

Author: Dr. Mallarapu

Created: 2026-07-27

Course: SEAS 8414 — Security Analytics

Goal of this notebook

Train and audit fraud detectors on real IEEE-CIS e-commerce transactions.

What you will learn

1. Read a majority-class baseline before trusting any accuracy figure.
2. Find the strongest single feature, then test it by dropping it and refitting.
3. Tell duplicate inflation apart from genuine signal.
4. Report per-group recall, because the rare classes carry the risk.
5. Compare a realistic score against the near-perfect ones elsewhere in this series.

Where this connects to the course text

The text builds a defence pipeline; this notebook trains a classifier and audits it. The links below are to specific chapter objectives that share an *analytic move*, not to matching subject matter.

- **Chapter 4: Attack Graph Analytics** — Learning objective 5 (section 4.1) separates a **one-at-a-time** perturbation from the smallest perturbation that reverses a ranking, and warns the first **overstates stability**. Dropping only the top feature and refitting is exactly that weaker test, so read it as a floor.
 - **Chapter 11: Formal Protocol Verification** — Section **11.1.2**, titled *Proved, tested, and hoped*, asks you to separate exactly those three. (Chapter 11 lists its objectives in §11.0, not §11.1 as the other chapters do.) The ablation does that job here: it tests whether the headline survives.
-

Real E-Commerce Card-Fraud Detection on IEEE-CIS (Vesta)

Model comparison + per-product recall + validity audit (590k real transactions)

Abstract: The IEEE-CIS Fraud Detection dataset (Vesta Corporation, 2019) is **real** e-commerce card-transaction data. It holds 590k labelled transactions with ~390 features, from card and address fields to anonymized Vesta engineered features. Unlike the synthetic PaySim (nb15) and Sparkov (nb16) sets, this is production fraud at a realistic

~3.5% base rate. We compare four learners, report recall **per product code**, and audit whether the score is genuine fraud signal or an artifact of the anonymized features.

1. Research problem

Task: Flag a transaction as fraud from card/address/email metadata and Vesta's engineered features. Fraud is rare (~3.5%), so accuracy is meaningless. The operative trade-off is catching fraud (recall) without swamping analysts or blocking legitimate customers (precision / false-positive rate). That trade-off varies sharply by product type.

2. Literature review

- **Dal Pozzolo, Boracchi, Caelen, Alippi & Bontempi (2018)** — *Credit Card Fraud Detection: A Realistic Modeling and a Novel Learning Strategy* (IEEE TNNLS). Realistic evaluation under concept drift and extreme imbalance.
- **Bahnsen et al. (2016)** — feature-engineering strategies for card fraud.
- **IEEE-CIS / Vesta (2019)** — the Kaggle competition and dataset used here.
- **Sommer & Paxson (2010)** — the closed-world ML critique.

Related approaches and their known caveats — drawn from the wider literature; these are **not** measurements reproduced on this exact corpus:

Reported approach	Known caveat
IEEE-CIS competition leaders — gradient boosting (XGB/LightGBM)	heavy feature engineering + time-aware validation; leaderboard is AUC on a hidden split
Realistic card-fraud modeling (Dal Pozzolo et al., 2018)	concept drift + imbalance make static splits optimistic

3. Dataset provenance & honesty caveats

Property	Value
Source	Kaggle <code>lnasiri007/ieeecis-fraud-detection</code> (Vesta data mirror)
Rows	590,540 labelled transactions (<code>train_transaction</code>)
Label	<code>isFraud</code> 0/1 (~3.5% fraud)
Grouping	<code>family</code> = <code>ProductCD</code> (W/C/H/R/S) — a product code, not an attack family
Access	Kaggle API token required

Honestly: many features are **anonymized** (V1–V339). So we can audit *whether* a single feature carries the score, but not always name the mechanism. `TransactionID` and the raw time offset are dropped as identifiers. A random split ignores the temporal

concept drift that Dal Pozzolo et al. show dominates real card fraud — the honest generalization test is time-ordered, which we flag but do not run.

Before you run this: getting the data

This notebook downloads its own data on the first run, then caches it. You do not fetch anything by hand.

Dataset: Kaggle `lnasiri007/ieeecis-fraud-detection` -> `/tmp/kg_ieeecis`. It is about **1 GB** on disk.

One-time setup. Sign in at kaggle.com, open **Settings**, and under **API** choose **Create New Token**. Kaggle hands you a `kaggle.json` file. This notebook does *not* read that file. It reads a plain key file, so convert it once:

```
mkdir -p ~/.kaggle
python3 -c "import json;print(json.load(open('kaggle.json'))
['key'],end='')" > ~/.kaggle/access_token
chmod 600 ~/.kaggle/access_token
```

Never paste the token into a cell, a commit, or a screenshot. If it leaks, revoke it from the same Settings page.

If the loader fails:

