

Prediction Without Training: Google TabFM

and the Rise of Tabular Foundation Models

Zero-shot classification and regression on enterprise tables.
A new era of foundation models for structured data.



Prashant Akhawat
akhawat.com
Building AI-Native Enterprises



Single
Forward Pass



Works Across
Tables



Private. Secure.
Scalable.



Built for the
Enterprise

CLASSIFICATION
(Zero-Shot)



No Training
Needed

REGRESSION
(Zero-Shot)



No Training
Needed

Feature 01	Feature 02	Feature 03	Feature 04	Feature ...	Target
1.2	7.8	3.4	9.1	...	2.7
9.7	5.6	8.9	0.3	5.4	8.12
9.7	6.3	7.1	2.2	1.8	3.75
7.3	...	...	...	...	6.94

CXO INTELLIGENCE SERIES • A TABULAR INTELLIGENCE BRIEFING

Prediction Without Training: Google TabFM and the Rise of Tabular Foundation Models

An AI model that reads a business table and fills in the missing answers instantly, with no training required. What it is, where it helps, how to plug it into your enterprise, and what it costs to run.

Prashant Akhawat

Building AI-Native Enterprises

July 2026 · 26 min read

EXECUTIVE SUMMARY

The last major data type gets its foundation model

Picture a Monday next quarter, 9:14 am. A retention analyst types one query into the data warehouse and gets back churn scores for every active customer. No project was raised, no model was built, no data scientist was booked. Twelve months ago, the same request took a quarter of a year. This briefing explains the release that makes that morning possible.

On 30 June 2026, Google Research released TabFM, an AI model that makes predictions on business data tables it has never seen before. There is no training step, no tuning, and no lengthy data preparation. You show it your table, and it answers. The model is free to download on Hugging Face and GitHub, and Google has announced that it will soon be built directly into BigQuery, its data warehouse, so that asking for a prediction becomes one line of SQL through a simple `AI.PREDICT` command.

IN PLAIN WORDS

Imagine a spreadsheet where some rows are complete and some have one blank column: which customers left, which loans went bad, which machines failed. Until now, filling those blanks meant a specialist team building a custom model for that one spreadsheet, a project of weeks. TabFM is a ready-made brain that reads the completed rows and fills in the blanks immediately. One model, any table, answers in seconds.

WHY THIS MATTERS TO A CXO

Cost. Predictive use cases stop being projects. The per-model build, tuning, and maintenance bill largely disappears.

Speed. Answers arrive at the speed of a database query, so decisions stop waiting in the data science queue.

Access. Anyone on the team who can write SQL will be able to run predictions, not only specialists.

The significance is larger than one model release. Tables, the rows and columns inside ERP, CRM, core banking, claims, lab, and billing systems, are where most enterprise value actually sits. AI has already transformed how we work with text, images, and code. Tables were the holdout: every prediction still needed its own hand-built, hand-tuned model. TabFM, alongside TabPFN from Prior Labs (now backed by SAP) and Mitra from Amazon, signals that this era is closing.

Table A. The release at a glance.

At a glance	
1 pass	One pass through the model per batch of predictions. No training, no tuning, no retraining schedule
51 datasets	Public benchmark datasets evaluated: 38 category tasks, 13 number tasks, from 700 to 150,000 rows
100M+ tables	Computer-generated practice tables used to pretrain the model
Rank #1	TabFM-Ensemble ranked first on the public TabArena leaderboard for both prediction types

This briefing is arranged business first: the edge, the use cases, and the financial impact come before the technology. The how-it-works, architecture, and hardware sections follow for readers who want to go deeper, and the playbook and takeaways close with what to do next.

WHY IT MATTERS

The edge for businesses

The advantage TabFM offers is not primarily accuracy. It is the removal of an entire cost structure that sits between a business question and a prediction.

- **Time to prediction.** The traditional cycle of scoping, feature engineering, training, tuning, validation, and deployment collapses into a single inference call. Prediction becomes available at the speed of a query, not the speed of a project.
- **Democratization through SQL.** Once TabFM ships inside BigQuery through AI.PREDICT, any analyst who can write SQL can run classification and regression on warehouse tables. Predictive analytics stops being gated on scarce data science capacity.

- **A one-to-many operating model.** One pretrained model replaces a portfolio of per-dataset models, which means less machinery to maintain: fewer pipelines, fewer retraining jobs, fewer systems to watch, and fewer silent failures from stale models.
- **Strength on small data.** In-context learners are strongest exactly where classical ML is weakest, on small and medium datasets where there is not enough data to tune a bespoke model well. Many high-value enterprise problems, from clinical cohorts to niche product lines, live in this regime.
- **Low-cost experimentation.** Because the model was never trained on any customer dataset, evaluation is honest by construction. Teams can benchmark TabFM against incumbent models in days and adopt it only where it wins.



Illustrative orders of magnitude. The point is not the exact numbers; it is that the unit changes from months to moments.

Illustrative orders of magnitude: the unit of delivery changes from months to moments.

Wherever an analyst today would export a table and ask the data science team to build a scoring model, TabFM converts that request into a query.

```
-- Conceptual illustration of the announced BigQuery integration.
-- Exact syntax will be confirmed when AI.PREDICT reaches general availability.
```

```
SELECT
  customer_id,
  AI.PREDICT(
    MODEL tabfm,
    TABLE curated.churn_features,      -- labeled history as context
    INPUT curated.active_customers     -- rows to score
  ) AS churn_probability
FROM curated.active_customers;
```

The pattern will be familiar to leaders who watched TimesFM change time-series forecasting. Google frames TabFM explicitly as the tabular counterpart of that shift: the same zero-shot logic, applied to the data type that carries the most enterprise weight.

BUSINESS TAKEAWAY When every SQL-literate analyst can run predictions, demand for prediction expands rather than merely getting cheaper. Plan governance and intake for volume, not just savings.

