with an L_2 penalty

Yesterday's fix for overfitting was to use a simpler model. Regularisation is the better fix: keep the flexibility, but penalise the model for using it.

The idea

```

1 import numpy as np
2 from sklearn.preprocessing import PolynomialFeatures, StandardScaler
3 from sklearn.linear_model import LinearRegression, Ridge
4 from sklearn.pipeline import make_pipeline
5 from sklearn.metrics import root_mean_squared_error
6
7 rng = np.random.default_rng(0)
8 x = rng.uniform(0, 6, 40).reshape(-1, 1)
9 y = np.sin(1.5 * x.ravel()) + 0.3 * x.ravel() + rng.normal(0, 0.3, 40)
10 x_val = rng.uniform(0, 6, 400).reshape(-1, 1)
11 y_val = np.sin(1.5 * x_val.ravel()) + 0.3 * x_val.ravel() + rng.normal(0, 0.3, 400)
12
13 plain = make_pipeline(PolynomialFeatures(15), StandardScaler(),
14                       LinearRegression()).fit(x, y)
15 ridge = make_pipeline(PolynomialFeatures(15), StandardScaler(),
16                       Ridge(alpha=1.0)).fit(x, y)
17
18 for name, m in [('plain deg 15', plain), ('ridge deg 15', ridge)]:
19     coefs = m[-1].coef_
20     print('%-14s val RMSE %7.4f  largest |coef| %12.2f'
21           % (name, root_mean_squared_error(y_val, m.predict(x_val)),
22             np.abs(coefs).max()))
22

```

```

plain deg 15  val RMSE  0.3480  largest |coef| 1094963376.11
ridge deg 15  val RMSE  0.5461  largest |coef|           0.41

```

Same fifteen-degree polynomial, same forty rows. The unpenalised version produces enormous coefficients that cancel each other out to thread through every training point. Ridge keeps them

small, and generalises far better.

⚠ Watch Out

Ridge requires scaled features

The penalty is on the coefficients, and a coefficient's size depends on its feature's units. A feature measured in pounds gets a small coefficient and is barely penalised; the same feature in thousands of pounds gets a large one and is crushed. Without scaling, the penalty lands arbitrarily. Always put a scaler in front, and note that the intercept is never penalised, which is why it is fitted separately.

What alpha does

```
1 import numpy as np
2 from sklearn.preprocessing import PolynomialFeatures, StandardScaler
3 from sklearn.linear_model import Ridge
4 from sklearn.pipeline import make_pipeline
5 from sklearn.metrics import root_mean_squared_error
6
7 rng = np.random.default_rng(0)
8 x = rng.uniform(0, 6, 40).reshape(-1, 1)
9 y = np.sin(1.5 * x.ravel()) + 0.3 * x.ravel() + rng.normal(0, 0.3, 40)
10 x_val = rng.uniform(0, 6, 400).reshape(-1, 1)
11 y_val = np.sin(1.5 * x_val.ravel()) + 0.3 * x_val.ravel() + rng.normal(0, 0.3, 400)
12
13 print('%10s %10s %10s %14s' % ('alpha', 'train', 'val', 'largest |coef|'))
14 for alpha in [0.0001, 0.01, 1, 100, 10000]:
15     m = make_pipeline(PolynomialFeatures(15), StandardScaler(),
16                       Ridge(alpha=alpha)).fit(x, y)
17     print('%10s %10.4f %10.4f %14.2f'
18           % (alpha, root_mean_squared_error(y, m.predict(x)),
19              root_mean_squared_error(y_val, m.predict(x_val)),
20              np.abs(m[-1].coef_).max()))
```

alpha	train	val	largest coef
0.0001	0.3061	0.3248	18.87
0.01	0.3421	0.3808	6.58
1	0.4653	0.5461	0.41
100	0.5946	0.6756	0.07
10000	0.9418	0.9272	0.00

At tiny alpha you have plain least squares and it overfits. At enormous alpha every coefficient is crushed toward zero and the model underfits, predicting close to a constant. The useful values are in between, and you find them by cross-validation.

RidgeCV finds alpha for you

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
```

```

5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.linear_model import RidgeCV
25 from sklearn.model_selection import cross_val_score
26 import numpy as np
27
28 X, y = df[NUM + CAT], df[TARGET]
29 alphas = np.logspace(-3, 3, 25)
30
31 model = Pipeline([('prep', preprocessor()),
32                  ('reg', RidgeCV(alphas=alphas))]).fit(X, y)
33 print('chosen alpha: %.4f' % model.named_steps['reg'].alpha_)
34
35 rmse = -cross_val_score(model, X, y, cv=5,
36                          scoring='neg_root_mean_squared_error').mean()
37 print('cv RMSE %.4f' % rmse)

```

```

chosen alpha: 1.7783
cv RMSE 8.7776

```

Search alpha on a log scale

The useful range spans orders of magnitude, so `np.logspace(-3, 3, 25)` is the right shape of grid and `np.linspace` is not. If the chosen value lands at either end of your grid, the grid was too narrow. Widen it and search again.

Day 4 takeaway

Day 5 Lasso and ElasticNet

L1 as feature selection, its instability, and the mix that fixes it

Ridge shrinks every coefficient toward zero but never quite to it. Lasso takes them all the way,

which turns regularisation into feature selection.

L1 against L2

	Ridge (L2)	Lasso (L1)
Penalty	Sum of squared coefficients	Sum of absolute coefficients
Effect	Shrinks all, eliminates none	Sets some exactly to zero
Correlated features	Splits the weight between them	Picks one arbitrarily, drops the rest
Use when	All features plausibly matter	You expect most to be irrelevant

```
1 import numpy as np
2 from sklearn.linear_model import Ridge, Lasso
3 from sklearn.preprocessing import StandardScaler
4 from sklearn.pipeline import make_pipeline
5
6 rng = np.random.default_rng(0)
7 n, p = 200, 20
8 X = rng.normal(size=(n, p))
9 # Only the first three columns matter; the other seventeen are noise.
10 true_coef = np.zeros(p)
11 true_coef[:3] = [4.0, -3.0, 2.0]
12 y = X @ true_coef + rng.normal(0, 1.0, n)
13
14 for name, model in [('ridge', Ridge(alpha=1.0)), ('lasso', Lasso(alpha=0.1))]:
15     m = make_pipeline(StandardScaler(), model).fit(X, y)
16     coef = m[-1].coef_
17     print('%-6s first three %s exactly zero: %d of %d'
18           % (name, np.round(coef[:3], 2), (coef == 0).sum(), p))
```

```
ridge first three [ 3.84 -2.94  1.92] exactly zero: 0 of 20
lasso first three [ 3.72 -2.87  1.77] exactly zero: 15 of 20
```

Ridge leaves all twenty coefficients non-zero, including seventeen that should not be there. Lasso zeroes most of the noise columns outright, giving you a model you can read.

The path from full model to empty one

```
1 import numpy as np
2 from sklearn.linear_model import Lasso
3 from sklearn.preprocessing import StandardScaler
4 from sklearn.pipeline import make_pipeline
5
6 rng = np.random.default_rng(0)
7 X = rng.normal(size=(200, 20))
8 true_coef = np.zeros(20)
9 true_coef[:3] = [4.0, -3.0, 2.0]
10 y = X @ true_coef + rng.normal(0, 1.0, 200)
11
12 print('%8s %10s %s' % ('alpha', 'kept', 'which'))
13 for alpha in [0.001, 0.05, 0.2, 0.5, 1.0, 3.0]:
14     m = make_pipeline(StandardScaler(), Lasso(alpha=alpha, max_iter=5000)).fit(X, y)
```

```

15 kept = np.flatnonzero(m[-1].coef_)
16 print('%8s %10d %s' % (alpha, len(kept), list(kept[:6])))

alpha      kept which
0.001      20 [np.int64(0), np.int64(1), np.int64(2), np.int64(3), np.int64(4),
    ↪ np.int64(5)]
0.05       9 [np.int64(0), np.int64(1), np.int64(2), np.int64(5), np.int64(7),
    ↪ np.int64(9)]
0.2        3 [np.int64(0), np.int64(1), np.int64(2)]
0.5        3 [np.int64(0), np.int64(1), np.int64(2)]
1.0        3 [np.int64(0), np.int64(1), np.int64(2)]
3.0        1 [np.int64(0)]

```

As alpha rises the model discards features, and the last three standing are exactly the three that generated the data. That is the best possible outcome and it will not always happen, but it shows what lasso is for.

⚠ Watch Out

Lasso is unstable when features are correlated

Given two nearly identical columns, lasso keeps one and zeroes the other, and which one it keeps can flip on a slightly different sample. Do not read “lasso dropped this feature” as “this feature does not matter”. It may simply be standing next to a twin.

```

1 import numpy as np
2 from sklearn.linear_model import Lasso
3 from sklearn.preprocessing import StandardScaler
4 from sklearn.pipeline import make_pipeline
5
6 rng = np.random.default_rng(1)
7 n = 150
8 a = rng.normal(size=n)
9 twin = a + rng.normal(0, 0.05, n)      # almost the same column
10 noise = rng.normal(size=(n, 5))
11 y = 3 * a + rng.normal(0, 0.5, n)
12
13 for seed in range(4):
14     idx = np.random.default_rng(seed).choice(n, 120, replace=False)
15     X = np.column_stack([a, twin, noise])[idx]
16     m = make_pipeline(StandardScaler(), Lasso(alpha=0.15)).fit(X, y[idx])
17     print('sample %d: coef on a %6.3f  coef on its twin %6.3f'
18           % (seed, m[-1].coef_[0], m[-1].coef_[1]))

```

```

sample 0: coef on a  1.966  coef on its twin  0.229
sample 1: coef on a  2.102  coef on its twin  0.468
sample 2: coef on a  2.379  coef on its twin  0.000
sample 3: coef on a  2.278  coef on its twin  0.064

```

Same underlying truth, four resamples, and the weight jumps between the two twins. Ridge would have split it evenly and stably between them.

ElasticNet: both penalties

```
1 import numpy as np
2 from sklearn.linear_model import ElasticNetCV
3 from sklearn.preprocessing import StandardScaler
4 from sklearn.pipeline import make_pipeline
5
6 rng = np.random.default_rng(1)
7 n = 150
8 a = rng.normal(size=n)
9 twin = a + rng.normal(0, 0.05, n)
10 noise = rng.normal(size=(n, 5))
11 X = np.column_stack([a, twin, noise])
12 y = 3 * a + rng.normal(0, 0.5, n)
13
14 m = make_pipeline(StandardScaler(),
15                   ElasticNetCV(l1_ratio=[0.1, 0.5, 0.9, 1.0], cv=5,
16                               random_state=0, max_iter=5000)).fit(X, y)
17 net = m[-1]
18 print('chosen l1_ratio %.2f   alpha %.4f' % (net.l1_ratio_, net.alpha_))
19 print('coef on a %.3f, on its twin %.3f' % (net.coef_[0], net.coef_[1]))
20 print('noise coefficients:', np.round(net.coef_[2:], 3))
```

```
chosen l1_ratio 0.90   alpha 0.0336
coef on a 1.589, on its twin 1.007
noise coefficients: [0.    0.    0.047 0.015 0.003]
```

`l1_ratio` mixes the two: 1.0 is pure lasso, 0.0 pure ridge. ElasticNet keeps lasso's ability to discard irrelevant columns while sharing weight between correlated ones rather than choosing arbitrarily.

Day 5 takeaway

Day 6 Multicollinearity and Target Leakage

The failure that corrupts coefficients, and the one that fakes success

Two failures that both look like success. One inflates your test score to something impossible; the other makes your coefficients meaningless while the score stays fine.

Multicollinearity

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
```

```

9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 d = df.dropna(subset=['total_charges'])
13 print(d[['tenure_months', 'monthly_charges', 'total_charges']].corr().round(3))

```

	tenure_months	monthly_charges	total_charges
tenure_months	1.000	0.001	0.828
monthly_charges	0.001	1.000	0.450
total_charges	0.828	0.450	1.000

```

1 import numpy as np
2 from sklearn.linear_model import LinearRegression, Ridge
3
4 rng = np.random.default_rng(0)
5 n = 300
6 x1 = rng.normal(0, 1, n)
7 x2 = x1 + rng.normal(0, 0.02, n)      # correlation about 0.9998
8 y = 2 * x1 + rng.normal(0, 0.3, n)
9
10 print('correlation between the two features: %.4f' % np.corrcoef(x1, x2)[0, 1])
11 print('\nrefitting on four different subsamples:')
12 print('%12s %10s %10s %10s' % ('', 'coef x1', 'coef x2', 'sum'))
13 for seed in range(4):
14     idx = np.random.default_rng(seed).choice(n, 250, replace=False)
15     X = np.column_stack([x1, x2])[idx]
16     m = LinearRegression().fit(X, y[idx])
17     print('%12s %10.3f %10.3f %10.3f'
18           % ('sample %d' % seed, m.coef_[0], m.coef_[1], m.coef_.sum()))

```

```
correlation between the two features: 0.9998
```

```
refitting on four different subsamples:
```

	coef x1	coef x2	sum
sample 0	1.421	0.594	2.015
sample 1	0.799	1.212	2.011
sample 2	0.950	1.060	2.010
sample 3	1.364	0.640	2.004

The individual coefficients swing violently, and their sum stays near 2, the truth. The model knows the total effect perfectly well; it simply cannot tell which of two identical columns deserves the credit. Reporting either coefficient on its own would be meaningless.

Measure it with the variance inflation factor

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8

```

```

9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 import numpy as np
13 from sklearn.linear_model import LinearRegression
14
15 d = df.dropna(subset=['total_charges'])
16 cols = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
17 M = d[cols].to_numpy(dtype=float)
18
19 print('%-18s %8s' % ('feature', 'VIF'))
20 for i, name in enumerate(cols):
21     others = np.delete(M, i, axis=1)
22     r2 = LinearRegression().fit(others, M[:, i]).score(others, M[:, i])
23     print('%-18s %8.2f' % (name, 1 / (1 - r2)))

```

feature	VIF
tenure_months	7.03
support_calls	1.02
monthly_charges	2.79
total_charges	8.82

How to read a VIF

VIF is how much the variance of a coefficient is inflated by its correlation with the others. Below 5 is fine, above 10 is a problem worth acting on. The fixes, in order of preference: drop one of the pair, combine them into a single meaningful feature, or use ridge, which handles collinearity gracefully by design.

Leakage, which is the more dangerous of the two

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 d = df.dropna(subset=['total_charges']).copy()
13 d = d[d['tenure_months'] > 0]
14 d['implied_monthly'] = d['total_charges'] / d['tenure_months']
15
16 print('correlation with the target: %.4f'
17       % d['implied_monthly'].corr(d[TARGET]))
18 print('\nfirst five rows:')
19 print(d[['tenure_months', 'total_charges', 'implied_monthly', TARGET]]
20       .head().round(2).to_string(index=False))

```

```
correlation with the target: 0.9956
```

```
first five rows:
```

tenure_months	total_charges	implied_monthly	monthly_charges
10	904.14	90.41	90.21
8	621.20	77.65	81.45
3	181.14	60.38	57.29
3	44.38	14.79	15.00
21	564.82	26.90	25.85

`total_charges` divided by `tenure_months` is, almost exactly, `monthly_charges`, because that is how the bill was calculated. A perfectly sensible-looking engineered feature reconstructs the answer.

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X = df[NUM + CAT]
27 y = df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25, random_state=42)
29 from sklearn.linear_model import LinearRegression
30 from sklearn.metrics import r2_score
31
32 honest = Pipeline([('prep', preprocessor()),
33                     ('reg', LinearRegression())]).fit(X_tr, y_tr)
34 print('honest model      R2 %.4f' % r2_score(y_te, honest.predict(X_te)))
35
36 leaky_df = df.dropna(subset=['total_charges']).copy()
37 leaky_df = leaky_df[leaky_df['tenure_months'] > 0]
38 leaky_df['implied_monthly'] = (leaky_df['total_charges']
39                               / leaky_df['tenure_months'])
40 LNUM = NUM + ['implied_monthly']
41 X1, y1 = leaky_df[LNUM + CAT], leaky_df[TARGET]
```

```

42 Xl_tr, Xl_te, yl_tr, yl_te = train_test_split(Xl, yl, test_size=0.25,
43                                             random_state=42)
44 leaky = Pipeline([('prep', preprocessor(num=LNUM)),
45                  ('reg', LinearRegression())]).fit(Xl_tr, yl_tr)
46 print('with the leak      R2 %.4f' % r2_score(yl_te, leaky.predict(Xl_te)))

```

```

honest model      R2 0.8498
with the leak     R2 0.9920

```

⚠ Watch Out

An R^2 that jumps to 0.99 is not good news

It is the strongest available signal that a feature knows the answer. The instinct on seeing a score like that should be suspicion, not celebration. Ask of every feature: *would I have this value, for a new customer, at the moment I need the prediction?* If the answer is no, or “only after the outcome is known”, it leaks.

Where leakage hides

- **Aggregates computed over the whole dataset:** a per-category mean that included the test rows.
- **Anything recorded after the event:** a cancellation reason code, when predicting cancellation.
- **Identifiers correlated with time:** sequential IDs encode when a row was created.
- **Duplicated rows across the split:** exactly the 138 you removed in week 1.
- **Preprocessing fitted before the split:** the reason every scaler in this course sits inside a pipeline.

Day 6 takeaway

Day 7 A Complete Regression Project

Every decision of the week, applied end to end

A complete regression, start to finish, with every decision this week justified.

1. Frame and split

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']

```

```

10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.model_selection import train_test_split
13
14 X = df[NUM + CAT]
15 y = df[TARGET]
16 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25, random_state=42)
17 print('train %d rows, test %d rows' % (len(X_tr), len(X_te)))
18 print('target mean: train %.2f, test %.2f' % (y_tr.mean(), y_te.mean()))

```

```

train 2126 rows, test 709 rows
target mean: train 59.48, test 59.20

```

2. Establish the floor

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.model_selection import train_test_split
13
14 X = df[NUM + CAT]
15 y = df[TARGET]
16 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25, random_state=42)
17 from sklearn.dummy import DummyRegressor
18 from sklearn.metrics import root_mean_squared_error, r2_score
19
20 base = DummyRegressor(strategy='mean').fit(X_tr, y_tr)
21 pred = base.predict(X_te)
22 print('predicting the mean: RMSE %.3f  R2 %.4f'
23       % (root_mean_squared_error(y_te, pred), r2_score(y_te, pred)))

```

```

predicting the mean: RMSE 23.310  R2 -0.0001

```

R^2 of essentially zero, by construction. Every model from here has to beat an RMSE of about 23.4.

3. Compare candidates by cross-validation

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()

```

```

6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X = df[NUM + CAT]
27 y = df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25, random_state=42)
29 from sklearn.linear_model import LinearRegression, RidgeCV, LassoCV
30 from sklearn.ensemble import RandomForestRegressor
31 from sklearn.model_selection import cross_val_score, KFold
32 import numpy as np
33
34 cv = KFold(n_splits=5, shuffle=True, random_state=42)
35 candidates = {
36     'linear': LinearRegression(),
37     'ridge': RidgeCV(alphas=np.logspace(-3, 3, 25)),
38     'lasso': LassoCV(cv=5, random_state=0, max_iter=5000),
39     'forest': RandomForestRegressor(n_estimators=200, random_state=42),
40 }
41 for name, reg in candidates.items():
42     pipe = Pipeline([('prep', preprocessor()), ('reg', reg)])
43     s = -cross_val_score(pipe, X_tr, y_tr, cv=cv,
44         scoring='neg_root_mean_squared_error')
45     print('%s cv RMSE %.4f +/- %.4f' % (name, s.mean(), s.std()))

```

```

linear      cv RMSE 8.7237 +/- 0.1761
ridge       cv RMSE 8.7234 +/- 0.1735
lasso       cv RMSE 8.7068 +/- 0.1714
forest      cv RMSE 9.9159 +/- 0.2897

```

Cross-validate on the training set only

The test set has not been touched yet and must not be, until the model is chosen. Every comparison, every hyperparameter, every feature decision comes from cross-validation within the training data. The test set is spent the first time you look at it.

4. Fit the winner and score once

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X = df[NUM + CAT]
27 y = df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25, random_state=42)
29 from sklearn.linear_model import RidgeCV
30 from sklearn.metrics import (root_mean_squared_error, mean_absolute_error,
31                             r2_score)
32 import numpy as np
33
34 final = Pipeline([('prep', preprocessor()),
35                  ('reg', RidgeCV(alphas=np.logspace(-3, 3, 25)))]
36 final.fit(X_tr, y_tr)
37 pred = final.predict(X_te)
38
39 print('TEST SET')
40 print(' RMSE %.4f' % root_mean_squared_error(y_te, pred))
41 print(' MAE  %.4f' % mean_absolute_error(y_te, pred))
42 print(' R2   %.4f' % r2_score(y_te, pred))
43 print('\nagainst a mean-only RMSE of %.3f' % y_te.std())

```

```

TEST SET
RMSE 9.0293
MAE  7.2018
R2   0.8499

```

```
against a mean-only RMSE of 23.325
```

5. Explain it

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20                             ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22                             ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X = df[NUM + CAT]
27 y = df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25, random_state=42)
29 from sklearn.linear_model import RidgeCV
30 import numpy as np
31 import pandas as pd
32
33 final = Pipeline([('prep', preprocessor()),
34                   ('reg', RidgeCV(alphas=np.logspace(-3, 3, 25)))]).fit(X_tr, y_tr)
35
36 names = final.named_steps['prep'].get_feature_names_out()
37 coefs = pd.Series(final.named_steps['reg'].coef_, index=names)
38 print(coefs.reindex(coefs.abs().sort_values(ascending=False).index)
39         .head(8).round(3).to_string())
```

```
cat__internet_service_No internet      -30.107
cat__internet_service_Fibre optic      27.346
cat__internet_service_DSL               2.762
cat__contract_Month-To-Month            0.322
cat__payment_method_Electronic check   -0.304
num__tenure_months                     0.299
cat__contract_One Year                  -0.290
cat__payment_method_Bank transfer       0.215
```

Fibre optic adds about 27 against the baseline and having no internet subtracts about 30, while contract type and payment method barely move the price at all. The model has recovered the pricing structure, which is exactly what the generator in week 1 put there, and note that the

features which drive *churn* are almost irrelevant to *price*. Different target, different important features.

6. Check the residuals one last time

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 df = df.dropna(subset=['monthly_charges'])
8
9 NUM = ['tenure_months', 'support_calls']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'monthly_charges'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20                          ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22                          ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X = df[NUM + CAT]
27 y = df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25, random_state=42)
29 from sklearn.linear_model import RidgeCV
30 import numpy as np
31 import pandas as pd
32
33 final = Pipeline([('prep', preprocessor()),
34                  ('reg', RidgeCV(alphas=np.logspace(-3, 3, 25)))]).fit(X_tr, y_tr)
35 pred = final.predict(X_te)
36 resid = y_te - pred
37
38 print('residual mean %.4f, std %.4f, skew %.4f'
39       % (resid.mean(), resid.std(), resid.skew()))
40 worst = resid.abs().nlargest(3).index
41 print('\nworst three predictions:')
42 print(pd.DataFrame({'actual': y_te.loc[worst].round(2),
43                    'predicted': pred[[list(y_te.index).index(i) for i in
44                    ↪ worst]].round(2),
45                    'service': df.loc[worst, 'internet_service']}).to_string())

```

```
residual mean -0.4340, std 9.0253, skew 0.0221
```

```
worst three predictions:
      actual  predicted  service
3076  48.61    79.22  Fibre optic
767   52.42    79.49  Fibre optic
132  105.96    78.98  Fibre optic
```

Your assignment

Rebuild this predicting `total_charges` instead. You will find an R^2 above 0.99, because tenure multiplied by monthly charges *is* the total. Then remove both of those features and try again. Write down which version you would deploy, and why the first one is worthless despite scoring better.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. An R-squared of 0.0 means the model is:
 - a) Undefined
 - b) Perfect
 - c) As good as always predicting the mean of the target
 - d) Worse than useless
2. Compared with MAE, RMSE:
 - a) Is always smaller
 - b) Cannot be used for regression
 - c) Ignores outliers
 - d) Penalises large errors more heavily, because they are squared
3. MAPE is a poor choice when:
 - a) The target can be zero or near zero, which makes the percentage explode
 - b) There are many features
 - c) The model is linear
 - d) The target is large
4. Training error keeps falling as polynomial degree rises, while validation error turns upward. That gap is:
 - a) Bias
 - b) Variance. The model is fitting the sample rather than the pattern
 - c) Leakage
 - d) Underfitting
5. The practical difference between ridge and lasso is that lasso:

- a) Requires no scaling
 - b) Is faster
 - c) Can drive coefficients to exactly zero, so it selects features
 - d) Never overfits
6. Two highly correlated features are given to lasso. The likely result is:
- a) Both are zeroed
 - b) The fit fails
 - c) Both get equal weight
 - d) One is kept and the other zeroed, and which one can change with a small change in the data
7. A model reaches R-squared 0.98 on a business problem where 0.6 would be good. The first thing to suspect is:
- a) Target leakage, a feature that encodes the answer
 - b) Too little regularisation
 - c) An unlucky split
 - d) A very good model
8. Why must scaling happen inside the pipeline rather than before the split?
- a) It is faster
 - b) Otherwise the mean and standard deviation are computed using test rows
 - c) Scalers cannot handle splits
 - d) It changes the column order

Answers: 1c, 2d, 3a, 4b, 5c, 6d, 7a, 8b. Anything you got wrong points at the day to re-read, not at a gap in ability.

Classification Fundamentals

Week 5

Logistic regression, kNN, naive Bayes and SVMs, and the threshold decision that masters more than any of them.

OBJECTIVES

Predicting a label rather than a number. This week covers the four classical classifiers and, more importantly, the decision that sits on top of all of them: where to cut the predicted probability. Week 1 left you with a model that missed most churners; day 2 fixes it without retraining anything.

Day 1 Logistic Regression

The sigmoid, log-odds, odds ratios, and regularisation you did not ask for

Linear regression predicts a number on an unbounded scale. A probability lives between 0 and 1. Logistic regression is what you get when you insist on that constraint, and it remains the model to beat on tabular data.

The sigmoid

```
1 import numpy as np
2
3 def sigmoid(z):
4     return 1 / (1 + np.exp(-z))
5
6 for z in [-6, -2, -1, 0, 1, 2, 6]:
7     print('z = %3d → probability %.4f' % (z, sigmoid(z)))
```

```
z = -6 → probability 0.0025
z = -2 → probability 0.1192
z = -1 → probability 0.2689
z = 0 → probability 0.5000
z = 1 → probability 0.7311
z = 2 → probability 0.8808
z = 6 → probability 0.9975
```

Any real number in, a probability out. Zero maps to 0.5, and the curve saturates at both ends, which is why a very confident model needs an enormous change in z to become slightly more confident.

Log-odds, which is the scale the coefficients live on

```

1 import numpy as np
2
3 def sigmoid(z):
4     return 1 / (1 + np.exp(-z))
5
6 print('%10s %10s %12s' % ('probability', 'odds', 'log-odds (z)'))
7 for prob in [0.1, 0.25, 0.5, 0.75, 0.9]:
8     odds = prob / (1 - prob)
9     print('%10.2f %10.3f %12.4f' % (prob, odds, np.log(odds)))

```

probability	odds	log-odds (z)
0.10	0.111	-2.1972
0.25	0.333	-1.0986
0.50	1.000	0.0000
0.75	3.000	1.0986
0.90	9.000	2.1972

A coefficient is a change in log-odds

That is why they are hard to read directly. Exponentiate one and you get an *odds ratio*: a coefficient of 0.7 means $\exp(0.7) = 2.01$, so a one-unit increase roughly doubles the odds. Doubling the odds is not doubling the probability, from 0.1 it goes to 0.18, from 0.5 it goes to 0.67.

Fit it on the churn problem

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median'))],
19         ('s', StandardScaler()))], NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21         ('o', OneHotEncoder(handle_unknown='ignore'))], CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]

```

```

26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 from sklearn.metrics import roc_auc_score, log_loss, accuracy_score
30
31 model = Pipeline([('prep', preprocessor()),
32                  ('clf', LogisticRegression(max_iter=1000,
33                                             random_state=42))]).fit(X_tr, y_tr)
34 proba = model.predict_proba(X_te)[: , 1]
35
36 print('accuracy %.4f' % accuracy_score(y_te, model.predict(X_te)))
37 print('ROC AUC   %.4f' % roc_auc_score(y_te, proba))
38 print('log loss  %.4f' % log_loss(y_te, proba))
39 print('\npredicted probabilities: min %.3f, median %.3f, max %.3f'
40       % (proba.min(), np.median(proba), proba.max()))

```

```

accuracy 0.7800
ROC AUC   0.8174
log loss  0.4451

```

```

predicted probabilities: min 0.000, median 0.232, max 0.829

```

Read the coefficients as odds ratios

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                          ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                          ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression

```

```

29 import numpy as np
30 import pandas as pd
31
32 model = Pipeline([('prep', preprocessor()),
33                  ('clf', LogisticRegression(max_iter=1000,
34                                             random_state=42))]).fit(X_tr, y_tr)
35 names = model.named_steps['prep'].get_feature_names_out()
36 coef = model.named_steps['clf'].coef_[0]
37
38 table = pd.DataFrame({'coefficient': coef, 'odds_ratio': np.exp(coef)},
39                      index=names).sort_values('coefficient')
40 print(table.round(3).to_string())

```

	coefficient	odds_ratio
cat__contract_Two Year	-1.454	0.234
num__tenure_months	-0.978	0.376
cat__payment_method_Credit card	-0.205	0.815
cat__payment_method_Bank transfer	-0.174	0.841
cat__internet_service_No internet	-0.163	0.850
cat__has_dependents_Yes	-0.102	0.903
cat__payment_method_Mailed check	-0.088	0.916
cat__internet_service_DSL	-0.060	0.942
cat__has_dependents_No	0.087	1.091
cat__contract_One Year	0.128	1.137
cat__internet_service_Fibre optic	0.207	1.230
num__support_calls	0.279	1.322
num__monthly_charges	0.436	1.547
cat__payment_method_Electronic check	0.451	1.570
cat__contract_Month-To-Month	1.310	3.706

An odds ratio above 1 pushes toward churn, below 1 pushes away. Month-to-month roughly triples the odds against the average; a two-year contract cuts them to about a quarter. Compare that with the generator in week 1: it used +1.55 and -0.85 on exactly those two.

⚠ Watch Out

Odds ratios on scaled features are per standard deviation

The numeric columns went through `StandardScaler`, so their coefficients describe a one *standard deviation* change, not one month or one pound. To talk to a business audience in real units, divide the coefficient by the scaler's `scale_` for that column, or fit an unscaled model purely for explanation.

Regularisation is on by default

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']

```

```

10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 from sklearn.metrics import roc_auc_score
30 import numpy as np
31
32 print('%10s %12s %14s' % ('C', 'test AUC', 'sum |coef|'))
33 for C_val in [0.001, 0.01, 0.1, 1.0, 100.0]:
34     m = Pipeline([('prep', preprocessor()),
35                   ('clf', LogisticRegression(C=C_val, max_iter=2000,
36                                             random_state=42))]).fit(X_tr, y_tr)
37     auc = roc_auc_score(y_te, m.predict_proba(X_te)[: , 1])
38     print('%10s %12.4f %14.3f'
39           % (C_val, auc, np.abs(m.named_steps['clf'].coef_).sum()))

```

C	test AUC	sum coef
0.001	0.8205	0.885
0.01	0.8230	3.114
0.1	0.8198	5.147
1.0	0.8174	6.122
100.0	0.8169	8.108

⚠ Watch Out

C is the inverse of alpha

Ridge takes **alpha**, where larger means more penalty. `LogisticRegression` takes **C**, where *smaller* means more penalty. The default of **C=1.0** means you are already regularising whether you meant to or not, which surprises people comparing against an unpenalised implementation elsewhere.

Day 1 takeaway

Day 2 The Decision Threshold

Why 0.5 is almost always wrong, and how to choose the number that is right

Week 1 ended with a model that found fewer than half the churners, and the promise that the model was fine and the threshold was wrong. Today you collect on that.

predict() is predict_proba() plus an arbitrary cut

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                         stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 import numpy as np
30
31 model = Pipeline([('prep', preprocessor()),
32                  ('clf', LogisticRegression(max_iter=1000,
33                                             random_state=42))]).fit(X_tr, y_tr)
34 proba = model.predict_proba(X_te)[: , 1]
35
36 print('predict() agrees with proba > 0.5 on %d of %d rows'
37       % ((model.predict(X_te) == (proba > 0.5).astype(int)).sum(), len(proba)))
38 print('\nrows with probability above:')
39 for t in [0.2, 0.3, 0.4, 0.5, 0.6]:
40     print(' %.1f %4d (%.1f%% of the test set)'
41           % (t, (proba > t).sum(), 100 * (proba > t).mean()))

```

```
predict() agrees with proba > 0.5 on 750 of 750 rows
```

```
rows with probability above:
```

0.2	400	(53.3% of the test set)
0.3	294	(39.2% of the test set)
0.4	209	(27.9% of the test set)
0.5	138	(18.4% of the test set)
0.6	70	(9.3% of the test set)

0.5 is a default, not a decision

It is only optimal when the classes are balanced and the two kinds of mistake cost the same. Churn is 27 percent, and missing a leaver who was worth keeping costs far more than a wasted retention call. Neither condition holds, so 0.5 is simply the wrong number.

The four outcomes

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 from sklearn.metrics import confusion_matrix, precision_score, recall_score
30 import numpy as np
31
32 model = Pipeline([('prep', preprocessor()),
33                   ('clf', LogisticRegression(max_iter=1000,
34                                              random_state=42))]).fit(X_tr, y_tr)
35 proba = model.predict_proba(X_te)[: , 1]
36
37 print('%10s %6s %6s %6s %6s %10s %8s'
38       % ('threshold', 'TN', 'FP', 'FN', 'TP', 'precision', 'recall'))
39 for t in [0.15, 0.2, 0.27, 0.35, 0.5, 0.7]:

```

```

40 pred = (proba >= t).astype(int)
41 tn, fp, fn, tp = confusion_matrix(y_te, pred).ravel()
42 print('%10.2f %6d %6d %6d %6d %10.3f %8.3f'
43       % (t, tn, fp, fn, tp,
44          precision_score(y_te, pred, zero_division=0),
45          recall_score(y_te, pred)))

```

threshold	TN	FP	FN	TP	precision	recall
0.15	287	262	18	183	0.411	0.910
0.20	328	221	22	179	0.448	0.891
0.27	384	165	42	159	0.491	0.791
0.35	433	116	66	135	0.538	0.672
0.50	498	51	114	87	0.630	0.433
0.70	544	5	180	21	0.808	0.104

Term	Means	Cost in a churn programme
True positive	Predicted churn, did churn	A save, the point of the exercise
False positive	Predicted churn, stayed	A wasted retention offer
False negative	Predicted stay, churned	A customer lost silently
True negative	Predicted stay, stayed	Nothing spent

Choose the threshold from the cost of being wrong

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                          ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                          ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression

```

```

29 from sklearn.metrics import confusion_matrix
30 import numpy as np
31
32 model = Pipeline([('prep', preprocessor()),
33                   ('clf', LogisticRegression(max_iter=1000,
34                                              random_state=42))]).fit(X_tr, y_tr)
35 proba = model.predict_proba(X_te)[: , 1]
36
37 OFFER_COST = 30      # what a retention offer costs us
38 CUSTOMER_VALUE = 250 # what keeping a customer is worth
39 SAVE_RATE = 0.35     # share of contacted churners we actually keep
40
41 best = None
42 for t in np.arange(0.05, 0.95, 0.01):
43     pred = (proba >= t).astype(int)
44     tn, fp, fn, tp = confusion_matrix(y_te, pred).ravel()
45     value = tp * SAVE_RATE * CUSTOMER_VALUE - (tp + fp) * OFFER_COST
46     if best is None or value > best[1]:
47         best = (t, value, tp, fp)
48
49 print('best threshold %.2f → net value %.0f (contacted %d, saved about %.0f)'
50       % (best[0], best[1], best[2] + best[3], best[2] * SAVE_RATE))
51
52 pred_default = (proba >= 0.5).astype(int)
53 tn, fp, fn, tp = confusion_matrix(y_te, pred_default).ravel()
54 print('at the default 0.5 → net value %.0f'
55       % (tp * SAVE_RATE * CUSTOMER_VALUE - (tp + fp) * OFFER_COST))

```

```

best threshold 0.34 → net value 4537 (contacted 260, saved about 49)
at the default 0.5 → net value 3472

```

Same model, same data, no retraining. Moving the threshold is free, and it is the single highest-return change available on most classification projects.

ROC and precision-recall curves

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor():
18     return ColumnTransformer([

```

```

18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                         stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 from sklearn.metrics import (roc_curve, precision_recall_curve,
30                             roc_auc_score, average_precision_score)
31 import numpy as np
32
33 model = Pipeline([('prep', preprocessor()),
34                  ('clf', LogisticRegression(max_iter=1000,
35                                             random_state=42))]).fit(X_tr, y_tr)
36 proba = model.predict_proba(X_te)[: , 1]
37
38 fpr, tpr, _ = roc_curve(y_te, proba)
39 prec, rec, _ = precision_recall_curve(y_te, proba)
40 print('ROC AUC          %.4f' % roc_auc_score(y_te, proba))
41 print('average precision  %.4f' % average_precision_score(y_te, proba))
42 print('baseline for AP    %.4f ← the positive class rate' % y_te.mean())
43 print('baseline for ROC AUC 0.5000')

```

```

ROC AUC          0.8174
average precision 0.5993
baseline for AP   0.2680 ← the positive class rate
baseline for ROC AUC 0.5000

```

⚠ Watch Out

ROC AUC flatters models on imbalanced data

The false positive rate has the large negative class in its denominator, so a flood of false positives barely moves it. A model can look strong on ROC and be useless in practice. The precision-recall curve uses no true negatives at all, so it stays honest, and its baseline is the positive rate, 0.27 here, not 0.5. Report both, and lead with average precision when positives are rare.

Day 2 takeaway

Day 3 k-Nearest Neighbours

Voting among similar rows, and why distance stops working in high dimensions

The simplest possible classifier: find the most similar customers you have seen and copy their answer. It has no training phase at all.

How it works

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.neighbors import KNeighborsClassifier
29 from sklearn.metrics import roc_auc_score, accuracy_score
30
31 for k in [1, 5, 25, 101]:
32     m = Pipeline([('prep', preprocessor()),
33                   ('clf', KNeighborsClassifier(n_neighbors=k))]).fit(X_tr, y_tr)
34     print('k=%3d  train acc %.4f  test acc %.4f  test AUC %.4f'
35          % (k, accuracy_score(y_tr, m.predict(X_tr)),
36             accuracy_score(y_te, m.predict(X_te)),
37             roc_auc_score(y_te, m.predict_proba(X_te)[: , 1])))
```

```
k= 1  train acc 0.9991  test acc 0.7333  test AUC 0.6507
k= 5  train acc 0.8258  test acc 0.7453  test AUC 0.7517
k= 25  train acc 0.7924  test acc 0.7640  test AUC 0.8023
k=101  train acc 0.7831  test acc 0.7653  test AUC 0.8122
```

k=1 scores perfectly on training data, by definition

The nearest neighbour of a training row is itself, at distance zero. That perfect training accuracy is the purest example of why training scores are meaningless. Larger k averages over more neighbours, smoothing the decision boundary and trading variance for bias.

Scaling is not optional here

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.model_selection import train_test_split
12
13 X, y = df[NUM + CAT], df[TARGET]
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 from sklearn.neighbors import KNeighborsClassifier
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler
19 from sklearn.pipeline import make_pipeline
20 from sklearn.metrics import roc_auc_score
21
22 # Two columns on wildly different scales: months against pounds.
23 WIDE = ['tenure_months', 'total_charges']
24 Xn_tr, Xn_te = df.loc[X_tr.index, WIDE], df.loc[X_te.index, WIDE]
25
26 unscaled = make_pipeline(SimpleImputer(strategy='median'),
27                          KNeighborsClassifier(25)).fit(Xn_tr, y_tr)
28 scaled = make_pipeline(SimpleImputer(strategy='median'), StandardScaler(),
29                        KNeighborsClassifier(25)).fit(Xn_tr, y_tr)
30
31 print('unscaled AUC %.4f' % roc_auc_score(y_te, unscaled.predict_proba(Xn_te)[: ,
32                                     ↪ 1]))
33 print('scaled AUC %.4f' % roc_auc_score(y_te, scaled.predict_proba(Xn_te)[: , 1]))
34 print('\ncolumn ranges:')
35 print((Xn_tr.max() - Xn_tr.min()).round(1).to_string())

```

```

unscaled AUC 0.6316
scaled AUC 0.7411

column ranges:
tenure_months      71.0
total_charges      7591.6

```

Eleven points of AUC, from one line. `total_charges` spans 7,592 while `tenure` spans 71, so unscaled the distance calculation is effectively *only* about total charges, tenure contributes about one percent of the typical gap between two rows and is ignored. Scaling gives each column an equal say. No other algorithm in this course is punished this hard for skipping it.

The curse of dimensionality

```

1 import numpy as np

```

```

2
3 rng = np.random.default_rng(0)
4 print('%6s %14s %14s %10s' % ('dims', 'nearest', 'furthest', 'ratio'))
5 for d in [2, 5, 20, 100, 500]:
6     pts = rng.uniform(size=(1000, d))
7     query = rng.uniform(size=d)
8     dist = np.linalg.norm(pts - query, axis=1)
9     print('%6d %14.4f %14.4f %10.3f'
10           % (d, dist.min(), dist.max(), dist.max() / dist.min()))

```

dims	nearest	furthest	ratio
2	0.0098	1.3341	135.618
5	0.1443	1.6328	11.317
20	1.0474	2.2671	2.164
100	3.0126	4.8993	1.626
500	8.1975	9.7549	1.190

⚠ Watch Out

In high dimensions, everything is equally far away

At two dimensions the furthest point is several times further than the nearest. At 500, the ratio approaches 1: every point is roughly the same distance from every other. “Nearest neighbour” stops meaning anything, and kNN degrades toward guessing. This is why kNN is a poor choice after one-hot encoding produces a wide sparse matrix, and why week 10’s dimensionality reduction exists.

Distance weighting

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                          ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                          ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split

```

```

24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                         stratify=y, random_state=42)
28 from sklearn.neighbors import KNeighborsClassifier
29 from sklearn.metrics import roc_auc_score
30
31 for weights in ['uniform', 'distance']:
32     m = Pipeline([('prep', preprocessor()),
33                  ('clf', KNeighborsClassifier(25, weights=weights))]).fit(X_tr,
34                               ↪ y_tr)
35     print('%-9s AUC %.4f'
36           % (weights, roc_auc_score(y_te, m.predict_proba(X_te)[: , 1])))

```

```

uniform    AUC 0.8023
distance   AUC 0.8009

```

`weights='distance'` gives closer neighbours more say, which usually helps when k is large.

	kNN	Logistic regression
Training cost	None	One optimisation
Prediction cost	Scans the training set	One dot product
Memory	Stores all training data	Stores $k+1$ numbers
Boundary	Arbitrarily shaped	A straight hyperplane
Interpretable	No	Yes, as odds ratios
Needs scaling	Absolutely	For regularisation to be fair

Day 3 takeaway

Day 4 Naive Bayes

A false assumption that ranks well and calibrates badly

Naive Bayes applies the theorem from week 3 directly, under an assumption that is almost always false. It works anyway, and understanding why is worth more than the algorithm itself.

The assumption

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11

```

```

12 # Are contract and tenure independent given the class? Week 3 said no.
13 sub = df[df['churned'] == 1]
14 a = (sub['contract'] == 'Month-To-Month')
15 b = (sub['tenure_months'] > 24)
16 print('among churners:')
17 print(' P(month-to-month)           %.4f' % a.mean())
18 print(' P(tenure > 24)              %.4f' % b.mean())
19 print(' product (if independent)    %.4f' % (a.mean() * b.mean()))
20 print(' actual joint probability    %.4f' % (a & b).mean())

```

```

among churners:
  P(month-to-month)           0.8545
  P(tenure > 24)              0.1182
  product (if independent)    0.1010
  actual joint probability    0.0945

```

Not independent, not even close. Naive Bayes will assume they are. The surprise is that this often costs less than it should.

Fit it anyway

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.naive_bayes import GaussianNB
29 from sklearn.linear_model import LogisticRegression
30 from sklearn.metrics import roc_auc_score, accuracy_score, log_loss
31
32 for name, clf in [('naive bayes', GaussianNB()),

```

```

33         ('logistic ', LogisticRegression(max_iter=1000,
34                                           random_state=42))]:
35     m = Pipeline([('prep', preprocessor()), ('clf', clf)].fit(X_tr, y_tr)
36     proba = m.predict_proba(X_te)[: , 1]
37     print('%s acc %.4f   AUC %.4f   log loss %.4f'
38           % (name, accuracy_score(y_te, m.predict(X_te)),
39             roc_auc_score(y_te, proba), log_loss(y_te, proba)))

```

```

naive bayes acc 0.7173   AUC 0.8044   log loss 1.1953
logistic    acc 0.7800   AUC 0.8174   log loss 0.4451

```

Good ranking, bad probabilities

Look at the gap between AUC and log loss. Naive Bayes ranks customers by risk reasonably well, but its actual probability estimates are poor, because multiplying correlated probabilities as though they were independent pushes the result toward 0 or 1. If you only need an ordering it is fine. If you need a number a business can act on, “this customer has a 30 percent chance”. It is not.

```

1  import numpy as np
2  import pandas as pd
3
4  df = pd.read_csv('customers.csv').drop_duplicates().copy()
5  df['contract'] = df['contract'].str.strip().str.title()
6  df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8  NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9  CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                           ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                           ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.naive_bayes import GaussianNB
29 import numpy as np
30
31 m = Pipeline([('prep', preprocessor()), ('clf', GaussianNB())].fit(X_tr, y_tr)
32 proba = m.predict_proba(X_te)[: , 1]
33 print('share of predictions below 0.01: %.3f' % (proba < 0.01).mean())
34 print('share of predictions above 0.99: %.3f' % (proba > 0.99).mean())

```

```

35 print('share between 0.2 and 0.8:      %.3f'
36       % ((proba > 0.2) & (proba < 0.8)).mean())

```

```

share of predictions below 0.01: 0.320
share of predictions above 0.99: 0.068
share between 0.2 and 0.8:      0.131

```

Most predictions are pinned at the extremes. A well-calibrated model on a 27 percent problem should produce a spread of intermediate probabilities. Week 7 shows how to measure and repair calibration.

Which variant for which data

Variant	Features	Typical use
GaussianNB	Continuous	Numeric tabular data
MultinomialNB	Counts	Word counts in text classification
BernoulliNB	Binary	Presence or absence flags
CategoricalNB	Categorical codes	Survey responses

```

1 from sklearn.feature_extraction.text import CountVectorizer
2 from sklearn.naive_bayes import MultinomialNB
3 from sklearn.pipeline import make_pipeline
4
5 notes = ['billing error on my invoice', 'charged twice this month',
6         'internet connection keeps dropping', 'speed is very slow',
7         'wrong amount on my bill', 'internet drops every evening']
8 labels = ['billing', 'billing', 'technical', 'technical',
9          'billing', 'technical']
10
11 m = make_pipeline(CountVectorizer(), MultinomialNB()).fit(notes, labels)
12 for text in ['my bill is wrong again', 'internet connection drops',
13             'charged the wrong amount']:
14     print('%-32s → %s' % (text, m.predict([text])[0]))

```

```

my bill is wrong again      → billing
internet connection drops   → technical
charged the wrong amount    → billing

```

This is what naive Bayes is genuinely good at. Text has thousands of sparse features, the independence assumption is less damaging when each word contributes a little, and training takes milliseconds. It remains a sensible baseline for text classification.

Day 4 takeaway

Day 5 Support Vector Machines

Margins, kernels, and the two parameters that must be tuned together

Logistic regression draws a boundary that separates the classes. A support vector machine

asks a better question: of all the boundaries that separate them, which one leaves the most room on either side?

The margin

```

1 import numpy as np
2 from sklearn.svm import SVC
3 from sklearn.datasets import make_blobs
4
5 X, y = make_blobs(n_samples=200, centers=2, cluster_std=1.2, random_state=7)
6 m = SVC(kernel='linear', C=1.0).fit(X, y)
7
8 print('training points %d' % len(X))
9 print('support vectors %d' % len(m.support_))
10 print('so the boundary depends on %.1f%% of the data'
11       % (100 * len(m.support_) / len(X)))

```

```

training points 200
support vectors 3
so the boundary depends on 1.5% of the data

```

The kernel trick

Some classes cannot be separated by a straight line at all. A kernel maps the data into a higher-dimensional space where they can be, without ever computing the coordinates in that space.

```

1 import numpy as np
2 from sklearn.svm import SVC
3 from sklearn.datasets import make_circles
4 from sklearn.model_selection import cross_val_score
5
6 X, y = make_circles(n_samples=400, factor=0.4, noise=0.12, random_state=0)
7 print('one class forms a ring around the other -- no line can split it')
8
9 for kernel in ['linear', 'poly', 'rbf']:
10     acc = cross_val_score(SVC(kernel=kernel, random_state=0), X, y, cv=5).mean()
11     print(' %-8s accuracy %.4f' % (kernel, acc))

```

```

one class forms a ring around the other -- no line can split it
linear    accuracy 0.5450
poly      accuracy 0.6250
rbf       accuracy 0.9975

```

The linear kernel manages barely better than guessing, because no straight line can separate a ring from its centre. The RBF kernel handles it almost perfectly.

C and gamma

Parameter	Low value	High value
C	Wide margin, tolerates mistakes (more bias)	Narrow margin, fits training data hard (more variance)
gamma	Each point influences a wide area (smoother)	Influence is local, boundary wraps individual points (overfits)

```
1 import numpy as np
2 from sklearn.svm import SVC
3 from sklearn.datasets import make_circles
4 from sklearn.model_selection import train_test_split
5 from sklearn.metrics import accuracy_score
6
7 X, y = make_circles(n_samples=400, factor=0.4, noise=0.12, random_state=0)
8 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.3, random_state=0)
9
10 print('%8s %8s %10s %10s' % ('C', 'gamma', 'train', 'test'))
11 for C_val, g in [(1, 0.1), (1, 1), (1, 100), (1000, 100)]:
12     m = SVC(C=C_val, gamma=g, random_state=0).fit(X_tr, y_tr)
13     print('%8s %8s %10.4f %10.4f'
14           % (C_val, g, accuracy_score(y_tr, m.predict(X_tr)),
15             accuracy_score(y_te, m.predict(X_te))))
```

C	gamma	train	test
1	0.1	0.9750	0.9583
1	1	1.0000	1.0000
1	100	1.0000	0.9917
1000	100	1.0000	0.9917

At gamma 100 the model scores near-perfectly on training data and worse on test: each support vector now influences only its immediate neighbourhood, so the boundary wraps individual points. Classic overfitting, and visible only because we looked at both columns.

On the churn problem

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
```

```

17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                         stratify=y, random_state=42)
28 from sklearn.svm import SVC
29 from sklearn.linear_model import LogisticRegression
30 from sklearn.metrics import roc_auc_score
31 import time
32
33 for name, clf in [('svm rbf ', SVC(probability=True, random_state=42)),
34                  ('logistic ', LogisticRegression(max_iter=1000,
35                                                    random_state=42))]:
36     t = time.perf_counter()
37     m = Pipeline([('prep', preprocessor()), ('clf', clf)]).fit(X_tr, y_tr)
38     elapsed = time.perf_counter() - t
39     auc = roc_auc_score(y_te, m.predict_proba(X_te)[:, 1])
40     print('%s AUC %.4f   fitted in %.2fs' % (name, auc, elapsed))

```

```

svm rbf      AUC 0.7804   fitted in 0.50s
logistic     AUC 0.8174   fitted in 0.02s

```

⚠ Watch Out

SVMs scale badly

Training cost grows between quadratically and cubically with the number of rows. At three thousand rows that is a second; at a million it is impractical. `probability=True` makes it worse still, because it fits an extra calibration model by internal cross-validation. For large tabular datasets, reach for the gradient boosting of week 6 instead.

Day 5 takeaway

Day 6 More Than Two Classes

One-vs-rest, confusion matrices, averaging traps and softmax

Everything so far has had two classes. Most real problems have more: which of nine defect types, which of five support queues, which of three risk bands.

Two strategies

Strategy	Models fitted	Each model learns	Cost
One-vs-rest	One per class	This class against all others	n models
One-vs-one	One per pair	This class against that one	$n(n-1)/2$ models

```
1 from sklearn.datasets import load_wine
2 from sklearn.linear_model import LogisticRegression
3 from sklearn.multiclass import OneVsRestClassifier, OneVsOneClassifier
4 from sklearn.preprocessing import StandardScaler
5 from sklearn.pipeline import make_pipeline
6 from sklearn.model_selection import cross_val_score
7
8 X, y = load_wine(return_X_y=True)
9 print('classes %d, rows %d, features %d' % (len(set(y)), *X.shape))
10
11 base = LogisticRegression(max_iter=5000, random_state=0)
12 for name, clf in [('built-in', base),
13                  ('one-vs-rest', OneVsRestClassifier(base)),
14                  ('one-vs-one', OneVsOneClassifier(base))]:
15     acc = cross_val_score(make_pipeline(StandardScaler(), clf), X, y, cv=5).mean()
16     print('%s accuracy %.4f' % (name, acc))
```

```
classes 3, rows 178, features 13
built-in      accuracy 0.9832
one-vs-rest   accuracy 0.9889
one-vs-one    accuracy 0.9833
```

scikit-learn handles multi-class automatically for most estimators, so you rarely wrap anything by hand. Knowing the strategies matters when you need to explain why a model produced an inconsistent set of scores.

The confusion matrix earns its name here

```
1 from sklearn.datasets import load_wine
2 from sklearn.linear_model import LogisticRegression
3 from sklearn.preprocessing import StandardScaler
4 from sklearn.pipeline import make_pipeline
5 from sklearn.model_selection import train_test_split
6 from sklearn.metrics import confusion_matrix, classification_report
7
8 X, y = load_wine(return_X_y=True)
9 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.35,
10                                          stratify=y, random_state=0)
11 m = make_pipeline(StandardScaler(),
12                  LogisticRegression(max_iter=5000, random_state=0)).fit(X_tr, y_tr)
13
14 print('rows = truth, columns = prediction')
15 print(confusion_matrix(y_te, m.predict(X_te)))
16 print()
17 print(classification_report(y_te, m.predict(X_te), digits=3))
```

```
rows = truth, columns = prediction
[[21  0  0]
 [ 1 24  0]
 [ 0  0 17]]
```

		precision	recall	f1-score	support
	0	0.955	1.000	0.977	21
	1	1.000	0.960	0.980	25
	2	1.000	1.000	1.000	17
	accuracy			0.984	63
	macro avg	0.985	0.987	0.985	63
	weighted avg	0.985	0.984	0.984	63

The off-diagonal cells tell you which classes the model confuses, and that is usually more actionable than the overall accuracy. Two classes that get mixed up may need a feature that distinguishes them specifically.

Averaging, and why the choice matters

```
1 import numpy as np
2 from sklearn.metrics import f1_score, precision_score, recall_score
3
4 # 100 of class 0, 100 of class 1, 10 of class 2 (the rare, important one)
5 truth = np.array([0] * 100 + [1] * 100 + [2] * 10)
6 pred = truth.copy()
7 pred[200:] = 0 # every single class-2 row misclassified
8
9 for avg in ['micro', 'macro', 'weighted']:
10     print('%-9s F1 %.4f' % (avg, f1_score(truth, pred, average=avg)))
```

```
micro      F1 0.9524
macro      F1 0.6508
weighted   F1 0.9297
```

Watch Out

micro and weighted hide a total failure on a rare class

The model gets every one of the rare class wrong, and micro-averaged F1 still reads 0.95 because the rare class is 5 percent of the rows. Macro-averaging treats every class equally regardless of size and drops to 0.65, which is the honest figure. When the rare class is the one you care about, fraud, defects, disease, use **macro**, or report per-class figures.

Probabilities across several classes

```
1 import numpy as np
2
3 def softmax(z):
4     z = z - z.max() # subtract the max for numerical stability
5     e = np.exp(z)
```

```

6     return e / e.sum()
7
8 scores = np.array([2.0, 1.0, 0.1])
9 probs = softmax(scores)
10 print('raw scores ', scores)
11 print('probabilities', probs.round(4))
12 print('they sum to  %.4f' % probs.sum())
13
14 big = np.array([1000.0, 999.0, 998.0])
15 print('\nwithout the max subtraction exp(1000) overflows;')
16 print('with it:', softmax(big).round(4))

```

```

raw scores  [2.  1.  0.1]
probabilities [0.659  0.2424 0.0986]
they sum to  1.0000

without the max subtraction exp(1000) overflows;
with it: [0.6652 0.2447 0.09  ]

```

Softmax is the multi-class sigmoid, and it appears again as the output layer of every classification network in week 12. Subtracting the maximum changes nothing mathematically and prevents overflow, which is why every real implementation does it.

Day 6 takeaway

Day 7 Comparing Classifiers Honestly

Five algorithms, cross-validation, and whether the differences are real

Five algorithms, one problem, one honest comparison, including how long each takes and whether the differences mean anything.

Compare by cross-validation

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([

```

```

18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                         stratify=y, random_state=42)
28 from sklearn.dummy import DummyClassifier
29 from sklearn.linear_model import LogisticRegression
30 from sklearn.neighbors import KNeighborsClassifier
31 from sklearn.naive_bayes import GaussianNB
32 from sklearn.svm import SVC
33 from sklearn.tree import DecisionTreeClassifier
34 from sklearn.model_selection import cross_validate, StratifiedKFold
35 import numpy as np
36
37 cv = StratifiedKFold(n_splits=5, shuffle=True, random_state=42)
38 candidates = {
39     'baseline': DummyClassifier(strategy='most_frequent'),
40     'logistic': LogisticRegression(max_iter=1000, random_state=42),
41     'kNN k=25': KNeighborsClassifier(25),
42     'naive bayes': GaussianNB(),
43     'svm rbf': SVC(probability=True, random_state=42),
44     'tree d=5': DecisionTreeClassifier(max_depth=5, random_state=42),
45 }
46
47 print('%-12s %14s %16s %10s' % ('model', 'ROC AUC', 'avg precision', 'fit secs'))
48 for name, clf in candidates.items():
49     pipe = Pipeline([('prep', preprocessor()), ('clf', clf)])
50     r = cross_validate(pipe, X_tr, y_tr, cv=cv,
51                       scoring=['roc_auc', 'average_precision'])
52     print('%-12s %.4f +/-%.3f %10.4f +/-%.3f %10.2f'
53           % (name, r['test_roc_auc'].mean(), r['test_roc_auc'].std(),
54              r['test_average_precision'].mean(),
55              r['test_average_precision'].std(), r['fit_time'].mean()))

```

model	ROC AUC	avg precision	fit secs
baseline	0.5000 +/-0.000	0.2680 +/-0.001	0.01
logistic	0.8208 +/-0.014	0.5974 +/-0.029	0.03
kNN k=25	0.8001 +/-0.010	0.5659 +/-0.022	0.01
naive bayes	0.8098 +/-0.006	0.5774 +/-0.027	0.01
svm rbf	0.7944 +/-0.020	0.6017 +/-0.034	0.34
tree d=5	0.7886 +/-0.014	0.5161 +/-0.025	0.01

Read the standard deviations before declaring a winner

If two models differ by less than the spread across folds, you have not shown one is better. Week 3's bootstrap made the same point about a single test set. The honest summary is usually “these three are equivalent, and one of them is ten times faster”.

Is the difference real?

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 from sklearn.svm import SVC
30 from sklearn.model_selection import cross_val_score, StratifiedKFold
31 from scipy import stats
32 import numpy as np
33
34 cv = StratifiedKFold(n_splits=10, shuffle=True, random_state=42)
35 a = cross_val_score(Pipeline([('prep', preprocessor()),
36                               ('clf', LogisticRegression(max_iter=1000,
37                                                           random_state=42))]),
38                     X_tr, y_tr, cv=cv, scoring='roc_auc')
39 b = cross_val_score(Pipeline([('prep', preprocessor()),
40                               ('clf', SVC(probability=True, random_state=42))]),
41                     X_tr, y_tr, cv=cv, scoring='roc_auc')
42
43 print('logistic %.4f +/- %.4f' % (a.mean(), a.std()))
44 print('svm       %.4f +/- %.4f' % (b.mean(), b.std()))
45 t, pval = stats.ttest_rel(a, b)
46 print('\npaired t-test on the same folds: p = %.4f' % pval)
```

```
logistic 0.8228 +/- 0.0175
svm       0.7926 +/- 0.0247
```

```
paired t-test on the same folds: p = 0.0002
```

A paired test compares the two models fold by fold, which removes the variation caused by some folds simply being harder. It is the right test here precisely because both models saw identical splits.

Choosing

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 from sklearn.metrics import (roc_auc_score, average_precision_score,
30                              confusion_matrix, classification_report)
31 import numpy as np
32
33 final = Pipeline([('prep', preprocessor()),
34                  ('clf', LogisticRegression(max_iter=1000,
35                                             random_state=42))]).fit(X_tr, y_tr)
36 proba = final.predict_proba(X_te)[:, 1]
37
38 # The threshold chosen on day 2, applied once to held-out data.
39 OFFER_COST, VALUE, SAVE_RATE = 30, 250, 0.35
40 best_t = max(np.arange(0.05, 0.95, 0.01),
41              key=lambda t: (lambda c: c[3] * SAVE_RATE * VALUE
42                             - (c[3] + c[1]) * OFFER_COST)(
43                  confusion_matrix(y_te, (proba >= t).astype(int)).ravel()))
44
45 print('TEST SET')
46 print(' ROC AUC          %.4f' % roc_auc_score(y_te, proba))
47 print(' average precision %.4f' % average_precision_score(y_te, proba))

```

```

48 print('  chosen threshold  %.2f' % best_t)
49 print()
50 print(classification_report(y_te, (proba >= best_t).astype(int),
51                             target_names=['stayed', 'churned'], digits=3))

```

```

TEST SET
ROC AUC          0.8174
average precision 0.5993
chosen threshold  0.34

              precision    recall  f1-score   support

   stayed       0.878        0.783     0.828       549
   churned       0.542        0.701     0.612       201

 accuracy                   0.761       750
 macro avg       0.710        0.742     0.720       750
 weighted avg     0.788        0.761     0.770       750

```

Recall on the churned class is now far above what the default threshold gave in week 1, at the cost of precision, which is exactly the trade the business case asked for.

Your assignment

Take the two best models from the comparison and average their predicted probabilities. Score the average. You will often find it beats both. That is ensembling, and week 6 explains why it works. Then check whether the improvement is larger than the fold-to-fold spread before you believe it.

Day 7 takeaway

Self-Assessment

Try It Yourself

- Logistic regression's coefficients are additive on the scale of:
 - Probability
 - Log-odds
 - Percentage points
 - Counts
- Why is a decision threshold of 0.5 usually wrong?
 - It favours the majority class
 - Probabilities are never exactly 0.5
 - It implicitly assumes a false positive and a false negative cost the same
 - The model is not calibrated
- k-nearest neighbours degrades in high dimensions because:
 - It cannot handle categories

- b) k must grow with dimension
 - c) It runs out of memory
 - d) Distances between all points become similar, so "nearest" stops meaning anything
4. Naive Bayes assumes features are independent given the class. In practice this makes it:
- a) Often a decent ranker but a poor estimator of probability
 - b) Perfectly calibrated
 - c) Slow
 - d) Useless
5. For an SVM with an RBF kernel, **C** and **gamma**:
- a) Can be tuned one at a time
 - b) Interact, so they must be tuned together
 - c) Have no effect on accuracy
 - d) Are set automatically
6. Macro-averaged F1 differs from micro-averaged F1 in that macro:
- a) Only applies to binary problems
 - b) Weights each class by its size
 - c) Gives every class equal weight, so a small class can drag the score down
 - d) Is always higher
7. Which model family is most affected by failing to scale the features?
- a) Random forest
 - b) Gradient boosting
 - c) Decision tree
 - d) k-nearest neighbours
8. Five classifiers score within 0.01 of each other in cross-validation. The correct conclusion is:
- a) The differences may well be noise; check the spread across folds before declaring a winner
 - b) More folds are needed
 - c) The data is broken
 - d) The best one wins

Answers: 1b, 2c, 3d, 4a, 5b, 6c, 7d, 8a. Anything you got wrong points at the day to re-read, not at a gap in ability.

Trees, Forests and Boosting

Week 6

Decision trees, bagging, random forests and gradient boosting, and an honest look at when they lose to a straight line.

OBJECTIVES

Tree ensembles win most tabular competitions, and this week builds them up from a single greedy split to boosted forests. It also ends by settling the question week 1 left open: why the random forest scored worse than logistic regression on our churn data, and what that tells you about choosing an algorithm at all.

Day 1 How a Decision Tree Decides

Impurity, greedy splitting, and the only model you can read end to end

Every model so far draws a single boundary through the whole space. A decision tree does something different: it asks a question, splits the data, and asks a different question of each half.

How a split gets chosen

The tree considers every feature and every possible cut point, and keeps the one that most reduces impurity.

```

1 import numpy as np
2
3 def gini(labels):
4     labels = np.asarray(labels)
5     if len(labels) == 0:
6         return 0.0
7     p1 = labels.mean()
8     return 1 - (p1 ** 2 + (1 - p1) ** 2)
9
10 print('all one class      %.4f' % gini([1, 1, 1, 1]))
11 print('three to one      %.4f' % gini([1, 1, 1, 0]))
12 print('even mix          %.4f' % gini([1, 1, 0, 0]))

```

```

all one class      0.0000
three to one       0.3750
even mix           0.5000

```

```

1 import numpy as np
2 import pandas as pd

```

```

3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 import numpy as np
12
13 def gini(labels):
14     labels = np.asarray(labels)
15     if len(labels) == 0:
16         return 0.0
17     p1 = labels.mean()
18     return 1 - (p1 ** 2 + (1 - p1) ** 2)
19
20 y = df[TARGET].to_numpy()
21 x = df['tenure_months'].to_numpy()
22 parent = gini(y)
23
24 print('impurity before any split: %.4f' % parent)
25 print('\n%12s %10s %10s %12s' % ('split at', 'left n', 'right n', 'gain'))
26 best = None
27 for cut in [6, 12, 18, 24, 36, 48]:
28     left, right = y[x <= cut], y[x > cut]
29     weighted = (len(left) * gini(left) + len(right) * gini(right)) / len(y)
30     gain = parent - weighted
31     if best is None or gain > best[1]:
32         best = (cut, gain)
33     print('%12d %10d %10d %12.5f' % (cut, len(left), len(right), gain))
34
35 print('\nbest of these: tenure <= %d, gain %.5f' % best)

```

```
impurity before any split: 0.3924
```

split at	left n	right n	gain
6	345	2655	0.00723
12	949	2051	0.02653
18	1545	1455	0.03244
24	1918	1082	0.03664
36	2420	580	0.02537
48	2686	314	0.01562

```
best of these: tenure <= 24, gain 0.03664
```

That is the entire algorithm, applied recursively. Split, then repeat on each side, until a stopping rule fires.

Let scikit-learn do it and read the result

```

1 import numpy as np
2 import pandas as pd
3

```

```

4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)
30 from sklearn.tree import DecisionTreeClassifier, export_text
31
32 tree = Pipeline([('prep', preprocessor(scale=False)),
33                  ('clf', DecisionTreeClassifier(max_depth=3,
34                                                  random_state=42))]).fit(X_tr, y_tr)
35 names = list(tree.named_steps['prep'].get_feature_names_out())
36 print(export_text(tree.named_steps['clf'], feature_names=names))

```

```

|--- cat__contract_Month-To-Month <= 0.50
|   |--- num__tenure_months <= 17.50
|   |   |--- cat__contract_One Year <= 0.50
|   |   |   |--- class: 0
|   |   |--- cat__contract_One Year > 0.50
|   |   |   |--- class: 0
|   |--- num__tenure_months > 17.50
|   |   |--- cat__contract_Two Year <= 0.50
|   |   |   |--- class: 0
|   |   |--- cat__contract_Two Year > 0.50
|   |   |   |--- class: 0
|--- cat__contract_Month-To-Month > 0.50
|   |--- num__monthly_charges <= 51.02
|   |   |--- num__tenure_months <= 9.50
|   |   |   |--- class: 0
|   |   |--- num__tenure_months > 9.50
|   |   |   |--- class: 0
|   |--- num__monthly_charges > 51.02
|   |   |--- num__tenure_months <= 32.50
|   |   |   |--- class: 1
|   |   |--- num__tenure_months > 32.50

```

```
| | | |--- class: 0
```

This is why people like trees

You can read the whole model. Every prediction is a path from the root to a leaf, and you can print that path for any customer and hand it to someone who has never heard of machine learning. No other model in this course offers that.

Trees do not need scaling, and barely care about outliers

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)
30 from sklearn.tree import DecisionTreeClassifier
31 from sklearn.metrics import roc_auc_score
32
33 for scale in [True, False]:
34     m = Pipeline([('prep', preprocessor(scale=scale)),
35                  ('clf', DecisionTreeClassifier(max_depth=5,
36                                                  random_state=42))]).fit(X_tr, y_tr)
37     print('scaled=%-6s AUC %.6f'
38           % (scale, roc_auc_score(y_te, m.predict_proba(X_te)[: , 1])))
```

```
scaled=True    AUC 0.774524
scaled=False   AUC 0.774524
```

Identical to six decimal places. A split asks “is this value above the threshold”, and scaling preserves order, so it changes nothing. The same reasoning is why the 97-call outlier in week 2 moved a tree not at all.

Gini or entropy

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)
30 from sklearn.tree import DecisionTreeClassifier
31 from sklearn.metrics import roc_auc_score
32
33 for criterion in ['gini', 'entropy', 'log_loss']:
34     m = Pipeline([('prep', preprocessor(scale=False)),
35                  ('clf', DecisionTreeClassifier(criterion=criterion,
36                                                  max_depth=5,
37                                                  random_state=42))]).fit(X_tr, y_tr)
38     print('%-9s AUC %.4f'
39           % (criterion, roc_auc_score(y_te, m.predict_proba(X_te)[: , 1])))
```

```
gini      AUC 0.7745
entropy   AUC 0.7684
log_loss   AUC 0.7684
```

Practically indistinguishable, which is the usual result. Do not spend time tuning the criterion; spend it on depth, which is tomorrow.

Day 1 takeaway**Day 2 Depth, Pruning and Instability**

Stopping a tree memorising, and the weakness that becomes tomorrow's strength

Left alone, a tree will keep splitting until every leaf is pure. That means it memorises the training set perfectly and generalises terribly.

Watch it happen

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)
30 from sklearn.tree import DecisionTreeClassifier
31 from sklearn.metrics import roc_auc_score
32
33 print('%8s %10s %10s %8s %8s' % ('depth', 'train', 'test', 'gap', 'leaves'))
34 for depth in [1, 2, 3, 5, 8, 12, None]:
35     m = Pipeline([('prep', preprocessor(scale=False)),
36                  ('clf', DecisionTreeClassifier(max_depth=depth,
37                                                  random_state=42))]).fit(X_tr, y_tr)
38     tr = roc_auc_score(y_tr, m.predict_proba(X_tr)[:, 1])
39     te = roc_auc_score(y_te, m.predict_proba(X_te)[:, 1])
40     print('%8s %10.4f %10.4f %8.4f %8d'

```

```
41 | % (depth, tr, te, tr - te, m.named_steps['clf'].get_n_leaves()))
```

depth	train	test	gap	leaves
1	0.7057	0.7056	0.0001	2
2	0.7660	0.7733	-0.0072	4
3	0.7980	0.7774	0.0206	8
5	0.8409	0.7745	0.0664	29
8	0.9047	0.7309	0.1738	116
12	0.9779	0.6763	0.3016	311
None	1.0000	0.6291	0.3709	539

Unlimited depth gives a perfect training score and the worst test score of the lot. Hundreds of leaves, most holding a single customer. The tree has recorded the training set rather than learned from it.

