



Inside AutoML Reproducibility Hub



Static Mode
Fast • Deterministic
Shareable



Full Mode
Extended Models
Advanced Features

A Visual Guide to Reproducible Browser-Based Machine Learning

By **Faramarz Kowsari**



Built with



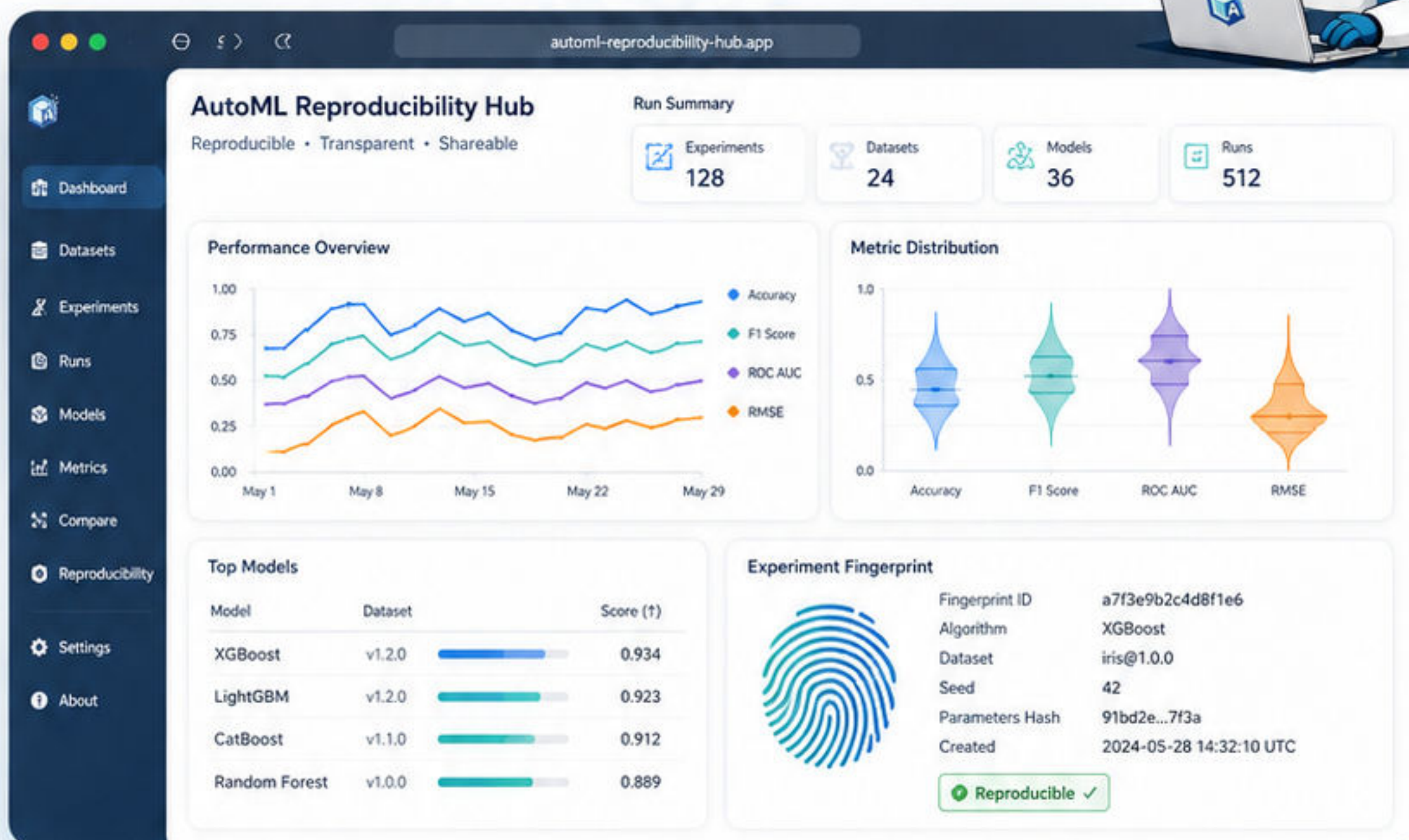
TypeScript



Pyodide
(Python in Web)



DuckDB-WASM
(In-Process Analytics)



Seeds



Deterministic
Randomness

seed = 42

Dataset
Versioning



v1.2.0

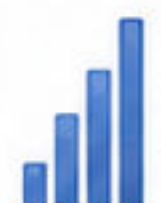
Immutable &
Trackable

Parameters

```
{  
  max_depth: 6,  
  n_estimators: 200,  
  learning_rate: 0.05  
}
```

Tracked &
Hashed

Metrics



Accuracy
F1 Score
ROC AUC
RMSE

Comprehensive
Evaluation

Experiment
Fingerprinting



Unique Identity for
Every Experiment

Reproducibility



Rerun Anywhere.
Same Result.

End-to-End
AutoML
Pipeline



Load Data



Preprocess



Feature
Engineering



Model
Selection



Train & Tune



Evaluate



Log & Version



Reproduce
& Share



Explore the project on GitHub

github.com/FaramarzKowsari/automl-reproducibility-hub





What is AutoML Reproducibility Hub?

AutoML Reproducibility Hub is an interactive, browser-based platform that helps data scientists run, track, and reproduce AutoML experiments with complete transparency.

It captures every critical detail—data, parameters, environment, and results—so you can verify, compare, and share experiments with confidence.

Browser-Based Transparent Shareable Reproducible

Mini Dashboard Snapshot



At a Glance

Run experiments.
Capture everything.
Reproduce anywhere.
Share with confidence.

The Reproducibility Problem in ML

- Results vary across runs
- Hidden changes break comparisons
- Hard to trust or verify models
- Difficult to share and collaborate
- "It worked yesterday, not today!"