APPLICATIONS

Where businesses can use it

TabFM addresses the broad middle of enterprise prediction: supervised classification and regression on structured tables. Ten concrete use cases below show the shape of the opportunity; Table 1 then maps the pattern by sector.

1. **Customer churn early warning.** Classify which customers are likely to cancel in the next 90 days from usage, billing, and support tables. Segment-level churn models that were never economical to build become a weekly query.
2. **Credit scoring for thin-file segments.** Score default probability for MSME, new-to-credit, and niche portfolios whose history is too small to tune a bespoke model well. Small-data regimes are exactly where in-context learning is strongest.
3. **Fraud and anomaly triage.** Rank flagged transactions or claims by risk so investigators work the highest-risk items first. Re-scoring adapts to shifting patterns through fresh context, with no retraining pipeline to rebuild.
4. **Claim severity at first notice of loss.** Regress the expected cost of an insurance claim on the day it is reported, from structured FNOL tables. Early severity estimates sharpen reserving and route simple claims to fast-track handling.
5. **Clinical trial dropout prediction.** Flag participants at risk of discontinuation from visit, adherence, and demographic tables so sites can intervene early. Trial cohorts are archetypal small data; per-study models were rarely justified.
6. **Predictive maintenance from sensor snapshots.** Classify which machines, lines, or vehicles are likely to fail before the next service window from aggregated sensor tables. One model covers every plant and route instead of a model per asset class.
7. **SKU-level stock-out and demand risk.** Predict which store-and-SKU combinations will stock out or overstock next week. The long tail of SKUs never justified bespoke models; a single query now covers all of them.
8. **Lead and propensity scoring.** Score open pipeline and campaign audiences for conversion likelihood directly from CRM tables. Marketing analysts self-serve on the managed path without waiting for a data science sprint.
9. **Invoice payment delay prediction.** Regress expected days-to-pay for every open invoice to prioritize collections and forecast working capital. A finance-team prediction the moment AI.PREDICT reaches the warehouse.
10. **Employee attrition risk.** Flag flight-risk roles and teams from HR tables to guide retention conversations. A consequential-decision use case by design: run it with human-in-loop review and the full governance spine of Figure 6.

The long tail of prediction demand (illustrative)

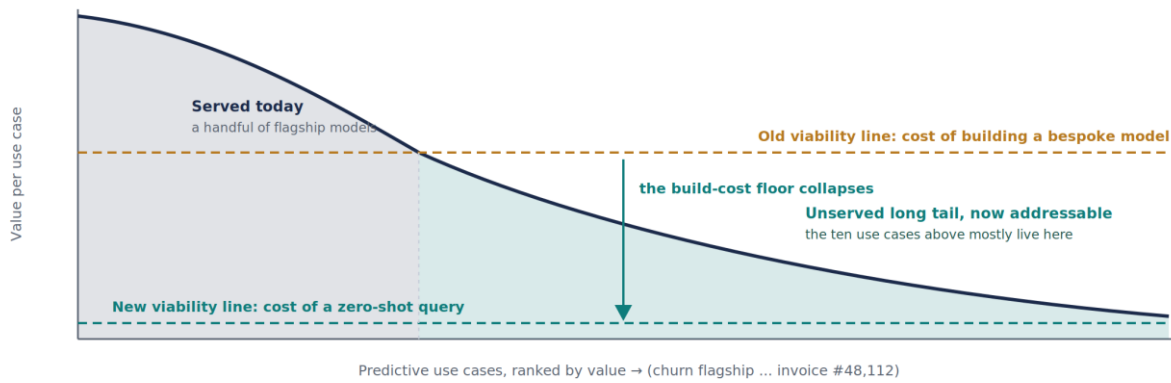


Figure 1. The thought experiment behind the ten use cases: most prediction demand was never met, because every prediction cost a project. When the viability line drops from the cost of a bespoke model to the cost of a query, the addressable surface does not grow at the margin; it multiplies. Curve and thresholds are illustrative.

Table 1. Representative applications of zero-shot tabular prediction by sector.

Sector	Representative predictions	Why zero-shot changes the economics
Banking and financial services	Credit default scoring, fraud triage, transaction anomaly ranking, collections prioritization, insurance claim severity	Hundreds of portfolio-level and segment-level models can be replaced by one callable model, cutting development queues from months to days
Life sciences and pharma	Clinical trial dropout risk, adverse event likelihood ranking, batch quality deviation prediction, site enrollment forecasting	Small, high-value datasets are common in this sector, and in-context learners are strongest precisely where training data is scarce
Retail and consumer	Customer churn, propensity to buy, return likelihood, promotion response, stock-out risk	Merchandising and CRM analysts can score tables directly in the warehouse without waiting for data science capacity
Manufacturing and logistics	Predictive maintenance classification, defect detection from sensor tables, delivery delay prediction, supplier risk scoring	Every plant, line, and route no longer needs its own trained model; one foundation model adapts to each table it is shown
Media and information services	Subscriber churn, content demand scoring, advertising yield prediction, licensing revenue forecasting inputs	Rapid experimentation on new segments becomes a query rather than a project
Government and public sector	Benefit fraud detection, tax anomaly flagging, service demand forecasting, infrastructure failure risk	Agencies with limited ML staffing gain access to competitive predictive accuracy through SQL-level tooling

In regulated environments such as pharmacovigilance and BFSI document operations, where platforms such as AOTM already apply confidence-tiered automation, zero-shot tabular predictors slot in as scoring components that can be validated per use rather than retrained per dataset, which materially simplifies computerised system validation workloads.

Early real-world evidence

TabFM itself is only weeks old, but the model class it belongs to is already producing named, public enterprise evidence. Creditplus Bank AG, a German consumer bank, has publicly endorsed tabular foundation models

for credit decisioning. Taktile, a decision-intelligence platform used by financial services firms, benchmarked TabPFN on a real account-opening fraud dataset and found the advantage over boosted trees grew as the data got sparser, exactly the regime where fraud labels are scarce. Mission Lane, a US credit card company, published an engineering evaluation of the TabICL architecture against its production LightGBM credit-risk baseline and concluded that tabular foundation models reach competitive accuracy with zero task-specific training. Prior Labs additionally lists live application areas spanning automotive finance fraud, mortgage and SME credit risk, and 180-day churn prediction.