The four ways to stop it

Parameter	Controls	Sensible starting point
max_depth	How many questions deep	3 to 8 for a single tree
min_samples_leaf	Smallest allowed leaf	1 to 5 percent of rows
min_samples_split	Smallest node worth splitting	Twice the leaf minimum
max_leaf_nodes	Total leaves, grown best-first	An alternative to depth

```
1 | import numpy as np
2 | import pandas as pd
3 |
4 | df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 | df['contract'] = df['contract'].str.strip().str.title()
6 | df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7 |
8 | NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 | CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 | TARGET = 'churned'
11 | from sklearn.compose import ColumnTransformer
12 | from sklearn.pipeline import Pipeline
13 | from sklearn.impute import SimpleImputer
14 | from sklearn.preprocessing import StandardScaler, OneHotEncoder
15 |
16 | def preprocessor(scale=True):
17 |     num_steps = [('i', SimpleImputer(strategy='median'))]
18 |     if scale:
19 |         num_steps.append(('s', StandardScaler()))
20 |     return ColumnTransformer([
21 |         ('num', Pipeline(num_steps), NUM),
22 |         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23 |                             ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24 |     ])
25 | from sklearn.model_selection import train_test_split
26 |
27 | X, y = df[NUM + CAT], df[TARGET]
28 | X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29 |                                           stratify=y, random_state=42)
```

```

30 from sklearn.tree import DecisionTreeClassifier
31 from sklearn.metrics import roc_auc_score
32
33 print('%18s %10s %10s' % ('min_samples_leaf', 'test AUC', 'leaves'))
34 for leaf in [1, 5, 20, 50, 150]:
35     m = Pipeline([('prep', preprocessor(scale=False)),
36                  ('clf', DecisionTreeClassifier(min_samples_leaf=leaf,
37                                                random_state=42))]).fit(X_tr, y_tr)
38     print('%18d %10.4f %10d'
39           % (leaf, roc_auc_score(y_te, m.predict_proba(X_te)[: , 1]),
40             m.named_steps['clf'].get_n_leaves()))

```

min_samples_leaf	test AUC	leaves
1	0.6291	539
5	0.7260	232
20	0.7630	75
50	0.7915	31
150	0.7974	12

min_samples_leaf is the better knob

Depth limits every branch equally, whether or not that branch had enough data to justify going deeper. A minimum leaf size adapts: dense regions of the data get to split further, sparse ones stop early. If you only tune one parameter on a tree, tune this one.

Cost-complexity pruning

Rather than guessing a limit in advance, grow the tree fully and then cut back the branches that do not pay for themselves.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                          ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
23

```

```

24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                         stratify=y, random_state=42)
30 from sklearn.tree import DecisionTreeClassifier
31 from sklearn.metrics import roc_auc_score
32 import numpy as np
33
34 prep = preprocessor(scale=False).fit(X_tr)
35 Xt_tr, Xt_te = prep.transform(X_tr), prep.transform(X_te)
36
37 path = DecisionTreeClassifier(random_state=42).cost_complexity_pruning_path(
38     Xt_tr, y_tr)
39 alphas = path.ccp_alphas[:max(1, len(path.ccp_alphas) // 8)][:8]
40
41 print('%12s %10s %10s %8s' % ('ccp_alpha', 'train', 'test', 'leaves'))
42 for a in alphas:
43     t = DecisionTreeClassifier(random_state=42, ccp_alpha=a).fit(Xt_tr, y_tr)
44     print('%12.6f %10.4f %10.4f %8d'
45           % (a, roc_auc_score(y_tr, t.predict_proba(Xt_tr)[: , 1]),
46              roc_auc_score(y_te, t.predict_proba(Xt_te)[: , 1]),
47              t.get_n_leaves()))

```

ccp_alpha	train	test	leaves
0.000000	1.0000	0.6291	539
0.000111	1.0000	0.6308	537
0.000221	0.9999	0.6071	533
0.000254	0.9999	0.6093	530
0.000267	0.9999	0.6141	521
0.000267	0.9999	0.6141	521
0.000274	0.9998	0.6143	518
0.000279	0.9998	0.6199	515

As alpha rises the tree shrinks and the training score falls, while the test score improves up to a point. Pruning finds a smaller and better tree than growing an unrestricted one ever could.

A single tree is unstable

The measure that matters is not whether the printed tree looks different, but whether it *predicts* differently. Fit trees on overlapping samples and count how often two of them disagree about the same customer.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'

```

```

11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25
26 from sklearn.model_selection import train_test_split
27
28 X, y = df[NUM + CAT], df[TARGET]
29 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
30                                           stratify=y, random_state=42)
31
32 from sklearn.tree import DecisionTreeClassifier
33 from sklearn.model_selection import train_test_split
34 import numpy as np
35 from itertools import combinations
36
37 def disagreement(depth, n_models=6):
38     probas = []
39     for seed in range(n_models):
40         Xs, _, ys, _ = train_test_split(X_tr, y_tr, train_size=0.7,
41                                         random_state=seed)
42         m = Pipeline([('prep', preprocessor(scale=False)),
43                       ('clf', DecisionTreeClassifier(max_depth=depth,
44                                                       random_state=42))]).fit(Xs, ys)
45         probas.append(m.predict_proba(X_te)[:, 1])
46     label = np.mean([(a > 0.5) != (b > 0.5)].mean()
47                     for a, b in combinations(probas, 2))
48     prob = np.mean([np.abs(a - b).mean() for a, b in combinations(probas, 2)])
49     return label, prob
50
51 print('%10s %16s %16s' % ('max_depth', 'labels differ', 'mean |p1 - p2|'))
52 for depth in [2, 4, 8, 12, None]:
53     label, prob = disagreement(depth)
54     print('%10s %15.1f%% %16.4f' % (depth, 100 * label, prob))

```

max_depth	labels differ	mean p1 - p2
2	19.9%	0.0342
4	15.2%	0.0899
8	18.8%	0.1741
12	23.5%	0.2371
None	23.6%	0.2366

⚠ Watch Out

Two trees on 70 percent of the same data, disagreeing on a quarter of customers

Read the probability column first: it rises steadily with depth, from 0.03 to 0.24. The greedy algorithm commits at each step with no way to reconsider, so once two candidate splits are close, a handful of rows decides which wins and every branch below inherits that coin toss. The deeper the tree, the more of those coin tosses accumulate.

The label column is not monotone, and the exception is worth understanding. A depth-2 tree has four leaves, so its probabilities are four numbers, extremely stable. But if one of those numbers sits near 0.5, a small shift flips an entire leaf's worth of customers at once. Stable probabilities, unstable labels. It is a compact reminder that thresholding throws information away.

Either way, this variance is what tomorrow exploits. If a deep tree disagrees with itself depending on which rows it saw, build hundreds on different samples and average them: the disagreement cancels, the agreement survives.

Day 2 takeaway

Day 3 Bagging and Random Forests

Averaging away instability, and the validation set you get for free

Yesterday ended on instability. Today that becomes the mechanism: average many unstable models and the noise cancels while the signal survives.

Bagging

```
1 import numpy as np
2
3 rng = np.random.default_rng(0)
4 n = 1000
5 rows = np.arange(n)
6 sample = rng.choice(rows, size=n, replace=True)
7
8 unique = len(np.unique(sample))
9 print('bootstrap sample of %d drawn from %d rows' % (n, n))
10 print('distinct rows included: %d (%.1f%%)' % (unique, 100 * unique / n))
11 print('left out entirely: %d (%.1f%%)' % (n - unique, 100 * (n - unique) / n))
12 print('\ntheory says 1 - 1/e = %.3f are included' % (1 - np.exp(-1)))
```

```
bootstrap sample of 1000 drawn from 1000 rows
distinct rows included: 622 (62.2%)
left out entirely: 378 (37.8%)

theory says 1 - 1/e = 0.632 are included
```

About 63 percent of rows appear in any given bootstrap sample. The 37 percent left out are that tree's *out-of-bag* rows, and they give you a free validation set.

Random forests add a second source of randomness

Bagging alone leaves the trees correlated: if one feature is much stronger than the rest, every tree splits on it first. A random forest fixes that by offering each split only a random subset of features.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)
30 from sklearn.ensemble import BaggingClassifier, RandomForestClassifier
31 from sklearn.tree import DecisionTreeClassifier
32 from sklearn.metrics import roc_auc_score
33
34 models = {
35     'single tree': DecisionTreeClassifier(random_state=42),
36     'bagging x200': BaggingClassifier(DecisionTreeClassifier(random_state=42),
37                                     n_estimators=200, random_state=42),
38     'forest x200': RandomForestClassifier(n_estimators=200, random_state=42),
39 }
40 for name, clf in models.items():
41     m = Pipeline([('prep', preprocessor(scale=False)), ('clf', clf)]).fit(X_tr,
42     ↪ y_tr)
43     print('%s AUC %.4f'
44           % (name, roc_auc_score(y_te, m.predict_proba(X_te)[: , 1])))

```

```

single tree  AUC 0.6291
bagging x200 AUC 0.7699
forest x200  AUC 0.7788

```

The out-of-bag score

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)
30 from sklearn.ensemble import RandomForestClassifier
31 from sklearn.metrics import roc_auc_score, accuracy_score
32
33 m = Pipeline([('prep', preprocessor(scale=False)),
34              ('clf', RandomForestClassifier(n_estimators=300, oob_score=True,
35                                           random_state=42))]).fit(X_tr, y_tr)
36 forest = m.named_steps['clf']
37 print('out-of-bag accuracy %.4f' % forest.oob_score_)
38 print('test accuracy      %.4f' % accuracy_score(y_te, m.predict(X_te)))
```

```
out-of-bag accuracy 0.7560
test accuracy      0.7587
```

A validation score for free

Each tree is scored on the rows it never saw, and those scores are pooled. You get something close to cross-validation without fitting anything extra. Useful when a full cross-validation would be expensive, though it is slightly pessimistic because each row is judged by only about a third of the forest.

What to tune

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)
30 from sklearn.ensemble import RandomForestClassifier
31 from sklearn.metrics import roc_auc_score
32
33 print('%12s %10s' % ('n_estimators', 'test AUC'))
34 for n in [5, 25, 100, 300, 800]:
35     m = Pipeline([('prep', preprocessor(scale=False)),
36                  ('clf', RandomForestClassifier(n_estimators=n,
37                                                random_state=42))]).fit(X_tr, y_tr)
38     print('%12d %10.4f'
39           % (n, roc_auc_score(y_te, m.predict_proba(X_te)[: , 1])))
40
41 print('\n%18s %10s' % ('min_samples_leaf', 'test AUC'))
42 for leaf in [1, 5, 20, 50]:
43     m = Pipeline([('prep', preprocessor(scale=False)),
44                  ('clf', RandomForestClassifier(n_estimators=300,
45                                                min_samples_leaf=leaf,
46                                                random_state=42))]).fit(X_tr, y_tr)
47     print('%18d %10.4f'
48           % (leaf, roc_auc_score(y_te, m.predict_proba(X_te)[: , 1])))

```

n_estimators	test AUC
5	0.7205

25	0.7671
100	0.7728
300	0.7776
800	0.7779

min_samples_leaf	test AUC
1	0.7776
5	0.7989
20	0.8098
50	0.8158

⚠ Watch Out

More trees never overfit; deeper trees still can

Adding trees to a forest cannot make it worse, the score plateaus and stays there, so `n_estimators` is a budget decision, not a tuning decision. The parameters that genuinely matter are `min_samples_leaf` and `max_features`. This is the opposite of boosting, where more rounds absolutely can overfit.

Day 3 takeaway

Day 4 Boosting

Fitting each model to the last one's mistakes, implemented then industrialised

A forest builds hundreds of independent trees and averages them. Boosting builds them in sequence, each one working on what the previous ones got wrong.

The idea, implemented

```

1 import numpy as np
2 from sklearn.tree import DecisionTreeRegressor
3
4 rng = np.random.default_rng(0)
5 x = rng.uniform(0, 6, 300).reshape(-1, 1)
6 y = np.sin(1.5 * x.ravel()) + 0.3 * x.ravel() + rng.normal(0, 0.2, 300)
7
8 prediction = np.full(len(y), y.mean())    # start with the mean
9 lr = 0.3
10
11 print('%8s %12s' % ('round', 'train MSE'))
12 print('%8s %12.5f' % ('start', ((y - prediction) ** 2).mean()))
13 for r in range(1, 41):
14     residual = y - prediction              # what we still get wrong
15     stump = DecisionTreeRegressor(max_depth=2).fit(x, residual)
16     prediction = prediction + lr * stump.predict(x)
17     if r in (1, 5, 10, 20, 40):
18         print('%8d %12.5f' % (r, ((y - prediction) ** 2).mean()))

```

round	train MSE
start	0.86697
1	0.49189
5	0.10329
10	0.04882
20	0.03073
40	0.02490

Thirty lines, and it is the core of the algorithm that wins most tabular competitions.

The library versions

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)
30 from sklearn.ensemble import (AdaBoostClassifier, GradientBoostingClassifier,
31                               HistGradientBoostingClassifier,
32                               RandomForestClassifier)
33 from sklearn.metrics import roc_auc_score
34 import time
35
36 models = {
37     'random forest': RandomForestClassifier(n_estimators=300, random_state=42),
38     'adaboost': AdaBoostClassifier(n_estimators=200, random_state=42),
39     'grad boosting': GradientBoostingClassifier(random_state=42),
40     'hist grad': HistGradientBoostingClassifier(random_state=42),
41 }
42 for name, clf in models.items():

```

```

43     t = time.perf_counter()
44     m = Pipeline([('prep', preprocessor(scale=False)), ('clf', clf)]).fit(X_tr,
    ↪ y_tr)
45     el = time.perf_counter() - t
46     print('%s AUC %.4f    %.2fs'
47           % (name, roc_auc_score(y_te, m.predict_proba(X_te)[: , 1]), el))

```

```

random forest AUC 0.7776    2.00s
adaboost      AUC 0.8098    1.19s
grad boosting AUC 0.8034    0.74s
hist grad     AUC 0.7797    5.20s

```

Learning rate against number of rounds

```

1  import numpy as np
2  import pandas as pd
3
4  df = pd.read_csv('customers.csv').drop_duplicates().copy()
5  df['contract'] = df['contract'].str.strip().str.title()
6  df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8  NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9  CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)
30 from sklearn.ensemble import GradientBoostingClassifier
31 from sklearn.metrics import roc_auc_score
32
33 print('%8s %8s %10s %10s' % ('lr', 'rounds', 'train', 'test'))
34 for lr, n in [(1.0, 200), (0.3, 200), (0.1, 200), (0.05, 200), (0.05, 800)]:
35     m = Pipeline([('prep', preprocessor(scale=False)),
36                   ('clf', GradientBoostingClassifier(learning_rate=lr,
37                                                       n_estimators=n,
38                                                       random_state=42))]).fit(X_tr,

```

```

39     ↪ y_tr)
40     print('%8s %8d %10.4f %10.4f'
41           % (lr, n, roc_auc_score(y_tr, m.predict_proba(X_tr)[: , 1]),
              roc_auc_score(y_te, m.predict_proba(X_te)[: , 1])))

```

lr	rounds	train	test
1.0	200	0.9998	0.7410
0.3	200	0.9797	0.7774
0.1	200	0.9157	0.7968
0.05	200	0.8857	0.8082
0.05	800	0.9547	0.7843

⚠ Watch Out

Boosting can and will overfit

At a learning rate of 1.0 the training AUC runs away toward 1.0 while the test score falls. Unlike a forest, more rounds is not free. The two parameters trade against each other: halve the learning rate and you need roughly twice the rounds, which costs time but usually generalises better. A small rate with early stopping is the standard recipe.

Early stopping

```

1  import numpy as np
2  import pandas as pd
3
4  df = pd.read_csv('customers.csv').drop_duplicates().copy()
5  df['contract'] = df['contract'].str.strip().str.title()
6  df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8  NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9  CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)

```

```

30 from sklearn.ensemble import HistGradientBoostingClassifier
31 from sklearn.metrics import roc_auc_score
32
33 m = Pipeline([('prep', preprocessor(scale=False)),
34               ('clf', HistGradientBoostingClassifier(
35                   max_iter=1000, learning_rate=0.05,
36                   early_stopping=True, validation_fraction=0.15,
37                   n_iter_no_change=20, random_state=42))]).fit(X_tr, y_tr)
38 clf = m.named_steps['clf']
39 print('allowed up to 1000 rounds, stopped at %d' % clf.n_iter_)
40 print('test AUC %.4f' % roc_auc_score(y_te, m.predict_proba(X_te)[: , 1]))

```

```

allowed up to 1000 rounds, stopped at 80
test AUC 0.7924

```

It holds back 15 percent of the training data, watches the score on it, and stops when twenty rounds pass without improvement. You set a generous ceiling and let the model decide, which is more reliable than tuning the round count by hand.

	Random forest	Gradient boosting
Trees built	Independently, in parallel	Sequentially, each on the last errors
Individual trees	Deep, overfitting	Shallow, underfitting
More trees	Never hurts	Eventually overfits
Key parameters	<code>min_samples_leaf</code> , <code>max_features</code>	<code>learning_rate</code> , rounds, depth
Tuning effort	Low	Higher, and it pays
Typical winner	When you have no time to tune	When you do

Day 4 takeaway

Day 5 XGBoost and LightGBM

Regularised, histogram-based boosting that handles gaps natively

XGBoost and LightGBM are gradient boosting done properly, faster, regularised, and with native handling of the things that make real data awkward.

What they add

- **Regularisation on the trees themselves**, so leaf values cannot grow unchecked.
- **Histogram-based splitting**: bucket each feature into 256 bins instead of testing every value, which is dramatically faster.
- **Native missing-value handling**: each split learns which way a missing value should go.
- **Serious engineering**: multi-core, cache-aware, and able to work out of memory.

```

1 import numpy as np
2 import pandas as pd
3

```

```

4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                         ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                         stratify=y, random_state=42)
30 import xgboost as xgb
31 import lightgbm as lgb
32 from sklearn.ensemble import HistGradientBoostingClassifier
33 from sklearn.metrics import roc_auc_score
34 import time
35
36 models = {
37     'sklearn hist': HistGradientBoostingClassifier(random_state=42),
38     'xgboost': xgb.XGBClassifier(n_estimators=300, learning_rate=0.1,
39                                max_depth=4, random_state=42,
40                                eval_metric='logloss'),
41     'lightgbm': lgb.LGBMClassifier(n_estimators=300, learning_rate=0.1,
42                                   max_depth=4, random_state=42,
43                                   verbose=-1),
44 }
45 for name, clf in models.items():
46     t = time.perf_counter()
47     m = Pipeline([('prep', preprocessor(scale=False)), ('clf', clf)]).fit(X_tr,
48                               ↪ y_tr)
49     el = time.perf_counter() - t
50     print('%s AUC %.4f   %.2fs'
51           % (name, roc_auc_score(y_te, m.predict_proba(X_te)[: , 1]), el))

```

```

sklearn hist AUC 0.7797   4.45s
xgboost      AUC 0.7765   3.21s
lightgbm     AUC 0.7755   1.40s

```

Missing values need no imputation

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.model_selection import train_test_split
12
13 X, y = df[NUM + CAT], df[TARGET]
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 import xgboost as xgb
17 from sklearn.metrics import roc_auc_score
18
19 # No imputer at all -- the raw numeric columns, gaps included.
20 Xn_tr = X_tr[NUM]
21 Xn_te = X_te[NUM]
22 print('missing values passed straight in: %d' % Xn_tr.isna().sum().sum())
23
24 m = xgb.XGBClassifier(n_estimators=300, max_depth=4, learning_rate=0.1,
25                      random_state=42, eval_metric='logloss').fit(Xn_tr, y_tr)
26 print('AUC %.4f' % roc_auc_score(y_te, m.predict_proba(Xn_te)[: , 1]))
```

```
missing values passed straight in: 133
AUC 0.7237
```

Learned, not guessed

At each split the algorithm tries sending missing values left and sending them right, and keeps whichever reduces the loss more. That is strictly better than imputing a median, because it lets missingness carry its own meaning, the point week 2 made about informative gaps, handled automatically.

The parameters that matter

Parameter	Does	Try
n_estimators	Boosting rounds	300 to 2000 with early stopping
learning_rate	Step size per round	0.01 to 0.1
max_depth	Tree depth	3 to 8
subsample	Row fraction per tree	0.6 to 1.0
colsample_bytree	Feature fraction per tree	0.6 to 1.0
reg_lambda	L2 on leaf weights	1 to 10
min_child_weight	Minimum weight in a leaf	1 to 20

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)
30 import xgboost as xgb
31 from sklearn.metrics import roc_auc_score
32
33 configs = [
34     ('defaults', {}),
35     ('shallow, slow', dict(max_depth=3, learning_rate=0.03, n_estimators=600)),
36     ('subsampled', dict(max_depth=4, learning_rate=0.05, n_estimators=500,
37                          subsample=0.8, colsample_bytree=0.8)),
38     ('regularised', dict(max_depth=5, learning_rate=0.05, n_estimators=500,
39                          reg_lambda=10, min_child_weight=10)),
40 ]
41 for name, kw in configs:
42     clf = xgb.XGBClassifier(random_state=42, eval_metric='logloss', **kw)
43     m = Pipeline([('prep', preprocessor(scale=False)), ('clf', clf)]).fit(X_tr,
44     ↪ y_tr)
45     tr = roc_auc_score(y_tr, m.predict_proba(X_tr)[:, 1])
46     te = roc_auc_score(y_te, m.predict_proba(X_te)[:, 1])
47     print('%s train %.4f test %.4f gap %.4f' % (name, tr, te, tr - te))

```

defaults	train 0.9940	test 0.7574	gap 0.2366
shallow, slow	train 0.8913	test 0.8020	gap 0.0893
subsampled	train 0.9507	test 0.7804	gap 0.1703
regularised	train 0.9071	test 0.7905	gap 0.1166

Read the gap column. The default configuration overfits hardest; the shallow, slow and regularised

versions trade training score for a smaller gap. The best test score of the four, 0.8020, comes from the slowest and shallowest configuration, and tomorrow shows that even that does not settle the question.

Day 5 takeaway

Day 6 Feature Importance, and Why the Default Lies

Impurity importance ranking pure noise first, and the honest alternative

Every tree model offers a `feature_importances_` attribute. It is the most-quoted and least-trustworthy number in applied machine learning.

What the default importance measures

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)
30 from sklearn.ensemble import RandomForestClassifier
31 import pandas as pd
32
33 m = Pipeline([('prep', preprocessor(scale=False)),
34              ('clf', RandomForestClassifier(n_estimators=300,
35                                             random_state=42))]).fit(X_tr, y_tr)
```

```

36 names = m.named_steps['prep'].get_feature_names_out()
37 imp = pd.Series(m.named_steps['clf'].feature_importances_, index=names)
38 print(imp.sort_values(ascending=False).head(8).round(4).to_string())

```

```

num__monthly_charges      0.3216
num__tenure_months        0.3167
num__support_calls        0.1009
cat__contract_Month-To-Month 0.0701
cat__contract_Two Year    0.0359
cat__payment_method_Electronic check 0.0213
cat__contract_One Year    0.0185
cat__internet_service_Fibre optic 0.0180

```

The bias, demonstrated

```

1 import numpy as np
2 import pandas as pd
3 from sklearn.ensemble import RandomForestClassifier
4
5 rng = np.random.default_rng(0)
6 n = 2000
7 y = rng.integers(0, 2, n)
8
9 X = pd.DataFrame({
10     'real_signal': y + rng.normal(0, 1.2, n),    # genuinely predictive
11     'binary_noise': rng.integers(0, 2, n),      # pure noise, 2 values
12     'id_noise': rng.permutation(n),             # pure noise, 2000 values
13 })
14
15 m = RandomForestClassifier(n_estimators=300, random_state=42).fit(X, y)
16 print(pd.Series(m.feature_importances_, index=X.columns).round(4).to_string())

```

```

real_signal      0.5739
binary_noise     0.0035
id_noise         0.4225

```

Watch Out

It rewards high-cardinality columns for being noisy

`id_noise` is a random permutation. It contains no information whatsoever, and impurity importance ranks it comparably to the column that genuinely predicts the target, far above the binary noise column. A feature with thousands of distinct values offers thousands of candidate split points, so by chance some of them separate the training rows well. The importance measures opportunity to overfit, not usefulness.

Permutation importance is the honest one

```

1 import numpy as np
2 import pandas as pd

```

```

3 from sklearn.ensemble import RandomForestClassifier
4 from sklearn.inspection import permutation_importance
5 from sklearn.model_selection import train_test_split
6
7 rng = np.random.default_rng(0)
8 n = 2000
9 y = rng.integers(0, 2, n)
10 X = pd.DataFrame({
11     'real_signal': y + rng.normal(0, 1.2, n),
12     'binary_noise': rng.integers(0, 2, n),
13     'id_noise': rng.permutation(n),
14 })
15 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.3, random_state=0)
16 m = RandomForestClassifier(n_estimators=300, random_state=42).fit(X_tr, y_tr)
17
18 r = permutation_importance(m, X_te, y_te, n_repeats=20, random_state=0,
19                             scoring='roc_auc')
20 print(pd.DataFrame({'importance': r.importances_mean,
21                     'std': r.importances_std},
22                     index=X.columns).round(4).to_string())

```

	importance	std
real_signal	0.1876	0.0209
binary_noise	-0.0009	0.0128
id_noise	0.0107	0.0107

Now both noise columns sit within one standard deviation of zero, where they belong, while the real signal is an order of magnitude above them. Note that `id_noise` has not vanished entirely, the model did latch onto it slightly, and permutation importance correctly reports that as small rather than pretending it is absent. Same model, same data, a different and better question.

On the churn model

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([

```

```

21     ('num', Pipeline(num_steps), NUM),
22     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                      ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24 ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                          stratify=y, random_state=42)
30 from sklearn.ensemble import RandomForestClassifier
31 from sklearn.inspection import permutation_importance
32 import pandas as pd
33
34 m = Pipeline([('prep', preprocessor(scale=False)),
35              ('clf', RandomForestClassifier(n_estimators=300,
36                                          random_state=42))]).fit(X_tr, y_tr)
37 r = permutation_importance(m, X_te, y_te, n_repeats=15, random_state=0,
38                           scoring='roc_auc')
39 print(pd.Series(r.importances_mean, index=X_te.columns)
40       .sort_values(ascending=False).round(4).to_string())

```

```

contract          0.1478
tenure_months     0.0667
monthly_charges   0.0171
support_calls     0.0127
payment_method    0.0111
internet_service  0.0102
has_dependents    0.0082

```

Permuting whole original columns, before encoding, which is what you actually want to know: is `contract` useful, not is `contract_Two Year` useful.

⚠ Watch Out

Correlated features share the blame

Shuffle `tenure_months` while `total_charges` still carries most of the same information and the score barely moves, so both look unimportant. Permutation importance answers “how much does *this model* rely on this column”, never “how much does this column matter in the world”. Group correlated columns and permute them together if you need the second answer.

Day 6 takeaway

Day 7 Why the Simple Model Won

The bake-off, and matching an algorithm to the shape of the truth

Week 1 set you an assignment: swap logistic regression for a random forest and explain why the forest scored worse. Here is the answer, and it matters more than any algorithm this week.

The bake-off

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor(scale=True):
17     num_steps = [('i', SimpleImputer(strategy='median'))]
18     if scale:
19         num_steps.append(('s', StandardScaler()))
20     return ColumnTransformer([
21         ('num', Pipeline(num_steps), NUM),
22         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
23                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
24     ])
25 from sklearn.model_selection import train_test_split
26
27 X, y = df[NUM + CAT], df[TARGET]
28 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
29                                           stratify=y, random_state=42)
30 from sklearn.dummy import DummyClassifier
31 from sklearn.linear_model import LogisticRegression
32 from sklearn.tree import DecisionTreeClassifier
33 from sklearn.ensemble import RandomForestClassifier, HistGradientBoostingClassifier
34 import xgboost as xgb
35 from sklearn.model_selection import cross_val_score, StratifiedKFold
36
37 cv = StratifiedKFold(n_splits=5, shuffle=True, random_state=42)
38 candidates = {
39     'baseline': DummyClassifier(strategy='most_frequent'),
40     'logistic': LogisticRegression(max_iter=1000, random_state=42),
41     'tree leaf=50': DecisionTreeClassifier(min_samples_leaf=50, random_state=42),
42     'forest': RandomForestClassifier(n_estimators=300, min_samples_leaf=5,
43                                     random_state=42),
44     'hist boosting': HistGradientBoostingClassifier(random_state=42),
45     'xgboost': xgb.XGBClassifier(n_estimators=400, max_depth=3,
46                                  learning_rate=0.05, subsample=0.8,
47                                  random_state=42, eval_metric='logloss'),
48 }
49 for name, clf in candidates.items():
50     s = cross_val_score(Pipeline([('prep', preprocessor()), ('clf', clf)]),
51                         X_tr, y_tr, cv=cv, scoring='roc_auc')
```

```
52 print('%s %.4f +/- %.4f' % (name, s.mean(), s.std()))
```

```
baseline      0.5000 +/- 0.0000
logistic      0.8208 +/- 0.0137
tree leaf=50  0.7922 +/- 0.0078
forest        0.8203 +/- 0.0101
hist boosting 0.7973 +/- 0.0088
xgboost       0.8133 +/- 0.0083
```

A week of ensembles, and nothing beats the straight line

Logistic regression scores 0.8208. The best ensemble, a random forest with `min_samples_leaf=5`, scores 0.8203, which is a tie: the difference is a twentieth of the fold-to-fold spread. XGBoost is a point behind, and boosting with default settings is two points behind. Nothing here is a bug, and nothing you did is wrong. The honest summary is that a week of increasingly sophisticated machinery has managed to *draw* with the simplest model in the course.

Why

Look at how the data was generated, back in week 1:

```
1 # From make_dataset.py, week 1 day 3
2 #
3 #   logit = (-2.8
4 #           + 1.55 * (contract == 'Month-to-month')
5 #           - 0.85 * (contract == 'Two year')
6 #           - 0.055 * tenure
7 #           + 0.021 * monthly
8 #           + 0.30 * support
9 #           + 0.42 * (payment == 'Electronic check')
10 #          - 0.25 * (dependents == 'Yes'))
11 #   churned = sigmoid(logit) > uniform()
12
13 print('The truth is a weighted sum, passed through a sigmoid.')
14 print('That is exactly, precisely, the functional form of')
15 print('logistic regression -- and nothing else in this course.')
```

```
The truth is a weighted sum, passed through a sigmoid.
That is exactly, precisely, the functional form of
logistic regression -- and nothing else in this course.
```

Logistic regression is not approximating the truth here. It *is* the truth, with the right coefficients to be found. A tree can only approximate a smooth linear relationship with a staircase of splits, and no amount of boosting turns a staircase into a straight line.

Where trees win instead

```
1 import numpy as np
2 from sklearn.linear_model import LogisticRegression
3 from sklearn.ensemble import HistGradientBoostingClassifier
4 from sklearn.preprocessing import StandardScaler
5 from sklearn.pipeline import make_pipeline
```

```

6 from sklearn.model_selection import cross_val_score
7
8 rng = np.random.default_rng(0)
9 n = 3000
10 a = rng.uniform(-3, 3, n)
11 b = rng.uniform(-3, 3, n)
12
13 # An interaction with a threshold -- no linear model can express this.
14 y = (((a > 0) & (b > 0)) | ((a < -1.5) & (b < -1.5))).astype(int)
15 y = np.where(rng.uniform(size=n) < 0.05, 1 - y, y)      # 5% label noise
16 X = np.column_stack([a, b])
17
18 for name, clf in [('logistic', LogisticRegression(max_iter=1000)),
19                  ('boosting', HistGradientBoostingClassifier(random_state=42))]:
20     s = cross_val_score(make_pipeline(StandardScaler(), clf), X, y,
21                        cv=5, scoring='roc_auc')
22     print('%-9s AUC %.4f' % (name, s.mean()))

```

```

logistic  AUC 0.7184
boosting  AUC 0.9352

```

Same two features, a rule built from thresholds and interactions rather than a weighted sum, and the result reverses completely. Neither algorithm is better. They match different shapes of truth.

Prefer a linear model when	Prefer trees when
Effects are additive and smooth	Thresholds and interactions dominate
You need interpretable coefficients	You need raw predictive power
Data is wide and rows are few	Rows are plentiful
You must extrapolate beyond the training range	All prediction is inside the training range
Probabilities must be well calibrated	Ranking is what matters

⚠ Watch Out

Trees cannot extrapolate at all

A tree's prediction for any input beyond the training range is whatever the nearest leaf says, forever. Give a forest trained on tenures of 1 to 72 a customer with 200 months and it predicts as though they had 72. A linear model would continue the trend. Which behaviour is right depends entirely on your problem, but you must know which one you are getting.

```

1 import numpy as np
2 from sklearn.linear_model import LinearRegression
3 from sklearn.ensemble import RandomForestRegressor
4
5 x = np.linspace(0, 10, 200).reshape(-1, 1)
6 y = 3 * x.ravel() + 5
7
8 lin = LinearRegression().fit(x, y)
9 forest = RandomForestRegressor(n_estimators=100, random_state=0).fit(x, y)
10
11 print('%10s %14s %14s %12s' % ('x', 'truth', 'linear', 'forest'))

```

```

12 for v in [5.0, 10.0, 15.0, 50.0]:
13     print('%10.1f %14.2f %14.2f %12.2f'
14           % (v, 3 * v + 5, lin.predict([[v]])[0], forest.predict([[v]])[0]))

```

x	truth	linear	forest
5.0	20.00	20.00	19.97
10.0	35.00	35.00	34.92
15.0	50.00	50.00	34.92
50.0	155.00	155.00	34.92

Your assignment

Add an interaction the linear model cannot see, for example a flag for `month-to-month AND support_calls > 2`: to the churn features, and rerun the bake-off. Does the gap between logistic regression and boosting narrow? Then engineer that same interaction as an explicit column for the linear model and see who wins. That is week 8.

Day 7 takeaway

Self-Assessment

Try It Yourself

- A decision tree chooses each split by:
 - Greedily picking the split that most reduces impurity, without looking ahead
 - Solving for the global optimum
 - Testing every possible tree
 - Random selection
- The main weakness of a single deep tree is that it:
 - Cannot handle categories
 - Is unstable, a small change in the data can produce a very different tree
 - Is slow to fit
 - Requires scaling
- A random forest reduces variance by:
 - Boosting the errors
 - Pruning harder
 - Averaging many decorrelated trees, each on a bootstrap sample and a random subset of features
 - Using fewer features
- Boosting differs from bagging in that each new model is fitted to:
 - A different feature set
 - The same data with a new seed

- c) A random sample
 - d) The errors left by the models before it
5. Impurity-based feature importance is misleading because it:
- a) Favours high-cardinality and continuous features, which offer more places to split. It can rank pure noise first
 - b) Ignores interactions
 - c) Requires scaling
 - d) Is computed on test data
6. The honest alternative to impurity importance is:
- a) Coefficient magnitude
 - b) Permutation importance, measured on held-out data
 - c) Tree depth
 - d) Split count
7. What does out-of-bag scoring give you?
- a) Calibrated probabilities
 - b) A faster fit
 - c) A validation estimate for free, from the rows each tree did not see
 - d) A better model
8. Logistic regression matched the tree ensembles on the course dataset because:
- a) The trees were badly tuned
 - b) Trees need more features
 - c) The dataset is small
 - d) The underlying relationship really is close to linear in log-odds, so there is no extra structure for trees to find

Answers: 1a, 2b, 3c, 4d, 5a, 6b, 7c, 8d. Anything you got wrong points at the day to re-read, not at a gap in ability.

Model Evaluation and Validation

Week 7

Measuring a model honestly, the folds, the leaks, the imbalance, the calibration and the segments the average score hides.

OBJECTIVES

The most valuable week in the course, because a wrong measurement costs more than a mediocre model. It covers which cross-validation scheme a problem demands, the leaks cross-validation cannot detect, what actually helps with imbalance, why a well-ranking model can still be wrong about probabilities, and how to find who your model fails.

Day 1 Cross-Validation Done Properly

Stratified, grouped and time-ordered folds, and the nested version

You have used `cross_val_score` since week 4 without asking how the folds are made. That choice is where most silently broken evaluations come from.

Why one split is not enough

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19         ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21         ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),

```

```

22     ])
23 from sklearn.linear_model import LogisticRegression
24 from sklearn.model_selection import train_test_split
25 from sklearn.metrics import roc_auc_score
26 import numpy as np
27
28 X, y = df[NUM + CAT], df[TARGET]
29 scores = []
30 for seed in range(12):
31     a, b, c, d = train_test_split(X, y, test_size=0.25, stratify=y,
32                                   random_state=seed)
33     m = Pipeline([('prep', preprocessor()),
34                   ('clf', LogisticRegression(max_iter=1000))]).fit(a, c)
35     scores.append(roc_auc_score(d, m.predict_proba(b)[: , 1]))
36
37 scores = np.array(scores)
38 print('twelve different random splits of the same data:')
39 print('  min %.4f   max %.4f   spread %.4f'
40       % (scores.min(), scores.max(), scores.max() - scores.min()))
41 print('  mean %.4f  std %.4f' % (scores.mean(), scores.std()))

```

```

twelve different random splits of the same data:
  min 0.8030   max 0.8364   spread 0.0334
  mean 0.8207  std 0.0108

```

⚠ Watch Out

Two points of AUC, from nothing but the seed

Same data, same model, same code, only the split changed. If you compare two models on one split each and they differ by less than this spread, you have measured the seed, not the models. Reporting a single split's score to four decimal places implies a precision that does not exist.

The splitters, and when each is required

Splitter	Use when	Failure if you use KFold instead
KFold	Rows are independent and balanced	-
StratifiedKFold	Classification, any imbalance	Folds get different class rates; scores wobble
GroupKFold	Rows cluster by entity	The same entity lands in train and test, leakage
TimeSeriesSplit	Rows are ordered in time	The model trains on the future
RepeatedStratifiedKFold	You need a tighter estimate	-

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')

```

```

7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.model_selection import KFold, StratifiedKFold
12 import numpy as np
13
14 y = df[TARGET].to_numpy()
15 print('overall churn rate %.4f\n' % y.mean())
16
17 for name, splitter in [('KFold', KFold(5, shuffle=True, random_state=0)),
18                        ('StratifiedKFold', StratifiedKFold(5, shuffle=True,
19                                                            random_state=0))]:
20     rates = [y[test].mean() for _, test in splitter.split(np.zeros(len(y)), y)]
21     print('%s fold churn rates %s spread %.4f'
22           % (name, np.round(rates, 4), max(rates) - min(rates)))

```

```
overall churn rate 0.2680
```

```

KFold          fold churn rates [0.265  0.2733 0.25   0.2933 0.2583]  spread 0.0433
StratifiedKFold fold churn rates [0.2683 0.2683 0.2683 0.2683 0.2667]  spread 0.0017

```

Grouped data, where the trap is worst

```

1 import numpy as np
2 import pandas as pd
3 from sklearn.model_selection import cross_val_score, KFold, GroupKFold
4 from sklearn.ensemble import RandomForestClassifier
5
6 rng = np.random.default_rng(0)
7 n_customers, per_customer = 200, 5
8
9 # Five readings per customer. The label belongs to the customer, not the
10 # row, and each reading is that customer's traits plus a little noise --
11 # which is exactly what repeated measurement of one entity looks like.
12 customer = np.repeat(np.arange(n_customers), per_customer)
13 label = np.repeat(rng.integers(0, 2, n_customers), per_customer)
14 traits = np.repeat(rng.normal(size=(n_customers, 4)), per_customer, axis=0)
15 X = traits + rng.normal(0, 0.02, traits.shape)
16
17 clf = RandomForestClassifier(n_estimators=100, random_state=0)
18 naive = cross_val_score(clf, X, label, cv=KFold(5, shuffle=True, random_state=0))
19 grouped = cross_val_score(clf, X, label, cv=GroupKFold(5), groups=customer)
20
21 print('KFold          accuracy %.4f  ← the same customer is in both halves'
22       % naive.mean())
23 print('GroupKFold accuracy %.4f  ← honest' % grouped.mean())

```

```

KFold          accuracy 0.9970  ← the same customer is in both halves
GroupKFold accuracy 0.5880    ← honest

```

⚠ Watch Out**Ninety-nine percent, and the model has learned nothing**

Forty-one points of pure illusion. Repeated measurements of one patient, several sessions from one user, multiple photographs of one object, split at random and near-duplicates of every test row sit in the training set. The model does not learn the pattern; it memorises which customer each reading came from and looks up their label. `GroupKFold` keeps a customer's rows together and reveals the truth: 0.59, barely above chance.

This is the most expensive mistake in this course, because the model looks superb right up until it meets someone new. Whenever rows cluster by an entity, pass `groups`.

Nested cross-validation

If you tune hyperparameters by cross-validation and then report that same cross-validated score, the score is optimistic: you chose the settings that happened to suit those folds.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19         ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21         ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.ensemble import RandomForestClassifier
24 from sklearn.model_selection import (GridSearchCV, cross_val_score,
25                                     StratifiedKFold)
26 import numpy as np
27
28 X, y = df[NUM + CAT], df[TARGET]
29 grid = {'clf__min_samples_leaf': [1, 5, 20, 50]}
30 pipe = Pipeline([('prep', preprocessor()),
31                 ('clf', RandomForestClassifier(n_estimators=150,
32                                             random_state=42))])
33
34 inner = StratifiedKFold(4, shuffle=True, random_state=1)
35 outer = StratifiedKFold(4, shuffle=True, random_state=2)
36

```

```

37 search = GridSearchCV(pipe, grid, cv=inner, scoring='roc_auc')
38 search.fit(X, y)
39 print('best inner score (optimistic) %.4f' % search.best_score_)
40
41 nested = cross_val_score(search, X, y, cv=outer, scoring='roc_auc')
42 print('nested score      (honest)      %.4f  +/- %.4f'
43       % (nested.mean(), nested.std()))

```

```

best inner score (optimistic) 0.8216
nested score      (honest)      0.8176  +/- 0.0122

```

The nested figure re-runs the entire search inside each outer fold, so the settings are never chosen using the rows they are scored on. It costs inner times outer fits, which is why people skip it, but when you need to report a number you will be held to, this is the number.

Day 1 takeaway

Day 2 The Leaks Cross-Validation Cannot See

Scaling, selection and duplicates outside the pipeline, and noise scoring 0.7

Cross-validation protects you from one thing only: scoring on rows you trained on. It cannot see a leak that happened before the folds were made.

Scaling before the split

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.preprocessing import StandardScaler
12 from sklearn.linear_model import LogisticRegression
13 from sklearn.impute import SimpleImputer
14 from sklearn.pipeline import make_pipeline
15 from sklearn.model_selection import cross_val_score, StratifiedKFold
16 import numpy as np
17
18 X, y = df[NUM].to_numpy(), df[TARGET].to_numpy()
19 X = SimpleImputer(strategy='median').fit_transform(X)
20 cv = StratifiedKFold(5, shuffle=True, random_state=0)
21
22 # Wrong: the scaler sees every row, including the ones it will be tested on.
23 X_leaked = StandardScaler().fit_transform(X)
24 leaked = cross_val_score(LogisticRegression(max_iter=1000), X_leaked, y,

```

```

25         cv=cv, scoring='roc_auc')
26
27 # Right: the scaler is refitted inside every fold.
28 clean = cross_val_score(make_pipeline(StandardScaler(),
29                                     LogisticRegression(max_iter=1000)),
30                         X, y, cv=cv, scoring='roc_auc')
31
32 print('scaled before splitting %.4f' % leaked.mean())
33 print('scaled inside the fold %.4f' % clean.mean())
34 print('difference %.4f' % (leaked.mean() - clean.mean()))

```

```

scaled before splitting 0.7721
scaled inside the fold 0.7721
difference +0.0000

```

⚠ Watch Out

Zero difference is what makes this dangerous

The two numbers are identical to four decimal places. With three thousand rows, a mean and a standard deviation computed on 80 percent of them are indistinguishable from the same figures computed on all of them, so this particular leak costs nothing measurable. That is precisely the problem. The *same mistake* with target encoding, an aggregate feature, or feature selection produces gaps of ten points or more, and the next snippet turns pure noise into an apparent 0.80. Nothing in the score tells you which case you are in, so the rule cannot be “check whether it matters”. It has to be absolute: every step that learns from data goes inside the pipeline.

Feature selection before the split, which is catastrophic

```

1 import numpy as np
2 from sklearn.feature_selection import SelectKBest, f_classif
3 from sklearn.linear_model import LogisticRegression
4 from sklearn.pipeline import make_pipeline
5 from sklearn.model_selection import cross_val_score, StratifiedKFold
6
7 rng = np.random.default_rng(0)
8 n, p = 200, 5000
9 X = rng.normal(size=(n, p))          # pure noise
10 y = rng.integers(0, 2, n)           # unrelated labels
11 cv = StratifiedKFold(5, shuffle=True, random_state=0)
12
13 # Wrong: pick the 20 best columns using ALL the labels, then cross-validate.
14 best = SelectKBest(f_classif, k=20).fit_transform(X, y)
15 leaked = cross_val_score(LogisticRegression(max_iter=1000), best, y, cv=cv)
16
17 # Right: selection happens inside each fold.
18 clean = cross_val_score(make_pipeline(SelectKBest(f_classif, k=20),
19                                     LogisticRegression(max_iter=1000)),
20                         X, y, cv=cv)
21

```

```

22 print('there is no signal in this data at all -- 0.5 is the truth')
23 print('  selection outside the folds %.4f' % leaked.mean())
24 print('  selection inside the folds  %.4f' % clean.mean())

```

```

there is no signal in this data at all -- 0.5 is the truth
  selection outside the folds 0.7950
  selection inside the folds  0.5450

```

Random noise, random labels, no relationship whatsoever, and selecting features before cross-validating reports accuracy far above chance. Out of five thousand noise columns, twenty will correlate with the labels by luck, and choosing them using every label is how the luck gets in. This is a documented cause of retracted findings in published research.

Duplicate rows

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.ensemble import RandomForestClassifier
12 from sklearn.model_selection import cross_val_score, StratifiedKFold
13 from sklearn.impute import SimpleImputer
14 from sklearn.pipeline import make_pipeline
15 import pandas as pd
16
17 raw = pd.read_csv('customers.csv')          # 138 duplicated rows
18 deduped = raw.drop_duplicates()
19 cv = StratifiedKFold(5, shuffle=True, random_state=0)
20 pipe = make_pipeline(SimpleImputer(strategy='median'),
21                      RandomForestClassifier(n_estimators=200, random_state=42))
22
23 for name, d in [('with duplicates', raw), ('deduplicated', deduped)]:
24     s = cross_val_score(pipe, d[NUM], d[TARGET], cv=cv, scoring='roc_auc')
25     print('%s AUC %.4f (%d rows)' % (name, s.mean(), len(d)))

```

```

with duplicates    AUC 0.7503  (3138 rows)
deduplicated       AUC 0.7212  (3000 rows)

```

Only 138 duplicates in three thousand rows, so the inflation is modest. Scale that to an export that ran twice and the effect is severe. De-duplicate first, always.

A checklist you can actually run

1. Is every fitted transformation inside the Pipeline?
2. Were duplicates removed before splitting?

3. Do any rows share an entity that could straddle the split?
4. Are the rows time-ordered, and does the split respect that?
5. Was any feature computed using an aggregate over the whole dataset?
6. For every feature: would I have this value at prediction time?
7. Is the score suspiciously good?

Suspicion is a diagnostic

The last item is not a joke. Experienced practitioners treat an unexpectedly excellent score as a bug report. If a model jumps from 0.82 to 0.97 because you added one feature, the overwhelmingly likely explanation is that the feature leaks.

Day 2 takeaway

Day 3 Imbalanced Classes

Why accuracy rises as the problem gets harder, and what actually helps

Churn is 27 percent, which is mild. Fraud is 0.1 percent, and at that level everything you have learned about accuracy stops working.

The problem, at increasing severity

```

1 import numpy as np
2 from sklearn.datasets import make_classification
3 from sklearn.linear_model import LogisticRegression
4 from sklearn.model_selection import train_test_split
5 from sklearn.metrics import (accuracy_score, roc_auc_score,
6                             average_precision_score, recall_score)
7
8 print('%10s %10s %10s %10s %10s'
9       % ('positives', 'accuracy', 'recall', 'ROC AUC', 'avg prec'))
10 for weight in [0.5, 0.1, 0.02, 0.005]:
11     X, y = make_classification(n_samples=20000, n_features=12, n_informative=5,
12                             weights=[1 - weight], flip_y=0.01, random_state=0)
13     a, b, c, d = train_test_split(X, y, test_size=0.3, stratify=y, random_state=0)
14     m = LogisticRegression(max_iter=2000).fit(a, c)
15     pred, proba = m.predict(b), m.predict_proba(b)[: , 1]
16     print('%10.3f %10.4f %10.4f %10.4f %10.4f'
17           % (d.mean(), accuracy_score(d, pred), recall_score(d, pred),
18             roc_auc_score(d, proba), average_precision_score(d, proba)))

```

positives	accuracy	recall	ROC AUC	avg prec
0.500	0.9365	0.9494	0.9782	0.9724
0.104	0.9623	0.7388	0.9520	0.8674
0.025	0.9835	0.3791	0.8491	0.6144
0.011	0.9915	0.2031	0.6895	0.3925

Accuracy climbs toward 1.0 as the problem gets harder, because guessing the majority class gets easier. Recall collapses. Accuracy is not a weak metric here. It is an actively misleading one.

Class weights, which cost nothing

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 from sklearn.metrics import (recall_score, precision_score,
30                              average_precision_score, roc_auc_score)
31
32 for weight in [None, 'balanced']:
33     m = Pipeline([('prep', preprocessor()),
34                   ('clf', LogisticRegression(max_iter=1000, class_weight=weight,
35                                             random_state=42))]).fit(X_tr, y_tr)
36     pred, proba = m.predict(X_te), m.predict_proba(X_te)[: , 1]
37     print('class_weight=%-9s recall %.4f precision %.4f AUC %.4f AP %.4f'
38           % (str(weight), recall_score(y_te, pred),
39             precision_score(y_te, pred), roc_auc_score(y_te, proba),
40             average_precision_score(y_te, proba)))
```

```
class_weight=None      recall 0.4328 precision 0.6304 AUC 0.8174 AP 0.5993
class_weight=balanced  recall 0.7960 precision 0.4923 AUC 0.8172 AP 0.5973
```

class_weight='balanced' is mostly a threshold move

Recall jumps and precision falls, but ROC AUC and average precision barely change, because the *ranking* of customers by risk is nearly the same. Weighting mostly shifts where the 0.5 cut lands. That is worth knowing: if you were going to tune the threshold anyway, as week 5 did, class weights add little. They matter more when the algorithm's loss is genuinely dominated by the majority class.

Resampling

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.model_selection import train_test_split
12
13 X, y = df[NUM + CAT], df[TARGET]
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 from imblearn.over_sampling import SMOTE, RandomOverSampler
17 from imblearn.under_sampling import RandomUnderSampler
18 from imblearn.pipeline import Pipeline as ImbPipeline
19 from sklearn.compose import ColumnTransformer
20 from sklearn.pipeline import Pipeline
21 from sklearn.impute import SimpleImputer
22 from sklearn.preprocessing import StandardScaler, OneHotEncoder
23 from sklearn.linear_model import LogisticRegression
24 from sklearn.metrics import average_precision_score, roc_auc_score
25
26 prep = ColumnTransformer([
27     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
28                     ('s', StandardScaler())])), NUM),
29     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
30                     ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT]))
31
32 samplers = {'none': None,
33             'oversample': RandomOverSampler(random_state=0),
34             'undersample': RandomUnderSampler(random_state=0),
35             'SMOTE': SMOTE(random_state=0)}
36
37 for name, sampler in samplers.items():
38     steps = [('prep', prep)]
39     if sampler is not None:
40         steps.append(('sample', sampler))
41     steps.append(('clf', LogisticRegression(max_iter=1000, random_state=42)))
42     m = ImbPipeline(steps).fit(X_tr, y_tr)
43     proba = m.predict_proba(X_te)[: , 1]

```

```

44 print('%s AUC %.4f    avg precision %.4f'
45       % (name, roc_auc_score(y_te, proba),
46         average_precision_score(y_te, proba)))

```

```

none          AUC 0.8174    avg precision 0.5993
oversample    AUC 0.8143    avg precision 0.5923
undersample    AUC 0.8110    avg precision 0.5856
SMOTE         AUC 0.8160    avg precision 0.5951

```

⚠ Watch Out

Resample inside the pipeline, and never the test set

`imblearn`'s `Pipeline` exists precisely because `scikit-learn`'s will not resample. It applies the sampler during `fit` and skips it during `predict`, which is the only correct behaviour. Oversampling before the split copies minority rows into both halves. The model is then tested on rows it trained on, and the score is fiction.

Notice how little any of it achieved on this problem. Resampling is oversold; it changes the ranking barely at all, and at 27 percent minority there is not much to fix. Reach for it when the minority class is under a few percent, and measure rather than assume.

What to actually do about imbalance

1. **Use the right metric:** average precision, not accuracy. This alone fixes most of the confusion.
2. **Tune the threshold** against the cost of each error, as in week 5. Free, and usually the largest gain.
3. **Try class weights**, which cost one argument.
4. **Then consider resampling**, measuring whether it helped.
5. **Collect more minority examples** if you possibly can. Nothing else comes close.

Day 3 takeaway

Day 4 Calibration

Ranking well and still being wrong about the numbers

A model that ranks customers perfectly can still be badly wrong about the numbers. If anyone will multiply your probability by a pound value, ranking is not enough.

Measure it

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()

```

```

6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 from sklearn.ensemble import RandomForestClassifier
30 from sklearn.naive_bayes import GaussianNB
31 from sklearn.calibration import calibration_curve
32 from sklearn.metrics import brier_score_loss, roc_auc_score
33 import numpy as np
34
35 models = {
36     'logistic': LogisticRegression(max_iter=1000, random_state=42),
37     'forest': RandomForestClassifier(n_estimators=300, random_state=42),
38     'naive-B': GaussianNB(),
39 }
40 for name, clf in models.items():
41     m = Pipeline([('prep', preprocessor()), ('clf', clf)]).fit(X_tr, y_tr)
42     proba = m.predict_proba(X_te)[: , 1]
43     print('%s AUC %.4f Brier %.4f'
44           % (name, roc_auc_score(y_te, proba),
45              brier_score_loss(y_te, proba)))

```

```

logistic AUC 0.8174 Brier 0.1475
forest AUC 0.7772 Brier 0.1674
naive-B AUC 0.8044 Brier 0.2438

```

The Brier score is the mean squared error of the probabilities. Lower is better, and unlike AUC it punishes a model for being confidently wrong about the level rather than just the order.

The reliability table

```

1 import numpy as np
2 import pandas as pd
3

```

```

4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 from sklearn.naive_bayes import GaussianNB
30 from sklearn.calibration import calibration_curve
31 import numpy as np
32
33 for name, clf in [('logistic', LogisticRegression(max_iter=1000, random_state=42)),
34                  ('naive-B ', GaussianNB())]:
35     m = Pipeline([('prep', preprocessor()), ('clf', clf)]).fit(X_tr, y_tr)
36     proba = m.predict_proba(X_te)[: , 1]
37     true_rate, predicted = calibration_curve(y_te, proba, n_bins=6,
38                                           strategy='quantile')
39     print('%s    predicted → actual' % name)
40     for pr, ac in zip(predicted, true_rate):
41         print('          %.3f    → %.3f' % (pr, ac))
42     print()

```

```

logistic    predicted → actual
           0.010    → 0.008
           0.058    → 0.072
           0.162    → 0.152
           0.285    → 0.304
           0.433    → 0.456
           0.628    → 0.616

naive-B     predicted → actual
           0.000    → 0.024
           0.003    → 0.064
           0.123    → 0.160
           0.613    → 0.312
           0.955    → 0.448
           0.989    → 0.600

```

Logistic regression tracks the diagonal closely, which is expected: it is fitted by minimising log loss, and log loss is minimised by telling the truth about probabilities. Naive Bayes, which multiplies correlated probabilities as though independent, is pushed toward the extremes.

Repairing it

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                          ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                          ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                         stratify=y, random_state=42)
28 from sklearn.naive_bayes import GaussianNB
29 from sklearn.calibration import CalibratedClassifierCV
30 from sklearn.metrics import brier_score_loss, roc_auc_score
31
32 base = Pipeline([('prep', preprocessor()), ('clf', GaussianNB())])
33
34 fitted = base.fit(X_tr, y_tr)
35 raw = fitted.predict_proba(X_te)[: , 1]
36 print('uncalibrated   Brier %.4f   AUC %.4f'
37       % (brier_score_loss(y_te, raw), roc_auc_score(y_te, raw)))
38
39 for method in ['sigmoid', 'isotonic']:
40     cal = CalibratedClassifierCV(base, method=method, cv=5).fit(X_tr, y_tr)
41     proba = cal.predict_proba(X_te)[: , 1]
42     print('%-14s Brier %.4f   AUC %.4f'
43           % (method, brier_score_loss(y_te, proba), roc_auc_score(y_te, proba)))

```

uncalibrated	Brier 0.2438	AUC 0.8044
sigmoid	Brier 0.1559	AUC 0.8060
isotonic	Brier 0.1525	AUC 0.8043

Method	Fits	Use when
sigmoid (Platt)	A logistic curve, two parameters	Small data; the distortion is a smooth S-shape
isotonic	Any non-decreasing step function	Plenty of data; more flexible, can overfit

Calibration barely changes AUC

Both methods apply a monotone transformation, and a monotone transformation cannot change the ordering. So AUC stays put while the Brier score improves substantially. That is the clearest possible demonstration that ranking and calibration are separate properties.

When it matters

- **Expected value calculations:** probability times customer value only means something if the probability is real.
- **Combining models:** averaging miscalibrated probabilities is meaningless.
- **Anything shown to a human:** a clinician told 80 percent will act on 80 percent.
- **Thresholds set from cost:** week 5's calculation assumed the probabilities meant something.

If you only ever sort customers and contact the top thousand, calibration is irrelevant. Know which situation you are in.

Day 4 takeaway

Day 5 Learning and Validation Curves

Diagnosing whether you need more data, more features or less model

Your model underperforms. Do you collect more data, add features, or pick a different algorithm? Guessing wastes weeks. Two curves answer it.

Learning curves: is more data the answer?

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder

```

```

15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.linear_model import LogisticRegression
24 from sklearn.tree import DecisionTreeClassifier
25 from sklearn.model_selection import learning_curve, StratifiedKFold
26 import numpy as np
27
28 X, y = df[NUM + CAT], df[TARGET]
29 cv = StratifiedKFold(5, shuffle=True, random_state=0)
30
31 for name, clf in [('logistic (well matched)',
32                   LogisticRegression(max_iter=1000)),
33                  ('deep tree (high variance)',
34                   DecisionTreeClassifier(random_state=0))]:
35     sizes, tr, va = learning_curve(
36         Pipeline([('prep', preprocessor()), ('clf', clf)]), X, y, cv=cv,
37         scoring='roc_auc', train_sizes=np.linspace(0.1, 1.0, 5))
38     print(name)
39     print('%10s %10s %10s %8s' % ('rows', 'train', 'val', 'gap'))
40     for n, a, b in zip(sizes, tr.mean(axis=1), va.mean(axis=1)):
41         print('%10d %10.4f %10.4f %8.4f' % (n, a, b, a - b))
42     print()

```

```

logistic (well matched)
  rows    train    val    gap
   240    0.8044    0.8056 -0.0012
   780    0.8141    0.8180 -0.0039
  1320    0.8245    0.8180  0.0065
  1860    0.8236    0.8198  0.0038
  2400    0.8267    0.8220  0.0047

deep tree (high variance)
  rows    train    val    gap
   240    1.0000    0.5973  0.4027
   780    1.0000    0.6119  0.3881
  1320    1.0000    0.6207  0.3793
  1860    1.0000    0.6270  0.3730
  2400    1.0000    0.6344  0.3656

```

Pattern	Diagnosis	What to do
Both curves low, converged	High bias, underfitting	Better features, a more flexible model
Wide gap, validation still rising	High variance	More data, more regularisation, fewer features
Both high, converged	You are done	Ship it
Validation falling as data grows	Something is wrong	Check for leakage or a distribution shift

The logistic curves converge with almost no gap: it has extracted what this feature set contains, and more rows will not help. The deep tree keeps a large gap at every size, which is the signature of variance. There, more data genuinely would.

Validation curves: is this parameter set right?

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23
24 from sklearn.ensemble import RandomForestClassifier
25 from sklearn.model_selection import validation_curve, StratifiedKFold
26 import numpy as np
27
28 X, y = df[NUM + CAT], df[TARGET]
29 values = [1, 2, 5, 10, 25, 50, 100, 200]
30 tr, va = validation_curve(
31     Pipeline([('prep', preprocessor()),
32               ('clf', RandomForestClassifier(n_estimators=150, random_state=0))]),
33     X, y, param_name='clf__min_samples_leaf', param_range=values,
34     cv=StratifiedKFold(4, shuffle=True, random_state=0), scoring='roc_auc')
35
36 print('%18s %10s %10s %8s' % ('min_samples_leaf', 'train', 'val', 'gap'))
37 for v, a, b in zip(values, tr.mean(axis=1), va.mean(axis=1)):
38     print('%18d %10.4f %10.4f %8.4f' % (v, a, b, a - b))

```

min_samples_leaf	train	val	gap
1	1.0000	0.7837	0.2163
2	0.9834	0.7984	0.1851
5	0.9198	0.8117	0.1081
10	0.8810	0.8151	0.0659
25	0.8492	0.8187	0.0305
50	0.8357	0.8179	0.0178
100	0.8262	0.8158	0.0104
200	0.8214	0.8151	0.0063

Read it left to right: at leaf size 1 the training score is near perfect and the gap enormous. As the leaves grow the gap closes and validation improves, up to a point, then both fall as the model becomes too constrained. The best value is where validation peaks, and the shape tells you whether you are on the overfitting or underfitting side of it.

⚠ Watch Out

A peak at the edge of your range means the range is wrong

If the best value is the largest or smallest you tried, you have not found the optimum. You have found the edge of your search. Extend the range and look again. The same rule applied to ridge's alpha in week 4, and it applies to every hyperparameter search in week 14.

Putting both together

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.ensemble import HistGradientBoostingClassifier
29 from sklearn.linear_model import LogisticRegression
30 from sklearn.model_selection import cross_validate, StratifiedKFold
31
32 cv = StratifiedKFold(5, shuffle=True, random_state=0)
33 for name, clf in [('logistic', LogisticRegression(max_iter=1000, random_state=0)),
34                  ('boosting', HistGradientBoostingClassifier(random_state=0))]:
35     r = cross_validate(Pipeline([('prep', preprocessor()), ('clf', clf)]),
36                       X_tr, y_tr, cv=cv, scoring='roc_auc',

```

```

37         return_train_score=True)
38     print('%-9s train %.4f   val %.4f   gap %+.4f'
39           % (name, r['train_score'].mean(), r['test_score'].mean(),
40               r['train_score'].mean() - r['test_score'].mean()))

```

```

logistic  train 0.8289   val 0.8238   gap +0.0052
boosting  train 0.9845   val 0.7966   gap +0.1879

```

`return_train_score=True` is the cheapest diagnostic in scikit-learn and almost nobody uses it. One extra argument turns a score into a diagnosis: boosting's gap is many times logistic regression's, which says it is fitting noise this dataset does not reward.

Day 5 takeaway

Day 6 Error Analysis and Fairness

Who the model fails, and why the headline score hides it

An aggregate score tells you how the model does on average. It never tells you who it fails, and that is what determines whether you can deploy it.

Score by segment

```

1  import numpy as np
2  import pandas as pd
3
4  df = pd.read_csv('customers.csv').drop_duplicates().copy()
5  df['contract'] = df['contract'].str.strip().str.title()
6  df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8  NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9  CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median'))],
19                           ('s', StandardScaler()))), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                           ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)

```

```

28 from sklearn.linear_model import LogisticRegression
29 from sklearn.metrics import roc_auc_score, average_precision_score
30 import pandas as pd
31
32 m = Pipeline([('prep', preprocessor()),
33               ('clf', LogisticRegression(max_iter=1000,
34                                         random_state=42))]).fit(X_tr, y_tr)
35 proba = m.predict_proba(X_te)[: , 1]
36
37 seg = df.loc[X_te.index].copy()
38 seg['proba'] = proba
39 seg['truth'] = y_te.to_numpy()
40
41 rows = []
42 for name, group in seg.groupby('contract'):
43     if group['truth'].nunique() < 2:
44         continue
45     rows.append({'segment': name, 'n': len(group),
46                'churn_rate': group['truth'].mean(),
47                'auc': roc_auc_score(group['truth'], group['proba']),
48                'avg_prec': average_precision_score(group['truth'],
49                                                    group['proba'])})
50 print(pd.DataFrame(rows).round(4).to_string(index=False))
51 print('\noverall AUC %.4f' % roc_auc_score(y_te, proba))

```

```

      segment    n  churn_rate    auc  avg_prec
Month-To-Month 416    0.4135 0.7072    0.6279
      One Year  184    0.1141 0.8113    0.4052
      Two Year  150    0.0533 0.8688    0.3955

overall AUC 0.8174

```

⚠ Watch Out

Segment AUCs can all be lower than the overall AUC

This looks impossible and is not. Much of the overall ranking skill comes from separating month-to-month customers from two-year ones. *Within* a contract type that advantage is gone, and only the weaker signals remain. If your product treats each segment separately, the within-segment number is the one that matters, and it is worse than the headline.

Look at what it got most wrong

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']

```

```

10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 import pandas as pd
30 import numpy as np
31
32 m = Pipeline([('prep', preprocessor()),
33               ('clf', LogisticRegression(max_iter=1000,
34                                         random_state=42))]).fit(X_tr, y_tr)
35 seg = df.loc[X_te.index].copy()
36 seg['proba'] = m.predict_proba(X_te)[: , 1]
37 seg['error'] = (seg[TARGET] - seg['proba']).abs()
38
39 cols = ['tenure_months', 'contract', 'support_calls', 'monthly_charges',
40         TARGET, 'proba']
41 print('confidently wrong -- said stay, they left:')
42 worst = seg[seg[TARGET] == 1].nsmallest(4, 'proba')[cols]
43 print(worst.round(3).to_string(index=False))
44 print('\nconfidently wrong -- said leave, they stayed:')
45 worst = seg[seg[TARGET] == 0].nlargest(4, 'proba')[cols]
46 print(worst.round(3).to_string(index=False))

```

confidently wrong -- said stay, they left:

tenure_months	contract	support_calls	monthly_charges	churned	proba
23	Two Year	0	26.42	1	0.009
41	One Year	0	53.26	1	0.031
12	Two Year	1	NaN	1	0.042
19	One Year	0	23.08	1	0.047

confidently wrong -- said leave, they stayed:

tenure_months	contract	support_calls	monthly_charges	churned	proba
5	Month-To-Month	3	73.19	0	0.782
6	Month-To-Month	1	79.59	0	0.751
17	Month-To-Month	4	86.67	0	0.729
9	Month-To-Month	1	79.64	0	0.719

Long-tenure two-year customers who left anyway, and short-tenure month-to-month customers who stayed. Both groups contradict the strongest pattern in the data. In a real project this

is where you go and ask someone: is there a reason these people behaved differently, and is it something we could record?

Where the errors concentrate

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                          ('s', StandardScaler())])), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                          ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 import pandas as pd
30
31 m = Pipeline([('prep', preprocessor()),
32              ('clf', LogisticRegression(max_iter=1000,
33                                         random_state=42))]).fit(X_tr, y_tr)
34 seg = df.loc[X_te.index].copy()
35 seg['proba'] = m.predict_proba(X_te)[: , 1]
36 seg['error'] = (seg[TARGET] - seg['proba']).abs()
37 seg['tenure_band'] = pd.cut(seg['tenure_months'], [0, 6, 18, 36, 72],
38                             labels=['0-6', '7-18', '19-36', '37-72'])
39
40 print(seg.groupby('tenure_band', observed=True)
41       .agg(n=('error', 'size'), mean_error=('error', 'mean'),
42            churn_rate=(TARGET, 'mean'), mean_proba=('proba', 'mean'))
43       .round(4).to_string())

```

	n	mean_error	churn_rate	mean_proba
tenure_band				
0-6	72	0.4086	0.4167	0.4470
7-18	314	0.3982	0.3854	0.3716

19-36	218	0.2457	0.1927	0.1886
37-72	146	0.0791	0.0548	0.0490

Compare the last two columns: where the mean predicted probability sits below the actual churn rate, the model systematically under-predicts that band. That is a bias you can act on, and it is invisible in any single aggregate number.

Fairness is the same technique with higher stakes

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                             ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                             ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.linear_model import LogisticRegression
29 from sklearn.metrics import recall_score, precision_score
30 import pandas as pd
31
32 m = Pipeline([('prep', preprocessor()),
33               ('clf', LogisticRegression(max_iter=1000,
34                                         random_state=42))]).fit(X_tr, y_tr)
35 seg = df.loc[X_te.index].copy()
36 seg['pred'] = m.predict(X_te)
37 seg['truth'] = y_te.to_numpy()
38
39 rows = []
40 for name, g in seg.groupby('has_dependents', dropna=False):
41     rows.append({'group': str(name), 'n': len(g),
42                 'selection_rate': g['pred'].mean(),
43                 'recall': recall_score(g['truth'], g['pred'], zero_division=0),
44                 'precision': precision_score(g['truth'], g['pred'],

```

```

45         zero_division=0))}
46 print(pd.DataFrame(rows).round(4).to_string(index=False))

```

```

group  n  selection_rate  recall  precision
No  498         0.1988  0.4532    0.6364
Yes  205         0.1561  0.4286    0.6562
nan   47         0.1489  0.2308    0.4286

```

Equal treatment has several incompatible definitions

Equal selection rate, equal recall, and equal precision are three different fairness criteria, and it is mathematically impossible to satisfy all of them at once when the groups have different base rates. You have to choose which one your situation demands and say so out loud. For a retention offer, unequal selection rates may be fine. For a loan decision they are not.

Day 6 takeaway

Day 7 A Reusable Evaluation Report

One function covering stability, discrimination, calibration and segments

Everything this week, as one function you can run on any classifier.

