

TREE-BASED METHODS

FROM BEGINNER TO EXPERT

A complete, visual, application-driven course — from a single decision tree through random forests and gradient boosting (XGBoost, LightGBM, CatBoost), with a capstone on how trees power Machine Learning, Deep Learning, and AI.

What you will be able to do

Build and prune decision trees for classification and regression; understand impurity, information gain, and variance reduction; tame variance with bagging and random forests; master boosting and the modern gradient-boosting libraries; interpret models with importance, permutation, and SHAP; and explain why trees still dominate tabular machine learning and how they meet deep learning and AI.

Format: self-paced reference & course | **Prerequisites:** basic statistics, a little probability
Part of the Comprehensive Machine Learning & Deep Learning Curriculum.

Preface: how to use this book

There is a quiet truth in applied machine learning: while deep networks dominate images, audio, and text, the workhorse for the *tables* that run businesses — spreadsheets of customers, transactions, sensors, and records — is a forest of decision trees. Gradient-boosted trees win the majority of tabular competitions and power countless production systems precisely because they fit the messy, irregular structure of real-world tabular data better than almost anything else. This book builds tree-based methods from a single “yes/no” question all the way to the libraries (XGBoost, LightGBM, CatBoost) behind that dominance.

Who this is for

This book starts from the very beginning — what it means to split data with a question — and climbs to ensembles, gradient boosting, model interpretation, and the research frontier where trees meet deep learning. It is for the self-learner who wants both **fluency** (“I can train, tune, and interpret a gradient-boosting model and trust it”) and **understanding** (“I know why a single tree overfits, why averaging and boosting fix it, and why trees beat neural nets on tables”). Every major idea gets a picture, a worked example, and a forward link to machine learning.

The central idea

All of tree-based learning grows from one move and its consequences:

- **Partition, then predict.** A tree recursively splits the feature space into boxes with simple yes/no questions, then predicts a constant in each box — the conditional mean or majority class.
- **One tree is weak; many are strong.** A single deep tree overfits (high variance). *Averaging* independent trees (bagging, random forests) cuts variance; *adding* corrective trees (boosting) cuts bias.
- **Trees fit tabular reality.** Axis-aligned splits handle mixed types, missing values, outliers, and irregular, non-smooth relationships with almost no preprocessing.

The structure

- **Part I — Foundations** (beginner): what tree-based methods are; classification trees; regression trees.
- **Part II — Building Trees Well** (core): growing, splitting, and pruning; the strengths, weaknesses, and bias–variance of trees.
- **Part III — Ensembles** (advanced): bagging and random forests; boosting and gradient boosting; the modern libraries.
- **Part IV — Interpretation & Practice** (advanced/expert): feature importance and SHAP; practical modeling end to end.
- **Part V — Frontier:** specialized trees; trees versus the rest of ML; and a capstone on trees in ML/DL/AI.

How to read it

Read with a pen and a computer. When you meet a *Definition*, restate it; when you meet an *Example*, work it before reading on; when you meet a *Try it in code* box, run it. The colored boxes recur throughout:

- **Key idea** — the one sentence to remember.
- **Intuition** — the picture or analogy behind the math.
- **Common pitfall** — the mistakes that bite everyone.
- **How this powers ML / AI** — the forward link to applications.
- **Try it in code** — a hands-on check in Python (scikit-learn / XGBoost).

Key idea

Hold one picture above all others: a tree **carves the feature space into rectangles with simple questions and predicts a constant in each**. A single tree is interpretable but unstable; the power comes from *combining many trees*. Bagging and random forests average trees to kill variance; boosting stacks trees to kill bias. Master the single tree, and forests, XGBoost, and the tabular state of the art are the same idea, scaled.

Contents

Preface: how to use this book	1
I Foundations	6
1 What Are Tree-Based Methods?	7
1.1 The flowchart model	7
1.2 Recursive partitioning	7
1.3 A short history	8
1.4 Why trees matter	8
2 Decision Trees for Classification	10
2.1 Impurity: how mixed is a node?	10
2.2 Information gain: scoring a split	10
2.3 Greedy recursive growth	11
2.4 Probabilities, not just labels	11
3 Decision Trees for Regression	13
3.1 Splitting to reduce variance	13
3.2 Reading a regression tree	14
3.3 The extrapolation weakness	14
II Building Trees Well	16
4 Growing, Splitting, and Pruning	17
4.1 The CART growth algorithm	17
4.2 Pre-pruning: stop early	17
4.3 Cost-complexity (post-)pruning	18
4.4 Categorical splits and missing values	18
5 Strengths, Weaknesses, and the Bias–Variance of Trees	20
5.1 What trees do well	20
5.2 The core weakness: high variance	20
5.3 Other limitations	21
5.4 From one tree to many	21
III Ensembles	23
6 Bagging and Random Forests	24
6.1 Bagging: bootstrap aggregating	24
6.2 Random forests: decorrelate the trees	24
6.3 Out-of-bag (OOB) error: free validation	25
6.4 Why forests are such a good default	25

7	Boosting: AdaBoost and Gradient Boosting	27
7.1	The boosting idea	27
7.2	AdaBoost: reweight the hard cases	27
7.3	Gradient boosting: fit the residuals	28
7.4	Regularization: learning rate, subsampling, depth	28
8	Modern Gradient Boosting: XGBoost, LightGBM, CatBoost	30
8.1	XGBoost: regularized, second-order boosting	30
8.2	LightGBM: histograms, leaf-wise, GOSS, EFB	31
8.3	CatBoost: ordered boosting and native categoricals	31
8.4	Choosing a library	32
IV	Interpretation and Practice	34
9	Feature Importance and Interpretability	35
9.1	Built-in (impurity / gain) importance	35
9.2	Permutation importance	35
9.3	SHAP: consistent, local attributions	36
9.4	Partial dependence and ICE	36
10	Practical Tree Modeling	38
10.1	Categorical features	38
10.2	Missing values	38
10.3	Class imbalance	38
10.4	Hyperparameter tuning	39
10.5	Pipelines, validation, and leakage	39
10.6	Metrics and calibration	39
V	Frontier and Applications	41
11	Specialized and Advanced Trees	42
11.1	Extremely randomized trees (ExtraTrees)	42
11.2	Isolation Forests: anomaly detection	42
11.3	Gradient-boosted trees for ranking	43
11.4	Survival, quantile, and probabilistic trees	43
11.5	Bayesian and oblique trees	43
11.6	Stacking and blending	43
12	Trees versus the World: When to Use What	45
12.1	Trees vs linear models	45
12.2	Why trees beat deep nets on tabular data	45
12.3	When deep learning wins on tables	46
12.4	A decision guide	46
13	Tree-Based Methods in Machine Learning, Deep Learning, and AI	48
13.1	The reign over tabular machine learning	48
13.2	Trees meet deep learning: hybrids	48
13.3	Differentiable and “deep” trees	49
13.4	Tabular deep learning and foundation models	49
13.5	Trees as the interpretability and reasoning layer	49
13.6	Where trees sit in the AI stack	50

A	Algorithm and Hyperparameter Reference	51
A.1	Impurity and split criteria	51
A.2	Gradient boosting in one page	51
A.3	Hyperparameter cheat-sheet	52
A.4	Method selection at a glance	52
A.5	Key references	52

PART I

Foundations

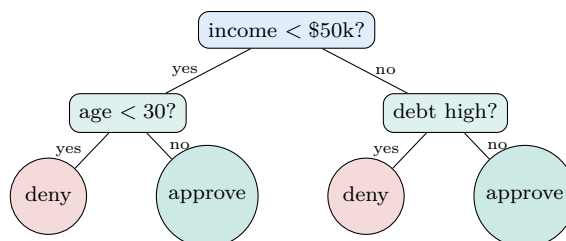
What Are Tree-Based Methods?

Level: Beginner

A decision tree is the most human of machine-learning models: it is a flowchart of yes/no questions, exactly how a doctor triages a patient or a loan officer screens an application. This chapter introduces the core idea — *recursive partitioning* — and the vocabulary of nodes, splits, and leaves, then previews how stacking and averaging many trees produces the methods that dominate tabular machine learning.

1.1 The flowchart model

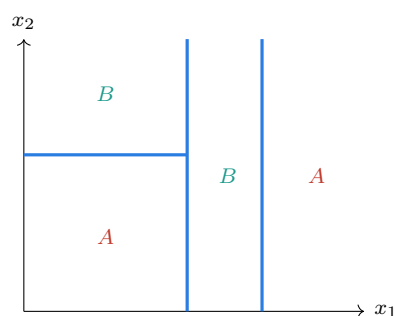
Definition 1.1. A **decision tree** predicts an output by routing an input down a series of tests. Each **internal node** asks a yes/no question about one feature (“is age < 45?”); each **branch** is an answer; each **leaf** holds a prediction — a class label (classification) or a number (regression). The path from **root** to leaf is a chain of conditions that defines a region of feature space.



1.2 Recursive partitioning

Key idea

A tree **recursively partitions** the feature space into axis-aligned rectangles (“boxes”), then predicts a *constant* within each box — the majority class or the average target. Growing a tree means repeatedly asking: *which single yes/no question splits this region into the two purest possible sub-regions?* Each split is a straight cut perpendicular to one feature axis; stacking cuts builds an increasingly fine staircase that approximates the true relationship.



each region $\rightarrow$ one constant prediction

Intuition

The tree and the partition are two views of the same object. Every leaf corresponds to one rectangle; every split corresponds to one cut. This is why trees handle *interactions* naturally: a split on x_2 that occurs only inside the “ $x_1 < 2.4$ ” branch means the effect of x_2 depends on x_1 — something a linear model must be told about explicitly, but a tree discovers automatically.

1.3 A short history

- **ID3 (1986)** and **C4.5 (1993)**, by Quinlan, grew trees using *information gain* (entropy); C4.5 added pruning and continuous features.
- **CART (1984)**, by Breiman, Friedman, Olshen, and Stone, unified classification and regression with *binary* splits, the *Gini* criterion, and *cost-complexity pruning*. This is what scikit-learn, XGBoost, and LightGBM implement — when practitioners say “decision tree,” they almost always mean CART.
- **Ensembles**: bagging (1996) and random forests (2001) by Breiman; AdaBoost (1997) by Freund and Schapire; gradient boosting (2001) by Friedman; and the modern libraries **XGBoost** (2016), **LightGBM** (2017), and **CatBoost** (2018).

1.4 Why trees matter

Key idea

A single tree is interpretable but **unstable** and weak — a small data change can reshuffle the whole tree, and a shallow tree underfits while a deep one overfits. The breakthrough is **ensembling**: combine many trees so their errors cancel. *Bagging* and *random forests* average many independent trees to cut variance; *boosting* adds trees sequentially, each fixing the last one’s mistakes, to cut bias. Ensembles of trees are, for tabular data, the most reliably accurate models we have.