“

An exciting opportunity to adopt tabular foundation models and unlock their potential.

Dr. Dietrich Eherler

Head of Statistical Analysis and Development, Creditplus Bank AG

Source: Prior Labs, 2026

BUSINESS TAKEAWAY Start where models are many, small, and expensive to maintain, not where a single flagship model already performs. The ten use cases above sit almost entirely in the long tail that Figure 1 shows becoming viable.

ECONOMICS & OPERATING MODEL

What changes on the P&L, and in the team

The financial impact of zero-shot tabular prediction lands on three lines: the cost to build a predictive use case, the cost to keep it alive, and the cost of the infrastructure underneath. The table below contrasts the two operating models on the drivers that CFOs and CDOs actually budget for. Figures are directional patterns, not vendor quotes; every organization should price its own baseline in the benchmarking step of the playbook.

Table 2. Cost-driver comparison: traditional supervised models versus the zero-shot foundation model path.

Cost driver	Traditional supervised model	Zero-shot foundation model path
Build effort per use case	Typically six to twelve weeks of data science effort: feature engineering, training, tuning, validation	Days: prediction view design, benchmark run, validation evidence. No training or tuning stage exists
Specialist dependency	Every use case queues on scarce data scientists	SQL-literate analysts self-serve on the managed path; specialists focus on validation and edge cases
Infrastructure	Training compute per model, plus serving infrastructure per model	Serverless per-query cost on the managed path, or a small shared GPU pool (Table 3) on the self-hosted path
Maintenance	Retraining pipelines, drift-triggered rebuilds, per-model monitoring across the portfolio	One pinned model version, input drift monitoring, periodic revalidation. No retraining pipelines at all
Portfolio shape	N models, N pipelines, N registry entries, N silent failure modes	One model, N validated use cases. Governance effort scales with use cases, not with models
Cost of experimentation	High enough that weak ideas are never tested	Near zero, so the idea funnel widens and portfolio decisions improve

The KPIs worth tracking

A CXO sponsoring this shift should instrument it. Six measures capture whether the economics are actually materializing: time to first prediction for a new use case (target: under one week); the share of new predictive requests served through the zero-shot path; the net change in the maintained model portfolio count; cost per thousand predictions on each path; the number of analyst-initiated predictions per month, which measures democratization; and the challenger win rate, the fraction of use cases where the foundation model beats the incumbent, which tells you whether adoption should accelerate or pause.

The first dividend of zero-shot prediction is cost. The second, larger dividend is what your specialists do with the time it returns.

Where the effort goes: the migration up the stack (illustrative share of team effort)

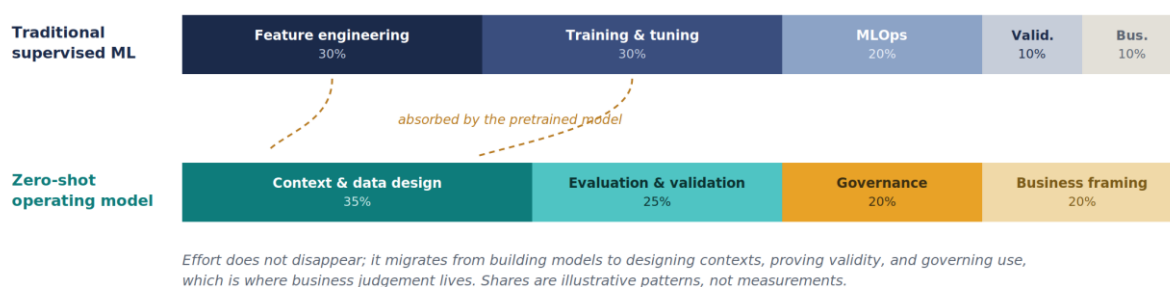


Figure 2. The talent inversion. The two largest effort blocks of the traditional cycle are absorbed into the pretrained model, and the remaining work shifts up the stack. This is why redeployment, not reduction, captures the second dividend. Shares are illustrative.

The talent implication

The roles do not disappear; they move up the stack. Feature engineers become context designers, deciding which rows and columns constitute an informative, leakage-free prompt. MLOps engineers become model governance engineers, owning version pinning, evaluation harnesses, and drift monitoring. Data scientists shift from fitting models to the work that zero-shot cannot do: causal analysis, experiment design, and the judgement calls in the champion-challenger process. Organizations that redeploy along these lines convert a cost reduction into a capability expansion; those that simply cut headcount forfeit the second, larger dividend. How AI redraws these organizational lines is a theme the author examines at length in a forthcoming book on the AI-native firm, publishing later this year.

BUSINESS TAKEAWAY Fund this as an operating model change, not a tooling purchase. The savings come from the maintenance line and the queue, and the upside comes from redeploying specialists to work that training-based ML never left time for.

BACKGROUND

Why tabular data was the last frontier

Enterprises do not run on paragraphs. They run on tables: customer lists, transaction ledgers, sensor logs, claims registers, trial data, and inventory snapshots. Predicting things from this data, who will leave, which transaction is fraud, who will repay, what will sell, has for twenty years been the job of specialist statistical models. The industry workhorses are called gradient-boosted trees, and XGBoost is the best-known name.

These methods are accurate, but expensive to live with. Every new dataset needs its own model, built by hand. The useful signals must be crafted from raw columns (the trade calls this feature engineering), the settings must be searched and adjusted, and the whole thing must be rebuilt regularly as the business changes. For a large enterprise this multiplies into hundreds of custom models, each with its own upkeep and its own queue of waiting business teams.