Hidden Sources of Inconsistency

- Random seed not fixed
- Dataset version changes
- Different hyperparameters
- Library / version mismatches
- Hardware / environment drift
- Preprocessing variations
- Time & timezone differences

What the Hub Standardizes

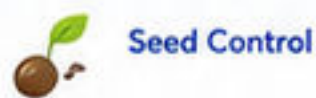
- Seed control and tracking
- Dataset versioning & hashing
- Parameter logging & tracking
- Comprehensive metrics capture
- Environment & dependency info
- Fingerprint for unique identity
- Shareable experiment packages

The Core Promise

Same setup.
Same evidence.
Clearer repeatability.



Why Reproducibility Matters



Seed Control

Fix randomness at the source to ensure runs are repeatable and comparable.

Seed



Dataset Versioning

Track immutable dataset versions and hashes to preserve data lineage.

Dataset Version



Parameter Tracking

Log every parameter and configuration for full experiment transparency.

Parameters



Metric Transparency

Capture all important metrics consistently for honest and apples-to-apples comparison.

Metrics



Environment Awareness

Record libraries, versions, OS, and environment to explain every result.

Fingerprint

Reliable science starts with reproducible evidence.

Without Reproducibility

- Results change, hard to trust
- Lost context & missing details
- Slow debugging & root cause
- Poor collaboration & sharing
- No proof, hard to validate

VS

With AutoML Reproducibility Hub

- Consistent, verifiable results
- Complete context captured
- Fast debugging & root cause
- Seamless sharing & collaboration
- Reproducible anywhere, anytime



Key Concepts Captured

- Seed
- Dataset Version
- Parameters
- Metrics
- Fingerprint
- Shareability

The Reproducible Experiment Lifecycle



Every run has a unique identity (Fingerprint) that proves what was run, where, and how.

Share experiments like a link. Anyone can reproduce the same result.



Reproducibility turns experiments into trustworthy knowledge.



Trust

Build confidence in every result.



Efficiency

Save time, reduce rework.



Collaboration

Share, review, and build together.



Compliance

Meet standards, ensure accountability.



Better science.
Better models.
Better decisions.



Inside AutoML

Reproducibility Hub

A Visual Guide to Reproducible Browser-Based Machine Learning



Architecture & Technology Stack

100% Browser-Native • Reproducible • Transparent • Shareable



Two Operating Modes



Static Reference Mode

Fast • Deterministic • Shareable

- ✓ Pre-computed demo results
- ✓ Zero code execution
- ✓ Instant loads, perfect for sharing
- ✓ Best for documentation & demos



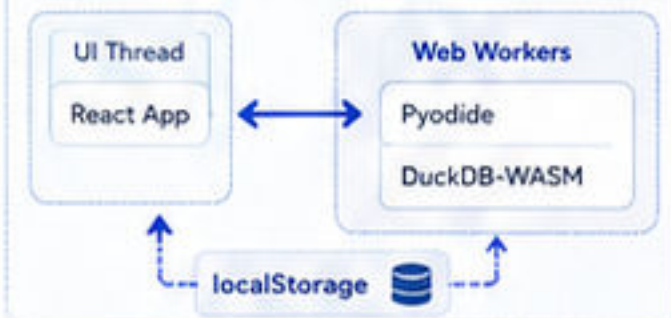
Full Browser Execution Mode

Extended Models • Advanced Features

- ✓ Runs Python via Pyodide in-browser
- ✓ Full pipelines & custom datasets
- ✓ Reproducible on any modern browser
- ✓ Best for exploration & learning

Data & Execution Isolation

Your Browser (Local)



Everything stays on your device.
No data leaves your browser.

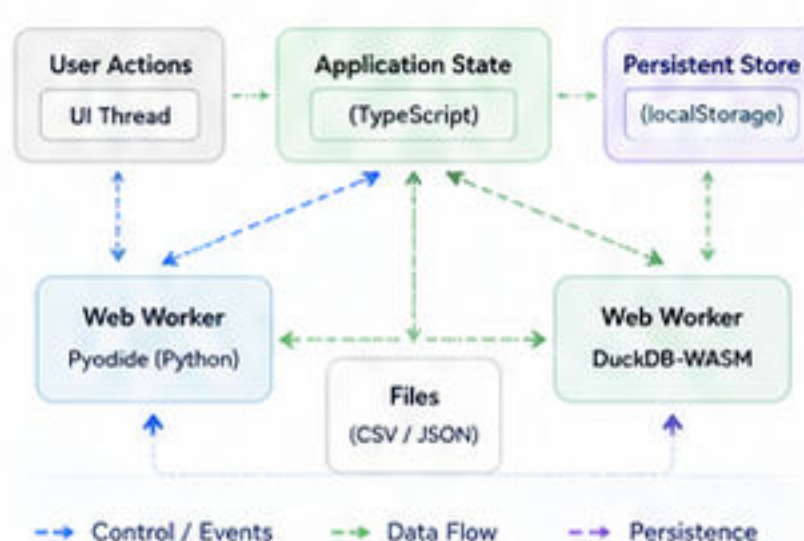
Key Data Artifacts

TS	TypeScript	Application code & logic
Python	Pyodide (Python)	ML algorithms & execution
DuckDB	DuckDB-WASM	SQL analytics & profiling
localStorage	localStorage	Persistent runs & settings
JSON	JSON Exports	Manifests, runs, configs
CSV	CSV Datasets	Tabular data & results
Web Workers	Web Workers	Isolated, parallel execution

System Flow (High Level)



Architecture at a Glance (Data & Execution Paths)



Why Browser-First?

- No project-owned backend**
Everything runs in the browser. Easier to maintain, deploy, and scale.
- No hidden server execution**
Transparent, inspectable, and trustworthy results.
- Portable demos**
Share a link, anyone can explore immediately.
- Reproducible education**
Same environment, same results—every time.
- Low deployment friction**
Static site on GitHub Pages. Fast, secure, global.