How this powers ML / AI

Tree ensembles are the quiet champions of *tabular* machine learning — the spreadsheets and databases that run finance, healthcare, advertising, fraud detection, and recommendation. Gradient-boosted decision trees (XGBoost, LightGBM, CatBoost) win the majority of tabular Kaggle competitions and remain the pragmatic default in industry, often outperforming deep neural networks on structured data while training far faster (Chapter 12, 13). They need almost no preprocessing — no scaling, native handling of mixed types and missing values — which makes them the first model most practitioners reach for on a new

table.

Try it in NumPy

```
1 from sklearn.tree import DecisionTreeClassifier, export_text
2 from sklearn.datasets import load_iris
3 X, y = load_iris(return_X_y=True)
4 tree = DecisionTreeClassifier(max_depth=2).fit(X, y)
5 print(export_text(tree, feature_names=load_iris().feature_names)) #
   read the flowchart
6 print("accuracy:", round(tree.score(X, y), 3))
```

Chapter summary

- A **decision tree** is a flowchart of yes/no feature tests; root-to-leaf paths define regions, and each leaf predicts a constant.
- Trees work by **recursive partitioning** into axis-aligned boxes; splits capture feature **interactions** automatically.
- The dominant algorithm is **CART** (binary splits, Gini/variance, pruning); ID3/C4.5 are the entropy-based ancestors.
- A single tree is interpretable but unstable; the power comes from **ensembles** — bagging/-forests (variance) and boosting (bias).
- Tree ensembles dominate **tabular** ML, with minimal preprocessing — the workhorse behind much of applied data science.

Decision Trees for Classification

Level: Beginner

To grow a classification tree we need a precise answer to one question: *what makes a split good?* The answer is **purity** — a good split produces child nodes whose samples are as close to a single class as possible. This chapter defines the two purity measures (Gini and entropy), shows how information gain chooses splits, and walks through building a tree by hand.

2.1 Impurity: how mixed is a node?

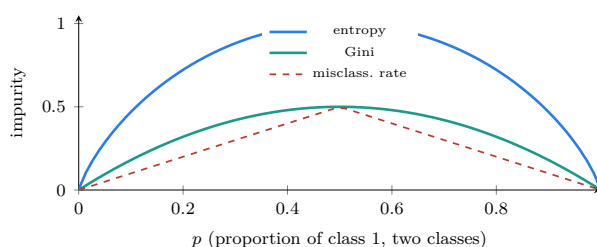
Definition 2.1. For a node with class proportions $p_1, \dots, p_K$, the two standard **impurity** measures are

$$\text{Gini} = 1 - \sum_{k=1}^K p_k^2, \quad H = - \sum_{k=1}^K p_k \log_2 p_k.$$

Both are 0 when the node is pure (one class) and maximal when classes are evenly mixed.

Intuition

Impurity measures disorder. A node of all one class has nothing to disagree about (impurity 0); a 50/50 node is maximally confused. **Gini** is the probability that two random draws from the node have different labels; **entropy** is the average number of bits to encode the label. They almost always pick the same splits — Gini is slightly faster (no logarithm) and is scikit-learn's default; entropy comes from information theory and underlies ID3/C4.5.



2.2 Information gain: scoring a split

Key idea

A split sends fractions of the parent's samples to child nodes. Its **information gain** is the

impurity removed — parent impurity minus the weighted average of child impurities:

$$\text{Gain} = I(\text{parent}) - \sum_{c \in \text{children}} \frac{n_c}{n} I(c),$$

where I is Gini or entropy and n_c/n is the fraction of samples in child c . The tree-growing algorithm tries every feature and every threshold and **greedily picks the split with the largest gain**.

Example 2.2. A node has 40 positives and 10 negatives, so $p = 0.8$ and $\text{Gini} = 1 - (0.8^2 + 0.2^2) = 0.32$. A candidate split yields children $\{30^+, 2^-\}$ and $\{10^+, 8^-\}$, with Gini $1 - ((\frac{30}{32})^2 + (\frac{2}{32})^2) = 0.117$ and $1 - ((\frac{10}{18})^2 + (\frac{8}{18})^2) = 0.494$. The weighted child impurity is $\frac{32}{50}(0.117) + \frac{18}{50}(0.494) = 0.253 < 0.32$, so the gain is 0.067 and the split is worth making.

2.3 Greedy recursive growth

Intuition

The algorithm is simple and greedy: at each node, scan all features and split points, choose the highest-gain split, partition the data, and recurse on each child — stopping when a node is pure, too small, or a depth limit is hit. “Greedy” means it never looks ahead; it takes the locally best split, which is fast but not globally optimal (finding the optimal tree is NP-hard). In practice the greedy choice works remarkably well.

2.4 Probabilities, not just labels

A leaf can report class *proportions*, not just the majority label, giving predicted probabilities $\hat{p}_k$. These power ROC/AUC analysis and threshold tuning — but, as the pitfall warns, raw tree probabilities are often poorly calibrated.

Common pitfall

Three traps. (1) **Greedy myopia**: the best first split is not always part of the best tree; a feature that only helps *after* another split can be overlooked. (2) **Bias toward many-valued features**: information gain favors features with many distinct values (an ID column gives “perfect” but useless splits) — C4.5’s *gain ratio* and CART’s binary splits mitigate this. (3) **Uncalibrated probabilities**: a pure leaf claims probability 1.0 even from two samples; trust the predicted class more than the exact probability, and calibrate (Platt/isotonic) if you need reliable probabilities.

How this powers ML / AI

The impurity measures here are the same ones that score splits inside every random forest and gradient-boosting library — billions of Gini/entropy evaluations underpin tabular ML. Entropy and information gain are the bridge to information theory used throughout deep learning (cross-entropy loss, KL divergence). And the “predict class proportions in a region” idea reappears as the leaf values in gradient-boosted trees, where leaves hold not class votes but real-valued scores fit to gradients (Chapter 7).

Try it in NumPy

```

1 import numpy as np
2 def gini(y):
3     _, c = np.unique(y, return_counts=True); p = c/c.sum()
4     return 1 - (p**2).sum()
5 def best_split(x, y):
6     best = (None, -1)
7     for t in np.unique(x):
8         left, right = y[x <= t], y[x > t]
9         if len(left)==0 or len(right)==0: continue
10        g = (len(left)*gini(left)+len(right)*gini(right))/len(y)    #
        weighted child Gini
11        gain = gini(y) - g
12        if gain > best[1]: best = (t, gain)
13    return best
14 x = np.array([1,2,3,4,5,6]); y = np.array([0,0,0,1,1,1])
15 print(best_split(x, y))    # threshold ~3 gives the perfect split (
    gain = 0.5)

```

Chapter summary

- A split is good if it increases **purity**; measure impurity with **Gini** $= 1 - \sum p_k^2$ or **entropy** $= -\sum p_k \log_2 p_k$.
- **Information gain** = parent impurity – weighted child impurity; the tree greedily picks the highest-gain split.
- Growth is **greedy and recursive** — fast, locally optimal, not globally optimal (optimal trees are NP-hard).
- Leaves can report class **proportions** for probabilities, but these are often **uncalibrated**.
- Watch greedy myopia and bias toward many-valued features; these impurity measures drive all tree ensembles.

Decision Trees for Regression

Level: Core

Trees predict numbers as easily as labels. A **regression tree** partitions the feature space the same way, but each leaf predicts a *constant value* — the average target in that region — and splits are chosen to reduce *variance* rather than class impurity. The result is a piecewise-constant surface: a staircase that approximates any function, with characteristic strengths and one notable weakness.

3.1 Splitting to reduce variance

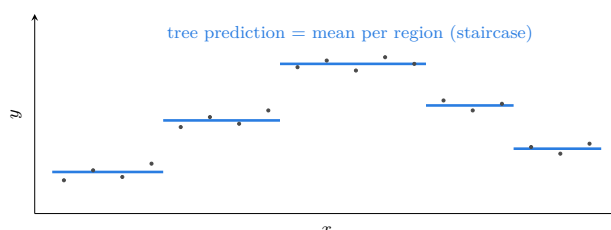
Definition 3.1. For a regression tree, a node's impurity is the **variance** (mean squared error) of its targets around their mean $\bar{y}$:

$$I(\text{node}) = \frac{1}{n} \sum_{i \in \text{node}} (y_i - \bar{y})^2.$$

A split is chosen to maximize the **variance reduction** — parent MSE minus the weighted average child MSE — exactly the information-gain formula with variance as the impurity.

Key idea

The optimal constant prediction in a region, under squared error, is the **mean** of the targets there. So a regression tree fits a step function: the feature space is carved into boxes, and each box predicts its average y . Minimizing within-node variance is the same as making each box's points as flat (similar in y) as possible — the regression analog of making each box pure in class.



3.2 Reading a regression tree

Intuition

Prediction is identical to classification: route the input down the questions to a leaf, then output that leaf's stored average. Deeper trees have more, smaller boxes and finer steps; a tree of depth d can have up to 2^d leaves. Because each region is a flat constant, a regression tree captures *nonlinearity and interactions* through the *arrangement* of cuts, not through any smooth formula.

3.3 The extrapolation weakness

Common pitfall

A regression tree **cannot extrapolate**. Its prediction is always an average of *training* targets, so for inputs beyond the training range it just repeats the nearest leaf's constant — a flat line, never a rising trend. Feed a tree trained on house sizes 50–200 m² a 400 m² mansion and it predicts the same price as a 200 m² house. Linear models extrapolate (sometimes wrongly, but they try); trees refuse. For trending data (time series, growth), detrend first or use a model that can extrapolate.

Intuition

The flip side is robustness: because predictions are bounded by observed targets and based on order (not magnitude), regression trees are **insensitive to outliers in features** and to monotonic feature transforms — taking logs of a predictor does not change a single split. This is part of why trees need so little preprocessing.

How this powers ML / AI

Regression trees are the base learners of **gradient boosting** — and there, crucially, the trees do not fit the target y directly but the *negative gradient of a loss* (for squared error, the residuals). Each boosting round adds a small regression tree that nudges predictions toward the truth (Chapter 7). So the humble “predict the mean in each box” rule, repeated thousands of times on residuals, becomes XGBoost — the model behind a large share of winning tabular solutions. The extrapolation weakness also explains a known failure mode of boosted trees on trending targets.