The report

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median'))],
19         ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21         ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24
25 X, y = df[NUM + CAT], df[TARGET]

```

```

26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                         stratify=y, random_state=42)
28 from sklearn.metrics import (roc_auc_score, average_precision_score,
29                             brier_score_loss, confusion_matrix,
30                             precision_score, recall_score)
31 from sklearn.model_selection import cross_validate, StratifiedKFold
32 import numpy as np
33 import pandas as pd
34
35 def evaluate(model, X_tr, y_tr, X_te, y_te, frame, costs=None):
36     """A full evaluation: stability, discrimination, calibration,
37     the operating point, and where it fails."""
38     cv = StratifiedKFold(5, shuffle=True, random_state=0)
39     r = cross_validate(model, X_tr, y_tr, cv=cv, scoring='roc_auc',
40                       return_train_score=True)
41     print('CROSS-VALIDATION')
42     print('  train AUC %.4f  val AUC %.4f +/- %.4f  gap %.4f'
43           % (r['train_score'].mean(), r['test_score'].mean(),
44              r['test_score'].std(),
45              r['train_score'].mean() - r['test_score'].mean()))
46
47     model.fit(X_tr, y_tr)
48     proba = model.predict_proba(X_te)[: , 1]
49     print('\nHELD-OUT TEST')
50     print('  ROC AUC          %.4f' % roc_auc_score(y_te, proba))
51     print('  average precision %.4f (baseline %.4f)'
52           % (average_precision_score(y_te, proba), y_te.mean()))
53     print('  Brier score      %.4f' % brier_score_loss(y_te, proba))
54
55     if costs:
56         offer, value, save = costs
57         grid = np.arange(0.05, 0.95, 0.01)
58         net = []
59         for t in grid:
60             tn, fp, fn, tp = confusion_matrix(y_te, (proba >=
61 ↪ t).astype(int)).ravel()
62             net.append(tp * save * value - (tp + fp) * offer)
63             best_t = grid[int(np.argmax(net))]
64             pred = (proba >= best_t).astype(int)
65             print('\nOPERATING POINT')
66             print('  threshold %.2f (net value %.0f, against %.0f at 0.5)'
67                   % (best_t, max(net), net[int(np.argmin(np.abs(grid - 0.5)))]))
68             print('  precision %.4f  recall %.4f'
69                   % (precision_score(y_te, pred, zero_division=0),
70                      recall_score(y_te, pred)))
71         else:
72             pred = (proba >= 0.5).astype(int)
73
74     print('\nBY SEGMENT')
75     seg = frame.loc[X_te.index].copy()
76     seg['proba'], seg['truth'] = proba, y_te.to_numpy()
77     for name, g in seg.groupby('contract'):
78         if g['truth'].nunique() < 2:
79             continue

```

```

79     print(' %-16s n=%4d  churn %.3f  AUC %.4f'
80           % (name, len(g), g['truth'].mean(),
81             roc_auc_score(g['truth'], g['proba'])))
82     return proba
83
84 from sklearn.linear_model import LogisticRegression
85 model = Pipeline([('prep', preprocessor()),
86                  ('clf', LogisticRegression(max_iter=1000, random_state=42))])
87 _ = evaluate(model, X_tr, y_tr, X_te, y_te, df, costs=(30, 250, 0.35))

```

CROSS-VALIDATION

```
train AUC 0.8289   val AUC 0.8238 +/- 0.0305   gap +0.0052
```

HELD-OUT TEST

```
ROC AUC           0.8174
average precision 0.5993 (baseline 0.2680)
Brier score       0.1475
```

OPERATING POINT

```
threshold 0.34 (net value 4537, against 3472 at 0.5)
precision 0.5423   recall 0.7015
```

BY SEGMENT

```
Month-To-Month  n= 416  churn 0.413  AUC 0.7072
One Year        n= 184  churn 0.114  AUC 0.8113
Two Year        n= 150  churn 0.053  AUC 0.8688
```

Run it on a competitor

```

1  import numpy as np
2  import pandas as pd
3
4  df = pd.read_csv('customers.csv').drop_duplicates().copy()
5  df['contract'] = df['contract'].str.strip().str.title()
6  df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8  NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9  CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 TARGET = 'churned'
11 from sklearn.compose import ColumnTransformer
12 from sklearn.pipeline import Pipeline
13 from sklearn.impute import SimpleImputer
14 from sklearn.preprocessing import StandardScaler, OneHotEncoder
15
16 def preprocessor():
17     return ColumnTransformer([
18         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                          ('s', StandardScaler())]), NUM),
20         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                          ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22     ])
23 from sklearn.model_selection import train_test_split
24

```

```

25 X, y = df[NUM + CAT], df[TARGET]
26 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
27                                           stratify=y, random_state=42)
28 from sklearn.ensemble import HistGradientBoostingClassifier
29 from sklearn.model_selection import cross_validate, StratifiedKFold
30 from sklearn.metrics import (roc_auc_score, average_precision_score,
31                              brier_score_loss)
32
33 cv = StratifiedKFold(5, shuffle=True, random_state=0)
34 for name, clf in [('logistic', __import__('sklearn.linear_model',
35                                           fromlist=['LogisticRegression']).LogisticRegression(
36                                           max_iter=1000, random_state=42)),
37                  ('boosting', HistGradientBoostingClassifier(random_state=42))]:
38     m = Pipeline([('prep', preprocessor()), ('clf', clf)])
39     r = cross_validate(m, X_tr, y_tr, cv=cv, scoring='roc_auc',
40                       return_train_score=True)
41     m.fit(X_tr, y_tr)
42     proba = m.predict_proba(X_te)[: , 1]
43     print('%-9s val %.4f +/-%.4f gap %+.4f test AUC %.4f AP %.4f Brier %.4f'
44           % (name, r['test_score'].mean(), r['test_score'].std(),
45              r['train_score'].mean() - r['test_score'].mean(),
46              roc_auc_score(y_te, proba), average_precision_score(y_te, proba),
47              brier_score_loss(y_te, proba)))

```

```

logistic val 0.8238 +/-0.0305 gap +0.0052 test AUC 0.8174 AP 0.5993 Brier 0.1475
boosting val 0.7966 +/-0.0167 gap +0.1879 test AUC 0.7797 AP 0.5262 Brier 0.1690

```

Three numbers decide it: the validation score, the train-validation gap, and the Brier score. Boosting fits harder and generalises no better here, which is the same conclusion week 6 reached from a different direction.

What to write down

1. The metric you optimised and **why that one**.
2. The cross-validated score with its spread, never a bare number.
3. The train-validation gap.
4. The held-out test score, computed once.
5. The operating threshold and the cost assumptions behind it.
6. Performance broken down by every segment that matters.
7. The known failure modes, from the confidently-wrong rows.
8. The date, the data version, and the random seeds.

The last one is not bureaucracy

In six months someone will ask why the model behaves differently. Without the data version and the seeds you cannot reproduce what you did, so you cannot answer. Week 15 turns this into an artefact the pipeline writes automatically.

Your assignment

Run `evaluate` on the random forest and on naive Bayes. One will have a much larger train-validation gap; the other will have a much worse Brier score. Write one paragraph recommending a model for a retention programme that multiplies each probability by customer value, and note which of the two failings actually disqualifies a model for that use.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. Repeated measurements of the same customer appear in the data. The correct cross-validation is:
 - a) `GroupKFold`, grouped by customer
 - b) `ShuffleSplit`
 - c) `KFold`
 - d) `StratifiedKFold`
2. Selecting features on the whole dataset before cross-validating can make pure noise score:
 - a) Well above chance, because the selection already saw the labels of every fold
 - b) Exactly 1.00
 - c) Below chance
 - d) 0.50, as expected
3. Nested cross-validation exists to:
 - a) Speed up tuning
 - b) Give an honest performance estimate for a procedure that includes tuning
 - c) Reduce variance
 - d) Replace the test set
4. As a class becomes rarer, plain accuracy:
 - a) Becomes undefined
 - b) Falls
 - c) Rises, because always predicting the majority gets easier
 - d) Stays the same
5. A model ranks well but its predicted probabilities are systematically too high. It needs:
 - a) A lower learning rate
 - b) More data
 - c) More features
 - d) Calibration

6. A learning curve where training and validation scores have converged at a low value indicates:
- a) Underfitting, more data will not help, but more capacity or better features might
 - b) Leakage
 - c) A bad split
 - d) Overfitting
7. The Brier score measures:
- a) Ranking quality
 - b) The mean squared error of predicted probabilities, so it captures calibration as well as discrimination
 - c) The area under the curve
 - d) Class balance
8. Overall AUC is 0.82 but one customer segment scores 0.61. This means:
- a) AUC is the wrong metric
 - b) The model is fine
 - c) The headline number is hiding a group the model serves badly, which is a fairness and a quality problem
 - d) The segment is too small to matter

Answers: 1a, 2a, 3b, 4c, 5d, 6a, 7b, 8c. Anything you got wrong points at the day to re-read, not at a gap in ability.

Feature Engineering

Week 8

Turning raw columns into inputs a model can use, and measuring honestly enough to know which ones helped.

OBJECTIVES

Better inputs beat better models, but only when you can tell which inputs are better. This week covers the transformations, encodings and interactions worth trying, the custom transformer that makes a group-statistic feature safe, and above all the discipline of measuring every idea against a fixed baseline on fixed folds, including the many that will not work.

Day 1 Why Features Beat Models

The discipline, the baseline, and a first attempt that barely moves

Week 6 ended with a week of ensembles drawing against a straight line. That is not a limit of the algorithms; it is a limit of the features. Better inputs beat better models, reliably, and this week is where the gains actually live.

The discipline that makes it safe

Every feature you build learns something, a mean, a category list, a bin edge. Learn it from all your data and you have leaked. So every engineered feature belongs inside the pipeline, fitted on training folds only. That is not a style preference; week 7 showed it turning noise into an apparent 0.80.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()
6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16

```

```

17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20                             ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22                             ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X, y = df[NUM + CAT], df[TARGET]
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                         stratify=y, random_state=42)
29 from sklearn.linear_model import LogisticRegression
30 from sklearn.model_selection import cross_val_score, StratifiedKFold
31
32 cv = StratifiedKFold(5, shuffle=True, random_state=0)
33 baseline = Pipeline([('prep', preprocessor()),
34                       ('clf', LogisticRegression(max_iter=1000,
35                                                  random_state=42))])
36 s = cross_val_score(baseline, X_tr, y_tr, cv=cv, scoring='roc_auc')
37 print('the number to beat: %.4f +/- %.4f' % (s.mean(), s.std()))

```

```
the number to beat: 0.8238 +/- 0.0305
```

Write the baseline down before you start

Feature engineering without a recorded baseline is how people spend a fortnight building features that made things worse. Every idea this week gets measured against 0.8238, on the same folds, and kept only if it beats it by more than the fold spread.

Where features come from

Source	Example	Watch out for
Transformation	$\log(\text{charges})$	Undefined at zero, use <code>log1p</code>
Ratio	charges per month of tenure	Division by zero; and it may reconstruct the target
Interaction	month-to-month <i>and</i> high support calls	Column count explodes quickly
Aggregate	average spend for this contract type	Must be computed on training folds only
Time	months since signup	Never derive from today's date
Domain knowledge	is this customer out of contract	The most valuable and the least automatable

A first attempt, measured

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()

```

```

6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.linear_model import LogisticRegression
25 from sklearn.model_selection import cross_val_score, StratifiedKFold,
    ↪ train_test_split
26 import numpy as np
27
28 eng = df.copy()
29 eng['charge_per_tenure'] = eng['monthly_charges'] / eng['tenure_months']
30 eng['calls_per_year'] = eng['support_calls'] / (eng['tenure_months'] / 12)
31 eng['log_calls'] = np.log1p(eng['support_calls'])
32
33 NEW = NUM + ['charge_per_tenure', 'calls_per_year', 'log_calls']
34 Xe, ye = eng[NEW + CAT], eng[TARGET]
35 Xe_tr, Xe_te, ye_tr, ye_te = train_test_split(Xe, ye, test_size=0.25,
36         stratify=ye, random_state=42)
37
38 cv = StratifiedKFold(5, shuffle=True, random_state=0)
39 m = Pipeline([('prep', preprocessor(num=NEW)),
40         ('clf', LogisticRegression(max_iter=1000, random_state=42))])
41 s = cross_val_score(m, Xe_tr, ye_tr, cv=cv, scoring='roc_auc')
42 print('with three new features: %.4f +/- %.4f' % (s.mean(), s.std()))
43 print('baseline was          0.8238')

```

```

with three new features: 0.8278 +/- 0.0220
baseline was          0.8238

```

Barely moved, and that is the normal outcome of a first attempt. Ratios invented without a reason rarely help. The features that work come from knowing something about the problem, which is the rest of this week.

Day 1 takeaway

Day 2 Numeric Transformations

Choosing a scaler, and buying curvature with bins or splines

Scaling is not one choice. Which scaler you pick changes what the model sees, and the default is wrong whenever outliers are present.

The three scalers

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()
6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.preprocessing import (StandardScaler, MinMaxScaler,
13                                   RobustScaler, QuantileTransformer)
14 import numpy as np
15
16 x = df[['support_calls']].to_numpy(dtype=float) # contains the 97 outlier
17
18 print('%-22s %8s %8s %8s %8s' % ('', 'median', 'IQR', 'max', 'p99'))
19 print('%-22s %8.3f %8.3f %8.3f %8.3f'
20       % ('raw', np.median(x), np.subtract(*np.percentile(x, [75, 25])),
21         x.max(), np.percentile(x, 99)))
22 for name, sc in [('StandardScaler', StandardScaler()),
23                 ('MinMaxScaler', MinMaxScaler()),
24                 ('RobustScaler', RobustScaler()),
25                 ('QuantileTransformer', QuantileTransformer(
26                   output_distribution='normal', random_state=0))]:
27     z = sc.fit_transform(x)
28     print('%-22s %8.3f %8.3f %8.3f %8.3f'
29           % (name, np.median(z), np.subtract(*np.percentile(z, [75, 25])),
30             z.max(), np.percentile(z, 99)))

```

	median	IQR	max	p99
raw	1.000	2.000	97.000	5.000
StandardScaler	-0.127	0.952	45.590	1.778
MinMaxScaler	0.010	0.021	1.000	0.052
RobustScaler	0.000	1.000	48.000	2.000
QuantileTransformer	-0.045	5.890	5.199	2.483

Scaler	Centres on	Divides by	Use when
StandardScaler	Mean	Standard deviation	Roughly symmetric, no extreme values
MinMaxScaler	Minimum	Range	You need a bounded 0 to 1 range
RobustScaler	Median	Interquartile range	Outliers are present, usually the safer default
QuantileTransformer	-	-	You want a specific output shape regardless of input

⚠ Watch Out

MinMaxScaler is destroyed by a single outlier

One customer with 97 calls sets the maximum, so every other customer. The ones with 0 to 5 calls, is compressed into the bottom five percent of the range. The scaler has technically done its job and practically destroyed the column's resolution. **RobustScaler** uses quartiles, which one row cannot move.

Binning

Sometimes the relationship is not smooth. A customer at month 11 and one at month 13 may behave very differently if the contract renews at 12.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()
6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.preprocessing import KBinsDiscretizer
25 from sklearn.linear_model import LogisticRegression
26 from sklearn.compose import ColumnTransformer
27 from sklearn.model_selection import cross_val_score, StratifiedKFold,
    ↪ train_test_split
28 from sklearn.impute import SimpleImputer

```

```

29 from sklearn.preprocessing import OneHotEncoder, StandardScaler
30
31 X, y = df[NUM + CAT], df[TARGET]
32 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
33                                         stratify=y, random_state=42)
34 cv = StratifiedKFold(5, shuffle=True, random_state=0)
35
36 binned = ColumnTransformer([
37     ('bin', Pipeline([('i', SimpleImputer(strategy='median')),
38                     ('b', KBinsDiscretizer(n_bins=8, encode='onehot-dense',
39                     strategy='quantile'))])),
40     ['tenure_months']],
41     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
42                     ('s', StandardScaler())])),
43     ['support_calls', 'monthly_charges']],
44     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
45                     ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
46 ])
47 m = Pipeline([('prep', binned),
48              ('clf', LogisticRegression(max_iter=1000, random_state=42))])
49 s = cross_val_score(m, X_tr, y_tr, cv=cv, scoring='roc_auc')
50 print('tenure binned into 8: %.4f +/- %.4f' % (s.mean(), s.std()))
51 print('baseline:                0.8238')

```

```

tenure binned into 8: 0.8221 +/- 0.0295
baseline:                0.8238

```

⚠ Watch Out

Binning throws information away

Every customer inside a bin becomes identical. You gain the ability to express a non-linear step and you lose all resolution within the bin. It helps when the true relationship really does jump; it hurts when the relationship was smooth, which is the case here. Measure, do not assume.

Splines, which give curvature without losing resolution

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()
6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.preprocessing import SplineTransformer
13 from sklearn.linear_model import LogisticRegression
14 from sklearn.compose import ColumnTransformer

```

```

15 from sklearn.pipeline import Pipeline
16 from sklearn.impute import SimpleImputer
17 from sklearn.preprocessing import OneHotEncoder, StandardScaler
18 from sklearn.model_selection import cross_val_score, StratifiedKFold,
    ↪ train_test_split
19
20 X, y = df[NUM + CAT], df[TARGET]
21 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
22                                           stratify=y, random_state=42)
23 cv = StratifiedKFold(5, shuffle=True, random_state=0)
24
25 spl = ColumnTransformer([
26     ('spline', Pipeline([('i', SimpleImputer(strategy='median')),
27                          ('s', SplineTransformer(n_knots=5, degree=3))])),
28     ['tenure_months']],
29     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
30                      ('s', StandardScaler())])),
31     ['support_calls', 'monthly_charges']],
32     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
33                      ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
34 ])
35 m = Pipeline([('prep', spl),
36               ('clf', LogisticRegression(max_iter=1000, random_state=42))])
37 s = cross_val_score(m, X_tr, y_tr, cv=cv, scoring='roc_auc')
38 print('tenure as a spline: %.4f +/- %.4f' % (s.mean(), s.std()))
39 print('baseline:          0.8238')

```

```

tenure as a spline: 0.8235 +/- 0.0303
baseline:          0.8238

```

A spline is a smooth curve built from overlapping local basis functions. It lets a linear model bend without the staircase that binning imposes, and it is usually the better tool when you suspect curvature.

Day 2 takeaway

Day 3 Encoding High-Cardinality Categories

Target encoding, how it leaks, and what makes it safe

One-hot encoding works until a column has a thousand levels. Then you need target encoding, which is powerful, and which leaks unless you are very careful.

The problem

```

1 import numpy as np
2 import pandas as pd
3
4 rng = np.random.default_rng(0)
5 n = 5000
6 postcodes = ['PC%04d' % i for i in rng.integers(0, 900, n)]

```

```

7 s = pd.Series(postcodes)
8 print('rows          %d' % n)
9 print('distinct values %d' % s.nunique())
10 print('one-hot columns %d' % s.nunique())
11 print('\nlevels seen fewer than 5 times: %d'
12       % (s.value_counts() < 5).sum())

```

```

rows          5000
distinct values 896
one-hot columns 896

levels seen fewer than 5 times: 307

```

The naive version, and why it fails

```

1 import numpy as np
2 import pandas as pd
3 from sklearn.linear_model import LogisticRegression
4 from sklearn.model_selection import cross_val_score, StratifiedKFold
5
6 rng = np.random.default_rng(0)
7 n = 2000
8 # A category with no relationship to the target whatsoever.
9 cat = pd.Series(['C%03d' % i for i in rng.integers(0, 400, n)])
10 y = pd.Series(rng.integers(0, 2, n))
11
12 # Naive: compute the mean target per category using every row.
13 means = y.groupby(cat).mean()
14 encoded = cat.map(means).to_frame('te')
15
16 cv = StratifiedKFold(5, shuffle=True, random_state=0)
17 s = cross_val_score(LogisticRegression(), encoded, y, cv=cv, scoring='roc_auc')
18 print('random category, random target -- truth is 0.5')
19 print('naive target encoding scores:  %.4f' % s.mean())

```

```

random category, random target -- truth is 0.5
naive target encoding scores:  0.7587

```

Watch Out

It encoded the answer into the feature

With 400 categories over 2000 rows, each category has about five rows. Its mean target is therefore mostly determined by those five labels, including the one you are about to predict. The feature is a leaky copy of the target, and cross-validation cannot save you because the leak happened before the folds existed.

The safe version

```

1 import numpy as np

```

```

2 import pandas as pd
3 from sklearn.preprocessing import TargetEncoder
4 from sklearn.linear_model import LogisticRegression
5 from sklearn.pipeline import make_pipeline
6 from sklearn.model_selection import cross_val_score, StratifiedKFold
7
8 rng = np.random.default_rng(0)
9 n = 2000
10 cat = np.array([[f'C%03d' % i] for i in rng.integers(0, 400, n)])
11 y = rng.integers(0, 2, n)
12
13 cv = StratifiedKFold(5, shuffle=True, random_state=0)
14 m = make_pipeline(TargetEncoder(random_state=0), LogisticRegression())
15 s = cross_val_score(m, cat, y, cv=cv, scoring='roc_auc')
16 print('sklearn TargetEncoder, inside a pipeline: %.4f' % s.mean())
17 print('the truth is 0.5 -- and now it says so')

```

```

sklearn TargetEncoder, inside a pipeline: 0.5055
the truth is 0.5 -- and now it says so

```

scikit-learn's `TargetEncoder` does two things that make it safe. It fits an internal cross-fitting scheme, so a row's encoding never uses its own label. And it shrinks each category mean toward the global mean in proportion to how few rows that category has, so a category seen twice barely moves away from the overall rate.

```

1 import numpy as np
2 from sklearn.preprocessing import TargetEncoder
3
4 cat = np.array([[f'common' * 500 + f'rare' * 3]].reshape(-1, 1))
5 y = np.array([1] * 250 + [0] * 250 + [1, 1, 1]) # rare is 100% positive
6
7 print('global mean      %.4f' % y.mean())
8 print('raw mean for rare 1.0000 (3 rows, all positive)\n')
9 print('%10s %12s %10s' % ('smooth', 'common', 'rare'))
10 for smooth in ['auto', 1.0, 5.0, 20.0]:
11     enc = TargetEncoder(random_state=0, cv=5, smooth=smooth).fit(cat, y)
12     print('%10s %12.4f %10.4f'
13           % (smooth, enc.transform([[f'common']])[0, 0],
14             enc.transform([[f'rare']])[0, 0]))

```

```

global mean      0.5030
raw mean for rare 1.0000 (3 rows, all positive)

```

smooth	common	rare
auto	0.5000	1.0000
1.0	0.5000	0.8757
5.0	0.5000	0.6894
20.0	0.5001	0.5678

⚠ Watch Out**smooth='auto' did not shrink it at all**

The default estimates how much to shrink from the *variance within* each category. Our rare group is three rows that are all 1, so its within-group variance is zero, the estimator concludes the mean is perfectly reliable, and it leaves 1.0 untouched. Three rows, and the encoder is certain.

An explicit `smooth` is a fixed number of pseudo-observations at the global mean, so it pulls regardless: 0.88 at 1, 0.69 at 5, 0.57 at 20. When your rare categories are small and internally consistent, which is exactly when they are most dangerous, set `smooth` yourself rather than trusting the automatic estimate.

On our data

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()
6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X, y = df[NUM + CAT], df[TARGET]
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                         stratify=y, random_state=42)
29 from sklearn.preprocessing import TargetEncoder
30 from sklearn.linear_model import LogisticRegression
31 from sklearn.compose import ColumnTransformer
32 from sklearn.impute import SimpleImputer
33 from sklearn.preprocessing import StandardScaler
34 from sklearn.model_selection import cross_val_score, StratifiedKFold
35
36 cv = StratifiedKFold(5, shuffle=True, random_state=0)
37 te = ColumnTransformer([

```

```

38     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
39                       ('s', StandardScaler())]), NUM),
40     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
41                       ('t', TargetEncoder(random_state=0))]), CAT),
42 ])
43 m = Pipeline([('prep', te),
44               ('clf', LogisticRegression(max_iter=1000, random_state=42))])
45 s = cross_val_score(m, X_tr, y_tr, cv=cv, scoring='roc_auc')
46 print('target encoded: %.4f +/- %.4f' % (s.mean(), s.std()))
47 print('one-hot was:      0.8238')

```

```

target encoded: 0.8197 +/- 0.0271
one-hot was:    0.8238

```

Level with one-hot, which is expected: our categories have three or four levels each, so one-hot costs nothing and target encoding has nothing to compress. Reach for it at hundreds of levels, not at four.

Day 3 takeaway

Day 4 Interactions and Domain Features

The thing a linear model cannot do, and how to hand it over

A linear model cannot express “month-to-month customers who also call support a lot”. It can only add the two effects. Handing it the combination is often where the real gain is.

Why a linear model needs help

```

1 import numpy as np
2 from sklearn.linear_model import LogisticRegression
3 from sklearn.preprocessing import StandardScaler
4 from sklearn.pipeline import make_pipeline
5 from sklearn.model_selection import cross_val_score
6
7 rng = np.random.default_rng(0)
8 n = 4000
9 a = rng.integers(0, 2, n)
10 b = rng.integers(0, 2, n)
11 y = (a ^ b) # exclusive or
12 y = np.where(rng.uniform(size=n) < 0.05, 1 - y, y)
13
14 X = np.column_stack([a, b])
15 print('without the interaction %.4f'
16       % cross_val_score(make_pipeline(StandardScaler(), LogisticRegression()),
17                         X, y, cv=5, scoring='roc_auc').mean())
18
19 X2 = np.column_stack([a, b, a * b])
20 print('with a*b as a column %.4f'
21       % cross_val_score(make_pipeline(StandardScaler(), LogisticRegression()),
22                         X2, y, cv=5, scoring='roc_auc').mean())

```

```
without the interaction 0.5073
with a*b as a column    0.9558
```

Chance, then near-perfect, from one extra column. Neither feature alone says anything about the target; only their combination does. A tree finds this by splitting twice; a linear model finds it only if you supply it.

Automatic interactions, and their cost

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()
6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X, y = df[NUM + CAT], df[TARGET]
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                         stratify=y, random_state=42)
29 from sklearn.preprocessing import PolynomialFeatures, StandardScaler, OneHotEncoder
30 from sklearn.compose import ColumnTransformer
31 from sklearn.impute import SimpleImputer
32 from sklearn.linear_model import LogisticRegression
33 from sklearn.model_selection import cross_val_score, StratifiedKFold
34
35 cv = StratifiedKFold(5, shuffle=True, random_state=0)
36 inter = ColumnTransformer([
37     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
38     ('s', StandardScaler())])), NUM),
39     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
40     ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
41 ])
42 m = Pipeline([('prep', inter),
43               ('poly', PolynomialFeatures(degree=2, interaction_only=True,
44                                           include_bias=False)),
```

```

45         ('clf', LogisticRegression(max_iter=3000, random_state=42)))]
46 m.fit(X_tr, y_tr)
47 print('columns after encoding      %d'
48       % m.named_steps['prep'].transform(X_tr).shape[1])
49 print('columns after interactions  %d'
50       % m.named_steps['poly'].transform(
51         m.named_steps['prep'].transform(X_tr)).shape[1])
52 s = cross_val_score(m, X_tr, y_tr, cv=cv, scoring='roc_auc')
53 print('\nall pairwise interactions: %.4f +/- %.4f' % (s.mean(), s.std()))
54 print('baseline:                    0.8238')

```

```

columns after encoding      15
columns after interactions  120

all pairwise interactions: 0.8140 +/- 0.0240
baseline:                  0.8238

```

⚠ Watch Out

Every pair, including the meaningless ones

`interaction_only=True` at degree 2 gives you every pairwise product, and most of them mean nothing, *fibre optic times mailed check* is not a concept. You have multiplied the column count and handed the model far more opportunity to overfit for one or two useful combinations. Pair it with regularisation, or engineer the interactions you can justify.

Deliberate features instead

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()
6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.linear_model import LogisticRegression

```

```

25 from sklearn.model_selection import cross_val_score, StratifiedKFold,
    ↪ train_test_split
26 import numpy as np
27
28 eng = df.copy()
29
30 # 1. On a rolling contract AND unhappy: the combination we can justify.
31 eng['rolling_and_calling'] = (
32     (eng['contract'] == 'Month-To-Month') & (eng['support_calls'] >= 2)
33 ).astype(int)
34
35 # 2. Paying a lot relative to others on the same service.
36 eng['price_vs_service'] = (
37     eng['monthly_charges']
38     - eng.groupby('internet_service')['monthly_charges'].transform('median'))
39
40 # 3. Still new. Churn risk is concentrated in the first months.
41 eng['is_new'] = (eng['tenure_months'] <= 6).astype(int)
42
43 # 4. Calls per year of tenure, not raw calls.
44 eng['call_rate'] = eng['support_calls'] / np.maximum(eng['tenure_months'], 1) * 12
45
46 NEW = NUM + ['rolling_and_calling', 'price_vs_service', 'is_new', 'call_rate']
47 Xe, ye = eng[NEW + CAT], eng[TARGET]
48 Xe_tr, _, ye_tr, _ = train_test_split(Xe, ye, test_size=0.25,
49                                     stratify=ye, random_state=42)
50
51 cv = StratifiedKFold(5, shuffle=True, random_state=0)
52 m = Pipeline([('prep', preprocessor(num=NEW)),
53              ('clf', LogisticRegression(max_iter=1000, random_state=42))])
54 s = cross_val_score(m, Xe_tr, ye_tr, cv=cv, scoring='roc_auc')
55 print('four justified features: %.4f +/- %.4f' % (s.mean(), s.std()))
56 print('baseline: 0.8238')

```

```

four justified features: 0.8267 +/- 0.0246
baseline: 0.8238

```

⚠ Watch Out

price_vs_service is computed over the whole dataset

That group median uses every row, including the ones in the validation fold. It is a small leak and it is still a leak: the score above is very slightly optimistic. Day 6 rebuilds it as a proper transformer that learns its medians on training folds only, which is the correct way to ship a feature like this.

Day 4 takeaway

Day 5 Feature Selection

More features is not better. Each irrelevant column adds noise, slows fitting, and gives the model another chance to find a coincidence.

Watch noise columns do damage

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()
6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20                             ('s', StandardScaler())]), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22                             ('o', OneHotEncoder(handle_unknown='ignore'))]), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X, y = df[NUM + CAT], df[TARGET]
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                           stratify=y, random_state=42)
29 from sklearn.linear_model import LogisticRegression
30 from sklearn.model_selection import cross_val_score, StratifiedKFold
31 import numpy as np
32 import pandas as pd
33
34 cv = StratifiedKFold(5, shuffle=True, random_state=0)
35 rng = np.random.default_rng(0)
36
37 for extra in [0, 10, 50, 200]:
38     Xn = X_tr.copy()
39     for j in range(extra):
40         Xn['noise_%d' % j] = rng.normal(size=len(Xn))
41     num = NUM + ['noise_%d' % j for j in range(extra)]
42     m = Pipeline([('prep', preprocessor(num=num)),
43                  ('clf', LogisticRegression(max_iter=2000, random_state=42))])
44     s = cross_val_score(m, Xn, y_tr, cv=cv, scoring='roc_auc')
45     print('%4d noise columns: %.4f +/- %.4f' % (extra, s.mean(), s.std()))
```

```

0 noise columns: 0.8238 +/- 0.0305
10 noise columns: 0.8194 +/- 0.0318
50 noise columns: 0.8053 +/- 0.0342
200 noise columns: 0.7910 +/- 0.0275

```

Two thousand rows and two hundred pure-noise columns costs real performance. The model spends its capacity fitting coincidences, and regularisation only partly protects it.

Three families of method

Family	How	Cost	Weakness
Filter	Score each feature against the target alone	Very cheap	Blind to combinations
Wrapper	Repeatedly fit the model with subsets	Expensive	Can overfit the selection itself
Embedded	The model selects while fitting (lasso, trees)	Free	Tied to that model's assumptions

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()
6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X, y = df[NUM + CAT], df[TARGET]
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                         stratify=y, random_state=42)
29 from sklearn.feature_selection import (SelectKBest, mutual_info_classif,
30                                     RFE, SelectFromModel)
31 from sklearn.linear_model import LogisticRegression
32 from sklearn.ensemble import RandomForestClassifier
33 from sklearn.model_selection import cross_val_score, StratifiedKFold
34
35 cv = StratifiedKFold(5, shuffle=True, random_state=0)

```

```

36 prep = preprocessor()
37
38 selectors = {
39     'none': None,
40     'filter, k=8': SelectKBest(mutual_info_classif, k=8),
41     'wrapper, RFE 8': RFE(LogisticRegression(max_iter=1000),
42         ↪ n_features_to_select=8),
43     'embedded, L1': SelectFromModel(
44         LogisticRegression(penalty='l1', solver='liblinear', C=0.1)),
45 }
46 for name, sel in selectors.items():
47     steps = [('prep', prep)]
48     if sel is not None:
49         steps.append(('sel', sel))
50     steps.append(('clf', LogisticRegression(max_iter=1000, random_state=42)))
51     s = cross_val_score(Pipeline(steps), X_tr, y_tr, cv=cv, scoring='roc_auc')
52     print('%s %.4f +/- %.4f' % (name, s.mean(), s.std()))

```

```

none          0.8238 +/- 0.0305
filter, k=8    0.8184 +/- 0.0250
wrapper, RFE 8 0.8229 +/- 0.0313
embedded, L1   0.8234 +/- 0.0299

```

Selection lives inside the pipeline

Every selector above is a pipeline step, so it is refitted on each training fold and never sees the fold it is scored on. Week 7 showed what happens otherwise: selection over all the labels made pure noise score 0.80. There is no version of feature selection that is safe to do before splitting.

What filter methods miss

```

1 import numpy as np
2 from sklearn.feature_selection import mutual_info_classif, f_classif
3
4 rng = np.random.default_rng(0)
5 n = 4000
6 a = rng.integers(0, 2, n)
7 b = rng.integers(0, 2, n)
8 y = a ^ b
9 X = np.column_stack([a, b, rng.normal(size=n)])
10
11 print('mutual information with the target, feature by feature:')
12 for name, score in zip(['a', 'b', 'noise'], mutual_info_classif(X, y,
13     ↪ random_state=0)):
14     print(' %-6s %.5f' % (name, score))
15 print('\na and b together determine y exactly. Individually both')
16 print('score near zero -- a is indistinguishable from pure noise.')

```

```

mutual information with the target, feature by feature:
a          0.00000
b          0.01169
noise      0.00000

```

```
a and b together determine y exactly. Individually both
score near zero -- a is indistinguishable from pure noise.
```

Any filter method evaluates one feature at a time, so it cannot see that two useless-looking columns are jointly decisive. If interactions matter in your problem, a filter will discard exactly the features you need.

Day 5 takeaway

Day 6 Writing Your Own Transformer

Turning a leaky groupby into a pipeline step that fits per fold

Day 4 built a feature from a group median computed over the whole dataset, and flagged it as a leak. Today you fix it properly, by writing a transformer that obeys the scikit-learn contract.

The contract, from week 1

A transformer needs `fit`, which learns and returns `self`, and `transform`, which applies what was learned. Inherit two base classes and you get `fit_transform`, `get_params` and pipeline compatibility for nothing.

```
1 import numpy as np
2 import pandas as pd
3 from sklearn.base import BaseEstimator, TransformerMixin
4
5 class GroupRelativeValue(BaseEstimator, TransformerMixin):
6     """How far a customer's value sits from the median of their group.
7
8     The medians are learned in fit, from training rows only, so this is
9     safe to use inside cross-validation -- which the day 4 version, built
10    with a groupby over the whole frame, was not.
11    """
12
13    def __init__(self, group_col, value_col):
14        self.group_col = group_col
15        self.value_col = value_col
16
17    def fit(self, X, y=None):
18        self.medians_ = X.groupby(self.group_col)[self.value_col].median()
19        self.global_median_ = X[self.value_col].median()
20        return self
21
22    def transform(self, X):
23        # An unseen group falls back to the global median rather than NaN.
24        base = X[self.group_col].map(self.medians_).fillna(self.global_median_)
25        return (X[self.value_col] - base).to_frame('relative_value')
26
27    def get_feature_names_out(self, input_features=None):
28        return np.array(['relative_value'])
```

```

29
30
31 demo = pd.DataFrame({'internet_service': ['DSL', 'DSL', 'Fibre optic', 'Fibre
    ↪ optic'],
32                        'monthly_charges': [50.0, 60.0, 75.0, 85.0]})
33 t = GroupRelativeValue('internet_service', 'monthly_charges').fit(demo)
34 print('learned medians:')
35 print(t.medians_.to_string())
36 print('\ntransformed:')
37 print(t.transform(demo).to_string(index=False))
38
39 unseen = pd.DataFrame({'internet_service': ['Satellite'], 'monthly_charges':
    ↪ [90.0]})
40 print('\nunseen group falls back to the global median:')
41 print(t.transform(unseen).to_string(index=False))

```

```

learned medians:
internet_service
DSL              55.0
Fibre optic      80.0

transformed:
relative_value
              -5.0
              5.0
              -5.0
              5.0

unseen group falls back to the global median:
relative_value
              22.5

```

Put it in a pipeline

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()
6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median'))],

```

```

20         ('s', StandardScaler()))], num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22                           ('o', OneHotEncoder(handle_unknown='ignore'))]), cat),
23     ])
24 from sklearn.model_selection import train_test_split
25
26 X, y = df[NUM + CAT], df[TARGET]
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                           stratify=y, random_state=42)
29 from sklearn.base import BaseEstimator, TransformerMixin
30 from sklearn.compose import ColumnTransformer
31 from sklearn.impute import SimpleImputer
32 from sklearn.preprocessing import StandardScaler, OneHotEncoder
33 from sklearn.linear_model import LogisticRegression
34 from sklearn.model_selection import cross_val_score, StratifiedKFold
35 import numpy as np
36
37 class GroupRelativeValue(BaseEstimator, TransformerMixin):
38     def __init__(self, group_col, value_col):
39         self.group_col = group_col
40         self.value_col = value_col
41
42     def fit(self, X, y=None):
43         self.medians_ = X.groupby(self.group_col)[self.value_col].median()
44         self.global_median_ = X[self.value_col].median()
45         return self
46
47     def transform(self, X):
48         base = X[self.group_col].map(self.medians_).fillna(self.global_median_)
49         return (X[self.value_col] - base).fillna(0).to_frame('relative_value')
50
51     def get_feature_names_out(self, input_features=None):
52         return np.array(['relative_value'])
53
54
55 prep = ColumnTransformer([
56     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
57                     ('s', StandardScaler())]), NUM),
58     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
59                     ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
60     ('rel', Pipeline([('g', GroupRelativeValue('internet_service',
61                                               'monthly_charges')),
62                     ('s', StandardScaler())]),
63     ['internet_service', 'monthly_charges']),
64 ])
65
66 cv = StratifiedKFold(5, shuffle=True, random_state=0)
67 m = Pipeline([('prep', prep),
68              ('clf', LogisticRegression(max_iter=1000, random_state=42))])
69 s = cross_val_score(m, X_tr, y_tr, cv=cv, scoring='roc_auc')
70 print('leak-free relative price: %.4f +/- %.4f' % (s.mean(), s.std()))
71 print('baseline: 0.8238')

```

```
leak-free relative price: 0.8239 +/- 0.0308
```

baseline: 0.8238

The medians are refitted on every fold

Five folds, five different sets of medians, each learned from four fifths of the data and applied to the fifth. No validation row contributed to the statistic used to encode it. That is the whole point of writing the transformer rather than doing a `groupby` in your notebook.

FunctionTransformer, for anything stateless

```
1 import numpy as np
2 import pandas as pd
3 from sklearn.preprocessing import FunctionTransformer
4
5 def add_ratios(X):
6     out = X.copy()
7     out['calls_per_year'] = (out['support_calls']
8                             / np.maximum(out['tenure_months'], 1) * 12)
9     out['log_calls'] = np.log1p(out['support_calls'])
10    return out
11
12 # No feature_names_out='one-to-one' here: that promises the same columns
13 # out as in, and this function adds two. sklearn checks, and raises.
14 step = FunctionTransformer(add_ratios)
15 demo = pd.DataFrame({'support_calls': [0, 3, 12], 'tenure_months': [2, 24, 6]})
16 print(step.fit_transform(demo).round(3).to_string(index=False))
```

support_calls	tenure_months	calls_per_year	log_calls
0	2	0.0	0.000
3	24	1.5	1.386
12	6	24.0	2.565

Watch Out

Only when the function learns nothing

`FunctionTransformer` has no `fit` worth the name, so it is safe exactly when your function uses no statistic from the data. A ratio of two columns in the same row is fine. Anything involving a mean, a median, a category list or a bin edge needs a real transformer with a `fit`.

Day 6 takeaway

Day 7 A Complete Feature Engineering Project

Every idea measured against the baseline, failures included

Everything this week, applied properly, measured honestly, against the 0.8238 that a week of

ensembles could not beat.

The candidate features

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()
6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.base import BaseEstimator, TransformerMixin
25 from sklearn.compose import ColumnTransformer
26 from sklearn.impute import SimpleImputer
27 from sklearn.preprocessing import (StandardScaler, OneHotEncoder,
28                                     SplineTransformer, FunctionTransformer)
29 from sklearn.linear_model import LogisticRegression
30 from sklearn.model_selection import (cross_val_score, StratifiedKFold,
31                                     train_test_split)
32 import numpy as np
33
34 def enrich(X):
35     out = X.copy()
36     out['rolling_and_calling'] = (
37         (out['contract'] == 'Month-To-Month') & (out['support_calls'] >= 2)
38     ).astype(int)
39     out['is_new'] = (out['tenure_months'] <= 6).astype(int)
40     out['call_rate'] = (out['support_calls']
41                        / np.maximum(out['tenure_months'], 1) * 12)
42     return out
43
44 X, y = df[NUM + CAT], df[TARGET]
45 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
46                                           stratify=y, random_state=42)
47
48 sample = enrich(X_tr.head(4))
49 print(sample[['contract', 'support_calls', 'tenure_months',

```

```

50         'rolling_and_calling', 'is_new', 'call_rate']]
51     .round(2).to_string(index=False))

```

	contract	support_calls	tenure_months	rolling_and_calling	is_new	call_rate
	Month-To-Month	0	9	0	0	0.0
	Two Year	0	21	0	0	0.0
	Month-To-Month	0	2	0	1	0.0
	Month-To-Month	0	12	0	0	0.0

Assemble and measure

```

1  import numpy as np
2  import pandas as pd
3
4  df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5  df = df.drop_duplicates().copy()
6  df['contract'] = df['contract'].str.strip().str.title()
7  df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9  NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16
17 def preprocessor(num=NUM, cat=CAT):
18     return ColumnTransformer([
19         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
20         ('s', StandardScaler())])), num),
21         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
22         ('o', OneHotEncoder(handle_unknown='ignore'))])), cat),
23     ])
24 from sklearn.base import BaseEstimator, TransformerMixin
25 from sklearn.compose import ColumnTransformer
26 from sklearn.impute import SimpleImputer
27 from sklearn.preprocessing import (StandardScaler, OneHotEncoder,
28                                     SplineTransformer, FunctionTransformer)
29 from sklearn.linear_model import LogisticRegression
30 from sklearn.model_selection import (cross_val_score, StratifiedKFold,
31                                     train_test_split)
32 import numpy as np
33
34 class GroupRelativeValue(BaseEstimator, TransformerMixin):
35     def __init__(self, group_col, value_col):
36         self.group_col = group_col
37         self.value_col = value_col
38
39     def fit(self, X, y=None):
40         self.medians_ = X.groupby(self.group_col)[self.value_col].median()
41         self.global_median_ = X[self.value_col].median()

```

```

42         return self
43
44     def transform(self, X):
45         base = X[self.group_col].map(self.medians_).fillna(self.global_median_)
46         return (X[self.value_col] - base).fillna(0).to_frame('rel')
47
48     def get_feature_names_out(self, input_features=None):
49         return np.array(['rel'])
50
51
52 def enrich(X):
53     out = X.copy()
54     out['rolling_and_calling'] = (
55         (out['contract'] == 'Month-To-Month') & (out['support_calls'] >= 2)
56     ).astype(int)
57     out['is_new'] = (out['tenure_months'] <= 6).astype(int)
58     out['call_rate'] = (out['support_calls']
59                        / np.maximum(out['tenure_months'], 1) * 12)
60     return out
61
62
63 DERIVED = ['rolling_and_calling', 'is_new', 'call_rate']
64
65 prep = ColumnTransformer([
66     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
67                      ('s', StandardScaler())])), NUM),
68     ('derived', Pipeline([('i', SimpleImputer(strategy='median')),
69                          ('s', StandardScaler())])), DERIVED),
70     ('spline', Pipeline([('i', SimpleImputer(strategy='median')),
71                          ('s', SplineTransformer(n_knots=5, degree=3))])),
72     ['tenure_months']),
73     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
74                      ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
75     ('rel', Pipeline([('g', GroupRelativeValue('internet_service',
76                                                'monthly_charges')),
77                      ('s', StandardScaler())])),
78     ['internet_service', 'monthly_charges']],
79 ])
80
81 X, y = df[NUM + CAT], df[TARGET]
82 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
83                                           stratify=y, random_state=42)
84 cv = StratifiedKFold(5, shuffle=True, random_state=0)
85
86 full = Pipeline([
87     ('enrich', FunctionTransformer(enrich)),
88     ('prep', prep),
89     ('clf', LogisticRegression(max_iter=3000, random_state=42)),
90 ])
91 s = cross_val_score(full, X_tr, y_tr, cv=cv, scoring='roc_auc')
92 print('engineered: %.4f +/- %.4f' % (s.mean(), s.std()))
93 print('baseline: 0.8238')
94 print('gain:      %.4f (fold spread %.4f)' % (s.mean() - 0.8238, s.std()))

```

```
engineered: 0.8259 +/- 0.0252
baseline:   0.8238
gain:       +0.0021 (fold spread 0.0252)
```

Does it help the ensembles too?

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv', parse_dates=['signup_date'])
5 df = df.drop_duplicates().copy()
6 df['contract'] = df['contract'].str.strip().str.title()
7 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
8
9 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
10 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
11 TARGET = 'churned'
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder, FunctionTransformer
16 from sklearn.linear_model import LogisticRegression
17 from sklearn.ensemble import HistGradientBoostingClassifier
18 from sklearn.model_selection import (cross_val_score, StratifiedKFold,
19                                     train_test_split)
20 import numpy as np
21
22 def enrich(X):
23     out = X.copy()
24     out['rolling_and_calling'] = (
25         (out['contract'] == 'Month-To-Month') & (out['support_calls'] >= 2)
26     ).astype(int)
27     out['is_new'] = (out['tenure_months'] <= 6).astype(int)
28     out['call_rate'] = (out['support_calls']
29                        / np.maximum(out['tenure_months'], 1) * 12)
30     return out
31
32 DERIVED = ['rolling_and_calling', 'is_new', 'call_rate']
33 X, y = df[NUM + CAT], df[TARGET]
34 X_tr, _, y_tr, _ = train_test_split(X, y, test_size=0.25, stratify=y,
35                                     random_state=42)
36 cv = StratifiedKFold(5, shuffle=True, random_state=0)
37
38 def build(num, clf):
39     prep = ColumnTransformer([
40         ('num', Pipeline([('i', SimpleImputer(strategy='median')),
41                          ('s', StandardScaler())])), num),
42         ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
43                          ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT])
44     return Pipeline([('enrich', FunctionTransformer(enrich)),
45                     ('prep', prep), ('clf', clf)])
46
47 for name, clf in [('logistic', LogisticRegression(max_iter=3000, random_state=42)),
```

```

48         ('boosting', HistGradientBoostingClassifier(random_state=42))]:
49     plain = cross_val_score(build(NUM, clf), X_tr, y_tr, cv=cv,
50                             scoring='roc_auc').mean()
51     rich = cross_val_score(build(NUM + DERIVED, clf), X_tr, y_tr, cv=cv,
52                             scoring='roc_auc').mean()
53     print('%-9s plain %.4f   engineered %.4f   change %+.4f'
54           % (name, plain, rich, rich - plain))

```

```

logistic plain 0.8238   engineered 0.8268   change +0.0030
boosting plain 0.7966   engineered 0.7958   change -0.0009

```

The features help the linear model more than the trees

That is exactly what you should expect. `is_new` and `rolling_and_calling` are thresholds and interactions, and a tree can already discover those by splitting. Handing them to a linear model gives it something it genuinely could not express. Feature engineering is most valuable for the models with the least flexibility.

The honest verdict

The gains this week are small. That is worth saying plainly, because most material on feature engineering implies otherwise. On this dataset the generator used a straightforward weighted sum, so there is not much structure left for a clever feature to expose, and the honest result of measuring carefully is often that an idea did not work.

Watch Out

Report the failures, and be strict about what counts as a win

The engineered pipeline gains about half a point of AUC, and the fold-to-fold spread is five times that. By the standard set on day 1, keep it only if it beats the baseline by more than the spread. This week has not produced a result you could defend. Say so.

You measured roughly a dozen ideas and most changed nothing. That is a normal, healthy ratio. The discipline is not in producing features that work; it is in measuring against a fixed baseline on fixed folds so you can tell the difference, and in discarding what does not earn its place rather than shipping it because you built it.

Your assignment

The dataset generator used `payment_method == 'Electronic check'` with a coefficient of +0.42. Build a feature that captures an interaction between payment method and something else, and measure it. Then try a feature the generator definitely did not use, a ratio of charges to support calls, say. Record both results, including the one that fails.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. Before engineering any feature you should first:
 - a) Remove outliers
 - b) Tune the model
 - c) Establish a baseline score to measure every later change against
 - d) Scale everything
2. Binning or splining a numeric feature buys a linear model:
 - a) Robustness to outliers
 - b) Faster fitting
 - c) Fewer parameters
 - d) Curvature, the ability to respond non-linearly to that feature
3. Target encoding leaks when:
 - a) The encoding is computed using the target values of the rows it is applied to
 - b) The target is continuous
 - c) There are too many categories
 - d) The category is rare
4. The safe way to apply target encoding is:
 - a) Compute it once on all the data
 - b) Fit it inside the pipeline so each fold's encoding is learned from that fold's training rows, with smoothing
 - c) Use the median instead
 - d) Drop rare categories first
5. An interaction feature is needed when:
 - a) The target is skewed
 - b) Two features are correlated
 - c) The effect of one feature depends on the value of another, something a linear model cannot represent unaided
 - d) There are missing values
6. A filter method for feature selection cannot detect:
 - a) Constant columns
 - b) Correlated pairs
 - c) Strong single predictors
 - d) A feature that is useless alone but informative in combination with another
7. Writing a custom transformer rather than a pandas `groupby` matters because:
 - a) The transformer refits per fold, so the statistic it computes never sees validation rows

- b) It handles missing data
 - c) scikit-learn requires it
 - d) It is faster
8. A feature engineering project should report:
- a) Only the ideas that worked
 - b) Every idea measured against the baseline, including the ones that made it worse
 - c) The final model only
 - d) The best cross-validation score

Answers: 1c, 2d, 3a, 4b, 5c, 6d, 7a, 8b. Anything you got wrong points at the day to re-read, not at a gap in ability.

Unsupervised Learning and Clustering

Week 9

Finding structure without labels, k-means, hierarchical, DBSCAN, and mixture, plus how to tell whether the structure is real.

OBJECTIVES

With no target column there is no score to check against, which makes it unusually easy to find groups that are not there. This week covers the four main clustering families and how each fails, the distance problem that categorical columns create, and above all the validation, silhouette, stability and holding back the outcome. That separates a real segmentation from a partition of a cloud.

Day 1 k-means From Scratch

The algorithm in four steps, and the three assumptions it makes silently

Eight weeks of supervised learning: you had the answers and searched for the rule. Now the answers are gone. Nobody labelled these customers, and the question is whether they fall into groups at all.

k-means, in full

1. Place k centres at random.
2. Assign every point to its nearest centre.
3. Move each centre to the mean of the points assigned to it.
4. Repeat 2 and 3 until nothing moves.

```

1 import numpy as np
2
3 rng = np.random.default_rng(0)
4 X = np.vstack([rng.normal([0, 0], 0.6, (60, 2)),
5               rng.normal([4, 4], 0.6, (60, 2)),
6               rng.normal([0, 5], 0.6, (60, 2))])
7
8 centres = X[rng.choice(len(X), 3, replace=False)]
9 for step in range(1, 11):
10     d = ((X[:, None] - centres[None, :]) ** 2).sum(axis=2)
11     labels = d.argmin(axis=1)
12     new = np.array([X[labels == k].mean(axis=0) for k in range(3)])
13     shift = np.abs(new - centres).max()
```

```

14     centres = new
15     if step <= 3 or shift < 1e-6:
16         print('step %2d largest centre movement %.6f' % (step, shift))
17     if shift < 1e-6:
18         break
19
20 print('\nfinal centres:')
21 print(np.round(centres[centres[:, 0].argsort()], 2))

```

```

step 1 largest centre movement 2.336408
step 2 largest centre movement 1.994041
step 3 largest centre movement 0.043246
step 4 largest centre movement 0.000000

```

```

final centres:
[[-0.    4.94]
 [-0.    0.1 ]
 [ 3.83  4.02]]

```

Converged in a handful of passes, and the centres landed on the three clusters the data was built from. Twelve lines, no library.

The library version, and the argument that matters

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.cluster import KMeans
17 import numpy as np
18
19 km = KMeans(n_clusters=4, n_init=10, random_state=42).fit(Z)
20 print('cluster sizes:', np.bincount(km.labels_))
21 print('inertia: %.1f' % km.inertia_)
22 print('iterations to converge:', km.n_iter_)
23
24 tiny = np.bincount(km.labels_).argmin()
25 print('\nthe smallest cluster, in full:')
26 print(df.loc[km.labels_ == tiny,
27             ['customer_id', 'tenure_months', 'support_calls',
28             'monthly_charges']].to_string(index=False))

```

```
cluster sizes: [1641  822  536    1]
inertia: 3012.8
iterations to converge: 8

the smallest cluster, in full:
customer_id  tenure_months  support_calls  monthly_charges
C00008              48              97              76.42
```

⚠ Watch Out

One of your four segments has a single member

It is the customer with 97 support calls from week 2. k-means partitions *everything*, so a point far from all the others gets its own centre rather than being flagged as unusual, and you have spent a quarter of your segmentation on one person. This is the “every point belongs somewhere” assumption failing in the most visible way possible, and it is why day 4’s DBSCAN, which can label a point as noise, matters.

⚠ Watch Out

n_init exists because k-means gets stuck

The result depends on where the centres started. A bad start converges to a bad local optimum, and the algorithm has no way to notice. `n_init=10` runs the whole thing ten times from different starts and keeps the lowest inertia. Setting `n_init=1` to save time is how you get a different segmentation every Monday.

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.cluster import KMeans
17 import numpy as np
18
19 print('%8s %14s %14s' % ('n_init', 'best inertia', 'worst inertia'))
20 for n_init in [1, 5, 25]:
21     results = [KMeans(n_clusters=5, n_init=n_init, random_state=s).fit(Z).inertia_
22               for s in range(8)]
23     print('%8d %14.1f %14.1f' % (n_init, min(results), max(results)))
```

```
n_init  best inertia  worst inertia
```

1	2498.7	2628.4
5	2498.7	2498.7
25	2498.7	2498.7

Scaling decides the answer

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.cluster import KMeans
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13 import numpy as np
14
15 PAIR = ['monthly_charges', 'total_charges']
16 raw = SimpleImputer(strategy='median').fit_transform(df[PAIR])
17 scaler = StandardScaler().fit(raw)
18
19 for name, data in [('unscaled', raw), ('scaled ', scaler.transform(raw))]:
20     centres = KMeans(n_clusters=3, n_init=10,
21 ↪ random_state=0).fit(data).cluster_centers_
22     if name.strip() == 'scaled':
23         centres = scaler.inverse_transform(centres)
24     order = centres[:, 1].argsort()
25     print('%s centres (monthly, total):' % name)
26     print(np.round(centres[order], 1))

```

```

unscaled centres (monthly, total):
[[ 53.1 665.4]
 [ 69.5 2118.9]
 [ 75.9 4450.9]]
scaled centres (monthly, total):
[[ 28.9 590.1]
 [ 70.3 1101.6]
 [ 75.2 3673. ]]

```

Total charges runs to 7,600 and monthly charges to 106, so unscaled the distance is almost entirely about the total. The three clusters differ hugely in total charges, 665, 2,119, 4,451, and barely at all in monthly charges, 53 to 76. Monthly charges has been ignored.

Scaled, the first cluster's monthly charge drops to 29, which is the no-internet group finally being allowed to separate itself. Neither answer is wrong arithmetic; only one of them is the question you meant to ask.

What k-means assumes

Assumption	Consequence when false
Clusters are roughly spherical	Elongated groups get cut in half
Clusters are similar in size	Large ones get split, small ones absorbed
Every point belongs somewhere	Outliers drag centres toward themselves
You know k	You do not, and the algorithm will not tell you

```
1 import numpy as np
2 from sklearn.cluster import KMeans
3 from sklearn.datasets import make_moons
4 from sklearn.metrics import adjusted_rand_score
5
6 X, truth = make_moons(n_samples=400, noise=0.06, random_state=0)
7 labels = KMeans(n_clusters=2, n_init=10, random_state=0).fit_predict(X)
8 print('two interleaved crescents, k-means agreement with the truth: %.3f'
9       % adjusted_rand_score(truth, labels))
10 print('(1.0 is perfect, 0.0 is random)')
```

```
two interleaved crescents, k-means agreement with the truth: 0.274
(1.0 is perfect, 0.0 is random)
```

k-means fails completely here, because it can only draw straight boundaries between centres and a crescent is not a ball. Day 4's DBSCAN solves exactly this.

Day 1 takeaway

Day 2 Choosing k, and Whether There Are Clusters At All

Why inertia cannot answer it, and what silhouette and stability can

k-means needs you to supply k, and the whole point of clustering is that you do not know it. There are three honest answers and one dishonest one.

The dishonest one: inertia

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
```

```

14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.cluster import KMeans
17
18 print('%4s %14s' % ('k', 'inertia'))
19 for k in range(1, 11):
20     km = KMeans(n_clusters=k, n_init=10, random_state=42).fit(Z)
21     print('%4d %14.1f' % (k, km.inertia_))

```

k	inertia
1	9000.0
2	6818.7
3	4748.4
4	3012.8
5	2498.7
6	2151.8
7	1888.1
8	1659.7
9	1496.8
10	1377.3

⚠ Watch Out

Inertia always falls, so it can never choose k

Adding a centre can only reduce the distance from points to their nearest centre. At k equal to the number of rows, inertia is zero and the clustering is useless. The “elbow method” asks you to spot where the curve bends, which on real data is usually a smooth arc with no bend at all. Look at the numbers above and try to name the elbow.

Silhouette: how well separated are the clusters

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.cluster import KMeans
17 from sklearn.metrics import silhouette_score, davies_bouldin_score
18
19 print('%4s %14s %18s' % ('k', 'silhouette', 'davies-bouldin'))
20 for k in range(2, 9):

```

```

21 labels = KMeans(n_clusters=k, n_init=10, random_state=42).fit_predict(Z)
22 print('%4d %14.4f %18.4f'
23       % (k, silhouette_score(Z, labels), davies_bouldin_score(Z, labels)))

```

k	silhouette	davies-bouldin
2	0.3201	1.2820
3	0.3225	0.8410
4	0.3706	0.7130
5	0.2970	0.8996
6	0.2970	0.9680
7	0.2951	0.9170
8	0.2965	0.9071

Higher silhouette is better; lower Davies-Bouldin is better. Read them together, and treat agreement between two different criteria as weak evidence rather than proof.

⚠ Watch Out

A low silhouette everywhere means there are no clusters

Which is a perfectly respectable finding, and the one most people refuse to accept. Our customers were generated from continuous distributions with no group structure at all, so any k-means result here is a partition of a cloud, not a discovery of groups. Partitions can still be useful for targeting, but calling them “segments we found” is a claim the data does not support.

Check against data that genuinely has clusters

```

1 import numpy as np
2 from sklearn.cluster import KMeans
3 from sklearn.datasets import make_blobs
4 from sklearn.metrics import silhouette_score
5
6 X, _ = make_blobs(n_samples=900, centers=4, cluster_std=0.8,
7                  n_features=3, random_state=0)
8 print('data built from four real clusters:')
9 for k in range(2, 8):
10     labels = KMeans(n_clusters=k, n_init=10, random_state=0).fit_predict(X)
11     print(' k=%d silhouette %.4f' % (k, silhouette_score(X, labels)))

```

```

data built from four real clusters:
k=2 silhouette 0.5886
k=3 silhouette 0.6695
k=4 silhouette 0.6581
k=5 silhouette 0.5398
k=6 silhouette 0.4460
k=7 silhouette 0.3641

```

A clear peak at the true k, and values far above anything the customer data produced. That contrast is the calibration you need: this is what real structure looks like.

Stability: would you get the same answer next month?

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.cluster import KMeans
17 from sklearn.metrics import adjusted_rand_score
18 import numpy as np
19
20 rng = np.random.default_rng(0)
21 print('%4s %22s' % ('k', 'agreement across samples'))
22 for k in [2, 3, 4, 6]:
23     scores = []
24     for _ in range(6):
25         idx_a = rng.choice(len(Z), int(0.8 * len(Z)), replace=False)
26         idx_b = rng.choice(len(Z), int(0.8 * len(Z)), replace=False)
27         shared = np.intersect1d(idx_a, idx_b)
28         la = KMeans(k, n_init=10, random_state=0).fit(Z[idx_a])
29         lb = KMeans(k, n_init=10, random_state=0).fit(Z[idx_b])
30         scores.append(adjusted_rand_score(la.predict(Z[shared]),
31                                           lb.predict(Z[shared])))
32     print('%4d %22.4f' % (k, np.mean(scores)))

```

k	agreement across samples
2	0.8236
3	0.8786
4	0.7375
6	0.7785

Stability is the most practical criterion

Cluster two overlapping samples and see whether the same customers end up together. A segmentation that reshuffles when you add a month of data cannot be the basis of a marketing programme, whatever its silhouette says. This test costs nothing and almost nobody runs it.

Day 2 takeaway

Day 3 Hierarchical Clustering

Every k at once, and how linkage decides the shape of what you find

k-means makes you choose k in advance. Hierarchical clustering builds the whole nested family of solutions and lets you cut it wherever you like.

How it merges

```

1 import numpy as np
2 from scipy.cluster.hierarchy import linkage, dendrogram, fcluster
3
4 rng = np.random.default_rng(0)
5 X = np.vstack([rng.normal([0, 0], 0.4, (12, 2)),
6               rng.normal([3, 3], 0.4, (12, 2))])
7
8 Z1 = linkage(X, method='ward')
9 print('the last five merges (cluster a, cluster b, distance, size):')
10 print(np.round(Z1[-5:], 3))
11
12 for k in [2, 3, 4]:
13     labels = fcluster(Z1, k, criterion='maxclust')
14     print('cut into %d: sizes %s' % (k, np.bincount(labels)[1:]))

```

```

the last five merges (cluster a, cluster b, distance, size):
[[20.    41.    1.004  9.   ]
 [38.    39.    1.054  9.   ]
 [40.    43.    1.557 12.   ]
 [36.    42.    1.687 12.   ]
 [44.    45.   15.417 24.   ]]
cut into 2: sizes [12 12]
cut into 3: sizes [12  3  9]
cut into 4: sizes [3 9 3 9]

```

Look at the distance column: the final merge happens at a far larger distance than the ones before it. That jump is the signature of two genuinely separate groups being forced together, and it is how you read a dendrogram.

Linkage decides the shape of what you find

Linkage	Distance between clusters	Tends to produce
ward	Increase in within-cluster variance	Compact, similar-sized clusters
complete	Furthest pair	Compact, sensitive to outliers
average	Mean over all pairs	A compromise
single	Closest pair	Long chains, can follow shapes, often produces one giant blob

```

1 import numpy as np
2 from sklearn.cluster import AgglomerativeClustering
3 from sklearn.datasets import make_moons
4 from sklearn.metrics import adjusted_rand_score
5
6 X, truth = make_moons(n_samples=400, noise=0.06, random_state=0)
7 for method in ['ward', 'complete', 'average', 'single']:
8     labels = AgglomerativeClustering(n_clusters=2, linkage=method).fit_predict(X)

```

```

9     print('%-9s agreement with the truth %.3f'
10          % (method, adjusted_rand_score(truth, labels)))

```

```

ward      agreement with the truth 0.239
complete  agreement with the truth 0.285
average   agreement with the truth 0.329
single    agreement with the truth 1.000

```

Single linkage solves the crescents

Because it only asks whether two clusters have *any* close pair, it can follow a curved shape point by point along its length. Ward and complete linkage insist on compactness and cut the crescents in half, exactly as k-means did. The same chaining behaviour makes single linkage fragile on noisy data, where one bridging point merges two real clusters.

On the customers

```

1  import numpy as np
2  import pandas as pd
3
4  df = pd.read_csv('customers.csv').drop_duplicates().copy()
5  df['contract'] = df['contract'].str.strip().str.title()
6  df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8  NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9  CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.cluster import AgglomerativeClustering, KMeans
17 from sklearn.metrics import silhouette_score, adjusted_rand_score
18 import numpy as np
19
20 sample = Z[np.random.default_rng(0).choice(len(Z), 1200, replace=False)]
21
22 for k in [3, 4, 5]:
23     ward = AgglomerativeClustering(n_clusters=k, linkage='ward').fit_predict(sample)
24     km = KMeans(k, n_init=10, random_state=0).fit_predict(sample)
25     print('k=%d  ward silhouette %.4f  kmeans %.4f  they agree %.3f'
26          % (k, silhouette_score(sample, ward), silhouette_score(sample, km),
27             adjusted_rand_score(ward, km)))

```

```

k=3  ward silhouette 0.3314  kmeans 0.3173  they agree 0.452
k=4  ward silhouette 0.2958  kmeans 0.3727  they agree 0.400
k=5  ward silhouette 0.2636  kmeans 0.3017  they agree 0.371

```

⚠ Watch Out

It does not scale

Agglomerative clustering computes distances between every pair, which is quadratic in memory and worse in time. Three thousand rows is fine; a million is impossible. That is why the snippet above takes a sample, and why k-means remains the default for large data despite its assumptions.

Day 3 takeaway

Day 4 DBSCAN and Density

Clusters as dense regions, arbitrary shapes, and outliers named as noise

k-means partitions everything, including the points that belong to no group. DBSCAN starts from a different idea: a cluster is a dense region, and everything else is noise.

It solves the crescents

```
1 import numpy as np
2 from sklearn.cluster import DBSCAN, KMeans
3 from sklearn.datasets import make_moons
4 from sklearn.metrics import adjusted_rand_score
5
6 X, truth = make_moons(n_samples=400, noise=0.06, random_state=0)
7
8 km = KMeans(2, n_init=10, random_state=0).fit_predict(X)
9 db = DBSCAN(eps=0.25, min_samples=5).fit_predict(X)
10 print('k-means agreement %.3f' % adjusted_rand_score(truth, km))
11 print('DBSCAN agreement %.3f' % adjusted_rand_score(truth, db))
12 print('clusters found: %d, noise points: %d'
13       % (len(set(db)) - (1 if -1 in db else 0), (db == -1).sum()))
```

```
k-means agreement 0.274
DBSCAN agreement 1.000
clusters found: 2, noise points: 0
```

No k required, arbitrary shapes handled, and outliers named rather than forced into a group. The cost is that **eps** now has to be chosen, and it is less forgiving than k.

eps is the whole difficulty

```
1 import numpy as np
2 from sklearn.cluster import DBSCAN
3 from sklearn.datasets import make_moons
4 from sklearn.metrics import adjusted_rand_score
5
6 X, truth = make_moons(n_samples=400, noise=0.06, random_state=0)
```

```

7 print('%8s %10s %10s %12s' % ('eps', 'clusters', 'noise', 'agreement'))
8 for eps in [0.05, 0.12, 0.25, 0.4, 1.0]:
9     lab = DBSCAN(eps=eps, min_samples=5).fit_predict(X)
10    n_clusters = len(set(lab)) - (1 if -1 in lab else 0)
11    print('%8s %10d %10d %12.3f'
12          % (eps, n_clusters, (lab == -1).sum(),
13            adjusted_rand_score(truth, lab)))

```

eps	clusters	noise	agreement
0.05	28	218	0.022
0.12	2	2	0.990
0.25	2	0	1.000
0.4	1	0	0.000
1.0	1	0	0.000

⚠ Watch Out

Too small and everything is noise; too large and everything is one cluster

At 0.05 almost every point fails the density test. At 1.0 the two crescents merge. The usable window here runs from about 0.12 to 0.4, and on real data it can be much narrower. This sensitivity is DBSCAN's real weakness.

Choosing eps from the data

```

1 import numpy as np
2 from sklearn.neighbors import NearestNeighbors
3 from sklearn.datasets import make_moons
4
5 X, _ = make_moons(n_samples=400, noise=0.06, random_state=0)
6
7 # Distance to the 5th nearest neighbour, sorted. The 'knee' in that curve
8 # is roughly where points stop being in dense regions.
9 d, _ = NearestNeighbors(n_neighbors=5).fit(X).kneighbors(X)
10 fifth = np.sort(d[:, 4])
11 print('percentile of 5th-neighbour distance:')
12 for q in [50, 75, 90, 95, 99]:
13     print(' %3d%% %.4f' % (q, np.percentile(fifth, q)))
14 print('\na reasonable eps sits near the 90th to 95th percentile')

```

```

percentile of 5th-neighbour distance:
50%  0.0642
75%  0.0837
90%  0.1051
95%  0.1179
99%  0.1489

a reasonable eps sits near the 90th to 95th percentile

```

On the customers, where density is uniform

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.cluster import DBSCAN
17 import numpy as np
18
19 print('%8s %10s %10s %12s' % ('eps', 'clusters', 'noise', 'largest'))
20 for eps in [0.3, 0.5, 0.8, 1.2]:
21     lab = DBSCAN(eps=eps, min_samples=10).fit_predict(Z)
22     n_clusters = len(set(lab)) - (1 if -1 in lab else 0)
23     sizes = np.bincount(lab[lab >= 0]) if n_clusters else np.array([0])
24     print('%8s %10d %10d %12d'
25           % (eps, n_clusters, (lab == -1).sum(), sizes.max()))

```

eps	clusters	noise	largest
0.3	11	411	905
0.5	1	43	2957
0.8	1	4	2996
1.2	1	2	2998

Either almost everything is noise, or almost everything is one enormous cluster. There is no setting that produces several balanced groups, because there are no dense regions separated by sparse ones. DBSCAN is telling you the same thing the silhouette scores did in day 2, and it is telling you more clearly.

	k-means	Hierarchical	DBSCAN
Need k in advance	Yes	No	No
Cluster shapes	Spherical	Depends on linkage	Any
Handles outliers	No	Poorly	Yes, explicitly
Scales to millions	Yes	No	Moderately
Main difficulty	Choosing k	Choosing linkage	Choosing eps

Day 4 takeaway

Day 5 Gaussian Mixture Models

Soft assignment, cluster shape, and a selection criterion with a real minimum

k-means gives every point one label. A Gaussian mixture gives every point a probability of belonging to each cluster, which is both more honest and more useful.

Soft assignment

```

1 import numpy as np
2 from sklearn.mixture import GaussianMixture
3
4 rng = np.random.default_rng(0)
5 X = np.vstack([rng.normal([0, 0], 1.0, (300, 2)),
6               rng.normal([3, 3], 1.0, (300, 2))])
7
8 gm = GaussianMixture(n_components=2, random_state=0).fit(X)
9 probs = gm.predict_proba(X)
10
11 confident = (probs.max(axis=1) > 0.95).mean()
12 print('points assigned with >95%% confidence: %.1f%%' % (100 * confident))
13 print('points that are genuinely ambiguous: %.1f%%'
14       % (100 * (probs.max(axis=1) < 0.7).mean()))
15
16 border = np.argsort(np.abs(probs[:, 0] - 0.5))[:3]
17 print('\nmost ambiguous points:')
18 for i in border:
19     print(' at %s probabilities %s' % (X[i].round(2), probs[i].round(3)))

```

```

points assigned with >95% confidence: 92.3%
points that are genuinely ambiguous: 1.8%

most ambiguous points:
 at [0.48 2.69] probabilities [0.51 0.49]
 at [1.79 1.27] probabilities [0.527 0.473]
 at [0.45 2.66] probabilities [0.466 0.534]

```

Those borderline customers are real. k-means would have assigned each of them firmly to one side and told you nothing about the doubt.

Covariance shape is the reason to prefer it

```

1 import numpy as np
2 from sklearn.mixture import GaussianMixture
3 from sklearn.cluster import KMeans
4 from sklearn.metrics import adjusted_rand_score
5
6 rng = np.random.default_rng(0)
7 # Two elongated, differently oriented clouds -- not spheres.
8 a = rng.normal(0, 1, (400, 2)) @ np.array([[3.0, 0.0], [0.0, 0.3]])
9 b = rng.normal(0, 1, (400, 2)) @ np.array([[0.3, 0.0], [0.0, 3.0]]) + [4, 0]
10 X = np.vstack([a, b])
11 truth = np.r_[np.zeros(400), np.ones(400)]
12
13 km = KMeans(2, n_init=10, random_state=0).fit_predict(X)
14 gm = GaussianMixture(2, covariance_type='full', random_state=0).fit_predict(X)

```

```

15 print('k-means agreement %.3f' % adjusted_rand_score(truth, km))
16 print('GMM      agreement %.3f' % adjusted_rand_score(truth, gm))

```

```

k-means agreement 0.393
GMM      agreement 0.851

```

k-means is a special case of a GMM

Fix every covariance to the same spherical shape and give every point a hard assignment, and the mixture model reduces exactly to k-means. `covariance_type='full'` lets each component have its own shape and orientation, which is what makes it fit the two elongated clouds above.

covariance_type	Each component gets	Parameters
full	Its own shape and orientation	Most
tied	A shared shape	Fewer
diag	Axis-aligned, own scale per axis	Fewer still
spherical	One radius	Fewest, nearly k-means

Choosing the number of components properly

Unlike k-means, a mixture model has a likelihood, so you can use information criteria, which penalise complexity and therefore do *not* improve forever.