There is also a simple reason AI arrived late to tables. Language models learn from text, and the internet supplied a near-unlimited amount of it. Tables are different: their order does not matter (swap two rows and nothing changes), and companies do not publish their real tables, because they contain confidential records. Both problems had to be solved before one general model for tables could exist.

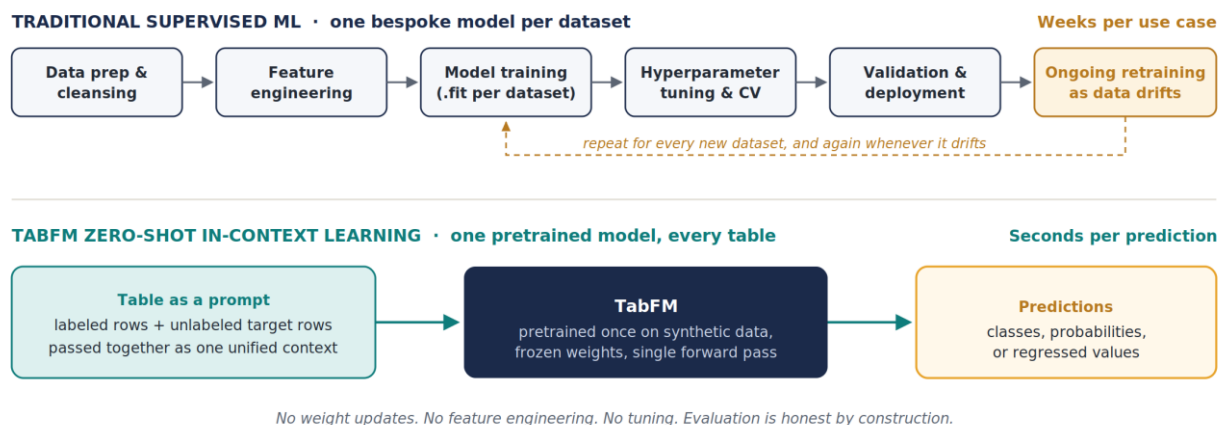


Figure 3. From pipeline to prompt. The traditional per-dataset supervised cycle, with its perpetual retraining loop, collapses into a single in-context inference call against one pretrained model.

BUSINESS TAKEAWAY The bottleneck in enterprise prediction has never been algorithms; it has been the per-dataset cost of building and maintaining them. That cost line, not accuracy, is what tabular foundation models attack.

INSIDE THE MODEL

What Google launched, and how TabFM works

TabFM works the way ChatGPT-style models work with text, but for tables. Instead of being trained on your data, it reads the whole table in one go: the rows where the answer is already known, and the rows where it is not. It then fills in the missing answers directly. Nothing inside the model changes in the process; it simply reads and responds. Researchers call this in-context learning, because the model learns from the examples contained in your request itself.

Under the hood, TabFM combines the best ideas of the two strongest earlier designs, TabPFN and TabICL, in three steps. First, it studies how the rows and columns of your table relate to one another, which is the understanding a data scientist would otherwise build by hand. Second, it compresses each row into a compact summary. Third, a final stage reasons over those summaries to produce the predictions. The compression step is what keeps it fast even on larger tables.

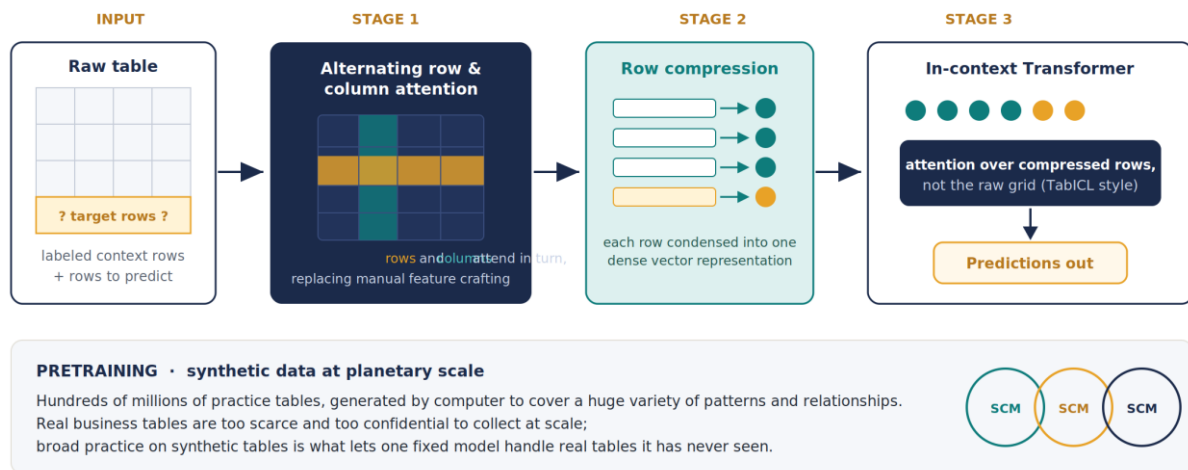


Figure 4. How TabFM reads a table: it first studies how the rows and columns relate, then summarizes each row, then predicts from the summaries. It learned this skill from hundreds of millions of computer-generated practice tables, because real business tables are too confidential to collect at scale.

Benchmark performance on TabArena

Google tested TabFM on TabArena, a public leaderboard that ranks prediction methods the way chess players are ranked: by who beats whom, head to head, across 38 category-prediction datasets and 13 number-prediction datasets ranging from 700 to 150,000 rows. Two versions were tested. The base TabFM answers straight out of the box. TabFM-Ensemble layers a set of accuracy-boosting tricks on top, including running 32 variations of the problem and blending their answers. In the published results the ensemble version ranked first on both boards, with the base version close behind, ahead of the heavily tuned models that expert teams spend weeks on.

TabArena Elo ranking, classification (illustrative)

Higher is better · bar lengths indicative of published rank order



(T+E) = tuned + ensembled · (D) = default. Zero effort outranking weeks of tuning is the headline.