Explore the project on GitHub

github.com/FaramarzKowsari/automl-reproducibility-hub





Static Mode (S)

Fast • Deterministic • Shareable

Static Reference Mode delivers curated, precomputed experiment examples instantly — without running a new model in this browser session.



Important: Reference result — no model was executed in this browser session.

You are viewing static, precomputed results for reference and learning.

What is Static Mode?

Static Mode lets you explore reproducible results that were generated previously and stored in the repository. It loads instantly, uses minimal resources, and provides a reliable reference for comparison, learning, and communication.



Instant Startup

Loads in milliseconds.
No pipelines, no waiting.



No Heavy Packages

No model, dataset, or training libraries loaded.



Ideal for Teaching & Demos

Explain results clearly without compute.



Deterministic Reference Examples

Curated, versioned, and fully reproducible.



Metric Previews

See key metrics immediately.



Dataset Cards

Understand data versions at a glance.



Manifest Inspection

View experiment config and lineage.

How Static Mode Works



Curated & Precomputed Experiments

Generated offline and stored.



Stored with Manifest

Config, metrics, artifacts saved.



Loaded Instantly

Static files fetched from repository.



Render Reference UI

Explore metrics, results, and configs.



No Execution

No model run in this session.



Startup Timeline (Typical)

0ms

Click Static Mode

~50–150ms

Fetch manifest & reference files

~150–300ms

Parse & prepare UI data

~300–600ms

Render charts, cards & tables

< 1s

Ready to explore reference results

Static Mode vs Full Mode



Static Mode
(Reference)



Full Mode
(Execution)



Instant load
(milliseconds)

Slower startup
(seconds–minutes)



No packages loaded

Loads ML/DL packages



No GPU/CPU intensive work

High compute usage possible



Uses precomputed reference data

Runs new model in this session



Deterministic results

Results may vary with changes



Easy to share and embed

Not shareable until completed



Best for learning, demos, docs

Best for research, new experiments

Best Use Cases



Classroom Explanation

Teach concepts using stable, known results.



Repository Preview

Preview what's inside without any setup.



Quick Inspection

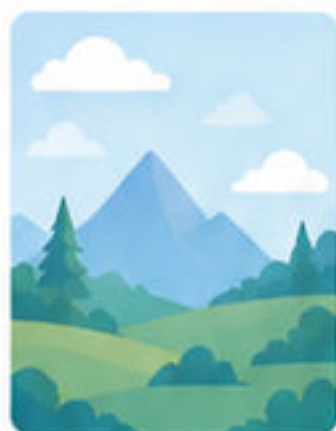
Check metrics, configs, and datasets fast.



Shareable Examples

Embed links to exact reference results.

Sample Static Experiment (Reference)



Experiment ID	a73fe9b2cd68f1e6
Seed	42
Dataset Version	v1.2.0
Dataset	XGBoost-Tabular
Model	XGBoost v1.2.0
Created	2024-05-28 14:32:10 UTC
Static Reference	No model executed in this session

Metrics (Preview)

Cross-Validation (5-Fold)

Accuracy	0.934
F1 Score	0.923
ROC AUC	0.912
RMSE	0.889

Average Rank
(Lower is better) 2.1 / 10

Manifest Snapshot

```
{
  "experiment_id": "a73fe9b2cd68f1e6",
  "seed": 42,
  "dataset": "XGBoost-Tabular",
  "dataset_version": "v1.2.0",
  "model": "XGBoost v1.2.0",
  "metrics": ["Accuracy", "F1 Score", "ROC AUC", "RMSE"],
  "cv": 5,
  "created": "2024-05-28T14:32:10Z",
  "mode": "static_reference"
}
```

[View Full Manifest](#)



Static Mode is perfect when you need speed, stability, and certainty.

Explore, learn, compare, and share — all without running a single model.



Explore the project on GitHub

github.com/FaramarzKowsari/automl-reproducibility-hub





Full Mode (F)

Full Browser Execution Mode

Full Mode runs the entire machine learning workflow inside your browser. AutoML loads **Pyodide** and a rich Python scientific stack to perform real training, evaluation, and comparison **locally**—no server required.



Your data stays in your browser.
100% private.

Powered by
Browser Python



Pyodide
Python in WebAssembly

NumPy

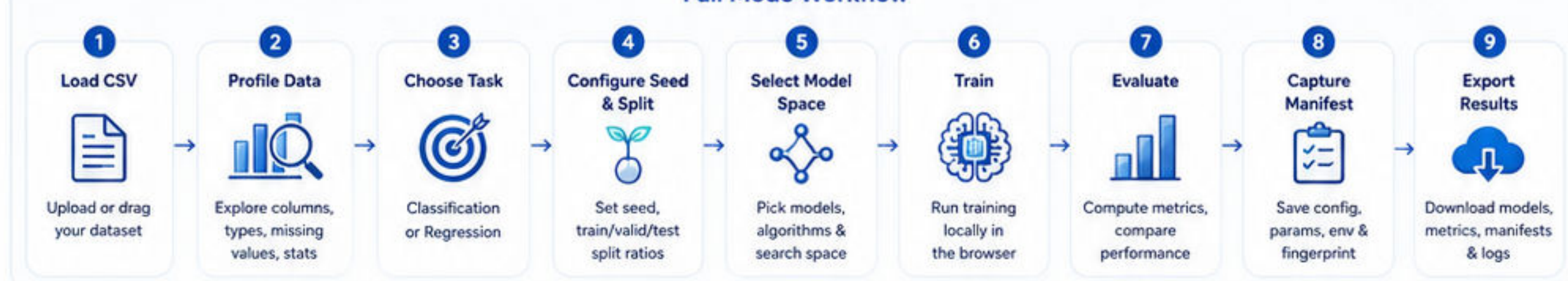
pandas

scikit-learn



DuckDB-WASM
(In-Browser Analytics)