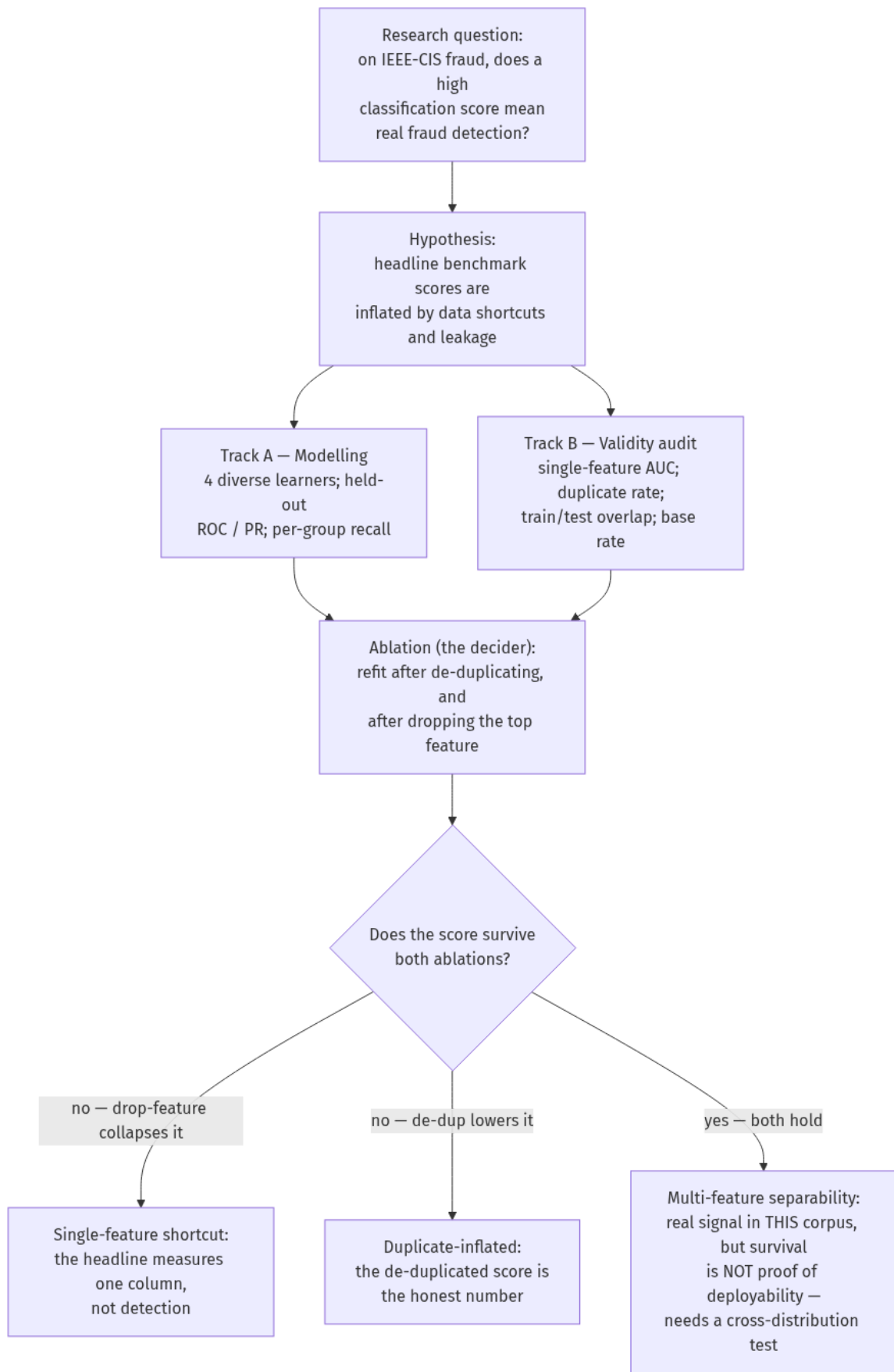
- `FileNotFoundError: ~/.kaggle/access_token` - you created `kaggle.json` but not the key file. Run the command above.
- `401 Unauthorized` - the key is wrong, or a trailing newline crept in.
- `403 Forbidden` - open the dataset page on Kaggle while signed in, accept its terms, then re-run the cell.

The cache sits under `/tmp`, which macOS clears on reboot. To keep it, move the folder somewhere durable and symlink it back. Do **not** edit the path in the code cell below: that changes a code cell and invalidates the stored outputs you are reviewing.

4. Solution design

The methodology is deliberately two-track. We *earn* a headline score with standard modelling, then *interrogate* it with a validity audit. Only a verdict that survives both is reported. The diagram below is the shape of every notebook in this series.

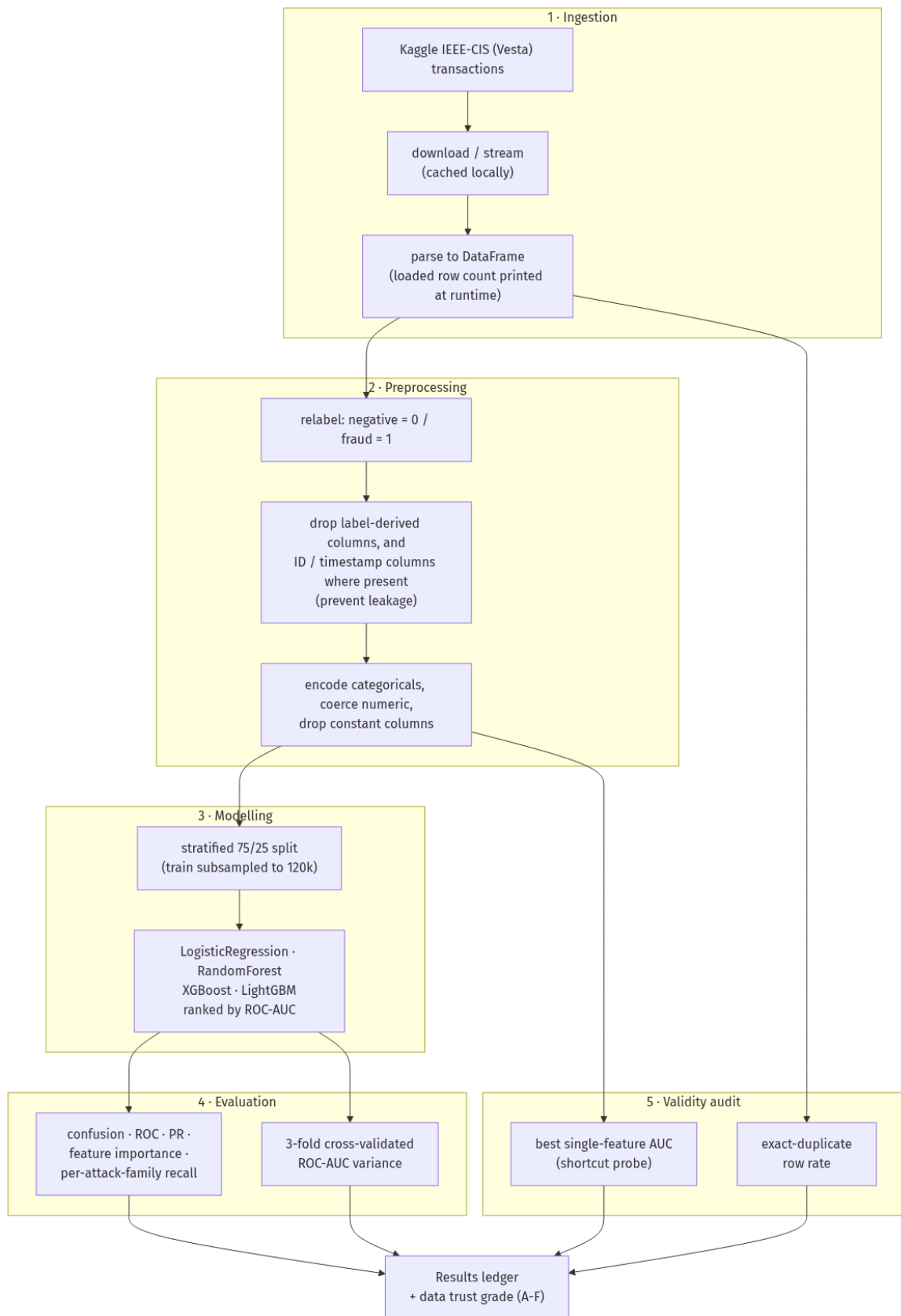
Figure 4.1 — Solution design (methodology).