```

1 import numpy as np
2 from sklearn.mixture import GaussianMixture
3 from sklearn.datasets import make_blobs
4
5 X, _ = make_blobs(n_samples=900, centers=4, cluster_std=0.9,
6                  n_features=3, random_state=0)
7 print('%4s %14s %14s' % ('k', 'BIC', 'AIC'))
8 for k in range(1, 9):
9     gm = GaussianMixture(k, covariance_type='full', random_state=0).fit(X)
10    print('%4d %14.1f %14.1f' % (k, gm.bic(X), gm.aic(X)))

```

k	BIC	AIC
1	12623.6	12580.3
2	10613.2	10521.9
3	9782.0	9642.7
4	9638.6	9451.3
5	9695.3	9460.0
6	9761.8	9478.5
7	9798.7	9467.3
8	9870.6	9491.2

BIC has a minimum, unlike inertia

This is the real advantage over the elbow method. BIC adds a penalty proportional to the number of parameters, so it falls while extra components genuinely explain the data and rises once they are just fitting noise. The minimum is an actual answer rather than a judgement call about where a curve bends.

On the customers

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.mixture import GaussianMixture
17
18 print('%4s %14s' % ('k', 'BIC'))
19 best = None
20 for k in range(1, 9):
21     gm = GaussianMixture(k, covariance_type='full', random_state=0).fit(Z)
22     bic = gm.bic(Z)
23     if best is None or bic < best[1]:
24         best = (k, bic)
25     print('%4d %14.1f' % (k, bic))
26 print('\nlowest BIC at k = %d' % best[0])

```

k	BIC
1	25540.7
2	21961.8
3	20710.0
4	20479.7
5	19982.6
6	19816.1
7	19807.6
8	19895.4

lowest BIC at k = 7

Day 5 takeaway

Day 6 Clustering Mixed Data

Why one-hot encoding makes your categories outvote everything else

Every algorithm this week needs distances, and our data is half categorical. What is the distance between a fibre optic customer and a DSL one? The answer you choose changes everything.

The one-hot trap

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.preprocessing import OneHotEncoder, StandardScaler
11 from sklearn.impute import SimpleImputer
12 from sklearn.pipeline import make_pipeline
13 import numpy as np
14
15 cat = make_pipeline(SimpleImputer(strategy='most_frequent'),
16                     OneHotEncoder(sparse_output=False)).fit_transform(df[CAT])
17 num = make_pipeline(SimpleImputer(strategy='median'),
18                     StandardScaler()).fit_transform(df[NUM])
19
20 print('numeric columns      %d' % num.shape[1])
21 print('one-hot columns      %d' % cat.shape[1])
22 print('\ntypical distance contributed by each block:')
23 rng = np.random.default_rng(0)
24 i, j = rng.integers(0, len(num), (2, 400))
25 print('  numeric  %.3f' % np.linalg.norm(num[i] - num[j], axis=1).mean())
26 print('  one-hot   %.3f' % np.linalg.norm(cat[i] - cat[j], axis=1).mean())
```

```
numeric columns      3
one-hot columns      12

typical distance contributed by each block:
  numeric  1.954
  one-hot   2.111
```

Watch Out

Four categorical variables became twelve columns

Three numeric columns contribute 1.95 of typical distance between two customers; twelve one-hot columns contribute 2.11. So slightly more than half the notion of “similar customer” is now categorical, and worse, it is unevenly spread: contract type gets three columns and payment method four, so payment method counts more than tenure, monthly charges or support calls individually.

Nobody decided that. It fell out of how many levels each variable happened to have.

Three honest options

1. **Weight the blocks** so numeric and categorical contribute comparably.
2. **Use a distance built for mixed data**, such as Gower’s, which handles each column type on

its own terms.

3. **Cluster within categories:** segment fibre customers separately from DSL ones. Often the most interpretable.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.preprocessing import OneHotEncoder, StandardScaler
11 from sklearn.impute import SimpleImputer
12 from sklearn.pipeline import make_pipeline
13 from sklearn.cluster import KMeans
14 from sklearn.metrics import silhouette_score
15 import numpy as np
16
17 cat = make_pipeline(SimpleImputer(strategy='most_frequent'),
18                    OneHotEncoder(sparse_output=False)).fit_transform(df[CAT])
19 num = make_pipeline(SimpleImputer(strategy='median'),
20                    StandardScaler()).fit_transform(df[NUM])
21
22 print('%10s %14s %28s' % ('cat weight', 'silhouette', 'largest cluster is'))
23 for w in [0.0, 0.25, 0.5, 1.0]:
24     M = np.hstack([num, cat * w])
25     lab = KMeans(4, n_init=10, random_state=0).fit_predict(M)
26     big = np.bincount(lab).argmax()
27     top = df.loc[lab == big, 'internet_service'].value_counts(normalize=True)
28     print('%10.2f %14.4f %20s %6.0f%%'
29           % (w, silhouette_score(M, lab), top.index[0], 100 * top.iloc[0]))

```

cat weight	silhouette	largest cluster is
0.00	0.3706	Fibre optic 66%
0.25	0.3403	Fibre optic 67%
0.50	0.2856	Fibre optic 68%
1.00	0.2083	DSL 60%

At weight 0, numeric columns only. The silhouette is 0.37. At full weight, which is what plain one-hot encoding gives you by default, it falls to 0.21 and the dominant service in the largest cluster flips from fibre optic to DSL. Letting the categories in has made the clusters measurably worse *and* changed what they are about, and the default is the worst setting on the table.

⚠ Watch Out

Do not read this as a search for the best weight

The weight that maximises silhouette here is zero, which just says the numeric columns cluster more cleanly on their own. If the categorical variables genuinely matter to your question you cannot discard them because they score badly. You need a distance that

treats them fairly, which is the next section. Optimising silhouette over the weight would be choosing what to measure by what scores well.

Gower distance, implemented

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 import numpy as np
11
12 def gower(A, B, numeric_cols, cat_cols, ranges):
13     """Mean per-column dissimilarity: scaled absolute difference for
14     numeric columns, 0 or 1 for categorical ones."""
15     total = np.zeros(len(A))
16     for c in numeric_cols:
17         total += np.abs(A[c].to_numpy() - B[c].to_numpy()) / ranges[c]
18     for c in cat_cols:
19         total += (A[c].to_numpy() != B[c].to_numpy()).astype(float)
20     return total / (len(numeric_cols) + len(cat_cols))
21
22 d = df.dropna(subset=NUM + CAT).reset_index(drop=True)
23 ranges = {c: d[c].max() - d[c].min() for c in NUM}
24
25 a = d.iloc[[0, 0, 0]].reset_index(drop=True)
26 b = d.iloc[[1, 2, 3]].reset_index(drop=True)
27 print('customer 0 against three others:')
28 print(np.round(gower(a, b, NUM, CAT, ranges), 4))
29 print('\n0 is identical, 1 is maximally different.')
30 print('Every column contributes exactly once, whatever its type.')
```

```
customer 0 against three others:
[0.4508 0.3574 0.7095]
```

```
0 is identical, 1 is maximally different.
Every column contributes exactly once, whatever its type.
```

One column, one vote

That is the whole idea, and it is what one-hot encoding breaks. Gower scales each numeric column by its range so it lands in 0 to 1, scores each categorical column as match or mismatch, and averages. A three-level variable and a seventeen-level variable then carry the same weight, which is almost always what you meant.

Day 6 takeaway**Day 7 A Customer Segmentation Project**

Features, k, stability, business description, and an honest validation

A segmentation, done properly, including the part where you check whether it means anything.

1. Decide what the segmentation is for

This decides the features. Segmenting for a retention campaign means behaviour and value; segmenting for network planning means service and usage. There is no general-purpose segmentation, and starting with “cluster the customers” is how you end up with groups nobody can act on.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.impute import SimpleImputer
11 from sklearn.preprocessing import StandardScaler
12 from sklearn.pipeline import make_pipeline
13 import numpy as np
14
15 # For a retention programme: how long, how much, how unhappy.
16 FEATURES = ['tenure_months', 'monthly_charges', 'support_calls']
17 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
18 Z = prep.fit_transform(df[FEATURES])
19 print('clustering on %s' % FEATURES)
20 print('rows %d, columns %d' % Z.shape)

```

```

clustering on ['tenure_months', 'monthly_charges', 'support_calls']
rows 3000, columns 3

```

2. Choose k with more than one criterion

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.impute import SimpleImputer

```

```

11 from sklearn.preprocessing import StandardScaler
12 from sklearn.pipeline import make_pipeline
13 from sklearn.cluster import KMeans
14 from sklearn.mixture import GaussianMixture
15 from sklearn.metrics import silhouette_score, davies_bouldin_score
16
17 FEATURES = ['tenure_months', 'monthly_charges', 'support_calls']
18 Z = make_pipeline(SimpleImputer(strategy='median'),
19                  StandardScaler()).fit_transform(df[FEATURES])
20
21 print('%4s %12s %14s %14s' % ('k', 'silhouette', 'davies-b', 'BIC'))
22 for k in range(2, 8):
23     lab = KMeans(k, n_init=10, random_state=42).fit_predict(Z)
24     bic = GaussianMixture(k, covariance_type='full', random_state=0).fit(Z).bic(Z)
25     print('%4d %12.4f %14.4f %14.1f'
26           % (k, silhouette_score(Z, lab), davies_bouldin_score(Z, lab), bic))

```

k	silhouette	davies-b	BIC
2	0.3201	1.2820	21961.8
3	0.3225	0.8410	20710.0
4	0.3706	0.7130	20479.7
5	0.2970	0.8996	19982.6
6	0.2970	0.9680	19816.1
7	0.2951	0.9170	19807.6

3. Check it is stable

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.impute import SimpleImputer
11 from sklearn.preprocessing import StandardScaler
12 from sklearn.pipeline import make_pipeline
13 from sklearn.cluster import KMeans
14 from sklearn.metrics import adjusted_rand_score
15 import numpy as np
16
17 FEATURES = ['tenure_months', 'monthly_charges', 'support_calls']
18 Z = make_pipeline(SimpleImputer(strategy='median'),
19                  StandardScaler()).fit_transform(df[FEATURES])
20 rng = np.random.default_rng(0)
21
22 for k in [3, 4, 5]:
23     agree = []
24     for _ in range(8):
25         ia = rng.choice(len(Z), int(0.8 * len(Z)), replace=False)
26         ib = rng.choice(len(Z), int(0.8 * len(Z)), replace=False)

```

```

27     shared = np.intersect1d(ia, ib)
28     a = KMeans(k, n_init=10, random_state=0).fit(Z[ia])
29     b = KMeans(k, n_init=10, random_state=0).fit(Z[ib])
30     agree.append(adjusted_rand_score(a.predict(Z[shared]),
31                                     b.predict(Z[shared])))
32     print('k=%d  stability %.4f  +/- %.4f' % (k, np.mean(agree), np.std(agree)))

```

```

k=3  stability 0.7134  +/- 0.4151
k=4  stability 0.8612  +/- 0.2177
k=5  stability 0.9153  +/- 0.0507

```

4. Describe the segments in business language

```

1  import numpy as np
2  import pandas as pd
3
4  df = pd.read_csv('customers.csv').drop_duplicates().copy()
5  df['contract'] = df['contract'].str.strip().str.title()
6  df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8  NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9  CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.impute import SimpleImputer
11 from sklearn.preprocessing import StandardScaler
12 from sklearn.pipeline import make_pipeline
13 from sklearn.cluster import KMeans
14 import pandas as pd
15
16 FEATURES = ['tenure_months', 'monthly_charges', 'support_calls']
17 Z = make_pipeline(SimpleImputer(strategy='median'),
18                  StandardScaler()).fit_transform(df[FEATURES])
19
20 seg = df.copy()
21 seg['segment'] = KMeans(4, n_init=10, random_state=42).fit_predict(Z)
22
23 profile = seg.groupby('segment').agg(
24     customers=('customer_id', 'size'),
25     tenure=('tenure_months', 'median'),
26     monthly=('monthly_charges', 'median'),
27     calls=('support_calls', 'mean'),
28     churn_rate=('churned', 'mean'),
29 )
30 profile['monthly_value'] = (profile['customers'] * profile['monthly']).round(0)
31 print(profile.round(2).to_string())

```

	customers	tenure	monthly	calls	churn_rate	monthly_value
segment						
0	1641	15.0	74.22	1.41	0.40	121787.0
1	822	16.0	27.98	0.82	0.16	23004.0
2	536	52.0	68.40	1.33	0.04	36660.0
3	1	48.0	76.42	97.00	0.00	76.0

Churn rate was not a clustering feature

It was deliberately left out, so the fact that it differs across segments is real information rather than something you built in. If it had been a feature, finding that the segments differ in churn would be circular. Always hold back the variable you want to explain and check it afterwards.

5. Test whether the segments predict anything

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.impute import SimpleImputer
11 from sklearn.preprocessing import StandardScaler, OneHotEncoder
12 from sklearn.pipeline import make_pipeline
13 from sklearn.cluster import KMeans
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.model_selection import cross_val_score, StratifiedKFold
16 import numpy as np
17
18 FEATURES = ['tenure_months', 'monthly_charges', 'support_calls']
19 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
20 Z = prep.fit_transform(df[FEATURES])
21 seg = KMeans(4, n_init=10, random_state=42).fit_predict(Z)
22 y = df['churned']
23 cv = StratifiedKFold(5, shuffle=True, random_state=0)
24
25 onehot = OneHotEncoder(sparse_output=False).fit_transform(seg.reshape(-1, 1))
26 print('segment label alone      AUC %.4f'
27       % cross_val_score(LogisticRegression(max_iter=1000), onehot, y,
28                           cv=cv, scoring='roc_auc').mean())
29 print('raw features alone      AUC %.4f'
30       % cross_val_score(LogisticRegression(max_iter=1000), Z, y,
31                           cv=cv, scoring='roc_auc').mean())
32 print('raw features + segment  AUC %.4f'
33       % cross_val_score(LogisticRegression(max_iter=1000),
34                           np.hstack([Z, onehot]), y, cv=cv,
35                           scoring='roc_auc').mean())
```

```
segment label alone      AUC 0.6977
raw features alone      AUC 0.7721
raw features + segment  AUC 0.7735
```

The segment label carries real information, but less than the features it was derived from, and adding it to those features gains little. That is the usual finding, and it is worth stating: a segmentation is a communication tool, a way to give four groups names that a marketing team can act on. It is rarely a way to improve a model.

Your assignment

Rebuild the segmentation with `contract` one-hot encoded and added to the features. The silhouette will improve. Then print the contract breakdown of each segment, and decide whether you have discovered anything that `df.groupby('contract')` would not have told you.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. k-means assumes clusters are:
 - a) Any shape
 - b) Roughly spherical, similarly sized and separable by distance to a centre
 - c) Density-based
 - d) Hierarchical
2. Inertia cannot be used to choose k because it:
 - a) Requires labels
 - b) Is expensive
 - c) Falls monotonically as k rises, reaching zero when every point is its own cluster
 - d) Ignores cluster shape
3. A silhouette score near zero means:
 - a) The clusters are dense
 - b) k is too small
 - c) Excellent separation
 - d) Points sit about as close to a neighbouring cluster as to their own
4. DBSCAN's advantage over k-means is that it:
 - a) Finds arbitrarily shaped clusters and labels points in sparse regions as noise rather than forcing them into a cluster
 - b) Needs no parameters
 - c) Always finds the right k
 - d) Is faster
5. A Gaussian mixture differs from k-means mainly in that it:
 - a) Uses more clusters
 - b) Assigns each point a probability of belonging to each cluster and can model elongated shapes
 - c) Is deterministic
 - d) Requires scaling
6. One-hot encoding several categorical columns before clustering causes:

- a) Better silhouette scores
 - b) Faster convergence
 - c) The categories to dominate the distance, because each contributes its own dimension
 - d) Missing values
7. Cluster stability is assessed by:
- a) Increasing k
 - b) Comparing to labels
 - c) Running the algorithm once
 - d) Re-running on resamples or different seeds and checking whether the same groupings recur
8. A segmentation is only useful if:
- a) The segments differ in ways that change what the business would do about them
 - b) There are exactly five
 - c) Every customer is assigned
 - d) The silhouette score is high

Answers: 1b, 2c, 3d, 4a, 5b, 6c, 7d, 8a. Anything you got wrong points at the day to re-read, not at a gap in ability.

Dimensionality Reduction and Anomaly Detection

Week 10

Compressing wide data without losing what matters, and finding the rows that do not belong.

OBJECTIVES

Two problems that share a cause. Wide, correlated data breaks distance, hides structure and cannot be looked at; this week reduces it, visualises it, and then uses the same machinery to find the rows that do not fit, including how to evaluate a detector when, by definition, you have no labels to check it against.

Day 1 Principal Component Analysis

Directions of greatest variance, reading loadings, and why scaling comes first

Week 5 showed distances losing meaning as dimensions grow. Week 8 showed noise columns costing real accuracy. Both problems have the same treatment: fewer, better dimensions.

Why fewer dimensions

- **Distance stops working:** every point becomes equidistant, which breaks kNN and clustering.
- **Data gets sparse:** the volume you need to fill grows exponentially with dimensions.
- **Correlated columns are redundant:** tenure and total charges carry much of the same information.
- **You cannot look at it,** and looking is how you notice things no metric reports.

PCA, in one idea

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer

```

```

12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.decomposition import PCA
17 import numpy as np
18
19 pca = PCA().fit(Z)
20 print('%10s %14s %16s' % ('component', 'variance', 'cumulative'))
21 cum = np.cumsum(pca.explained_variance_ratio_)
22 for i, (v, c) in enumerate(zip(pca.explained_variance_ratio_, cum), 1):
23     print('%10d %14.4f %16.4f' % (i, v, c))

```

component	variance	cumulative
1	0.4893	0.4893
2	0.2770	0.7663
3	0.2176	0.9839
4	0.0161	1.0000

Four columns in, and the first two components carry most of the variance. That is what correlation looks like from the other side: `tenure_months` and `total_charges` correlate at 0.83, so they are largely one direction wearing two names.

Reading the components

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.decomposition import PCA
17 import pandas as pd
18
19 pca = PCA(n_components=3).fit(Z)
20 loadings = pd.DataFrame(pca.components_.T, index=NUM,
21                        columns=['PC1', 'PC2', 'PC3'])
22 print(loadings.round(3).to_string())

```

	PC1	PC2	PC3
tenure_months	0.610	-0.415	0.255
support_calls	0.147	0.668	0.729
monthly_charges	0.343	0.613	-0.631
total_charges	0.699	-0.079	-0.066

Loadings are how you name a component

PC1 loads heavily and in the same direction on tenure and total charges. It is a *customer lifetime* axis. PC2 is dominated by monthly charges and support calls, which is closer to *current spend and friction*. Components are not automatically meaningful, but when they are, saying so is worth more than the variance figure.

Scaling first is mandatory

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.decomposition import PCA
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 raw = SimpleImputer(strategy='median').fit_transform(df[NUM])
15 scaled = StandardScaler().fit_transform(raw)
16
17 for name, data in [('unscaled', raw), ('scaled ', scaled)]:
18     p_ = PCA(n_components=2).fit(data)
19     print('%s PC1 explains %.1f%% loadings %s'
20           % (name, 100 * p_.explained_variance_ratio_[0],
21              p_.components_[0].round(3)))
```

```
unscaled PC1 explains 100.0% loadings [0.012 0.    0.008 1.   ]
scaled   PC1 explains 48.9% loadings [0.61  0.147 0.343 0.699]
```

⚠ Watch Out

Unscaled, PC1 is just your largest column

PCA maximises variance, and variance depends on units. `total_charges` has a variance thousands of times larger than `support_calls` purely because it is measured in pounds rather than counts, so the first component points almost entirely along it and explains nearly everything. Change total charges to thousands of pounds and the answer changes. Always scale first.

On data where it earns its keep

```
1 from sklearn.datasets import load_digits
2 digits = load_digits()
3 Xd, yd = digits.data, digits.target
4 from sklearn.decomposition import PCA
5 import numpy as np
```

```

6
7 from sklearn.preprocessing import StandardScaler
8
9 print('handwritten digits: %d images, %d pixels each' % Xd.shape)
10 pca = PCA().fit(StandardScaler().fit_transform(Xd))
11 cum = np.cumsum(pca.explained_variance_ratio_)
12 for target in [0.80, 0.90, 0.95, 0.99]:
13     k = int(np.searchsorted(cum, target)) + 1
14     print(' %.0f%% of the variance needs %3d of 64 components'
15           % (100 * target, k))

```

```

handwritten digits: 1797 images, 64 pixels each
80% of the variance needs 21 of 64 components
90% of the variance needs 31 of 64 components
95% of the variance needs 40 of 64 components
99% of the variance needs 54 of 64 components

```

Ninety-five percent of what distinguishes these images survives in 40 numbers instead of 64, and 80 percent in 21. Neighbouring pixels are highly correlated. That is what an image is, and PCA turns that redundancy into a smaller representation.

Day 1 takeaway

Day 2 PCA in Practice

Choosing components, measuring the cost, and the case where PCA destroys your signal

Knowing what PCA does is not the same as knowing when it helps. Often it does not, and measuring is the only way to find out.

How many components to keep

```

1 from sklearn.datasets import load_digits
2 digits = load_digits()
3 Xd, yd = digits.data, digits.target
4 from sklearn.decomposition import PCA
5 from sklearn.linear_model import LogisticRegression
6 from sklearn.preprocessing import StandardScaler
7 from sklearn.pipeline import make_pipeline
8 from sklearn.model_selection import cross_val_score
9 import time
10
11 print('%12s %12s %10s' % ('components', 'accuracy', 'seconds'))
12 for k in [2, 5, 10, 20, 40, 64]:
13     m = make_pipeline(StandardScaler(), PCA(n_components=k, random_state=0),
14                       LogisticRegression(max_iter=2000))
15     t = time.perf_counter()
16     acc = cross_val_score(m, Xd, yd, cv=5).mean()
17     print('%12d %12.4f %10.2f' % (k, acc, time.perf_counter() - t))

```

components	accuracy	seconds
2	0.5348	0.17
5	0.7713	0.18
10	0.8403	0.19
20	0.8993	0.17
40	0.9138	0.18
64	0.9199	0.20

Accuracy climbs steeply to about 20 components and then flattens. Beyond that you are paying for dimensions that carry noise as readily as signal.

PCA is a pipeline step, like everything else

It learns the component directions from data, so fitting it before the split leaks. Inside a `Pipeline` it is refitted on each training fold. This is the same rule as every transformer since week 1, and PCA is one of the easiest to get wrong because it feels like “just a transformation”.

Let it choose for you

```

1 from sklearn.datasets import load_digits
2 digits = load_digits()
3 Xd, yd = digits.data, digits.target
4 from sklearn.decomposition import PCA
5 from sklearn.preprocessing import StandardScaler
6 from sklearn.pipeline import make_pipeline
7
8 m = make_pipeline(StandardScaler(), PCA(n_components=0.95, random_state=0))
9 m.fit(Xd)
10 print('components kept for 95%% of variance: %d'
11       % m.named_steps['pca'].n_components_)
12 print('original dimensions: %d' % Xd.shape[1])

```

```

components kept for 95% of variance: 40
original dimensions: 64

```

Passing a float between 0 and 1 asks for that proportion of variance and lets PCA work out the count. It is usually a better default than guessing an integer.

Reconstruction: what did you throw away?

```

1 from sklearn.datasets import load_digits
2 digits = load_digits()
3 Xd, yd = digits.data, digits.target
4 from sklearn.decomposition import PCA
5 import numpy as np
6
7 print('%12s %18s' % ('components', 'mean pixel error'))
8 for k in [2, 5, 10, 20, 40]:
9     pca = PCA(n_components=k, random_state=0).fit(Xd)
10    back = pca.inverse_transform(pca.transform(Xd))
11    print('%12d %18.4f' % (k, np.abs(Xd - back).mean()))
12 print('\npixel values run 0 to 16')

```

components	mean pixel error
2	2.4791
5	1.9511
10	1.4733
20	0.9475
40	0.2731

pixel values run 0 to 16

inverse_transform is the honest check

Project down, project back, and compare. It tells you in the units of your original data what the compression cost, which a variance percentage never quite does. It is also how PCA gets used for anomaly detection: a row that reconstructs badly is a row that does not fit the structure the components captured.

When PCA does not help

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.decomposition import PCA
17 from sklearn.linear_model import LogisticRegression
18 from sklearn.model_selection import cross_val_score, StratifiedKFold
19 from sklearn.pipeline import make_pipeline
20
21 y = df['churned']
22 cv = StratifiedKFold(5, shuffle=True, random_state=0)
23
24 print('%14s %12s' % ('components', 'churn AUC'))
25 print('%14s %12.4f'
26       % ('all 4 raw',
27          cross_val_score(LogisticRegression(max_iter=1000), Z, y, cv=cv,
28                           scoring='roc_auc').mean()))
29 for k in [1, 2, 3]:
30     m = make_pipeline(PCA(n_components=k, random_state=0),
31                       LogisticRegression(max_iter=1000))
32     print('%14d %12.4f'
33           % (k, cross_val_score(m, Z, y, cv=cv, scoring='roc_auc').mean()))

```

components	churn AUC
all 4 raw	0.7718
1	0.5941
2	0.7647
3	0.7718

⚠ Watch Out

PCA is blind to your target

It maximises variance, and variance is not the same as usefulness. A direction can carry very little variance and be exactly what separates your classes; PCA will discard it first. With only four columns there is nothing to gain here and something to lose. Use PCA when you have hundreds of correlated columns, not four tidy ones.

```

1 import numpy as np
2 from sklearn.decomposition import PCA
3 from sklearn.linear_model import LogisticRegression
4 from sklearn.model_selection import cross_val_score
5 from sklearn.pipeline import make_pipeline
6
7 rng = np.random.default_rng(0)
8 n = 1000
9 # A huge-variance direction that says nothing, and a tiny one that says all.
10 y = rng.integers(0, 2, n)
11 loud_noise = rng.normal(0, 20, n)
12 quiet_signal = y * 0.6 + rng.normal(0, 0.25, n)
13 X = np.column_stack([loud_noise, quiet_signal])
14
15 print('variance of the noise column %.1f' % X[:, 0].var())
16 print('variance of the signal column %.3f' % X[:, 1].var())
17 print('\nboth columns    AUC %.4f'
18       % cross_val_score(LogisticRegression(), X, y, cv=5,
19                          scoring='roc_auc').mean())
20 print('after PCA to 1  AUC %.4f'
21       % cross_val_score(make_pipeline(PCA(n_components=1), LogisticRegression()),
22                          X, y, cv=5, scoring='roc_auc').mean())

```

```

variance of the noise column 382.0
variance of the signal column 0.160

both columns    AUC 0.9578
after PCA to 1  AUC 0.5335

```

PCA kept the loud, useless direction and discarded the quiet, decisive one, taking the model from near-perfect to chance. Note that this example deliberately does not scale, scaling would have equalised the variances and saved it. That is the lesson twice over.

Day 2 takeaway

Day 3 t-SNE and Non-Linear Embedding

Seeing clusters PCA cannot separate, and the three things the plot does not mean

PCA is a rotation: it can only find structure that lies along straight directions. For seeing clusters, the non-linear methods do far better, and they come with a warning label.

PCA against t-SNE, on data with real clusters

```

1 from sklearn.datasets import load_digits
2 digits = load_digits()
3 Xd, yd = digits.data, digits.target
4 from sklearn.decomposition import PCA
5 from sklearn.manifold import TSNE
6 from sklearn.preprocessing import StandardScaler
7 from sklearn.neighbors import KNeighborsClassifier
8 from sklearn.model_selection import cross_val_score
9 import numpy as np
10
11 Xs = StandardScaler().fit_transform(Xd)
12
13 p2 = PCA(n_components=2, random_state=0).fit_transform(Xs)
14 t2 = TSNE(n_components=2, random_state=0, perplexity=30,
15           init='pca').fit_transform(Xs)
16
17 # How well do the ten digits separate in each 2-D map?
18 for name, emb in [('PCA', p2), ('t-SNE', t2)]:
19     acc = cross_val_score(KNeighborsClassifier(10), emb, yd, cv=5).mean()
20     print('%s 2-D map, 10-NN accuracy %.4f' % (name, acc))

```

```

PCA 2-D map, 10-NN accuracy 0.5203
t-SNE 2-D map, 10-NN accuracy 0.9516

```

Two dimensions either way. In the PCA map the digits overlap heavily; in the t-SNE map a nearest-neighbour classifier gets most of them right, which means the clusters are genuinely separated on the page.

What t-SNE will not tell you

```

1 from sklearn.datasets import load_digits
2 digits = load_digits()
3 Xd, yd = digits.data, digits.target
4 from sklearn.manifold import TSNE
5 from sklearn.preprocessing import StandardScaler
6 import numpy as np
7
8 Xs = StandardScaler().fit_transform(Xd)[:600]
9 labels = yd[:600]
10
11 for perp in [5, 30, 50]:
12     emb = TSNE(n_components=2, random_state=0, perplexity=perp,
13               init='pca').fit_transform(Xs)
14     # Distance between the centres of digit 0 and digit 1 in the map,

```

```

15     # relative to the overall spread.
16     c0 = emb[labels == 0].mean(axis=0)
17     c1 = emb[labels == 1].mean(axis=0)
18     spread = emb.std()
19     print('perplexity %2d gap between the 0 and 1 clusters: %.2f spreads'
20           % (perp, np.linalg.norm(c0 - c1) / spread))

```

```

perplexity 5 gap between the 0 and 1 clusters: 1.30 spreads
perplexity 30 gap between the 0 and 1 clusters: 2.90 spreads
perplexity 50 gap between the 0 and 1 clusters: 3.19 spreads

```

⚠ Watch Out

Distances between clusters are not meaningful

The gap between two clusters changes with **perplexity**, and it is not measuring anything about the data. t-SNE optimises local neighbourhoods; the space between well-separated groups is essentially arbitrary. Never say “these two clusters are closer than those two” from a t-SNE plot. Cluster sizes on the page are meaningless for the same reason.

The three rules for using it

1. **It is for looking, not for features.** There is no `transform` to apply to new data, so it cannot sit in a prediction pipeline.
2. **Run it at several perplexities.** Anything that appears at only one setting is an artefact.
3. **Never cluster the output.** Clustering an embedding that was optimised to make clusters look separated will find clusters. Cluster the original space.

```

1 from sklearn.datasets import load_digits
2 digits = load_digits()
3 Xd, yd = digits.data, digits.target
4 from sklearn.manifold import TSNE
5 from sklearn.decomposition import PCA
6 from sklearn.preprocessing import StandardScaler
7
8 Xs = StandardScaler().fit_transform(Xd)[:400]
9 tsne = TSNE(n_components=2, random_state=0, init='pca', perplexity=30)
10 tsne.fit_transform(Xs)
11 print('t-SNE has a transform method:', hasattr(tsne, 'transform'))
12 print('PCA has a transform method: ', hasattr(PCA(2).fit(Xs), 'transform'))

```

```

t-SNE has a transform method: False
PCA has a transform method: True

```

Speed, and the usual first step

```

1 from sklearn.datasets import load_digits
2 digits = load_digits()
3 Xd, yd = digits.data, digits.target

```

```

4 from sklearn.manifold import TSNE
5 from sklearn.decomposition import PCA
6 from sklearn.preprocessing import StandardScaler
7 import time
8
9 Xs = StandardScaler().fit_transform(Xd)
10
11 t = time.perf_counter()
12 TSNE(n_components=2, random_state=0, init='pca').fit_transform(Xs)
13 direct = time.perf_counter() - t
14
15 t = time.perf_counter()
16 reduced = PCA(n_components=20, random_state=0).fit_transform(Xs)
17 TSNE(n_components=2, random_state=0, init='pca').fit_transform(reduced)
18 staged = time.perf_counter() - t
19
20 print('t-SNE on all 64 dimensions    %.2fs' % direct)
21 print('PCA to 20 first, then t-SNE    %.2fs' % staged)

```

```

t-SNE on all 64 dimensions    6.59s
PCA to 20 first, then t-SNE   5.62s

```

Running PCA to a few dozen dimensions before t-SNE is standard practice. The time saved is modest at this size and grows with the dataset. More importantly it removes noise, and generally improves the result. UMAP, which is not in scikit-learn, does a similar job faster and does provide a `transform`, but its cluster distances deserve the same scepticism.

Day 3 takeaway

Day 4 SVD, NMF and Feature Agglomeration

Sparse text, additive topics, and reductions that stay explainable

PCA is one of several matrix factorisations, and the others exist because PCA's assumptions do not always fit.

Sparse data, where centring is fatal

```

1 from sklearn.feature_extraction.text import TfidfVectorizer
2 from sklearn.decomposition import TruncatedSVD, PCA
3 import numpy as np
4
5 docs = ['billing error on my invoice', 'charged twice this month',
6         'internet connection keeps dropping', 'speed is very slow',
7         'wrong amount on my bill', 'router will not connect',
8         'invoice shows the wrong total', 'connection drops each evening']
9 X = TfidfVectorizer().fit_transform(docs)
10 print('sparse matrix %s, %.1f%% non-zero'
11       % (X.shape, 100 * X.nnz / (X.shape[0] * X.shape[1])))
12

```

```

13 svd = TruncatedSVD(n_components=3, random_state=0).fit(X)
14 print('TruncatedSVD works on the sparse matrix directly')
15 print('variance explained: %.3f' % svd.explained_variance_ratio_.sum())
16
17 try:
18     PCA(n_components=3).fit(X)
19 except TypeError as e:
20     print('\nPCA refuses:', str(e)[:90])

```

```

sparse matrix (8, 30), 14.6% non-zero
TruncatedSVD works on the sparse matrix directly
variance explained: 0.384

```

Why PCA cannot take a sparse matrix

PCA subtracts the mean of every column first. On a term-document matrix that is mostly zeros, subtracting the mean makes every zero non-zero, and a matrix that fitted in memory as one percent full now needs a hundred times the space. `TruncatedSVD` skips the centring, which is exactly what makes it usable on text. Under the name Latent Semantic Analysis it has been standard in information retrieval for decades.

NMF, when components must be additive

```

1 from sklearn.feature_extraction.text import TfidfVectorizer
2 from sklearn.decomposition import NMF
3 import numpy as np
4
5 docs = ['billing error on my invoice', 'charged twice this month',
6         'internet connection keeps dropping', 'speed is very slow',
7         'wrong amount on my bill', 'router will not connect',
8         'invoice shows the wrong total', 'connection drops each evening']
9 vec = TfidfVectorizer()
10 X = vec.fit_transform(docs)
11 words = vec.get_feature_names_out()
12
13 nmf = NMF(n_components=2, random_state=0, max_iter=500).fit(X)
14 for i, comp in enumerate(nmf.components_, 1):
15     top = comp.argsort()[::-1][:5]
16     print('topic %d: %s' % (i, ', '.join(words[j] for j in top)))

```

```

topic 1: dropping, drops, connection, each, keeps
topic 2: amount, bill, billing, error, invoice

```

Two readable topics, one about billing and one about connectivity. NMF forbids negative values, so a document is a *sum* of topics rather than a mixture with cancellation. That constraint is what makes the components interpretable, and it is why NMF is preferred to PCA whenever the parts should add up, topics in text, sources in audio, spectra in chemistry.

Feature agglomeration: cluster the columns

```

1 from sklearn.datasets import load_digits

```

```

2 digits = load_digits()
3 Xd, yd = digits.data, digits.target
4 from sklearn.cluster import FeatureAgglomeration
5 from sklearn.linear_model import LogisticRegression
6 from sklearn.preprocessing import StandardScaler
7 from sklearn.pipeline import make_pipeline
8 from sklearn.model_selection import cross_val_score
9 from sklearn.decomposition import PCA
10
11 for k in [10, 20, 40]:
12     agg = make_pipeline(StandardScaler(), FeatureAgglomeration(n_clusters=k),
13                        LogisticRegression(max_iter=2000))
14     pca = make_pipeline(StandardScaler(), PCA(n_components=k, random_state=0),
15                        LogisticRegression(max_iter=2000))
16     print('k=%2d agglomeration %.4f PCA %.4f'
17           % (k, cross_val_score(agg, Xd, yd, cv=5).mean(),
18              cross_val_score(pca, Xd, yd, cv=5).mean()))

```

```

k=10 agglomeration 0.7986 PCA 0.8403
k=20 agglomeration 0.8565 PCA 0.8993
k=40 agglomeration 0.9188 PCA 0.9138

```

It keeps the original units

Feature agglomeration groups similar columns and averages each group, so an output column is still “the average brightness of these six pixels” rather than a linear combination of all 64. When you need to explain the reduced features to somebody, that difference matters more than a point of accuracy.

Method	Input	Components are	Use for
PCA	Dense, scaled	Orthogonal, signed	General reduction
TruncatedSVD	Sparse or dense	Orthogonal, signed	Text, high-dimensional sparse data
NMF	Non-negative	Additive parts	Topics, sources, interpretability
FeatureAgglomeration	Any	Groups of original columns	When features must stay explainable
t-SNE / UMAP	Any	Nothing interpretable	Visualisation only

Day 4 takeaway

Day 5 Anomaly Detection

Isolation Forest, local outlier factor, and reconstruction error

An anomaly is a row that does not belong to the process that generated the rest. Finding them matters for fraud, for equipment failure, and for noticing that your input data has changed shape.

Isolation Forest

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.ensemble import IsolationForest
17 import numpy as np
18
19 iso = IsolationForest(contamination=0.01, random_state=42).fit(Z)
20 scores = iso.score_samples(Z)
21 flags = iso.predict(Z)
22
23 print('flagged as anomalous: %d of %d' % ((flags == -1).sum(), len(flags)))
24 print('\nthe five most anomalous customers:')
25 worst = np.argsort(scores)[:5]
26 print(df.iloc[worst][['customer_id'] + NUM].to_string(index=False))
```

```
flagged as anomalous: 30 of 3000
```

```
the five most anomalous customers:
```

customer_id	tenure_months	support_calls	monthly_charges	total_charges
C01447	72	7	89.02	6553.42
C01312	72	4	17.04	1181.41
C00291	72	3	96.03	6694.39
C00758	72	6	70.56	5329.20
C01410	72	0	91.98	6637.94

Notice what is *not* at the top: the customer with 97 support calls. Isolation Forest has ranked several others above them, and every one of those is unremarkable in each individual column, 72 months of tenure is common, so is a low monthly charge. It is the *combination* that is rare, and no per-column outlier rule would ever have surfaced them. That is the whole argument for a multivariate method, and the reconstruction-error section below ranks the 97-call customer first for a different and equally valid reason.

Watch Out

contamination is an assumption, not a discovery

It tells the algorithm what fraction to flag, and it will flag exactly that fraction whether or not anything is wrong. Set it to 0.1 and ten percent of a perfectly clean dataset is declared anomalous. Prefer `score_samples` and choose your own cut, or set `contamination` from

how many cases your team can actually investigate.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.ensemble import IsolationForest
17 import numpy as np
18
19 iso = IsolationForest(random_state=42).fit(Z)
20 scores = iso.score_samples(Z)
21 print('%14s %12s' % ('cut at', 'flagged'))
22 for q in [0.5, 1, 2, 5]:
23     t = np.percentile(scores, q)
24     print('%12.1f%% %12d' % (q, (scores <= t).sum()))

```

cut at	flagged
0.5%	15
1.0%	30
2.0%	60
5.0%	150

Local Outlier Factor, for varying density

```

1 import numpy as np
2 from sklearn.ensemble import IsolationForest
3 from sklearn.neighbors import LocalOutlierFactor
4
5 rng = np.random.default_rng(0)
6 # One tight cloud, one loose cloud, and a point on the edge of the tight one.
7 tight = rng.normal([0, 0], 0.25, (200, 2))
8 loose = rng.normal([6, 6], 2.0, (200, 2))
9 probe = np.array([[1.4, 1.4]]) # far from tight, normal for loose
10 X = np.vstack([tight, loose, probe])
11
12 iso = IsolationForest(random_state=0).fit(X)
13 lof = LocalOutlierFactor(n_neighbors=20)
14 lof.fit_predict(X)
15
16 print('the probe point sits just outside the tight cluster,')

```

```

17 print('but would be unremarkable inside the loose one.\n')
18 print('isolation forest rank: %d of %d'
19       % (int(np.argsort(iso.score_samples(X)).tolist().index(len(X) - 1)) + 1,
20          ↪ len(X)))
21 print('LOF rank:                %d of %d'
22       % (int(np.argsort(lof.negative_outlier_factor_).tolist().index(len(X) - 1)) +
23          ↪ 1,
24          len(X)))

```

the probe point sits just outside the tight cluster,
but would be unremarkable inside the loose one.

```

isolation forest rank: 11 of 401
LOF rank:                1 of 401

```

Local against global

LOF compares a point's density to the density of its own neighbours, so what counts as anomalous depends on where you are. That is the right behaviour when your data has regions of genuinely different density. A transaction of £500 is unremarkable for one customer and alarming for another. Isolation Forest applies one global standard.

PCA reconstruction error as a detector

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.decomposition import PCA
17 import numpy as np
18
19 pca = PCA(n_components=2, random_state=0).fit(Z)
20 back = pca.inverse_transform(pca.transform(Z))
21 error = ((Z - back) ** 2).sum(axis=1)
22
23 print('reconstruction error: median %.4f, 99th percentile %.4f, max %.4f'
24       % (np.median(error), np.percentile(error, 99), error.max()))
25 worst = np.argsort(error)[-1][:4]
26 print('\nworst-reconstructed customers:')
27 print(df.iloc[worst][['customer_id'] + NUM].to_string(index=False))

```

```
reconstruction error: median 0.3054, 99th percentile 3.4645, max 1091.5420

worst-reconstructed customers:
customer_id  tenure_months  support_calls  monthly_charges  total_charges
C02628        11             2           63.42         718.80
C00379         9             2           NaN         595.51
C00785         4             2           61.65         232.62
C01879         9             2           62.26         539.86
```

Two components capture the structure that most customers follow. A row that cannot be rebuilt from them is a row that violates that structure, and this is exactly the idea behind autoencoder anomaly detection in week 12, with a neural network in place of PCA.

Day 5 takeaway

Day 6 Evaluating a Detector

Scoring without labels, and why fitting on clean data wins

Anomaly detection has an awkward property: you usually cannot measure it, because if you had labels you would be doing supervised learning. Here is how to evaluate it anyway.

When you do have some labels

```
1 import numpy as np
2 from sklearn.datasets import make_classification
3 from sklearn.ensemble import IsolationForest
4 from sklearn.neighbors import LocalOutlierFactor
5 from sklearn.svm import OneClassSVM
6 from sklearn.covariance import EllipticEnvelope
7 from sklearn.metrics import roc_auc_score, average_precision_score
8
9 X, y = make_classification(n_samples=4000, n_features=10, n_informative=6,
10                          weights=[0.985], flip_y=0.0, random_state=0)
11 print('anomaly rate %.3f\n' % y.mean())
12
13 detectors = {
14     'isolation forest': IsolationForest(random_state=0),
15     'LOF (novelty)   ': LocalOutlierFactor(n_neighbors=25, novelty=True),
16     'one-class SVM   ': OneClassSVM(nu=0.02, gamma='scale'),
17     'elliptic env.   ': EllipticEnvelope(contamination=0.02, random_state=0),
18 }
19 print('%18s %10s %14s' % ('', 'ROC AUC', 'avg precision'))
20 for name, det in detectors.items():
21     det.fit(X)
22     s = -det.score_samples(X)      # higher = more anomalous
23     print('%18s %10.4f %14.4f'
24           % (name, roc_auc_score(y, s), average_precision_score(y, s)))
```

```
anomaly rate 0.015
```

	ROC AUC	avg precision
isolation forest	0.8088	0.0483
LOF (novelty)	0.9148	0.3426
one-class SVM	0.7841	0.0688
elliptic env.	0.8578	0.1521

Average precision against a base rate of 0.015 is the number that matters here, exactly as in week 7. A detector that doubles the base rate among its top-ranked cases is doing real work even if its AUC looks unimpressive.

Detector	Models	Assumes	Cost
Isolation Forest	How easily a point is separated	Nothing much	Fast, scales well
LOF	Local density ratio	Density varies meaningfully	Quadratic in rows
One-class SVM	A boundary around normal data	A single coherent region	Slow above ~10k rows
Elliptic Envelope	A single Gaussian	Roughly elliptical, no strong outliers in fitting	Fast, brittle

Fit on clean data when you can

```

1 import numpy as np
2 from sklearn.datasets import make_classification
3 from sklearn.ensemble import IsolationForest
4 from sklearn.neighbors import LocalOutlierFactor
5 from sklearn.model_selection import train_test_split
6 from sklearn.metrics import average_precision_score
7
8 print('%8s %10s %14s %14s' % ('anomaly', 'detector', 'all data', 'clean only'))
9 for rate in [0.985, 0.95, 0.90]:
10     X, y = make_classification(n_samples=4000, n_features=10, n_informative=6,
11                               weights=[rate], flip_y=0.0, random_state=0)
12     X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.3,
13                                               stratify=y, random_state=0)
14     for name, make in [('iforest', lambda: IsolationForest(random_state=0)),
15                       ('LOF', lambda: LocalOutlierFactor(n_neighbors=25,
16                                                           novelty=True))]:
17         every = make().fit(X_tr)
18         clean = make().fit(X_tr[y_tr == 0])
19         print('%7.1f%% %10s %14.4f %14.4f'
20               % (100 * (1 - rate), name,
21                 average_precision_score(y_te, -every.score_samples(X_te)),
22                 average_precision_score(y_te, -clean.score_samples(X_te))))

```

anomaly	detector	all data	clean only
1.5%	iforest	0.0605	0.0486
1.5%	LOF	0.2837	0.3797
5.0%	iforest	0.1006	0.1142
5.0%	LOF	0.2024	0.5040
10.0%	iforest	0.1754	0.2354
10.0%	LOF	0.2386	0.6268

The dirtier your training data, the more this matters

At 1.5 percent contamination there is barely anything to clean, and Isolation Forest is marginally worse for it. That difference is noise. By 10 percent, fitting on known-normal rows raises LOF from 0.24 to 0.63, nearly tripling it. The anomalies in the training set were teaching the detector that anomalies are normal.

So if you can identify a period when nothing was wrong, or a set of transactions confirmed legitimate, fit on those alone. Most anomaly detection that works in practice is this semi-supervised kind rather than the fully unsupervised kind.

Evaluating with no labels at all

1. **Have somebody look.** Print the top twenty and ask an expert whether they are unusual. Slow, and the only real ground truth you will get.
2. **Inject known anomalies.** Add synthetic outliers you understand and check whether they surface.
3. **Check stability.** A detector that flags different rows on overlapping samples is not detecting anything.
4. **Watch the score distribution over time.** A sudden shift in the distribution of anomaly scores usually means your input data changed, which is itself the alert you wanted.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges', 'total_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.pipeline import make_pipeline
11 from sklearn.impute import SimpleImputer
12 from sklearn.preprocessing import StandardScaler
13
14 prep = make_pipeline(SimpleImputer(strategy='median'), StandardScaler())
15 Z = prep.fit_transform(df[NUM])
16 from sklearn.ensemble import IsolationForest
17 import numpy as np
18
19 rng = np.random.default_rng(0)
20
21 # Inject twenty customers with plausible-looking but impossible profiles:
22 # very short tenure, very high total charges.
23 fake = Z[rng.choice(len(Z), 20, replace=False)].copy()
24 fake[:, 0] = -1.5      # tenure well below average
25 fake[:, 3] = 4.0      # total charges far above
26 mixed = np.vstack([Z, fake])
27 truth = np.r_[np.zeros(len(Z)), np.ones(20)]
28
29 iso = IsolationForest(random_state=0).fit(mixed)
30 scores = -iso.score_samples(mixed)

```

```

31 top = np.argsort(scores)[: -1][ :20]
32 print('injected anomalies found in the top 20: %d of 20'
33       % int(truth[top].sum()))
34 from sklearn.metrics import average_precision_score
35 print('average precision: %.4f' % average_precision_score(truth, scores))

```

```

injected anomalies found in the top 20: 0 of 20
average precision: 0.5345

```

Day 6 takeaway

Day 7 Reduction and Detection Together

Sixty correlated columns, a rare event, and the supervised ceiling

Reduction and detection, applied together, on the problem they were invented for: many correlated columns and a rare event.

1. A wide, correlated dataset

```

1 import numpy as np
2 from sklearn.datasets import make_classification
3
4 rng = np.random.default_rng(0)
5 X, y = make_classification(n_samples=5000, n_features=60, n_informative=8,
6                           n_redundant=30, weights=[0.98], flip_y=0.0,
7                           random_state=0)
8 print('shape %s, positive rate %.3f' % (X.shape, y.mean()))
9
10 corr = np.corrcoef(X.T)
11 off = corr[np.triu_indices_from(corr, k=1)]
12 print('feature pairs correlated above 0.8: %d' % (np.abs(off) > 0.8).sum()))

```

```

shape (5000, 60), positive rate 0.020
feature pairs correlated above 0.8: 13

```

2. How many dimensions does it really have?

```

1 import numpy as np
2 from sklearn.datasets import make_classification
3 from sklearn.decomposition import PCA
4 from sklearn.preprocessing import StandardScaler
5
6 X, y = make_classification(n_samples=5000, n_features=60, n_informative=8,
7                           n_redundant=30, weights=[0.98], flip_y=0.0,
8                           random_state=0)
9 Xs = StandardScaler().fit_transform(X)
10 cum = np.cumsum(PCA().fit(Xs).explained_variance_ratio_)

```

```

11 for target in [0.80, 0.90, 0.95, 0.99]:
12     print('%0.0f%% of variance in %2d of 60 components'
13           % (100 * target, int(np.searchsorted(cum, target)) + 1))

```

```

80% of variance in 17 of 60 components
90% of variance in 24 of 60 components
95% of variance in 27 of 60 components
99% of variance in 30 of 60 components

```

3. Detect, with and without reduction

```

1 import numpy as np
2 from sklearn.datasets import make_classification
3 from sklearn.decomposition import PCA
4 from sklearn.preprocessing import StandardScaler
5 from sklearn.ensemble import IsolationForest
6 from sklearn.pipeline import make_pipeline
7 from sklearn.metrics import average_precision_score, roc_auc_score
8
9 X, y = make_classification(n_samples=5000, n_features=60, n_informative=8,
10                           n_redundant=30, weights=[0.98], flip_y=0.0,
11                           random_state=0)
12 Xs = StandardScaler().fit_transform(X)
13
14 print('%22s %10s %14s' % ('', 'ROC AUC', 'avg precision'))
15 print('%22s %10.4f %14.4f'
16       % ('base rate', 0.5, y.mean()))
17
18 iso = IsolationForest(random_state=0).fit(Xs)
19 s = -iso.score_samples(Xs)
20 print('%22s %10.4f %14.4f' % ('all 60 dimensions', roc_auc_score(y, s),
21                               average_precision_score(y, s)))
22
23 for k in [5, 10, 20]:
24     Xr = PCA(n_components=k, random_state=0).fit_transform(Xs)
25     iso = IsolationForest(random_state=0).fit(Xr)
26     s = -iso.score_samples(Xr)
27     print('%22s %10.4f %14.4f' % ('PCA to %d' % k, roc_auc_score(y, s),
28                                   average_precision_score(y, s)))

```

	ROC AUC	avg precision
base rate	0.5000	0.0200
all 60 dimensions	0.7843	0.0718
PCA to 5	0.6409	0.0559
PCA to 10	0.8265	0.2426
PCA to 20	0.7948	0.1684

Reduction helps detection more than it helps prediction

Week 8 found feature engineering barely moved a supervised model. Here, cutting 60 correlated dimensions to 10 more than triples average precision, from 0.07 to 0.24. Note that 5 components is *worse* than doing nothing, compress too far and you discard the structure along with the

redundancy. The gain comes because Isolation Forest cuts on *random* features, and when fifty of your sixty columns are near-duplicates of each other, most random cuts are wasted on redundancy. Removing the redundancy concentrates the algorithm's effort where the structure is.

4. Compare against the supervised ceiling

```
1 import numpy as np
2 from sklearn.datasets import make_classification
3 from sklearn.preprocessing import StandardScaler
4 from sklearn.linear_model import LogisticRegression
5 from sklearn.model_selection import cross_val_predict, StratifiedKFold
6 from sklearn.metrics import average_precision_score
7
8 X, y = make_classification(n_samples=5000, n_features=60, n_informative=8,
9                           n_redundant=30, weights=[0.98], flip_y=0.0,
10                          random_state=0)
11 Xs = StandardScaler().fit_transform(X)
12 cv = StratifiedKFold(5, shuffle=True, random_state=0)
13 proba = cross_val_predict(LogisticRegression(max_iter=2000), Xs, y, cv=cv,
14                           method='predict_proba')[:, 1]
15 print('supervised, with labels: avg precision %.4f'
16       % average_precision_score(y, proba))
17 print('base rate:                %.4f' % y.mean())
```

```
supervised, with labels: avg precision 0.8278
base rate:                0.0200
```

Watch Out

Labels are worth more than any algorithm

The supervised model, given the same features, is in a different league. That comparison is worth running whenever you are tempted to reach for anomaly detection: if you can get labels for even a few hundred cases, doing so will beat any amount of unsupervised cleverness. Anomaly detection is for when labels genuinely do not exist. A new failure mode, a fraud pattern nobody has seen yet.

5. What to put into production

1. Fit the scaler and PCA on a period you believe was clean.
2. Store both, with the data version and the date, week 15 covers how.
3. Score new rows and alert on the top fraction your team can investigate.
4. Record every investigated case and its outcome. Within months you have labels, and you should switch to a supervised model.
5. Monitor the score distribution. A shift means the input data changed, which is worth knowing regardless of any individual alert.

Your assignment

Run the day 5 Isolation Forest on the customer data, take the top 20 and read them. Decide for each whether it is genuinely odd or merely uncommon. There is a difference, and only the first is worth anyone's time. Then inject 20 fake customers with an impossible combination of your own choosing and see how many surface.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. PCA finds the directions that:
 - a) Capture the greatest variance in the features
 - b) Separate the classes
 - c) Minimise correlation with the target
 - d) Best predict the target
2. PCA must be preceded by scaling because:
 - a) It is faster
 - b) Otherwise a feature measured in large units dominates the variance purely because of its units
 - c) It cannot handle negatives
 - d) The components must be positive
3. Which of these should you *not* read from a t-SNE plot?
 - a) That structure exists
 - b) That points formed groups
 - c) The distance between two clusters as a measure of how different they are
 - d) That a point sits apart from others
4. NMF is often preferred to PCA on text because:
 - a) It needs no scaling
 - b) It handles missing data
 - c) It is faster
 - d) Its non-negative, additive components read as topics
5. Isolation Forest scores a point as anomalous when:
 - a) It takes few random splits to isolate it
 - b) Its density is high
 - c) It has missing values
 - d) It is far from the mean
6. Local Outlier Factor differs from Isolation Forest in that it:

- a) Is supervised
 - b) Compares each point's density with that of its neighbours, so it can find local anomalies a global method misses
 - c) Requires labels
 - d) Only works in two dimensions
7. An anomaly detector generally performs better when fitted on:
- a) A random sample
 - b) All available data
 - c) Data known to be clean, so "normal" is defined by normal rows
 - d) Only the anomalies
8. Reconstruction error from PCA can be used as an anomaly score because:
- a) It equals the distance to the mean
 - b) It is normally distributed
 - c) It is always large
 - d) A row that cannot be rebuilt from the main components is a row that violates the structure most rows follow

Answers: 1a, 2b, 3c, 4d, 5a, 6b, 7c, 8d. Anything you got wrong points at the day to re-read, not at a gap in ability.

Neural Networks From Scratch, Then Keras

Week 11

Build a network in numpy, derive backpropagation, then do the same in TensorFlow and find out where it wins.

OBJECTIVES

Deep learning is week 3's matrix multiplication, stacked, with a nonlinearity in between. This week builds that up from a single neuron that cannot solve XOR to a trained Keras model, deriving backpropagation on the way so the framework is never a black box, and ends by measuring the result honestly against logistic regression.

Day 1 From a Neuron to a Network

The perceptron, why XOR broke it, and what a hidden layer actually does

A neural network is not a new idea bolted onto machine learning. It is week 3's $X @ w + b$, stacked, with a nonlinearity between the layers. Everything else is engineering.

The perceptron

```

1 import numpy as np
2
3 def step(z):
4     return (z > 0).astype(float)
5
6 # AND: fire only when both inputs are 1.
7 X = np.array([[0, 0], [0, 1], [1, 0], [1, 1]], dtype=float)
8 y_and = np.array([0, 0, 0, 1], dtype=float)
9
10 w = np.zeros(2)
11 b = 0.0
12 for epoch in range(20):
13     for xi, target in zip(X, y_and):
14         error = target - step(xi @ w + b)
15         w += 0.1 * error * xi
16         b += 0.1 * error
17
18 print('weights %s, bias %.2f' % (w.round(2), b))
19 print('predictions %s' % step(X @ w + b))
20 print('targets %s' % y_and)

```

```
weights [0.2 0.1], bias -0.20
predictions [0. 0. 0. 1.]
targets      [0. 0. 0. 1.]
```

A single neuron, trained by nudging the weights whenever it is wrong. It learned AND in a few passes. It will learn OR just as easily.

And then it fails

```
1 import numpy as np
2
3 def step(z):
4     return (z > 0).astype(float)
5
6 X = np.array([[0, 0], [0, 1], [1, 0], [1, 1]], dtype=float)
7 y_xor = np.array([0, 1, 1, 0], dtype=float) # exclusive or
8
9 w, b = np.zeros(2), 0.0
10 for epoch in range(500):
11     for xi, target in zip(X, y_xor):
12         error = target - step(xi @ w + b)
13         w += 0.1 * error * xi
14         b += 0.1 * error
15
16 print('after 500 epochs: predictions %s' % step(X @ w + b))
17 print('targets %s' % y_xor)
18 print('\nit never converges, and it never will')
```

```
after 500 epochs: predictions [1. 1. 0. 0.]
targets                    [0. 1. 1. 0.]

it never converges, and it never will
```

Watch Out

One neuron draws one straight line

XOR needs the output to be 1 in two opposite corners of the square and 0 in the other two. No single straight line separates those. This limitation, published in 1969, stopped neural network research for roughly fifteen years, and the fix, a second layer, was already known. What was missing was an efficient way to train it.

Two layers solve it

```
1 import numpy as np
2
3 def relu(z):
4     return np.maximum(0, z)
5
6 X = np.array([[0, 0], [0, 1], [1, 0], [1, 1]], dtype=float)
7
```

```

8 # Hand-chosen weights, to show that a solution exists.
9 W1 = np.array([[1.0, 1.0], [1.0, 1.0]])
10 b1 = np.array([0.0, -1.0])
11 W2 = np.array([[1.0], [-2.0]])
12 b2 = np.array([0.0])
13
14 hidden = relu(X @ W1 + b1)
15 out = hidden @ W2 + b2
16
17 print('%12s %18s %10s' % ('input', 'hidden', 'output'))
18 for xi, h, o in zip(X, hidden, out.ravel()):
19     print('%12s %18s %10.1f' % (xi, h, o))
20 print('\ntarget                                [0, 1, 1, 0]')

```

input	hidden	output
[0. 0.]	[0. 0.]	0.0
[0. 1.]	[1. 0.]	1.0
[1. 0.]	[1. 0.]	1.0
[1. 1.]	[2. 1.]	0.0

target	[0, 1, 1, 0]
--------	--------------

The first hidden unit computes “at least one input is on”, the second computes “both are on”, and the output subtracts twice the second from the first. The hidden layer built a representation in which the problem became linearly separable. That is what hidden layers are for.

Why the nonlinearity is essential

```

1 import numpy as np
2
3 rng = np.random.default_rng(0)
4 X = rng.normal(size=(5, 4))
5 W1 = rng.normal(size=(4, 8))
6 W2 = rng.normal(size=(8, 3))
7
8 two_layers = (X @ W1) @ W2      # no activation between them
9 one_layer = X @ (W1 @ W2)      # collapsed into a single matrix
10
11 print('identical:', np.allclose(two_layers, one_layer))
12 print('so without an activation, depth buys you nothing at all.')

```



```

identical: True
so without an activation, depth buys you nothing at all.

```

Stacked linear layers are one linear layer

Matrix multiplication is associative, so a hundred layers with no activation collapse to a single matrix. The activation function is not a detail or a tuning choice. It is the entire reason depth means anything.

The activation functions

```

1 import numpy as np
2
3 def sigmoid(z):
4     return 1 / (1 + np.exp(-z))
5
6 def tanh(z):
7     return np.tanh(z)
8
9 def relu(z):
10    return np.maximum(0, z)
11
12 def leaky_relu(z, a=0.01):
13    return np.where(z > 0, z, a * z)
14
15 zs = np.array([-3.0, -0.5, 0.0, 0.5, 3.0])
16 print('%8s %10s %10s %8s %12s' % ('z', 'sigmoid', 'tanh', 'relu', 'leaky_relu'))
17 for z in zs:
18     print('%8.1f %10.4f %10.4f %8.1f %12.4f'
19           % (z, sigmoid(z), tanh(z), relu(z), leaky_relu(z)))

```

z	sigmoid	tanh	relu	leaky_relu
-3.0	0.0474	-0.9951	0.0	-0.0300
-0.5	0.3775	-0.4621	0.0	-0.0050
0.0	0.5000	0.0000	0.0	0.0000
0.5	0.6225	0.4621	0.5	0.5000
3.0	0.9526	0.9951	3.0	3.0000

Activation	Range	Gradient issue	Use for
Sigmoid	0 to 1	Saturates; derivative caps at 0.25	Binary output layer only
Tanh	-1 to 1	Saturates, but centred on zero	Rarely; recurrent layers sometimes
ReLU	0 upward	Dead units when input stays negative	The default for hidden layers
Leaky ReLU	All reals	Fixes dead units	When ReLU units are dying

Week 3 showed the sigmoid's derivative never exceeding 0.25. Multiply several of those together through a deep network and the gradient reaching the first layer effectively vanishes. ReLU's derivative is exactly 1 wherever the input is positive, which is why replacing sigmoid with ReLU in hidden layers made deep networks trainable.

Day 1 takeaway

Day 2 Backpropagation, Implemented

The chain rule backwards, gradient checking, and why initialisation scale matters

Yesterday's two-layer network had weights chosen by hand. Today you derive how to learn them, which is the chain rule from week 3, applied layer by layer.

Forward, then backward

1. **Forward pass:** push the input through each layer and compute the loss.
2. **Backward pass:** work out how much each weight contributed to that loss, from the output backwards.
3. **Update:** step every weight against its gradient.

A complete network, in numpy

```
1 import numpy as np
2
3 def sigmoid(z):
4     return 1 / (1 + np.exp(-np.clip(z, -60, 60)))
5
6 rng = np.random.default_rng(0)
7 X = np.array([[0, 0], [0, 1], [1, 0], [1, 1]], dtype=float)
8 y = np.array([[0.0], [1.0], [1.0], [0.0]]) # XOR again
9
10 # He initialisation: scale by sqrt(2/fan_in), which keeps ReLU signals alive.
11 W1 = rng.normal(0, np.sqrt(2 / 2), (2, 4))
12 b1 = np.zeros(4)
13 W2 = rng.normal(0, np.sqrt(2 / 4), (4, 1))
14 b2 = np.zeros(1)
15 lr = 0.5
16
17 for epoch in range(1, 4001):
18     # ---- forward
19     z1 = X @ W1 + b1
20     a1 = np.maximum(0, z1) # relu
21     z2 = a1 @ W2 + b2
22     out = sigmoid(z2)
23     loss = -np.mean(y * np.log(out + 1e-9) + (1 - y) * np.log(1 - out + 1e-9))
24
25     # ---- backward
26     # For sigmoid output with log loss, the gradient simplifies to (out - y).
27     d2 = (out - y) / len(X)
28     dW2 = a1.T @ d2
29     db2 = d2.sum(axis=0)
30     d1 = (d2 @ W2.T) * (z1 > 0) # chain rule through the relu
31     dW1 = X.T @ d1
32     db1 = d1.sum(axis=0)
33
34     # ---- update
35     W2 -= lr * dW2; b2 -= lr * db2
36     W1 -= lr * dW1; b1 -= lr * db1
37
38     if epoch in (1, 100, 1000, 4000):
39         print('epoch %5d loss %.6f' % (epoch, loss))
40
41 print('\npredictions %s' % out.ravel().round(3))
42 print('targets %s' % y.ravel())
```

```

epoch      1  loss 0.774224
epoch     100  loss 0.210313
epoch    1000  loss 0.004712
epoch    4000  loss 0.000931

predictions [0.003 1.    1.    0.   ]
targets      [0.  1.  1.  0.]

```

Forty lines, no framework, and it learned XOR from random initial weights. Every deep learning library is this loop, made fast and general.

Why (out - y) and not something messier

The derivative of log loss with respect to the output, multiplied by the derivative of the sigmoid, simplifies to exactly $\text{out} - y$. The two awkward terms cancel. This is not a coincidence. It is why log loss is paired with sigmoid, and cross-entropy with softmax. Pair them differently and you get a messier gradient that trains worse.

Check your gradients numerically

The most useful debugging technique for anyone implementing this: compare your analytical gradient against a numerical one.

```

1 import numpy as np
2
3 rng = np.random.default_rng(0)
4 X = rng.normal(size=(20, 3))
5 y = (X[:, 0] + X[:, 1] > 0).astype(float).reshape(-1, 1)
6 W = rng.normal(size=(3, 1)) * 0.1
7
8 def loss_of(W):
9     out = 1 / (1 + np.exp(-(X @ W)))
10    return -np.mean(y * np.log(out + 1e-9) + (1 - y) * np.log(1 - out + 1e-9))
11
12 out = 1 / (1 + np.exp(-(X @ W)))
13 analytic = X.T @ (out - y) / len(X)
14
15 numeric = np.zeros_like(W)
16 h = 1e-6
17 for i in range(W.shape[0]):
18     up, down = W.copy(), W.copy()
19     up[i, 0] += h
20     down[i, 0] -= h
21     numeric[i, 0] = (loss_of(up) - loss_of(down)) / (2 * h)
22
23 print('analytic %s' % analytic.ravel().round(8))
24 print('numeric %s' % numeric.ravel().round(8))
25 print('max difference %.2e' % np.abs(analytic - numeric).max())

```

```

analytic [-0.44631288 -0.20424764 -0.04042924]
numeric  [-0.44631288 -0.20424763 -0.04042924]
max difference 9.48e-10

```

⚠ Watch Out

Anything above about 1e-6 means a bug

A gradient that is subtly wrong still trains, just badly, so this class of bug is invisible without the check. Run it once on a tiny network with random data before you trust any hand-written backpropagation. Frameworks compute gradients automatically, which is the main reason to use one.

Initialisation matters more than it looks

```
1 import numpy as np
2
3 rng = np.random.default_rng(0)
4
5 def signal_through(scale, depth=12, width=100):
6     a = rng.normal(size=(200, width))
7     for _ in range(depth):
8         W = rng.normal(0, scale, (width, width))
9         a = np.maximum(0, a @ W)
10    return a.std()
11
12 width = 100
13 print('%-26s %16s' % ('initialisation', 'std after 12 layers'))
14 for name, scale in [('too small (0.01)', 0.01),
15                     ('Xavier sqrt(1/n)', np.sqrt(1 / width)),
16                     ('He sqrt(2/n)', np.sqrt(2 / width)),
17                     ('too large (0.3)', 0.3)]:
18     print('%-26s %16.6g' % (name, signal_through(scale)))
```

initialisation	std after 12 layers
too small (0.01)	1.20237e-14
Xavier sqrt(1/n)	0.0231427
He sqrt(2/n)	0.720329
too large (0.3)	7450.56

Too small and the signal decays to nothing by the twelfth layer; too large and it explodes. He initialisation, which scales by the square root of two over the number of inputs, keeps the variance roughly constant through ReLU layers, the factor of two compensates for ReLU discarding half the values. Keras uses it by default, which is why you rarely think about it.

Day 2 takeaway

Day 3 Your First Keras Model

The same network in eight lines, and how to read a model summary

You have written a network. Now use a framework, and see what it gives you in exchange for the loop you just wrote.

The same XOR, in Keras

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7
8 X = np.array([[0, 0], [0, 1], [1, 0], [1, 1]], dtype='float32')
9 y = np.array([0, 1, 1, 0], dtype='float32')
10
11 model = tf.keras.Sequential([
12     tf.keras.layers.Input(shape=(2,)),
13     tf.keras.layers.Dense(4, activation='relu'),
14     tf.keras.layers.Dense(1, activation='sigmoid'),
15 ])
16 model.compile(optimizer=tf.keras.optimizers.Adam(0.05),
17               loss='binary_crossentropy', metrics=['accuracy'])
18
19 hist = model.fit(X, y, epochs=500, verbose=0)
20 print('final loss %.6f' % hist.history['loss'][-1])
21 print('predictions %s' % model.predict(X, verbose=0).ravel().round(3))
22 print('targets      %s' % y)

```

```

final loss 0.477525
predictions [0.334 0.999 0.334 0.334]
targets      [0. 1. 1. 0.]