Figure 5. Illustrative representation of Google's published TabArena rank order. Exact values are published on the TabFM GitHub repository and the TabArena leaderboard, and should be consulted before procurement decisions.

The plain takeaway: a model that needs no training now matches or beats models that consume weeks of expert effort per dataset, at least on public benchmarks. An independent reviewer has already reproduced the core claims, and also found an early bug in multi-GPU set-ups that Google fixed within days, a healthy sign of an actively maintained open-source release.

Ten terms, translated

- **Foundation model:** one large pre-trained model reused for many tasks, instead of building one model per task.

- **Zero-shot:** works on data it has never seen before, with no training step.
- **In-context learning:** the model learns from the examples inside your request, at the moment you ask.
- **Inference:** the act of asking a model for an answer, as opposed to training it.
- **Tabular data:** information in rows and columns, like spreadsheets and database tables.
- **Classification and regression:** predicting a category (will this customer leave: yes or no) and predicting a number (what will this claim cost).
- **Feature engineering:** hand-crafting the input signals a traditional model needs. TabFM skips this step.
- **GPU:** the specialized chip AI runs on. This model class needs only modest ones, or none at all on the managed path.
- **Data warehouse:** the central store of a company's analytical data; BigQuery is Google's.
- **Champion-challenger:** running the new model against the current one on the same data, and keeping whichever wins.

BUSINESS TAKEAWAY The architecture details matter to leaders for one reason: they eliminate the two largest effort line items in tabular ML, feature engineering and hyperparameter tuning. Competitive advantage shifts from model building to data quality and context design.

BLUEPRINT

An integrated reference architecture

Adopting a tabular foundation model is not a model decision; it is a platform decision. The architecture below shows a complete, production-oriented enterprise integration with two deliberate prediction paths: a managed path where prediction runs serverlessly inside BigQuery through AI.PREDICT, and a self-hosted path where the open TabFM weights run as a GPU-backed inference service for organizations with data residency, sovereignty, latency, or air-gap requirements. A governance spine runs beneath both paths, because a model that never retrain still needs version pinning, evaluation, drift monitoring, and audit.

TABULAR FOUNDATION MODEL · ENTERPRISE REFERENCE ARCHITECTURE

Managed and self-hosted prediction paths over a governed data platform · Google Cloud reference, portable by pattern

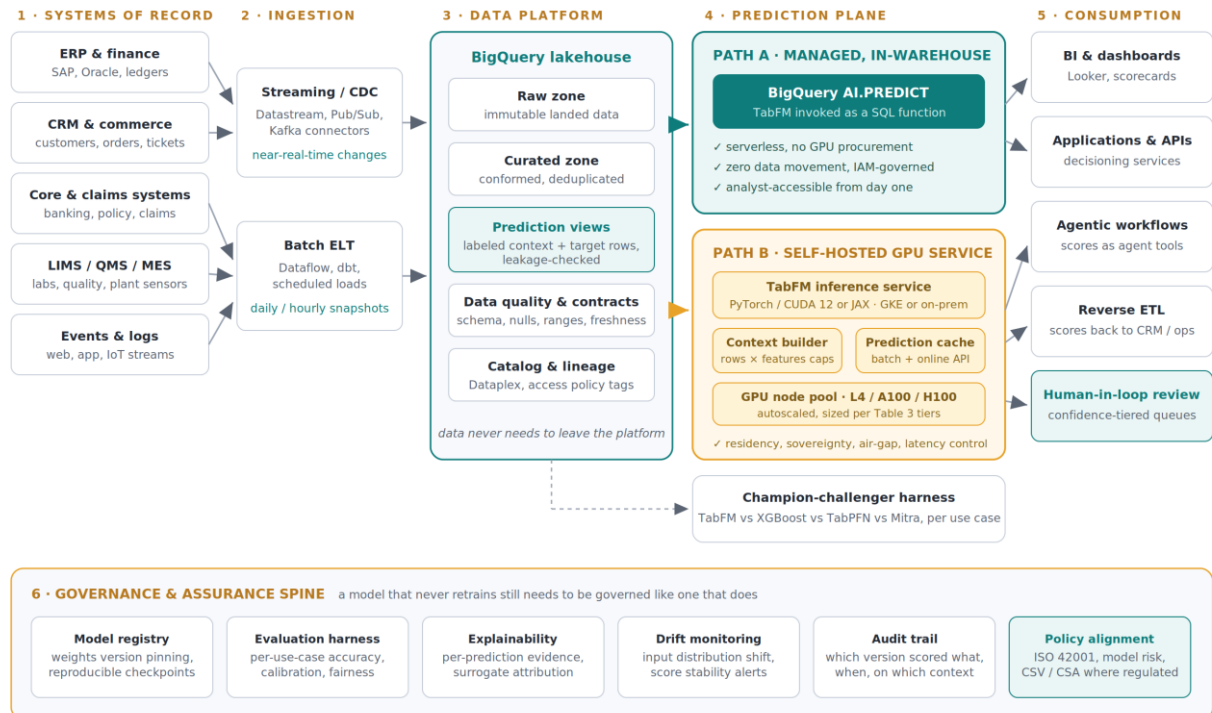


Figure 6. Integrated reference architecture. Two prediction paths, managed AI.PREDICT inside BigQuery and a self-hosted GPU inference service, share one governed data platform and one assurance spine. The champion-challenger harness keeps the foundation model honest against XGBoost, TabPFN, and Mitra per use case. The pattern is drawn on Google Cloud but is portable to any lakehouse estate.