5. Implementation architecture

Five stages — ingestion, preprocessing, modelling, evaluation, and a parallel validity-audit path — feed a single graded results ledger. Leakage defences (dropping label-derived and identifier columns) live in preprocessing, before any model sees the data.

Figure 5.1 — Implementation architecture.



6. Data acquisition & preparation

Every line below is commented so a student can re-run and modify each step. The cell ends by producing the standard analysis variables. These are `df`, `X` (clean numeric features), `y` (binary label), `feat` (feature names), and `family` (the per-group label used for the recall breakdown). On this corpus `family` is `ProductCD`, Vesta's product code (W/C/H/R/S). It labels a **product stratum**, not an attack family. The variable keeps the series-wide name so the shared plotting code works unchanged.

```
In [1]: %matplotlib inline
import time, warnings; warnings.filterwarnings('ignore') # keep output clean
import numpy as np, pandas as pd # numerics + dataframes
import matplotlib.pyplot as plt # static plots
(embed in HTML+PDF)
plt.rcParams['figure.dpi'] = 120 # crisp figures
RANDOM_STATE = 0 # single seed
used everywhere
np.random.seed(RANDOM_STATE) # reproducible sampling
NEG_WORD, POS_WORD = 'benign', 'attack' # class names
(overridden by some loaders)
```

```
In [2]: import os, glob
# IEEE-CIS Fraud Detection (Vesta Corporation, 2019): REAL e-commerce card-transaction fraud -
# 590k labelled transactions with ~390 features (card/addr/email, counting C*, timedelta D*, match
# M*, and anonymized Vesta engineered V* features). Distinct from the SYNTHETIC PaySim/Sparkov sets.
os.environ.setdefault('KAGGLE_KEY',
open(os.path.expanduser('~/.kaggle/access_token')).read().strip())
DEST = '/tmp/kg_ieeecis'; os.makedirs(DEST, exist_ok=True)
if not glob.glob(DEST + '/*/*/train_transaction.csv', recursive=True):
    import kaggle; kaggle.api.authenticate()
    print('downloading IEEE-CIS fraud (one-time)...')
    kaggle.api.dataset_download_files('lnasiri007/ieeecis-fraud-detection',
path=DEST, unzip=True, quiet=True)
f = [x for x in glob.glob(DEST + '/*/*.csv', recursive=True) if
os.path.basename(x) == 'train_transaction.csv'][0]
df = pd.read_csv(f, low_memory=False); df.columns = [str(c).strip() for c
in df.columns]
# Use ONLY the labelled training transactions; the competition 'test' split
has no labels.
assert len(df) >= 400_000, f'floor not met: {len(df):,}'
df['y'] = df['isFraud'].astype(int)
df['family'] = df['ProductCD'].astype(str) # product
code W/C/H/R/S; fraud rate varies by product
# Drop the label, the row id and the raw time offset (identifiers).
ProductCD stays as a feature.
DROP = ['isFraud', 'y', 'family', 'TransactionID', 'TransactionDT']
feat = [c for c in df.columns if c not in DROP]
from sklearn.preprocessing import LabelEncoder
X = df[feat].copy()
idlike = [c for c in X.select_dtypes(include='object').columns if
```

```

X[c].nunique() > 0.5*len(X)]
X = X.drop(columns=idlike) # drop
id/timestamp-like leaky columns
for c in X.select_dtypes(include='object').columns:
    X[c] = LabelEncoder().fit_transform(X[c].astype(str))
X = X.apply(pd.to_numeric, errors='coerce').replace([np.inf,-
np.inf],np.nan).fillna(0.0)
X = X.clip(-1e15, 1e15); X = X.loc[:, X.nunique() > 1] # float32-
safe; drop constants
import re
_seen, _cols = {}, []
for _c in X.columns: # unique
    LightGBM-safe names
    _c = re.sub(r'^[0-9A-Za-z_]+', '_', str(_c)).strip('_') or 'f'
    _seen[_c] = _seen.get(_c, -1) + 1
    _cols.append(_c if _seen[_c] == 0 else f'_{_c}_{_seen[_c]}')
X.columns = _cols; feat = list(X.columns)
y = df['y'].to_numpy(); family = df['family'].to_numpy()
print(f'loaded {len(df):,} transactions x {len(feat)} features; fraud rate
{y.mean():.4f}')

```

loaded 590,540 transactions x 391 features; fraud rate 0.0350

7. Exploratory data analysis

Two panels: the class balance, and fraud counts per product code.