```

Eight lines against forty, automatic gradients, and a better optimiser. That is the trade: you give up seeing the arithmetic and you gain everything else.

Reading a model summary

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 model = tf.keras.Sequential([
7     tf.keras.layers.Input(shape=(15,)),
8     tf.keras.layers.Dense(32, activation='relu'),
9     tf.keras.layers.Dense(16, activation='relu'),
10    tf.keras.layers.Dense(1, activation='sigmoid'),
11 ])
12 model.summary()

```

```

Model: "sequential"
+-----+-----+-----+
| Layer (type)                | Output Shape          | Param # |
+-----+-----+-----+
| dense (Dense)                | (None, 32)            | 512      |
+-----+-----+-----+

```

```

| dense_1 (Dense)          | (None, 16)          | 528 |
+-----+-----+-----+
| dense_2 (Dense)          | (None, 1)           | 17  |
+-----+-----+-----+
Total params: 1,057 (4.13 KB)
Trainable params: 1,057 (4.13 KB)
Non-trainable params: 0 (0.00 B)

```

```

1 print('first layer: 15 inputs x 32 units + 32 biases = %d' % (15 * 32 + 32))
2 print('second layer: 32 inputs x 16 units + 16 biases = %d' % (32 * 16 + 16))
3 print('output layer: 16 inputs x 1 unit + 1 bias = %d' % (16 * 1 + 1))
4 print('total'                                     = %d'
5       % (15 * 32 + 32 + 32 * 16 + 16 + 17))

```

```

first layer: 15 inputs x 32 units + 32 biases = 512
second layer: 32 inputs x 16 units + 16 biases = 528
output layer: 16 inputs x 1 unit + 1 bias = 17
total = 1057

```

Count the parameters against your row count

1,073 parameters and 2,250 training rows is about two rows per parameter, which is thin. A rough sanity rule for tabular data is that you want at least ten rows per parameter before you start trusting the network, and far more if the signal is weak. This calculation takes ten seconds and saves days.

The three arguments to compile

Argument	Choose	For
loss	binary_crossentropy	Two classes, sigmoid output
	categorical_crossentropy	Several classes, one-hot labels, softmax
	sparse_categorical_crossentropy	Several classes, integer labels
	mse or mae	Regression
optimizer	adam	The default that almost always works
metrics	AUC, accuracy	Reported but never optimised

⚠ Watch Out

A metric is not a loss

Metrics are computed and printed; the loss is what gradients come from. Adding `metrics=['AUC']` does not make the network optimise AUC. If you need a different objective you have to change the loss, and most ranking metrics are not differentiable, which is why everyone optimises cross-entropy and then tunes the threshold afterwards, exactly as week 5 did.

On the churn data

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7 import pandas as pd
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 from sklearn.metrics import roc_auc_score
37
38 model = tf.keras.Sequential([
39     tf.keras.layers.Input(shape=(Xt.shape[1],)),
40     tf.keras.layers.Dense(32, activation='relu'),
41     tf.keras.layers.Dense(16, activation='relu'),
42     tf.keras.layers.Dense(1, activation='sigmoid'),
43 ])
44 model.compile(optimizer='adam', loss='binary_crossentropy')
45 hist = model.fit(Xt, yt, epochs=40, batch_size=64,
46                 validation_split=0.2, verbose=0)
47
48 proba = model.predict(Xv, verbose=0).ravel()
49 print('training loss    %.4f' % hist.history['loss'][-1])
50 print('validation loss  %.4f' % hist.history['val_loss'][-1])
51 print('test ROC AUC     %.4f' % roc_auc_score(yv, proba))
52 print('\nlogistic regression got 0.8174 on this split.')

```

```
training loss    0.4037
```

```
validation loss 0.4526
test ROC AUC    0.8050

logistic regression got 0.8174 on this split.
```

A neural network, several hundred times the parameters, landing in the same place. Hold that thought until week 12 day 7, which takes it seriously.

Day 3 takeaway

Day 4 Optimisers and Learning Rates

Momentum, Adam, and what adaptive really means

Plain gradient descent works. Every optimiser since exists because it works slowly, and because the learning rate is painful to choose.

Momentum

```
1 import numpy as np
2
3 # A ravine: steep across, shallow along. Gradient descent zig-zags.
4 def loss(w):
5     return 0.5 * (w[0] ** 2 * 20 + w[1] ** 2)
6
7 def grad(w):
8     return np.array([20 * w[0], w[1]])
9
10 for name, momentum in [('plain', 0.0), ('momentum', 0.9)]:
11     w = np.array([1.0, 1.0])
12     v = np.zeros(2)
13     for _ in range(60):
14         v = momentum * v - 0.05 * grad(w)
15         w = w + v
16     print('%s after 60 steps: loss %.6f position %s'
17           % (name, loss(w), w.round(4)))
```

```
plain      after 60 steps: loss 0.001061 position [0.      0.0461]
momentum   after 60 steps: loss 0.009800 position [-0.0305  0.0312]
```

Momentum accumulates a running average of the gradient. Across the ravine the gradient keeps flipping sign and cancels; along it the gradient is consistent and builds up. The result is less zig-zag and faster progress.

Adam

```
1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
```

```

5  tf.get_logger().setLevel('ERROR')
6  import numpy as np
7  import pandas as pd
8
9  df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 from sklearn.metrics import roc_auc_score
37
38 def build():
39     return tf.keras.Sequential([
40         tf.keras.layers.Input(shape=(Xt.shape[1],)),
41         tf.keras.layers.Dense(32, activation='relu'),
42         tf.keras.layers.Dense(16, activation='relu'),
43         tf.keras.layers.Dense(1, activation='sigmoid'),
44     ])
45
46 optimisers = {
47     'SGD': tf.keras.optimizers.SGD(0.01),
48     'SGD+momentum': tf.keras.optimizers.SGD(0.01, momentum=0.9),
49     'RMSprop': tf.keras.optimizers.RMSprop(0.001),
50     'Adam': tf.keras.optimizers.Adam(0.001),
51 }
52 for name, opt in optimisers.items():
53     tf.random.set_seed(42)
54     m = build()
55     m.compile(optimizer=opt, loss='binary_crossentropy')
56     h = m.fit(Xt, yt, epochs=25, batch_size=64, verbose=0)
57     print('%s final loss %.4f test AUC %.4f'
58           % (name, h.history['loss'][-1],

```

```
59 roc_auc_score(yv, m.predict(Xv, verbose=0).ravel()))
```

SGD	final loss	0.4492	test AUC	0.8100
SGD+momentum	final loss	0.4210	test AUC	0.8157
RMSprop	final loss	0.4193	test AUC	0.8154
Adam	final loss	0.4180	test AUC	0.8137

The learning rate still matters

```
1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7 import pandas as pd
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 from sklearn.metrics import roc_auc_score
37
38 print('%10s %14s %12s' % ('lr', 'final loss', 'test AUC'))
39 for lr in [0.0001, 0.001, 0.01, 0.1, 1.0]:
40     tf.random.set_seed(42)
41     m = tf.keras.Sequential([
42         tf.keras.layers.Input(shape=(Xt.shape[1],)),
43         tf.keras.layers.Dense(32, activation='relu'),
```

```

44     tf.keras.layers.Dense(1, activation='sigmoid'),
45 ]
46 m.compile(optimizer=tf.keras.optimizers.Adam(lr),
47           loss='binary_crossentropy')
48 h = m.fit(Xt, yt, epochs=25, batch_size=64, verbose=0)
49 print('%10s %14.4f %12.4f'
50       % (lr, h.history['loss'][-1],
51          roc_auc_score(yv, m.predict(Xv, verbose=0).ravel())))

```

lr	final loss	test AUC
0.0001	0.5349	0.7365
0.001	0.4218	0.8121
0.01	0.4128	0.8099
0.1	0.4235	0.8041
1.0	0.5750	0.5519

⚠ Watch Out

Adam adapts per parameter, not the global scale

It is often described as removing the need to tune the learning rate. It does not. It removes the need to tune it *per parameter*. Set it to 1.0 and the model still falls apart. Adam's default of 0.001 is a good starting point and the value most worth trying next is one order of magnitude either side.

Schedules

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7 import pandas as pd
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                      ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                      ('o', OneHotEncoder(handle_unknown='ignore'),

```

```

26         sparse_output=False))]), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 from sklearn.metrics import roc_auc_score
37 import numpy as np
38
39 tf.random.set_seed(42)
40 schedule = tf.keras.optimizers.schedules.ExponentialDecay(
41     initial_learning_rate=0.01, decay_steps=100, decay_rate=0.9)
42
43 m = tf.keras.Sequential([
44     tf.keras.layers.Input(shape=(Xt.shape[1],)),
45     tf.keras.layers.Dense(32, activation='relu'),
46     tf.keras.layers.Dense(1, activation='sigmoid'),
47 ])
48 m.compile(optimizer=tf.keras.optimizers.Adam(schedule),
49           loss='binary_crossentropy')
50 h = m.fit(Xt, yt, epochs=40, batch_size=64, verbose=0)
51
52 print('decaying schedule: final loss %.4f    test AUC %.4f'
53       % (h.history['loss'][-1],
54          roc_auc_score(yv, m.predict(Xv, verbose=0).ravel())))
55 print('\nlearning rate at step 0, 200, 800:')
56 for step in [0, 200, 800]:
57     print('  %4d  %.6f' % (step, float(schedule(step))))

```

```

decaying schedule: final loss 0.4039    test AUC 0.8148

learning rate at step 0, 200, 800:
  0  0.010000
 200 0.008100
 800 0.004305

```

Start large to cover ground, finish small to settle. A decaying rate is the standard recipe for long training runs, and `ReduceLROnPlateau`, which day 6 covers, does the same thing reactively, cutting the rate whenever progress stalls.

Day 4 takeaway

Day 5 Stopping a Network Memorising

Dropout, weight decay and the early stopping argument everyone forgets

A network with enough parameters will memorise its training set. Today, the four things that

stop it.

Watch it overfit first

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7 import pandas as pd
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 from sklearn.metrics import roc_auc_score
37
38 tf.random.set_seed(42)
39 big = tf.keras.Sequential([
40     tf.keras.layers.Input(shape=(Xt.shape[1],)),
41     tf.keras.layers.Dense(256, activation='relu'),
42     tf.keras.layers.Dense(256, activation='relu'),
43     tf.keras.layers.Dense(256, activation='relu'),
44     tf.keras.layers.Dense(1, activation='sigmoid'),
45 ])
46 big.compile(optimizer='adam', loss='binary_crossentropy')
47 h = big.fit(Xt, yt, epochs=60, batch_size=32, validation_split=0.2, verbose=0)
48
49 print('parameters: %d for %d training rows' % (big.count_params(), len(Xt)))

```

```

50 print('%8s %12s %12s' % ('epoch', 'train loss', 'val loss'))
51 for e in [0, 9, 29, 59]:
52     print('%8d %12.4f %12.4f'
53           % (e + 1, h.history['loss'][e], h.history['val_loss'][e]))
54 print('\ntest AUC %.4f' % roc_auc_score(yv, big.predict(Xv, verbose=0).ravel()))

```

```
parameters: 135937 for 2250 training rows
```

epoch	train loss	val loss
1	0.4699	0.4433
10	0.3929	0.5115
30	0.2770	0.8956
60	0.1921	1.4788

```
test AUC 0.7818
```

Training loss falls steadily; validation loss bottoms out early and then climbs. More parameters than training rows, and the network is recording the data rather than learning from it.

Dropout

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7 import pandas as pd
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')

```

```

34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 from sklearn.metrics import roc_auc_score
37
38 print('%10s %12s %12s %10s' % ('dropout', 'train loss', 'val loss', 'test AUC'))
39 for rate in [0.0, 0.2, 0.5]:
40     tf.random.set_seed(42)
41     m = tf.keras.Sequential([
42         tf.keras.layers.Input(shape=(Xt.shape[1],)),
43         tf.keras.layers.Dense(256, activation='relu'),
44         tf.keras.layers.Dropout(rate),
45         tf.keras.layers.Dense(256, activation='relu'),
46         tf.keras.layers.Dropout(rate),
47         tf.keras.layers.Dense(1, activation='sigmoid'),
48     ])
49     m.compile(optimizer='adam', loss='binary_crossentropy')
50     h = m.fit(Xt, yt, epochs=60, batch_size=32, validation_split=0.2, verbose=0)
51     print('%10.1f %12.4f %12.4f %10.4f'
52           % (rate, h.history['loss'][-1], h.history['val_loss'][-1],
53              roc_auc_score(yv, m.predict(Xv, verbose=0).ravel())))

```

dropout	train loss	val loss	test AUC
0.0	0.2726	1.1482	0.7817
0.2	0.3219	0.6458	0.7842
0.5	0.3861	0.5365	0.8009

⚠ Watch Out

Dropout is active in training and off in prediction

Keras handles the switch for you. It is worth knowing because it explains something that confuses people: your training loss can be *worse* than your validation loss with heavy dropout, because the training figure is measured on a crippled network and the validation figure on the whole one.

Weight decay

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7 import pandas as pd
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer

```

```

16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 from sklearn.metrics import roc_auc_score
37
38 print('%12s %12s %12s %10s' % ('l2', 'train loss', 'val loss', 'test AUC'))
39 for l2 in [0.0, 0.001, 0.01]:
40     tf.random.set_seed(42)
41     reg = tf.keras.regularizers.l2(l2) if l2 else None
42     m = tf.keras.Sequential([
43         tf.keras.layers.Input(shape=(Xt.shape[1],)),
44         tf.keras.layers.Dense(256, activation='relu', kernel_regularizer=reg),
45         tf.keras.layers.Dense(256, activation='relu', kernel_regularizer=reg),
46         tf.keras.layers.Dense(1, activation='sigmoid'),
47     ])
48     m.compile(optimizer='adam', loss='binary_crossentropy')
49     h = m.fit(Xt, yt, epochs=60, batch_size=32, validation_split=0.2, verbose=0)
50     print('%12s %12.4f %12.4f %10.4f'
51           % (l2, h.history['loss'][-1], h.history['val_loss'][-1],
52             roc_auc_score(yv, m.predict(Xv, verbose=0).ravel())))

```

l2	train loss	val loss	test AUC
0.0	0.2686	0.7303	0.7802
0.001	0.4138	0.5613	0.8106
0.01	0.4426	0.4747	0.8161

The same L2 penalty as ridge regression in week 4, applied to every weight matrix. The mechanism is identical: a weight has to reduce the loss by more than it adds to the penalty.

Early stopping, which is the one you will always use

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')

```

```

6 import numpy as np
7 import pandas as pd
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 from sklearn.metrics import roc_auc_score
37
38 tf.random.set_seed(42)
39 m = tf.keras.Sequential([
40     tf.keras.layers.Input(shape=(Xt.shape[1],)),
41     tf.keras.layers.Dense(256, activation='relu'),
42     tf.keras.layers.Dense(256, activation='relu'),
43     tf.keras.layers.Dense(1, activation='sigmoid'),
44 ])
45 m.compile(optimizer='adam', loss='binary_crossentropy')
46 stop = tf.keras.callbacks.EarlyStopping(monitor='val_loss', patience=8,
47                                         restore_best_weights=True)
48 h = m.fit(Xt, yt, epochs=300, batch_size=32, validation_split=0.2,
49         callbacks=[stop], verbose=0)
50
51 print('allowed 300 epochs, ran %d' % len(h.history['loss']))
52 print('best epoch was %d' % (int(min(range(len(h.history['val_loss'])),
53                                       key=lambda i: h.history['val_loss'][i])) + 1))
54 print('test AUC %.4f' % roc_auc_score(yv, m.predict(Xv, verbose=0).ravel()))

```

```

allowed 300 epochs, ran 9
best epoch was 1
test AUC 0.8174

```

restore_best_weights is not the default

Without it, training stops after the patience runs out and you keep the weights from that *last* epoch, which is several epochs past the best one, and therefore worse. Setting it to `True` rewinds to the best validation loss. It is a one-word change that most tutorials omit.

Day 5 takeaway

Day 6 Batch Normalisation and Callbacks

Stable layer inputs, reacting to plateaus, and reading a training curve

Two more techniques that made deep networks practical, and the tools for watching training rather than guessing at it.

Batch normalisation

```
1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7 import pandas as pd
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
```

```

35 yv = y_te.to_numpy().astype('float32')
36 from sklearn.metrics import roc_auc_score
37
38 def build(batchnorm):
39     layers = [tf.keras.layers.Input(shape=(Xt.shape[1],))]
40     for _ in range(6):
41         layers.append(tf.keras.layers.Dense(64, use_bias=not batchnorm))
42         if batchnorm:
43             layers.append(tf.keras.layers.BatchNormalization())
44             layers.append(tf.keras.layers.Activation('relu'))
45         layers.append(tf.keras.layers.Dense(1, activation='sigmoid'))
46     return tf.keras.Sequential(layers)
47
48 print('%14s %10s %14s %10s' % ('', 'lr', 'final loss', 'test AUC'))
49 for bn in [False, True]:
50     for lr in [0.001, 0.02]:
51         tf.random.set_seed(42)
52         m = build(bn)
53         m.compile(optimizer=tf.keras.optimizers.Adam(lr),
54                   loss='binary_crossentropy')
55         h = m.fit(Xt, yt, epochs=30, batch_size=64, verbose=0)
56         print('%14s %10s %14.4f %10.4f'
57               % ('batchnorm' if bn else 'plain', lr,
58                 h.history['loss'][-1],
59                 roc_auc_score(yv, m.predict(Xv, verbose=0).ravel()))

```

	lr	final loss	test AUC
plain	0.001	0.3024	0.7707
plain	0.02	0.3947	0.8026
batchnorm	0.001	0.1113	0.7355
batchnorm	0.02	0.3341	0.7530

⚠ Watch Out

Batch normalisation behaves differently in training and prediction

In training it uses the current batch's statistics; at prediction time it uses a running average accumulated during training. That is why predicting on a single row works at all. It also means a very small batch size makes the statistics noisy and the layer unreliable, below about 16 rows per batch, prefer `LayerNormalization`.

Callbacks: watching, saving and reacting

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7 import pandas as pd
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()

```

```

10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 from sklearn.metrics import roc_auc_score
37 import os
38
39 tf.random.set_seed(42)
40 m = tf.keras.Sequential([
41     tf.keras.layers.Input(shape=(Xt.shape[1],)),
42     tf.keras.layers.Dense(64, activation='relu'),
43     tf.keras.layers.Dense(32, activation='relu'),
44     tf.keras.layers.Dense(1, activation='sigmoid'),
45 ])
46 m.compile(optimizer=tf.keras.optimizers.Adam(0.01),
47           loss='binary_crossentropy')
48
49 callbacks = [
50     tf.keras.callbacks.EarlyStopping(patience=12, restore_best_weights=True),
51     tf.keras.callbacks.ReduceLROnPlateau(factor=0.5, patience=5, min_lr=1e-5),
52     tf.keras.callbacks.ModelCheckpoint('best.keras', save_best_only=True),
53 ]
54 h = m.fit(Xt, yt, epochs=200, batch_size=64, validation_split=0.2,
55         callbacks=callbacks, verbose=0)
56
57 print('epochs run: %d of 200' % len(h.history['loss']))
58 print('learning rate went from %.5f to %.5f'
59       % (h.history['learning_rate'][0], h.history['learning_rate'][-1]))
60 print('checkpoint written:', os.path.exists('best.keras'))
61 print('test AUC %.4f' % roc_auc_score(yv, m.predict(Xv, verbose=0).ravel()))

```

epochs run: 15 of 200

```
learning rate went from 0.01000 to 0.00250
checkpoint written: True
test AUC 0.8134
```

Callback	Does
EarlyStopping	Stops when validation stops improving
ReduceLROnPlateau	Cuts the learning rate when progress stalls
ModelCheckpoint	Writes the best model to disk as it goes
TensorBoard	Logs everything for the live dashboard
CSVLogger	Appends per-epoch metrics to a file

Reading the training curve

```
1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7 import pandas as pd
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 import numpy as np
37
38 tf.random.set_seed(42)
```

```

39 m = tf.keras.Sequential([
40     tf.keras.layers.Input(shape=(Xt.shape[1],)),
41     tf.keras.layers.Dense(128, activation='relu'),
42     tf.keras.layers.Dense(64, activation='relu'),
43     tf.keras.layers.Dense(1, activation='sigmoid'),
44 ])
45 m.compile(optimizer='adam', loss='binary_crossentropy')
46 h = m.fit(Xt, yt, epochs=50, batch_size=32, validation_split=0.2, verbose=0)
47
48 tr = np.array(h.history['loss'])
49 va = np.array(h.history['val_loss'])
50 print('%8s %12s %12s %10s' % ('epoch', 'train', 'val', 'gap'))
51 for e in [0, 4, 9, 19, 34, 49]:
52     print('%8d %12.4f %12.4f %10.4f' % (e + 1, tr[e], va[e], va[e] - tr[e]))
53 print('\nbest validation loss at epoch %d' % (int(va.argmin()) + 1))
54 print('everything after that is overfitting')

```

epoch	train	val	gap
1	0.5349	0.4457	-0.0892
5	0.4210	0.4559	0.0349
10	0.4122	0.4628	0.0505
20	0.3964	0.4713	0.0749
35	0.3687	0.5238	0.1551
50	0.3420	0.5999	0.2579

```

best validation loss at epoch 2
everything after that is overfitting

```

Pattern	Means	Do
Both falling, gap small	Still learning	Train longer
Train falls, validation rises	Overfitting	Early stopping, dropout, fewer units
Both flat and high	Underfitting or the rate is too low	Bigger network, higher rate
Loss becomes nan	Exploded	Lower the rate; check for infinities in your inputs
Validation wildly jumpy	Batch or validation split too small	Increase both

Day 6 takeaway

Day 7 A Complete Keras Workflow

End to end on churn, and an honest comparison against a straight line

A complete Keras workflow on the churn problem, and an honest comparison against everything you built in weeks 5 and 6.

Preprocessing belongs to scikit-learn

Keras has preprocessing layers, but a `ColumnTransformer` is clearer for tabular data and it is what the rest of your pipeline already speaks. Fit it on training data, and keep it. You will need it at prediction time.

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7 import pandas as pd
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 print('features after encoding: %d' % Xt.shape[1])
37 print('train %s test %s' % (Xt.shape, Xv.shape))
38 print('positive rate: train %.3f, test %.3f' % (yt.mean(), yv.mean()))

```

```

features after encoding: 15
train (2250, 15) test (750, 15)
positive rate: train 0.268, test 0.268

```

Build, train, evaluate

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7 import pandas as pd

```

```

8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                           stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 from sklearn.metrics import (roc_auc_score, average_precision_score,
37                              brier_score_loss)
38
39 tf.random.set_seed(42)
40 model = tf.keras.Sequential([
41     tf.keras.layers.Input(shape=(Xt.shape[1],)),
42     tf.keras.layers.Dense(64, activation='relu'),
43     tf.keras.layers.Dropout(0.3),
44     tf.keras.layers.Dense(32, activation='relu'),
45     tf.keras.layers.Dropout(0.3),
46     tf.keras.layers.Dense(1, activation='sigmoid'),
47 ])
48 model.compile(optimizer=tf.keras.optimizers.Adam(0.003),
49              loss='binary_crossentropy',
50              metrics=[tf.keras.metrics.AUC(name='auc')])
51
52 callbacks = [
53     tf.keras.callbacks.EarlyStopping(monitor='val_auc', mode='max',
54                                     patience=20, restore_best_weights=True),
55     tf.keras.callbacks.ReduceLROnPlateau(monitor='val_loss', factor=0.5,
56                                         patience=8, min_lr=1e-5),
57 ]
58 h = model.fit(Xt, yt, epochs=300, batch_size=64, validation_split=0.2,
59              callbacks=callbacks, verbose=0)
60
61 proba = model.predict(Xv, verbose=0).ravel()

```

```

62 print('epochs run      %d' % len(h.history['loss']))
63 print('parameters      %d' % model.count_params())
64 print('test ROC AUC     %.4f' % roc_auc_score(yv, proba))
65 print('average precision %.4f' % average_precision_score(yv, proba))
66 print('Brier score      %.4f' % brier_score_loss(yv, proba))

```

```

epochs run      22
parameters      3137
test ROC AUC     0.8152
average precision 0.5824
Brier score      0.1490

```

Against the alternatives

```

1  import os
2  os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3  import tensorflow as tf
4  tf.random.set_seed(42)
5  tf.get_logger().setLevel('ERROR')
6  import numpy as np
7  import pandas as pd
8
9  df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                      ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                      ('o', OneHotEncoder(handle_unknown='ignore',
26                                         sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                         stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 from sklearn.linear_model import LogisticRegression
37 from sklearn.ensemble import HistGradientBoostingClassifier
38 from sklearn.metrics import (roc_auc_score, average_precision_score,

```

```

39         brier_score_loss)
40 import time
41
42 print('%-16s %10s %14s %10s %10s'
43       % ('model', 'ROC AUC', 'avg precision', 'Brier', 'seconds'))
44
45 for name, clf in [('logistic', LogisticRegression(max_iter=1000, random_state=42)),
46                  ('boosting', HistGradientBoostingClassifier(random_state=42))]:
47     t = time.perf_counter()
48     clf.fit(Xt, yt)
49     el = time.perf_counter() - t
50     pr = clf.predict_proba(Xv)[: , 1]
51     print('%-16s %10.4f %14.4f %10.4f %10.2f'
52          % (name, roc_auc_score(yv, pr), average_precision_score(yv, pr),
53             brier_score_loss(yv, pr), el))
54
55 tf.random.set_seed(42)
56 net = tf.keras.Sequential([
57     tf.keras.layers.Input(shape=(Xt.shape[1],)),
58     tf.keras.layers.Dense(64, activation='relu'),
59     tf.keras.layers.Dropout(0.3),
60     tf.keras.layers.Dense(32, activation='relu'),
61     tf.keras.layers.Dense(1, activation='sigmoid'),
62 ])
63 net.compile(optimizer=tf.keras.optimizers.Adam(0.003), loss='binary_crossentropy')
64 t = time.perf_counter()
65 net.fit(Xt, yt, epochs=100, batch_size=64, validation_split=0.2,
66        callbacks=[tf.keras.callbacks.EarlyStopping(patience=15,
67                                                    restore_best_weights=True)],
68        verbose=0)
69 el = time.perf_counter() - t
70 pr = net.predict(Xv, verbose=0).ravel()
71 print('%-16s %10.4f %14.4f %10.4f %10.2f'
72       % ('neural net', roc_auc_score(yv, pr), average_precision_score(yv, pr),
73          brier_score_loss(yv, pr), el))

```

model	ROC AUC	avg precision	Brier	seconds
logistic	0.8174	0.5993	0.1475	0.02
boosting	0.7796	0.5259	0.1691	2.44
neural net	0.8135	0.5949	0.1491	3.26

Watch Out

Read the last column before you celebrate

Logistic regression fits in a fraction of a second and matches everything else on this data. The network takes orders of magnitude longer, has thousands of parameters instead of fifteen, cannot be explained to a stakeholder, and needs a GPU to train at any serious scale. None of that means networks are bad. It means *tabular data with three thousand rows* is not where they win. Weeks 12 and 13 show what they are actually for.

Saving a model properly

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import tensorflow as tf
4 tf.random.set_seed(42)
5 tf.get_logger().setLevel('ERROR')
6 import numpy as np
7 import pandas as pd
8
9 df = pd.read_csv('customers.csv').drop_duplicates().copy()
10 df['contract'] = df['contract'].str.strip().str.title()
11 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
12
13 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
14 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19 from sklearn.model_selection import train_test_split
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                      ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                      ('o', OneHotEncoder(handle_unknown='ignore',
26                                         sparse_output=False))])), CAT),
27 ])
28
29 X, y = df[NUM + CAT], df['churned']
30 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
31                                          stratify=y, random_state=42)
32 Xt = prep.fit_transform(X_tr).astype('float32')
33 Xv = prep.transform(X_te).astype('float32')
34 yt = y_tr.to_numpy().astype('float32')
35 yv = y_te.to_numpy().astype('float32')
36 import joblib
37 import os
38 import numpy as np
39
40 tf.random.set_seed(42)
41 model = tf.keras.Sequential([
42     tf.keras.layers.Input(shape=(Xt.shape[1],)),
43     tf.keras.layers.Dense(32, activation='relu'),
44     tf.keras.layers.Dense(1, activation='sigmoid'),
45 ])
46 model.compile(optimizer='adam', loss='binary_crossentropy')
47 model.fit(Xt, yt, epochs=15, batch_size=64, verbose=0)
48
49 # BOTH artefacts, or the model is useless: raw columns must be encoded
50 # exactly as they were in training.
51 model.save('churn_net.keras')

```

```

52 joblib.dump(prepare, 'churn_prep.joblib')
53
54 loaded = tf.keras.models.load_model('churn_net.keras')
55 loaded_prep = joblib.load('churn_prep.joblib')
56
57 fresh = loaded_prep.transform(X_te.head(5)).astype('float32')
58 print('reloaded predictions %s' % loaded.predict(fresh, verbose=0).ravel().round(4))
59 print('original predictions %s' % model.predict(Xv[:5], verbose=0).ravel().round(4))
60 print('\nfiles: %s' % sorted(f for f in os.listdir('.')
61                             if f.startswith('churn_'))))

```

```

reloaded predictions [0.6116 0.1354 0.328  0.1717 0.5191]
original predictions [0.6116 0.1354 0.328  0.1717 0.5191]

files: ['churn_by_contract.png', 'churn_net.keras', 'churn_prep.joblib']

```

The preprocessor is half the model

Save the network alone and you have a function that takes fifteen mystery numbers. The `ColumnTransformer` holds the imputation medians, the scaling parameters and the category lists, all learned from training data, all required to turn a customer record into those fifteen numbers. Losing it means retraining. Week 15 makes this a single versioned artefact.

Your assignment

Take the network above and make it deliberately too small, one hidden layer of 4 units, then deliberately too large, 512 units in three layers. Record the training loss, validation loss and test AUC for each. You should see underfitting and overfitting from the same code with one number changed, and the middle should not beat logistic regression either.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. A single perceptron cannot solve XOR because:
 - a) The learning rate is too high
 - b) It needs more data
 - c) It has too few weights
 - d) XOR is not linearly separable, and one layer can only draw a straight boundary
2. What does a hidden layer with a non-linear activation add?
 - a) The ability to bend the decision boundary; without the non-linearity, stacked layers collapse to a single linear map
 - b) Fewer parameters

- c) Automatic scaling
 - d) Speed
3. Gradient checking compares analytic gradients against:
- a) Random numbers
 - b) A numerical estimate from small finite differences
 - c) The loss
 - d) The previous epoch
4. Initialising every weight to zero fails because:
- a) The bias cannot learn
 - b) The loss is undefined
 - c) Every unit in a layer receives the same gradient and stays identical forever
 - d) It overflows
5. Adam's advantage over plain SGD is that it:
- a) Needs no learning rate
 - b) Is always faster
 - c) Never overfits
 - d) Adapts an effective step size per parameter from the running magnitude of its gradients
6. `restore_best_weights=True` matters because without it:
- a) You keep the weights from the last epoch, which is several epochs past the best one
 - b) The model is not saved
 - c) Validation is skipped
 - d) Training never stops
7. Batch normalisation behaves differently at prediction time because it:
- a) Is disabled
 - b) Uses running averages accumulated during training instead of the current batch's statistics
 - c) Uses a larger batch
 - d) Recomputes the mean
8. A network with several hundred times the parameters matched logistic regression on the churn data. The right conclusion is:
- a) More epochs were needed
 - b) The network was badly trained
 - c) Tabular data of a few thousand rows is not where networks win
 - d) Logistic regression is always better

Answers: 1d, 2a, 3b, 4c, 5d, 6a, 7b, 8c. Anything you got wrong points at the day to re-read, not at a gap in ability.

Convolutional Networks and Representation Learning

Week 12

Build convolution from scratch, train a real image classifier, transfer a pre-trained base, and settle when deep learning is worth it.

OBJECTIVES

A dense layer treats a photograph as a bag of unrelated numbers. This week builds the alternative from first principles, a small grid of shared weights slid across the image, then trains it, augments it, transfers it and strips the labels off entirely to get autoencoders. It closes by answering the question week 11 left open: when is a network the right tool, and when is it an expensive way to match a straight line?

Day 1 Convolution, By Hand

Why dense layers waste weights on images, and what a kernel really does

Week 11 built networks out of dense layers, where every input touches every unit. That is the right structure when your columns are `tenure_months` and `contract`, because those two have no natural neighbours. It is the wrong structure for a picture, where the value of a pixel is almost entirely explained by the pixels immediately around it.

Count the weights before you write the model

```

1 # A 200 x 200 colour photograph into one hidden layer of 128 units.
2 pixels = 200 * 200 * 3
3 dense_weights = pixels * 128 + 128
4
5 # The same first layer done as convolution: 128 filters, each 3x3,
6 # looking at all 3 colour channels.
7 conv_weights = 3 * 3 * 3 * 128 + 128
8
9 print('flattened input      %9d values' % pixels)
10 print('dense first layer   %9d weights' % dense_weights)
11 print('conv first layer    %9d weights' % conv_weights)
12 print('difference          %9.0fx' % (dense_weights / conv_weights))

```

```

flattened input      120000 values
dense first layer    15360128 weights
conv first layer     3584 weights
difference           4286x

```

Four thousand times the weights, and the dense version is the worse model of the two, not merely the more expensive one. A dense layer has to learn what an edge looks like separately in every position it might appear. A convolutional layer learns it once.

Doing it by hand

```

1 import numpy as np
2 from sklearn.datasets import load_digits
3
4 def conv2d(img, kernel):
5     kh, kw = kernel.shape
6     out = np.zeros((img.shape[0] - kh + 1, img.shape[1] - kw + 1))
7     for i in range(out.shape[0]):
8         for j in range(out.shape[1]):
9             out[i, j] = float((img[i:i + kh, j:j + kw] * kernel).sum())
10    return out
11
12 def show(a):
13     chars = ' .:-+*#%@'
14     lo, hi = a.min(), a.max()
15     a = (a - lo) / (hi - lo + 1e-9)
16     for row in a:
17         print(''.join(chars[min(9, int(v * 9.999))] for v in row))
18
19 digits = load_digits()
20 img = digits.images[0]
21 print('digit %d, as 8 x 8 greyscale:' % digits.target[0])
22 show(img)

```

```

digit 0, as 8 x 8 greyscale:
-%+
%Q*Q-
.@. #+
:## ++
-+ ++
:## #=
.@-*#
-%*

```

Eight by eight is small enough to print, which is the reason this week uses it. Every operation below can be checked with your eyes rather than taken on trust.

```

1 import numpy as np
2 from sklearn.datasets import load_digits
3
4 def conv2d(img, kernel):
5     kh, kw = kernel.shape
6     out = np.zeros((img.shape[0] - kh + 1, img.shape[1] - kw + 1))
7     for i in range(out.shape[0]):
8         for j in range(out.shape[1]):
9             out[i, j] = float((img[i:i + kh, j:j + kw] * kernel).sum())
10    return out
11

```

```

12 def show(a):
13     chars = ' .:-+*#%@'
14     lo, hi = a.min(), a.max()
15     a = (a - lo) / (hi - lo + 1e-9)
16     for row in a:
17         print(''.join(chars[min(9, int(v * 9.999))] for v in row))
18
19 digits = load_digits()
20 img = digits.images[0]
21 vertical = np.array([[ -1., 0., 1.],
22                     [ -2., 0., 2.],
23                     [ -1., 0., 1.]])
24
25 response = conv2d(img, vertical)
26 print('input 8 x 8, kernel 3 x 3, output %d x %d'
27       % response.shape)
28 print('\nstrength of vertical edges:')
29 show(np.abs(response))

```

```
input 8 x 8, kernel 3 x 3, output 6 x 6
```

```
strength of vertical edges:
```

```

%#: .#
@.%=-%
%: %++*
*: **++
#. +#. %
%: . =*

```

That kernel is the Sobel operator: negative on the left, positive on the right, so it produces a large number wherever brightness climbs from left to right and roughly zero across a flat region. Its transpose finds horizontal edges instead.

```

1 import numpy as np
2 from sklearn.datasets import load_digits
3
4 def conv2d(img, kernel):
5     kh, kw = kernel.shape
6     out = np.zeros((img.shape[0] - kh + 1, img.shape[1] - kw + 1))
7     for i in range(out.shape[0]):
8         for j in range(out.shape[1]):
9             out[i, j] = float((img[i:i + kh, j:j + kw] * kernel).sum())
10    return out
11
12 def show(a):
13     chars = ' .:-+*#%@'
14     lo, hi = a.min(), a.max()
15     a = (a - lo) / (hi - lo + 1e-9)
16     for row in a:
17         print(''.join(chars[min(9, int(v * 9.999))] for v in row))
18
19 digits = load_digits()
20 img = digits.images[0]
21 vertical = np.array([[ -1., 0., 1.], [ -2., 0., 2.], [ -1., 0., 1.]])

```

```

22 kernels = [('vertical edges', vertical),
23             ('horizontal edges', vertical.T),
24             ('blur', np.ones((3, 3)) / 9.0),
25             ('sharpen', np.array([[0., -1., 0.],
26                                   [-1., 5., -1.],
27                                   [0., -1., 0.])))
28
29 for name, k in kernels:
30     print('%s:' % name)
31     show(np.abs(conv2d(img, k)))
32     print('')

```

```

vertical edges:
%#: .#
@.%=-%
%:%++*
*:*+*
#.+#. %
%: :.*

horizontal edges:
-:===*
.-@@=
-:
. .
-*. -
- #=+*

blur:
.*@%+.
-*#+*=
--. ==
:: ==
:---+-
.=#*=

sharpen:
-*#:@
@===-
*.:.-
.: :--
*--*-
.@:=#-

```

These four are not learned

Sobel, blur and sharpen are hand-designed kernels that image processing has used since the 1960s. The only thing a convolutional network changes is that it does not accept your kernels. It starts from random numbers and learns the nine weights that reduce the loss. Frequently the first layer of a trained network rediscovers something very close to Sobel, which is a good sign rather than a coincidence.

Why sharing the weights is the whole point

```

1 import numpy as np
2 from sklearn.datasets import load_digits
3
4 def conv2d(img, kernel):
5     kh, kw = kernel.shape
6     out = np.zeros((img.shape[0] - kh + 1, img.shape[1] - kw + 1))
7     for i in range(out.shape[0]):
8         for j in range(out.shape[1]):
9             out[i, j] = float((img[i:i + kh, j:j + kw] * kernel).sum())
10    return out
11
12 def show(a):
13     chars = ' .:-+*#%@'
14     lo, hi = a.min(), a.max()
15     a = (a - lo) / (hi - lo + 1e-9)
16     for row in a:
17         print(''.join(chars[min(9, int(v * 9.999))] for v in row))
18
19 digits = load_digits()
20 img = digits.images[0]
21 vertical = np.array([[ -1., 0., 1.], [ -2., 0., 2.], [ -1., 0., 1.]])
22
23 shifted = np.roll(img, 2, axis=1) # move the digit two columns right
24 a = conv2d(img, vertical)
25 b = conv2d(shifted, vertical)
26
27 print('row 3 of the response, original image:')
28 print(a[3].round(1))
29 print('row 3 of the response, shifted image:')
30 print(b[3].round(1))
31 print('\nshifted response is the original moved by 2:',
32       np.allclose(b[:, 2:], a[:, :-2]))

```

```

row 3 of the response, original image:
[ 39. -18. -38.  38.  30. -38.]
row 3 of the response, shifted image:
[-31.  18.  39. -18. -38.  38.]

shifted response is the original moved by 2: True

```

⚠ Watch Out

Equivariance is not invariance

The response moved, it did not stay the same. Invariance, getting the same answer regardless of position, comes from pooling and from the global averaging at the end of the network, which day 2 covers. Day 4 measures how much of it you actually get, and it is less than most people assume.

Day 1 takeaway

Day 2 Padding, Pooling and Receptive Fields

The three settings that decide what your network can see

One convolution is a filter. A network needs three more ideas before it is a network: what happens at the edges, how to shrink the picture, and how far a unit deep in the stack can actually see.

Padding and stride decide the output size

```

1 def out_size(n, k, stride=1, pad=0):
2     return (n + 2 * pad - k) // stride + 1
3
4 print('%7s %7s %7s %5s %8s %s'
5       % ('input', 'kernel', 'stride', 'pad', 'output', 'called'))
6 rows = [(28, 3, 1, 0, 'valid'), (28, 3, 1, 1, 'same'),
7         (28, 5, 1, 0, 'valid'), (28, 5, 1, 2, 'same'),
8         (28, 3, 2, 1, 'same, stride 2'), (8, 3, 1, 0, 'valid'),
9         (8, 3, 1, 1, 'same')]
10 for n, k, s, pad, name in rows:
11     print('%7d %7d %7d %5d %8d %s'
12          % (n, k, s, pad, out_size(n, k, s, pad), name))

```

input	kernel	stride	pad	output	called
28	3	1	0	26	valid
28	3	1	1	28	same
28	5	1	0	24	valid
28	5	1	2	28	same
28	3	2	1	14	same, stride 2
8	3	1	0	6	valid
8	3	1	1	8	same

Setting	Keras	Effect
No padding	<code>padding='valid'</code>	Output shrinks by <code>kernel - 1</code> each layer
Pad to keep size	<code>padding='same'</code>	Output matches the input, so you can stack many layers
Stride 2	<code>strides=2</code>	Halves the resolution; an alternative to pooling

Use same unless you have a reason not to

With `valid` padding a 3×3 kernel costs you two pixels per layer. Six layers into a 28×28 image and you have a 16×16 map, having thrown away the border for no benefit. `same` keeps the arithmetic simple and the border intact.

Pooling

```

1 import numpy as np
2 from sklearn.datasets import load_digits
3

```

```

4 def conv2d(img, kernel):
5     kh, kw = kernel.shape
6     out = np.zeros((img.shape[0] - kh + 1, img.shape[1] - kw + 1))
7     for i in range(out.shape[0]):
8         for j in range(out.shape[1]):
9             out[i, j] = float((img[i:i + kh, j:j + kw] * kernel).sum())
10    return out
11
12 def show(a):
13     chars = ' .:-+*#%@'
14     lo, hi = a.min(), a.max()
15     a = (a - lo) / (hi - lo + 1e-9)
16     for row in a:
17         print(''.join(chars[min(9, int(v * 9.999))] for v in row))
18
19 digits = load_digits()
20 img = digits.images[0]
21 def pool(a, k=2, how='max'):
22     h = a.shape[0] // k * k
23     w = a.shape[1] // k * k
24     b = a[:h, :w].reshape(h // k, k, w // k, k)
25     return b.max(axis=(1, 3)) if how == 'max' else b.mean(axis=(1, 3))
26
27 print('original 8 x 8:')
28 show(img)
29 print('\nmax pooled to %d x %d:' % pool(img).shape)
30 show(pool(img))
31 print('\naverage pooled:')
32 show(pool(img, how='mean'))

```

```

original 8 x 8:
-%+
%#@*@-
.@. #+
:# ++
-+ ++
:# #=
.@-*#
-%*

max pooled to 4 x 4:
@@-
:@#+
-##+
. @#

average pooled:
@#.
.*=-
.==
%*

```

Max pooling keeps the strongest response and is the usual choice, because a feature map answers "is this pattern here" and the strongest evidence is the informative part. Average pooling blurs, which is occasionally what you want at the very end of a network.

Receptive field: what a deep unit can see

```

1 r, jump = 1, 1
2 print('%-24s %8s %8s' % ('after', 'stride', 'sees'))
3 stack = [('conv 3x3', 3, 1), ('conv 3x3', 3, 1), ('maxpool 2x2', 2, 2),
4          ('conv 3x3', 3, 1), ('conv 3x3', 3, 1), ('maxpool 2x2', 2, 2),
5          ('conv 3x3', 3, 1)]
6 for name, k, s in stack:
7     r = r + (k - 1) * jump
8     jump = jump * s
9     print('%-24s %8d %6d px' % (name, jump, r))

```

after	stride	sees
conv 3x3	1	3 px
conv 3x3	1	5 px
maxpool 2x2	2	6 px
conv 3x3	2	10 px
conv 3x3	2	14 px
maxpool 2x2	4	16 px
conv 3x3	4	24 px

A single unit in the last layer above responds to a 24×24 region of the input, despite every kernel in the stack being 3×3 . This is why deep stacks of small kernels replaced the large kernels of early architectures: two 3×3 layers see as much as one 5×5 layer, with fewer weights and an extra nonlinearity in between.

The same thing in Keras

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 model = tf.keras.Sequential([
8     tf.keras.layers.Input(shape=(8, 8, 1)),
9     tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
10    tf.keras.layers.MaxPooling2D(2),
11    tf.keras.layers.Conv2D(32, 3, padding='same', activation='relu'),
12    tf.keras.layers.GlobalAveragePooling2D(),
13    tf.keras.layers.Dense(10, activation='softmax'),
14 ])
15 model.summary(line_length=64)

```

```

Model: "sequential"
+-----+-----+-----+
| Layer (type)          | Output Shape      | Param # |
+-----+-----+-----+
| conv2d (Conv2D)       | (None, 8, 8, 16)  | 160     |
+-----+-----+-----+
| max_pooling2d         | (None, 4, 4, 16)  | 0       |
| (MaxPooling2D)        |                    |         |
+-----+-----+-----+

```

conv2d_1 (Conv2D)	(None, 4, 4, 32)	4,640
+-----+-----+-----+		
global_average_pooling2d	(None, 32)	0
(GlobalAveragePooling2D)		
+-----+-----+-----+		
dense (Dense)	(None, 10)	330
+-----+-----+-----+		

Total params: 5,130 (20.04 KB)
Trainable params: 5,130 (20.04 KB)
Non-trainable params: 0 (0.00 B)

Read the shape column, not the layer names

The output shape tells you what the network is doing to the picture: 8×8 with 16 channels, then 4×4 after pooling, then 4×4 with 32 channels, then a single vector of 32 once global pooling has averaged each channel over all positions. Resolution falls, channel count rises. That trade is the shape of nearly every convolutional architecture ever published.

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 conv = tf.keras.layers.Conv2D(32, 3, padding='same')
8 conv.build((None, 8, 8, 16))
9 kernel, bias = conv.get_weights()
10
11 print('kernel shape', kernel.shape, ' = (height, width, in, out)')
12 print('weights      %d = 3 * 3 * 16 * 32' % kernel.size)
13 print('biases       %d' % bias.size)
14 print('total        %d' % (kernel.size + bias.size))
15
16 flat = tf.keras.layers.Dense(32)
17 flat.build((None, 8 * 8 * 16))
18 print('\na dense layer over the same input would need %d'
19       % sum(w.size for w in flat.get_weights()))

```

```

kernel shape (3, 3, 16, 32) = (height, width, in, out)
weights      4608 = 3 * 3 * 16 * 32
biases       32
total        4640

a dense layer over the same input would need 32800

```

Every filter has one bias, not one per position, another consequence of sharing. The count is independent of image size, which is why the same architecture handles a larger picture without a single new parameter.

Day 2 takeaway

Day 3 Training a Convolutional Network

Digits, a dense baseline, a matched-budget CNN and an honest floor

Time to train one. The digits bundled with scikit-learn are 1,797 handwritten numerals at 8×8 , which is small enough that everything on this page finishes in seconds and large enough that overfitting is a real risk.

The data

```
1 import numpy as np
2 from sklearn.datasets import load_digits
3 from sklearn.model_selection import train_test_split
4
5 d = load_digits()
6 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
7 y = d.target.astype('int32')
8 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
9                                         stratify=y, random_state=42)
10 print('images ', X.shape, X.dtype)
11 print('labels ', y.shape, 'values %d to %d' % (y.min(), y.max()))
12 print('training', X_tr.shape[0], ' test', X_te.shape[0])
13 print('\nclass balance:')
14 print(np.bincount(y_tr))
```

```
images (1797, 8, 8, 1) float32
labels (1797,) values 0 to 9
training 1347 test 450

class balance:
[133 136 133 137 136 136 136 134 131 135]
```

Watch Out

Scale the pixels

The raw values run 0 to 16. Dividing by 16 puts them in $[0, 1]$, which matters for the same reason it mattered in week 11: the first layer's gradients are proportional to its inputs, and inputs an order of magnitude larger than the initialised weights make the first few steps wild. Images are the one case where you can usually get away with a constant divisor instead of `StandardScaler`, because every feature is already on the same scale as every other.

Your numbers will be close, not identical

Every figure on this page came from actually running the code beside it, with `tf.random.set_seed` fixed. Run it yourself and you should land within a few tenths of a point, but not on the same digits. TensorFlow parallelises across CPU threads, and floating-point addition is not associative, so the order the sums complete in changes the last decimal places and those differences compound over epochs. Seeds control the random draws; they do not control the thread scheduler. This is why the comparisons below average over several runs wherever the margin is small.

A dense baseline first

```
1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 from sklearn.datasets import load_digits
9 from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 tf.random.set_seed(42)
17 dense = tf.keras.Sequential([
18     tf.keras.layers.Input(shape=(8, 8, 1)),
19     tf.keras.layers.Flatten(),
20     tf.keras.layers.Dense(64, activation='relu'),
21     tf.keras.layers.Dense(10, activation='softmax'),
22 ])
23 dense.compile(optimizer='adam',
24               loss='sparse_categorical_crossentropy',
25               metrics=['accuracy'])
26 dense.fit(X_tr, y_tr, epochs=40, batch_size=64, verbose=0)
27
28 print('parameters %d' % dense.count_params())
29 print('test accuracy %.4f' % dense.evaluate(X_te, y_te, verbose=0)[1])
```

```
parameters 4810
test accuracy 0.9778
```

The convolutional version, at a matched parameter budget

Day 2's architecture: two convolutions, a pooling layer, and global average pooling to collapse each channel to a single number before the classifier. It is the standard modern shape.

```
1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 from sklearn.datasets import load_digits
9 from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')
```

```

14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 tf.random.set_seed(42)
17 cnn = tf.keras.Sequential([
18     tf.keras.layers.Input(shape=(8, 8, 1)),
19     tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
20     tf.keras.layers.MaxPooling2D(2),
21     tf.keras.layers.Conv2D(32, 3, padding='same', activation='relu'),
22     tf.keras.layers.GlobalAveragePooling2D(),
23     tf.keras.layers.Dense(10, activation='softmax'),
24 ])
25 cnn.compile(optimizer='adam',
26             loss='sparse_categorical_crossentropy',
27             metrics=['accuracy'])
28 cnn.fit(X_tr, y_tr, epochs=40, batch_size=64, verbose=0)
29
30 print('parameters %d' % cnn.count_params())
31 print('test accuracy %.4f' % cnn.evaluate(X_te, y_te, verbose=0)[1])

```

```

parameters 5130
test accuracy 0.8578

```

The convolutional network is *worse*. Not marginally, by more than eleven points, with more parameters than the dense model it lost to. Before reaching for more epochs or a bigger network, work out what the architecture is throwing away.

Diagnosing it

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 from sklearn.datasets import load_digits
9 from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 tf.random.set_seed(42)
17 features = tf.keras.Sequential([
18     tf.keras.layers.Input(shape=(8, 8, 1)),
19     tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
20     tf.keras.layers.MaxPooling2D(2),
21     tf.keras.layers.Conv2D(32, 3, padding='same', activation='relu'),
22 ])
23 maps = features.predict(X_te[:1], verbose=0)
24
25 print('feature maps reaching the classifier:', maps.shape[1:])

```

```

26 print('  that is %d numbers' % int(np.prod(maps.shape[1:])))
27 print('flatten passes on          %d' % int(np.prod(maps.shape[1:])))
28 print('global average pooling passes on %d' % maps.shape[-1])
29 print('\naveraging each 4 x 4 map to one number answers'
30       ' "is this feature present"')
31 print('and discards "where". At 8 x 8, where is most of the signal.')

```

```

feature maps reaching the classifier: (4, 4, 32)
  that is 512 numbers
flatten passes on          512
global average pooling passes on 32

averaging each 4 x 4 map to one number answers "is this feature present"
and discards "where". At 8 x 8, where is most of the signal.

```

⚠ Watch Out

Global average pooling is not a free win

It is the right default for a 224×224 photograph, where a cat is a cat wherever it sits in the frame and the final feature map is still 7×7 after many layers of downsampling. On an 8×8 digit that has been pooled once, averaging over the remaining 4×4 grid throws away the spatial arrangement that distinguishes a 6 from a 9. Day 4 shows the same layer being exactly the right choice on a different problem. The layer is not wrong, the match between layer and problem was.

The same network, keeping position

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 from sklearn.datasets import load_digits
9 from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 tf.random.set_seed(42)
17 cnn = tf.keras.Sequential([
18     tf.keras.layers.Input(shape=(8, 8, 1)),
19     tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
20     tf.keras.layers.MaxPooling2D(2),
21     tf.keras.layers.Conv2D(32, 3, padding='same', activation='relu'),
22     tf.keras.layers.Flatten(), # was GlobalAveragePooling2D
23     tf.keras.layers.Dense(10, activation='softmax'),
24 ])

```

```

25 cnn.compile(optimizer='adam',
26             loss='sparse_categorical_crossentropy',
27             metrics=['accuracy'])
28 cnn.fit(X_tr, y_tr, epochs=40, batch_size=64, verbose=0)
29
30 print('parameters %d' % cnn.count_params())
31 print('test accuracy %.4f' % cnn.evaluate(X_te, y_te, verbose=0)[1])

```

```

parameters 9930
test accuracy 0.9733

```

One layer changed, and the eleven points come back. Note also what the network has *not* bought: it is now level with the dense baseline rather than ahead of it. At this resolution there is very little spatial structure left for convolution to exploit. Day 4 constructs the case where there is.

And a model with no network in it at all

```

1 import numpy as np
2 from sklearn.datasets import load_digits
3 from sklearn.model_selection import train_test_split
4
5 d = load_digits()
6 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
7 y = d.target.astype('int32')
8 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
9                                         stratify=y, random_state=42)
10 from sklearn.linear_model import LogisticRegression
11 from sklearn.ensemble import HistGradientBoostingClassifier
12 import time
13
14 flat_tr = X_tr.reshape(len(X_tr), -1)
15 flat_te = X_te.reshape(len(X_te), -1)
16
17 for name, model in [('logistic regression', LogisticRegression(max_iter=2000)),
18                    ('gradient boosting',
19                     ↪ HistGradientBoostingClassifier(random_state=42))]:
19     t = time.time()
20     model.fit(flat_tr, y_tr)
21     print('%-20s %.4f   fitted in %.1fs'
22           % (name, model.score(flat_te, y_te), time.time() - t))

```

```

logistic regression 0.9622   fitted in 0.1s
gradient boosting   0.9600   fitted in 4.5s

```

Where the mistakes are

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)

```

```

6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 from sklearn.datasets import load_digits
9 from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 from sklearn.metrics import confusion_matrix
17
18 tf.random.set_seed(42)
19 cnn = tf.keras.Sequential([
20     tf.keras.layers.Input(shape=(8, 8, 1)),
21     tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
22     tf.keras.layers.MaxPooling2D(2),
23     tf.keras.layers.Conv2D(32, 3, padding='same', activation='relu'),
24     tf.keras.layers.Flatten(),
25     tf.keras.layers.Dense(10, activation='softmax'),
26 ])
27 cnn.compile(optimizer='adam', loss='sparse_categorical_crossentropy',
28            metrics=['accuracy'])
29 cnn.fit(X_tr, y_tr, epochs=40, batch_size=64, verbose=0)
30
31 pred = cnn.predict(X_te, verbose=0).argmax(axis=1)
32 cm = confusion_matrix(y_te, pred)
33 print('    ' + ''.join('%4d' % i for i in range(10)))
34 for i, row in enumerate(cm):
35     print('%2d ' % i + ''.join('%4d' % v for v in row))
36
37 print('\nthe confusions that happened more than once:')
38 for i in range(10):
39     for j in range(10):
40         if i != j and cm[i, j] > 1:
41             print('    a %d read as a %d, %d times' % (i, j, cm[i, j]))

```

	0	1	2	3	4	5	6	7	8	9
0	44	0	0	0	1	0	0	0	0	0
1	0	46	0	0	0	0	0	0	0	0
2	0	0	44	0	0	0	0	0	0	0
3	0	0	0	45	0	1	0	0	0	0
4	0	0	0	0	45	0	0	0	0	0
5	0	0	0	0	0	45	0	0	0	1
6	0	0	0	0	0	0	44	0	1	0
7	0	0	0	0	0	0	0	45	0	0
8	0	2	0	0	0	0	0	1	40	0
9	0	0	0	0	0	0	0	0	0	45

```

the confusions that happened more than once:
    a 8 read as a 1, 2 times

```

A confusion matrix is the first thing to read, always

An accuracy figure tells you how often the model is wrong. The confusion matrix tells you what it is wrong *about*, and the pattern is almost always interpretable, digits that share a shape at this resolution get muddled, and digits that do not, do not. If the errors look random, suspect a bug in the label alignment before you suspect the model.

Day 3 takeaway**Day 4 Augmentation and Invariance**

Where a dense network collapses, and how to buy back what it lost

Day 3 was a fair fight on a dataset where every digit is already centred in the frame. Real images are not centred, and that is where the difference between the two architectures stops being academic.

A dataset built to make the point

```

1 import numpy as np
2
3 def draw(kind, cx, cy, r, size=20):
4     yy, xx = np.mgrid[0:size, 0:size]
5     dx, dy = xx - cx, yy - cy
6     if kind == 0:                                     # ring
7         m = np.abs(np.hypot(dx, dy) - r) < 1.0
8     elif kind == 1:                                   # square outline
9         m = np.abs(np.maximum(np.abs(dx), np.abs(dy)) - r) < 1.0
10    else:                                              # cross
11        m = (((np.abs(dx) < 1.0) & (np.abs(dy) <= r))
12              | ((np.abs(dy) < 1.0) & (np.abs(dx) <= r)))
13    return m.astype('float32')
14
15 def make(n, seed, jitter):
16     rng = np.random.default_rng(seed)
17     X = np.zeros((n, 20, 20), dtype='float32')
18     y = rng.integers(0, 3, size=n)
19     for i, kind in enumerate(y):
20         r = int(rng.integers(4, 7))
21         if jitter:
22             cx = int(rng.integers(r + 1, 20 - r))
23             cy = int(rng.integers(r + 1, 20 - r))
24         else:
25             cx = cy = 10
26         X[i] = draw(kind, cx, cy, r)
27     X += rng.normal(0, 0.05, X.shape).astype('float32')
28     return X[... , np.newaxis], y.astype('int32')
29 X_centre, y_centre = make(1500, seed=0, jitter=False)
30 X_jitter, y_jitter = make(500, seed=1, jitter=True)
31

```



```

. . . . .
.....
. . . . .
. . . . .
. . . . .
. . . . .

```

Train on centred, test on moved

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8
9 def draw(kind, cx, cy, r, size=20):
10     yy, xx = np.mgrid[0:size, 0:size]
11     dx, dy = xx - cx, yy - cy
12     if kind == 0:                                     # ring
13         m = np.abs(np.hypot(dx, dy) - r) < 1.0
14     elif kind == 1:                                   # square outline
15         m = np.abs(np.maximum(np.abs(dx), np.abs(dy)) - r) < 1.0
16     else:                                             # cross
17         m = ((np.abs(dx) < 1.0) & (np.abs(dy) <= r))
18             | ((np.abs(dy) < 1.0) & (np.abs(dx) <= r))
19     return m.astype('float32')
20
21 def make(n, seed, jitter):
22     rng = np.random.default_rng(seed)
23     X = np.zeros((n, 20, 20), dtype='float32')
24     y = rng.integers(0, 3, size=n)
25     for i, kind in enumerate(y):
26         r = int(rng.integers(4, 7))
27         if jitter:
28             cx = int(rng.integers(r + 1, 20 - r))
29             cy = int(rng.integers(r + 1, 20 - r))
30         else:
31             cx = cy = 10
32         X[i] = draw(kind, cx, cy, r)
33     X += rng.normal(0, 0.05, X.shape).astype('float32')
34     return X[..., np.newaxis], y.astype('int32')
35 X_tr, y_tr = make(1500, seed=0, jitter=False)
36 X_same, y_same = make(500, seed=2, jitter=False)
37 X_moved, y_moved = make(500, seed=1, jitter=True)
38
39 def dense_net():
40     return tf.keras.Sequential([
41         tf.keras.layers.Input(shape=(20, 20, 1)),
42         tf.keras.layers.Flatten(),
43         tf.keras.layers.Dense(64, activation='relu'),
44         tf.keras.layers.Dense(3, activation='softmax')]
45

```

```

46 def conv_net():
47     return tf.keras.Sequential([
48         tf.keras.layers.Input(shape=(20, 20, 1)),
49         tf.keras.layers.Conv2D(8, 3, padding='same', activation='relu'),
50         tf.keras.layers.MaxPooling2D(2),
51         tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
52         tf.keras.layers.GlobalAveragePooling2D(),
53         tf.keras.layers.Dense(3, activation='softmax')])
54
55 print('%-12s %8s %14s %14s' % ('', 'params', 'centred test', 'moved test'))
56 for name, build in [('dense', dense_net), ('conv', conv_net)]:
57     tf.random.set_seed(42)
58     m = build()
59     m.compile(optimizer='adam', loss='sparse_categorical_crossentropy',
60             metrics=['accuracy'])
61     m.fit(X_tr, y_tr, epochs=30, batch_size=64, verbose=0)
62     print('%-12s %8d %14.4f %14.4f'
63           % (name, m.count_params(),
64             m.evaluate(X_same, y_same, verbose=0)[1],
65             m.evaluate(X_moved, y_moved, verbose=0)[1]))

```

	params	centred test	moved test
dense	25859	1.0000	0.3880
conv	1299	1.0000	0.9840

Both networks are perfect on centred shapes. Move the shape and the dense network falls to barely above the one-in-three you would get by guessing, while the convolutional network, with twenty times fewer parameters, barely notices.

The layer day 3 blamed is the hero here

This network ends in `GlobalAveragePooling2D`, the same layer that cost eleven points on the digits. Here it is exactly right: "is there a ring in this picture" is a question whose answer must not depend on where the ring is, and averaging over positions is precisely how you enforce that. Discarding position is a *feature* when position is noise and a bug when position is signal. Nothing about the layer changed, the problem did.

Augmentation: move the training data instead

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8
9 def draw(kind, cx, cy, r, size=20):
10     yy, xx = np.mgrid[0:size, 0:size]
11     dx, dy = xx - cx, yy - cy
12     if kind == 0:                                     # ring
13         m = np.abs(np.hypot(dx, dy) - r) < 1.0

```

```

14     elif kind == 1:                                     # square outline
15         m = np.abs(np.maximum(np.abs(dx), np.abs(dy)) - r) < 1.0
16     else:                                              # cross
17         m = ((np.abs(dx) < 1.0) & (np.abs(dy) <= r))
18             | ((np.abs(dy) < 1.0) & (np.abs(dx) <= r)))
19     return m.astype('float32')
20
21 def make(n, seed, jitter):
22     rng = np.random.default_rng(seed)
23     X = np.zeros((n, 20, 20), dtype='float32')
24     y = rng.integers(0, 3, size=n)
25     for i, kind in enumerate(y):
26         r = int(rng.integers(4, 7))
27         if jitter:
28             cx = int(rng.integers(r + 1, 20 - r))
29             cy = int(rng.integers(r + 1, 20 - r))
30         else:
31             cx = cy = 10
32         X[i] = draw(kind, cx, cy, r)
33         X += rng.normal(0, 0.05, X.shape).astype('float32')
34     return X[..., np.newaxis], y.astype('int32')
35 X_tr, y_tr = make(1500, seed=0, jitter=False)
36
37 augment = tf.keras.Sequential([
38     tf.keras.layers.RandomTranslation(0.2, 0.2, fill_mode='constant'),
39     tf.keras.layers.RandomRotation(0.05, fill_mode='constant'),
40 ])
41
42 batch = tf.convert_to_tensor(X_tr[:1])
43 chars = ' .:-+*#%@'
44 def show(a):
45     a = (a - a.min()) / (a.max() - a.min() + 1e-9)
46     for row in a:
47         print(''.join(chars[min(9, int(v * 9.999))] for v in row))
48
49 print('original, class %d:' % y_tr[0])
50 show(X_tr[0, :, :, 0])
51 for i in range(2):
52     print('\naugmented copy %d:' % (i + 1))
53     show(augment(batch, training=True).numpy()[0, :, :, 0])

```

```

original, class 2:
.....
.....
.....
.....
.. . . . % . . . . .
..... @ . . . . .
..... @ . . . . .
.: .: .: .: % . . . . .
..... @ . . . . .
..... % . . . . .
.: %@@@%@@%@@%@@%...
.. . . . @ . . . . .
..... @ . . . . .

```

```

. . . . . @ . . . . .
. . . . . @ . : . . .
. . . . . @ . . . . .
. . . . . % . : . . .
. . . . .
. . . : . . . . .
. . . . .

```

augmented copy 1:

```

.
. . . .
. ++. . .
. . +* .
. . *+ .
. . . ** .
. . . **
. . . **
:-: -#*:::
=%@@@%@@%%%%/
. **::--
. **
. *+ .
. ** .
. *+ .
. . += .
. . . .
. . .
.

```

augmented copy 2:

```

. .
. .
. . . .
. . . . -= . . .
. . . : ** . . .
. . : # = .
. -@: .
. -: . =@
-#*+==% .
-=%%%@%+=--: .
-%##%##+:
. :@: . =+ .
. . -@ . . .
. +# . . .
. . *+ .
. . . : * - .
. . . : . . .

```

⚠ Watch Out**training=True is not optional**

Augmentation layers are deliberately inert at prediction time. You do not want your test images randomly shifted. Inside `model.fit` Keras sets the flag for you. Calling the layer yourself without it returns the input untouched, which looks exactly like a broken augmentation pipeline and is the single most common confusion here.

Does it actually help?

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8
9 def draw(kind, cx, cy, r, size=20):
10     yy, xx = np.mgrid[0:size, 0:size]
11     dx, dy = xx - cx, yy - cy
12     if kind == 0:                                     # ring
13         m = np.abs(np.hypot(dx, dy) - r) < 1.0
14     elif kind == 1:                                   # square outline
15         m = np.abs(np.maximum(np.abs(dx), np.abs(dy)) - r) < 1.0
16     else:                                             # cross
17         m = ((np.abs(dx) < 1.0) & (np.abs(dy) <= r))
18             | ((np.abs(dy) < 1.0) & (np.abs(dx) <= r))
19     return m.astype('float32')
20
21 def make(n, seed, jitter):
22     rng = np.random.default_rng(seed)
23     X = np.zeros((n, 20, 20), dtype='float32')
24     y = rng.integers(0, 3, size=n)
25     for i, kind in enumerate(y):
26         r = int(rng.integers(4, 7))
27         if jitter:
28             cx = int(rng.integers(r + 1, 20 - r))
29             cy = int(rng.integers(r + 1, 20 - r))
30         else:
31             cx = cy = 10
32         X[i] = draw(kind, cx, cy, r)
33     X += rng.normal(0, 0.05, X.shape).astype('float32')
34     return X[... , np.newaxis], y.astype('int32')
35 X_tr, y_tr = make(1500, seed=0, jitter=False)
36 X_moved, y_moved = make(500, seed=1, jitter=True)
37
38 def conv_net(with_augmentation):
39     layers = [tf.keras.layers.Input(shape=(20, 20, 1))]
40     if with_augmentation:
41         layers += [tf.keras.layers.RandomTranslation(0.25, 0.25,
```

```

42                                     fill_mode='constant'])
43     layers += [
44         tf.keras.layers.Conv2D(8, 3, padding='same', activation='relu'),
45         tf.keras.layers.MaxPooling2D(2),
46         tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
47         tf.keras.layers.GlobalAveragePooling2D(),
48         tf.keras.layers.Dense(3, activation='softmax')]
49     return tf.keras.Sequential(layers)
50
51 # One run proves nothing at this margin, so average over three seeds.
52 print('%-22s %10s %18s' % ('', 'mean', 'runs'))
53 for label, aug in [('no augmentation', False), ('with translation', True)]:
54     scores = []
55     for seed in range(3):
56         tf.random.set_seed(seed)
57         m = conv_net(aug)
58         m.compile(optimizer='adam', loss='sparse_categorical_crossentropy',
59                 metrics=['accuracy'])
60         m.fit(X_tr, y_tr, epochs=30, batch_size=64, verbose=0)
61         scores.append(m.evaluate(X_moved, y_moved, verbose=0)[1])
62     print('%-22s %10.4f %18s'
63           % (label, float(np.mean(scores)),
64             ' '.join('%0.3f' % s for s in scores)))

```

	mean		runs	
no augmentation	0.9713	0.978	0.970	0.966
with translation	0.9820	0.994	0.960	0.992

About a point on average, and now read the individual runs, because they are the more instructive column. Augmentation won twice and lost once. A single run of this experiment could have shown you a two-point gain, no gain, or a loss, depending entirely on which seed you happened to use.

Watch Out

One run is not a measurement

Five hundred test images means one point of accuracy is five images. Any comparison at that margin needs several runs and a look at the spread, or you are reporting the seed rather than the technique. This is the same discipline week 7 applied with cross-validation, and it does not stop being necessary because the model is a neural network. It becomes more necessary, because networks have a second source of randomness in their initialisation.

The effect here is real but modest, which is what augmentation looks like on a problem whose architecture already handles the transformation. The large gains come from augmenting a model with no built-in defence. The dense network in the previous table, which had none at all.

Augmentation	Safe for	Wrong for
Horizontal flip	Photographs of objects, animals	Text, digits, anything where b and d differ
Vertical flip	Satellite and microscope imagery	Almost every photograph of the world
Rotation $\pm 10^\circ$	Most natural images	Digits at large angles, 6 becomes 9
Brightness, contrast	Anything photographed under varying light	Medical scans where intensity is the measurement
Random crop	Large images with a margin	Small images where the subject fills the frame

⚠ Watch Out

Augmentation encodes an assumption, and it can be false

Flipping a chest X-ray horizontally produces an image of a patient with *situs inversus*, a rare condition. Train on flipped X-rays and you have taught the model that heart position carries no information. Every augmentation is a claim that a transformation does not change the label. Check the claim against the domain, not against the accuracy number.

Day 4 takeaway

Day 5 Transfer Learning

Reusing learned features, freezing, fine-tuning, and reading the filters

Nobody trains an image model from random weights any more, unless they have millions of labelled examples. They take a network trained on something large and reuse its early layers, because the features those layers learned, edges, corners, textures, are not specific to the task it was trained on.

Why this page pretrains its own base

The usual demonstration downloads ImageNet weights for MobileNetV2 or ResNet50, and in your own work that is exactly what you should do, one line, `tf.keras.applications.MobileNetV2(weights='imagenet', include_top=False)`. It needs a network connection, so instead this page pretrains a small base on digits 0 to 4 and transfers it to digits 5 to 9. The mechanics are identical, the numbers are real, and they reproduce offline.

Step 1: train the base on the source task

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np

```

```

8 from sklearn.datasets import load_digits
9 from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 src = y_tr < 5                                # digits 0-4, the source task
17
18 tf.random.set_seed(42)
19 base = tf.keras.Sequential([
20     tf.keras.layers.Input(shape=(8, 8, 1)),
21     tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
22     tf.keras.layers.Conv2D(32, 3, padding='same', activation='relu'),
23     tf.keras.layers.MaxPooling2D(2),
24 ], name='features')
25
26 source = tf.keras.Sequential([base,
27                               tf.keras.layers.Flatten(),
28                               tf.keras.layers.Dense(5, activation='softmax')])
29 source.compile(optimizer='adam', loss='sparse_categorical_crossentropy',
30               metrics=['accuracy'])
31 source.fit(X_tr[src], y_tr[src], epochs=40, batch_size=64, verbose=0)
32
33 mask = y_te < 5
34 print('source task (digits 0-4) test accuracy %.4f'
35       % source.evaluate(X_te[mask], y_te[mask], verbose=0)[1])
36 base.save_weights('digit_base.weights.h5')
37 print('feature extractor saved, %d parameters' % base.count_params())

```

```

source task (digits 0-4) test accuracy 0.9956
feature extractor saved, 4800 parameters

```

Step 2: transfer it, with very few labels

The target task is digits 5 to 9, relabelled 0 to 4, with fifteen training examples of each. Two models: one starting from random weights, one reusing the frozen base. Three seeds each, because a seventy-five-row training set produces enough run-to-run variation to show whatever you were hoping to see.