Try it in NumPy

```
1 import numpy as np
2 from sklearn.tree import DecisionTreeRegressor
3 rng = np.random.default_rng(0)
4 X = np.sort(rng.uniform(0, 10, 120)).reshape(-1, 1)
5 y = np.sin(X).ravel() + rng.normal(0, 0.1, 120)
6 for d in [2, 5]:
7     t = DecisionTreeRegressor(max_depth=d).fit(X, y)
8     print(f"depth {d}: leaves={t.get_n_leaves()}, train R^2={t.score(
9         X, y):.3f}")
10 print("extrapolation at x=20:", t.predict([[20]]).round(3), # flat,
11       = nearest leaf mean)
```

```
10 " vs at x=10:", t.predict([[10]]).round(3))
```

Chapter summary

- A **regression tree** predicts the **mean target** in each leaf region; splits maximize **variance reduction** (MSE).
- The model is **piecewise-constant** — a staircase that captures nonlinearity and interactions via the arrangement of cuts.
- Trees **cannot extrapolate**: predictions are bounded by training targets and flatten outside the training range.
- They are robust to outliers and invariant to monotonic feature transforms — hence minimal preprocessing.
- Regression trees on *residuals/gradients* are the base learners of gradient boosting (XGBoost, LightGBM, CatBoost).

PART II

Building Trees Well

Growing, Splitting, and Pruning

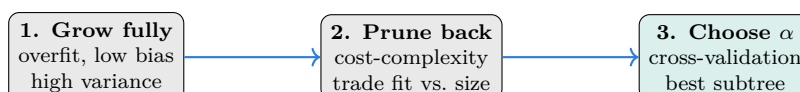
Level: Core

A tree left to grow freely will keep splitting until every leaf is pure — memorizing the training set, noise and all. The art of a single good tree is **controlling its size**: how to grow it, when to stop, and how to cut it back. This chapter covers the greedy growth algorithm, the stopping (pre-pruning) hyperparameters, and the principled alternative — cost-complexity (post-)pruning.

4.1 The CART growth algorithm

Key idea

CART grows a tree by **greedy recursive binary splitting**: starting at the root with all data, find the single feature–threshold split that most reduces impurity, partition into two children, and recurse on each. It is greedy (best split now, no lookahead) and produces strictly *binary* trees. Left unchecked it grows until leaves are pure or contain one point — a perfect, useless overfit.



4.2 Pre-pruning: stop early

Intuition

Pre-pruning (early stopping) halts growth using hyperparameters — it is fast and the everyday way to control trees:

- **max_depth** — the maximum number of question levels (the single biggest knob).
- **min_samples_split** — don't split a node with too few samples.
- **min_samples_leaf** — require each leaf to keep at least this many samples (smooths predictions).
- **max_leaf_nodes** / **min_impurity_decrease** — cap total leaves / require a minimum gain to split.

The danger of pre-pruning is *horizon effect*: a weak split that unlocks a great split below it gets blocked, because the algorithm cannot see past the next step.

4.3 Cost-complexity (post-)pruning

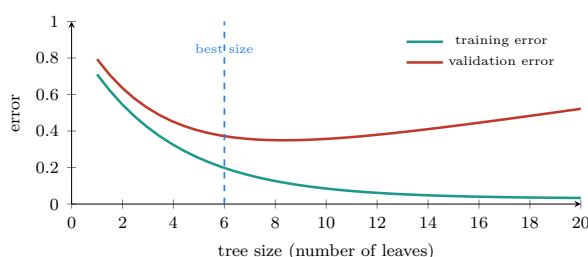
Definition 4.1. Cost-complexity pruning grows a large tree T_0 , then prunes it back to minimize a penalized error

$$R_\alpha(T) = R(T) + \alpha |T|,$$

where $R(T)$ is the training error (misclassification or MSE), $|T|$ is the number of leaves, and $\alpha \geq 0$ is the **complexity parameter**. Larger α favors smaller trees.

Key idea

As α increases from 0, cost-complexity pruning produces a *nested sequence* of ever-smaller subtrees — and remarkably, the best subtree of each size lies on this sequence (weakest-link pruning). You then pick α (hence tree size) by **cross-validation**, choosing the simplest tree within one standard error of the best. This separates *growing* (greedy, fast) from *model selection* (principled, validated) — generally beating pre-pruning at the cost of more computation.



4.4 Categorical splits and missing values

Intuition

For categorical features, CART considers splits that send a *subset* of categories left and the rest right; with q categories there are $2^q - 1$ possible binary groupings, though efficient algorithms avoid brute force for ordered targets. Missing values are handled either by *surrogate splits* (a backup feature that mimics the primary split) in classic CART, or by learning a default direction (XGBoost/LightGBM) — more in Chapter 10.

Common pitfall

Don't tune by training error — it always prefers the biggest tree. Always select depth/ α on a validation set or by cross-validation. And remember: *pruning matters far less once you ensemble*. In a random forest you deliberately grow deep, unpruned trees (variance is killed by averaging, Chapter 6); in gradient boosting you grow shallow trees (depth 3–8) and control complexity with the learning rate and number of trees. Heavy single-tree pruning is mostly for when you need *one interpretable tree*.

How this powers ML / AI

The growth/stop/prune trio is the original **regularization** story of machine learning, pre-dating weight decay and dropout: $R_\alpha(T) = R(T) + \alpha|T|$ is loss-plus-complexity-penalty,

exactly the template of modern regularized objectives. XGBoost makes this explicit — its objective adds $\gamma T + \frac{1}{2}\lambda \sum w_j^2$, penalizing the number of leaves T and leaf weights, a direct descendant of cost-complexity pruning baked into the training loss (Chapter 8).

Try it in NumPy

```

1 from sklearn.tree import DecisionTreeClassifier
2 from sklearn.datasets import load_breast_cancer
3 from sklearn.model_selection import cross_val_score
4 X, y = load_breast_cancer(return_X_y=True)
5 path = DecisionTreeClassifier(random_state=0).
    cost_complexity_pruning_path(X, y)
6 for a in path.ccp_alphas[::-6]:           # sweep the complexity
    parameter
7     clf = DecisionTreeClassifier(ccp_alpha=a, random_state=0)
8     print(f"alpha={a:.4f}  cv acc={cross_val_score(clf, X, y, cv=5).
    mean():.3f}")

```

Chapter summary

- CART grows by **greedy recursive binary splitting**; unchecked, it overfits to pure leaves.
- **Pre-pruning** stops early via `max_depth`, `min_samples_leaf`, etc. — fast, but suffers the horizon effect.
- **Cost-complexity pruning** minimizes $R(T) + \alpha|T|$, giving a nested subtree sequence; pick α by cross-validation.
- Never select tree size by training error; use validation/CV (and the one-standard-error rule for simplicity).
- In ensembles, pruning matters less: forests grow deep, boosting grows shallow — the penalty $R + \alpha|T|$ foreshadows regularized boosting objectives.

5

Strengths, Weaknesses, and the Bias–Variance of Trees

Level: Core

Before we multiply trees into forests and boosters, it pays to understand a single tree’s character honestly: what it does brilliantly, where it fails, and — most importantly — the bias–variance behavior that motivates every ensemble method to come. The headline: a deep tree has *low bias but high variance*, and that single fact explains both bagging and boosting.

5.1 What trees do well

Key idea

Decision trees earn their place because they are **low-maintenance and flexible**:

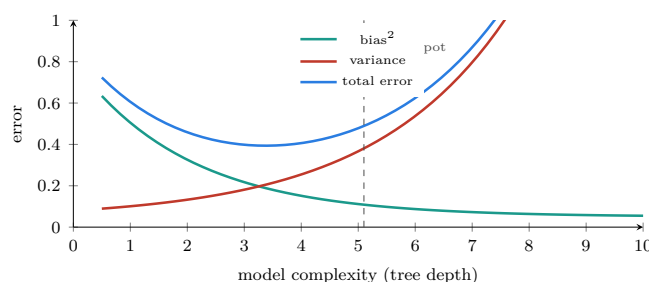
- **No feature scaling.** Splits depend only on order, so standardization and normalization are irrelevant.
- **Mixed data, minimal prep.** Numeric and categorical features, different scales, outliers — handled with little fuss.
- **Nonlinearity and interactions for free.** Captured by the arrangement of splits, no manual feature engineering.
- **Interpretable (when small).** A shallow tree is a readable flowchart; feature importance is built in.
- **Fast** to train and predict.

5.2 The core weakness: high variance

Definition 5.1. A model has high **variance** if small changes in the training data produce large changes in the fitted model. Deep decision trees are the textbook example: because each split is chosen greedily and all descendant splits depend on it, perturbing a few points can change the root split and *reshuffle the entire tree*.

Intuition

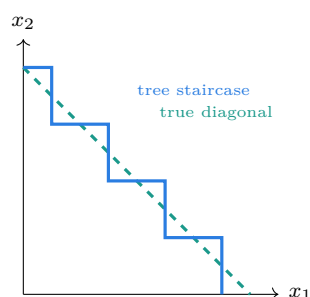
Picture the bias–variance tradeoff as tree depth grows. A *stump* (depth 1) is stable but crude — high bias, low variance, it underfits. A *fully grown* tree fits the training data perfectly — low bias, but wildly high variance; the staircase chases noise. Somewhere between lies the best single tree, but its test error is still limited by that variance. The key realization: *trees are unstable, low-bias learners*. That is exactly the raw material averaging loves.



5.3 Other limitations

Common pitfall

Beyond variance, single trees have real weaknesses: (1) **no extrapolation** (Chapter 3); (2) **axis-aligned only** — a diagonal boundary must be approximated by a clumsy staircase of horizontal/vertical cuts; (3) **biased splits** toward high-cardinality features; (4) **instability** makes “the” tree hard to trust for inference; and (5) a tendency to create **unbalanced trees** on skewed data. Most of these are cured or muted by ensembling — the price is interpretability.



5.4 From one tree to many

Key idea

Here is the pivot of the whole book. Averaging B independent, identically distributed predictions with variance σ^2 yields variance σ^2/B ; if they are only *correlated* with correlation ρ , the averaged variance is

$$\rho \sigma^2 + \frac{1 - \rho}{B} \sigma^2.$$

So to slash variance you want *many* trees (B large) that are *decorrelated* (ρ small). High-variance, low-bias trees are the ideal ingredient: averaging them barely touches bias but crushes variance. This single formula motivates **bagging** (grow many trees on resampled data) and **random forests** (also decorrelate them by random feature subsets) — Chapter 6. *Boosting* attacks the problem from the other side, reducing bias by adding trees sequentially — Chapter 7.

How this powers ML / AI

The bias–variance lens is the through-line from trees to all of ML. Random forests exploit “low bias, high variance, decorrelate and average”; gradient boosting exploits “start high-bias, reduce bias additively, regularize against variance.” Deep learning lives in the same

coordinates — big networks are low-bias/high-variance and are tamed by data augmentation, dropout, and ensembling, the neural echoes of bagging. Understanding why one tree is unstable is understanding why every powerful model needs a variance-control strategy.