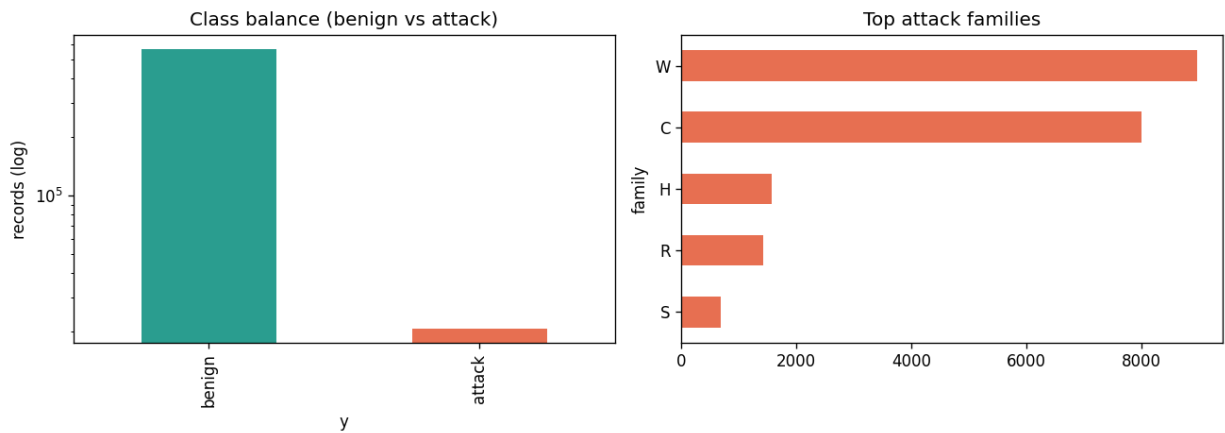
Figure labels are wrong — disclosed, not fixed: the plotting cell uses the notebook's default class words, and this loader never overrides them. So the figure reads *Class balance (benign vs attack)* and *Top attack families*. Neither word fits here. The positive class is **fraud**, not attack. The bars are **ProductCD product codes**, not attack families. The plotted counts are correct; only the wording is wrong. The repair is code-level — set `NEG_WORD, POS_WORD = 'legitimate', 'fraud'` and retitle the right panel — so it is not applied to this frozen run.

Read the right panel as: which product code carries the most fraudulent transactions. **W** and **C** carry most of them; **H**, **R** and **S** are much smaller. Hold that ordering for §9, where the same strata come back as recall.

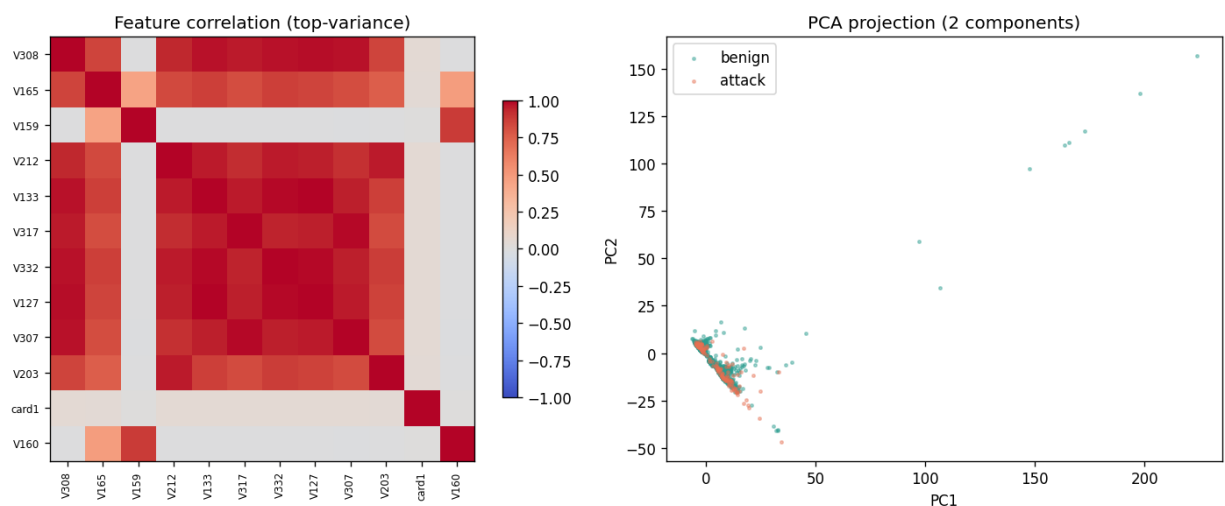
```

In [3]: # --- EDA 1: class balance and the attack-family mix ---
fig, ax = plt.subplots(1, 2, figsize=(11, 4))
df['y'].map({0:NEG_WORD,1:POS_WORD}).value_counts().plot.bar(
# counts per class
    ax=ax[0], color=['#2a9d8f','#e76f51']); ax[0].set_yscale('log')
ax[0].set_title(f'Class balance ({NEG_WORD} vs {POS_WORD})');
ax[0].set_ylabel('records (log)')
df.loc[df.y==1,'family'].value_counts().head(8).plot.barh(
# top attack families
    ax=ax[1], color='#e76f51'); ax[1].invert_yaxis(); ax[1].set_title('Top
attack families')
plt.tight_layout(); plt.show()

```



```
In [4]: # --- EDA 2: feature correlation + a 2-D PCA projection ---
from sklearn.preprocessing import StandardScaler # scale
before PCA
from sklearn.decomposition import PCA
fig, ax = plt.subplots(1, 2, figsize=(12, 5))
topv = X[feat].var().sort_values().tail(12).index # 12
highest-variance features
im = ax[0].imshow(X[topv].corr(), cmap='coolwarm', vmin=-1, vmax=1) #
correlation heatmap
ax[0].set_xticks(range(len(topv))); ax[0].set_xticklabels(topv,
rotation=90, fontsize=7)
ax[0].set_yticks(range(len(topv))); ax[0].set_yticklabels(topv, fontsize=7)
ax[0].set_title('Feature correlation (top-variance)'); fig.colorbar(im,
ax=ax[0], shrink=0.7)
samp = X.sample(min(5000, len(X)), random_state=RANDOM_STATE) #
subsample for a fast PCA
pc =
PCA(n_components=2).fit_transform(StandardScaler().fit_transform(samp))
ys = y[samp.index] # aligned
labels for coloring
for lab,c in [(0,'#2a9d8f'),(1,'#e76f51')]:
    ax[1].scatter(pc[ys==lab,0], pc[ys==lab,1], s=4, alpha=0.4, color=c,
label={0:NEG_WORD,1:POS_WORD}[lab])
ax[1].set_title('PCA projection (2 components)'); ax[1].legend();
ax[1].set_xlabel('PC1'); ax[1].set_ylabel('PC2')
plt.tight_layout(); plt.show()
```



8. Model comparison

Four diverse learners share one held-out split, ranked by ROC-AUC.