Design decisions that matter

- **Prediction views, not feature stores, become the contract.** The unit of input to an in-context model is a leakage-checked view (checked so that no column accidentally reveals the answer) containing labeled example rows plus the rows to score. Data quality checks shift from feature pipelines to these views.
- **The context builder is the new performance lever.** Because inference cost and memory scale with rows times features in the context, a service that selects the most informative context rows (sampling, stratification, recency windows) governs both accuracy and GPU spend.
- **Champion-challenger is permanent, not a phase.** Zero-shot evaluation is cheap, so every scoring use case should continuously compare TabFM against the incumbent gradient-boosted model and at least one rival foundation model.
- **Governance survives the disappearance of training.** There is no retraining pipeline to audit, so assurance shifts to weights version pinning, per-use-case validation evidence, input drift monitoring, and confidence-tiered human review on consequential decisions.

BUSINESS TAKEAWAY Treat adoption as a platform decision with two paths. Default to the managed warehouse path; reserve GPU self-hosting for residency, sovereignty, air-gap, and latency constraints that money cannot route around.

INFRASTRUCTURE

What it takes to run: hardware and prerequisites

The infrastructure story is one of TabFM’s quiet advantages, because there are three ways to pay for inference, and only one of them involves buying GPUs.

First, the facts as published. TabFM ships with two backends: a JAX backend whose base install is CPU-only, with a CUDA extra that pulls the CUDA 12 plugin and NVIDIA libraries for GPU runs, and a PyTorch backend supporting CPU and GPU. Weights download automatically from Hugging Face, Python 3.11 or higher is required, and the model plugs into Python through the familiar scikit-learn style interface. Google has not yet published official GPU sizing for TabFM, so the tiers below are indicative, grounded in the published guidance for the same in-context model class, where TabPFN documentation states that GPUs with roughly 8 GB of VRAM work well, 16 GB is needed for some large datasets, and CPU inference is feasible only for small tables of around a thousand samples. The governing rule is simple: the bigger the table you show it, rows times columns, the more memory it needs, because the whole dataset is the prompt.

PREREQUISITES: WHAT YOU NEED BEFORE DAY ONE

Technical. Either a BigQuery estate (the managed path: nothing to install and no GPU to buy), or for self-hosting: Python 3.11 or higher, the PyTorch backend with CUDA 12 for GPU or the JAX backend for CPU-only evaluation, model weights pulled automatically from Hugging Face, and a GPU tier chosen from Table 3.

Data. A flat, single table with a clearly defined target column; a labeled history, where even a few hundred completed rows is enough to start; a leakage check so that no column accidentally reveals the answer; and a size inside the proven envelope, benchmarked to roughly 150,000 rows today.

People and process. An SQL-literate analyst to own the first pilot; a named validator responsible for accuracy evidence; a governance sign-off route agreed before any production use; and an incumbent baseline model to beat in the champion-challenger bake-off.

Choosing a deployment path

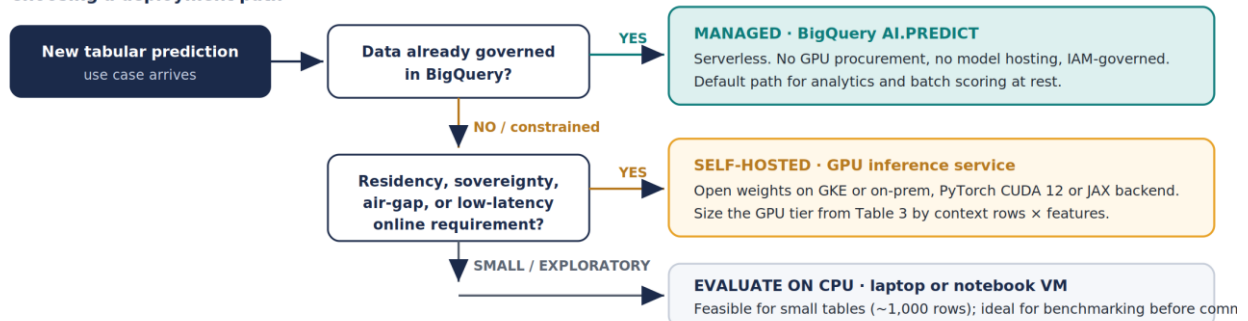


Figure 7. Deployment decision flow. Most organizations should exhaust the managed path before procuring GPUs; the self-hosted path exists for constraints that money cannot route around.

Table 3. Indicative GPU sizing tiers for self-hosted TabFM inference, pending official guidance.

Tier	Typical workload	Indicative GPU	Cloud examples	Notes
T0 Evaluate	Benchmarking, notebooks, tables up to ~1,000 rows	None (CPU), or any workstation GPU	Any VM; Colab / Vertex Workbench	CPU-only JAX base install; enough to run the champion-challenger bake-off
T1 Departmental	Ad hoc scoring, small to mid tables, tens of thousands of rows	NVIDIA T4 16 GB or L4 24 GB, single card	GCP G2 (L4); AWS g4dn / g6	Aligns with published guidance for this model class: ~8 GB VRAM works, 16

				GB for larger datasets
T2 Production service	Scheduled batch scoring plus online API, ~100K-row contexts, concurrent use cases	L40S 48 GB or A100 40 GB, 1 to 2 cards with autoscaling	GCP A2; AWS g6e / p4d	Context builder caps rows times features; prediction cache absorbs repeat queries
T3 Heavy ensemble	TabFM-Ensemble (32-way) runs, wide tables, largest benchmarked contexts (~150K samples)	A100 80 GB or H100 80 GB	GCP A3 (H100); AWS p5	Ensembles multiply forward passes; batch windows matter more than raw card count
Managed	Everything the warehouse can host	None procured by you	BigQuery AI.PREDICT (serverless)	Google operates the accelerators; watch pricing and rate limits at GA

SIZING DISCLAIMER

Google has not published official TabFM hardware requirements at the time of writing. Tiers T1 to T3 above are indicative estimates derived from the published behavior of the same in-context tabular model class, where memory scales with context size (rows times features). Validate on representative workloads before procurement, and note that ensemble configurations multiply compute roughly linearly with ensemble width.