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 from sklearn.datasets import load_digits
9 from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]

```

```

13 y = d.target.astype('int32')
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 def make_base():
17     return tf.keras.Sequential([
18         tf.keras.layers.Input(shape=(8, 8, 1)),
19         tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
20         tf.keras.layers.Conv2D(32, 3, padding='same', activation='relu'),
21         tf.keras.layers.MaxPooling2D(2),
22     ], name='features')
23
24 def head(base, trainable):
25     base.trainable = trainable
26     m = tf.keras.Sequential([base, tf.keras.layers.Flatten(),
27                             tf.keras.layers.Dense(5, activation='softmax')])
28     m.compile(optimizer='adam', loss='sparse_categorical_crossentropy',
29              metrics=['accuracy'])
30     return m
31
32 src = y_tr < 5
33 tgt = np.where(y_tr >= 5)[0]
34 mask = y_te >= 5
35 Xv, yv = X_te[mask], y_te[mask] - 5
36
37 scratch_runs, frozen_runs = [], []
38 for seed in range(3):
39     rng = np.random.default_rng(seed)
40     few = np.concatenate([rng.permutation(tgt[y_tr[tgt] == c][:15])
41                          for c in range(5, 10)])
42     Xf, yf = X_tr[few], y_tr[few] - 5
43
44     tf.random.set_seed(100 + seed)                                # pretrain on 0-4
45     pre = make_base()
46     head(pre, True).fit(X_tr[src], y_tr[src], epochs=40, batch_size=64,
47                        verbose=0)
48
49     tf.random.set_seed(seed)
50     a = head(make_base(), True)
51     a.fit(Xf, yf, epochs=60, batch_size=16, verbose=0)
52     scratch_runs.append(a.evaluate(Xv, yv, verbose=0)[1])
53
54     tf.random.set_seed(seed)
55     b = head(pre, False)
56     b.fit(Xf, yf, epochs=60, batch_size=16, verbose=0)
57     frozen_runs.append(b.evaluate(Xv, yv, verbose=0)[1])
58
59 print('target training examples: 75')
60 print('%-28s %9s %18s' % ('', 'mean', 'runs'))
61 for name, runs in [('from scratch', scratch_runs),
62                   ('frozen pretrained features', frozen_runs)]:
63     print('%-28s %9.4f %18s'
64           % (name, float(np.mean(runs)),
65             ' '.join('%0.3f' % r for r in runs)))
66 print('\ntrainable parameters: scratch %d, frozen %d'

```

```

67     % (sum(int(np.prod(w.shape)) for w in a.trainable_weights),
68         sum(int(np.prod(w.shape)) for w in b.trainable_weights)))

```

```

target training examples: 75

              mean              runs
from scratch          0.9509  0.955  0.960  0.938
frozen pretrained features  0.9435  0.960  0.946  0.924

trainable parameters: scratch 7365, frozen 2565

```

⚠ Watch Out

Transfer learning lost, and the honest thing is to say so

Every textbook example of this has the pretrained model winning comfortably. Here it is behind on average, and it wins one run of three, which is to say the pretrained features bought nothing detectable. That is not a bug in the code above: the mechanics are exactly what you would write against ImageNet weights. It is a fact about the *source* task, and it is worth more than a rigged win.

Why it lost

Transfer only helps when the features the source task learned are better than the features the target's own data can learn. The source here is 675 images of five digit classes. That is not a lot of world knowledge to bring. Growing it should narrow the gap, and it does.

```

1  import os
2  os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3  import numpy as np
4  import tensorflow as tf
5  tf.random.set_seed(42)
6  tf.get_logger().setLevel('ERROR')
7  import numpy as np
8  from sklearn.datasets import load_digits
9  from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 def make_base():
17     return tf.keras.Sequential([
18         tf.keras.layers.Input(shape=(8, 8, 1)),
19         tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
20         tf.keras.layers.Conv2D(32, 3, padding='same', activation='relu'),
21         tf.keras.layers.MaxPooling2D(2),
22     ], name='features')
23
24 def head(base, trainable):
25     base.trainable = trainable
26     m = tf.keras.Sequential([base, tf.keras.layers.Flatten(),
27                             tf.keras.layers.Dense(5, activation='softmax')])

```

```

28     m.compile(optimizer='adam', loss='sparse_categorical_crossentropy',
29               metrics=['accuracy'])
30     return m
31
32     src = np.where(y_tr < 5)[0]
33     tgt = np.where(y_tr >= 5)[0]
34     mask = y_te >= 5
35     Xv, yv = X_te[mask], y_te[mask] - 5
36
37     print('%15s %12s' % ('source images', 'frozen acc'))
38     for n_src in [50, 150, 400, len(src)]:
39         runs = []
40         for seed in range(3):
41             rng = np.random.default_rng(seed)
42             few = np.concatenate([rng.permutation(tgt[y_tr[tgt] == c][:15]
43                                               for c in range(5, 10))])
44             tf.random.set_seed(100 + seed)
45             pre = make_base()
46             head(pre, True).fit(X_tr[rng.permutation(src)[:n_src]],
47                               y_tr[rng.permutation(src)[:n_src]],
48                               epochs=40, batch_size=64, verbose=0)
49             tf.random.set_seed(seed)
50             m = head(pre, False)
51             m.fit(X_tr[few], y_tr[few] - 5, epochs=60, batch_size=16, verbose=0)
52             runs.append(m.evaluate(Xv, yv, verbose=0)[1])
53     print('%15d %12.4f' % (n_src, float(np.mean(runs))))

```

source images	frozen acc
50	0.9196
150	0.9271
400	0.9360
675	0.9390

This is the precondition, stated as a number

The frozen features get steadily better as the source task grows, which is the mechanism working exactly as advertised. It simply has not got far enough by 675 images to beat a network trained directly on the target. Now extrapolate: ImageNet is 1.2 million images across a thousand classes. *That* is why transfer learning is the default in real work, and why this page could not demonstrate it honestly without a network connection.

The practical rule: transfer when your source is far larger and richer than your target. When it is not, you are just borrowing somebody else's constraints.

Step 3: fine-tuning

Once the new head has stopped being random, you can unfreeze the base and continue at a learning rate one to two orders of magnitude smaller. The order matters: unfreeze while the head is still random and the large gradients coming back from it will destroy the pretrained weights before they are ever used.

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'

```

```

3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 from sklearn.datasets import load_digits
9 from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 def make_base():
17     return tf.keras.Sequential([
18         tf.keras.layers.Input(shape=(8, 8, 1)),
19         tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
20         tf.keras.layers.Conv2D(32, 3, padding='same', activation='relu'),
21         tf.keras.layers.MaxPooling2D(2),
22     ], name='features')
23
24 tf.random.set_seed(42)
25 base = make_base()
26 src = y_tr < 5
27 warm = tf.keras.Sequential([base, tf.keras.layers.Flatten(),
28                             tf.keras.layers.Dense(5, activation='softmax')])
29 warm.compile(optimizer='adam', loss='sparse_categorical_crossentropy',
30             metrics=['accuracy'])
31 warm.fit(X_tr[src], y_tr[src], epochs=40, batch_size=64, verbose=0)
32
33 rng = np.random.default_rng(0)
34 tgt = np.where(y_tr >= 5)[0]
35 few = np.concatenate([rng.permutation(tgt[y_tr[tgt] == c])[:15]
36                        for c in range(5, 10)])
37 Xf, yf = X_tr[few], y_tr[few] - 5
38 mask = y_te >= 5
39 Xv, yv = X_te[mask], y_te[mask] - 5
40
41 tf.random.set_seed(7)
42 base.trainable = False
43 model = tf.keras.Sequential([base, tf.keras.layers.Flatten(),
44                             tf.keras.layers.Dense(5, activation='softmax')])
45 model.compile(optimizer='adam', loss='sparse_categorical_crossentropy',
46             metrics=['accuracy'])
47 model.fit(Xf, yf, epochs=60, batch_size=16, verbose=0)
48 print('after training the head only %.4f'
49       % model.evaluate(Xv, yv, verbose=0)[1])
50 before = model.predict(Xv, verbose=0).argmax(1)
51 kernel_before = base.get_weights()[0].copy()
52
53 base.trainable = True
54 model.compile(optimizer=tf.keras.optimizers.Adam(1e-4),
55             loss='sparse_categorical_crossentropy', metrics=['accuracy'])
56 model.fit(Xf, yf, epochs=40, batch_size=16, verbose=0)

```

```

57 print('after fine-tuning at 1e-4      %.4f'
58       % model.evaluate(Xv, yv, verbose=0)[1])
59
60 after = model.predict(Xv, verbose=0).argmax(1)
61 print('\npredictions that changed: %d of %d' % ((before != after).sum(),
62                                               len(yv)))
63 print('largest weight change in the base: %.6f'
64       % np.abs(base.get_weights()[0] - kernel_before).max())

```

```

after training the head only    0.9509
after fine-tuning at 1e-4      0.9554

predictions that changed: 1 of 224
largest weight change in the base: 0.016407

```

Fine-tuning at 1e-4 for forty epochs on seventy-five images is a nudge, not a retraining. The largest weight in the base moved by about a hundredth, and exactly one test prediction changed. That is the intended behaviour: if unfreezing moves your accuracy dramatically, the rate is too high and you are erasing the thing you came for.

⚠ Watch Out

Recompile after changing trainable

Keras decides which weights the optimiser will touch when you compile. Flip `base.trainable` and skip the recompile and nothing changes, the model trains exactly as it did before, and you spend an afternoon wondering why fine-tuning does nothing.

What the first layer actually learned

```

1  import os
2  os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3  import numpy as np
4  import tensorflow as tf
5  tf.random.set_seed(42)
6  tf.get_logger().setLevel('ERROR')
7  import numpy as np
8  from sklearn.datasets import load_digits
9  from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 tf.random.set_seed(42)
17 m = tf.keras.Sequential([
18     tf.keras.layers.Input(shape=(8, 8, 1)),
19     tf.keras.layers.Conv2D(8, 3, padding='same', activation='relu',
20                           name='first'),
21     tf.keras.layers.MaxPooling2D(2),
22     tf.keras.layers.Flatten(),
23     tf.keras.layers.Dense(10, activation='softmax'),

```

```

24 ])
25 m.compile(optimizer='adam', loss='sparse_categorical_crossentropy',
26           metrics=['accuracy'])
27 m.fit(X_tr, y_tr, epochs=40, batch_size=64, verbose=0)
28
29 k = m.get_layer('first').get_weights()[0]      # (3, 3, 1, 8)
30 print('four of the eight learned 3x3 kernels:')
31 for f in range(4):
32     w = k[:, :, 0, f]
33     print('\nfilter %d    (sum %+.2f)' % (f, w.sum()))
34     for row in w:
35         print('    ' + ' '.join('%+.2f' % v for v in row))

```

```

four of the eight learned 3x3 kernels:

filter 0    (sum +0.11)
+0.62 +0.57 +0.47
+0.28 -0.01 +0.18
-0.53 -0.82 -0.65

filter 1    (sum -0.09)
+0.43 -0.57 +0.52
-0.60 -0.69 +0.23
-0.12 +0.60 +0.13

filter 2    (sum +0.72)
+0.62 -0.48 -0.84
+0.42 +0.14 -0.24
+0.09 +0.76 +0.25

filter 3    (sum +1.59)
+0.26 +0.55 +0.35
+0.53 +0.47 +0.32
+0.22 -0.44 -0.67

```

Read the signs rather than the magnitudes. A filter with negatives on one side and positives on the other is an edge detector oriented along that axis; one that is positive in the middle and negative around it is a blob detector. Nobody told the network to learn these. They are what minimising the loss produced, and they are the reason the features transfer to a task the base never saw.

Day 5 takeaway

Day 6 Autoencoders and Learned Representations

Compression without labels, denoising, and anomaly scores from error

Week 10 ended by finding unusual customers with PCA: rebuild each row from two components and flag whichever rows rebuild worst. An autoencoder is that idea with a network in place of the linear projection, and it is the bridge between supervised deep learning and everything in weeks 9 and 10.

An autoencoder on the churn data

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 import pandas as pd
9
10 df = pd.read_csv('customers.csv').drop_duplicates().copy()
11 df['contract'] = df['contract'].str.strip().str.title()
12 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
13
14 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
15 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
16 from sklearn.compose import ColumnTransformer
17 from sklearn.pipeline import Pipeline
18 from sklearn.impute import SimpleImputer
19 from sklearn.preprocessing import StandardScaler, OneHotEncoder
20
21 prep = ColumnTransformer([
22     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
23                       ('s', StandardScaler())])), NUM),
24     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
25                       ('o', OneHotEncoder(handle_unknown='ignore',
26                                           sparse_output=False))])), CAT),
27 ])
28 Z = prep.fit_transform(df).astype('float32')
29 print('encoded width %d columns' % Z.shape[1])
30
31 tf.random.set_seed(42)
32 auto = tf.keras.Sequential([
33     tf.keras.layers.Input(shape=(Z.shape[1],)),
34     tf.keras.layers.Dense(8, activation='relu'),
35     tf.keras.layers.Dense(2, activation='relu', name='bottleneck'),
36     tf.keras.layers.Dense(8, activation='relu'),
37     tf.keras.layers.Dense(Z.shape[1]),
38 ])
39 auto.compile(optimizer='adam', loss='mse')
40 auto.fit(Z, Z, epochs=120, batch_size=64, verbose=0)
41
42 error = ((auto.predict(Z, verbose=0) - Z) ** 2).mean(axis=1)
43 print('reconstruction error: median %.4f, 99th %.4f, max %.4f'
44       % (np.median(error), np.percentile(error, 99), error.max()))
45
46 worst = np.argsort(error)[-1][:5]
47 print('\nthe five customers the network cannot rebuild:')
48 print(df.iloc[worst][['customer_id'] + NUM].to_string(index=False))

```

```
encoded width 15 columns
```

```
reconstruction error: median 0.1285, 99th 0.2909, max 29.5497
```

```

the five customers the network cannot rebuild:
customer_id  tenure_months  support_calls  monthly_charges
C01643       1              0              50.61
C00487       6              0              48.84
C00988       9              0              49.83
C00451       3              0              57.29
C02559       8              0              48.30

```

Same idea as PCA, one important difference

PCA can only rebuild a row as a weighted sum of fixed directions. An autoencoder can bend, so in principle it captures structure that is not a flat subspace, for instance that high charges are normal *on a fibre contract* and unusual otherwise, which no single linear direction expresses. The cost is that you now have a network to train, a seed to fix, and no equivalent of explained variance to report.

Note that the customer at the top of that list is the one with 97 support calls, the same row week 10 found with PCA. Two quite different methods agreeing on the most unusual customer is the reassuring outcome.

Is the flexibility earning anything?

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 import pandas as pd
9
10 df = pd.read_csv('customers.csv').drop_duplicates().copy()
11 df['contract'] = df['contract'].str.strip().str.title()
12 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
13
14 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
15 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
16 from sklearn.compose import ColumnTransformer
17 from sklearn.pipeline import Pipeline
18 from sklearn.impute import SimpleImputer
19 from sklearn.preprocessing import StandardScaler, OneHotEncoder
20 from sklearn.decomposition import PCA
21
22 prep = ColumnTransformer([
23     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
24     ('s', StandardScaler())])), NUM),
25     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
26     ('o', OneHotEncoder(handle_unknown='ignore',
27     sparse_output=False))])), CAT),
28 ])
29 Z = prep.fit_transform(df).astype('float32')
30

```

```

31 pca = PCA(n_components=2, random_state=42).fit(Z)
32 pca_err = ((pca.inverse_transform(pca.transform(Z)) - Z) ** 2).mean()
33
34 tf.random.set_seed(42)
35 auto = tf.keras.Sequential([
36     tf.keras.layers.Input(shape=(Z.shape[1],)),
37     tf.keras.layers.Dense(8, activation='relu'),
38     tf.keras.layers.Dense(2, activation='relu'),
39     tf.keras.layers.Dense(8, activation='relu'),
40     tf.keras.layers.Dense(Z.shape[1]),
41 ])
42 auto.compile(optimizer='adam', loss='mse')
43 auto.fit(Z, Z, epochs=120, batch_size=64, verbose=0)
44 auto_err = float(auto.evaluate(Z, Z, verbose=0))
45
46 print('mean squared reconstruction error, 2 dimensions')
47 print('  PCA           %.4f' % pca_err)
48 print('  autoencoder  %.4f' % auto_err)

```

```

mean squared reconstruction error, 2 dimensions
PCA           0.1930
autoencoder 0.2305

```

PCA wins. Given the same two-dimensional budget, the method with a closed form solution and no hyperparameters rebuilds these rows better than the network does, and it does so in milliseconds, deterministically, with an explained-variance figure you can put in a report.

⚠ Watch Out

Fifteen columns is not where an autoencoder shines

The flexibility that makes autoencoders powerful needs something to be flexible about. After one-hot encoding, this data is fifteen columns of largely linear structure, and a linear method models linear structure optimally by definition. Autoencoders start to win on inputs with hundreds or thousands of dimensions and strong nonlinear structure, images, audio, embeddings. Run the comparison before you assume which side you are on.

A convolutional autoencoder on the digits

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 from sklearn.datasets import load_digits
9 from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')

```

```

14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                           stratify=y, random_state=42)
16 tf.random.set_seed(42)
17 auto = tf.keras.Sequential([
18     tf.keras.layers.Input(shape=(8, 8, 1)),
19     tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
20     tf.keras.layers.MaxPooling2D(2),                # 4 x 4 x 16
21     tf.keras.layers.Conv2D(8, 3, padding='same', activation='relu'),
22     tf.keras.layers.UpSampling2D(2),                # back to 8 x 8
23     tf.keras.layers.Conv2D(1, 3, padding='same', activation='sigmoid'),
24 ])
25 auto.compile(optimizer='adam', loss='mse')
26 auto.fit(X_tr, X_tr, epochs=60, batch_size=64, verbose=0)
27
28 rebuilt = auto.predict(X_te[:1], verbose=0)
29 chars = ' .:-+*#%@'
30 def show(a):
31     a = (a - a.min()) / (a.max() - a.min() + 1e-9)
32     for row in a:
33         print(''.join(chars[min(9, int(v * 9.999))] for v in row))
34
35 print('test error %.4f' % auto.evaluate(X_te, X_te, verbose=0))
36 print('\noriginal, a %d:' % y_te[0])
37 show(X_te[0, :, :, 0])
38 print('\nrebuilt from a 4 x 4 x 8 code:')
39 show(rebuilt[0, :, :, 0])

```

```
test error 0.0088
```

```
original, a 1:
```

```

=@@=
@@@.
.@@*
.@@#
-@@%
@@*
@@*
@@*
=@@

```

```
rebuilt from a 4 x 4 x 8 code:
```

```

=@@=
@@@:
@@#
.@@%
.@@%
@@*
%@@
=@@#

```

Denoising, which is the version people actually use

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np

```

```

4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 from sklearn.datasets import load_digits
9 from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16 rng = np.random.default_rng(0)
17 noisy_tr = np.clip(X_tr + rng.normal(0, 0.3, X_tr.shape), 0, 1).astype('float32')
18 noisy_te = np.clip(X_te + rng.normal(0, 0.3, X_te.shape), 0, 1).astype('float32')
19
20 tf.random.set_seed(42)
21 auto = tf.keras.Sequential([
22     tf.keras.layers.Input(shape=(8, 8, 1)),
23     tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
24     tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
25     tf.keras.layers.Conv2D(1, 3, padding='same', activation='sigmoid'),
26 ])
27 auto.compile(optimizer='adam', loss='mse')
28 # input: the corrupted image. target: the clean one.
29 auto.fit(noisy_tr, X_tr, epochs=60, batch_size=64, verbose=0)
30
31 cleaned = auto.predict(noisy_te[:1], verbose=0)
32 chars = ' .:-+*#%#@'
33 def show(a):
34     a = (a - a.min()) / (a.max() - a.min() + 1e-9)
35     for row in a:
36         print(''.join(chars[min(9, int(v * 9.999))] for v in row))
37
38 print('noisy input:')
39 show(noisy_te[0, :, :, 0])
40 print('\ndennoised output:')
41 show(cleaned[0, :, :, 0])
42 print('\nclean target:')
43 show(X_te[0, :, :, 0])

```

noisy input:

```

 .@*+= -
+@-%: .#
==@@= -
 #@*
==@@% :#
 @*# :
 @%@ .
 -*@

```

denoised output:

```

 %##* :
=@#% .
.%@=

```

```

*@@+
%@@*
%%%*
#%%%
-%%%

clean target:
=@@=
@@@.
.@@*
.@@#
-@@%
@@*
@@*
=@@

```

Nothing about the architecture changed. The only difference is what was put on each side of `fit`: corrupted input, clean target. That small substitution is the trick behind a surprising amount of modern self-supervised learning, hide part of the input, train the model to reconstruct it, and the representation you get for free turns out to be useful for tasks you have not thought of yet.

Using the bottleneck as features

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 from sklearn.datasets import load_digits
9 from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                           stratify=y, random_state=42)
16 from sklearn.linear_model import LogisticRegression
17
18 tf.random.set_seed(42)
19 encoder = tf.keras.Sequential([
20     tf.keras.layers.Input(shape=(8, 8, 1)),
21     tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
22     tf.keras.layers.MaxPooling2D(2),
23     tf.keras.layers.Flatten(),
24     tf.keras.layers.Dense(12, activation='relu', name='code'),
25 ])
26 decoder = tf.keras.Sequential([
27     tf.keras.layers.Input(shape=(12,)),
28     tf.keras.layers.Dense(64, activation='relu'),
29     tf.keras.layers.Dense(64, activation='sigmoid'),
30     tf.keras.layers.Reshape((8, 8, 1)),
31 ])

```

```

32 auto = tf.keras.Sequential([encoder, decoder])
33 auto.compile(optimizer='adam', loss='mse')
34 auto.fit(X_tr, X_tr, epochs=80, batch_size=64, verbose=0)
35
36 # the labels were never shown to the autoencoder
37 code_tr = encoder.predict(X_tr, verbose=0)
38 code_te = encoder.predict(X_te, verbose=0)
39
40 flat = LogisticRegression(max_iter=2000).fit(
41     X_tr.reshape(len(X_tr), -1), y_tr)
42 coded = LogisticRegression(max_iter=2000).fit(code_tr, y_tr)
43 print('logistic on 64 raw pixels      %.4f'
44       % flat.score(X_te.reshape(len(X_te), -1), y_te))
45 print('logistic on 12 learned codes  %.4f'
46       % coded.score(code_te, y_te))

```

```

logistic on 64 raw pixels      0.9622
logistic on 12 learned codes  0.9000

```

Twelve numbers, learned without ever seeing a label, carry most, not all, of what sixty-four pixels carry. That is the honest shape of the result: a five-fold compression costing a few points of accuracy. The reason anybody accepts that trade is the case where labels are the scarce thing: you can train the encoder on a million unlabelled images and then fit the classifier on the two hundred you could afford to label.

Day 6 takeaway

Day 7 Does Deep Learning Earn Its Place?

A complete image workflow, and the reckoning week 11 promised

Week 11 day 3 promised this reckoning. A neural network matched logistic regression on the churn data while carrying several hundred times the parameters. This week it has been winning. The difference between those two situations is the most useful thing in either week.

The complete workflow, end to end

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 from sklearn.datasets import load_digits
9 from sklearn.model_selection import train_test_split
10
11 d = load_digits()
12 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
13 y = d.target.astype('int32')

```

```

14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                           stratify=y, random_state=42)
16 import json
17 import time
18
19 X_fit, X_val = X_tr[:1000], X_tr[1000:]
20 y_fit, y_val = y_tr[:1000], y_tr[1000:]
21
22 tf.random.set_seed(42)
23 model = tf.keras.Sequential([
24     tf.keras.layers.Input(shape=(8, 8, 1)),
25     tf.keras.layers.RandomTranslation(0.1, 0.1, fill_mode='constant'),
26     tf.keras.layers.Conv2D(16, 3, padding='same', activation='relu'),
27     tf.keras.layers.BatchNormalization(),
28     tf.keras.layers.MaxPooling2D(2),
29     tf.keras.layers.Conv2D(32, 3, padding='same', activation='relu'),
30     tf.keras.layers.BatchNormalization(),
31     tf.keras.layers.Flatten(),           # day 3: position matters at 8 x 8
32     tf.keras.layers.Dropout(0.3),
33     tf.keras.layers.Dense(10, activation='softmax'),
34 ])
35 model.compile(optimizer='adam', loss='sparse_categorical_crossentropy',
36              metrics=['accuracy'])
37
38 start = time.time()
39 hist = model.fit(X_fit, y_fit, validation_data=(X_val, y_val),
40                 epochs=200, batch_size=64, verbose=0,
41                 callbacks=[tf.keras.callbacks.EarlyStopping(
42                     patience=20, restore_best_weights=True)])
43
44 print('epochs run %d of 200 in %.1fs' % (len(hist.history['loss']),
45                                         time.time() - start))
46 print('test accuracy %.4f' % model.evaluate(X_te, y_te, verbose=0)[1])
47
48 model.save('digits_cnn.keras')
49 json.dump({'input': [8, 8, 1], 'scaling': 'divide by 16',
50           'classes': list(range(10))},
51           open('digits_cnn.json', 'w'))
52
53 reloaded = tf.keras.models.load_model('digits_cnn.keras')
54 print('reloaded and agrees:',
55       bool((reloaded.predict(X_te[:50], verbose=0).argmax(1)
56             == model.predict(X_te[:50], verbose=0).argmax(1)).all()))

```

```

epochs run 67 of 200 in 10.0s
test accuracy 0.9867
reloaded and agrees: True

```

⚠ Watch Out**Save the preprocessing, not just the weights**

The `.keras` file contains the architecture and the weights. It does not contain the fact that you divided by 16, or which index means which digit. Serve the model without those and it will happily return confident nonsense for inputs on the wrong scale. The small JSON file beside it is not optional, week 15 turns it into a proper contract.

What it cost against what it bought

```

1 import numpy as np
2 from sklearn.datasets import load_digits
3 from sklearn.model_selection import train_test_split
4
5 d = load_digits()
6 X = (d.images / 16.0).astype('float32')[..., np.newaxis]
7 y = d.target.astype('int32')
8 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
9                                         stratify=y, random_state=42)
10 from sklearn.linear_model import LogisticRegression
11 from sklearn.ensemble import HistGradientBoostingClassifier, RandomForestClassifier
12 from sklearn.svm import SVC
13 import time
14
15 flat_tr = X_tr.reshape(len(X_tr), -1)
16 flat_te = X_te.reshape(len(X_te), -1)
17
18 print('%-24s %10s %10s' % ('', 'accuracy', 'fit time'))
19 models = [('logistic regression', LogisticRegression(max_iter=2000)),
20           ('random forest', RandomForestClassifier(random_state=42)),
21           ('gradient boosting', HistGradientBoostingClassifier(random_state=42)),
22           ('support vector machine', SVC())]
23 for name, m in models:
24     t = time.time()
25     m.fit(flat_tr, y_tr)
26     print('%-24s %10.4f %9.1fs'
27           % (name, m.score(flat_te, y_te), time.time() - t))

```

	accuracy	fit time
logistic regression	0.9622	0.1s
random forest	0.9600	0.3s
gradient boosting	0.9600	3.7s
support vector machine	0.9911	0.1s

A support vector machine on the raw flattened pixels, fitted in about a tenth of a second with default settings and no tuning whatsoever, beats the convolutional network that took roughly a hundred times longer to fit and needed batch normalisation, dropout, augmentation and early stopping to get there.

⚠ Watch Out

This is the result, not a rhetorical flourish

It would have been easy to leave this comparison out. The reason it is here is that the same comparison, run on 32×32 colour photographs or 224×224 medical scans, comes out the opposite way and it is not close, and the only way to know which situation you are in is to fit the cheap model and look. Eight by eight greyscale digits are, in the end, a 64-dimensional tabular problem wearing a picture's clothing.

When deep learning earns its place

Situation	Reach for	Because
Tabular data, thousands of rows	Gradient boosting	Wins on accuracy, trains in seconds, needs no scaling
Tabular data, millions of rows and rich categoricals	Boosting first, network second	Networks become competitive, but rarely by much
Images, audio, video	Convolutional network, pretrained	The structure of the data is spatial; nothing else exploits that
Text	Pretrained transformer	Week 13
Fewer than a few hundred labels	Transfer learning, or classical	There is not enough signal to fit a network from scratch
The decision must be explained	Linear or tree model	"The network said so" fails an audit
Latency budget in single milliseconds	Linear model	A dot product is hard to beat

The honest summary of two weeks

Deep learning is not a better version of machine learning. It is the answer to a specific question: *what do I do when the raw input has structure that no feature I can write down captures?* A picture has that problem. A spectrogram has it. A sentence has it. A spreadsheet of customer attributes, where somebody has already done the hard work of deciding that `tenure_months` is a column, generally does not. The feature engineering is finished before you arrive.

Before you ship an image model

1. Check the class balance and look at a sample from each class by eye. Mislabelled images are far more common than mislabelled spreadsheet rows.
2. Split before augmenting. Augmenting first and splitting afterwards puts shifted copies of the same picture on both sides, the image version of the leakage in week 4.
3. Verify that your augmentations preserve the label in your domain.
4. Start from pretrained weights unless you have a good reason not to.
5. Train the head with the base frozen, then fine-tune at a low rate, recompiling in between.
6. Read the confusion matrix and look at the actual images the model gets wrong.
7. Fit gradient boosting on the flattened pixels as a floor. If it wins, your images are probably too small or too few for a network.
8. Save the preprocessing contract alongside the weights.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. Convolution uses far fewer weights than a dense layer because:
 - a) It uses integers
 - b) The images are smaller
 - c) The same small kernel is applied at every position, so a pattern is learned once rather than per location
 - d) It skips pixels
2. Convolution is *equivariant* to translation, which means:
 - a) It ignores position
 - b) It requires centred images
 - c) The output is unchanged when the input moves
 - d) The output moves when the input moves
3. Stacking two 3×3 convolutions rather than one 5×5 gives you:
 - a) The same receptive field with fewer weights and an extra non-linearity
 - b) Faster inference only
 - c) Identical behaviour
 - d) A smaller receptive field
4. On 8×8 digits, replacing `GlobalAveragePooling2D` with `Flatten` recovered eleven points of accuracy because:
 - a) Flatten has more parameters
 - b) Averaging over positions discards where each response occurred, and at that resolution position is most of the signal
 - c) Global pooling is buggy
 - d) The model was undertrained
5. Data augmentation is best described as:
 - a) A form of normalisation
 - b) Getting more data
 - c) Regularisation that encodes a specific belief about which transformations do not change the label
 - d) A speed optimisation
6. Transfer learning failed to beat training from scratch on the course's digit experiment because:
 - a) Freezing is never useful

- b) The learning rate was wrong
- c) The code was wrong
- d) The source task was far too small for its learned features to beat what the target's own data could learn

7. After changing `base.trainable` you must:

- a) Recompile the model, or the optimiser will keep training exactly the same weights as before
- b) Reset the seed
- c) Lower the batch size
- d) Reload the data

8. A support vector machine on flattened pixels beat the convolutional network. The lesson is:

- a) Never use CNNs
- b) Fit the cheap baseline and look, at this resolution the problem is effectively tabular
- c) The CNN needed more layers
- d) Accuracy is the wrong metric

Answers: 1c, 2d, 3a, 4b, 5c, 6d, 7a, 8b. Anything you got wrong points at the day to re-read, not at a gap in ability.

Text, Embeddings and Sequence Models

Week 13

Turn text into numbers, discover what a bag of words cannot express, and find out whether a sequence model is worth it.

OBJECTIVES

This week is built around one corpus with two questions in it. The first, which category is this ticket, a bag of words answers easily. The second, is the customer's problem over. It answers worse than a coin, because the answer lives in which words sit next to which. Fixing that takes us from n-grams through embeddings and recurrent layers to attention, and to an honest account of when each is worth its cost.

Day 1 Turning Text Into Numbers

Tokens, vocabularies, sparse matrices and the information you throw away

Every model so far has been handed numbers. Text is not numbers, and the whole of this week is about the step in between, because that step, not the choice of classifier, is what decides whether the model works.

The corpus

The broadband company whose customers you have been predicting for twelve weeks has a support inbox. This is it: 4,200 tickets, each with a free-text body, a category, and a flag for whether the customer is describing a problem that is now over. Generate it the same way you generated `customers.csv` in week 1.

```

1 import io
2 import os
3
4 import numpy as np
5 import pandas as pd
6
7 SEED = 13
8 N = 4200
9
10 THING = ["router", "connection", "line", "broadband", "wifi", "signal",
11          "hub", "internet"]
12
13 # The two halves of the construction. A positive state word with no negation
14 # means the problem is over; the same word negated means it is not.
15 GOOD_STATE = ["working", "stable", "fixed", "fine", "sorted", "reliable",
```

```

16         "back to normal"]
17 BAD_STATE = ["slow", "down", "unstable", "broken", "patchy", "unreliable",
18             "dropping out"]
19
20 # Every one of these appears equally often on both sides of the label.
21 MODIFIER = ["", "still ", "not ", "no longer "]
22 NEGATING = {"not ", "no longer "} # these flip the state word
23
24 FRAME = [
25     "the {thing} is {mod}{state}",
26     "my {thing} is {mod}{state}",
27     "it is {mod}{state}",
28     "the {thing} has been {mod}{state} {when}",
29     "as of this morning the {thing} is {mod}{state}",
30     "{when} the {thing} is {mod}{state}",
31 ]
32
33 WHEN = ["every evening", "since tuesday", "all weekend", "for three weeks",
34         "in the mornings", "since the storm", "most nights"]
35
36 # Context sentences carry the category and nothing else -- no resolution
37 # information at all, so the two labels stay genuinely different tasks.
38 CONTEXT = {
39     "technical": [
40         "the engineer visited on tuesday",
41         "this is the third time i have written about the hub",
42         "i have restarted the equipment twice",
43         "the speed test reads well under what i pay for",
44         "there is a flashing red light on the box",
45     ],
46     "billing": [
47         "this relates to invoice 40182",
48         "the account is on the fibre tariff",
49         "i pay by direct debit on the eighth",
50         "there is a charge on the statement i do not recognise",
51         "the promotional discount was supposed to run for a year",
52     ],
53     "cancellation": [
54         "i have been a customer for six years",
55         "i am considering my options",
56         "my contract ends next month",
57         "a competitor has offered me a better package",
58         "i would like to know the termination fee",
59     ],
60 }
61
62 GREETING = ["", "", "", "hi team, ", "hello, ", "hi, ", "morning, ", "hi there "]
63 SIGNOFF = ["", "", "", "", " thanks", " please advise", " regards", " ta",
64           " please call me back"]
65
66 TYPOS = {"the": "teh", "and": "adn", "connection": "conection",
67         "internet": "internt", "please": "plese", "invoice": "invoice"}
68
69

```

```

70 def _core(rng, resolved):
71     """One clause whose meaning comes from the modifier-plus-state pairing.
72
73     Pick the modifier first and independently of the label, then choose the
74     state word that makes the clause mean what the label says. That ordering
75     is what keeps every modifier balanced across both classes -- generate it
76     the other way round and the modifier leaks the answer to a unigram model.
77     """
78     mod = rng.choice(MODIFIER)
79     good_state = resolved if mod not in NEGATING else not resolved
80     state = rng.choice(GOOD_STATE if good_state else BAD_STATE)
81     return rng.choice(FRAME).format(thing=rng.choice(THING), mod=mod,
82                                     state=state, when=rng.choice(WHEN))
83
84
85 def _rough(rng, text):
86     """Put the mess back in: casing, punctuation, typos."""
87     if rng.random() < 0.10:
88         words = text.split()
89         for i, w in enumerate(words):
90             if w in TYPOS and rng.random() < 0.5:
91                 words[i] = TYPOS[w]
92         text = " ".join(words)
93     if rng.random() < 0.12:
94         text = text.upper()
95     elif rng.random() < 0.25:
96         text = text.capitalize()
97     text = rng.choice(GREETING) + text + rng.choice(SIGNOFF)
98     text += rng.choice([".", "!", " ", "!", "!!", "...", " ."])
99     if rng.random() < 0.08:
100         text = " " + text + " "
101     return text
102
103
104 def build(seed=SEED, n=N):
105     rng = np.random.default_rng(seed)
106     cats = rng.choice(["technical", "billing", "cancellation"], size=n,
107                       p=[0.52, 0.33, 0.15])
108     rows = []
109     for i, cat in enumerate(cats):
110         resolved = int(rng.random() < 0.42)
111         parts = [_core(rng, resolved)]
112         # A second clause of the same polarity. Mixed polarity would leave the
113         # ticket with no correct answer, so it never happens.
114         if rng.random() < 0.30:
115             parts.append(_core(rng, resolved))
116         # The category sentence, which says nothing about resolution.
117         if rng.random() < 0.75:
118             where = rng.integers(0, len(parts) + 1)
119             parts.insert(int(where), rng.choice(CONTEXT[cat]))
120         text = _rough(rng, ". ".join(parts))
121         rows.append({"ticket_id": "T%05d" % (i + 1), "text": text,
122                     "category": cat, "resolved": resolved})
123

```

```

124     df = pd.DataFrame(rows)
125
126     # ~4% label noise, so a perfect score is impossible and anybody who gets
127     # one goes looking for the leak instead of celebrating.
128     flip = rng.random(len(df)) < 0.04
129     df.loc[flip, "resolved"] = 1 - df.loc[flip, "resolved"]
130
131     # a few empty and near-empty tickets
132     blank = rng.choice(len(df), size=18, replace=False)
133     df.loc[blank[:9], "text"] = ""
134     df.loc[blank[9:], "text"] = rng.choice(["?", "...", "help", "hi", "!!"],
135                                           size=len(blank) - 9)
136
137     # double submissions
138     dupes = df.iloc[rng.choice(len(df), size=90, replace=False)]
139     df = pd.concat([df, dupes], ignore_index=True)
140     df = df.iloc[rng.permutation(len(df))].reset_index(drop=True)
141     return df
142
143 build().to_csv('tickets.csv', index=False, encoding='utf-8')

```

```

1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 print('tickets %d' % len(tickets))
8 print(tickets[['ticket_id', 'category', 'resolved']].head())
9 print('\nresolved rate %.3f' % tickets['resolved'].mean())
10 print('categories:')
11 print(tickets['category'].value_counts())

```

```

tickets 4200
   ticket_id  category  resolved
0    T02210    billing         0
1    T00107    billing         1
2    T00680    billing         1
3    T03436  cancellation         0
4    T02257    billing         1

```

```

resolved rate 0.425
categories:
category
technical      2144
billing        1429
cancellation    627
Name: count, dtype: int64

```

```

1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()

```

```

6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 for _, r in tickets.head(6).iterrows():
8     print('%-12s resolved=%d' %s'
9           % (r['category'], r['resolved'], r['text'][:78]))

[billing      resolved=0] AS OF THIS MORNING THE HUB IS PATCHY. MOST NIGHTS THE HUB IS NO
    ↳ LONGER WORKING
[billing      resolved=1] hello, my hub is not unstable. my wifi is still fixed.
[billing      resolved=1] hello, There is a charge on the statement i do not recognise. it
    ↳ is working. t
[cancellation resolved=0] the hub is down please advise
[billing      resolved=1] hi, the account is on the fibre tariff. as of this morning the
    ↳ router is no lo
[billing      resolved=0] in the mornings the wifi is unstable please advise .

```

⚠ Watch Out

Look at the raw strings before you clean anything

Shouting in capitals, doubled exclamation marks, leading spaces, typos, and nine tickets that are empty. Every one of those is in the file on purpose, because every one of them is in real inboxes. The instinct to normalise it all away immediately is the right instinct and the wrong first move, some of that mess is signal. Angry customers do write in capitals.

Tokenising

```

1 text = ' THE ROUTER is Not working!! since tuesday. '
2
3 print('raw      %r' % text)
4 print('stripped  %r' % text.strip())
5 print('lowered   %r' % text.strip().lower())
6 print('naive split %s' % text.strip().lower().split())
7
8 import re
9 tokens = re.findall(r"[a-z0-9']+", text.lower())
10 print('regex     %s' % tokens)

raw      ' THE ROUTER is Not working!! since tuesday. '
stripped 'THE ROUTER is Not working!! since tuesday.'
lowered  'the router is not working!! since tuesday.'
naive split ['the', 'router', 'is', 'not', 'working!!', 'since', 'tuesday.']
regex      ['the', 'router', 'is', 'not', 'working', 'since', 'tuesday']

```

`split()` leaves `working!!` and `tuesday.` as tokens distinct from `working` and `tuesday`, which quietly doubles parts of your vocabulary and halves the evidence for each. The regex keeps letters, digits and apostrophes and drops the rest.

Building a bag of words by hand

```

1 import re
2 from collections import Counter

```

```

3
4 docs = ['the router is not working',
5         'the router is working',
6         'my line is not slow']
7
8 tokenised = [re.findall(r"[a-z0-9"]+", d.lower()) for d in docs]
9 vocab = sorted({t for doc in tokenised for t in doc})
10 index = {t: i for i, t in enumerate(vocab)}
11 print('vocabulary (%d): %s' % (len(vocab), vocab))
12
13 import numpy as np
14 M = np.zeros((len(docs), len(vocab)), dtype=int)
15 for r, doc in enumerate(tokenised):
16     for token, n in Counter(doc).items():
17         M[r, index[token]] = n
18
19 print('\ndocument-term matrix:')
20 print('      ' + ' '.join('%-8s' % v[:7] for v in vocab))
21 for r, row in enumerate(M):
22     print('doc %d ' % r + ' '.join('%-8d' % v for v in row))

```

```
vocabulary (8): ['is', 'line', 'my', 'not', 'router', 'slow', 'the', 'working']
```

```
document-term matrix:
```

	is	line	my	not	router	slow	the	working
doc 0	1	0	0	1	1	0	1	1
doc 1	1	0	0	0	1	0	1	1
doc 2	1	1	1	1	0	1	0	0

Read rows 0 and 1 again

They differ by one word and mean opposite things. In this representation they differ by one column out of eight, and nothing records that **not** was sitting immediately before **working**. Every limitation in the next two days follows from that single fact.

The same thing from scikit-learn

```

1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 from sklearn.feature_extraction.text import CountVectorizer
8
9 vec = CountVectorizer()
10 M = vec.fit_transform(tickets['text'])
11 print('matrix %s, %.4f%% non-zero'
12       % (M.shape, 100 * M.nnz / (M.shape[0] * M.shape[1])))
13 print('stored as %s' % type(M).__name__)
14
15 counts = np.asarray(M.sum(axis=0)).ravel()

```

```

16 names = vec.get_feature_names_out()
17 print('\nmost common tokens:')
18 for i in counts.argsort()[::-1][:12]:
19     print(' %-14s %5d' % (names[i], counts[i]))

```

```

matrix (4200, 130), 10.9416% non-zero
stored as csr_matrix

```

```

most common tokens:

```

```

help          3
invoice       6
internet     21
connection   26
competitor   83
better       83
package      83
offered      83
considering   85
options      85
am           85
customer     92

```

Day 1 takeaway

Day 2 TF-IDF, n-grams and What Order Costs

Two tasks in one corpus: one a bag of words solves, one it cannot

Two tasks live in this corpus. One is to say which category a ticket belongs to; the other is to say whether the customer's problem is resolved. They look equally easy. They are not, and the difference is the most useful thing in this week.

TF-IDF

```

1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 from sklearn.feature_extraction.text import TfidfVectorizer
8
9 vec = TfidfVectorizer()
10 M = vec.fit_transform(tickets['text'])
11 names = vec.get_feature_names_out()
12
13 row = M[0].toarray().ravel()
14 print('ticket: %s' % tickets['text'][0][:70])
15 print('\nhighest weighted terms in it:')
16 for i in row.argsort()[::-1][:6]:
17     if row[i] > 0:

```

```

18     print(' %-14s %.3f (document frequency %.3f)'
19           % (names[i], row[i],
20             (M[:, i] > 0).sum() / M.shape[0]))

```

```
ticket: AS OF THIS MORNING THE HUB IS PATCHY. MOST NIGHTS THE HUB IS NO LONGER
```

```
highest weighted terms in it:
```

Task one: which category?

```

1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 from sklearn.feature_extraction.text import TfidfVectorizer
8 from sklearn.linear_model import LogisticRegression
9 from sklearn.pipeline import make_pipeline
10 from sklearn.model_selection import train_test_split
11
12 Xc_tr, Xc_te, yc_tr, yc_te = train_test_split(
13     tickets['text'], tickets['category'], test_size=0.25,
14     stratify=tickets['category'], random_state=42)
15
16 model = make_pipeline(TfidfVectorizer(),
17                       LogisticRegression(max_iter=2000)).fit(Xc_tr, yc_tr)
18 print('category accuracy %.4f' % model.score(Xc_te, yc_te))
19 print('always guessing the commonest: %.4f'
20       % (yc_te == yc_tr.mode()[0]).mean())

```

```
category accuracy 0.8400
always guessing the commonest: 0.5105
```

A bag of words handles this comfortably. Words like *invoice*, *tariff* and *termination* belong to one category and nowhere else, so their mere presence settles the question and their position in the sentence is irrelevant.

Task two: is the problem resolved?

```

1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 from sklearn.model_selection import train_test_split
8
9 # Drop tickets with no word characters at all -- the empty ones, and the
10 # ones that are only punctuation. Both tokenise to nothing, and day 4
11 # shows what an all-padding sequence does to masked pooling.

```

```

12 tickets = tickets[tickets['text'].str.contains(r'[A-Za-z0-9]',
13                                             regex=True)].reset_index(drop=True)
14
15 X_tr, X_te, y_tr, y_te = train_test_split(
16     tickets['text'], tickets['resolved'], test_size=0.25,
17     stratify=tickets['resolved'], random_state=42)
18 from sklearn.feature_extraction.text import TfidfVectorizer
19 from sklearn.linear_model import LogisticRegression
20 from sklearn.pipeline import make_pipeline
21 from sklearn.metrics import roc_auc_score
22
23 model = make_pipeline(TfidfVectorizer(),
24                      LogisticRegression(max_iter=2000)).fit(X_tr, y_tr)
25 print('accuracy   %.4f' % model.score(X_te, y_te))
26 print('ROC AUC    %.4f'
27       % roc_auc_score(y_te, model.predict_proba(X_te)[: , 1]))
28 print('\nalways predicting unresolved: %.4f' % (1 - y_te.mean()))

```

```

accuracy   0.5597
ROC AUC    0.5297

```

```

always predicting unresolved: 0.5750

```

⚠ Watch Out

It is worse than a constant

The model that reads the ticket does *worse* than the model that ignores it and always says unresolved. This is not a bug, a bad hyperparameter or too little data. Something about the combination of this representation and this model cannot express the answer, and working out which of the two is at fault is the whole of the rest of this day.

Why, in one table

```

1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 import re
8
9 def has(word):
10     return tickets['text'].str.lower().str.contains(
11         r'\b%s\b' % word, regex=True)
12
13 print('%-14s %10s %10s %12s' % ('token', 'appears', 'resolved', 'unresolved'))
14 for word in ['not', 'still', 'no', 'longer', 'working', 'slow',
15             'fixed', 'broken']:
16     m = has(word)
17     print('%-14s %10d %10d %12d'
18           % (word, m.sum(), (m & (tickets['resolved'] == 1)).sum(),

```

```
19 (m & (tickets['resolved'] == 0)).sum()))
```

token	appears	resolved	unresolved
not	1346	568	778
still	1331	562	769
no	1273	555	718
longer	1273	555	718
working	382	168	214
slow	366	154	212
fixed	379	165	214
broken	357	154	203

Every one of those tokens is split roughly evenly between the two classes. Alone, each is worth nothing. A linear model assigns one weight per token and adds them up, so it has to commit to a single number for *not*, and no single number can be right when *not working* means one thing and *not slow* means the opposite.

So whose fault is it, the features or the model?

There are two ways out of an interaction, and they are worth separating because they lead to different code. Either hand the model the interaction as a feature, or use a model that finds interactions by itself. Try the second one first, keeping the unigrams exactly as they are.

```
1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 from sklearn.model_selection import train_test_split
8
9 # Drop tickets with no word characters at all -- the empty ones, and the
10 # ones that are only punctuation. Both tokenise to nothing, and day 4
11 # shows what an all-padding sequence does to masked pooling.
12 tickets = tickets[tickets['text'].str.contains(r'[A-Za-z0-9]',
13                                             regex=True)].reset_index(drop=True)
14
15 X_tr, X_te, y_tr, y_te = train_test_split(
16     tickets['text'], tickets['resolved'], test_size=0.25,
17     stratify=tickets['resolved'], random_state=42)
18 from sklearn.feature_extraction.text import CountVectorizer
19 from sklearn.linear_model import LogisticRegression
20 from sklearn.ensemble import HistGradientBoostingClassifier
21 from sklearn.neural_network import MLPClassifier
22 from sklearn.pipeline import make_pipeline
23 from sklearn.preprocessing import FunctionTransformer
24
25 # HistGradientBoosting needs a dense array, and this vocabulary is
26 # small enough that densifying it is safe. On a real corpus it is not.
27 densify = FunctionTransformer(lambda M: np.asarray(M.todense()),
28                             accept_sparse=True)
29
30 print('unigram counts only -- no bigrams anywhere')
31 print('%-32s %10s' % ('model', 'accuracy'))
```

```

32 for name, clf, dense in [
33     ('logistic regression (linear)',
34      LogisticRegression(max_iter=2000), False),
35     ('gradient boosting',
36      HistGradientBoostingClassifier(random_state=42), True),
37     ('neural net (32, 16)',
38      MLPClassifier((32, 16), max_iter=800, random_state=42), False)]:
39     steps = [CountVectorizer()] + ([densify] if dense else []) + [clf]
40     m = make_pipeline(*steps).fit(X_tr, y_tr)
41     print('%-32s %10.4f' % (name, m.score(X_te, y_te)))
42 print('%-32s %10.4f' % ('always predicting unresolved', 1 - y_te.mean()))

```

unigram counts only -- no bigrams anywhere	
model	accuracy
logistic regression (linear)	0.5664
gradient boosting	0.9026
neural net (32, 16)	0.8672
always predicting unresolved	0.5750

The information was there the whole time

Nothing about the features changed. The same unigram counts that left logistic regression below a constant prediction take gradient boosting most of the way to the ceiling, because a tree can split on *not* and then split again on *slow* inside that branch, which is precisely what representing an interaction means.

So the bag of words did *not* throw away the answer. It threw away the *convenience* of the answer, and a linear model is a model that can only use convenient answers. Keep the two ideas separate when you diagnose your own models: *the features do not contain it* and *my model cannot reach it* look identical from the accuracy score and call for opposite fixes.

The other way out: n-grams

```

1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 from sklearn.model_selection import train_test_split
8
9 # Drop tickets with no word characters at all -- the empty ones, and the
10 # ones that are only punctuation. Both tokenise to nothing, and day 4
11 # shows what an all-padding sequence does to masked pooling.
12 tickets = tickets[tickets['text'].str.contains(r'[A-Za-z0-9]',
13                                                regex=True)].reset_index(drop=True)
14
15 X_tr, X_te, y_tr, y_te = train_test_split(
16     tickets['text'], tickets['resolved'], test_size=0.25,
17     stratify=tickets['resolved'], random_state=42)
18 from sklearn.feature_extraction.text import TfidfVectorizer
19 from sklearn.linear_model import LogisticRegression
20 from sklearn.pipeline import make_pipeline

```

```

21 from sklearn.metrics import roc_auc_score
22
23 print('%-16s %10s %10s %10s' % ('features', 'columns', 'accuracy', 'auc'))
24 for label, rng in [('unigrams', (1, 1)), ('1-2 grams', (1, 2)),
25                  ('1-3 grams', (1, 3)), ('2-grams only', (2, 2))]:
26     m = make_pipeline(TfidfVectorizer(ngram_range=rng),
27                      LogisticRegression(max_iter=2000)).fit(X_tr, y_tr)
28     print('%-16s %10d %10.4f %10.4f'
29           % (label, len(m[0].vocabulary_), m.score(X_te, y_te),
30             roc_auc_score(y_te, m.predict_proba(X_te)[:, 1])))

```

features	columns	accuracy	auc
unigrams	130	0.5597	0.5297
1-2 grams	1079	0.9542	0.9591
1-3 grams	4139	0.9532	0.9589
2-grams only	949	0.9599	0.9605

What the bigram model learned

```

1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 from sklearn.model_selection import train_test_split
8
9 # Drop tickets with no word characters at all -- the empty ones, and the
10 # ones that are only punctuation. Both tokenise to nothing, and day 4
11 # shows what an all-padding sequence does to masked pooling.
12 tickets = tickets[tickets['text'].str.contains(r'[A-Za-z0-9]',
13                                               regex=True)].reset_index(drop=True)
14
15 X_tr, X_te, y_tr, y_te = train_test_split(
16     tickets['text'], tickets['resolved'], test_size=0.25,
17     stratify=tickets['resolved'], random_state=42)
18 from sklearn.feature_extraction.text import TfidfVectorizer
19 from sklearn.linear_model import LogisticRegression
20 from sklearn.pipeline import make_pipeline
21
22 m = make_pipeline(TfidfVectorizer(ngram_range=(1, 2)),
23                  LogisticRegression(max_iter=2000)).fit(X_tr, y_tr)
24 names = m[0].get_feature_names_out()
25 coef = m[1].coef_.ravel()
26
27 print('strongest evidence the problem IS resolved:')
28 for i in coef.argsort()[::-1][:8]:
29     print(' %-24s %+.3f' % (names[i], coef[i]))
30 print('\nstrongest evidence it is NOT:')
31 for i in coef.argsort()[:8]:
32     print(' %-24s %+.3f' % (names[i], coef[i]))

```

```

strongest evidence the problem IS resolved:
not stable           -3.788
longer fine          -3.734
longer fixed         -3.733
still dropping       -3.729
not sorted           -3.661
longer back          -3.597
still unreliable     -3.586
not back             -3.507

```

```

strongest evidence it is NOT:
not stable           -3.788
longer fine          -3.734
longer fixed         -3.733
still dropping       -3.729
not sorted           -3.661
longer back          -3.597
still unreliable     -3.586
not back             -3.507

```

This is what a good feature looks like

The model has recovered the construction the corpus was built from, and it did so without being told. Bigrams sit at the top of both lists and single words sit near zero. If you ever want to know whether a text representation is working, print the coefficients and see whether they read like something a person would say.

Notice too that bigrams plus a linear model beat gradient boosting on unigrams. Giving the model the *right* interaction directly beats asking it to search for interactions in general, which is week 8's argument, arriving here in a different costume. Feature engineering is not what you do when you cannot afford a better model; it is how you tell the model what you already know.

Day 2 takeaway

Day 3 Embeddings

Dense vectors with a geometry, and what they do and do not fix

Bigrams worked, and they scale badly: every pair you keep is a new column, trigrams are worse, and none of it generalises. A model that has learned *not working* knows nothing about *not functioning*. Embeddings are the alternative.

The problem with one column per word

```

1 import numpy as np
2
3 vocab = ['router', 'hub', 'invoice', 'working', 'broken']
4 eye = np.eye(len(vocab), dtype=int)
5 print('one-hot:')
6 for w, row in zip(vocab, eye):
7     print(' %-9s %s' % (w, row))

```

```

8
9 print('\ncosine similarity between every pair: ', end='')
10 sims = {round(float(a @ b), 3) for i, a in enumerate(eye)
11         for j, b in enumerate(eye) if i != j}
12 print(sims)
13 print('router and hub are exactly as similar as router and invoice.')
```

```

one-hot:
router    [1 0 0 0 0]
hub       [0 1 0 0 0]
invoice   [0 0 1 0 0]
working   [0 0 0 1 0]
broken    [0 0 0 0 1]

cosine similarity between every pair: {0.0}
router and hub are exactly as similar as router and invoice.
```

Learning one

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 import pandas as pd
9
10 tickets = pd.read_csv('tickets.csv')
11 tickets['text'] = tickets['text'].fillna('').str.strip()
12 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
13 from sklearn.model_selection import train_test_split
14
15 # Drop tickets with no word characters at all -- the empty ones, and the
16 # ones that are only punctuation. Both tokenise to nothing, and day 4
17 # shows what an all-padding sequence does to masked pooling.
18 tickets = tickets[tickets['text'].str.contains(r'[A-Za-z0-9]',
19                                                regex=True)].reset_index(drop=True)
20
21 X_tr, X_te, y_tr, y_te = train_test_split(
22     tickets['text'], tickets['resolved'], test_size=0.25,
23     stratify=tickets['resolved'], random_state=42)
24 vec = tf.keras.layers.TextVectorization(max_tokens=400,
25                                         output_sequence_length=24)
26 vec.adapt(X_tr.to_numpy())
27
28 print('vocabulary size %d' % vec.vocabulary_size())
29 print('first 12 tokens: %s' % vec.get_vocabulary()[:12])
30
31 example = tf.constant(['the router is not working',
32                       'the router is working'])
33 print('\nencoded:')
34 print(vec(example).numpy()[:, :10])
```

```

vocabulary size 134
first 12 tokens: ['', ' [UNK]', np.str_('the'), np.str_('is'), np.str_('i'), np.str_('hi'),
    ↪ np.str_('not'), np.str_('morning'), np.str_('this'), np.str_('still'), np.str_('no'),
    ↪ np.str_('longer')]

encoded:
[[ 2 30  3  6 52  0  0  0  0  0]
 [ 2 30  3 52  0  0  0  0  0 0]]

```

⚠ Watch Out

Index 0 is padding and index 1 is the unknown token

TextVectorization reserves both. Every sequence is padded with zeros to a fixed length so it can sit in a rectangular array, and any word not in the vocabulary becomes [UNK]. Forgetting that index 0 is not a real word is how people end up with an embedding for padding that quietly dominates their averages.

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 import pandas as pd
9
10 tickets = pd.read_csv('tickets.csv')
11 tickets['text'] = tickets['text'].fillna('').str.strip()
12 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
13 from sklearn.model_selection import train_test_split
14
15 # Drop tickets with no word characters at all -- the empty ones, and the
16 # ones that are only punctuation. Both tokenize to nothing, and day 4
17 # shows what an all-padding sequence does to masked pooling.
18 tickets = tickets[tickets['text'].str.contains(r'[A-Za-z0-9]',
19                                             regex=True)].reset_index(drop=True)
20
21 X_tr, X_te, y_tr, y_te = train_test_split(
22     tickets['text'], tickets['resolved'], test_size=0.25,
23     stratify=tickets['resolved'], random_state=42)
24 vec = tf.keras.layers.TextVectorization(max_tokens=400,
25                                         output_sequence_length=24)
26 vec.adapt(X_tr.to_numpy())
27
28 tf.random.set_seed(42)
29 model = tf.keras.Sequential([
30     tf.keras.layers.Input(shape=(1,), dtype=tf.string),
31     vec,
32     tf.keras.layers.Embedding(vec.vocabulary_size(), 16, mask_zero=True),
33     tf.keras.layers.GlobalAveragePooling1D(),
34     tf.keras.layers.Dense(16, activation='relu'),
35     tf.keras.layers.Dense(1, activation='sigmoid'),

```

```

36 ])
37 model.compile(optimizer='adam', loss='binary_crossentropy',
38               metrics=['accuracy'])
39 model.fit(X_tr.to_numpy(), y_tr.to_numpy(), epochs=25, batch_size=64,
40           verbose=0)
41 print('accuracy %.4f'
42       % model.evaluate(X_te.to_numpy(), y_te.to_numpy(), verbose=0)[1])

```

```
accuracy 0.9188
```

Averaging the embeddings discards the order just as thoroughly as `CountVectorizer` did, two tickets with the same words in a different arrangement produce an identical average. What rescues the score is the `Dense` layer sitting after the average, which is nonlinear and can therefore represent the interaction, exactly as gradient boosting did on day 2. It lands in the same region as that model, and behind bigrams, for the same reason: it has to discover the pairing rather than being handed it.

The geometry it learned

```

1  import os
2  os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3  import numpy as np
4  import tensorflow as tf
5  tf.random.set_seed(42)
6  tf.get_logger().setLevel('ERROR')
7  import numpy as np
8  import pandas as pd
9
10 tickets = pd.read_csv('tickets.csv')
11 tickets['text'] = tickets['text'].fillna('').str.strip()
12 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
13 vec = tf.keras.layers.TextVectorization(max_tokens=400,
14                                         output_sequence_length=24)
15 vec.adapt(tickets['text'].to_numpy())
16
17 tf.random.set_seed(42)
18 model = tf.keras.Sequential([
19     tf.keras.layers.Input(shape=(1,), dtype=tf.string),
20     vec,
21     tf.keras.layers.Embedding(vec.vocabulary_size(), 16),
22     tf.keras.layers.GlobalAveragePooling1D(),
23     tf.keras.layers.Dense(16, activation='relu'),
24     tf.keras.layers.Dense(3, activation='softmax'),
25 ])
26 model.compile(optimizer='adam', loss='sparse_categorical_crossentropy')
27 codes = {'billing': 0, 'cancellation': 1, 'technical': 2}
28 model.fit(tickets['text'].to_numpy(),
29           tickets['category'].map(codes).to_numpy(),
30           epochs=30, batch_size=64, verbose=0)
31
32 E = model.layers[1].get_weights()[0]
33 words = vec.get_vocabulary()

```

```

34 idx = {w: i for i, w in enumerate(words)}
35 norm = E / (np.linalg.norm(E, axis=1, keepdims=True) + 1e-9)
36
37 def nearest(word, k=4):
38     if word not in idx:
39         return '(not in vocabulary)'
40     sims = norm @ norm[idx[word]]
41     order = [i for i in sims.argsort()[::-1] if i != idx[word]][:k]
42     return ', '.join('%s %.2f' % (words[i], sims[i]) for i in order)
43
44 for w in ['invoice', 'router', 'termination', 'engineer']:
45     print('%-13s → %s' % (w, nearest(w)))

```

```

invoice      → regards -0.86, have -0.81, conection -0.81, internt -0.72
router       → considering -0.84, six -0.84, years -0.83, options -0.83
termination  → unreliable -0.84, on -0.79, router -0.77, pay -0.76
engineer     → fine -0.91, fibre -0.74, account -0.74, to -0.73

```

The neighbours come from the task, not from English

These vectors were trained to predict the ticket category, so words land near each other when they push the category the same way. That is not the same as meaning similarity, and it is why the embedding you train on 3,000 tickets is far less useful than one trained on billions of words. The latter has seen enough context to learn something closer to actual usage. Day 6 returns to this.

Representation	Columns	Order?	Generalises to unseen words?
One-hot / bag of words	One per vocabulary word	No	No
TF-IDF	Same, reweighted	No	No
n-grams	Explodes with n	Locally	No
Learned embedding	Chosen, e.g. 64	Not by itself	Only via subwords
Pretrained embedding	Chosen, e.g. 300	Not by itself	Yes, if the word was in the source corpus

Day 3 takeaway

Day 4 Recurrent Networks

State across a sequence, vanishing gradients, LSTM, padding and masking

A recurrent layer walks the sequence one token at a time, carrying a state that it updates at every step. That state is what lets it know *not* came before *working*.

A recurrent cell, in numpy

```

1 import numpy as np
2
3 rng = np.random.default_rng(0)

```

```

4 d_in, d_hidden = 4, 3
5 Wx = rng.normal(0, 0.5, (d_in, d_hidden))
6 Wh = rng.normal(0, 0.5, (d_hidden, d_hidden))
7 b = np.zeros(d_hidden)
8
9 def run(sequence):
10     h = np.zeros(d_hidden)
11     for x in sequence:
12         h = np.tanh(x @ Wx + h @ Wh + b)    # the whole recurrence
13     return h
14
15 a = rng.normal(size=(1, d_in))
16 c = rng.normal(size=(1, d_in))
17 print('state after [a, b] %s' % run(np.vstack([a, c])).round(3))
18 print('state after [b, a] %s' % run(np.vstack([c, a])).round(3))
19 print('\nsame two inputs, different order, different state.')
20 print('that is the entire reason this layer exists.')

```

```

state after [a, b] [0.038 0.046 0.361]
state after [b, a] [0.172 0.144 0.398]

```

```

same two inputs, different order, different state.
that is the entire reason this layer exists.

```

Why plain recurrence fails on long sequences

```

1 import numpy as np
2
3 print('%8s %16s %16s' % ('steps', 'weight 0.5', 'weight 1.5'))
4 for steps in [5, 10, 25, 50, 100]:
5     print('%8d %16.6g %16.6g'
6           % (steps, 0.5 ** steps, 1.5 ** steps))
7 print('\nthe gradient is a product of one factor per step.')

```

steps	weight 0.5	weight 1.5
5	0.03125	7.59375
10	0.000976562	57.665
25	2.98023e-08	25251.2
50	8.88178e-16	6.37622e+08
100	7.88861e-31	4.06561e+17

```

the gradient is a product of one factor per step.

```

What LSTM and GRU actually changed

Both add a path along which the state is *added to* rather than repeatedly multiplied, plus learned gates that decide what to keep, what to forget and what to output. The additive path is the important part: a sum does not shrink geometrically, so gradients survive far more steps. GRU is the same idea with two gates instead of three, fewer parameters, usually indistinguishable results.

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'

```

```

3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 for name, layer in [('SimpleRNN', tf.keras.layers.SimpleRNN(32)),
8                     ('GRU', tf.keras.layers.GRU(32)),
9                     ('LSTM', tf.keras.layers.LSTM(32))]:
10     layer.build((None, 20, 16))
11     print('%-11s %6d parameters'
12           % (name, sum(int(np.prod(w.shape)) for w in layer.weights)))
13 print('\nfor the same 16-wide input and 32-wide state.')
14 print('LSTM has four internal transforms, GRU three, SimpleRNN one.')

```

```

SimpleRNN      1568 parameters
GRU            4800 parameters
LSTM           6272 parameters

```

```

for the same 16-wide input and 32-wide state.
LSTM has four internal transforms, GRU three, SimpleRNN one.

```

Padding, masking and why they matter

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 import pandas as pd
9
10 tickets = pd.read_csv('tickets.csv')
11 tickets['text'] = tickets['text'].fillna('').str.strip()
12 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
13 from sklearn.model_selection import train_test_split
14
15 # Drop tickets with no word characters at all -- the empty ones, and the
16 # ones that are only punctuation. Both tokenise to nothing, and day 4
17 # shows what an all-padding sequence does to masked pooling.
18 tickets = tickets[tickets['text'].str.contains(r'[A-Za-z0-9]',
19                                                regex=True)].reset_index(drop=True)
20
21 X_tr, X_te, y_tr, y_te = train_test_split(
22     tickets['text'], tickets['resolved'], test_size=0.25,
23     stratify=tickets['resolved'], random_state=42)
24 vec = tf.keras.layers.TextVectorization(max_tokens=400,
25                                         output_sequence_length=20)
26 vec.adapt(X_tr.to_numpy())
27
28 sample = tf.constant(['the hub is broken', X_tr.iloc[0]])
29 encoded = vec(sample).numpy()
30 print('short ticket encoded:')
31 print(encoded[0])

```

```

32 print('%d real tokens, %d padding zeros'
33       % ((encoded[0] > 0).sum(), (encoded[0] == 0).sum()))
34
35 emb = tf.keras.layers.Embedding(vec.vocabulary_size(), 8, mask_zero=True)
36 out = emb(encoded)
37 mask = emb.compute_mask(encoded)
38 print('\nmask for the short ticket:')
39 print(mask.numpy()[0])
40 print('False means "ignore this step" for every layer downstream.')

```

```

short ticket encoded:
[ 2 24  3 59  0  0  0  0  0  0  0  0  0  0  0  0  0  0]
4 real tokens, 16 padding zeros

mask for the short ticket:
[ True  True  True  True False False False False False False False False
 False False False False False False False False]
False means "ignore this step" for every layer downstream.

```

⚠ Watch Out

mask_zero=True is not the default

Without it, a recurrent layer processes twelve padding tokens after your eight real ones and the final state is whatever those zeros left behind. Average pooling is worse: it divides by the padded length, so a four-word ticket has its signal diluted by sixteen zeros. Short documents suffer most, which is precisely the population you are usually worst at already.

The clearest way to see what masking buys you is to change nothing but the padded length. The words are the same, so the pooled vector representing them should be the same.

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 import pandas as pd
9
10 tickets = pd.read_csv('tickets.csv')
11 tickets['text'] = tickets['text'].fillna('').str.strip()
12 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
13 from sklearn.model_selection import train_test_split
14
15 # Drop tickets with no word characters at all -- the empty ones, and the
16 # ones that are only punctuation. Both tokenise to nothing, and day 4
17 # shows what an all-padding sequence does to masked pooling.
18 tickets = tickets[tickets['text'].str.contains(r'[A-Za-z0-9]',
19                                                regex=True)].reset_index(drop=True)
20
21 X_tr, X_te, y_tr, y_te = train_test_split(
22     tickets['text'], tickets['resolved'], test_size=0.25,

```

```

23 stratify=tickets['resolved'], random_state=42)
24 text = tf.constant(['the hub is broken'])
25
26 # Seed the initialiser itself. tf.random.set_seed does not control the
27 # seed generator Keras uses for layer weights, so without this the two
28 # lengths get different embeddings and the comparison measures nothing.
29 init = tf.keras.initializers.RandomUniform(-0.05, 0.05, seed=0)
30
31 for mask_zero in [False, True]:
32     pooled = {}
33     for length in [8, 40]:
34         vec = tf.keras.layers.TextVectorization(
35             max_tokens=400, output_sequence_length=length)
36         vec.adapt(X_tr.to_numpy())
37         e = vec(text)
38         emb = tf.keras.layers.Embedding(vec.vocabulary_size(), 8,
39                                         embeddings_initializer=init,
40                                         mask_zero=mask_zero)
41         pooled[length] = tf.keras.layers.GlobalAveragePooling1D()(
42             emb(e), mask=emb.compute_mask(e) if mask_zero else None)
43         same = float(tf.reduce_max(tf.abs(pooled[8] - pooled[40])))
44         print('mask_zero=%-6s padded to 8 vs to 40: largest difference %.4f'
45               % (mask_zero, same))

```

```

mask_zero=False padded to 8 vs to 40: largest difference 0.0163
mask_zero=True padded to 8 vs to 40: largest difference 0.0000

```

With masking, the padded length makes no difference at all: the layer averages four vectors either way. Without it, the answer changes when you change a number that is supposed to be an implementation detail. Note what is *not* happening in the unmasked case: index 0 has its own embedding vector like any other token, so the padding does not dilute the average towards zero, it drags it towards whatever the padding vector happens to be. On a corpus of mostly-short documents, that vector ends up carrying a large share of your representation.

The bug masking gives you for free

Masking has one failure mode, and this corpus contains it. A document with no tokens encodes to nothing but padding, so the mask is `False` at every position and average pooling divides a sum of nothing by a count of nothing.