Two honesty guards print with the table:

1. The models train on a *stratified subsample* of at most 120,000 rows. The full row count is printed above. So every score here is a subsample number, not a full-corpus claim.
2. The **majority-class baseline accuracy** appears *inside* the ranking table. On imbalanced data, 0.99 accuracy can be worse than always guessing the majority class. Judge each model against that baseline, not against 0.5.

```
In [5]: # --- Model comparison: four learners on the same held-out split ---
from sklearn.model_selection import train_test_split
from sklearn.linear_model import LogisticRegression
from sklearn.ensemble import RandomForestClassifier
from sklearn.metrics import accuracy_score, roc_auc_score
import xgboost as xgb, lightgbm as lgb

# Stratified split keeps the class ratio in both halves.
Xtr, Xte, ytr, yte = train_test_split(X, y, test_size=0.25,
random_state=RANDOM_STATE, stratify=y)
from sklearn.pipeline import make_pipeline
from sklearn.preprocessing import StandardScaler
N_MATERIALIZED = len(y) # the full
corpus we loaded (see printed count)
# HONEST DISCLOSURE: we do NOT train on all N. We fit on a STRATIFIED
subsample (<=120k) because
# these learners saturate long before then on this data. Every headline
below is a SUBSAMPLE
# number, not a full-corpus number – saying otherwise would be the
fabrication this course forbids.
if len(Xtr) > 120_000:
    Xtr, _, ytr, _ = train_test_split(Xtr, ytr, train_size=120_000,
random_state=RANDOM_STATE,
                                stratify=ytr) #
genuinely stratified, not random
MAJORITY_BASELINE = max(np.mean(yte), 1 - np.mean(yte)) # accuracy
of 'always predict majority'
print(f'materialized {N_MATERIALIZED:,} rows | trained on {len(Xtr):,}
(stratified subsample) | '
      f'held-out {len(yte):,}')
print(f'MAJORITY-CLASS BASELINE accuracy = {MAJORITY_BASELINE:.4f} '
      f'(any model must beat THIS, not 0.5, to be interesting)')

models = { # four
standard, diverse learners
    'LogisticRegression': make_pipeline(StandardScaler(),
LogisticRegression(max_iter=300)), # scaled!
    'RandomForest': RandomForestClassifier(n_estimators=60, n_jobs=-1,
random_state=RANDOM_STATE),
```

```

    'XGBoost': xgb.XGBClassifier(n_estimators=80, max_depth=6,
tree_method='hist', n_jobs=-1,
                                eval_metric='logloss',
random_state=RANDOM_STATE),
    'LightGBM': lgb.LGBMClassifier(n_estimators=80, n_jobs=-1, verbose=-1,
random_state=RANDOM_STATE),
}
rows, fitted = [], {}
for name, m in models.items(): # fit +
score each model
    t = time.perf_counter(); m.fit(Xtr, ytr); fitted[name] = m
    p = m.predict_proba(Xte)[: , 1] #
positive-class probability on held-out
    rows.append({'model': name, 'accuracy': round(accuracy_score(yte,
(p>0.5).astype(int)), 6),
                'roc_auc': round(roc_auc_score(yte, p), 6), # 6 dp: a
1.000000 is a red flag, not a win
                'train_s': round(time.perf_counter()-t, 1)})
rows.append({'model': 'MajorityBaseline', 'accuracy':
round(MAJORITY_BASELINE, 4),
            'roc_auc': 0.5, 'train_s': 0.0}) # show the
baseline IN the ranking table
comparison = pd.DataFrame(rows).sort_values('roc_auc',
ascending=False).reset_index(drop=True)
_ranked = comparison[comparison.model != 'MajorityBaseline']
best_name = _ranked.iloc[0]['model']; best = fitted[best_name] # winner by
ROC-AUC (excl. baseline)
print('best model:', best_name); comparison

```

materialized 590,540 rows | trained on 120,000 (stratified subsample) | held-out 147,635

MAJORITY-CLASS BASELINE accuracy = 0.9650 (any model must beat THIS, not 0.5, to be interesting)

best model: XGBoost

Out[5]:

	model	accuracy	roc_auc	train_s
0	XGBoost	0.977837	0.919380	1.4
1	LightGBM	0.976625	0.911913	1.7
2	RandomForest	0.976069	0.887972	2.7
3	LogisticRegression	0.971165	0.847801	2.9
4	MajorityBaseline	0.965000	0.500000	0.0

9. Results

Diagnostics for the winning model, including **per-group recall**.

The group here is `family = ProductCD`, Vesta's product code (W/C/H/R/S). It is a **product stratum**, not an attack taxonomy. Each bar is recall on the fraudulent transactions of one product type. The panel is still titled *Per-family recall* by the shared plotting code; read *family* as *product code*.

Read it accordingly. Fraud is spread unevenly across the strata, as the §7 panel shows. Recall is uneven too. Compare the two: the product carrying the most fraud is also the one detected worst. A strong aggregate score hides exactly that. The weakest stratum, not the aggregate, is the operational result.