Full Mode Workflow



Browser-Local Execution



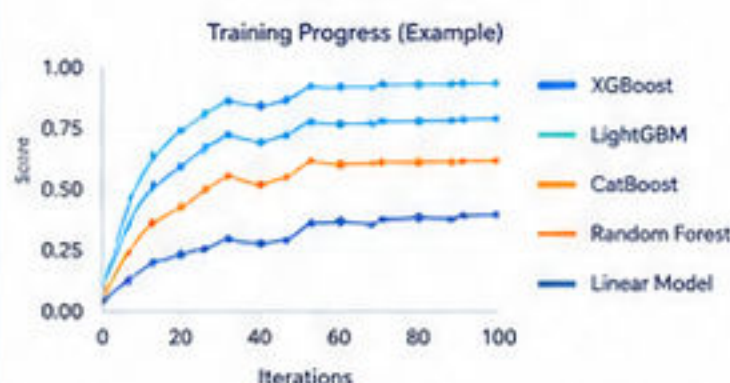
Loading Pyodide... ✓
Loading NumPy... ✓
Loading pandas... ✓
Loading scikit-learn... ✓
Loading DuckDB-WASM... ✓
Environment Ready! ✓

All computation happens locally using WebAssembly.
No data leaves your device.

Private • Secure • Offline-Ready

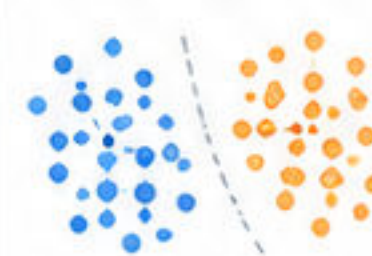
Real Training Runs

Full Mode performs real model training with cross-validation and hyperparameter search.



Classification & Regression Support

Classification (e.g. Titanic Survived)



Regression (e.g. House Prices)



Auto detects task type and evaluates with appropriate metrics.

AutoML Comparison Workflows

Compare many algorithms side-by-side with the same data split, seed and metrics.

Model	Search Type	Best Score (F1)
XGBoost	Bayesian	0.934
LightGBM	Bayesian	0.923
CatBoost	Bayesian	0.912
Random Forest	Randomized	0.889
Logistic Regression	Grid	0.812
KNN	Grid	0.735

✓ Fair Comparison: Same data, same seed, same metrics.

Metric Computation

Comprehensive metrics computed locally.

Classification		Regression	
Accuracy	0.934	RMSE	18.42
F1 Score	0.923	MAE	12.73
Precision	0.921	R ² Score	0.892
Recall	0.925	MAPE	11.34%
ROC AUC	0.971	MedAE	9.87
Log Loss	0.182	Expl. Var.	0.894

Local Export

Export everything you need, all generated in-browser.



Download All Artifacts

DuckDB-WASM Data Inspection

Fast SQL-powered exploration on your data—fully local.

<pre>SELECT age, COUNT(*) as cnt FROM df WHERE age IS NOT NULL GROUP BY age ORDER BY age;</pre>	<table><tr><th>age</th><th>cnt</th></tr><tr><td>20</td><td>52</td></tr><tr><td>21</td><td>63</td></tr><tr><td>22</td><td>47</td></tr><tr><td>23</td><td>59</td></tr><tr><td>24</td><td>61</td></tr><tr><td>25</td><td>58</td></tr><tr><td>...</td><td>...</td></tr></table>	age	cnt	20	52	21	63	22	47	23	59	24	61	25	58	...	...
age	cnt																
20	52																
21	63																
22	47																
23	59																
24	61																
25	58																
...	...																



Blazing fast
in-browser
analytics.

Reproducibility Fingerprint

A unique fingerprint is generated from your experiment to make results verifiable and reusable.



Fingerprint ID	f3e9a7b2c4d8e91f
Algorithm	XGBoost
Dataset	titanic-ml
Seed	42
Parameters Hash	a193f2...c8b1
Created	2024-05-28 14:32:10 UTC

✓ Reproducible

Trade-offs



Heavier Startup

Loading Python + libraries takes longer than Static Mode.



More Browser Memory

Large datasets or many models can use more RAM.



Richer Experimentation

Enables real training, more models, deep comparisons and exports.

Best for deep experimentation & full privacy.

Example: Train a Model in Full Mode

```
import pandas as pd
from sklearn.ensemble import RandomForestClassifier
from sklearn.model_selection import train_test_split
from sklearn.metrics import f1_score

df = pd.read_csv('titanic.csv')
X = df.drop('Survived', axis=1)
y = df['Survived']
X_train, X_test, y_train, y_test = train_test_split(
    X, y, test_size=0.2, random_state=42, stratify=y)
model = RandomForestClassifier(n_estimators=200, random_state=42)
model.fit(X_train, y_train)
preds = model.predict(X_test)
print('F1 Score:', f1_score(y_test, preds))
```

F1 Score:
0.932

Mode Reminder

Full Mode = Full Power



- ✓ Real training
- ✓ Rich metrics
- ✓ Deep comparisons
- ✓ Export everything
- ✓ 100% local & private

When to Use Full Mode

- You need real training and hyperparameter search
- You want to compare many models deeply
- Your data is sensitive and must stay local
- You want to export models, metrics & manifests
- You value reproducibility and transparency