Try it in NumPy

```

1 import numpy as np
2 from sklearn.tree import DecisionTreeRegressor
3 # Variance demo: refit a deep tree on bootstrap samples; predictions
  # at x=5 vary a lot
4 rng = np.random.default_rng(0)
5 X = np.sort(rng.uniform(0,10,200)).reshape(-1,1); y = np.sin(X).ravel
  ()+rng.normal(0,.2,200)
6 preds = []
7 for _ in range(50):
8     idx = rng.integers(0, 200, 200)           # bootstrap
  resample
9     t = DecisionTreeRegressor(max_depth=None).fit(X[idx], y[idx])
10    preds.append(t.predict([[5.0]])[0])
11 print("deep tree std at x=5:", round(np.std(preds), 3)) # high ->
  averaging will help

```

Chapter summary

- Trees shine at **minimal preprocessing**, mixed data, automatic nonlinearity/interactions, and (when small) interpretability.
- Their core flaw is **high variance**: deep trees are low-bias but unstable, reshuffling under small data changes.
- Other limits: no extrapolation, axis-aligned (staircase) boundaries, biased splits, instability.
- Averaging B trees with correlation ρ gives variance $\rho\sigma^2 + \frac{1-\rho}{B}\sigma^2$ — the case for many, decorrelated trees.
- This bias–variance picture motivates **bagging/forests** (cut variance) and **boosting** (cut bias) — the rest of the book.

PART III

Ensembles

Bagging and Random Forests

Level: Advanced

We now turn high-variance single trees into one of the most robust models in machine learning. The recipe has two ingredients: **bagging** (train many trees on bootstrap resamples and average them) and the **random forest** twist (also randomize the features at each split). Together they slash variance while keeping bias low, producing an excellent, low-tuning default for tabular data.

6.1 Bagging: bootstrap aggregating

Definition 6.1. Bagging (Breiman, 1996) trains B models on B **bootstrap samples** — datasets of size n drawn *with replacement* from the training set — and aggregates their predictions: *average* for regression, *majority vote* (or averaged probabilities) for classification.

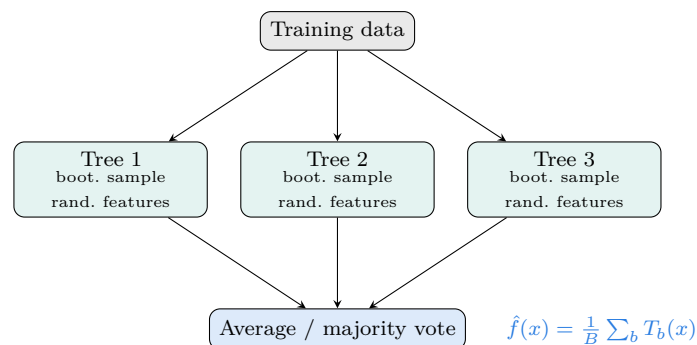
Intuition

Recall the averaging law: B trees each with variance σ^2 and pairwise correlation ρ have averaged variance $\rho\sigma^2 + \frac{1-\rho}{B}\sigma^2$. Bagging drives the second term toward zero by using many trees. But the trees are trained on overlapping bootstrap samples, so they remain *correlated* (ρ stays high), and the first term $\rho\sigma^2$ sets a floor on how much bagging alone can help. Removing that floor is exactly the random-forest idea.

6.2 Random forests: decorrelate the trees

Key idea

A **random forest** (Breiman, 2001) is bagging plus **feature subsampling**: at *each split*, only a random subset of m features (out of p) is considered. This stops every tree from latching onto the same one or two dominant features, **decorrelating** the trees ($\rho \downarrow$) so that averaging helps far more. Typical defaults: $m = \sqrt{p}$ for classification, $m = p/3$ for regression. Trees are grown deep and *unpruned* — variance is handled by the ensemble, not by each tree.



6.3 Out-of-bag (OOB) error: free validation

Key idea

Each bootstrap sample omits about $1/e \approx 37\%$ of the data — these are the **out-of-bag** samples for that tree. Predict each training point using only the trees that did *not* see it, and you get an almost-free, nearly unbiased estimate of test error — **no separate validation split needed**. OOB error is a hallmark convenience of random forests, great for quick model assessment and even feature-importance estimates.

Example 6.2. Why $\approx 37\%$? The probability a specific point is *not* chosen in one draw is $1 - \frac{1}{n}$; over n draws it is $(1 - \frac{1}{n})^n \rightarrow e^{-1} \approx 0.368$. So each tree trains on $\approx 63\%$ of the unique points and can be validated on the remaining $\approx 37\%$.

6.4 Why forests are such a good default

Intuition

Random forests are forgiving: they rarely overfit as you add trees (more trees only stabilize the average — you cannot overfit by growing the forest, only by other means), need little tuning, handle thousands of features, give OOB error and importances for free, and parallelize trivially (trees are independent). The main costs: the model is a *black box* (hundreds of deep trees), larger memory/inference, and they usually trail well-tuned gradient boosting by a small margin in accuracy.

Common pitfall

Misconceptions to avoid. (1) “*More trees can overfit*” — no; more trees reduce variance and plateau. Overfitting comes from too-deep trees on too-little data or leakage, not from B . (2) *Default impurity importances are trustworthy* — they are biased toward high-cardinality and correlated features; prefer permutation importance (Chapter 9). (3) *m doesn’t matter* — with many correlated/irrelevant features, tuning `max_features` noticeably affects decorrelation and accuracy. (4) *Forests extrapolate* — they still cannot, being averages of non-extrapolating trees.

How this powers ML / AI

Random forests are the canonical **variance-reduction-by-averaging** method, and that idea recurs everywhere: *deep ensembles* (averaging several independently trained neural

networks) are essentially bagging for deep learning, and remain a top method for accuracy and uncertainty estimation. *Bayesian model averaging*, *Monte-Carlo dropout*, and *snapshot ensembles* are all variations on “average decorrelated predictors.” Random forests also remain a go-to for quick, robust baselines, feature screening (importances), and embedded uses like *isolation forests* for anomaly detection (Chapter 11).

Try it in NumPy

```

1 from sklearn.ensemble import RandomForestClassifier
2 from sklearn.datasets import load_breast_cancer
3 X, y = load_breast_cancer(return_X_y=True)
4 rf = RandomForestClassifier(n_estimators=500, max_features="sqrt",
5                             oob_score=True, n_jobs=-1, random_state
6                             =0).fit(X, y)
7 print("OOB accuracy:", round(rf.oob_score_, 3))      # free validation
8                                                     estimate
9 import numpy as np
10 top = np.argsort(rf.feature_importances_)[:-1][:-3]
11 print("top feature indices:", top)

```

Chapter summary

- **Bagging** averages trees trained on **bootstrap resamples**, cutting variance — but correlated trees leave a floor $\rho\sigma^2$.
- **Random forests** add per-split **feature subsampling** ($m \approx \sqrt{p}$ or $p/3$) to **decorrelate** trees and average better.
- Trees are grown deep and unpruned; variance control comes from the ensemble.
- **OOB error** ($\approx 37\%$ held out per tree) gives near-free, unbiased test-error and importance estimates.
- Forests are robust, low-tuning, parallel defaults; the deep-learning analog is the **deep ensemble**.

Boosting: AdaBoost and Gradient Boosting

Level: Advanced

Bagging builds trees *in parallel* to cut variance. **Boosting** builds them *in sequence* to cut bias: each new tree focuses on what the ensemble got wrong so far. This chapter develops boosting from AdaBoost (reweight the mistakes) to the more general and powerful gradient boosting (fit the residuals — gradient descent in function space), the engine behind XGBoost, LightGBM, and CatBoost.

7.1 The boosting idea

Key idea

Boosting combines many *weak learners* (typically shallow trees — “stumps” or depth-3 trees) into a strong one by adding them sequentially. Each new learner is trained to correct the errors of the current ensemble, and predictions accumulate:

$$F_M(x) = \sum_{m=1}^M \nu h_m(x),$$

where h_m is the m -th tree and ν is a small **learning rate** (shrinkage). Unlike bagging’s independent trees, boosting’s trees are deeply dependent — order matters.

7.2 AdaBoost: reweight the hard cases

Intuition

AdaBoost (Freund & Schapire, 1997) keeps a weight on each training example. Start uniform; fit a weak classifier; *increase* the weights of misclassified points and *decrease* the weights of correct ones; fit the next classifier on the reweighted data; repeat. Final prediction is a weighted vote, where more accurate learners get larger say ($\alpha_m = \frac{1}{2} \ln \frac{1-\text{err}_m}{\text{err}_m}$). Each round literally pays more attention to the examples the ensemble keeps getting wrong.



7.3 Gradient boosting: fit the residuals

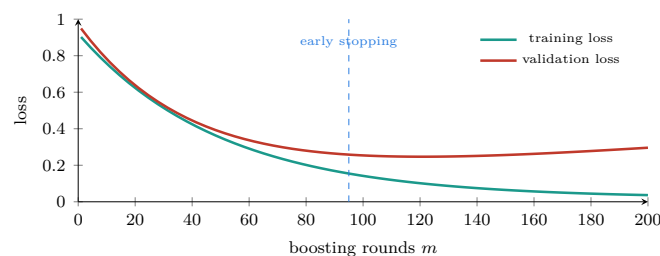
Key idea

Gradient boosting (Friedman, 2001) generalizes boosting to any differentiable loss $L(y, F)$ by viewing it as **gradient descent in function space**. At each step, compute the negative gradient of the loss with respect to the current predictions — the *pseudo-residuals* $r_i = -[\partial L(y_i, F(x_i))/\partial F(x_i)]$ — then fit a regression tree to those residuals and add it (scaled by ν):

$$F_m(x) = F_{m-1}(x) + \nu h_m(x), \quad h_m \approx \text{tree fit to } r_i.$$

For squared-error loss the pseudo-residuals are just the ordinary residuals $y_i - F_{m-1}(x_i)$, so gradient boosting *literally fits the errors* and walks the predictions downhill.

Example 7.1. Squared-error regression. Start with the mean $F_0 = \bar{y}$. Round 1: residuals $r_i = y_i - \bar{y}$; fit a small tree to them; update $F_1 = F_0 + \nu h_1$. Round 2: new residuals $y_i - F_1(x_i)$ are smaller; fit another tree; update. After M rounds the residuals (and training loss) shrink steadily. The learning rate ν controls step size — small ν with many trees generalizes best.