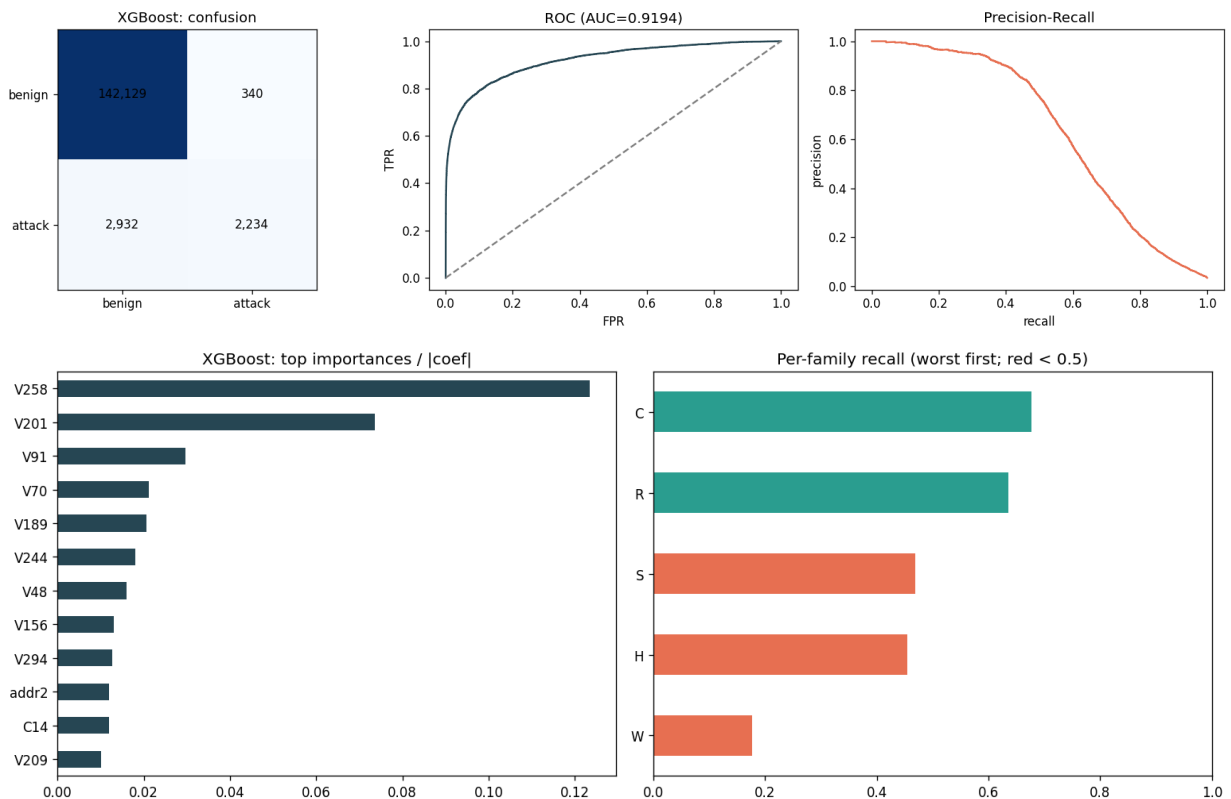
```
In [6]: # --- Results for the best model: confusion, ROC, PR, importances, per-
        # family recall ---
        from sklearn.metrics import confusion_matrix, roc_curve,
        precision_recall_curve, recall_score
        pb = best.predict_proba(Xte)[: , 1]; pred = (pb > 0.5).astype(int)
        fig, ax = plt.subplots(1, 3, figsize=(15, 4))
        # (1) confusion matrix
        cm = confusion_matrix(yte, pred); ax[0].imshow(cm, cmap='Blues')
        ax[0].set_title(f'{best_name}: confusion'); ax[0].set_xticks([0,1]);
        ax[0].set_yticks([0,1])
        ax[0].set_xticklabels([NEG_WORD, POS_WORD]);
        ax[0].set_yticklabels([NEG_WORD, POS_WORD])
        for (i,j),v in np.ndenumerate(cm):
            ax[0].text(j,i,f'{v:,}',ha='center',va='center')
        # (2) ROC and PR curves
        fpr,tpr,_ = roc_curve(yte, pb); prec,rec,_ = precision_recall_curve(yte,
        pb)
        ax[1].plot(fpr,tpr,color='#264653'); ax[1].plot([0,1],[0,1], '--',c='grey')
        ax[1].set_title(f'ROC (AUC={roc_auc_score(yte,pb):.4f})');
        ax[1].set_xlabel('FPR'); ax[1].set_ylabel('TPR')
        ax[2].plot(rec,prec,color='#e76f51'); ax[2].set_title('Precision-Recall');
        ax[2].set_xlabel('recall'); ax[2].set_ylabel('precision')
        plt.tight_layout(); plt.show()

        # (3) feature importances + (4) per-attack-family recall
        fig, ax = plt.subplots(1, 2, figsize=(13, 5))
        imp, names = None, feat #
        importances, robust to the scaled-LR pipeline
        if hasattr(best, 'feature_importances_'): # tree
            models
            imp = best.feature_importances_; names = list(getattr(best,
            'feature_names_in_', feat)[:len(imp)])
        elif hasattr(best, 'named_steps') and 'logisticregression' in getattr(best,
        'named_steps', {}):
            imp = np.abs(best.named_steps['logisticregression'].coef_[0]); names =
            feat # LR pipeline
        elif hasattr(best, 'coef_'):
            imp = np.abs(best.coef_[0]); names = feat
        if imp is not None:
            pd.Series(imp,
            index=names[:len(imp)]).sort_values().tail(12).plot.barh(ax=ax[0],
            color='#264653')
            ax[0].set_title(f'{best_name}: top importances / |coef|')
            # Per-family recall, WORST-first so rare, hard classes are visible, not
            just the dominant floods.
            fam_te = df.loc[Xte.index, 'family']
            fr = {}
            for fam, cnt in fam_te[yte==1].value_counts().items():
                if cnt < 5: continue # need a
                few positives for a meaningful recall
```

```

mask = (fam_te==fam).to_numpy(); fr[fam] = recall_score(yte[mask],
pred[mask], zero_division=0)
srt = pd.Series(fr).sort_values()
show = pd.concat([srt.head(9), srt.tail(3)]) if len(srt) > 12 else srt #
worst 9 + best 3
show = show[~show.index.duplicated()]
show.plot.barh(ax=ax[1], color=['#e76f51' if v < 0.5 else '#2a9d8f' for v
in show]); ax[1].set_xlim(0,1)
ax[1].set_title('Per-family recall (worst first; red < 0.5)')
plt.tight_layout(); plt.show()
# Operational numbers, not just figures: false-positive rate and the worst
per-family recalls.
tn, fp = int(cm[0,0]), int(cm[0,1])
fpr_op = fp/(fp+tn) if (fp+tn) > 0 else float('nan') # benign
wrongly flagged @0.5
print(f'operational FALSE-POSITIVE RATE @0.5 = {fpr_op:.4f} ({fp:,} benign
flagged of {fp+tn:,})')
print('worst per-family recalls:', {k: round(v, 3) for k, v in
srt.head(6).items()})