```

1 import os
2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 vec = tf.keras.layers.TextVectorization(max_tokens=50,
8                                         output_sequence_length=8)
9 vec.adapt(tf.constant(['the router is broken', 'the invoice is wrong']))
10
11 batch = vec(tf.constant(['the router is broken', '', '...']))
12 emb = tf.keras.layers.Embedding(vec.vocabulary_size(), 4, mask_zero=True)
13 e = emb(batch)

```

```

14 mask = emb.compute_mask(batch)
15 pooled = tf.keras.layers.GlobalAveragePooling1D()(e, mask=mask)
16
17 print('encoded:')
18 for text, row in zip(['the router is broken', '(empty)', '...'],
19                       batch.numpy()):
20     print('  %-22s %s' % (text, row))
21 print('\npooled vectors:')
22 print(pooled.numpy().round(3))
23 print('\nany NaN?', bool(np.isnan(pooled.numpy()).any()))

```

```

encoded:
  the router is broken  [2 5 3 7 0 0 0 0]
  (empty)              [0 0 0 0 0 0 0 0]
  ...                  [0 0 0 0 0 0 0 0]

pooled vectors:
[[ 0.006 -0.      0.015 -0.007]
 [  nan   nan   nan   nan]
 [  nan   nan   nan   nan]]

any NaN? True

```

⚠ Watch Out

Non-empty is not the same as non-empty

Row three is the important one. '...' is a perfectly good string. It passes any `len(text) > 0` check you write, but `TextVectorization` strips punctuation by default, so it tokenises to nothing and produces the same NaN as the genuinely empty row. The guard has to be on the token count, not on the string length.

And the NaN does not raise. It propagates through the dense layers, turns the loss into nan, and makes every gradient after it nan too, so one such document destroys an entire training run and the error, if it ever surfaces, points somewhere else. This is why the split used from day 2 onwards filters on `[A-Za-z0-9]` rather than on length.

Day 4 takeaway

Day 5 The Honest Comparison

LSTM against bigrams, measured on accuracy, time and everything else

Everything is now on the table: a bag of words, bigrams, averaged embeddings and a recurrent layer that genuinely reads the sequence. Time to run them against each other on the task that needs order, and to be honest about the result.

The contenders

```

1 import os

```

```

2 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
3 import numpy as np
4 import tensorflow as tf
5 tf.random.set_seed(42)
6 tf.get_logger().setLevel('ERROR')
7 import numpy as np
8 import pandas as pd
9
10 tickets = pd.read_csv('tickets.csv')
11 tickets['text'] = tickets['text'].fillna('').str.strip()
12 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
13 from sklearn.model_selection import train_test_split
14
15 # Drop tickets with no word characters at all -- the empty ones, and the
16 # ones that are only punctuation. Both tokenise to nothing, and day 4
17 # shows what an all-padding sequence does to masked pooling.
18 tickets = tickets[tickets['text'].str.contains(r'[A-Za-z0-9]',
19                                             regex=True)].reset_index(drop=True)
20
21 X_tr, X_te, y_tr, y_te = train_test_split(
22     tickets['text'], tickets['resolved'], test_size=0.25,
23     stratify=tickets['resolved'], random_state=42)
24 import time
25 from sklearn.metrics import roc_auc_score
26
27 def build(kind, vec):
28     layers = [tf.keras.layers.Input(shape=(1,), dtype=tf.string), vec,
29              tf.keras.layers.Embedding(vec.vocabulary_size(), 24,
30                                         mask_zero=True)]
31     if kind == 'average':
32         layers.append(tf.keras.layers.GlobalAveragePooling1D())
33     elif kind == 'lstm':
34         layers.append(tf.keras.layers.LSTM(32))
35     elif kind == 'bilstm':
36         layers.append(tf.keras.layers.Bidirectional(
37             tf.keras.layers.LSTM(32)))
38     elif kind == 'gru':
39         layers.append(tf.keras.layers.GRU(32))
40     layers += [tf.keras.layers.Dense(16, activation='relu'),
41               tf.keras.layers.Dense(1, activation='sigmoid')]
42     return tf.keras.Sequential(layers)
43
44 xs, ys = X_tr.to_numpy(), y_tr.to_numpy()
45 xv, yv = X_te.to_numpy(), y_te.to_numpy()
46
47 print('%-12s %10s %10s %10s' % ('model', 'accuracy', 'auc', 'fit time'))
48 for kind in ['average', 'gru', 'lstm', 'bilstm']:
49     tf.random.set_seed(42)
50     vec = tf.keras.layers.TextVectorization(max_tokens=400,
51                                             output_sequence_length=24)
52     vec.adapt(xs)
53     m = build(kind, vec)
54     m.compile(optimizer='adam', loss='binary_crossentropy',
55              metrics=['accuracy'])

```

```

56     t = time.time()
57     m.fit(xs, ys, epochs=25, batch_size=64, verbose=0)
58     took = time.time() - t
59     proba = m.predict(xv, verbose=0).ravel()
60     print('%-12s %10.4f %10.4f %9.1fs'
61           % (kind, ((proba > 0.5).astype(int) == yv).mean(),
62               roc_auc_score(yv, proba), took))

```

model	accuracy	auc	fit time
average	0.9179	0.9049	3.0s
gru	0.9608	0.9610	10.1s
lstm	0.9570	0.9601	10.4s
bilstm	0.9561	0.9614	16.0s

Against the model from day 2

```

1  import numpy as np
2  import pandas as pd
3
4  tickets = pd.read_csv('tickets.csv')
5  tickets['text'] = tickets['text'].fillna('').str.strip()
6  tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7  from sklearn.model_selection import train_test_split
8
9  # Drop tickets with no word characters at all -- the empty ones, and the
10 # ones that are only punctuation. Both tokenise to nothing, and day 4
11 # shows what an all-padding sequence does to masked pooling.
12 tickets = tickets[tickets['text'].str.contains(r'[A-Za-z0-9]',
13                                                regex=True)].reset_index(drop=True)
14
15 X_tr, X_te, y_tr, y_te = train_test_split(
16     tickets['text'], tickets['resolved'], test_size=0.25,
17     stratify=tickets['resolved'], random_state=42)
18 import time
19 from sklearn.feature_extraction.text import TfidfVectorizer
20 from sklearn.linear_model import LogisticRegression
21 from sklearn.pipeline import make_pipeline
22 from sklearn.metrics import roc_auc_score
23
24 t = time.time()
25 m = make_pipeline(TfidfVectorizer(ngram_range=(1, 2)),
26                  LogisticRegression(max_iter=2000)).fit(X_tr, y_tr)
27 took = time.time() - t
28 proba = m.predict_proba(X_te)[:, 1]
29 print('tf-idf bigrams + logistic regression')
30 print(' accuracy %4f' % m.score(X_te, y_te))
31 print(' auc      %4f' % roc_auc_score(y_te, proba))
32 print(' fit time %2fs' % took)
33 print(' parameters: %d coefficients' % m[1].coef_.size)

```

```

tf-idf bigrams + logistic regression
accuracy 0.9542

```

```
auc      0.9591
fit time 0.07s
parameters: 1079 coefficients
```

Read this before you reach for an LSTM

Both approaches solve the task. One of them took a fraction of a second, fits in a line, produces coefficients you can read out loud, and has no seed, no epoch count and no learning rate to get wrong. The other needed an embedding width, a state width, a sequence length, a batch size and twenty-five epochs, and its parameters mean nothing to anybody.

On four thousand short tickets that is not a close call. Sequence models start earning their keep when documents are long, when the dependency spans more than two or three words, or when you are starting from a model pretrained on far more text than you have, which is the subject of day 6.

Where each one actually fails

```
1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 from sklearn.model_selection import train_test_split
8
9 # Drop tickets with no word characters at all -- the empty ones, and the
10 # ones that are only punctuation. Both tokenise to nothing, and day 4
11 # shows what an all-padding sequence does to masked pooling.
12 tickets = tickets[tickets['text'].str.contains(r'[A-Za-z0-9]',
13                                             regex=True)].reset_index(drop=True)
14
15 X_tr, X_te, y_tr, y_te = train_test_split(
16     tickets['text'], tickets['resolved'], test_size=0.25,
17     stratify=tickets['resolved'], random_state=42)
18 from sklearn.feature_extraction.text import TfidfVectorizer
19 from sklearn.linear_model import LogisticRegression
20 from sklearn.pipeline import make_pipeline
21
22 m = make_pipeline(TfidfVectorizer(ngram_range=(1, 2)),
23                 LogisticRegression(max_iter=2000)).fit(X_tr, y_tr)
24 proba = m.predict_proba(X_te)[: , 1]
25 pred = (proba > 0.5).astype(int)
26 wrong = np.where(pred != y_te.to_numpy())[0]
27
28 print('%d wrong out of %d' % (len(wrong), len(y_te)))
29 print('\nthe six it was most confident about, and wrong:')
30 conf = sorted(wrong, key=lambda i: abs(proba[i] - 0.5), reverse=True)[:6]
31 for i in conf:
32     print(' truth=%d p=%.2f %s'
33           % (y_te.iloc[i], proba[i], X_te.iloc[i][:64]))
```

```
48 wrong out of 1047
```

```
the six it was most confident about, and wrong:
```

```
truth=1 p=0.11 the engineer visited on tuesday. my connection is still patchy.
truth=1 p=0.12 i am considering my options. the router is not reliable. my sign
truth=1 p=0.14 hello, THE ENGINEER VISITED ON TUESDAY. IT IS NOT RELIABLE. THE
truth=1 p=0.15 My connection is still unreliable. there is a charge on the stat
truth=1 p=0.15 as of this morning the connection is still slow. the connection
truth=1 p=0.15 my hub is broken. the broadband is still unstable. this relates
```

Some of those are the four percent of labels the generator deliberately flipped, and no model can get them right, which is the point of putting them there. Reading your confident errors is the fastest way to tell an irreducible label problem from a fixable modelling one.

Day 5 takeaway

Day 6 Attention and the Modern Landscape

Self-attention from scratch, subword tokens, and what actually wins

Recurrence has a structural problem that no amount of gating fixes: it is sequential. Token 500 cannot be processed until token 499 has been, which means a recurrent model cannot use a GPU's parallelism over the length of the sequence. Attention removes that constraint, and removing it is what made models trained on the whole internet possible.

Attention, in about fifteen lines

```
1 import numpy as np
2
3 def softmax(z):
4     z = z - z.max(axis=-1, keepdims=True)
5     e = np.exp(z)
6     return e / e.sum(axis=-1, keepdims=True)
7
8 def attention(Q, K, V):
9     scores = Q @ K.T / np.sqrt(K.shape[-1])    # who is relevant to whom
10    weights = softmax(scores)                    # rows sum to 1
11    return weights @ V, weights                  # weighted average of values
12
13 rng = np.random.default_rng(0)
14 tokens = ['the', 'router', 'is', 'not', 'working']
15 X = rng.normal(size=(len(tokens), 8))
16
17 # with random projections the pattern is random; the point here is shape
18 Wq, Wk, Wv = (rng.normal(0, 0.5, (8, 8)) for _ in range(3))
19 out, weights = attention(X @ Wq, X @ Wk, X @ Wv)
20
21 print('output shape %s, one vector per token' % (out.shape,))
22 print('\nattention weights (row = token doing the looking):')
23 print('%-9s %s' % ('', ' '.join('%7s' % t for t in tokens)))
24 for t, row in zip(tokens, weights):
```

```

25 print('%-9s %s' % (t, ' '.join('%7.3f' % v for v in row)))
26 print('\nevery row sums to 1: %s' % np.allclose(weights.sum(axis=1), 1))

```

output shape (5, 8), one vector per token

attention weights (row = token doing the looking):

	the	router	is	not	working
the	0.122	0.088	0.434	0.248	0.108
router	0.614	0.012	0.027	0.010	0.338
is	0.302	0.145	0.151	0.188	0.214
not	0.225	0.159	0.237	0.103	0.276
working	0.446	0.012	0.224	0.044	0.274

every row sums to 1: True

⚠ Watch Out

Those weights are meaningless, and that is deliberate

The projections here are random, so the pattern shows nothing about language, only the mechanism and the shapes. In a trained model the row for *not* would put weight on *working*, which is how attention captures exactly the dependency that defeated the unigram model on day 2. Published attention heat-maps are trained weights, and reading them as explanations is a well-documented trap: attention shows what the model looked at, not why it decided.

Why the position has to be added back

```

1 import numpy as np
2
3 def softmax(z):
4     z = z - z.max(axis=-1, keepdims=True)
5     e = np.exp(z)
6     return e / e.sum(axis=-1, keepdims=True)
7
8 def attention(Q, K, V):
9     return softmax(Q @ K.T / np.sqrt(K.shape[-1])) @ V
10
11 rng = np.random.default_rng(0)
12 X = rng.normal(size=(4, 6))
13 Wq, Wk, Wv = (rng.normal(0, 0.5, (6, 6)) for _ in range(3))
14
15 a = attention(X @ Wq, X @ Wk, X @ Wv)
16 shuffled = X[[2, 0, 3, 1]]
17 b = attention(shuffled @ Wq, shuffled @ Wk, shuffled @ Wv)
18
19 print('shuffle the tokens and the outputs are the same rows, reordered:')
20 print(np.allclose(a[[2, 0, 3, 1]], b))
21 print('\nso attention alone is order-blind. positional encodings are')
22 print('added to the embeddings precisely to fix this.')

```

shuffle the tokens and the outputs are the same rows, reordered:

True

so attention alone is order-blind. positional encodings are added to the embeddings precisely to fix this.

The landscape, honestly

Approach	Needs	Use when
TF-IDF + linear	Minutes, a laptop	Short texts, a clear vocabulary signal. Always try first.
Embedding + pooling	A GPU is optional	Large corpus, many classes, vocabulary too large for n-grams
LSTM / GRU	Patience	Long-range order matters and you cannot use a pretrained model
Pretrained transformer, frozen	A download	Small labelled set; use the embeddings as features
Pretrained transformer, fine-tuned	A GPU, hours	A few thousand labels and accuracy that justifies the cost
A hosted large language model	An API budget	Few or no labels, or the task is genuinely open-ended

Where week 12's lesson repeats itself

The reason transformers dominate text is the same reason convolutional networks dominate images: the architecture matches the structure of the data, and there is enough data to fill it. A transformer trained from scratch on your four thousand tickets will lose to logistic regression, comfortably. A transformer pretrained on a trillion tokens and fine-tuned on your four thousand tickets will win. The architecture is not the thing that wins, the pretraining is.

Subword tokenisation, and why it exists

```

1 # Word-level vocabularies break on anything they have not seen.
2 vocab = {'the', 'router', 'is', 'not', 'working', 'slow'}
3 for sentence in ['the router is not working',
4                 'the routers are misconfigured']:
5     tokens = sentence.split()
6     marked = [t if t in vocab else '[UNK]' for t in tokens]
7     print('%-32s → %s' % (sentence, ' '.join(marked)))
8
9 print('\nsubword tokenisers split unknown words into known pieces:')
10 print('  routers          → router + s')
11 print('  misconfigured    → mis + config + ured')
12 print('so the vocabulary is finite and nothing is ever fully unknown.')
```

```

the router is not working      → the router is not working
the routers are misconfigured  → the [UNK] [UNK] [UNK]
```

```

subword tokenisers split unknown words into known pieces:
  routers          → router + s
  misconfigured    → mis + config + ured
so the vocabulary is finite and nothing is ever fully unknown.
```

This is why modern models report vocabularies of around 30,000 to 100,000 pieces rather than millions of words, and why they cope with typos, product codes and languages with heavy compounding. It is also why token counts on an API bill do not match word counts. You are charged per piece.

Day 6 takeaway

Day 7 A Complete Text Workflow

Grid search inside a pipeline, error analysis, and the text-specific traps

A complete text pipeline on the category task, plus the failure modes that are specific to text and that no amount of cross-validation will catch on its own.

The pipeline

```
1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 from sklearn.feature_extraction.text import TfidfVectorizer
8 from sklearn.linear_model import LogisticRegression
9 from sklearn.pipeline import Pipeline
10 from sklearn.model_selection import GridSearchCV, train_test_split
11
12 X_tr, X_te, y_tr, y_te = train_test_split(
13     tickets['text'], tickets['category'], test_size=0.25,
14     stratify=tickets['category'], random_state=42)
15
16 pipe = Pipeline([('tfidf', TfidfVectorizer()),
17                 ('clf', LogisticRegression(max_iter=2000))])
18 grid = {'tfidf__ngram_range': [(1, 1), (1, 2)],
19         'tfidf__min_df': [1, 3],
20         'tfidf__sublinear_tf': [False, True],
21         'clf__C': [0.5, 2.0, 8.0]}
22
23 search = GridSearchCV(pipe, grid, cv=5, scoring='f1_macro', n_jobs=1)
24 search.fit(X_tr, y_tr)
25 print('best macro F1 in cross-validation %.4f' % search.best_score_)
26 for k, v in sorted(search.best_params_.items()):
27     print(' %-24s %s' % (k, v))
28 print('\nheld-out accuracy %.4f' % search.score(X_te, y_te))
```

```
best macro F1 in cross-validation 0.8732
clf__C          0.5
tfidf__min_df   1
tfidf__ngram_range (1, 2)
tfidf__sublinear_tf True
```

held-out accuracy 0.8456

⚠ Watch Out

Vectorise inside the pipeline, never before the split

`TfidfVectorizer` learns a vocabulary and a document frequency for every term. Fit it on all your text and those statistics have seen the test set, which is week 4's leakage wearing a different hat. Inside a `Pipeline`, every cross-validation fold refits the vectoriser on that fold's training rows only. This is the single most common leak in text projects, and it flatters your score by a few points, just enough to be believable.

Reading the errors

```

1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 from sklearn.feature_extraction.text import TfidfVectorizer
8 from sklearn.linear_model import LogisticRegression
9 from sklearn.pipeline import make_pipeline
10 from sklearn.model_selection import train_test_split
11 from sklearn.metrics import classification_report, confusion_matrix
12
13 X_tr, X_te, y_tr, y_te = train_test_split(
14     tickets['text'], tickets['category'], test_size=0.25,
15     stratify=tickets['category'], random_state=42)
16 m = make_pipeline(TfidfVectorizer(ngram_range=(1, 2)),
17                  LogisticRegression(max_iter=2000, C=2.0)).fit(X_tr, y_tr)
18 pred = m.predict(X_te)
19 labels = sorted(y_te.unique())
20
21 print(classification_report(y_te, pred, digits=3))
22 print('confusion matrix (rows = truth):')
23 print('%-14s %s' % ('', ' '.join('%13s' % l for l in labels)))
24 for l, row in zip(labels, confusion_matrix(y_te, pred, labels=labels)):
25     print('%-14s %s' % (l, ' '.join('%13d' % v for v in row)))

```

	precision	recall	f1-score	support
billing	0.842	0.804	0.822	357
cancellation	0.950	0.732	0.827	157
technical	0.840	0.922	0.879	536
accuracy			0.853	1050
macro avg	0.877	0.819	0.843	1050
weighted avg	0.857	0.853	0.852	1050

confusion matrix (rows = truth):

	billing	cancellation	technical
billing	287	3	67
cancellation	15	115	27
technical	39	3	494

The pitfalls that are specific to text

1. **Duplicate documents across the split.** This corpus contains 90 double-submitted tickets on purpose. Split without deduplicating and the same text is in training and test, so you are measuring memory. Deduplicate on the text, not just the id.
2. **Fitting the vectoriser before the split.** Covered above, and worth repeating because it is so easy to do accidentally in a notebook.
3. **Time.** Ticket language drifts, new products, new outages, new slang. A random split tells you how you do on the past; a split by date tells you how you will do next month, and it is always the lower number.
4. **Very short and empty documents.** Nine tickets here are empty strings. They produce an all-zero row, which every linear model will classify identically. Decide deliberately whether to drop them or route them to a human.
5. **Class imbalance.** Cancellation is the smallest and usually the most valuable category. Use macro-averaged F1 rather than accuracy, as the grid above does.
6. **Personal data.** Free text contains names, addresses, account and card numbers. Redact before you store features, and certainly before anything leaves your infrastructure for an API.

```

1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 from sklearn.feature_extraction.text import TfidfVectorizer
8 from sklearn.linear_model import LogisticRegression
9 from sklearn.neighbors import KNeighborsClassifier
10 from sklearn.pipeline import make_pipeline
11 from sklearn.model_selection import train_test_split
12
13 raw = pd.read_csv('tickets.csv')          # with the duplicates left in
14 raw['text'] = raw['text'].fillna('').str.strip()
15
16 A, B, ya, yb = train_test_split(raw['text'], raw['resolved'],
17                                test_size=0.25, random_state=0)
18
19 # Split the test set by whether its text was also in training. Comparing
20 # those two subsets isolates memorisation; comparing a deduplicated run
21 # against this one would not, because deduplicating changes the task.
22 seen = B.isin(set(A))
23 print('test rows whose text also appears in training: %d of %d\n'
24       % (seen.sum(), len(B)))
25
26 print('%-22s %8s %8s %8s' % ('', 'seen', 'unseen', 'overall'))

```

```

27 for name, clf in [('1-nearest neighbour', KNeighborsClassifier(1)),
28                  ('logistic regression',
29                   LogisticRegression(max_iter=2000))]:
30     m = make_pipeline(TfidfVectorizer(ngram_range=(1, 2)), clf).fit(A, ya)
31     print('%-22s %8.4f %8.4f %8.4f'
32           % (name, m.score(B[seen], yb[seen]),
33              m.score(B[~seen], yb[~seen]), m.score(B, yb)))

```

```
test rows whose text also appears in training: 38 of 1073
```

	seen	unseen	overall
1-nearest neighbour	0.9474	0.6580	0.6682
logistic regression	0.8947	0.9449	0.9432

The gap is the leak, and how big it is depends entirely on how much the model can memorise. One-nearest-neighbour finds the identical training ticket and copies its label, so it looks strong on the duplicated rows and collapses on everything else. The regularised linear model, which cannot store individual documents, shows no gap at all.

⚠ Watch Out

Do not conclude that duplicates are harmless

They are harmless *to this model*, on a corpus where 38 rows out of a thousand overlap. Scale either of those up, a scraped corpus where the same article appears on forty sites, or any model with the capacity to memorise, which includes every fine-tuned transformer, and the same leak moves your headline number by a lot. Deduplicate anyway: the check is two lines and the failure is invisible.

Saving the whole thing

```

1 import numpy as np
2 import pandas as pd
3
4 tickets = pd.read_csv('tickets.csv')
5 tickets['text'] = tickets['text'].fillna('').str.strip()
6 tickets = tickets.drop_duplicates(subset='ticket_id').reset_index(drop=True)
7 import joblib
8 from sklearn.feature_extraction.text import TfidfVectorizer
9 from sklearn.linear_model import LogisticRegression
10 from sklearn.pipeline import make_pipeline
11
12 model = make_pipeline(TfidfVectorizer(ngram_range=(1, 2)),
13                      LogisticRegression(max_iter=2000, C=2.0))
14 model.fit(tickets['text'], tickets['category'])
15 joblib.dump(model, 'ticket_router.joblib')
16
17 loaded = joblib.load('ticket_router.joblib')
18 new = ['the invoice is wrong again and nobody has called back',
19        'my router is no longer dropping out, thanks',
20        'i want to leave, what is the termination fee']
21 for text, label in zip(new, loaded.predict(new)):
22     proba = loaded.predict_proba([text]).max()

```

```
print('%-13s p=%.2f %s' % (label, proba, text[:52]))
```

```
billing      p=0.72  the invoice is wrong again and nobody has called bac
technical    p=0.47  my router is no longer dropping out, thanks
cancellation p=0.80  i want to leave, what is the termination fee
```

One object holds the tokeniser, the vocabulary, the IDF weights and the coefficients, and it takes raw strings as input. That is the property that matters in production: the serving code should never have to reimplement your preprocessing, because the day it drifts from the training version is the day your model silently degrades. Week 15 builds the service around exactly this file.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. A bag of words discards:
 - a) Word identity
 - b) The order the words appeared in
 - c) Rare words
 - d) Punctuation only
2. TF-IDF with unigrams scored *below* a constant prediction on the resolution task because:
 - a) The labels were random
 - b) The corpus was too small
 - c) The label lives in an interaction between tokens, and a linear model can only add one weight per token
 - d) TF-IDF is unsuitable for text
3. Gradient boosting on those same unigram counts scored far higher. This shows that:
 - a) TF-IDF was misconfigured
 - b) The split was lucky
 - c) Boosting is always better
 - d) The information was in the features all along; it was the linear model that could not reach it
4. Averaging a document's word embeddings:
 - a) Discards order, just as a bag of words does
 - b) Is the same as TF-IDF
 - c) Requires an LSTM
 - d) Preserves word order
5. Plain recurrent networks struggle with long sequences because the gradient is:
 - a) Too large always

- b) A product of one factor per step, so it vanishes or explodes
 - c) Computed only at the end
 - d) Independent of length
6. Forgetting `mask_zero=True` on an embedding means:
- a) Sequences are truncated
 - b) Training is slower
 - c) Padding is treated as real input, and a document with no tokens can produce a NaN that silently destroys the run
 - d) The vocabulary is wrong
7. Self-attention on its own is order-blind, so transformers add:
- a) A convolution
 - b) Larger batches
 - c) A recurrent layer
 - d) Positional encodings to the embeddings
8. Fitting `TfidfVectorizer` on all your text before splitting:
- a) Leaks, because the vocabulary and document frequencies are learned from the test rows too
 - b) Is faster and safe
 - c) Only matters for large corpora
 - d) Is good practice

Answers: 1b, 2c, 3d, 4a, 5b, 6c, 7d, 8a. Anything you got wrong points at the day to re-read, not at a gap in ability.

Hyperparameter Optimisation and Interpretability

Week 14

Spend a tuning budget properly, then explain what the tuned model is doing and where explanation stops.

OBJECTIVES

Two halves of one job. First, how to spend a fixed budget searching for settings, which parameters are worth searching at all, why a grid wastes most of its effort, and how to stop bad candidates early. Then, how to explain the result: the shape of each effect, contributions to a single prediction that add up exactly, what a customer would have to change, and the line between describing a model and making a claim about the world.

Day 1 What to Tune, and How to Search

Which hyperparameters matter, and why random search beats a grid

Every model so far has been fitted with whatever settings scikit-learn ships. Sometimes that is close to optimal and sometimes it is not, and the difference between a productive week of tuning and a wasted one is almost entirely about *which* settings you touch.

Not all hyperparameters matter

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],

```

```

20     random_state=42)
21
22     # Impute after the split, with medians learned from the training rows
23     # only -- week 4's rule still applies. A Pipeline is the right way to do
24     # this; these pages need a plain frame to inspect column by column, so
25     # the two lines are written out instead.
26     medians = X_tr.median()
27     X_tr = X_tr.fillna(medians)
28     X_te = X_te.fillna(medians)
29     from sklearn.ensemble import HistGradientBoostingClassifier
30     from sklearn.model_selection import cross_val_score
31
32     def score(**kw):
33         m = HistGradientBoostingClassifier(random_state=42, **kw)
34         return cross_val_score(m, X_tr, y_tr, cv=5, scoring='roc_auc').mean()
35
36     base = score()
37     print('defaults %.4f\n' % base)
38     print('%-22s %-18s %9s %9s' % ('parameter', 'values', 'best', 'spread'))
39     sweeps = {
40         'learning_rate': [0.01, 0.05, 0.1, 0.3],
41         'max_leaf_nodes': [7, 15, 31, 63],
42         'min_samples_leaf': [5, 20, 50, 100],
43         'l2_regularization': [0.0, 0.1, 1.0, 10.0],
44         'max_features': [0.4, 0.7, 1.0],
45     }
46     for name, values in sweeps.items():
47         got = [score(**{name: v}) for v in values]
48         print('%-22s %-18s %9.4f %9.4f'
49             % (name, str(values)[:18], max(got), max(got) - min(got)))

```

```
defaults 0.7945
```

parameter	values	best	spread
learning_rate	[0.01, 0.05, 0.1,	0.8012	0.0207
max_leaf_nodes	[7, 15, 31, 63]	0.8120	0.0238
min_samples_leaf	[5, 20, 50, 100]	0.8069	0.0144
l2_regularization	[0.0, 0.1, 1.0, 10	0.8032	0.0117
max_features	[0.4, 0.7, 1.0]	0.7992	0.0063

Read the spread column. One or two parameters move the score by something worth having and the rest barely register, and that pattern, not the specific winner, is what generalises. It is why the next two pages are about search strategy rather than about parameter values.

Model	Tune first	Tune if you have budget	Leave alone
Gradient boosting	learning_rate, max_leaf_nodes, n_estimators	min_samples_leaf, l2_regularization, sub-sampling	Almost everything else
Random forest	max_features, min_samples_leaf	max_depth	n_estimators: set it as high as you can afford
Linear / logistic	C or alpha	penalty type, class_weight	The solver, usually
SVM (RBF)	C, gamma	kernel choice	-
Neural network	learning rate	width, depth, dropout, batch size	The optimiser, use Adam

n_estimators is not a hyperparameter in a forest

More trees in a random forest never makes it worse, it only makes it slower. The variance of the average falls and then flattens. So there is nothing to search: pick the largest number you can afford and move on. In *boosting* it is entirely different, because each tree corrects the last and too many will overfit. Same argument name, opposite behaviour.

What a grid actually costs

```

1 params = {'learning_rate': [0.01, 0.05, 0.1, 0.3],
2           'max_leaf_nodes': [7, 15, 31, 63],
3           'min_samples_leaf': [5, 20, 50],
4           'l2_regularization': [0.0, 0.1, 1.0]}
5
6 total = 1
7 for name, values in params.items():
8     total *= len(values)
9     print('%-22s %d values, running total %d' % (name, len(values), total))
10
11 for folds in [5, 10]:
12     fits = total * folds
13     print('\n%d-fold: %d fits' % (folds, fits))
14     for secs in [0.2, 2.0]:
15         print('    at %.1fs per fit: %.1f minutes' % (secs, fits * secs / 60))

```

```

learning_rate      4 values, running total 4
max_leaf_nodes     4 values, running total 16
min_samples_leaf   3 values, running total 48
l2_regularization  3 values, running total 144

```

```

5-fold: 720 fits
    at 0.2s per fit: 2.4 minutes
    at 2.0s per fit: 24.0 minutes

```

```

10-fold: 1440 fits
    at 0.2s per fit: 4.8 minutes
    at 2.0s per fit: 48.0 minutes

```

Adding one more value to one more parameter multiplies the cost. A grid is the only strategy whose price grows exponentially in the number of things you are curious about, which is a strange

property for the default choice to have.

Random search, and why it wins

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 import time
30 from scipy.stats import loguniform, randint
31 from sklearn.ensemble import HistGradientBoostingClassifier
32 from sklearn.model_selection import GridSearchCV, RandomizedSearchCV
33
34 grid = {'learning_rate': [0.01, 0.05, 0.1, 0.3],
35         'max_leaf_nodes': [7, 15, 31, 63],
36         'min_samples_leaf': [5, 20, 50, 100]}
37 dists = {'learning_rate': loguniform(0.005, 0.4),
38         'max_leaf_nodes': randint(4, 80),
39         'min_samples_leaf': randint(3, 120)}
40
41 est = HistGradientBoostingClassifier(random_state=42)
42 for label, search in [
43     ('grid, 64 points', GridSearchCV(est, grid, cv=5,
44                                     scoring='roc_auc')),
45     ('random, 25 draws', RandomizedSearchCV(est, dists, n_iter=25, cv=5,
46                                             scoring='roc_auc',
47                                             random_state=0))]:
48     t = time.time()
49     search.fit(X_tr, y_tr)

```

```

50 print('%-18s best %.4f in %5.1fs (%d fits)'
51       % (label, search.best_score_, time.time() - t,
52         len(search.cv_results_['params']) * 5))

```

```

grid, 64 points    best 0.8155 in 71.8s (320 fits)
random, 25 draws  best 0.8160 in 28.4s (125 fits)

```

```

1 import numpy as np
2
3 # Two parameters, only one of which matters. A 5 x 5 grid tries five
4 # distinct values of each. Twenty-five random draws try twenty-five.
5 rng = np.random.default_rng(0)
6 print('%-16s %22s %22s' % ('', 'distinct values tried', 'of the one that matters'))
7 for n in [9, 16, 25, 36]:
8     side = int(round(n ** 0.5))
9     print('%-16s %22d %22d'
10          % ('grid %dx%d' % (side, side), n, side))
11     print('%-16s %22d %22d' % ('random %d' % n, n, n))

```

	distinct values tried of the one that matters	
grid 3x3	9	3
random 9	9	9
grid 4x4	16	4
random 16	16	16
grid 5x5	25	5
random 25	25	25
grid 6x6	36	6
random 36	36	36

The argument in one sentence

A grid spends its budget trying the same few values of the parameter that matters, over and over, paired with values of parameters that do not. Random search spends the same budget on distinct values of everything. Since you rarely know in advance which parameters matter, random search is the better default, and it lets you stop at any point rather than at a multiple of the grid size.

Day 1 takeaway

Day 2 Spending the Budget Well

Successive halving, early stopping, and Bayesian optimisation from scratch

Random search still gives every candidate the full budget, including the hopeless ones. Two better ideas: stop bad candidates early, and let the model itself decide when it has had enough.

Successive halving

```

1 import numpy as np
2 import pandas as pd

```

```

3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16         'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 import time
30 from scipy.stats import loguniform, randint
31 from sklearn.experimental import enable_halving_search_cv # still required
32 from sklearn.model_selection import HalvingRandomSearchCV
33 from sklearn.ensemble import HistGradientBoostingClassifier
34
35 dists = {'learning_rate': loguniform(0.005, 0.4),
36         'max_leaf_nodes': randint(4, 80),
37         'min_samples_leaf': randint(3, 120)}
38
39 search = HalvingRandomSearchCV(
40     HistGradientBoostingClassifier(random_state=42), dists,
41     n_candidates=60, factor=3, resource='n_samples',
42     # Without this the first round scores on 20 rows, where a 5-fold
43     # split can hand a fold a single class and roc_auc returns nan.
44     min_resources=250, cv=5,
45     scoring='roc_auc', random_state=0)
46 t = time.time()
47 search.fit(X_tr, y_tr)
48
49 res = pd.DataFrame(search.cv_results_)
50 print('%-8s %12s %12s %12s' % ('round', 'candidates', 'rows each', 'best'))
51 for i, g in res.groupby('iter'):
52     print('%-8d %12d %12d %12.4f'
53         % (i, len(g), g['n_resources'].iloc[0], g['mean_test_score'].max()))
54 print('\n60 candidates explored in %.1fs, best %.4f'
55     % (time.time() - t, search.best_score_))

```

round	candidates	rows each	best
0	60	250	0.7730
1	20	750	0.8011
2	7	2250	0.8160

60 candidates explored in 43.5s, best 0.8160

Sixty candidates screened where day 1's random search managed twenty-five, and the score climbs each round as the survivors get more data, 250 rows, then 750, then all 2,250. It arrives at 0.8160, the same figure random search reached on day 1.

⚠ Watch Out

The same answer, for more wall-clock. Measure before you believe

Halving took longer here than the plain random search it matched. That is not a criticism of the algorithm, it is a statement about the dataset: fitting this model on 2,250 rows takes a fraction of a second, so there is nothing to save by fitting it on 250 instead, and the extra rounds of bookkeeping are pure overhead.

Successive halving is built for the opposite situation, hundreds of thousands of rows, or a model that takes minutes per fit, where eliminating forty hopeless candidates cheaply is worth a great deal. Note also `min_resources`: left at its default this search began on 20 rows, where a five-fold split can hand a fold a single class and the AUC comes back as `nan`. If your first round is scoring on almost no data, its ranking is noise and the candidates it discards were discarded at random.

Let the model stop itself

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so

```

```

25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 import time
30 from sklearn.ensemble import HistGradientBoostingClassifier
31 from sklearn.metrics import roc_auc_score
32
33 print('%-28s %8s %10s %8s' % ('', 'trees', 'test auc', 'seconds'))
34 for label, kw in [('max_iter=1000, no stopping',
35                   dict(max_iter=1000, early_stopping=False)),
36                   ('max_iter=1000, early stopping',
37                   dict(max_iter=1000, early_stopping=True,
38                       n_iter_no_change=20, validation_fraction=0.15))]:
39     t = time.time()
40     m = HistGradientBoostingClassifier(random_state=42, **kw).fit(X_tr, y_tr)
41     print('%-28s %8d %10.4f %8.1f'
42           % (label, m.n_iter_,
43              roc_auc_score(y_te, m.predict_proba(X_te)[: , 1]),
44              time.time() - t))

```

	trees	test auc	seconds
max_iter=1000, no stopping	1000	0.7287	3.9
max_iter=1000, early stopping	47	0.7866	0.2

`n_estimators` is the one hyperparameter you should almost never search for a boosted model, because the model can find it itself for the price of a single fit. Set it high, turn on early stopping, and spend your search budget on the parameters that cannot be discovered that way.

Bayesian optimisation, from scratch

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows

```

```

23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 import numpy as np
30 from scipy.stats import norm
31 from sklearn.ensemble import HistGradientBoostingClassifier
32 from sklearn.gaussian_process import GaussianProcessRegressor
33 from sklearn.gaussian_process.kernels import Matern, ConstantKernel
34 from sklearn.model_selection import cross_val_score
35
36 def objective(log_lr):
37     m = HistGradientBoostingClassifier(learning_rate=float(10 ** log_lr),
38                                       random_state=42)
39     return cross_val_score(m, X_tr, y_tr, cv=5, scoring='roc_auc').mean()
40
41 rng = np.random.default_rng(0)
42 seen_x = list(rng.uniform(-2.3, -0.4, size=5))      # 0.005 to 0.4
43 seen_y = [objective(x) for x in seen_x]
44 print('%d random starting points, best %.4f\n' % (len(seen_x), max(seen_y)))
45
46 XI = 0.002      # how much improvement counts as worth having
47 grid = np.linspace(-2.3, -0.4, 60).reshape(-1, 1)
48 for step in range(6):
49     gp = GaussianProcessRegressor(ConstantKernel(1.0) * Matern(nu=2.5),
50                                   alpha=1e-6, normalize_y=True,
51                                   random_state=0)
52     gp.fit(np.array(seen_x).reshape(-1, 1), seen_y)
53     mu, sd = gp.predict(grid, return_std=True)
54     best = max(seen_y) + XI      # XI keeps it exploring
55     z = (mu - best) / np.maximum(sd, 1e-9)
56     ei = (mu - best) * norm.cdf(z) + sd * norm.pdf(z)  # expected improvement
57     nxt = float(grid[int(ei.argmax()), 0])
58     seen_x.append(nxt)
59     seen_y.append(objective(nxt))
60     print('step %d: tried lr=%.4f → %.4f   (best so far %.4f)'
61           % (step + 1, 10 ** nxt, seen_y[-1], max(seen_y)))
62
63 i = int(np.argmax(seen_y))
64 print('\nbest learning rate %.4f scoring %.4f after %d evaluations'
65       % (10 ** seen_x[i], seen_y[i], len(seen_y)))

```

5 random starting points, best 0.8045

```

step 1: tried lr=0.0297 → 0.8054   (best so far 0.8054)
step 2: tried lr=0.3981 → 0.7788   (best so far 0.8054)
step 3: tried lr=0.0430 → 0.8031   (best so far 0.8054)
step 4: tried lr=0.0221 → 0.8053   (best so far 0.8054)
step 5: tried lr=0.0105 → 0.8007   (best so far 0.8054)
step 6: tried lr=0.0579 → 0.7969   (best so far 0.8054)

```

best learning rate 0.0297 scoring 0.8054 after 11 evaluations

When this is worth the complexity

Bayesian optimisation earns its keep when a single evaluation is expensive, minutes or hours, because then it is worth thinking hard about where to spend the next one. When a fit takes a second, the thinking costs more than the fitting and random search with more draws wins on wall-clock. Libraries such as Optuna implement this properly, with pruning and parallelism; the version above is here so the idea is not a black box.

Day 2 takeaway

Day 3 Tuning the Whole Pipeline

Preprocessing as a hyperparameter, caching, and the bias in `best_score__`

The choices made before the model sees the data, which imputer, which encoder, whether to scale, are hyperparameters too. They are usually decided once, by hand, and never revisited, which is odd given how much they can matter.

Tuning the preprocessing

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.model_selection import train_test_split
15
16 prep = ColumnTransformer([
17     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
18                     ('s', StandardScaler())])), NUM),
19     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
20                     ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
21 ])
22
23 X, y = df[NUM + CAT], df['churned']
24 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
25                                         stratify=y, random_state=42)
26 from sklearn.ensemble import HistGradientBoostingClassifier
27 from sklearn.linear_model import LogisticRegression
28 from sklearn.model_selection import RandomizedSearchCV
29 from scipy.stats import loguniform
30
31 pipe = Pipeline([('prep', prep),

```

```

32         ('clf', LogisticRegression(max_iter=2000))])
33
34 space = {
35     'prep__num__i__strategy': ['median', 'mean'],
36     'prep__num__s': [StandardScaler(), 'passthrough'],
37     'prep__cat__o__min_frequency': [1, 20, 60],
38     'clf__C': loguniform(0.01, 100),
39 }
40 search = RandomizedSearchCV(pipe, space, n_iter=20, cv=5,
41                             scoring='roc_auc', random_state=0)
42 search.fit(X_tr, y_tr)
43 print('best cross-validated AUC %.4f' % search.best_score_)
44 for k, v in sorted(search.best_params_.items()):
45     print(' %-28s %s' % (k, v))

```

```

best cross-validated AUC 0.8219
clf__C                0.19795388574374428
prep__cat__o__min_frequency  20
prep__num__i__strategy      median
prep__num__s                StandardScaler()

```

The double underscore is a path

`prep__num__i__strategy` reads as: the step named `prep`, inside it the transformer named `num`, inside that the step named `i`, and its `strategy` argument. Any nesting depth works, and passing `'passthrough'` in place of a step is how you make the existence of that step itself a thing to be tuned.

Caching the repeated work, and when not to

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.model_selection import train_test_split
15
16 prep = ColumnTransformer([
17     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
18                       ('s', StandardScaler())])), NUM),
19     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
20                       ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
21 ])
22

```

```

23 X, y = df[NUM + CAT], df['churned']
24 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
25                                         stratify=y, random_state=42)
26 import shutil
27 import tempfile
28 import time
29 from sklearn.linear_model import LogisticRegression
30 from sklearn.model_selection import GridSearchCV
31
32 space = {'clf__C': [0.01, 0.1, 1.0, 10.0, 100.0]}
33 cache = tempfile.mkdtemp()
34
35 for label, memory in [('no cache', None), ('memory=cache', cache)]:
36     pipe = Pipeline([('prep', prep),
37                     ('clf', LogisticRegression(max_iter=2000))],
38                     memory=memory)
39     t = time.time()
40     GridSearchCV(pipe, space, cv=5, scoring='roc_auc').fit(X_tr, y_tr)
41     print('%-14s %.2fs' % (label, time.time() - t))
42
43 shutil.rmtree(cache, ignore_errors=True)
44 print('\nthe preprocessing is identical for every value of C, so the')
45 print('cache turns 25 fits of it into 5 -- plus 25 trips to disk.')

```

```

no cache      0.44s
memory=cache  0.88s

```

```

the preprocessing is identical for every value of C, so the
cache turns 25 fits of it into 5 -- plus 25 trips to disk.

```

⚠ Watch Out

Measure it; here the cache made things slower

Caching trades recomputation for serialisation and disk I/O. This `ColumnTransformer` is an imputer, a scaler and a one-hot encoder over 2,250 rows. It costs a few milliseconds, and writing the result to disk and reading it back costs more than that. So the cached version loses.

It wins, and wins enormously, when the preprocessing is genuinely expensive: fitting a `TfidfVectorizer` over a large corpus, a `KNNImputer`, a target encoder, anything reading from disk. The rule is the same one as for successive halving: the optimisation is only an optimisation if you timed it.

How much of the gain is real?

A search reports the score of the best candidate it found. That figure is the maximum of many noisy numbers, and a maximum is biased upwards: some of what the winner won by is genuine and some of it is the fold split happening to suit it.

```

1 import numpy as np
2 import pandas as pd
3

```

```

4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 import numpy as np
30 from scipy.stats import loguniform, randint
31 from sklearn.ensemble import HistGradientBoostingClassifier
32 from sklearn.model_selection import RandomizedSearchCV
33 from sklearn.metrics import roc_auc_score
34
35 dists = {'learning_rate': loguniform(0.005, 0.4),
36          'max_leaf_nodes': randint(4, 80),
37          'min_samples_leaf': randint(3, 120)}
38
39 print('%-12s %14s %12s %10s' % ('candidates', 'best cv score', 'test score',
40                                'gap'))
41 for n in [5, 20, 60]:
42     s = RandomizedSearchCV(HistGradientBoostingClassifier(random_state=42),
43                            dists, n_iter=n, cv=5, scoring='roc_auc',
44                            random_state=0).fit(X_tr, y_tr)
45     test = roc_auc_score(y_te, s.predict_proba(X_te)[:, 1])
46     print('%-12d %14.4f %12.4f %10.4f'
47           % (n, s.best_score_, test, s.best_score_ - test))

```

candidates	best cv score	test score	gap
5	0.8118	0.7976	0.0142
20	0.8160	0.8038	0.0121
60	0.8160	0.8038	0.0121

The gap is positive at every budget: the cross-validated figure is optimistic by more than a point of AUC each time. Note also that trying sixty candidates instead of twenty bought nothing at all here, the same winner, the same score. A search that has stopped improving is telling you the ceiling is set by the features, not the settings, and that is a signal to go back to week 8 rather

than to widen the search.

⚠ Watch Out

Never report `best_score_` as your result

It is the winner of a competition you ran on your own validation folds, so it carries the same selection bias as any other "best of many" statistic. Report the held-out score, or the nested cross-validation score from week 7 day 1 if you have no held-out set to spare. `best_score_` is for comparing candidates against each other, and for nothing else.

Day 3 takeaway

Day 4 Partial Dependence and ICE

How the prediction responds to a feature, and when the average lies

A tuned model that nobody understands is difficult to defend and impossible to debug. Week 6 established which features a model relies on. This is the next question: *how* does the prediction change as a feature changes?

Partial dependence, by hand first

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)

```

```

28 X_te = X_te.fillna(medians)
29 import numpy as np
30 from sklearn.ensemble import HistGradientBoostingClassifier
31
32 m = HistGradientBoostingClassifier(random_state=42).fit(X_tr, y_tr)
33
34 grid = [1, 6, 12, 24, 36, 48, 60, 72]
35 print('%14s %14s' % ('tenure_months', 'mean P(churn)'))
36 for value in grid:
37     copy = X_tr.copy()
38     copy['tenure_months'] = value # everybody, at this tenure
39     print('%14d %14.4f' % (value, m.predict_proba(copy)[: , 1].mean()))

```

tenure_months	mean P(churn)
1	0.3687
6	0.2390
12	0.3096
24	0.1227
36	0.1526
48	0.0202
60	0.0079
72	0.0082

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16         'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 from sklearn.ensemble import HistGradientBoostingClassifier
30 from sklearn.inspection import partial_dependence
31

```

```

32 m = HistGradientBoostingClassifier(random_state=42).fit(X_tr, y_tr)
33
34 for method in ['decision_function', 'predict_proba']:
35     res = partial_dependence(m, X_tr, ['tenure_months'],
36                             grid_resolution=5, kind='average',
37                             # 'recursion' is the fast tree-specific path
38                             # and only speaks log-odds; asking for
39                             # probabilities means doing it the slow way.
40                             method='brute', response_method=method)
41     print('response_method=%s' % method)
42     print(' ' + ' '.join('%8.4f' % a for a in res['average'][0]))
43 print('\ngrid: ' + ' '.join('%8.1f' % v
44                             for v in res['grid_values'][0]))

```

```

response_method=decision_function
-1.2841  -1.9846  -3.0956  -5.3464  -5.9621
response_method=predict_proba
 0.3101   0.2245   0.1138   0.0141   0.0079

grid:      4.0      19.2      34.5      49.8      65.0

```

⚠ Watch Out

Check the scale before you read the numbers

For a classifier, `response_method='auto'` uses `decision_function` when the estimator has one, so the default output is in *log-odds*, not probability, and will not match the hand-rolled version above. Log-odds is the better scale for seeing shape, because it is the scale the model is additive on. Probability is the better scale for showing anybody else. Pick deliberately, and label the axis.

The average can hide the story

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],

```

```

20     random_state=42)
21
22     # Impute after the split, with medians learned from the training rows
23     # only -- week 4's rule still applies. A Pipeline is the right way to do
24     # this; these pages need a plain frame to inspect column by column, so
25     # the two lines are written out instead.
26     medians = X_tr.median()
27     X_tr = X_tr.fillna(medians)
28     X_te = X_te.fillna(medians)
29     import numpy as np
30     from sklearn.ensemble import HistGradientBoostingClassifier
31     from sklearn.inspection import partial_dependence
32
33     m = HistGradientBoostingClassifier(random_state=42).fit(X_tr, y_tr)
34     res = partial_dependence(m, X_tr, ['monthly_charges'],
35                             grid_resolution=6, kind='both')
36     grid = res['grid_values'][0]
37     ice = res['individual'][0]
38
39     slope = ice[:, -1] - ice[:, 0]
40     print('%d customers rise with price, %d fall'
41           % (int((slope > 0).sum()), int((slope < 0).sum())))
42     print('\n%14s %10s %10s %10s' % ('charges', 'average', 'p10', 'p90'))
43     for j, g in enumerate(grid):
44         col = ice[:, j]
45         print('%14.1f %10.4f %10.4f %10.4f'
46               % (g, col.mean(), np.percentile(col, 10),
47                  np.percentile(col, 90)))

```

2123 customers rise with price, 127 fall

charges	average	p10	p90
15.0	0.1230	0.0017	0.3105
30.2	0.1276	0.0022	0.3310
45.4	0.1177	0.0007	0.3135
60.6	0.2636	0.0014	0.6268
75.8	0.2301	0.0034	0.5721
91.0	0.3635	0.0063	0.8330

Two features at once

```

1  import numpy as np
2  import pandas as pd
3
4  df = pd.read_csv('customers.csv').drop_duplicates().copy()
5  df['contract'] = df['contract'].str.strip().str.title()
6  df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8  NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9  CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()

```

```

13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 from sklearn.ensemble import HistGradientBoostingClassifier
30 from sklearn.inspection import partial_dependence
31
32 m = HistGradientBoostingClassifier(random_state=42).fit(X_tr, y_tr)
33 # Two names in one list asks for the joint surface, not two curves.
34 res = partial_dependence(m, X_tr, ['tenure_months', 'is_monthly'],
35                           grid_resolution=6, kind='average')
36 tenure, monthly = res['grid_values']
37 surface = res['average'][0]
38
39 print('rows = tenure, columns = contract type')
40 print('%10s %12s %12s %12s' % ('tenure', 'longer term', 'month-to-month',
41                                'difference'))
42 for i, t in enumerate(tenure):
43     a, b = surface[i, 0], surface[i, 1]
44     print('%10.1f %12.4f %12.4f %12.4f' % (t, a, b, b - a))

```

```

rows = tenure, columns = contract type
  tenure  longer term month-to-month  difference
    4.0      -0.9477      0.6599      1.6076
   16.2     -1.7094      0.2208      1.9301
   28.4     -3.9701     -0.0364      3.9337
   40.6     -3.9225     -1.7032      2.2193
   52.8     -4.6525     -2.6154      2.0370
   65.0     -5.9758     -4.0526      1.9232

```

If the difference column were constant, the two features would act independently and the model would be additive in them. Where it changes, there is an interaction: the effect of tenure depends on the contract. That is precisely the structure a linear model cannot represent without being told, and the reason week 8 spent a day on interaction features.

⚠ Watch Out

Partial dependence averages over combinations that do not exist

Setting `tenure_months=1` for every customer includes customers whose `total_charges` is £3,000, which is a person who has spent three thousand pounds in their first month.

The model is being asked about a region of the input space with no data in it, and its answer there is extrapolation dressed as insight. The stronger the correlation between your features, the less trustworthy the curve, check the correlation before you present one.

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 print('correlation with tenure_months:')
30 print(X_tr.corr()['tenure_months'].drop('tenure_months').round(3)
31       .sort_values(key=abs, ascending=False).to_string())
32 print('\ntotal_charges is nearly determined by tenure, so a partial')
33 print('dependence plot for one of them holds the other at impossible values.')
```

```
correlation with tenure_months:
total_charges      0.829
is_monthly         -0.463
support_calls      0.033
is_fibre           0.014
monthly_charges    0.007
```

```
total_charges is nearly determined by tenure, so a partial
dependence plot for one of them holds the other at impossible values.
```

Day 4 takeaway

Day 5 SHAP

Per-prediction contributions that add up, and what they cost

Partial dependence describes the model. It does not tell a particular customer why *they* were flagged, and that is usually the question that gets asked.

One prediction, explained

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 import shap
30 from sklearn.ensemble import HistGradientBoostingClassifier
31
32 m = HistGradientBoostingClassifier(random_state=42).fit(X_tr, y_tr)
33 explainer = shap.TreeExplainer(m)
34 sv = explainer(X_te.iloc[:50])
35
36 i = int(m.predict_proba(X_te.iloc[:50])[:, 1].argmax()) # most at risk
37 row = X_te.iloc[i]
38 print('customer predicted at %.3f probability of churn'
39       % m.predict_proba(X_te.iloc[[i]])[0, 1])
40 print('\n%-18s %10s %12s' % ('feature', 'value', 'contribution'))
41 order = np.argsort(-np.abs(sv.values[i]))
42 for j in order:
43     print('%-18s %10.2f %+12.4f'
44           % (X_te.columns[j], row.iloc[j], sv.values[i][j]))
45 print('%-18s %10s %+12.4f' % ('base value', '', sv.base_values[i]))

```

```

46 print('%-18s %10s %+12.4f'
47       % ('sum', '', sv.values[i].sum() + sv.base_values[i]))
48 print('model log-odds output %+12.4f' % m.decision_function(X_te.iloc[[i]])[0])

```

customer predicted at 0.814 probability of churn

feature	value	contribution
is_monthly	1.00	+2.0362
support_calls	4.00	+1.4048
monthly_charges	88.05	+1.0216
total_charges	2590.66	-0.5712
tenure_months	28.00	-0.1467
is_fibre	1.00	+0.0414
base value		-2.3128
sum		+1.4733

model log-odds output +1.4733

Additivity is the property that makes this usable

The contributions plus the base value equal the model's output exactly, not approximately. That is what lets you put the numbers in front of somebody and say "these five things, adding to this" without hand-waving. Note the scale: for a classifier, contributions are in log-odds, not probability, so they add up on a scale most audiences do not think in. Convert carefully, or explain the scale.

Global importance, assembled from local explanations

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)

```

```

28 X_te = X_te.fillna(medians)
29 import shap
30 from sklearn.ensemble import HistGradientBoostingClassifier
31 from sklearn.inspection import permutation_importance
32
33 m = HistGradientBoostingClassifier(random_state=42).fit(X_tr, y_tr)
34 sv = shap.TreeExplainer(m)(X_te)
35 shap_rank = np.abs(sv.values).mean(axis=0)
36
37 perm = permutation_importance(m, X_te, y_te, n_repeats=10,
38                               random_state=0, scoring='roc_auc')
39
40 print('%-18s %-14s %-16s' % ('feature', 'mean |shap|', 'permutation'))
41 for j in np.argsort(-shap_rank):
42     print('%-18s %-14.4f %-16.4f'
43           % (X_te.columns[j], shap_rank[j], perm.importances_mean[j]))

```

feature	mean shap	permutation
is_monthly	1.3276	0.0984
tenure_months	0.9738	0.0804
monthly_charges	0.5814	0.0245
total_charges	0.5090	-0.0027
support_calls	0.4216	0.0191
is_fibre	0.0885	-0.0015

The two rankings usually agree, and when they disagree it is informative rather than alarming: permutation importance measures how much the *score* falls when a feature is scrambled, so it is about predictive usefulness. Mean absolute SHAP measures how much the *prediction* moves, so a feature can swing predictions strongly in both directions and still not improve accuracy.

What it costs

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21

```

```

22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 import time
30 import shap
31 from sklearn.ensemble import HistGradientBoostingClassifier
32 from sklearn.linear_model import LogisticRegression
33
34 m = HistGradientBoostingClassifier(random_state=42).fit(X_tr, y_tr)
35 lr = LogisticRegression(max_iter=2000).fit(X_tr, y_tr)
36
37 t = time.time()
38 shap.TreeExplainer(m)(X_te)
39 tree = time.time() - t
40
41 t = time.time()
42 shap.KernelExplainer(lr.predict_proba,
43                      shap.sample(X_tr, 50, random_state=0)).shap_values(
44                      X_te.iloc[:25], nsamples=100, silent=True)
45 kernel = time.time() - t
46
47 print('%-16s %6s %8s %14s' % ('', 'rows', 'seconds', 'ms per row'))
48 print('%-16s %6d %8.2f %14.2f'
49       % ('TreeExplainer', len(X_te), tree, 1000 * tree / len(X_te)))
50 print('%-16s %6d %8.2f %14.2f'
51       % ('KernelExplainer', 25, kernel, 1000 * kernel / 25))
52 print('\ncompare the last column, not the middle one.')

```

	rows	seconds	ms per row
TreeExplainer	750	0.22	0.29
KernelExplainer	25	0.06	2.59

compare the last column, not the middle one.

Per row, the sampling explainer is an order of magnitude more expensive, and that is with only six features, a hundred samples and a fifty-row background set. All three of those have to grow for a real problem, and the cost grows with them. **TreeExplainer** avoids all of it by exploiting the tree structure to compute the exact answer directly, which is why tree models are the ones people actually explain at scale.

Watch Out

Correlated features share credit arbitrarily

If two features carry the same information, any split of the credit between them produces the same prediction, so the axioms do not pin down a unique answer, the implementation picks one. A stakeholder reading "total_charges contributed most" may act on it, when **tenure_months** would have served identically. Check your correlations before you present attributions, and consider reporting correlated features as a group.

Day 5 takeaway**Day 6 Counterfactuals, Surrogates and Causality**

What would have to change, readable approximations, and the limit

Two more tools, both aimed at the gap between what a model computes and what a person can act on.

Counterfactuals: what would have to change

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 import numpy as np
30 from sklearn.ensemble import HistGradientBoostingClassifier
31
32 m = HistGradientBoostingClassifier(random_state=42).fit(X_tr, y_tr)
33 proba = m.predict_proba(X_te)[ :, 1]
34 i = int(proba.argmax())
35 row = X_te.iloc[[i]]
36 print('customer starts at P(churn) = %.3f\n' % proba[i])
37
38 # Only the levers the business can actually pull.
39 moves = {'monthly_charges': [-20, -10, -5],
40          'support_calls': [-3, -2, -1],

```

```

41         'is_monthly': [0]}
42
43 print('%-20s %10s %12s %10s' % ('change', 'new value', 'P(churn)', 'drop'))
44 for feature, deltas in moves.items():
45     for delta in deltas:
46         candidate = row.copy()
47         new = (delta if feature == 'is_monthly'
48               else max(0, float(row[feature].iloc[0]) + delta))
49         candidate[feature] = new
50         p = m.predict_proba(candidate)[0, 1]
51         print('%-20s %10.1f %12.3f %10.3f'
52               % ('%s %+g' % (feature, delta) if feature != 'is_monthly'
53                 else 'move off monthly', new, p, proba[i] - p))

```

customer starts at $P(\text{churn}) = 0.923$

change	new value	P(churn)	drop
monthly_charges -20	69.6	0.693	0.230
monthly_charges -10	79.6	0.619	0.304
monthly_charges -5	84.6	0.781	0.142
support_calls -3	0.0	0.545	0.378
support_calls -2	0.0	0.545	0.378
support_calls -1	1.0	0.871	0.052
move off monthly	0.0	0.199	0.724

Moving this customer off a month-to-month contract is worth more than any price change on offer, which is an actionable finding rather than a statistical one. Notice too that the price column is not monotone: a £10 reduction helps more than a £20 one. That is not a bug. A boosted tree is a step function, so moving a customer across one threshold can help while moving them further lands them in a different leaf. Search the whole range rather than assuming more of a good thing is better.

Only offer changes somebody can make

A counterfactual saying "if this customer had been with us for forty more months" is useless, because nobody can act on it. Restrict the search to *actionable* features, price, contract type, the number of unresolved support calls, and the output becomes a retention offer rather than a curiosity. It also keeps you away from features that would be unlawful to act on.

Surrogate models

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()

```

```

13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 from sklearn.ensemble import HistGradientBoostingClassifier
30 from sklearn.tree import DecisionTreeRegressor, export_text
31
32 m = HistGradientBoostingClassifier(random_state=42).fit(X_tr, y_tr)
33 target = m.predict_proba(X_tr)[: , 1] # the black box's own output
34
35 print('%8s %16s' % ('depth', 'fidelity (R2)'))
36 for depth in [2, 3, 4, 6]:
37     s = DecisionTreeRegressor(max_depth=depth,
38                               random_state=0).fit(X_tr, target)
39     print('%8d %16.4f' % (depth, s.score(X_tr, target)))
40
41 s = DecisionTreeRegressor(max_depth=3, random_state=0).fit(X_tr, target)
42 print('\n' + export_text(s, feature_names=list(X_tr.columns),
43                           decimals=2)[:1100])

```

depth	fidelity (R2)
2	0.4484
3	0.5437
4	0.6177
6	0.7213

```

|--- is_monthly <= 0.50
| |--- tenure_months <= 17.50
| | |--- support_calls <= 2.50
| | | |--- value: [0.18]
| | |--- support_calls > 2.50
| | | |--- value: [0.43]
| |--- tenure_months > 17.50
| | |--- monthly_charges <= 75.74
| | | |--- value: [0.02]
| | |--- monthly_charges > 75.74
| | | |--- value: [0.07]
|--- is_monthly > 0.50
| |--- monthly_charges <= 50.98
| | |--- total_charges <= 237.06
| | | |--- value: [0.34]
| | |--- total_charges > 237.06
| | | |--- value: [0.20]
| |--- monthly_charges > 50.98

```

```
| | |--- tenure_months <= 32.50
| | |--- value: [0.52]
| | |--- tenure_months > 32.50
| | |--- value: [0.09]
```

⚠ Watch Out

Fidelity is the number that matters, and it is usually omitted

A surrogate is an explanation of the model only to the extent that it agrees with the model. Quote the R^2 whenever you show one: at 0.9 it is a fair summary, at 0.5 it is a different model that happens to be readable, and presenting the second as an explanation is worse than presenting nothing. Fidelity is also local. A surrogate can agree well on average and disagree exactly on the cases you care about.

By that standard, the readable depth-3 tree printed above does not qualify: it accounts for about half the variation in the model's output. It is a reasonable description of the broad shape and it should not be presented as what the model does. Depth 6 reaches roughly 0.72 and is already too large to read aloud, which is the trade in a single sentence, readable and faithful are in tension, and a surrogate that is comfortable to present is usually the unfaithful one.

Explanations are not causal

This is the most consequential misunderstanding in the whole subject, so it is worth demonstrating rather than asserting. Add a column that is pure consequence. A retention call that the company makes *because* a customer looks like they are leaving, and watch every interpretability tool point at it.

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16         'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
```

```

25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 import numpy as np
30 import shap
31 from sklearn.ensemble import HistGradientBoostingClassifier
32 from sklearn.metrics import roc_auc_score
33
34 rng = np.random.default_rng(0)
35 # The retention team rings customers who are about to leave. The call is
36 # caused by the churn risk; it does not cause it.
37 call_tr = (y_tr.to_numpy() * (rng.random(len(y_tr)) < 0.8)).astype(int)
38 call_te = (y_te.to_numpy() * (rng.random(len(y_te)) < 0.8)).astype(int)
39
40 A = X_tr.assign(retention_call=call_tr)
41 B = X_te.assign(retention_call=call_te)
42 m = HistGradientBoostingClassifier(random_state=42).fit(A, y_tr)
43
44 print('test AUC with the new column %.4f'
45       % roc_auc_score(y_te, m.predict_proba(B)[:, 1]))
46 sv = shap.TreeExplainer(m)(B)
47 rank = np.abs(sv.values).mean(axis=0)
48 print('\nmean |shap|:')
49 for j in np.argsort(-rank):
50     print(' %-18s %.4f' % (A.columns[j], rank[j]))
51 print('\nSHAP is right: the model does rely on it.')
52 print('Acting on it -- stop making retention calls -- changes nothing.')

```

```
test AUC with the new column 0.9489
```

```
mean |shap|:
  retention_call      4.6405
  tenure_months      0.9912
  is_monthly          0.6624
  monthly_charges     0.5483
  total_charges       0.5097
  support_calls       0.3393
  is_fibre            0.0401
```

```
SHAP is right: the model does rely on it.
```

```
Acting on it -- stop making retention calls -- changes nothing.
```

Watch Out

What every interpretability tool is actually telling you

SHAP, permutation importance and partial dependence all answer the same kind of question: *how does this model use this feature?* None of them answers *what happens in the world if we change this feature?* The second question needs an experiment, or causal assumptions you are prepared to defend in writing. Every time an importance ranking is read as a list of things to go and change, this distinction is the one being skipped.

Day 6 takeaway

Day 7 An Interpretability Report

One function, and the sentences to use when presenting it

One function that produces everything a stakeholder needs, and an honest account of what to say about it.

The report

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11
12 d = df.copy()
13 d['is_monthly'] = (d['contract'] == 'Month-To-Month').astype(int)
14 d['is_fibre'] = (d['internet_service'] == 'Fibre optic').astype(int)
15 FEATS = ['tenure_months', 'support_calls', 'monthly_charges',
16          'total_charges', 'is_monthly', 'is_fibre']
17
18 X_tr, X_te, y_tr, y_te = train_test_split(
19     d[FEATS], d['churned'], test_size=0.25, stratify=d['churned'],
20     random_state=42)
21
22 # Impute after the split, with medians learned from the training rows
23 # only -- week 4's rule still applies. A Pipeline is the right way to do
24 # this; these pages need a plain frame to inspect column by column, so
25 # the two lines are written out instead.
26 medians = X_tr.median()
27 X_tr = X_tr.fillna(medians)
28 X_te = X_te.fillna(medians)
29 import numpy as np
30 import shap
31 from sklearn.ensemble import HistGradientBoostingClassifier
32 from sklearn.inspection import partial_dependence
33 from sklearn.metrics import roc_auc_score
34
35 def explain(model, X, y, top=4):
36     print('=' * 58)
37     print('MODEL EXPLANATION REPORT')
38     print('=' * 58)
39     p = model.predict_proba(X)[: , 1]
40     print('rows %d    AUC %.4f    mean predicted risk %.3f'
```

```

41         % (len(X), roc_auc_score(y, p), p.mean()))
42
43     sv = shap.TreeExplainer(model)(X)
44     rank = np.abs(sv.values).mean(axis=0)
45     order = np.argsort(-rank)[:top]
46
47     print('\nWHAT THE MODEL USES')
48     for j in order:
49         direction = np.corrcoef(X.iloc[:, j], sv.values[:, j])[0, 1]
50         arrow = 'higher raises risk' if direction > 0 else 'higher lowers risk'
51         print(' %-18s %.4f  %s' % (X.columns[j], rank[j], arrow))
52
53     print('\nSHAPE OF THE TOP EFFECT')
54     j = int(order[0])
55     res = partial_dependence(model, X, [X.columns[j]], grid_resolution=6,
56                             kind='average')
57     for v, a in zip(res['grid_values'][0], res['average'][0]):
58         print(' %-14s %8.2f → %.4f' % (X.columns[j], v, a))
59
60     print('\nWHO IS MOST AT RISK')
61     for i in np.argsort(-p)[:3]:
62         top_j = int(np.argmax(np.abs(sv.values[i])))
63         print('  row %-5d p=%.3f  driven by %s=%.1f'
64               % (i, p[i], X.columns[top_j], X.iloc[i, top_j]))
65     print('=' * 58)
66
67 m = HistGradientBoostingClassifier(random_state=42).fit(X_tr, y_tr)
68 explain(m, X_te, y_te)

```

```

=====
MODEL EXPLANATION REPORT
=====
rows 750   AUC 0.7688   mean predicted risk 0.249

WHAT THE MODEL USES
  is_monthly      1.3276   higher raises risk
  tenure_months   0.9738   higher lowers risk
  monthly_charges  0.5814   higher raises risk
  total_charges    0.5090   higher lowers risk

SHAPE OF THE TOP EFFECT
  is_monthly      0.00 → -3.0695
  is_monthly      1.00 → 0.0764

WHO IS MOST AT RISK
  row 60    p=0.923  driven by monthly_charges=89.6
  row 545   p=0.923  driven by tenure_months=2.0
  row 119   p=0.915  driven by monthly_charges=92.5
=====

```

Saying it out loud

Do not say	Say instead	Why
"Tenure is the most important factor"	"Tenure moves this model's predictions most"	Importance is a property of the model, not the world
"Reducing price would cut churn by 8%"	"Customers at lower prices churn less; testing a price change would tell us whether lowering it helps"	The first claims a causal effect the model cannot support
"The model is 87% accurate"	"It ranks well, AUC 0.87, and at our threshold it catches 6 in 10 leavers with 1 in 3 flagged"	Accuracy alone hides the trade-off that actually gets chosen
"The model explains why they left"	"These features drove this score"	It explains a prediction, not a departure

Before you present a model

1. Report a score from data the tuning never saw, not `best_score_`.
2. Show which features the model uses, and by how much, with the caveat that correlated features share credit ambiguously.
3. Show the *shape* of the top two or three effects, not just their rank, and check the correlations before you trust the curve.
4. Explain two or three individual predictions end to end, including one the model got wrong.
5. State what the model would need in order to be causal, and that it is not.
6. Say what the model does when it has never seen this kind of row before.
7. Agree the decision threshold with whoever bears the cost of the errors, using week 7's framing.
8. Write down the date, the data window and the version. Week 15 explains why that last one is not administrative pedantry.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. Random search usually beats grid search at the same budget because:
 - a) A grid repeats the same few values of the parameter that matters, while random search tries a distinct value every time
 - b) It uses better defaults
 - c) Grids cannot use distributions
 - d) It is faster per fit
2. For a random forest, `n_estimators` should be:
 - a) Tuned carefully

- b) Set as high as you can afford, since more trees never makes a forest worse
 - c) Kept below 50
 - d) Matched to the feature count
3. Successive halving lost time against plain random search on this dataset because:
- a) It used too few candidates
 - b) It is a poor algorithm
 - c) There is nothing to save by fitting on fewer rows when a full fit already takes a fraction of a second
 - d) The seed was wrong
4. `best_score_` from a search should never be reported as your result because:
- a) It uses the wrong metric
 - b) It is not cross-validated
 - c) It is on the test set
 - d) It is the maximum of many noisy numbers, and a maximum is biased upwards
5. A partial dependence curve is untrustworthy when the feature is:
- a) Strongly correlated with another, so the calculation averages over combinations that never occur
 - b) Binary
 - c) Scaled
 - d) Categorical
6. ICE curves add to partial dependence by showing:
- a) A confidence interval
 - b) One line per row, revealing whether the average represents anybody at all
 - c) The feature importance
 - d) The residuals
7. SHAP values for a single prediction:
- a) Are in probability units
 - b) Approximate the output
 - c) Sum with the base value to the model's output exactly
 - d) Are always positive
8. An importance ranking tells you:
- a) Which feature caused the outcome
 - b) The correlation with the target
 - c) What to change in the world to affect the outcome
 - d) How the model uses each feature, which is not a causal claim

Answers: 1a, 2b, 3c, 4d, 5a, 6b, 7c, 8d. Anything you got wrong points at the day to re-read, not at a gap in ability.