7.4 Regularization: learning rate, subsampling, depth

Intuition

Boosting can overfit if over-trained, so it is regularized on several axes: a small **learning rate** ν (0.01–0.1) takes cautious steps (more trees needed but better generalization — “shrinkage”); **early stopping** on a validation set picks the number of trees M ; **stochastic gradient boosting** fits each tree on a random subsample of rows (and columns), adding bagging-style variance reduction; and **shallow trees** (depth 3–8) keep each learner weak. The big three — ν , M , depth — must be tuned *together*: halve ν and you roughly double the M you need.

Common pitfall

Boosting is more sensitive than random forests. (1) **Adding trees can overfit** — unlike forests, M is a real overfitting knob; use early stopping. (2) **Noisy labels / outliers hurt**: AdaBoost in particular keeps up-weighting mislabeled points and can chase them; robust losses (Huber) or gradient boosting help. (3) **Tune ν and M jointly**, not separately. (4) **Sequential = slower to train** and not as trivially parallel as a forest (though modern libraries parallelize within each tree).

How this powers ML / AI

“Gradient descent in function space” is the conceptual bridge from boosting to deep learning: both minimize a differentiable loss by following negative gradients — boosting adds a tree per step, a neural net updates weights per step. The pseudo-residual trick (fit the next learner to the gradient of any loss) lets boosted trees optimize log-loss, Poisson, ranking, and quantile objectives, just as autodiff lets networks optimize arbitrary losses. Gradient-boosted trees are also routinely *stacked with* neural networks (their outputs as features, or vice versa) in winning tabular pipelines (Chapter 13).

Try it in NumPy

```

1 from sklearn.ensemble import GradientBoostingClassifier
2 from sklearn.datasets import make_classification
3 from sklearn.model_selection import train_test_split
4 X, y = make_classification(n_samples=3000, n_informative=8,
5                           random_state=0)
6 Xtr, Xte, ytr, yte = train_test_split(X, y, random_state=0)
7 for lr in [0.5, 0.1, 0.02]:
8     gb = GradientBoostingClassifier(learning_rate=lr, n_estimators
9                                     =300,
10                                     max_depth=3).fit(Xtr, ytr)
11     print(f"lr={lr}: test acc={gb.score(Xte, yte):.3f}") # small lr
12                                     usually generalizes best

```

Chapter summary

- **Boosting** adds weak learners **sequentially**, each correcting the ensemble’s errors — it reduces **bias**.
- **AdaBoost** reweights misclassified points each round and takes a weighted vote.
- **Gradient boosting** fits each tree to the **negative-gradient (pseudo-residuals)** of any differentiable loss — gradient descent in function space.
- Regularize with small **learning rate**, **early stopping**, **row/column subsampling**, and **shallow trees**; tune ν and M together.
- Boosting can overfit (unlike forests) and is sensitive to noise — but tuned well it is the most accurate tabular method.

Modern Gradient Boosting: XGBoost, LightGBM, CatBoost

Level: Advanced

Friedman’s gradient boosting is the idea; **XGBoost**, **LightGBM**, and **CatBoost** are the engineering that made it dominate. Each adds algorithmic and systems innovations — regularized objectives, second-order optimization, histogram binning, clever tree growth, and native categorical handling — that deliver state-of-the-art accuracy on tabular data at remarkable speed. This chapter explains what each brings and how to choose between them.

8.1 XGBoost: regularized, second-order boosting

Key idea

XGBoost (Chen & Guestrin, 2016) sharpens gradient boosting in two ways. First, a **regularized objective** adds an explicit complexity penalty per tree:

$$\mathcal{L} = \sum_i L(y_i, \hat{y}_i) + \sum_m \left(\gamma T_m + \frac{1}{2} \lambda \|w_m\|^2 \right),$$

penalizing the number of leaves T_m and the leaf weights w_m . Second, it uses a **second-order Taylor expansion** of the loss (gradients g_i and Hessians h_i), giving a closed-form optimal leaf weight $w_j^* = -\frac{\sum_{i \in j} g_i}{\sum_{i \in j} h_i + \lambda}$ and a principled split-gain score. The result: faster convergence, built-in regularization, and support for any twice-differentiable loss.

Intuition

The second-order view is why XGBoost feels “smarter” than vanilla GBM: instead of just walking downhill on the gradient, it uses curvature (the Hessian) to size each step — a Newton-style update inside every leaf. The γ (minimum split gain) and λ (L2 on weights) terms are literally the cost-complexity pruning idea from Chapter 4, now folded into the training objective. XGBoost also adds engineering wins: sparsity-aware split finding (a learned default direction for missing values), cache-aware access, and out-of-core training.

8.2 LightGBM: histograms, leaf-wise, GOSS, EFB

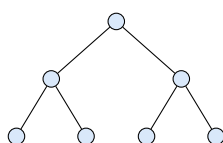
Key idea

LightGBM (Microsoft, 2017) is built for **speed on large data**:

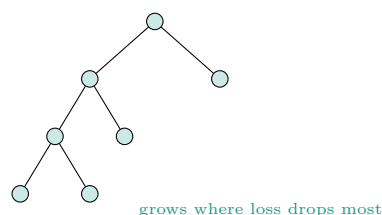
- **Histogram binning** — bucket continuous features into (e.g.) 255 bins so split finding scans bins, not sorted values: huge memory and time savings.
- **Leaf-wise growth** — grow the leaf with the largest loss reduction (best-first), rather than level-wise; deeper, more accurate trees for the same number of leaves (control with `num_leaves`).
- **GOSS** (gradient-based one-side sampling) — keep large-gradient (hard) examples, sub-sample small-gradient ones, to focus computation.
- **EFB** (exclusive feature bundling) — bundle mutually-exclusive sparse features into one, shrinking the effective feature count.

LightGBM is typically the fastest on millions of rows.

Level-wise (XGBoost default)



Leaf-wise (LightGBM)



8.3 CatBoost: ordered boosting and native categorical

Key idea

CatBoost (Yandex, 2018) targets **categorical features and small-data bias**. It encodes categories with **ordered target statistics** (each row's encoding uses only *prior* rows, never its own label) and uses **ordered boosting** to avoid *target leakage / prediction shift* — the subtle bias where a model peeks at the target it is trying to predict. It also builds **symmetric (oblivious) trees** (the same split across an entire level), which act as a regularizer and make inference extremely fast. CatBoost often wins with minimal tuning, especially on categorical-heavy data.

8.4 Choosing a library

	XGBoost	LightGBM	CatBoost
Tree growth	level-wise (depth)	leaf-wise (best-first)	symmetric trees
Speed	fast	fastest on big data	fast, fast inference
Categoricals	manual encoding	native (integer)	native, ordered
Best at	robust all-rounder	large datasets	categorical-heavy, low tuning
Watch out	tuning effort	overfits num_leaves large	if slower training

Common pitfall

Practical traps: (1) **LightGBM's num_leaves** is powerful but dangerous — leaf-wise growth overfits if it is large relative to data; keep it $< 2^{\text{max_depth}}$ and use **min_child_samples**. (2) **Don't one-hot high-cardinality categorical**s for LightGBM/CatBoost — use their native handling. (3) **Always use early stopping** with a validation set rather than guessing **n_estimators**. (4) **Tune learning rate with number of trees together**; lower LR + more trees + early stopping is the reliable recipe. (5) **Set seeds and log versions** — results differ across libraries and versions.

How this powers ML / AI

These three libraries *are* the state of the art for tabular machine learning in industry and competitions — the practical answer to “what model should I try first on a table?” Their innovations also cross-pollinate with deep learning: histogram binning resembles feature discretization in tabular nets; the regularized, second-order objective mirrors modern optimizers; and GBDTs are constantly benchmarked against (and combined with) tabular transformers like FT-Transformer and foundation models like TabPFN (Chapters 12, 13). Mastering them is, for tabular AI, as fundamental as knowing SGD is for deep learning.

Try it in NumPy

```

1 # pip install xgboost lightgbm catboost
2 import xgboost as xgb
3 from sklearn.datasets import make_classification
4 from sklearn.model_selection import train_test_split
5 X, y = make_classification(n_samples=5000, n_informative=10,
6                           random_state=0)
7 Xtr, Xva, ytr, yva = train_test_split(X, y, random_state=0)
8 clf = xgb.XGBClassifier(n_estimators=2000, learning_rate=0.03,
9                         max_depth=5,
10                          subsample=0.8, colsample_bytree=0.8,
11                          early_stopping_rounds=50, eval_metric="
12                          logloss")
13 clf.fit(Xtr, ytr, eval_set=[(Xva, yva)], verbose=False)
14 print("best #trees:", clf.best_iteration, " val acc:", round(clf.
15                       score(Xva, yva), 3))

```

Chapter summary

- **XGBoost**: regularized objective ($\gamma T + \frac{1}{2}\lambda\|w\|^2$) and **second-order** (gradient + Hessian) optimization with sparsity-aware splits.
- **LightGBM**: **histogram** binning, **leaf-wise** growth, **GOSS** and **EFB** — fastest on large data (mind `num_leaves`).
- **CatBoost**: **ordered boosting** + ordered target encoding to kill prediction shift; native categoricals; symmetric trees.
- Pick by data: XGBoost (robust default), LightGBM (large data), CatBoost (many categoricals, low tuning).
- Always use **early stopping**; tune learning rate and tree count together — these libraries are the tabular state of the art.

PART IV

Interpretation and Practice

Feature Importance and Interpretability

Level: Advanced

A single tree is a readable flowchart, but a forest of hundreds or a thousand boosted trees is a black box. The good news: tree ensembles support some of the best interpretability tools in machine learning. This chapter covers built-in (impurity) importance and its biases, the more trustworthy permutation importance, the gold-standard SHAP values, and partial-dependence plots — enough to explain a model to a stakeholder and to catch leakage.

9.1 Built-in (impurity / gain) importance

Definition 9.1. Impurity importance (Mean Decrease in Impurity, MDI) scores a feature by the total impurity reduction (Gini/variance, or split gain in boosting) it contributes, summed over all splits that use it and averaged over all trees.

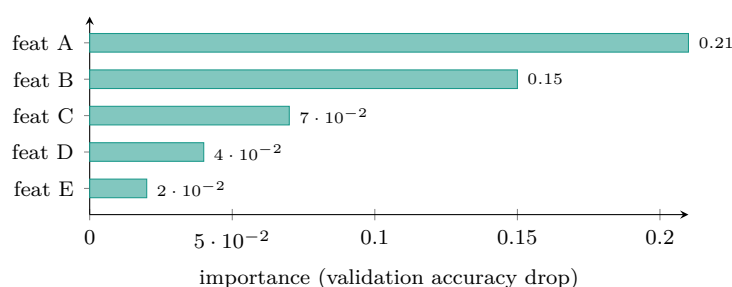
Common pitfall

Impurity importance is fast and free but **biased**: it inflates the importance of **high-cardinality** features (many split points $\Rightarrow$ more chances to reduce impurity, even by chance) and **continuous** features over binary ones, and it is computed on *training* data so it rewards overfitting. With **correlated** features, importance is split arbitrarily among them. Treat MDI as a rough first look, never as ground truth.