Explore the project on GitHub

github.com/FaramarzKowsari/automl-reproducibility-hub



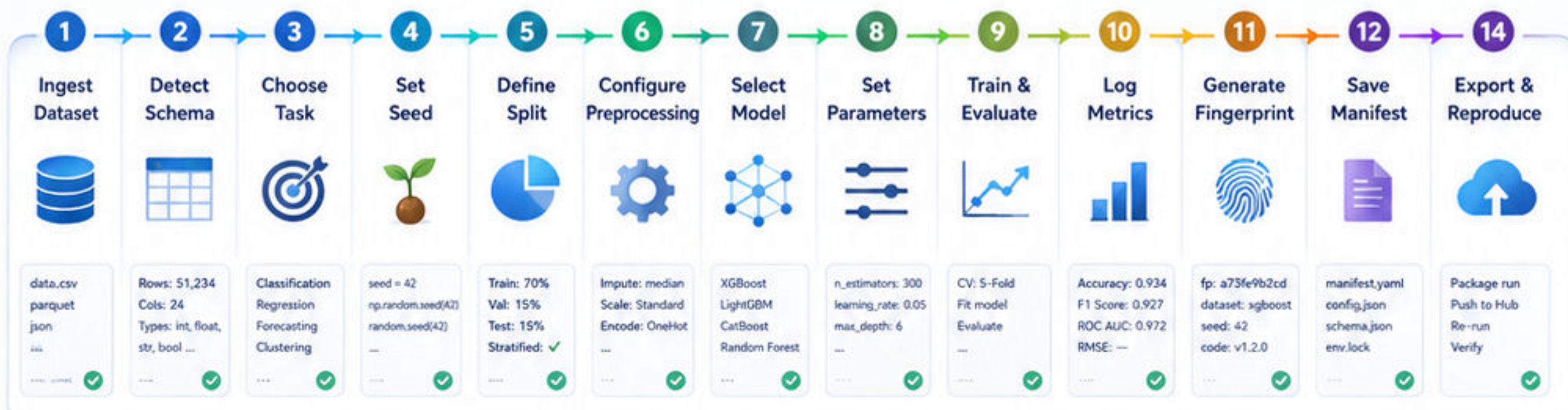


Experiment Lifecycle

End-to-End Reproducibility Pipeline



From raw data to reproducible results — every decision captured, every run repeatable.



What Happens at Each Stage

#	Stage	What Is Captured	Why It Matters
1	Ingest Dataset	Source path/URL, file checksum, format, row count	Proves the exact data you used
2	Detect Schema	Columns, types, nulls, stats, inferred schema	Locks structure & meaning
3	Choose Task	Problem type, target(s), metric intent	Clarifies goal & evaluation
4	Set Seed	Global seed for all libs (Python, NumPy, etc.)	Controls randomness
5	Define Split	Split strategy, ratios, stratification, groups	Ensures identical data partitions
6	Configure Preprocessing	Transform steps & fitted artifacts	Recreates identical feature pipeline
7	Select Model	Algorithm & library versions	Defines model family & implementation
8	Set Parameters	Hyperparameters & search space	Enables exact reconfiguration
9	Train & Evaluate	Training config, CV folds, early stopping, scores	Reproduces the training process
10	Log Metrics	All metrics, curves, confusion matrix, residuals	Transparent performance evidence
11	Generate Fingerprint	Content hash of data, code, params, env	Unique identity for the run
12	Save Manifest	Immutable manifest (YAML/JSON) of everything	Single source of reproducibility truth
13	Export & Reproduce	Package & publish run, re-run in clean env	Anyone can reproduce the same result



Classification Workflows

Data + Features → Model → Predict



Metrics

Accuracy F1 Score ROC AUC PR AUC Log Loss



Regression Workflows

Data + Features → Model → Predict



Metrics

RMSE MAE R² MAPE Expl. Var



Leaderboard Comparison

Rank	Model	Dataset	Metric (ROC AUC)	Score
1	XGBoost v1.2.0	CreditRisk	ROC AUC	0.972
2	LightGBM v1.2.0	CreditRisk	ROC AUC	0.965
3	CatBoost v1.1.0	CreditRisk	ROC AUC	0.958



Experiment History

2024-05-28 14:32	xgboost_v1.2.0	0.972
2024-05-27 09:11	lightgbm_v1.2.0	0.965
2024-05-26 16:45	catboost_v1.1.0	0.958
2024-05-25 10:02	rf_baseline_v1.0	0.912



What Makes a Run Reproducible?

- ✓ **Data Integrity**: Immutable data + checksums
- ✓ **Determinism**: All randomness controlled by seeds
- ✓ **Environment Lock**: Exact library & system versions
- ✓ **Complete Config**: Every parameter recorded
- ✓ **Traceability**: Unique fingerprint for every run
- ✓ **Verifiability**: Anyone can re-run & match results

★ If any one of the above is missing, the run is not truly reproducible.



Explore the project on GitHub

github.com/FaramarzKowsari/automl-reproducibility-hub





Manifests, Versions & Fingerprints

How AutoML Reproducibility Hub Records Reproducibility Metadata

Every experiment is captured as a manifest — a complete, machine-readable snapshot of data, parameters, environment, metrics, and lineage — and turned into a SHA-256 fingerprint that becomes its unique, tamper-evident identity.



Experiment Manifest (Reproducibility Record)

Seed	42	Parameters	{max_depth: 6, n_estimators: 200, learning_rate: 0.05, subsample: 0.8}
Dataset Name	Titanic	Metrics	Accuracy: 0.842 ROC AUC: 0.912
Dataset Version	v1.2.0	Runtime Versions	Python: 3.12.2, scikit-learn: 1.4.2, xgboost: 2.0.3, pandas: 2.2.1
SHA-256 (Dataset)	3f8c...a9d2	App Version	automl-hub: 1.2.0
Rows	891	Timestamp (UTC)	2024-05-28T14:32:10Z
Columns	12	Experiment Fingerprint	a73fe9b2c0dd8f1e6 98bde2...73b3
Target	Survived		Reproducible
Task Type	Binary Classification		
Split	Train: 80% Test: 20%		
Preprocessing	Impute(median) + OneHot		
Model	XGBoost		

From Manifest to Identity

1. Canonical JSON Representation

The manifest is serialized into a canonical JSON (sorted keys, stable formatting).

```
{
  "seed": 42,
  "dataset_name": "Titanic",
  "dataset_version": "v1.2.0",
  "dataset_sha256": "3f8c...a9d2",
  "rows": 891,
  "columns": 12,
  "target": "Survived",
  "task_type": "Binary Classification",
  "split": {"train": 0.8, "test": 0.2},
  "preprocessing": "Impute(median)+OneHot",
  "model": "XGBoost",
  "parameters": {"max_depth": 6,
    "n_estimators": 200, "learning_rate": 0.05,
    "subsample": 0.8},
  "metrics": {"accuracy": 0.842,
    "roc_auc": 0.912},
  "runtime_versions": {"python": "3.12.2",
    "xgboost": "2.0.3", "sklearn": "1.4.2",
    "pandas": "2.2.1"},
  "app_version": "1.2.0",
  "timestamp_utc": "2024-05-28T14:32:10Z"
}
```

2. SHA-256 Hash

A SHA-256 digest is computed from the canonical JSON.