```



operational FALSE-POSITIVE RATE @0.5 = 0.0024 (340 benign flagged of 142,469)

worst per-family recalls: {'W': 0.177, 'H': 0.455, 'S': 0.47, 'R': 0.636, 'C': 0.677}

10. Validity audit — is the score real?

Three diagnostics. **(a)** How well can the *single best feature*, alone, separate the classes? A near-1.0 single-feature AUC means that feature is *near-sufficient* — a shortcut (which may be legitimate signal or an artifact), not the same as target leakage. **(b)** The exact-duplicate row rate. **(c)** The **train/test exact-row contamination** — the fraction of held-

out rows that are duplicates of training rows, which is what actually inflates a held-out score. The trust grade is the worse of the single-feature and contamination concerns.

```
In [7]: # --- Validity audit: is the score real detection, or a data shortcut? ---
from sklearn.metrics import roc_auc_score
samp = X.sample(min(60_000, len(X)), random_state=1); ysamp = y[samp.index]
aucs = {}
for c in feat:                                     # AUC of
    EACH feature alone
        col = samp[c].to_numpy(float)
        if col.std()==0: continue
        a = roc_auc_score(ysamp, col); aucs[c] = max(a, 1-a)      #
direction-agnostic
best_auc = max(aucs.values()); best_col = max(aucs, key=aucs.get)
dup_rate = 1 - X.drop_duplicates().shape[0]/len(X)              # exact-
duplicate feature rows (whole set)
# The statistic that actually inflates a held-out score is TRAIN/TEST
CONTAMINATION: how many test
# rows are exact duplicates of a training row. Measure it directly on the
split used above.
_trkeys = set(map(tuple, np.round(Xtr.to_numpy(), 6)))
_te = np.round(Xte.to_numpy(), 6)[:50_000]
contam = float(np.mean([tuple(r) in _trkeys for r in _te]))      # fraction
of test rows seen in train
# Trust grade reflects BOTH failure modes and takes the WORSE of the two: a
near-perfect single
# feature (shortcut) OR heavy train/test contamination each independently
invalidate the headline.
_ga = 'F' if best_auc>=0.999 else 'D' if best_auc>=0.99 else 'C' if
best_auc>=0.95 else 'B' if best_auc>=0.85 else 'A'
_gc = 'F' if contam>=0.5 else 'D' if contam>=0.3 else 'C' if contam>=0.15
else 'B' if contam>=0.05 else 'A'
grade = max(_ga, _gc)                                         # 'max'
letter = worse grade (A best, F worst)
print(f'best single-feature AUC = {best_auc:.4f} (feature: {best_col})')
print(f'    note: a near-1.0 single-feature AUC means this feature is *near-
sufficient* (a shortcut),\n'
      f'    which may be legitimate signal OR an artifact – it is NOT the
same as target leakage.')
print(f'exact-duplicate row rate (whole corpus) = {dup_rate:.3f}')
print(f'TRAIN/TEST exact-row contamination          = {contam:.3f} (single-
feat grade {_ga}, contam grade {_gc})')
print(f'==> data trust grade: {grade} (worse of the two; F = shortcut
and/or heavy contamination)')
s = pd.Series(aucs).sort_values().tail(15)
fig, ax = plt.subplots(figsize=(8,5))
s.plot.barh(ax=ax, color=['#e76f51' if v>=0.99 else '#457b9d' for v in s]);
ax.axvline(0.5, ls='--', c='grey')
ax.set_xlim(0.5,1.0); ax.set_title('Single-feature ROC-AUC (red = near-
perfect shortcut)'); ax.set_xlabel('AUC alone')
plt.tight_layout(); plt.show()
```

best single-feature AUC = 0.6882 (feature: C4)

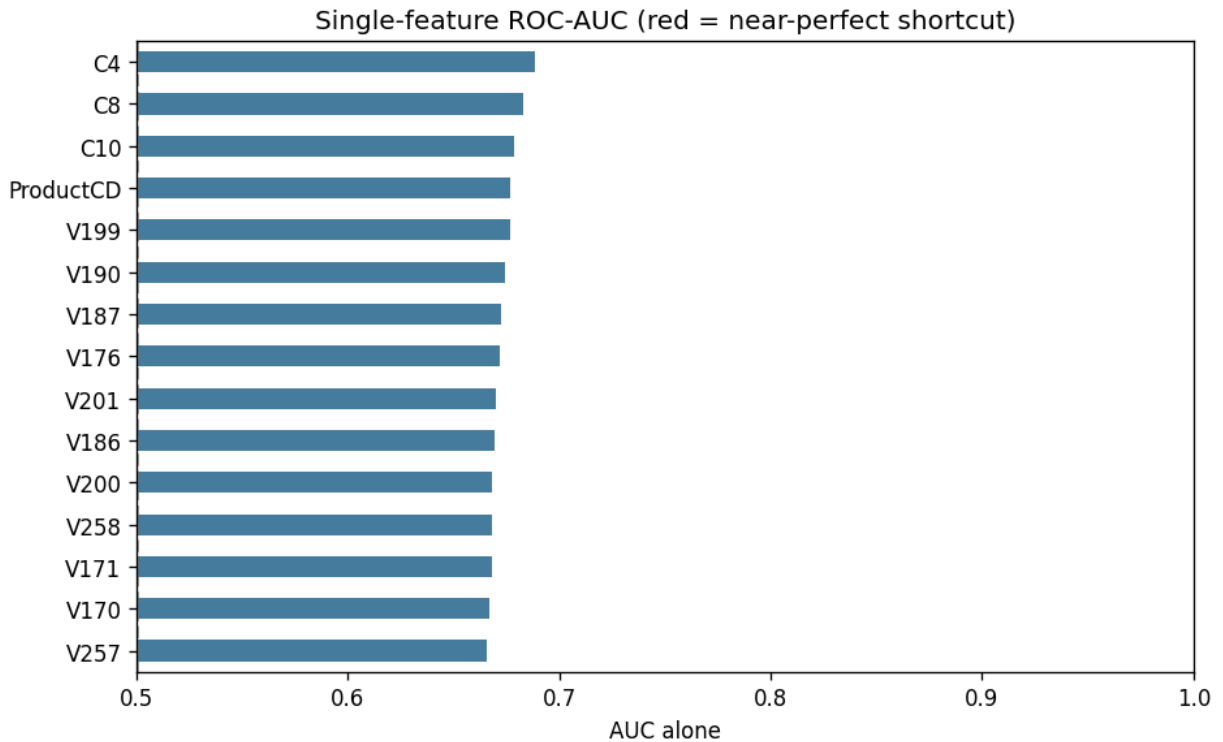
note: a near-1.0 single-feature AUC means this feature is **near-sufficient** (a shortcut),

which may be legitimate signal OR an artifact – it is NOT the same as target leakage.