9.2 Permutation importance

Key idea

Permutation importance measures how much a model's *validation* score drops when a feature's values are randomly shuffled, breaking its relationship with the target. It is **model-agnostic**, computed on held-out data, and directly answers “how much does the model rely on this feature for accuracy?” — avoiding MDI's cardinality bias. Its caveat: with correlated features, shuffling one is compensated by its correlate, so both can look unimportant; group correlated features or use conditional variants.



9.3 SHAP: consistent, local attributions

Key idea

SHAP (SHapley Additive exPlanations) assigns each feature a contribution to *each individual prediction*, grounded in cooperative game theory: a feature's SHAP value is its average marginal contribution over all orderings of features. SHAP is **locally accurate** (per-prediction contributions sum to the prediction minus the baseline) and **consistent** (if a model relies on a feature more, its SHAP value cannot decrease). **TreeSHAP** computes exact Shapley values for tree ensembles in polynomial time — making SHAP the modern standard for explaining tabular models both locally (one prediction) and globally (aggregated).

Intuition

Read SHAP two ways. A **waterfall/force plot** explains one prediction: start at the baseline (average output), then each feature pushes the prediction up or down by its SHAP value until you reach the final score — “this loan was denied mainly because of high debt (+0.3) despite good income (−0.1).” A **summary (beeswarm) plot** aggregates across all samples: features ranked by importance, with color showing whether high or low feature values push predictions up or down — revealing direction and interactions, not just magnitude.

9.4 Partial dependence and ICE

Intuition

A **partial dependence plot (PDP)** shows the average predicted output as one feature varies, marginalizing the others — the *shape* of the learned relationship (monotone? threshold? U-shaped?). **Individual Conditional Expectation (ICE)** curves show one line per sample, exposing heterogeneity and interactions that a single averaged PDP hides. PDPs assume feature independence, so they mislead when features are strongly correlated (they evaluate unrealistic combinations).

Common pitfall

Interpretation pitfalls: (1) **importance** $\neq$ **causation** — a feature can be important because it proxies a confounder or leaks the target; (2) a single feature with *implausibly dominant* importance/SHAP is a classic **leakage** red flag (e.g., an ID, a post-outcome timestamp); (3) PDP/permutation both **distort under correlation**; (4) explanations describe *the model*, not the world — a wrong model yields confident, wrong explanations.

Use these tools to debug and communicate, not to prove mechanisms.

How this powers ML / AI

SHAP began with trees but became a unifying language for explainable AI across *all* models, including deep networks (DeepSHAP, GradientSHAP) — the same Shapley-value foundation. Tree-based importances and SHAP are everywhere in production ML: detecting data leakage, satisfying regulatory “right to explanation” (credit, insurance), monitoring feature drift, and selecting features. Because TreeSHAP is exact and fast, gradient-boosted trees are often the most *auditable* high-accuracy models available — a major reason regulated industries prefer them over neural nets (Chapter 13).

Try it in NumPy

```

1 from sklearn.ensemble import RandomForestClassifier
2 from sklearn.inspection import permutation_importance
3 from sklearn.datasets import load_breast_cancer
4 from sklearn.model_selection import train_test_split
5 X, y = load_breast_cancer(return_X_y=True)
6 Xtr, Xte, ytr, yte = train_test_split(X, y, random_state=0)
7 rf = RandomForestClassifier(n_estimators=300, random_state=0).fit(Xtr
8     , ytr)
9 r = permutation_importance(rf, Xte, yte, n_repeats=20, random_state
10     =0)
11 import numpy as np
12 for i in np.argsort(r.importances_mean)[::-1][:5]:
13     print(f"feat {i}: {r.importances_mean[i]:.3f} +/- {r.
14         importances_std[i]:.3f}")
15 # import shap; shap.TreeExplainer(rf).shap_values(Xte) # exact per-
16     prediction attributions

```

Chapter summary

- **Impurity (MDI) importance** is free but biased toward high-cardinality/continuous features and computed on training data.
- **Permutation importance** (validation score drop under shuffling) is model-agnostic and more trustworthy — mind correlated features.
- **SHAP** gives consistent, locally accurate per-prediction attributions; **TreeSHAP** is exact and fast for tree ensembles.
- **PDP/ICE** reveal the shape of learned relationships but distort under feature correlation.
- Importance $\neq$ causation; a dominant feature often signals **leakage**. Trees + SHAP make ensembles highly auditable.

Practical Tree Modeling

Level: Expert

Knowing the algorithms is half the battle; the other half is the craft of applying them to messy real data. This chapter is a practitioner’s checklist: handling categoricals and missing values, dealing with class imbalance, tuning hyperparameters efficiently, avoiding leakage in pipelines, and reading the right metrics. Done well, a gradient-boosting model goes from “works on my laptop” to “trustworthy in production.”

10.1 Categorical features

Intuition

Options, best to worst by context: **native handling** (LightGBM integer-encoded, CatBoost ordered target stats) is usually best and avoids exploding dimensionality; **target/mean encoding** is powerful but *leaks* if not done inside cross-validation folds; **one-hot** is fine for low cardinality but creates sparse, deep-split-unfriendly features at high cardinality; **ordinal** encoding imposes a false order on nominal categories (acceptable for trees more than for linear models, but still risky). Never one-hot a 10,000-level ID — use native handling or hashing.

10.2 Missing values

Key idea

Trees handle missing data more gracefully than most models. Classic CART uses **surrogate splits** (a backup feature mimicking the primary split). XGBoost and LightGBM learn a **default direction**: at each split, missing values go whichever way reduces loss most — so missingness is used as signal, with no imputation required. This is a real advantage over neural nets and linear models, which need explicit imputation. Still, add a “*was-missing*” indicator when missingness itself is informative.

10.3 Class imbalance

Intuition

For rare-positive problems (fraud, disease), accuracy is useless (predict “all negative” scores 99%). Remedies: **class weights** / **scale_pos_weight** to up-weight the minority in the loss; **resampling** (SMOTE, undersampling) *inside* CV folds; and — most importantly — **evaluate with the right metric**: precision/recall, F1, PR-AUC, or business cost, not

raw accuracy. Then **tune the decision threshold** for your cost tradeoff rather than defaulting to 0.5.

10.4 Hyperparameter tuning

Key idea

For gradient boosting, tune in rough priority order:

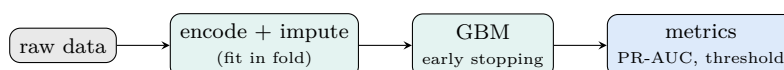
1. **Learning rate ν + number of trees** (jointly, with early stopping) — the master knobs.
2. **Tree complexity:** `max_depth` (3–8) or `num_leaves`; `min_child_samples/weight`.
3. **Subsampling:** `subsample` (rows), `colsample_bytree` (columns) — regularize and decorrelate.
4. **Regularization:** `reg_lambda` (L2), `reg_alpha` (L1), `gamma` (min split gain).

Use **random search** or **Bayesian optimization** (Optuna), not exhaustive grids. For random forests, defaults are often excellent — tune mainly `max_features` and `n_estimators` (more is just more stable).

10.5 Pipelines, validation, and leakage

Common pitfall

The deadliest production bug is **data leakage** — information from the target or the future sneaking into training. Guard rails: (1) do *all* fitting (encoders, scalers, imputers, feature selection) **inside cross-validation folds**, via a **Pipeline**, never on the full dataset first; (2) for time series, use **time-based splits**, never random shuffling; (3) beware target/mean encoding and any feature computed using the target; (4) a feature with implausibly high importance/SHAP is leakage until proven otherwise; (5) keep a truly untouched **test set** for the final estimate. Leakage produces gorgeous validation scores and humiliating production failures.



all steps fit *inside* CV folds $\Rightarrow$ no leakage

10.6 Metrics and calibration

Intuition

Match the metric to the goal: RMSE/MAE for regression (MAE if outliers matter); ROC-AUC for ranking ability, PR-AUC for imbalanced positives, log-loss/Brier for probability quality. If you need reliable *probabilities* (pricing, risk), **calibrate** the model (isotonic or Platt scaling) and check a reliability curve — boosted trees are often slightly miscalibrated out of the box.

How this powers ML / AI

This practical layer — pipelines, fold-safe preprocessing, Bayesian hyperparameter search, calibration, and threshold tuning — is identical in spirit to MLOps for deep learning, and tree models share the same tooling (scikit-learn pipelines, Optuna, MLflow). Crucially, the operational virtues here (no scaling, native missing/categorical handling, fast retraining, cheap inference, exact SHAP explanations) are exactly why gradient-boosted trees remain the **default production model** for tabular AI even at companies with deep-learning expertise.

Try it in NumPy

```

1 import optuna, xgboost as xgb
2 from sklearn.datasets import make_classification
3 from sklearn.model_selection import cross_val_score
4 X, y = make_classification(n_samples=4000, weights=[0.9, 0.1],
5                           random_state=0)
6 def objective(trial):
7     params = dict(
8         n_estimators=600, learning_rate=trial.suggest_float("lr",
9         0.01, 0.3, log=True),
10        max_depth=trial.suggest_int("depth", 3, 8),
11        subsample=trial.suggest_float("subsample", 0.6, 1.0),
12        scale_pos_weight=9) # handle 9:1
13    imbalance
14    m = xgb.XGBClassifier(**params, eval_metric="aucpr")
15    return cross_val_score(m, X, y, cv=4, scoring="average_precision")
16    ).mean()
17 study = optuna.create_study(direction="maximize")
18 study.optimize(objective, n_trials=20)
19 print("best PR-AUC:", round(study.best_value, 3), study.best_params)

```

Chapter summary

- Prefer **native categorical handling**; if target-encoding, do it *inside* CV folds to avoid leakage.
- Trees handle **missing values** via surrogate splits or learned default directions — often no imputation needed.
- For **imbalance**, use class weights/resampling-in-folds and judge with PR-AUC/F1 and a tuned threshold, not accuracy.
- Tune **learning rate** + **trees** first (early stopping), then depth, subsampling, regularization — via random/Bayesian search.
- Fit all preprocessing **inside CV** to prevent **leakage**; use time-based splits for temporal data; calibrate if you need probabilities.