SHA-256

a73fe9b2c0dd8f1e6
98bde2...73b3

3. Stable, Unique Identity

Any change in the manifest produces a new fingerprint.

Unique Experiment Identity

What Goes Into the Fingerprint?



Data

- Dataset name
- Version
- SHA-256
- Shape
- Target



Parameters

- Model
- Hyperparameters
- Preprocessing
- Split strategy
- Seed



Runtime

- Python version
- Library versions
- System info
- App version



Metrics

- Evaluation metrics
- Scores
- Cross-validation
- Details



SHA-256 over
canonical JSON

Unique Experiment Identity



Why Dataset Versioning Matters

- Data changes silently over time.
- Versioning + dataset hash guarantee you know exactly which data you used.
- Enables reliable reruns and audits.



Why Hashes Matter

- Hashes are tamper-evident and compact.
- They provide a stable, content-based identity across systems.
- Perfect for storage, indexing, and integrity checks.



How Fingerprints Help

- Quickly compare experiments.
- Detect identical runs instantly.
- Re-run anywhere, get the same result.
- Great for tracking, sharing, and publishing.



Any change to seed, dataset hash, parameters, or runtime versions creates a new identity.



Change Seed



Change Dataset
/ Hash



Change
Parameters



Change Runtime
Versions



New Experiment
Identity



Explore the project on GitHub

github.com/FaramarzKowsari/automl-reproducibility-hub



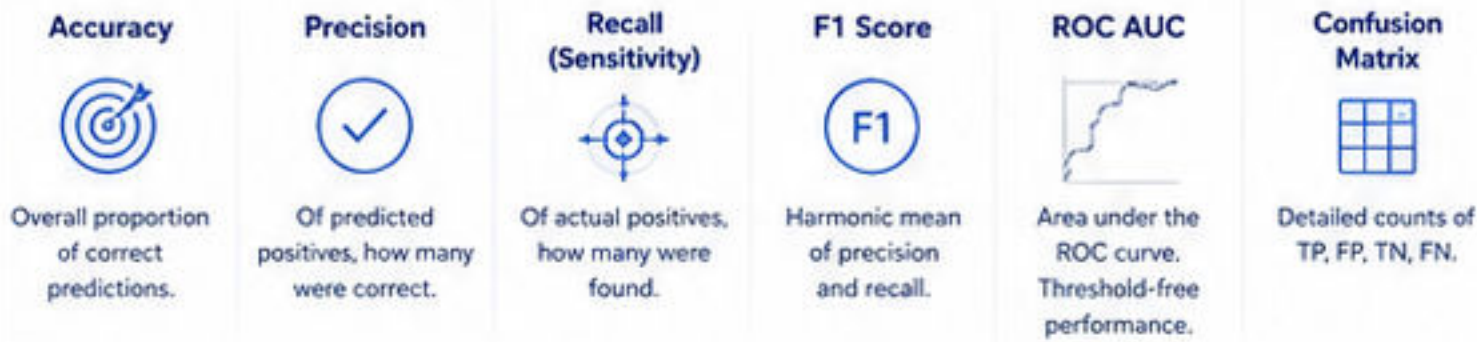


Metrics, Comparison & Outputs

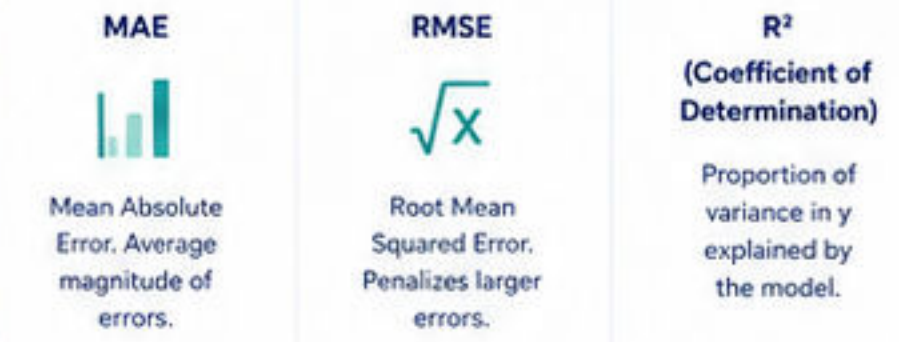
AutoML Reproducibility Hub evaluates experiments with standard metrics, compares models side-by-side, and provides rich visualizations and exportable artifacts so results are transparent, repeatable, and easy to share.



Classification Metrics



Regression Metrics



↓ Lower is better

↑ Higher is better



Comparison Dashboard

Compare experiments across metrics, datasets, and time.