Two cost observations follow. First, the GPU bill for a tabular foundation model is a fraction of LLM economics: this model class runs comfortably on single mid-range inference cards, not multi-node H100 clusters, and CPU evaluation is free. Second, the managed path inverts the capex question entirely: for BigQuery estates, the marginal cost of a prediction becomes a query cost, which is why the AI.PREDICT pricing model at general availability deserves close attention.

BUSINESS TAKEAWAY Tabular foundation models are inexpensive to run by AI standards: single mid-range GPUs, or none at all on the managed path. Budget scrutiny belongs on AI.PREDICT pricing and rate limits at GA, not on capex.

MARKET STRUCTURE

The competitive landscape

TabFM enters a field that has moved quickly since TabPFN appeared in 2022 and was published in Nature in January 2025. Three ecosystems now matter, aligned with three platform camps.

Three camps converge on one capability

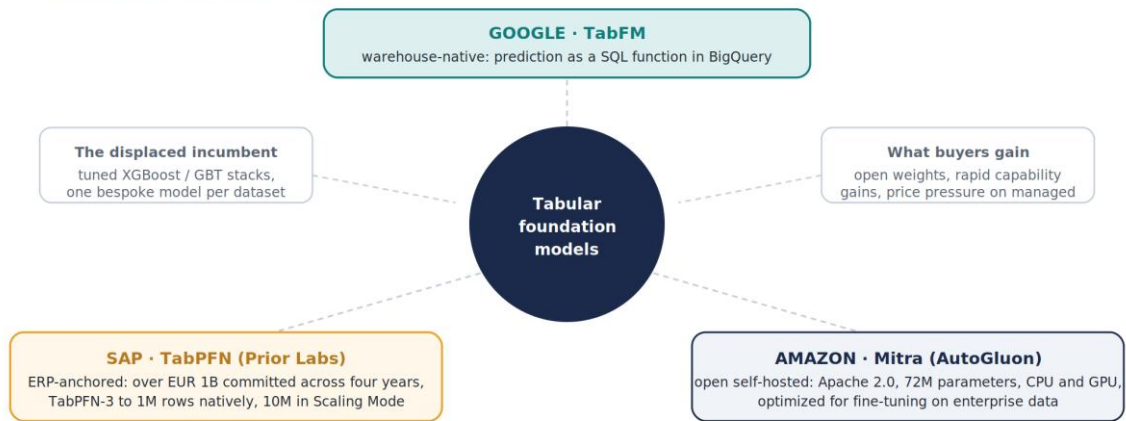


Figure 8. The tabular AI platform contest. SAP anchors tabular intelligence to the ERP layer, Amazon to the open self-hosted AutoML stack, and Google to the data warehouse. Three ecosystems converging on one capability within eighteen months means the capability is infrastructure, not curiosity.

TabPFN, from Prior Labs, is the incumbent pioneer. Its versions have scaled rapidly: TabPFN-3, released in May 2026, supports up to one million rows natively, and a Scaling Mode has been benchmarked to ten million rows. The strategic signal is stronger still: in May 2026, SAP announced it would acquire the eighteen-month-old company and invest over EUR 1 billion across four years to grow it into an independent frontier lab for structured business data. Prior Labs also offers an enterprise edition with a compression option for faster responses, on-premises deployment, and named adopters across credit risk and predictive maintenance.

Mitra, from Amazon, is the AWS-ecosystem answer. A compact model (72 million parameters, small by AI standards) trained purely on carefully designed synthetic data, Mitra ships inside the AutoGluon toolkit under a permissive open-source license, runs on ordinary CPUs as well as GPUs, and is strongest on small datasets. It is built to be fine-tuned, which suits enterprises that want to adapt a model to their own data while hosting it themselves.

TabFM, from Google, is the warehouse-native play. Its technical contribution is combining the accuracy of TabPFN’s design with the speed of TabICL’s, and its distribution advantage is BigQuery. Prediction inside the warehouse, invoked by SQL, with no data movement, is a distinct go-to-market position that neither rival currently matches.

Table 4. Head-to-head comparison, including the incumbent all three aim to displace: the tuned gradient-boosted tree stack.

Dimension	Google TabFM	TabPFN (Prior Labs, SAP)	Amazon Mitra	XGBoost / AutoGluon class
Paradigm	Zero-shot in-context learning, single forward pass	Zero-shot ICL with fine-tuning options	ICL, optimized for fine-tuning	Supervised training per dataset, tuned and ensembled
Architecture	Hybrid: alternating row/column attention plus compressed-row ICL Transformer	Alternating row and column attention over the full grid	12-layer Transformer, 72M parameters, mixed synthetic priors	Gradient-boosted trees and stacked ensembles
Training data	Hundreds of millions of synthetic datasets from structural causal models	Hundreds of millions of synthetic datasets	Purely synthetic data from curated prior mixtures	Customer data per task

Scale posture	Benchmarked from 700 to 150,000 samples; compressed-row design targets larger tables	TabPFN-3: up to 1M rows natively; Scaling Mode benchmarked to 10M rows	Strongest under ~5,000 samples and 100 features	Scales to very large data with engineering effort
Hardware	CPU-capable (JAX base), CUDA 12 GPU extra; single mid-range inference card typical	GPU recommended, ~8 GB VRAM typical, 16 GB for large data; hosted API and distillation available	Runs on CPU and GPU	CPU-native; GPUs optional
Distribution	Open weights on Hugging Face and GitHub; BigQuery AI.PREDICT announced	Open-source variants plus commercial Enterprise Edition; SAP investment of over EUR 1B	Apache 2.0 open weights, integrated in AutoGluon	Fully open-source, mature ecosystem
Enterprise angle	Warehouse-native prediction inside the Google Cloud data stack	On-premises and private cloud deployment, ERP alignment through SAP	Privacy-conscious self-hosting on the AWS stack	Total control, full explainability tooling, highest operational burden

BUSINESS TAKEAWAY Match the ecosystem to the estate: BigQuery estates lean TabFM, SAP estates lean TabPFN, AWS estates lean Mitra. Keep at least two in the champion-challenger harness to preserve negotiating leverage.