PART V

Frontier and Applications

Specialized and Advanced Trees

Level: Expert

Beyond classification, regression, forests, and boosting lies a rich ecosystem of tree variants tuned for special jobs: detecting anomalies, ranking search results, modeling time-to-event, quantifying uncertainty, and squeezing the last drop of accuracy by combining models. This chapter surveys the most useful advanced trees and ensembling tricks an expert reaches for.

11.1 Extremely randomized trees (ExtraTrees)

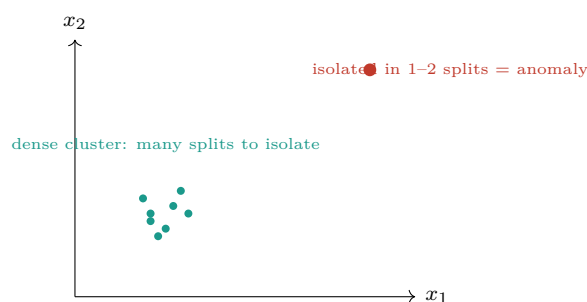
Intuition

ExtraTrees push randomization further than random forests: instead of searching for the best threshold on each candidate feature, they pick split thresholds *at random* and keep the best among those. This adds bias slightly but cuts variance and training time, sometimes beating random forests — and it uses the whole dataset (no bootstrap) by default. A cheap, strong alternative worth trying alongside RF.

11.2 Isolation Forests: anomaly detection

Key idea

An **Isolation Forest** detects outliers by exploiting a simple asymmetry: anomalies are “few and different,” so random splits isolate them in *fewer* steps than normal points. Build trees with random feature/threshold splits; the **average path length** to isolate a point becomes its anomaly score (short path = anomalous). It is unsupervised, linear-time, and scales to high dimensions — a workhorse for fraud, intrusion, and defect detection.



11.3 Gradient-boosted trees for ranking

Intuition

Search engines and recommenders need to *order* items, not score them in isolation. **Learning to rank** with boosted trees (LambdaMART — the basis of much of web search ranking) optimizes ranking metrics (NDCG, MAP) by defining gradients (“lambdas”) from pairwise preferences. XGBoost, LightGBM, and CatBoost all ship **rank** objectives. Tree ensembles remain dominant in production ranking because of speed, accuracy, and feature flexibility.

11.4 Survival, quantile, and probabilistic trees

Intuition

Trees flex to many output types by changing the loss: **survival trees** / **random survival forests** model time-to-event with censoring (medicine, churn, reliability); **quantile regression forests** and quantile-loss boosting predict *intervals* (e.g. 5th–95th percentile) rather than a point, giving calibrated uncertainty; and **NGBoost** (natural gradient boosting) predicts full probability distributions. These bring uncertainty quantification — usually a deep-learning selling point — to tabular trees.

11.5 Bayesian and oblique trees

Intuition

BART (Bayesian Additive Regression Trees) is a sum-of-trees model with priors that keep each tree weak, fit by MCMC — giving full posterior uncertainty and strong performance with little tuning. **Oblique trees** split on *linear combinations* of features ($a^\top x < t$) rather than single axes, capturing diagonal boundaries that standard trees approximate with staircases — at the cost of interpretability. **Explainable Boosting Machines (EBM)** are glass-box GA^2Ms : boosted trees restricted to one/two features at a time, yielding accuracy near full GBMs with full additive interpretability.

11.6 Stacking and blending

Key idea

Stacking trains a meta-learner on the out-of-fold predictions of several base models (e.g. a forest, an XGBoost, a neural net), letting it learn how to best combine them; **blending** is the simpler held-out-set version. Diverse base learners — especially trees *plus* a neural net — often stack into the best tabular models, and stacking is the staple of winning Kaggle solutions. The cost is complexity and inference latency.

Common pitfall

Advanced does not mean better by default. ExtraTrees/oblique/BART can underperform a well-tuned XGBoost while costing more. Stacking shines only with *diverse, decorrelated* base models and rigorous out-of-fold discipline — naive stacking leaks and overfits. And every added layer trades away the interpretability and operational simplicity that made trees

attractive. Reach for these when a measured gain justifies the complexity, not reflexively.

How this powers ML / AI

This menagerie shows trees adapting to jobs once thought to need bespoke models: **Isolation Forests** for unsupervised anomaly detection, **LambdaMART** for the ranking systems behind search and recommendation, **quantile/NGBoost/BART** for uncertainty quantification (the tabular answer to Bayesian deep learning), and **stacking** that fuses trees with neural networks. The recurring theme of the next chapters: trees are not a single algorithm but a flexible substrate that keeps absorbing tasks across machine learning and AI.

Try it in NumPy

```

1 from sklearn.ensemble import IsolationForest
2 import numpy as np
3 rng = np.random.default_rng(0)
4 X = np.r_[rng.normal(0, 0.4, (200, 2)), rng.uniform(-4, 4, (12, 2))]
   # cluster + outliers
5 iso = IsolationForest(contamination=0.05, random_state=0).fit(X)
6 scores = iso.score_samples(X)                # lower = more
   anomalous
7 print("flagged anomalies:", int((iso.predict(X) == -1).sum()))

```

Chapter summary

- **ExtraTrees** randomize thresholds for lower variance; **Isolation Forests** score anomalies by isolation path length.
- Boosted trees handle **ranking** (LambdaMART/NDCG), **survival**, **quantile**, and **distributional** (NGBoost) outputs.
- **BART** (Bayesian), **oblique** (linear splits), and **EBM** (glass-box additive) extend the tree idea in different directions.
- **Stacking/blending** fuse diverse base models (trees + nets) for top accuracy — with leakage risk and added complexity.
- Advanced variants pay off selectively; weigh the gain against lost interpretability and simplicity.

Trees versus the World: When to Use What

Level: Expert

When should you reach for a tree ensemble, and when for something else? This chapter places trees in the broader modeling landscape — against linear models and against deep learning — and explains the *inductive biases* that make trees so well-suited to tabular data. The punchline, backed by large benchmarks: for tables, tuned gradient boosting is still the model to beat.

12.1 Trees vs linear models

Intuition

Linear/logistic models assume a smooth, additive, roughly monotone relationship; they extrapolate, give signed coefficients, and shine when that assumption holds or data is scarce and high-dimensional (with regularization). Trees make no such assumption: they capture nonlinearity and interactions automatically, ignore feature scaling, and resist outliers — but cannot extrapolate and need more data to be stable. Rule of thumb: *simple, linear, well-understood signal* → linear; *complex interactions, mixed messy features* → trees.

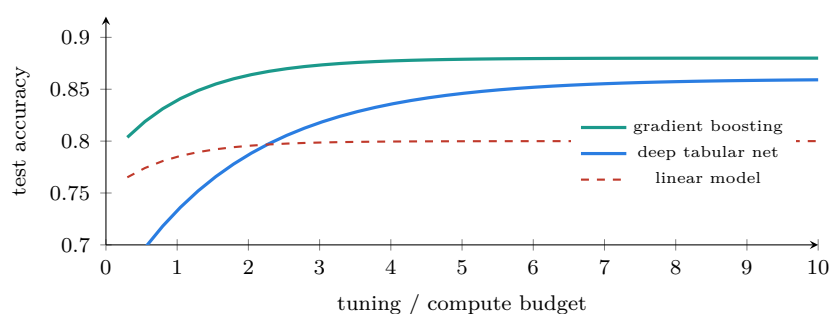
12.2 Why trees beat deep nets on tabular data

Key idea

A landmark benchmark (Grinsztajn et al., 2022) found tree-based models outperform deep nets across most tabular datasets, even after heavy tuning, for three inductive-bias reasons:

1. **Irregular, non-smooth targets.** Real tabular target functions have sharp steps and thresholds; trees fit them naturally, while neural nets are biased toward overly smooth functions.
2. **Uninformative features.** Tables carry many weak/irrelevant columns; trees ignore them via split selection, whereas MLPs are easily distracted.
3. **Non-rotationally-invariant data.** Each tabular column has its own meaning; trees are axis-aligned and respect that, while neural nets' rotational invariance mixes features and discards this structure.

Add lower compute, no scaling, native missing/categorical handling, and exact SHAP, and trees are the pragmatic winner.



12.3 When deep learning wins on tables

Intuition

The picture is nuanced and shifting (2024–2026). Deep models earn their keep when: data is **very large**; features are **high-cardinality categorical** learned as embeddings (recommenders); there is **unstructured signal** (text, images, sequences) mixed with the table, needing end-to-end learning; you want **transfer/foundation models** for small data (**TabPFN** excels there via in-context learning); or you need a single differentiable pipeline. Strong tabular nets (**FT-Transformer**, SAINT, RealMLP) and foundation models now match or beat GBDTs on parts of recent benchmarks (TabArena) — but usually at higher cost.

12.4 A decision guide

Situation	Reach for
New tabular problem, want a strong baseline fast	Gradient boosting (XGBoost/LightGBM/CatBoost)
Robust, low-tuning, free OOB + importances	Random forest
Small data, smooth signal, interpretability	Regularized linear/logistic, GAM, or EBM
Huge data, embeddings, text/image + table	Deep tabular net / multi-modal model
Very small data, want strong zero-tuning model	TabPFN (foundation model)
Maximum accuracy, competition	Stack GBDTs + a tuned neural net

Common pitfall

Don't fight the data type. Using a vanilla MLP on a plain 5,000-row table and expecting it to beat tuned XGBoost usually wastes weeks. Equally, don't force trees onto raw images, audio, or text — those have spatial/sequential structure that CNNs and transformers exploit and trees cannot. And beware benchmark cherry-picking: “method X beats GBDT” often hides dataset selection, tuning asymmetry, or ignored compute cost. Match the model's inductive bias to your data's structure.

How this powers ML / AI

The tree-vs-deep-learning debate *is* a debate about inductive bias — the central concept of modern ML. Trees encode “axis-aligned, piecewise-constant, scale-free” priors perfect for tables; CNNs encode locality/translation; transformers encode permutation-equivariant attention for sequences. Understanding why each wins on its home turf is what lets you choose architectures wisely. The next chapter closes the loop: rather than competing, trees and deep learning increasingly *combine* — the frontier of tabular AI.

Try it in NumPy

```

1 from sklearn.datasets import fetch_covtype
2 from sklearn.linear_model import LogisticRegression
3 from sklearn.ensemble import HistGradientBoostingClassifier
4 from sklearn.model_selection import train_test_split
5 X, y = fetch_covtype(return_X_y=True)
6 Xtr, Xte, ytr, yte = train_test_split(X[:30000], y[:30000],
7                                     random_state=0)
8 for name, m in [("logreg", LogisticRegression(max_iter=200)),
9               ("hist-GBM", HistGradientBoostingClassifier())]:
10     print(name, round(m.fit(Xtr, ytr).score(Xte, yte), 3)) # GBM
11                                     usually wins on tables

```

Chapter summary

- **Linear** models for smooth, additive, scarce-data signal; **trees** for nonlinear, interacting, messy mixed features.
- Trees beat deep nets on most tables due to inductive bias: **irregular targets**, **uninformative features**, **non-rotation-invariance**.
- **Deep learning** wins with huge data, embeddings, multi-modal signal, or foundation models (**TabPFN**) on small data.
- Use the decision guide; gradient boosting is the default first move on a new table.
- The real lesson is **inductive bias**: match the model’s assumptions to the data’s structure.