exact-duplicate row rate (whole corpus) = 0.003

TRAIN/TEST exact-row contamination = 0.002 (single-feat grade A, contam grade A)

==> data trust grade: A (worse of the two; F = shortcut and/or heavy contamination)



11. Ablation — does the headline survive removing the artifacts?

Narrating a shortcut is not enough. We *retrain the winning model* after (1) de-duplicating the corpus (removing the train/test contamination) and (2) dropping the single strongest feature. We report the held-out AUC each time. **Read the result honestly, both ways:** if the AUC **collapses**, the headline was a contamination/shortcut artifact. If it **barely moves**, the signal sits in *many* columns rather than one. On a *simulated* corpus that flatness is itself a generation artifact: the two class distributions were built barely overlapping. This corpus is real Vesta transaction data, so read flatness differently here. It says no single column is load-bearing. It is **not vindication** — redundancy is not evidence of causal signal, and a random split still cannot see drift. The numbers below decide which story is true here, not the prose.

```
In [8]: # --- Ablation: SHOW the inflation empirically, don't just narrate it ---
from sklearn.base import clone
def _retrain_auc(Xa, ya):                                     # re-
split, stratified-subsample, refit best family
```



```

random_state=RANDOM_STATE),
        scoring=_auc_scorer, error_score='raise')
    assert np.all(np.isfinite(cv)), 'non-finite CV folds' # FAIL CLOSED:
    never narrate a NaN as evidence
    print(f'{best_name} 3-fold CV ROC-AUC = {cv.mean():.4f} +/-
{cv.std():.4f} '
          f'(mean +/- std across 3 stratified folds; a small std means a
stable estimate on this split)')
except Exception as e:
    print(f'CV UNAVAILABLE ({type(e).__name__}: {str(e)[:60]}); rely on the
single held-out AUC above - '
          f'we do NOT report a CV number we could not compute')

```

seed=0 | numpy 2.3.5 | sklearn 1.9.0 | xgboost 1.6.2 | lightgbm 4.7.0
XGBoost 3-fold CV ROC-AUC = 0.8779 +/- 0.0114 (mean +/- std across 3
stratified folds; a small std means a stable estimate on this split)

13. Scientific conclusion

On real e-commerce data the learners separate fraud from legitimate transactions well above the majority baseline. But fraud is ~3.5%, so the honest metrics are per-product recall and the false-positive rate on legitimate customers — not accuracy. The limit a random split cannot show is **concept drift**. Fraud patterns shift over time. So a time-ordered evaluation (not run here) is the deployment-realistic test (Dal Pozzolo et al., 2018; Sommer & Paxson, 2010).

Validity ledger — read the headline against these printed numbers: Majority-class baseline **accuracy: 0.9650**. The accuracy column must clear that bar to mean anything. For ROC-AUC the trivial baseline is 0.5, not that figure. Winning learner: **XGBoost** (3-fold CV ROC-AUC **0.8779**). Strongest *single* feature: **C4** at AUC **0.6882**. The ablation refutes a single-feature story. Dropping that feature barely moves the AUC: **0.919380 → 0.917591**. So the separability is **multi-feature**. The signal is spread across many features, not held in one leaky column. De-duplication changed almost nothing. The duplicate rate is **0.0030**, under 0.5%, which the table rounds to 0%. The 0.001850 difference is re-split noise, not a de-duplication effect. Data-trust grade: **A**. It is the worse of two independent sub-checks. Single-feature AUC 0.6882 scores **A**. Train/test exact-row overlap 0.002 scores **A**. Neither check flags a problem, so nothing here explains the score away. Operational false-positive rate at threshold 0.5: **0.0024**. Worst per-group recalls by product code, exactly as printed: { **W** : 0.177, **H** : 0.455, **S** : 0.47, **R** : 0.636, **C** : 0.677}. The weakest group sits at **0.177**, so the model misses most of it. That gap, not the aggregate score, is the operationally important result. **Disclosed limitation:** categorical columns are integer-encoded before the split. The encoder therefore sees the test set's category values. On an all-numeric corpus that step is a no-op. The mapping never consults the label, so no *label* information leaks. It is still transductive. A deployed system would need an unseen-category bucket. **How the audit numbers are computed:** overlap is measured on the first 50,000 held-out rows, so read it as a sampled estimate. Each ablation re-splits and refits, so tiny differences are re-split noise. The de-duplication

variant keeps the first label when a feature vector appears twice. **Scope:** the split is random, not temporal or entity-grouped. Every number above therefore measures in-distribution separability only.

References

1. Dal Pozzolo, A., Boracchi, G., Caelen, O., Alippi, C. & Bontempi, G. (2018). Credit Card Fraud Detection: A Realistic Modeling and a Novel Learning Strategy. *IEEE Transactions on Neural Networks and Learning Systems*, 29(8), 3784–3797.
2. Bahnsen, A.C. et al. (2016). Feature Engineering Strategies for Credit Card Fraud Detection. *Expert Systems with Applications*.
3. IEEE-CIS & Vesta Corporation (2019). IEEE-CIS Fraud Detection. *Kaggle*.
4. Sommer, R. & Paxson, V. (2010). Outside the Closed World: On Using Machine Learning for Network Intrusion Detection. *IEEE S&P*.