Top Metric (Primary)
ROC AUC

Dataset
All

Models
All

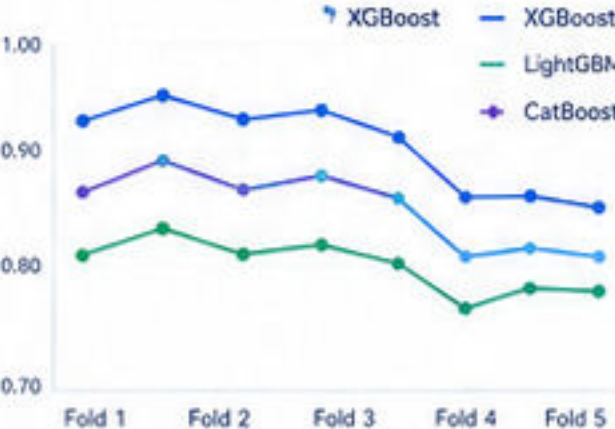
Runs
All

Sorted by
ROC AUC

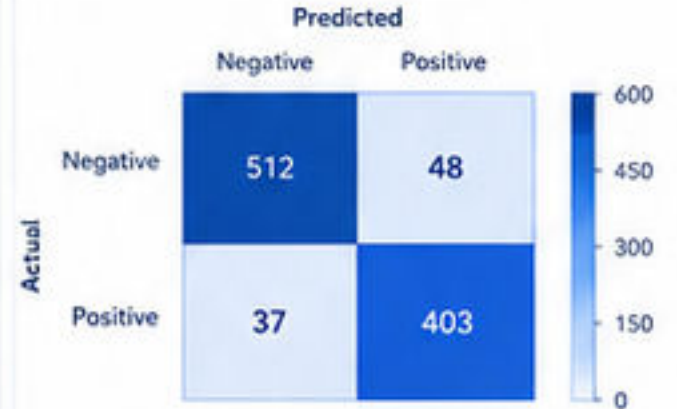
Leaderboard (Top 6)

Rank	Model	Dataset	ROC AUC ↑	F1 Score	Accuracy
1	XGBoost	v1.2.0	0.934	0.912	0.926
2	LightGBM	v1.2.0	0.923	0.901	0.915
3	CatBoost	v1.1.0	0.912	0.889	0.904
4	Random Forest	v1.0.0	0.889	0.862	0.879
5	Logistic Regression	v1.0.0	0.842	0.801	0.817
6	SVM (RBF)	v1.0.0	0.813	0.771	0.789

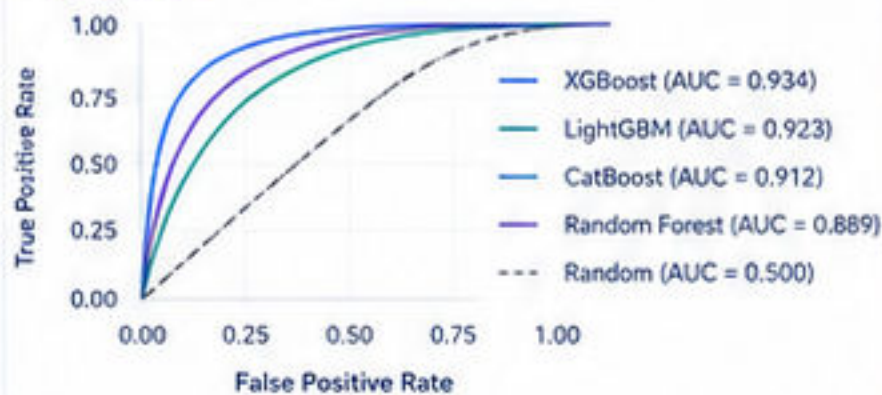
Metric Trends (Cross-Validation)



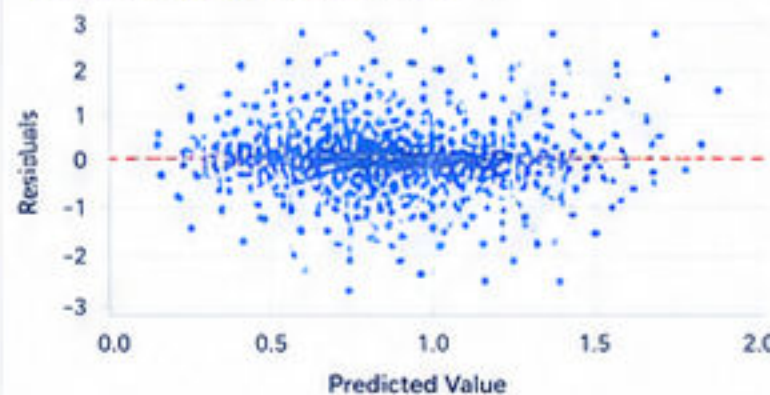
Confusion Matrix (Best Model: XGBoost)



ROC Curves



Residual Plot (Best Model: XGBoost)



Use the dashboard to quickly identify the best-performing model, understand errors, and make confident decisions.

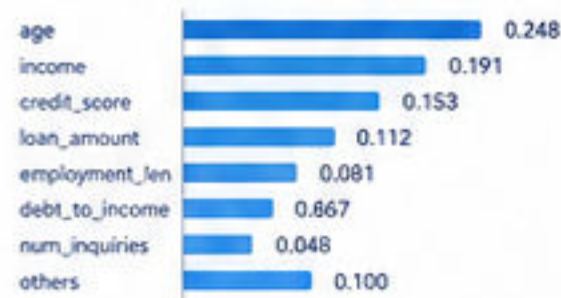
- ✓ Rank by any metric
- ✓ Filter by dataset, model, or date range
- ✓ Drill into runs for full experiment details
- ✓ Export charts and tables



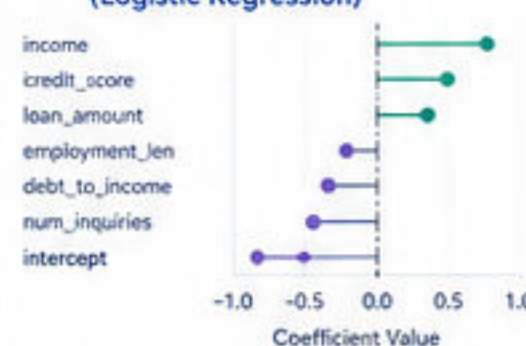
Feature Importance / Coefficient Interpretation

Understand what drives the model's predictions.