13

Tree-Based Methods in Machine Learning, Deep Learning, and AI

Level: Capstone

This capstone gathers the threads. Tree-based methods are not a quaint pre-deep-learning relic — they are a living, dominant branch of modern AI, and increasingly they *converge* with deep learning. We tour their reign over tabular ML, the hybrids that fuse trees with neural networks, the differentiable and “deep” trees, their role in interpretability and AutoML, and where the field is heading. Read this as the map that connects everything before it to the wider world of AI.

13.1 The reign over tabular machine learning

Key idea

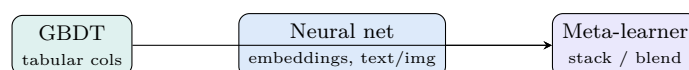
Most of the world’s high-value data is **tabular**: transactions, claims, clicks, sensor logs, electronic health records. On this data, gradient-boosted decision trees (XGBoost, LightGBM, CatBoost) are the dominant production and competition model — accurate, fast, cheap to serve, robust to messy features, and exactly interpretable via TreeSHAP. They win the majority of tabular Kaggle competitions and underpin credit scoring, fraud detection, ad click prediction, pricing, demand forecasting, and recommendation re-ranking. For the data that runs the economy, **trees are the default intelligence**.

13.2 Trees meet deep learning: hybrids

Intuition

The most effective tabular systems often *combine* trees and neural networks rather than choosing:

- **Stacking/ensembling** a GBDT with a neural net (their errors are decorrelated) routinely tops leaderboards and AutoML systems (AutoGluon).
- **Tree-derived features**: use leaf indices of a boosted ensemble as categorical features for a linear or deep model — the classic Facebook “GBDT + logistic regression” ad-CTR pipeline.
- **Neural embeddings** → **trees**: learn embeddings of high-cardinality categoricals (or text/images) with a network, then feed them to a GBDT for the final tabular decision.
- **Distillation**: train a fast, interpretable tree to mimic a large neural net (or vice versa) for deployment or explanation.



decorrelated strengths $\Rightarrow$ best tabular accuracy

13.3 Differentiable and “deep” trees

Key idea

Researchers have made trees *differentiable* so they can be trained by backpropagation and embedded in neural networks:

- **Soft / probabilistic trees** route inputs softly (each split is a sigmoid sending a fraction left/right), so the whole tree is differentiable end-to-end.
- **Deep Neural Decision Forests** graft differentiable trees onto CNN features, jointly learning representation and routing.
- **NODE** (Neural Oblivious Decision Ensembles) stacks differentiable oblivious trees into deep architectures.
- **Attention as soft routing**: the soft, content-based selection in trees parallels attention’s soft selection over tokens — both learn data-dependent paths.

These blur the tree/network boundary, aiming to combine trees’ tabular inductive bias with deep learning’s representation power and gradient training.

13.4 Tabular deep learning and foundation models

Intuition

A fast-moving frontier (2022–2026) builds neural nets *designed* for tables: **TabNet** (sequential attention over features), **FT-Transformer** and **SAINT** (transformers over feature tokens), and **RealMLP** (a strongly-tuned MLP). Most striking is **TabPFN** — a transformer pre-trained on synthetic tasks that does tabular prediction by *in-context learning* (no per-dataset training), excelling on small data. Benchmarks like **TabArena** show these now rival or beat GBDTs in places — yet boosted trees remain the efficiency and robustness baseline everyone measures against. The likely future is pluralistic: trees, tabular transformers, and foundation models coexisting and combining.

13.5 Trees as the interpretability and reasoning layer

Intuition

Trees serve AI even when they are not the predictor. **TreeSHAP** (exact, fast Shapley values) makes tree ensembles the most *auditable* accurate models — decisive in regulated credit, insurance, and healthcare. **Decision trees distilled from black boxes** (including neural nets and now LLM behaviors) give human-readable surrogate explanations. **Decision-tree-structured reasoning** — routing, mixtures of experts, and hierarchical policies — echoes the divide-and-conquer logic of trees inside large models. And classic AI search (game trees, Monte-Carlo Tree Search in AlphaGo) shares the tree’s recursive-decomposition DNA, a reminder of how foundational the structure is.

13.6 Where trees sit in the AI stack

Domain	Best-fit model	Role of trees
Tabular prediction	GBDT (often)	primary model
Recommendation / ranking	deep + GBDT	LambdaMART re-ranking, features
Vision / audio / text	CNNs, transformers	feature consumer, distillation target
Small-data tabular	TabPFN / GBDT	strong baseline / ensemble partner
Explainability / audit	any	TreeSHAP, surrogate trees
Anomaly detection	Isolation Forest	primary unsupervised model

Common pitfall

Two opposite errors. **Tree chauvinism:** insisting on GBDT for problems with rich spatial/sequential/linguistic structure, where deep models clearly win. **Tree dismissal:** assuming deep learning supersedes trees everywhere — on tables, tuned boosting still wins or ties at a fraction of the cost and with better interpretability. The expert keeps both, knows their inductive biases, benchmarks honestly (equal tuning, counting compute), and is happy to combine them.

How this powers ML / AI

The big picture. Trees contribute three enduring things to AI: (1) an *inductive bias* — axis-aligned, piecewise-constant, scale-free — that matches tabular data better than any other model family; (2) *ensembling principles* — bagging’s variance reduction and boosting’s bias reduction — that generalize to deep ensembles, gradient descent in function space, and beyond; and (3) *interpretability infrastructure* (SHAP, surrogates) now used across all of ML. Far from being replaced, trees are converging with deep learning into hybrid and differentiable forms. Master them and you hold both the most practical tool for real-world tabular AI and a deep lens on the principles — bias–variance, regularization, inductive bias, ensembling — that govern intelligent systems everywhere.

Chapter summary

- Gradient-boosted trees **dominate tabular ML** in industry and competitions — the default model for structured data.
- Trees and deep nets increasingly **combine**: stacking, leaf-index features, neural embeddings → trees, and distillation.
- **Differentiable/deep trees** (soft trees, Deep Neural Decision Forests, NODE) merge tree bias with backprop.
- **Tabular deep learning** (TabNet, FT-Transformer) and **foundation models** (TabPFN) now rival GBDTs, but trees remain the efficiency baseline.
- Trees power **interpretability** (TreeSHAP, surrogates), **anomaly detection**, and **ranking**; their legacy is inductive bias, ensembling, and explainability across AI.



Algorithm and Hyperparameter Reference

A compact reference for the formulas, algorithms, and knobs used throughout the book.

A.1 Impurity and split criteria

Criterion	Formula	Used by
Gini impurity	$1 - \sum_k p_k^2$	CART classification (default)
Entropy	$-\sum_k p_k \log_2 p_k$	ID3, C4.5
Information gain	$I(\text{parent}) - \sum_c \frac{n_c}{n} I(c)$	all classification trees
Variance / MSE	$\frac{1}{n} \sum_i (y_i - \bar{y})^2$	regression trees
Cost-complexity	$R_\alpha(T) = R(T) + \alpha T $	CART pruning

A.2 Gradient boosting in one page

1. Initialize $F_0(x) = \arg \min_c \sum_i L(y_i, c)$ (e.g. mean for MSE, log-odds for log-loss).
2. For $m = 1, \dots, M$:
 - (a) Pseudo-residuals $r_i = -[\partial L(y_i, F(x_i)) / \partial F(x_i)]_{F=F_{m-1}}$.
 - (b) Fit a regression tree h_m to $\{(x_i, r_i)\}$.
 - (c) (XGBoost) optimal leaf weight $w_j^* = -\frac{\sum_{i \in j} g_i}{\sum_{i \in j} h_i + \lambda}$ using gradients g_i , Hessians h_i .
 - (d) Update $F_m(x) = F_{m-1}(x) + \nu h_m(x)$ with learning rate ν .
3. Output $F_M(x)$; choose M by early stopping on a validation set.

A.3 Hyperparameter cheat-sheet

Parameter	Effect	Typical range
<code>n_estimators</code> / trees	more trees: RF stabilizes, GBM can overfit	RF 200–1000; GBM via early stop
<code>learning_rate</code> (ν)	step size; smaller = more trees	0.01–0.1 (GBM)
<code>max_depth</code>	tree complexity	3–8 (GBM), deeper/None (RF)
<code>num_leaves</code>	LightGBM complexity	$< 2^{\text{max_depth}}$
<code>min_child_samples</code>	min data per leaf; regularizes	10–100
<code>subsample</code>	row sampling per tree	0.6–1.0
<code>colsample_bytree</code> / <code>max_features</code>	column sampling; decorrelates	0.6–1.0; $\sqrt{p}$ or $p/3$ (RF)
<code>reg_lambda</code> / <code>reg_alpha</code>	L2 / L1 on leaf weights	0–10
<code>gamma</code> / min split gain	minimum gain to split	0–5

A.4 Method selection at a glance

Method	Strength	Use when
Single tree	interpretable flowchart	you need one readable rule set
Random forest	robust, low-tuning, OOB	strong baseline, little tuning time
Gradient boosting	highest tabular accuracy	you can tune + validate carefully
Isolation forest	unsupervised anomalies	outlier/fraud/defect detection
EBM / GA ² M	glass-box accuracy	accuracy <i>and</i> full interpretability
Stacking (trees+NN)	top accuracy	competitions, max performance

A.5 Key references

- Breiman, Friedman, Olshen, Stone, *Classification and Regression Trees* (1984).
- Breiman, “Bagging Predictors” (1996); “Random Forests” (2001).
- Freund & Schapire, AdaBoost (1997); Friedman, “Greedy Function Approximation: A Gradient Boosting Machine” (2001).
- Chen & Guestrin, “XGBoost” (2016); Ke et al., “LightGBM” (2017); Prokhorenkova et al., “CatBoost” (2018).
- Lundberg & Lee, “A Unified Approach to Interpreting Model Predictions” (SHAP, 2017).

- Grinsztajn et al., “Why do tree-based models still outperform deep learning on tabular data?” (2022).
- Hastie, Tibshirani, Friedman, *The Elements of Statistical Learning*; James et al., *An Introduction to Statistical Learning*.