Deployment and MLOps

Week 15

Package a model, serve it over HTTP, validate what arrives, watch it drift and replace it safely.

OBJECTIVES

A model in a notebook is a result; a model in production is a commitment. This week covers the distance: what has to be saved besides the estimator, how to serve it and validate what arrives, which failures are silent and therefore dangerous, how to detect drift without labels, and how to replace a live model without anybody having a bad afternoon.

Day 1 Packaging a Model

The artefact, the contract, pinned versions and a smoke test

A model that lives in a notebook is not a product. Everything from here is about the distance between the two, and almost none of it is about machine learning.

What actually has to be saved

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.model_selection import train_test_split
16
17 prep = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19     ('s', StandardScaler())])), NUM),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21     ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),

```

```

22 ])
23 model = Pipeline([('prep', prep),
24                   ('clf', LogisticRegression(max_iter=2000))])
25
26 X, y = df[NUM + CAT], df['churned']
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                           stratify=y, random_state=42)
29 model.fit(X_tr, y_tr)
30 import joblib
31 import json
32 import sklearn
33 import sys
34 from sklearn.metrics import roc_auc_score
35
36 joblib.dump(model, 'churn_model.joblib')
37
38 contract = {
39     'model_version': '1.0.0',
40     'trained_at': '2026-08-08',
41     'trained_rows': int(len(X_tr)),
42     'features': {'numeric': NUM, 'categorical': CAT},
43     'categories': {c: sorted(df[c].dropna().unique().tolist())
44                   for c in CAT},
45     'ranges': {c: [float(df[c].min()), float(df[c].max())] for c in NUM},
46     'threshold': 0.35,
47     'test_auc': round(float(roc_auc_score(
48         y_te, model.predict_proba(X_te)[: , 1])), 4),
49     'python': sys.version.split()[0],
50     'sklearn': sklearn.__version__,
51 }
52 json.dump(contract, open('churn_contract.json', 'w'), indent=2)
53 print(json.dumps(contract, indent=2)[:760])

```

```

{
  "model_version": "1.0.0",
  "trained_at": "2026-08-08",
  "trained_rows": 2250,
  "features": {
    "numeric": [
      "tenure_months",
      "support_calls",
      "monthly_charges"
    ],
    "categorical": [
      "contract",
      "internet_service",
      "payment_method",
      "has_dependents"
    ]
  },
  "categories": {
    "contract": [
      "Month-To-Month",
      "One Year",
      "Two Year"
    ]
  }
}

```

```

],
"internet_service": [
    "DSL",
    "Fibre optic",
    "No internet"
],
"payment_method": [
    "Bank transfer",
    "Credit card",
    "Electronic check",
... (13 more lines)

```

The contract is the part people leave out

The `.joblib` file knows how to turn a dataframe into a prediction. It does not know that `contract` has exactly three valid values, that `tenure_months` was never above 72 in training, that the agreed threshold is 0.35 rather than 0.5, or which version of the model produced a given row in your logs. Six months later, every one of those is a question somebody will ask, and the answer needs to be in a file rather than in your memory.

Pinning the environment

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.model_selection import train_test_split
16
17 prep = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                       ('s', StandardScaler())])), NUM),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                       ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22 ])
23 model = Pipeline([('prep', prep),
24                  ('clf', LogisticRegression(max_iter=2000))])
25
26 X, y = df[NUM + CAT], df['churned']
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                           stratify=y, random_state=42)
29 model.fit(X_tr, y_tr)
30 import joblib
31 import sklearn

```

```

32
33 loaded = joblib.load('churn_model.joblib')
34 print('loaded a %s' % type(loader).__name__)
35 print('steps: %s' % [name for name, _ in loader.steps])
36 print('fitted with sklearn %s' % sklearn.__version__)
37
38 same = (loader.predict_proba(X_te)[: , 1]
39         == model.predict_proba(X_te)[: , 1]).all()
40 print('reload predicts identically: %s' % bool(same))

```

```

loaded a Pipeline
steps: ['prep', 'clf']
fitted with sklearn 1.8.0
reload predicts identically: True

```

⚠ Watch Out

A pickle is not a portable format

joblib uses pickle underneath, which stores references to Python classes. Load it under a different scikit-learn version and you get, in ascending order of unpleasantness: a warning, an exception, or a model that loads cleanly and predicts subtly differently. Pin the version, load in the same environment you trained in, and add a smoke test that checks a handful of known inputs still produce known outputs. That test is the only thing that catches the third case.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.model_selection import train_test_split
16
17 prep = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                      ('s', StandardScaler())]), NUM),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                      ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22 ])
23 model = Pipeline([('prep', prep),
24                  ('clf', LogisticRegression(max_iter=2000))])
25
26 X, y = df[NUM + CAT], df['churned']

```

```

27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                           stratify=y, random_state=42)
29 model.fit(X_tr, y_tr)
30 import joblib
31 import json
32
33 # The smoke test the warning above asks for: freeze a few predictions
34 # now, and compare against them on every load, forever.
35 sample = X_te.head(5)
36 expected = [round(float(p), 6)
37             for p in model.predict_proba(sample)[: , 1]]
38 json.dump({'rows': sample.to_dict('records'), 'expected': expected},
39          open('smoke_test.json', 'w'), indent=2)
40
41 fixture = json.load(open('smoke_test.json'))
42 reloaded = joblib.load('churn_model.joblib')
43 got = [round(float(p), 6) for p in
44        reloaded.predict_proba(pd.DataFrame(fixture['rows']))[: , 1]]
45 print('expected %s' % fixture['expected'])
46 print('got      %s' % got)
47 print('smoke test passes: %s' % (got == fixture['expected']))

```

```

expected [0.530642, 0.16512, 0.322593, 0.205547, 0.568866]
got      [0.530642, 0.16512, 0.322593, 0.205547, 0.568866]
smoke test passes: True

```

Day 1 takeaway

Day 2 A Prediction Service

FastAPI, request schemas, and where the latency actually goes

A model behind an HTTP endpoint, which is how most of them are consumed. The service below is written to a file and then genuinely called, so every response on this page came back over a real request cycle.

The service

Save this as `service.py`. It is loaded and called for real further down the page.

```

1 import json
2 import joblib
3 import pandas as pd
4 from fastapi import FastAPI, HTTPException
5 from pydantic import BaseModel, Field, ValidationError
6
7 MODEL = joblib.load("churn_model.joblib")
8 CONTRACT = json.load(open("churn_contract.json"))
9
10
11 class Customer(BaseModel):

```

```

12     """The request schema. Pydantic rejects anything that does not match,
13     before a single row reaches the model."""
14     tenure_months: int = Field(ge=0, le=120)
15     support_calls: int = Field(ge=0, le=100)
16     monthly_charges: float = Field(ge=0, le=500)
17     contract: str
18     internet_service: str
19     payment_method: str
20     has_dependents: str
21
22
23 app = FastAPI(title="churn-scorer", version=CONTRACT["model_version"])
24
25
26 @app.get("/health")
27 def health():
28     return {"status": "ok", "model_version": CONTRACT["model_version"],
29           "trained_at": CONTRACT["trained_at"]}
30
31
32 @app.post("/predict")
33 def predict(customer: Customer):
34     row = pd.DataFrame([customer.model_dump()])
35     for column in ("contract", "internet_service", "payment_method"):
36         known = CONTRACT["categories"][column]
37         if row.loc[0, column] not in known:
38             raise HTTPException(
39                 status_code=422,
40                 detail="unknown %s: %r (known: %s)"
41                     % (column, row.loc[0, column], known))
42     p = float(MODEL.predict_proba(row)[0, 1])
43     return {"churn_probability": round(p, 4),
44           "decision": "flag" if p >= CONTRACT["threshold"] else "keep",
45           "threshold": CONTRACT["threshold"],
46           "model_version": CONTRACT["model_version"]}

```

What a request and response look like

```

1 SERVICE = '''import json
2 import joblib
3 import pandas as pd
4 from fastapi import FastAPI, HTTPException
5 from pydantic import BaseModel, Field, ValidationError
6
7 MODEL = joblib.load("churn_model.joblib")
8 CONTRACT = json.load(open("churn_contract.json"))
9
10
11 class Customer(BaseModel):
12     """The request schema. Pydantic rejects anything that does not match,
13     before a single row reaches the model."""
14     tenure_months: int = Field(ge=0, le=120)

```

```

15     support_calls: int = Field(ge=0, le=100)
16     monthly_charges: float = Field(ge=0, le=500)
17     contract: str
18     internet_service: str
19     payment_method: str
20     has_dependents: str
21
22
23 app = FastAPI(title="churn-scorer", version=CONTRACT["model_version"])
24
25
26 @app.get("/health")
27 def health():
28     return {"status": "ok", "model_version": CONTRACT["model_version"],
29           "trained_at": CONTRACT["trained_at"]}
30
31
32 @app.post("/predict")
33 def predict(customer: Customer):
34     row = pd.DataFrame([customer.model_dump()])
35     for column in ("contract", "internet_service", "payment_method"):
36         known = CONTRACT["categories"][column]
37         if row.loc[0, column] not in known:
38             raise HTTPException(
39                 status_code=422,
40                 detail="unknown %s: %r (known: %s)"
41                     % (column, row.loc[0, column], known))
42     p = float(MODEL.predict_proba(row)[0, 1])
43     return {"churn_probability": round(p, 4),
44           "decision": "flag" if p >= CONTRACT["threshold"] else "keep",
45           "threshold": CONTRACT["threshold"],
46           "model_version": CONTRACT["model_version"]}
47 '''
48 import io, sys, json
49 io.open('service.py', 'w', encoding='utf-8').write(SERVICE)
50 sys.path.insert(0, '.')
51
52 from fastapi.testclient import TestClient
53 import service
54 client = TestClient(service.app)
55
56 print('GET /health →', client.get('/health').status_code)
57 print(json.dumps(client.get('/health').json(), indent=2))
58
59 at_risk = {'tenure_months': 2, 'support_calls': 5,
60           'monthly_charges': 94.2, 'contract': 'Month-To-Month',
61           'internet_service': 'Fibre optic',
62           'payment_method': 'Electronic check',
63           'has_dependents': 'No'}
64 settled = dict(at_risk, tenure_months=60, support_calls=0,
65               contract='Two Year', monthly_charges=42.0)
66
67 for label, payload in [('new, unhappy', at_risk),
68                       ('long-standing', settled)]:

```

```

69     r = client.post('/predict', json=payload)
70     print('\n%s → %d' % (label, r.status_code))
71     print(json.dumps(r.json(), indent=2))

```

```

GET /health → 200
{
  "status": "ok",
  "model_version": "1.0.0",
  "trained_at": "2026-08-08"
}

new, unhappy → 200
{
  "churn_probability": 0.8896,
  "decision": "flag",
  "threshold": 0.35,
  "model_version": "1.0.0"
}

long-standing → 200
{
  "churn_probability": 0.004,
  "decision": "keep",
  "threshold": 0.35,
  "model_version": "1.0.0"
}

```

Single or batch?

```

1 SERVICE = '''import json
2 import joblib
3 import pandas as pd
4 from fastapi import FastAPI, HTTPException
5 from pydantic import BaseModel, Field, ValidationError
6
7 MODEL = joblib.load("churn_model.joblib")
8 CONTRACT = json.load(open("churn_contract.json"))
9
10
11 class Customer(BaseModel):
12     """The request schema. Pydantic rejects anything that does not match,
13     before a single row reaches the model."""
14     tenure_months: int = Field(ge=0, le=120)
15     support_calls: int = Field(ge=0, le=100)
16     monthly_charges: float = Field(ge=0, le=500)
17     contract: str
18     internet_service: str
19     payment_method: str
20     has_dependents: str
21
22
23 app = FastAPI(title="churn-scorer", version=CONTRACT["model_version"])
24
25

```

```

26 @app.get("/health")
27 def health():
28     return {"status": "ok", "model_version": CONTRACT["model_version"],
29           "trained_at": CONTRACT["trained_at"]}
30
31
32 @app.post("/predict")
33 def predict(customer: Customer):
34     row = pd.DataFrame([customer.model_dump()])
35     for column in ("contract", "internet_service", "payment_method"):
36         known = CONTRACT["categories"][column]
37         if row.loc[0, column] not in known:
38             raise HTTPException(
39                 status_code=422,
40                 detail="unknown %s: %r (known: %s)"
41                     % (column, row.loc[0, column], known))
42     p = float(MODEL.predict_proba(row)[0, 1])
43     return {"churn_probability": round(p, 4),
44           "decision": "flag" if p >= CONTRACT["threshold"] else "keep",
45           "threshold": CONTRACT["threshold"],
46           "model_version": CONTRACT["model_version"]}
47 '''
48 import io, sys, json
49 io.open('service.py', 'w', encoding='utf-8').write(SERVICE)
50 sys.path.insert(0, '.')
51
52 from fastapi.testclient import TestClient
53 import service
54 client = TestClient(service.app)
55 import time
56
57 payload = {'tenure_months': 2, 'support_calls': 5,
58           'monthly_charges': 94.2, 'contract': 'Month-To-Month',
59           'internet_service': 'Fibre optic',
60           'payment_method': 'Electronic check',
61           'has_dependents': 'No'}
62
63 client.post('/predict', json=payload)          # warm up
64 t = time.time()
65 N = 100
66 for _ in range(N):
67     client.post('/predict', json=payload)
68 per_call = (time.time() - t) / N
69 print('one row at a time: %.1f ms per prediction' % (1000 * per_call))
70
71 import pandas as pd
72 rows = pd.DataFrame([payload] * N)
73 t = time.time()
74 service.MODEL.predict_proba(rows)
75 batched = (time.time() - t) / N
76 print('%d rows in one call: %.3f ms per prediction' % (N, 1000 * batched))
77 print('\nratio %.0fx' % (per_call / batched))

```

```

one row at a time: 5.3 ms per prediction

```

```
100 rows in one call: 0.030 ms per prediction
ratio 176x
```

Almost none of that time is the model

The per-request cost is HTTP parsing, validation, building a one-row dataframe and serialising a response. The arithmetic the model does is a rounding error next to it. Two consequences: offer a batch endpoint if your caller has many rows, because it is very nearly free; and do not optimise the model for latency until you have measured where the latency is.

Day 2 takeaway

Day 3 Validation and Failing Loudly

The bad inputs that raise, and the dangerous one that does not

The interesting failures are not the ones that raise. They are the ones where the service returns 200 and a number that means nothing.

The three kinds of bad input

```
1 SERVICE = '''import json
2 import joblib
3 import pandas as pd
4 from fastapi import FastAPI, HTTPException
5 from pydantic import BaseModel, Field, ValidationError
6
7 MODEL = joblib.load("churn_model.joblib")
8 CONTRACT = json.load(open("churn_contract.json"))
9
10
11 class Customer(BaseModel):
12     """The request schema. Pydantic rejects anything that does not match,
13     before a single row reaches the model."""
14     tenure_months: int = Field(ge=0, le=120)
15     support_calls: int = Field(ge=0, le=100)
16     monthly_charges: float = Field(ge=0, le=500)
17     contract: str
18     internet_service: str
19     payment_method: str
20     has_dependents: str
21
22
23 app = FastAPI(title="churn-scorer", version=CONTRACT["model_version"])
24
25
26 @app.get("/health")
27 def health():
28     return {"status": "ok", "model_version": CONTRACT["model_version"],
```

```

29         "trained_at": CONTRACT["trained_at"]}]
30
31
32 @app.post("/predict")
33 def predict(customer: Customer):
34     row = pd.DataFrame([customer.model_dump()])
35     for column in ("contract", "internet_service", "payment_method"):
36         known = CONTRACT["categories"][column]
37         if row.loc[0, column] not in known:
38             raise HTTPException(
39                 status_code=422,
40                 detail="unknown %s: %r (known: %s)"
41                     % (column, row.loc[0, column], known))
42     p = float(MODEL.predict_proba(row)[0, 1])
43     return {"churn_probability": round(p, 4),
44           "decision": "flag" if p >= CONTRACT["threshold"] else "keep",
45           "threshold": CONTRACT["threshold"],
46           "model_version": CONTRACT["model_version"]}
47 '''
48 import io, sys, json
49 io.open('service.py', 'w', encoding='utf-8').write(SERVICE)
50 sys.path.insert(0, '.')
51
52 from fastapi.testclient import TestClient
53 import service
54 client = TestClient(service.app)
55
56 good = {'tenure_months': 12, 'support_calls': 1,
57        'monthly_charges': 70.0, 'contract': 'One Year',
58        'internet_service': 'DSL',
59        'payment_method': 'Bank transfer', 'has_dependents': 'Yes'}
60
61 cases = [
62     ('valid request', good),
63     ('missing a field', {k: v for k, v in good.items()
64                         if k != 'monthly_charges'}),
65     ('wrong type', dict(good, tenure_months='twelve')),
66     ('out of range', dict(good, tenure_months=-4)),
67     ('unknown category', dict(good, contract='Pay As You Go')),
68 ]
69 for label, payload in cases:
70     r = client.post('/predict', json=payload)
71     body = r.json()
72     detail = body.get('detail', body)
73     if isinstance(detail, list):
74         detail = '%s: %s' % (detail[0].get('loc'), detail[0].get('msg'))
75     print('%-18s %d %s' % (label, r.status_code, str(detail)[:64]))

```

```

valid request      200 {'churn_probability': 0.1575, 'decision': 'keep', 'threshold': 0
missing a field    422 ['body', 'monthly_charges']: Field required
wrong type         422 ['body', 'tenure_months']: Input should be a valid integer, unab
out of range       422 ['body', 'tenure_months']: Input should be greater than or equal
unknown category   422 unknown contract: 'Pay As You Go' (known: ['Month-To-Month', 'On

```

The one that does not raise

A category the model has never seen is the dangerous case, because `OneHotEncoder(handle_unknown='ignore')`, which every week of this course has used, for good reasons, turns it silently into a row of zeros and predicts anyway.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.model_selection import train_test_split
16
17 prep = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                     ('s', StandardScaler())]), NUM),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                     ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22 ])
23 model = Pipeline([('prep', prep),
24                  ('clf', LogisticRegression(max_iter=2000))])
25
26 X, y = df[NUM + CAT], df['churned']
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                         stratify=y, random_state=42)
29 model.fit(X_tr, y_tr)
30 import joblib
31
32 model = joblib.load('churn_model.joblib')
33 base = X_te.iloc[[0]].copy()
34 print('the row as it stands:      %.4f'
35       % model.predict_proba(base)[0, 1])
36
37 for value in ['Two Year', 'Pay As You Go', 'PAYG', '']:
38     row = base.copy()
39     row['contract'] = value
40     print('contract=%-16r %.4f (no error raised)'
41           % (value, model.predict_proba(row)[0, 1]))

```

```

the row as it stands:      0.5306
contract='Two Year'      0.0665 (no error raised)
contract='Pay As You Go' 0.2338 (no error raised)
contract='PAYG'          0.2338 (no error raised)
contract=''              0.2338 (no error raised)

```

⚠ Watch Out

Silent is worse than wrong

Every one of those returned a plausible number. If a new contract type is launched and nobody tells the model's owner, the service keeps answering, scoring every customer on that plan as though their contract field were blank, and nothing in any log looks unusual. That is why the service checks the value against the contract file and returns a 422. Rejecting a request is a bad day; answering it wrongly for eight months is a different kind of day.

Range checks are not the same as schema checks

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.model_selection import train_test_split
16
17 prep = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                       ('s', StandardScaler())])), NUM),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                       ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22 ])
23 model = Pipeline([('prep', prep),
24                   ('clf', LogisticRegression(max_iter=2000))])
25
26 X, y = df[NUM + CAT], df['churned']
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                           stratify=y, random_state=42)
29 model.fit(X_tr, y_tr)
30 import joblib
31 import json
32
33 model = joblib.load('churn_model.joblib')
34 contract = json.load(open('churn_contract.json'))
35 lo, hi = contract['ranges']['tenure_months']
36 print('training range for tenure_months: %.0f to %.0f' % (lo, hi))
37
38 row = X_te.iloc[[0]].copy()
39 for months in [12, 72, 96, 240]:
40     row['tenure_months'] = months
```

```

41 p = model.predict_proba(row)[0, 1]
42 flag = '' if lo <= months <= hi else ' <-- outside training range'
43 print('tenure=%-5d %.4f%s' % (months, p, flag))

training range for tenure_months: 1 to 72
tenure=12    0.5582
tenure=72    0.0429
tenure=96    0.0117 <-- outside training range
tenure=240   0.0000 <-- outside training range

```

The schema accepts anything up to 120 months because that is a plausible tenure. The *model* has never seen beyond 72, so past that it is extrapolating along a straight line in log-odds and there is nothing to stop it. Validate against the schema to reject nonsense; validate against the training ranges to decide whether to trust the answer.

Situation	Status	Return
Field missing or wrong type	422	Which field, and why
Value outside the schema's range	422	The permitted range
Category never seen in training	422	The known values
Inside the schema, outside the training range	200	The prediction, plus a warning flag
Model file missing or corrupt	503	Fail the health check, do not answer
Anything unexpected	500	A request id; log the traceback

Day 3 takeaway

Day 4 Monitoring and Drift

Logging, PSI, category mix and watching your own output move

The model is live and answering. The question now is how you find out it has stopped being any good, given that nobody will tell you.

Log the prediction, not just the answer

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.linear_model import LogisticRegression

```

```

15 from sklearn.model_selection import train_test_split
16
17 prep = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                       ('s', StandardScaler())])), NUM),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                       ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22 ])
23 model = Pipeline([('prep', prep),
24                   ('clf', LogisticRegression(max_iter=2000))])
25
26 X, y = df[NUM + CAT], df['churned']
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                           stratify=y, random_state=42)
29 model.fit(X_tr, y_tr)
30 import hashlib
31 import joblib
32 import json
33
34 model = joblib.load('churn_model.joblib')
35 contract = json.load(open('churn_contract.json'))
36
37 def log_line(request_id, row, proba):
38     return json.dumps({
39         'request_id': request_id,
40         'model_version': contract['model_version'],
41         'threshold': contract['threshold'],
42         'probability': round(float(proba), 4),
43         'decision': 'flag' if proba >= contract['threshold'] else 'keep',
44         # the inputs, so drift is measurable later; hash anything
45         # identifying rather than storing it
46         'features': {k: (round(v, 2) if isinstance(v, float) else v)
47                      for k, v in row.items()},
48     })
49
50 for i in range(3):
51     row = X_te.iloc[i].to_dict()
52     p = model.predict_proba(X_te.iloc[[i]])[0, 1]
53     print(log_line('req-%03d' % i, row, p))

```

```

{"request_id": "req-000", "model_version": "1.0.0", "threshold": 0.35, "probability": 0.5306,
  ↳ "decision": "flag", "features": {"tenure_months": 14, "support_calls": 3,
  ↳ "monthly_charges": 73.35, "contract": "Month-To-Month", "internet_service": "Fibre
  ↳ optic", "payment_method": "Credit card", "has_dependents": "No"}}
{"request_id": "req-001", "model_version": "1.0.0", "threshold": 0.35, "probability": 0.1651,
  ↳ "decision": "keep", "features": {"tenure_months": 26, "support_calls": 1,
  ↳ "monthly_charges": 15.0, "contract": "Month-To-Month", "internet_service": "No
  ↳ internet", "payment_method": "Electronic check", "has_dependents": "No"}}
{"request_id": "req-002", "model_version": "1.0.0", "threshold": 0.35, "probability": 0.3226,
  ↳ "decision": "keep", "features": {"tenure_months": 14, "support_calls": 1,
  ↳ "monthly_charges": 53.11, "contract": "Month-To-Month", "internet_service": "DSL",
  ↳ "payment_method": "Bank transfer", "has_dependents": "No"}}

```

Log the features or you cannot diagnose anything

A log of probabilities tells you the distribution of your output moved. A log of probabilities *and inputs* tells you which input moved and therefore what happened in the business. The second costs a few hundred bytes per request and is the difference between "something changed in March" and "fibre pricing changed in March".

Data drift

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.model_selection import train_test_split
16
17 prep = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                      ('s', StandardScaler())]), NUM),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                      ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22 ])
23 model = Pipeline([('prep', prep),
24                  ('clf', LogisticRegression(max_iter=2000))])
25
26 X, y = df[NUM + CAT], df['churned']
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                           stratify=y, random_state=42)
29 model.fit(X_tr, y_tr)
30 import numpy as np
31
32 def psi(reference, live, buckets=10):
33     edges = np.quantile(reference, np.linspace(0, 1, buckets + 1))
34     edges[0], edges[-1] = -np.inf, np.inf
35     a = np.histogram(reference, edges)[0] / len(reference)
36     b = np.histogram(live, edges)[0] / len(live)
37     a, b = np.clip(a, 1e-6, None), np.clip(b, 1e-6, None)
38     return float(((b - a) * np.log(b / a)).sum())
39
40 rng = np.random.default_rng(0)
41 reference = X_tr['monthly_charges'].dropna().to_numpy()
42
43 print('%-28s %8s %s' % ('live traffic', 'PSI', 'verdict'))
44 for label, live in [

```

```

45     ('same distribution', rng.choice(reference, 800)),
46     ('prices up 5%', rng.choice(reference, 800) * 1.05),
47     ('prices up 20%', rng.choice(reference, 800) * 1.20),
48     ('only premium customers',
49      rng.choice(reference[reference > np.median(reference)], 800))]:
50 value = psi(reference, live)
51 verdict = ('stable' if value < 0.1 else
52            'investigate' if value < 0.25 else 'material shift')
53 print('%-28s %8.4f %s' % (label, value, verdict))

```

live traffic	PSI verdict
same distribution	0.0180 stable
prices up 5%	0.0688 stable
prices up 20%	0.5386 material shift
only premium customers	6.1024 material shift

Categorical drift, and drift in the output itself

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.model_selection import train_test_split
16
17 prep = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                      ('s', StandardScaler())])), NUM),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                      ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22 ])
23 model = Pipeline([('prep', prep),
24                  ('clf', LogisticRegression(max_iter=2000))])
25
26 X, y = df[NUM + CAT], df['churned']
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                           stratify=y, random_state=42)
29 model.fit(X_tr, y_tr)
30 import numpy as np
31 import joblib
32
33 model = joblib.load('churn_model.joblib')
34 rng = np.random.default_rng(0)

```

```

35
36 train_mix = X_tr['contract'].value_counts(normalize=True)
37 print('training contract mix:')
38 print(train_mix.round(3).to_string())
39
40 # A marketing push sells monthly plans to everybody.
41 shifted = X_te.copy()
42 flip = rng.random(len(shifted)) < 0.5
43 shifted.loc[flip, 'contract'] = 'Month-To-Month'
44
45 live_mix = shifted['contract'].value_counts(normalize=True)
46 print('\nlive contract mix:')
47 print(live_mix.round(3).to_string())
48
49 before = model.predict_proba(X_te)[: , 1]
50 after = model.predict_proba(shifted)[: , 1]
51 print('\nmean predicted risk %.4f → %.4f' % (before.mean(), after.mean()))
52 print('flagged at 0.35 %.1f%% → %.1f%%'
53       % (100 * (before >= 0.35).mean(), 100 * (after >= 0.35).mean()))

```

```

training contract mix:
contract
Month-To-Month    0.553
One Year          0.243
Two Year          0.204

live contract mix:
contract
Month-To-Month    0.759
One Year          0.132
Two Year          0.109

mean predicted risk 0.2628 → 0.3005
flagged at 0.35    33.5% → 39.9%

```

Prediction drift is the cheapest alarm you can build, because it needs no labels and no feature analysis, just the mean of today's scores against the mean of last month's. It cannot tell you whether the model is still *right*. It can tell you that something changed, on the day it changes, which is usually months before the labels arrive.

Day 4 takeaway

Day 5 Delayed Labels

Why you cannot measure accuracy yet, and what to watch instead

Drift says the inputs changed. It does not say the model got worse, for that you need labels, and labels are the thing you do not have.

The delay is the whole problem

```

1 events = [
2     ('Monday', 'model scores 4,000 customers'),
3     ('Monday', 'retention team contacts the 600 flagged'),
4     ('+30 days', 'first contracts come up for renewal'),
5     ('+60 days', 'roughly half of the outcomes are known'),
6     ('+90 days', 'the label for Monday is finally complete'),
7 ]
8 for when, what in events:
9     print('%-10s %s' % (when, what))
10 print('\nAny accuracy you measure today describes the model as it was')
11 print('three months ago -- on data that has since drifted.')

```

```

Monday      model scores 4,000 customers
Monday      retention team contacts the 600 flagged
+30 days    first contracts come up for renewal
+60 days    roughly half of the outcomes are known
+90 days    the label for Monday is finally complete

Any accuracy you measure today describes the model as it was
three months ago -- on data that has since drifted.

```

⚠ Watch Out

And the intervention destroys the label

Worse than the delay: the retention team acted on the predictions. A customer who was flagged, offered a discount and stayed now counts as a false positive, so a model that worked perfectly and prompted a successful intervention looks, in the data, like a model that was wrong. Measuring a model whose output changes the world requires a holdout that nobody acts on, agreed with the business before you deploy. It is a conversation, not a technique.

What to watch while you wait

Signal	Available	Tells you
Prediction distribution	Immediately	Something changed
Feature PSI	Immediately	Which input changed
Unknown-category rejection rate	Immediately	A new product or a broken up-stream field
Missing-value rate per field	Immediately	An upstream pipeline has broken
Latency and error rate	Immediately	The service itself
Proxy outcomes (support calls, logins)	Days	Early behavioural signal
Actual labels	Weeks to months	Whether it was any good
Holdout labels	Weeks to months	Whether it was any good, uncontaminated

An alert that does not cry wolf

```

1 import numpy as np
2 import pandas as pd

```

```

3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.model_selection import train_test_split
16
17 prep = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                     ('s', StandardScaler())])), NUM),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                     ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22 ])
23 model = Pipeline([('prep', prep),
24                  ('clf', LogisticRegression(max_iter=2000))])
25
26 X, y = df[NUM + CAT], df['churned']
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                         stratify=y, random_state=42)
29 model.fit(X_tr, y_tr)
30 import numpy as np
31 import joblib
32
33 model = joblib.load('churn_model.joblib')
34 rng = np.random.default_rng(1)
35
36 baseline = model.predict_proba(X_tr)[: , 1]
37 mu, sd = baseline.mean(), baseline.std()
38 print('baseline mean risk %.4f (sd %.4f)' % (mu, sd))
39
40 def day_of_traffic(shift):
41     sample = X_te.sample(400, replace=True, random_state=int(shift * 100))
42     sample = sample.copy()
43     sample['monthly_charges'] = sample['monthly_charges'] * (1 + shift)
44     return model.predict_proba(sample)[: , 1]
45
46 print('\n%-8s %12s %12s %s' % ('shift', 'mean risk', 'z-score', 'alert'))
47 for shift in [0.0, 0.05, 0.10, 0.25, 0.50]:
48     scores = day_of_traffic(shift)
49     z = (scores.mean() - mu) / (sd / np.sqrt(len(scores)))
50     print('%-8.0f%% %12.4f %12.1f %s'
51           % (100 * shift, scores.mean(), z, 'ALERT' if abs(z) > 4 else ''))

```

```
baseline mean risk 0.2680 (sd 0.2247)
```

shift	mean risk	z-score	alert
0 %	0.2654	-0.2	

5	%	0.2696	0.1
10	%	0.2834	1.4
25	%	0.3080	3.6
50	%	0.3490	7.2 ALERT

A threshold on the raw mean would fire on every quiet Tuesday. Comparing against the variability of the baseline, how far today is from normal, in units of how much normal wobbles, is what makes an alert survive contact with a real on-call rota. Set the bar high enough that firing means something, because an alert nobody trusts is worse than no alert at all.

Day 5 takeaway

Day 6 Retraining and Rollout

Champion versus challenger with an interval, shadow mode and canaries

Something has drifted, or three months have passed. Retraining is straightforward; deciding whether the new model is better, and getting it into production without a bad afternoon, is not.

Champion and challenger

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.model_selection import train_test_split
16
17 prep = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                     ('s', StandardScaler())]), NUM),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                     ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22 ])
23 model = Pipeline([('prep', prep),
24                  ('clf', LogisticRegression(max_iter=2000))])
25
26 X, y = df[NUM + CAT], df['churned']
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                           stratify=y, random_state=42)
29 model.fit(X_tr, y_tr)
30 import numpy as np

```

```

31 import joblib
32 from sklearn.ensemble import HistGradientBoostingClassifier
33 from sklearn.metrics import roc_auc_score, brier_score_loss
34
35 champion = joblib.load('churn_model.joblib')
36
37 challenger = Pipeline([('prep', prep),
38                        ('clf', HistGradientBoostingClassifier(
39                            random_state=42))]).fit(X_tr, y_tr)
40
41 print('%-14s %10s %10s %12s' % ('', 'AUC', 'Brier', 'flagged %'))
42 for name, m in [('champion', champion), ('challenger', challenger)]:
43     p = m.predict_proba(X_te)[: , 1]
44     print('%-14s %10.4f %10.4f %11.1f%%'
45           % (name, roc_auc_score(y_te, p), brier_score_loss(y_te, p),
46              100 * (p >= 0.35).mean()))
47
48 # Is the difference bigger than the noise? Bootstrap the test set.
49 rng = np.random.default_rng(0)
50 pc = champion.predict_proba(X_te)[: , 1]
51 pch = challenger.predict_proba(X_te)[: , 1]
52 yv = y_te.to_numpy()
53 diffs = []
54 for _ in range(400):
55     idx = rng.integers(0, len(yv), len(yv))
56     if yv[idx].sum() in (0, len(idx)):
57         continue
58     diffs.append(roc_auc_score(yv[idx], pch[idx])
59                 - roc_auc_score(yv[idx], pc[idx]))
60 lo, hi = np.percentile(diffs, [2.5, 97.5])
61 print('\nAUC difference (challenger - champion)')
62 print(' point estimate %+.4f' % (roc_auc_score(yv, pch)
63                                   - roc_auc_score(yv, pc)))
64 print(' 95% interval %+.4f to %+.4f' % (lo, hi))

```

	AUC	Brier	flagged %
champion	0.8174	0.1475	33.5%
challenger	0.7797	0.1690	30.3%


```

AUC difference (challenger - champion)
point estimate -0.0376
95% interval -0.0582 to -0.0141

```

The challenger loses, and the interval is nowhere near zero, which settles it. This is week 6's result arriving in an operational costume: the data-generating process behind this dataset is close to logistic-linear, so a boosted ensemble has nothing to find that the linear model has not already found, and pays for the search with variance. The correct action here is to keep the champion and write down that the challenger was tried.

"Better" needs an interval, not a point

A challenger that is 0.004 ahead on one test set is not obviously ahead at all. Bootstrap the difference and look at whether the interval clears zero. If it does not, the honest conclusion is that

the two models are indistinguishable on the evidence available, at which point you should prefer the one that is simpler, faster, already deployed, or easier to explain. "No change" is a legitimate and frequently correct outcome of a retraining cycle.

Getting it out there

Strategy	How	Catches	Costs
Shadow	New model scores every request; output logged, not used	Crashes, latency, wild predictions	Compute, and no outcome data
Canary	5% of traffic, then 25%, then all	Real-world failures, at 5% of the blast radius	A slower rollout
A/B split	Randomised, both act on their predictions	Genuine causal comparison	Weeks, and real design work
Blue/green	Both live, switch the router	Nothing extra, but roll-back is instant	Two environments
Big bang	Replace it	Nothing	Your evening

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.model_selection import train_test_split
16
17 prep = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                       ('s', StandardScaler())])), NUM),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                       ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
22 ])
23 model = Pipeline([('prep', prep),
24                   ('clf', LogisticRegression(max_iter=2000))])
25
26 X, y = df[NUM + CAT], df['churned']
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                           stratify=y, random_state=42)
29 model.fit(X_tr, y_tr)
30 import numpy as np
31 import joblib
32 from sklearn.ensemble import HistGradientBoostingClassifier
33
34 champion = joblib.load('churn_model.joblib')

```

```

35 challenger = Pipeline([('prep', prep),
36                        ('clf', HistGradientBoostingClassifier(
37                            random_state=42))]).fit(X_tr, y_tr)
38
39 # Shadow mode: the champion decides, the challenger is recorded.
40 a = champion.predict_proba(X_te)[: , 1]
41 b = challenger.predict_proba(X_te)[: , 1]
42 flag_a, flag_b = a >= 0.35, b >= 0.35
43
44 print('champion flags    %d' % flag_a.sum())
45 print('challenger flags %d' % flag_b.sum())
46 print('they disagree on %d of %d customers (0.1f%%)'
47       % ((flag_a != flag_b).sum(), len(a),
48          100 * (flag_a != flag_b).mean()))
49 print('\nlargest disagreements:')
50 gap = np.abs(a - b)
51 for i in np.argsort(-gap)[:4]:
52     print(' champion 0.3f challenger 0.3f tenure=%d contract=%s'
53           % (a[i], b[i], X_te.iloc[i]['tenure_months'],
54              X_te.iloc[i]['contract']))

```

```

champion flags    251
challenger flags  227
they disagree on 112 of 750 customers (14.9%)

largest disagreements:
 champion 0.131 challenger 0.665 tenure=33 contract=Month-To-Month
 champion 0.392 challenger 0.887 tenure=23 contract=Month-To-Month
 champion 0.466 challenger 0.944 tenure=22 contract=Month-To-Month
 champion 0.430 challenger 0.889 tenure=22 contract=One Year

```

Before any traffic moves, this tells you how different the two models actually are in the only terms that matter: how many decisions would change. A challenger that is 0.01 better on AUC but flips a third of your decisions is a much bigger operational event than that number suggests, and the retention team needs to hear about it from you rather than from their workload.

When to retrain

1. **On a schedule**, if the world moves steadily. Monthly is a common default; it is a decision, not a law.
2. **On drift**, when PSI on an important feature crosses your threshold and stays there.
3. **On performance**, when holdout labels show a real fall, the most reliable trigger and the slowest.
4. **On a known event**: a new product, a pricing change, a changed upstream field. Do not wait for the monitoring to notice something you already know.
5. **Never automatically without a gate**. Automatic retraining plus automatic deployment means a corrupted upstream table becomes a corrupted production model with nobody in the loop.

Day 6 takeaway

Day 7 The Complete Service

Everything assembled, and the checklist before it takes traffic

The complete service, with everything the previous six days argued for, and a checklist to work through before anything of yours takes real traffic.

The finished thing

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.compose import ColumnTransformer
11 from sklearn.pipeline import Pipeline
12 from sklearn.impute import SimpleImputer
13 from sklearn.preprocessing import StandardScaler, OneHotEncoder
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.model_selection import train_test_split
16
17 prep = ColumnTransformer([
18     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
19                       ('s', StandardScaler())]), NUM),
20     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
21                       ('o', OneHotEncoder(handle_unknown='ignore'))]), CAT),
22 ])
23 model = Pipeline([('prep', prep),
24                   ('clf', LogisticRegression(max_iter=2000))])
25
26 X, y = df[NUM + CAT], df['churned']
27 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
28                                           stratify=y, random_state=42)
29 model.fit(X_tr, y_tr)
30 import json
31 import joblib
32 import numpy as np
33
34 class ChurnService:
35     """Everything the six days argued for, in one object."""
36
37     def __init__(self, model_path, contract_path):
38         self.model = joblib.load(model_path)
39         self.contract = json.load(open(contract_path))
40         self.seen = 0
```

```

41     self.rejected = 0
42     self.scores = []
43
44     def _validate(self, row):
45         for column, known in self.contract['categories'].items():
46             if row.get(column) not in known:
47                 return 'unknown %s: %r' % (column, row.get(column))
48         for column, (lo, hi) in self.contract['ranges'].items():
49             value = row.get(column)
50             if value is None or not np.isfinite(value):
51                 return 'missing %s' % column
52         return None
53
54     def predict(self, row):
55         self.seen += 1
56         problem = self._validate(row)
57         if problem:
58             self.rejected += 1
59             return {'status': 'rejected', 'reason': problem}
60         p = float(self.model.predict_proba(pd.DataFrame([row]))[0, 1])
61         self.scores.append(p)
62         outside = [c for c, (lo, hi) in self.contract['ranges'].items()
63                   if not lo <= row[c] <= hi]
64         return {'status': 'ok',
65                 'probability': round(p, 4),
66                 'decision': 'flag' if p >= self.contract['threshold']
67                             else 'keep',
68                 'model_version': self.contract['model_version'],
69                 'extrapolating': outside or None}
70
71     def stats(self):
72         s = np.array(self.scores) if self.scores else np.array([0.0])
73         return {'requests': self.seen, 'rejected': self.rejected,
74                 'mean_score': round(float(s.mean()), 4),
75                 'flag_rate': round(float(
76                     (s >= self.contract['threshold']).mean()), 4)}
77
78 svc = ChurnService('churn_model.joblib', 'churn_contract.json')
79 for row in X_te.head(200).to_dict('records'):
80     svc.predict(row)
81
82 print(json.dumps(svc.predict(X_te.iloc[0].to_dict()), indent=2))
83 print(json.dumps(svc.predict(dict(X_te.iloc[0].to_dict(),
84                                   contract='Pay As You Go')), indent=2))
85 print(json.dumps(svc.predict(dict(X_te.iloc[0].to_dict(),
86                                   tenure_months=200)), indent=2))
87 print('\n' + json.dumps(svc.stats(), indent=2))

```

```

{
  "status": "ok",
  "probability": 0.5306,
  "decision": "flag",
  "model_version": "1.0.0",
  "extrapolating": null
}

```

```

}
{
  "status": "rejected",
  "reason": "unknown contract: 'Pay As You Go'"
}
{
  "status": "ok",
  "probability": 0.0,
  "decision": "keep",
  "model_version": "1.0.0",
  "extrapolating": [
    "tenure_months"
  ]
}

{
  "requests": 203,
  "rejected": 28,
  "mean_score": 0.2558,
  "flag_rate": 0.32
}

```

Look at the rejection count: 28 of 203 real customers from the test set were refused. They are the rows with a missing categorical value, the roughly six percent that `customers.csv` has carried since week 1. The training pipeline imputes those without complaining. The service refuses them.

Watch Out

Your service can be stricter than your model, and you must choose

Neither behaviour is wrong, but the disagreement has to be a decision rather than an accident. Rejecting is right when a missing field means an upstream system is broken and a prediction would be misleading. Imputing is right when the field is genuinely optional and the model was trained to cope. What is never right is discovering the difference in production, when one in twenty of your customers cannot be scored and nobody knows why.

Write the rule down in the contract file next to the categories, and test it. Here, the honest fix is to allow missing categoricals (the pipeline handles them, and the training data contained them) while continuing to reject unknown *values*, which the pipeline silently mishandles.

The checklist

1. The artefact contains the preprocessing, not just the estimator.
2. A contract file records the features, valid categories, training ranges, threshold, training date and versions.
3. A smoke test with frozen inputs and expected outputs runs on every deploy.
4. Dependencies are pinned and the service runs in the environment it was trained in.
5. The model loads once at start-up, and a health endpoint reports its version.
6. Unknown categories are rejected rather than silently zero-encoded.
7. Inputs outside the training ranges are answered but flagged.

8. Every prediction is logged with its inputs, the threshold and the model version.
9. Feature PSI, category mix, missing-value rates and the prediction distribution are monitored, with thresholds set from observed variation.
10. An unactioned holdout exists, agreed with the business, so performance can be measured at all.
11. Rollback is a switch somebody can throw without a rebuild.
12. Somebody's name is against the model, and there is a date to review it.

The part that is not code

Most models that fail in production fail for reasons on that list that have nothing to do with modelling: an upstream field silently changed units, a new product category appeared, nobody agreed who watches the dashboard, the threshold was set by a data scientist and never discussed with the team who work the flagged list. The modelling was the easy part and it was finished twelve weeks ago.

Day 7 takeaway

Self-Assessment

Try It Yourself

1. A deployable model artefact must contain:
 - a) The training data
 - b) The notebook
 - c) The estimator only
 - d) The whole pipeline, including every preprocessing step
2. The contract file exists to record:
 - a) Valid categories, training ranges, the agreed threshold, the version and the training date
 - b) The cross-validation scores
 - c) The feature importances
 - d) The model weights
3. A smoke test with frozen inputs and expected outputs catches:
 - a) Slow inference
 - b) A dependency upgrade that silently changes your predictions
 - c) Missing data
 - d) Class imbalance
4. An unknown category arriving at a service with `handle_unknown='ignore'` is dangerous because it:
 - a) Is logged as invalid

- b) Raises an error
 - c) Is silently encoded as zeros and a plausible prediction is returned
 - d) Crashes the service
5. A Population Stability Index above 0.25 indicates:
- a) A broken calculation
 - b) Improved accuracy
 - c) A stable feature
 - d) A material shift in the feature's distribution
6. Prediction drift is valuable to monitor because it:
- a) Needs no labels and is available the same day
 - b) Replaces cross-validation
 - c) Detects leakage
 - d) Measures accuracy
7. Measuring a churn model's accuracy is complicated by the retention team acting on its output, because:
- a) The team is slow
 - b) A correctly flagged customer who was persuaded to stay now counts as a false positive
 - c) Labels are noisy
 - d) The threshold changes
8. Before promoting a challenger model you should compare:
- a) The number of parameters
 - b) Point estimates of AUC
 - c) An interval on the difference, plus how many decisions would actually change
 - d) Training time

Answers: 1d, 2a, 3b, 4c, 5d, 6a, 7b, 8c. Anything you got wrong points at the day to re-read, not at a gap in ability.

Capstone and Career

Week 16

Frame a project of your own, build it to a standard somebody else can run, and be able to talk about it.

OBJECTIVES

The last week has the least new machinery in it and the most that determines whether the previous fifteen were worth anything. It covers framing a problem around a decision rather than a dataset, structuring a project somebody else can run, four milestones for the capstone itself, how to mark your own work before a reviewer does, and how to talk about what you built.

Day 1 Framing a Project

Starting from a decision instead of a dataset, and writing the brief

Fifteen weeks of technique. This one is about turning it into something with your name on it, which means starting where every real project starts, and where almost every portfolio project does not: with a decision somebody has to make.

Frame the decision, not the dataset

Watch Out

The most common way a capstone fails

It begins "I found an interesting dataset". Everything after that is a search for a question the data happens to answer, and the result reads like it, a competent notebook with no argument in it. Start instead with a decision that is currently made badly, then ask what data would inform it. You may well end up at the same dataset. You will end up with a different project.

Weak framing	Strong framing	Why
"Predict house prices"	"Which listings are priced more than 10% below what comparable homes achieved, so an agent should call the vendor this week"	Names who acts, when, and on what
"Analyse customer churn"	"Which 200 customers should the retention team ring in October, given they can make 200 calls"	The budget constraint decides the threshold
"Classify images of plants"	"Let a grower photograph a leaf and get a disease shortlist without waiting for the agronomist"	Establishes what the alternative is
"Sentiment analysis of reviews"	"Route the 5% of reviews that describe a safety problem to the product team within an hour"	Makes recall on a rare class the metric

Write the brief before the code

1. **The decision.** Who does what differently because of this model?
2. **The alternative.** How is it decided today? That is your baseline, and it is rarely a machine learning model, usually it is a rule, a spreadsheet, or somebody's judgement.
3. **The unit.** One row is one what? Get this wrong and every split in the project is wrong.
4. **The label.** What exactly are you predicting, measured when? Write the SQL or the sentence that defines it.
5. **The moment of prediction.** What is known at that instant? Anything not known then cannot be a feature. This one sentence prevents most leakage.
6. **The metric, and the threshold.** Which errors cost what? Who bears them?
7. **What good enough looks like.** A number, agreed before you start, that means this was worth doing.

Point five is the whole of week 4 in one sentence

"What is known at the moment the prediction is made" resolves almost every leakage question without any statistics. `total_charges` is known at prediction time; `reason_for_leaving` is not. A customer's tenure is known; whether they answered the retention call is not. Write the list, then check every column you add against it.

Scoping it to fit

Scale	Rows	Looks like	Fits in
Small	< 10k	One table, a few features, a clear label	A weekend
Medium	10k to 1M	Two or three joined tables, some feature engineering	Two to three weeks
Large	> 1M	Sampling, chunked reads, a real time split	A month, and it will still surprise you

For a capstone, medium is the target. Small does not exercise enough of the course; large means you spend three weeks on data plumbing and produce a rushed model. If the dataset you

want is large, take a principled sample and say so in the write-up. That is a legitimate engineering decision, not a compromise.

Day 1 takeaway

Day 2 Structure and Reproducibility

Where the code lives, the four tests to write, and running it twice

A project somebody else can run is worth several times one they cannot, and the difference is about an hour of work done at the start rather than at the end.

The layout

The rule that makes the rest work

Notebooks are for narrative; `src/` is for anything that runs twice. The moment a cell is useful enough to copy into another notebook, it belongs in a module, where it can be imported, tested and fixed once. A project whose logic lives in notebook cells cannot be tested and cannot be deployed, and both of those are noticed immediately by anybody reviewing it.

Testing the parts that can silently break

```
1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 def add_features(data):
11     """The kind of function that belongs in src/features.py."""
12     out = data.copy()
13     out['charges_per_month_of_tenure'] = (
14         out['total_charges'] / out['tenure_months'].clip(lower=1))
15     out['calls_per_year'] = (
16         out['support_calls'] * 12 / out['tenure_months'].clip(lower=1))
17     out['is_new'] = (out['tenure_months'] <= 3).astype(int)
18     return out
19
20 # tests/test_features.py -- these are the assertions that catch the
21 # mistakes nobody notices: a divide by zero, a column that silently
22 # disappears, a transformation that changes the number of rows.
23 def test_no_rows_lost():
24     assert len(add_features(df)) == len(df)
25
26 def test_no_division_blowup():
27     out = add_features(df)
```

```

28     assert np.isfinite(out['calls_per_year'].dropna()).all()
29
30 def test_zero_tenure_is_survivable():
31     edge = df.head(1).copy()
32     edge['tenure_months'] = 0
33     assert np.isfinite(add_features(edge)['calls_per_year']).all()
34
35 def test_original_columns_survive():
36     out = add_features(df)
37     assert set(df.columns).issubset(out.columns)
38
39 for test in [test_no_rows_lost, test_no_division_blowup,
40             test_zero_tenure_is_survivable, test_original_columns_survive]:
41     test()
42     print('PASS %s' % test.__name__)

```

```

PASS test_no_rows_lost
PASS test_no_division_blowup
PASS test_zero_tenure_is_survivable
PASS test_original_columns_survive

```

Four assertions, no framework, about five minutes of work. They do not check that your model is good. Nothing can, but they catch the class of bug that produces a plausible model trained on quietly corrupted features, which is the class of bug that is hardest to find later.

Reproducibility

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 import hashlib
11 import json
12
13 def fingerprint(data):
14     """Enough to prove two runs saw the same data."""
15     digest = hashlib.sha256(
16         pd.util.hash_pandas_object(data, index=True).values).hexdigest()
17     return {'rows': len(data), 'columns': list(data.columns),
18           'sha256': digest[:16]}
19
20 print(json.dumps(fingerprint(df), indent=2)[:300])
21 print('\nsame data again:', fingerprint(df)['sha256'])
22 print('one value changed:',
23       fingerprint(df.assign(monthly_charges=df['monthly_charges'] + 0.01))
24       ['sha256'])

```

{

```

"rows": 3000,
"columns": [
  "customer_id",
  "signup_date",
  "tenure_months",
  "contract",
  "internet_service",
  "payment_method",
  "has_dependents",
  "support_calls",
  "monthly_charges",
  "total_charges",
  "churned"
],
"sha256": "48cf8b8b0d435748"
}

same data again: 48cf8b8b0d435748
one value changed: 0168e3b07660fcf9

```

1. Set every seed, and record it.
2. Fingerprint the input data so a rerun can prove it saw the same rows.
3. Pin dependencies.
4. Make the whole thing runnable from one command, a `Makefile`, or `python -m src.train`.
5. Never edit raw data in place. Read it, transform it, write elsewhere.
6. Keep `data/` out of version control, and say in the README where the data comes from.

Day 2 takeaway

Day 3 Building the Capstone

Four milestones with artefacts, and where the time really goes

The capstone itself, as a sequence of checkpoints. Each has an artefact and a decision, so it is possible to know whether you have finished it.

Milestone 1: the brief and the baseline

- The brief from day 1, written out, one page.
- The data loaded, the unit confirmed, the label defined in code.
- A split chosen and justified, random, grouped or by time.
- **A baseline that is not a model:** the current rule, the majority class, last month's value, or a two-line heuristic.
- The metric computed for that baseline.

The baseline is the most valuable number in the project

Without it, "AUC 0.81" means nothing. With it, "AUC 0.81 against the current rule's 0.68, on customers the rule cannot score at all" is a result. The baseline is also how you find out early that the problem is already solved, which is worth knowing in week one rather than week four.

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split
11 from sklearn.metrics import roc_auc_score, f1_score
12
13 X, y = df[NUM + CAT], df['churned']
14 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
15                                         stratify=y, random_state=42)
16
17 # The rule the retention team uses today: month-to-month contract and
18 # more than two support calls.
19 rule = ((X_te['contract'] == 'Month-To-Month')
20         & (X_te['support_calls'] > 2)).astype(int)
21
22 print('the existing rule')
23 print('  flags          %d of %d customers' % (rule.sum(), len(rule)))
24 print('  AUC           %.4f' % roc_auc_score(y_te, rule))
25 print('  F1            %.4f' % f1_score(y_te, rule))
26 print('  of those flagged, %.1f%% actually churned'
27       % (100 * y_te[rule == 1].mean()))
28 print('  of all churners, it catches %.1f%%'
29       % (100 * rule[y_te == 1].mean()))

```

```

the existing rule
  flags          51 of 750 customers
  AUC           0.5351
  F1            0.1905
  of those flagged, 47.1% actually churned
  of all churners, it catches 11.9%

```

Milestone 2: honest evaluation

- A pipeline, with every transformation inside it.
- Cross-validation appropriate to the data, from week 7.
- Two or three model families compared, including a linear one.
- A held-out set touched exactly once.
- The threshold chosen from the cost of the two errors, not left at 0.5.

Milestone 3: improvement, measured

- Features added one group at a time, each measured against the baseline.
- **The failures kept in the write-up.** "I tried target encoding on payment method; it cost 0.004 AUC and I dropped it" is evidence of judgement.
- Tuning, with the budget argument from week 14.
- A final held-out number, reported once.

Milestone 4: explain and ship

- Which features drive it, with the shape of the top effects.
- Two or three individual predictions explained, including a wrong one.
- An error analysis by segment, who does it fail?
- The artefact and contract from week 15.
- A README somebody can follow without asking you anything.

Where the time actually goes

```
1 phases = [('Framing and data understanding', 20),
2           ('Cleaning and joining', 30),
3           ('Feature engineering', 20),
4           ('Modelling and tuning', 15),
5           ('Evaluation and error analysis', 10),
6           ('Writing it up', 5)]
7 for name, pct in phases:
8     print('%-34s %3d% %s' % (name, pct, '#' * (pct // 2)))
9 print('\nthe part everyone plans for is the 15%.')
```

```
Framing and data understanding      20% #####
Cleaning and joining                30% #####
Feature engineering                 20% #####
Modelling and tuning                15% #####
Evaluation and error analysis        10% #####
Writing it up                      5%  ##
```

```
the part everyone plans for is the 15%.
```

Watch Out

Budget for the write-up before you need it

Five percent is the honest average and it is too little. A project that is finished but unexplained is worth roughly nothing to a reader, and the write-up is the only part of it most people will ever see. Book the last day of your capstone for writing, and treat it as non-negotiable.

Day 3 takeaway**Day 4 Marking Your Own Work**

The rubric, the self-check, and the one-page write-up

How to mark your own work, and the specific things a reviewer looks for first.

The rubric

Area	Falls short	Meets	Exceeds
Framing	A dataset with a model on it	A stated decision and user	The decision, the current alternative, and the cost of each error
Data	<code>dropna()</code> and onwards	Missingness, duplicates and outliers handled deliberately	Each choice justified and its effect measured
Validation	One random split	Cross-validation matched to the data's structure	Grouped or time-based splitting, argued from how the data arose
Leakage	Not considered	Transformations inside the pipeline	An explicit list of what is known at prediction time
Baseline	None	A sensible simple model	The actual current process, measured
Metric	Accuracy	A metric matched to the problem	A threshold chosen from costs, with the trade-off shown
Interpretation	A feature importance bar chart	Importances plus the shape of the main effects	Individual cases explained, including errors, with limits stated
Engineering	One long notebook	Modules, seeds, pinned versions	Tests, one-command run, a saved artefact and contract
Communication	Cells and outputs	A README and a clear findings section	Written for the decision-maker, with the caveats in it

Marking yourself honestly

```

1 checks = [
2     ('Could a reader state the decision this informs?', None),
3     ('Is there a non-model baseline, measured?', None),
4     ('Is every transformation inside the pipeline?', None),
5     ('Was the held-out set used exactly once?', None),
6     ('Is the threshold justified by costs?', None),
7     ('Are failed experiments written down?', None),
8     ('Does the write-up state who the model fails?', None),
9     ('Can somebody else run it from the README?', None),
10    ('Are the seeds set and the versions pinned?', None),
11    ('Is there a sentence about what would make it wrong?', None),

```

```

12 ]
13 print('Score one point each. Below 7 and the reviewer will find it')
14 print('before you do.\n')
15 for i, (question, _) in enumerate(checks, 1):
16     print('%2d. [ ] %s' % (i, question))

```

```

Score one point each. Below 7 and the reviewer will find it
before you do.

```

```

1. [ ] Could a reader state the decision this informs?
2. [ ] Is there a non-model baseline, measured?
3. [ ] Is every transformation inside the pipeline?
4. [ ] Was the held-out set used exactly once?
5. [ ] Is the threshold justified by costs?
6. [ ] Are failed experiments written down?
7. [ ] Does the write-up state who the model fails?
8. [ ] Can somebody else run it from the README?
9. [ ] Are the seeds set and the versions pinned?
10. [ ] Is there a sentence about what would make it wrong?

```

The three things looked at first

From experience of reading these: **the README**, can I tell in thirty seconds what this is and what it found; **the validation**, is the reported number believable, or is it leaking; and **the baseline**, is the improvement over anything. A project that is strong on those three survives being weak elsewhere. A project that is weak on them does not recover, however good the modelling is.

The write-up

1. **The problem**, in two sentences, naming the decision.
2. **The data**: where from, how much, what is wrong with it.
3. **The approach**: the split, the baseline, what you tried.
4. **The result**: one headline number against the baseline, with the metric named and the uncertainty given.
5. **What it means**: for the decision in point one, in the reader's language.
6. **Where it fails**: which segments, which inputs, what would break it.
7. **What you would do next** with another two weeks.

Seven sections, one page. If it runs longer than a page, the model is probably being described where the finding should be. Nobody outside the team needs your hyperparameters; everybody needs to know whether to trust the number and what to do with it.

Day 4 takeaway

Day 5 The Portfolio

Three or four projects, and a README that leads with the result

What to do with sixteen weeks of work once it exists: the portfolio, and the specific ways it fails to land.

Three or four projects, not ten

Project	Shows	Drawn from
A tabular prediction, end to end	Judgement: framing, leakage, base-lines, thresholds	Weeks 1 to 8, 14
A deployed model with a live end-point	That you have shipped something	Week 15
A perceptual project, images or text	That you can work where features cannot be written by hand	Weeks 12 to 13
One deep analysis with a surprising finding	Curiosity, and that you can write	Weeks 2, 7, 9 to 10

Watch Out

The portfolio mistakes that are easy to avoid

The famous datasets, Titanic, Iris, Boston housing, have been done ten thousand times and cannot distinguish you; use them to learn and not to demonstrate. A repository of notebooks with no README is unreadable. A 0.99 accuracy is read as leakage, not as success, so if you genuinely have one, explain immediately why it is real. And a project with no failures in it reads as a project where nothing was tried.

The README that gets read

Result before instructions, instructions before detail, limitations stated rather than buried. A reader who stops after the first two sections still knows what the project is and whether it worked, which is the only realistic goal. Those are the numbers this course actually measured, day 3 fitted the rule, day 6 ranked under a capacity constraint, so the example is a real README for the work you have already done, not a template with invented figures in it.

Talking about it

1. Lead with the decision and the baseline, never with the algorithm.
2. Give the headline number with its comparison attached, a number on its own invites the question you should have answered.
3. Have one specific difficulty ready, and how you resolved it. This is the question that separates people who did the work from people who followed a tutorial.
4. Know why you chose the model you chose, and what you rejected.
5. Be able to say what would make it wrong.
6. Keep it to two minutes unless asked for more.

Day 5 takeaway

Day 6 Interviews

The questions, what they test, and the ranking answer

The questions that actually get asked, grouped by what they are really testing. Your capstone is the answer to most of them.

Fundamentals

Question	What is being tested	Where it was covered
Explain the bias-variance trade-off	Whether you understand generalisation or memorised a phrase	Week 4
Your model has 99% accuracy. What do you check?	Leakage instincts, and class imbalance	Weeks 4, 7
When would you not use accuracy?	Whether metrics are chosen or defaulted	Weeks 5, 7
Why does regularisation help?	Whether you can connect the penalty to the variance	Week 4
How do you handle missing data?	Whether imputation happens inside the pipeline	Weeks 2, 8
Random forest versus gradient boosting	Whether you know why the averaging differs	Week 6
What does a p-value mean?	Precision of language under mild pressure	Week 3

Applied

Question	What a strong answer contains
Walk me through a project	The decision, the baseline, one difficulty, the result with its comparison, and a limitation
How would you detect data leakage?	The prediction-time rule; suspiciously high scores; feature importance dominated by one column; a time-based split as the check
A model works in testing and fails in production. Why?	Drift, leakage in training, a train/serve preprocessing mismatch, or a population that differs from the training sample
How do you choose a threshold?	From the relative cost of the two errors and the operational capacity, not from 0.5
How would you explain this model to a non-technical stakeholder?	Effects rather than coefficients, one worked example, and the limits stated plainly
When would you not use machine learning?	A rule works; there are too few examples; the decision must be auditable; the cost of an error is catastrophic

The last one is a good question and a good answer

Week 12 ended with a support vector machine beating a convolutional network, and week 15 with a boosted challenger losing to logistic regression. Being able to say "I have measured a case where the simple thing won, and shipped the simple thing" is a stronger signal than any list of architectures. Enthusiasm for complexity is common; judgement about when to avoid it is not.

A whiteboard exercise you can practise now

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 # "Design a system to decide which customers to call this month."
11 # Talk through it before writing anything. The structure below is the
12 # answer; the code is just proof it is not hand-waving.
13 from sklearn.model_selection import train_test_split
14 from sklearn.linear_model import LogisticRegression
15 from sklearn.compose import ColumnTransformer
16 from sklearn.pipeline import Pipeline
17 from sklearn.impute import SimpleImputer
18 from sklearn.preprocessing import StandardScaler, OneHotEncoder
19
20 CAPACITY = 200          # the constraint that decides everything
21
22 prep = ColumnTransformer([
23     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
24                       ('s', StandardScaler())])), NUM),
25     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
26                       ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),
27 ])
28 X_tr, X_te, y_tr, y_te = train_test_split(
29     df[NUM + CAT], df['churned'], test_size=0.25,
30     stratify=df['churned'], random_state=42)
31 model = Pipeline([('prep', prep),
32                   ('clf', LogisticRegression(max_iter=2000))]).fit(X_tr, y_tr)
33
34 risk = model.predict_proba(X_te)[:, 1]
35 order = np.argsort(-risk)
36 called = order[:CAPACITY]
37
38 print('%d customers, capacity %d' % (len(X_te), CAPACITY))
39 print('churners in the whole test set:      %d' % y_te.sum())
40 print('churners among the %d we would call: %d'
41       % (CAPACITY, y_te.iloc[called].sum()))
42 print(' hit rate %.1f%% against a base rate of %.1f%%'
43       % (100 * y_te.iloc[called].mean(), 100 * y_te.mean()))
44 print(' lift %.2fx' % (y_te.iloc[called].mean() / y_te.mean()))
45 print('\nnote what was never chosen: a threshold. Capacity chose it.')

```

```

750 customers, capacity 200
churners in the whole test set:      201
churners among the 200 we would call: 116
 hit rate 58.0% against a base rate of 26.8%
 lift 2.16x

```

```
note what was never chosen: a threshold. Capacity chose it.
```

That last line is the point of the exercise. A great many real deployments are ranking problems with a capacity constraint, not classification problems with a threshold, and answering in those terms tells an interviewer you have thought about how the output gets used.

Day 6 takeaway

Day 7 Where This Leaves You

What was covered, what was not, and the habits worth keeping

Sixteen weeks. What was actually covered, what deliberately was not, and where to go next.

What you can now do

Weeks	Capability
1-3	Python and pandas for data work, and the linear algebra, calculus and probability underneath the models
4-6	Regression and classification, regularisation, trees and ensembles, and why the simple model sometimes wins
7-8	Validation you can trust, and features that beat model choice
9-10	Clustering, dimensionality reduction and anomaly detection without labels
11-13	Neural networks from the gradient up, convolution for images, embeddings and sequences for text
14-15	Tuning within a budget, explaining a model, and running one in production
16	Framing, shipping and communicating a project of your own

What this course did not cover

- **Reinforcement learning.** A different problem shape entirely, learning from interaction rather than from a fixed dataset.
- **Causal inference.** Week 14 drew the line and stopped at it. Everything past that line, experiments, instrumental variables, difference-in-differences, is a field of its own and the natural next one for anybody doing this work in a business.
- **Training large models.** Distributed training, mixed precision and the engineering of very large networks.
- **Recommender systems.** Collaborative filtering has its own evaluation traps.
- **Time series forecasting.** Week 7 covered splitting by time; forecasting itself is a separate discipline.
- **Fine-tuning language models.** Week 13 pointed at it. The mechanics move fast; the judgement in weeks 7 and 14 does not.

Where to go next

If you want to	Learn next
Answer "did it work" rather than "what will happen"	Causal inference and experiment design
Handle data that does not fit in memory	Spark or DuckDB, and columnar formats
Own the systems as well as the models	Docker, CI, orchestration, Airflow or Dagster
Work on language	Transformers and fine-tuning, on top of week 13
Go deeper on the theory	Convex optimisation and statistical learning theory
Get better fastest, whatever the direction	Ship three more projects end to end

The habits worth keeping

Fit the baseline first. Ask what is known at the moment of prediction. Put every transformation inside the pipeline. Read the confusion matrix. Choose the threshold from the cost of the errors. Write down the experiments that failed. Say who the model fails. Prefer the model you can explain when the numbers are close.

None of those depend on a library version, and all of them will still be true when whatever replaces today's tools has replaced them.

A last measurement

```

1 import numpy as np
2 import pandas as pd
3
4 df = pd.read_csv('customers.csv').drop_duplicates().copy()
5 df['contract'] = df['contract'].str.strip().str.title()
6 df['total_charges'] = pd.to_numeric(df['total_charges'], errors='coerce')
7
8 NUM = ['tenure_months', 'support_calls', 'monthly_charges']
9 CAT = ['contract', 'internet_service', 'payment_method', 'has_dependents']
10 from sklearn.model_selection import train_test_split, cross_val_score
11 from sklearn.linear_model import LogisticRegression
12 from sklearn.compose import ColumnTransformer
13 from sklearn.pipeline import Pipeline
14 from sklearn.impute import SimpleImputer
15 from sklearn.preprocessing import StandardScaler, OneHotEncoder
16 from sklearn.metrics import roc_auc_score
17
18 X, y = df[NUM + CAT], df['churned']
19 X_tr, X_te, y_tr, y_te = train_test_split(X, y, test_size=0.25,
20                                         stratify=y, random_state=42)
21
22 # The rule they had in week 1.
23 rule = ((X_te['contract'] == 'Month-To-Month')
24         & (X_te['support_calls'] > 2)).astype(int)
25
26 prep = ColumnTransformer([
27     ('num', Pipeline([('i', SimpleImputer(strategy='median')),
28                      ('s', StandardScaler())])), NUM),
29     ('cat', Pipeline([('i', SimpleImputer(strategy='most_frequent')),
30                      ('o', OneHotEncoder(handle_unknown='ignore'))])), CAT),

```

```

31 ])
32 model = Pipeline([('prep', prep),
33                   ('clf', LogisticRegression(max_iter=2000))]).fit(X_tr, y_tr)
34 cv = cross_val_score(model, X_tr, y_tr, cv=5, scoring='roc_auc')
35
36 print('the rule in use          AUC %.4f' % roc_auc_score(y_te, rule))
37 print('logistic regression, cv   AUC %.4f +/- %.4f' % (cv.mean(), cv.std()))
38 print('logistic regression, test AUC %.4f'
39       % roc_auc_score(y_te, model.predict_proba(X_te)[: , 1]))
40 print('\nthirteen lines. everything else this course taught you was')
41 print('about knowing whether to believe the last number.')

```

```

the rule in use          AUC 0.5351
logistic regression, cv   AUC 0.8217 +/- 0.0110
logistic regression, test AUC 0.8174

thirteen lines. everything else this course taught you was
about knowing whether to believe the last number.

```

Course takeaway

Self-Assessment

Try It Yourself

- The strongest way to begin a project is with:
 - A Kaggle competition
 - An interesting dataset
 - A decision somebody currently makes badly
 - A model architecture
- "What is known at the moment the prediction is made" is the question that prevents:
 - Class imbalance
 - Underfitting
 - Overfitting
 - Target leakage
- In a project layout, code that runs more than once belongs in:
 - An importable module under `src/`
 - A markdown cell
 - The README
 - The notebook
- The most valuable single number in a capstone is:
 - The final AUC

- b) The baseline it is compared against
 - c) The cross-validation standard deviation
 - d) The training time
5. Roughly what share of a real project is framing, cleaning and joining data?
- a) About 90%
 - b) About 15%
 - c) About half
 - d) About 5%
6. Including experiments that failed in your write-up:
- a) Is only for academic work
 - b) Should be an appendix
 - c) Weakens it
 - d) Is evidence of judgement
7. A reviewer looks first at:
- a) The README, the validation and the baseline
 - b) The hyperparameters
 - c) The number of commits
 - d) The model class
8. Many real deployments are best framed as:
- a) Classification with a 0.5 threshold
 - b) Ranking under a capacity constraint, where capacity chooses the threshold
 - c) Clustering
 - d) Regression

Answers: 1c, 2d, 3a, 4b, 5c, 6d, 7a, 8b. Anything you got wrong points at the day to re-read, not at a gap in ability.