Tree-Based (XGBoost)



Linear Model Coefficients (Logistic Regression)



Positive coefficient/importance increases the prediction.



Negative coefficient/importance decreases the prediction.



Magnitude shows strength of impact.



Outputs You Can Export



Reproducibility Score



How it works

- ✓ Documentation completeness
- ✓ Code availability
- ✓ Data accessibility
- ✓ Environment capture
- ✓ Parameter logging
- ✓ Metric tracking

Higher scores mean higher trust and easier reuse.

Well-documented experiments are easier to verify, share, and build upon.



Explore the project on GitHub

github.com/FaramarzKowsari/automl-reproducibility-hub





Research Value, Privacy & Deployment

Why This Repository Matters

Transparent • Shareable • Reproducible • Browser-First

S Static Mode
Fast • Deterministic
Shareable

F Full Mode
Extended Models
Advanced Features

1 Research & Teaching Value

- Reproducible experiments in a transparent environment
- Compare models, metrics, and datasets side by side
- A living resource for lectures, assignments & workshops
- Bridge theory to practice with interactive demos

3 Browser-Local Workflow

- Load Data (CSV/JSON)
 - Preprocess
 - Train & Tune Models
 - Evaluate & Compare
 - Export Results (Local)
- All computation happens in your browser.
No installs. No servers. No upload.

Central Ecosystem: From Code to Reproducible Impact



2 Privacy-by-Design

- 100% browser-local processing
 - No data leaves your device
 - No analytics, no tracking
 - Datasets stay local (uploaded files)
 - Share results, not data
- Your data. Your device.
Your control.

4 Limitations (By Design)

- Browser memory constraints**
Very large datasets may not fit in memory.
- Full Mode startup is heavier**
Loads more models and modules for advanced features.
- Floating-point tolerance**
Minor numeric differences may occur across environments.
- Educational transparency**
Built for learning & research, not for production-scale serving.

GitHub Actions CI

- Runs on every push & PR
- TypeScript build & lint
- Python (Pyodide) tests
- Unit & integration tests
- Artifacts & logs for visibility

Status **passing**

GitHub Pages Deployment

- Automatic deployment from main branch
- Live demos & guidebook
- Fast, global CDN
- Always up-to-date
- Custom domain ready

Live Site **online**

Open-Source Documentation

- README for quick start
- Docs for deep reference
- Inline code comments
- Architecture & design notes
- Issue templates for help

Docs **available**

Citation Readiness

- CITATION.cff included
- Zenodo-ready metadata
- DOI minting in one click
- Clear authorship & license
- Reproducible references

CITATION.cff **ready**

Reproducible Sharing

- Share results, not data
- Export metrics & charts
- Config JSON export
- Versioned releases
- Re-run anywhere, same outcome

Reproducible **by design**

Why This Repository Matters

Empowering reproducible ML experiments for everyone.

Scholars

- Reproduce results with confidence
- Compare methods fairly
- Publish with transparency

Instructors

- Interactive demos for classes
- Assignments with instant feedback
- Lower barrier to experimentation

Developers

- Clean, modular, typed codebase
- Extensible & well-documented
- CI, tests & modern tooling

Students

- Learn by doing
- Explore models visually
- Build real ML intuition

Publication & Citation Ready

README.md
Overview, setup, features, usage

CITATION.cff
Standard citation metadata

Zenodo-Ready
One-click DOI & metadata

Guidebook
Visual guide series (this project)

Cite. Share. Advance Science.





About the Guidebook

Browser-Based ML Stack



TypeScript
Robust Frontend Logic



Pyodide
Python in Web (WASM)



DuckDB-WASM
In-Process Analytics

All running in your browser. No install. No server. 100% local.



Inside AutoML Reproducibility Hub is a ten-page visual handbook that serves as the official guidebook for the GitHub repository [AutoML Reproducibility Hub](#).

This guide explains the project's purpose, architecture, Static mode (S) and Full mode (F), seed control, dataset versioning, parameter tracking, experiment fingerprints, metrics, and reproducibility workflows. It also introduces the browser-based machine learning stack built with TypeScript, Pyodide, and DuckDB-WASM for transparent, shareable, and reproducible AutoML experiments.



Deterministic Randomness



Versioned Datasets



Tracked Parameters



Experiment Fingerprints



Comprehensive Metrics



Reproducible Results



“

From experiment to explanation — reproducible by design.

”



About the Author



Faramarz Kowsari is an author, Software Engineer and AI researcher based in Istanbul. Focusing on the intersection of technology, education, and personal growth, he has published over 80 digital titles on international platforms. His areas of expertise span Artificial Intelligence, prompt engineering, modern trading strategies (Smart Money Concepts & algorithmic trading), as well as classical literature and mindfulness. In addition to writing, he develops web-based educational tools and creates specialized instructional video content.



Official Profiles



ORCID: 0000-0003-1692-0453



Google Scholar scholar.google.com/citations?user=FaramarzKowsari



GitHub github.com/FaramarzKowsari



LinkedIn linkedin.com/in/faramarzkowsari



Google Books [books.google.com /Faramarz_Kowsari](https://books.google.com/Faramarz_Kowsari)



Personal Website FaramarzKowsari.github.io



Zenodo [zenodo.org /search?ln=en&cc=Faramarz%20Kowsari](https://zenodo.org/search?ln=en&cc=Faramarz%20Kowsari)

Reproducible by Design



1. Run Experiment
Configure data, models, parameters, and seed.



2. Capture Fingerprint
System, data, parameters, and environment recorded.



3. Evaluate & Compare
Metrics and visual analytics for clear comparison.



4. Reproduce Anywhere
Same setup, same result — on any supported browser.



Official guidebook for:

github.com/FaramarzKowsari/automl-reproducibility-hub