FOUR DATES AND SIGNALS TO WATCH

AI.PREDICT general availability, and above all its pricing and rate limits, which will set the marginal cost of a prediction for BigQuery estates.

TabArena leaderboard movements over the next two quarters, as all three camps respond to each other's releases.

Prior Labs' first enterprise releases under SAP ownership, which will show how tightly tabular intelligence gets tied to the ERP layer.

Mitra and AutoGluon updates, the signal for how far the open, self-hosted, fine-tunable path advances.

RISK & GOVERNANCE

Where caution is warranted

Enthusiasm should be paired with discipline, particularly in regulated industries.

- **Benchmarks are not your data.** Leadership on TabArena does not guarantee superiority on a specific enterprise table. Messy production data, extreme class imbalance, leakage risks, and very wide schemas all deserve side-by-side evaluation against incumbent models before any switch.
- **Explainability.** Business owners, model risk teams, and regulators will ask why a row was scored high. Tree ensembles have mature explainability tooling; foundation-model explanations for tabular predictions are younger. Firms subject to model risk management or GxP validation obligations should demand per-use-case explainability evidence.
- **Licensing and reproducibility.** Verify license terms for commercial production use, confirm how weights are versioned, and establish which model version produced which prediction, since reproducibility is a validation requirement in life sciences and BFSI contexts.

- **Data governance.** Warehouse-native inference is convenient, but firms with residency, sovereignty, or air-gap constraints should note that open weights permit self-hosted deployment, and should decide deliberately between the managed and self-hosted paths.
- **Scope limits.** TabFM addresses classification and regression on flat tables. Multi-table relational reasoning, streaming feature context, and very large-scale ranking systems remain outside the zero-shot envelope for now, and gradient-boosted trees will retain a role at extreme scale.

Which tool for which prediction problem

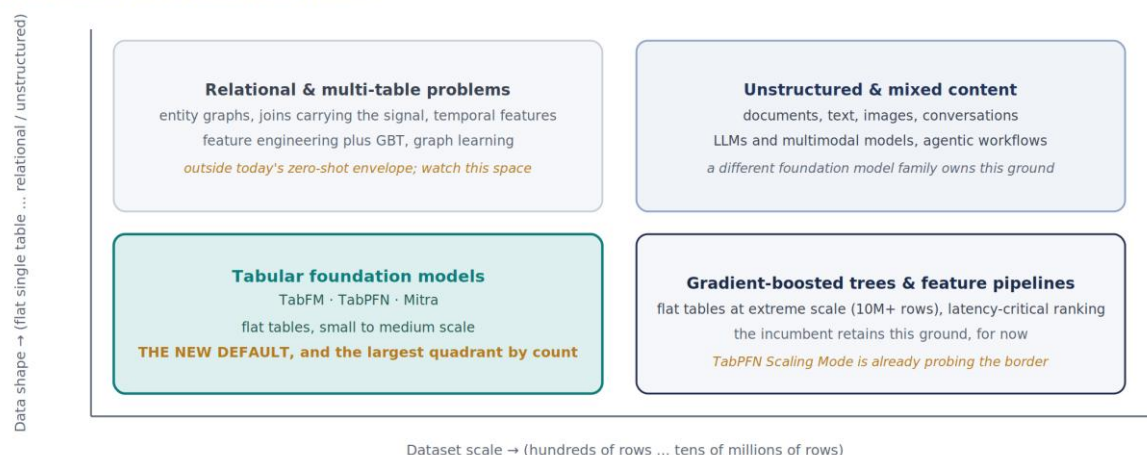


Figure 9. A regime map for CXOs. Zero-shot tabular models do not replace everything; they claim the quadrant that contains the most use cases by count: flat tables at small to medium scale. The strategic question is not whether to adopt, but how fast the borders of that quadrant move.

BUSINESS TAKEAWAY Zero-shot removes training, not accountability. Validation evidence, explainability, and version pinning become the new control points, and they must exist before the first production score is consumed.

The skeptic's corner: four objections, answered

- **"We already tried AutoML and it disappointed."** AutoML automated the search for a model, but it still built and maintained one model per dataset. Zero-shot removes the build entirely, which is a different economic object. The honest answer is empirical: run the bake-off on your own tables and let the numbers decide.
- **"Our data is too messy for this."** True, and no model fixes bad data. But the effort this approach saves on training is exactly the capacity you can redeploy into prediction views and data quality, which is where messy data actually gets fixed.
- **"Is this just Google lock-in?"** The weights are open and runnable anywhere; only the managed convenience is BigQuery-specific. Keeping TabPFN or Mitra alive in the champion-challenger harness preserves your exit options by design.
- **"Our flagship models already beat everything."** Keep them. The case here is the long tail of Figure 1, the hundreds of predictions you never built, not the three you have polished for years.

ACTION

An adoption playbook

1. **Map the portfolio.** Inventory the current model estate and identify candidates: small to medium supervised classification and regression models with high maintenance cost and moderate criticality.
2. **Benchmark on your own tables.** Run TabFM, TabPFN, and Mitra against incumbent models on five to ten representative internal datasets, measuring accuracy, calibration, latency, and total effort. Zero-shot evaluation makes this a matter of days, on CPU or a single T4-class GPU.
3. **Change the intake process.** Route new low-risk predictive requests from analysts through the zero-shot path first, reserving custom model builds for cases where the foundation model demonstrably underperforms.
4. **Update governance.** Extend model risk, validation, and AI governance frameworks to cover pretrained tabular models, including version pinning, explainability evidence, and monitoring for input drift, since the model itself no longer retrains.
5. **Align to the data platform.** If the organization runs on BigQuery, plan for AI.PREDICT early, because the marginal cost of experimentation will fall to near zero once prediction is a SQL function. If it runs on SAP or AWS estates, track the Prior Labs and Mitra roadmaps with equal attention.

CXO SUMMARY

Business takeaways

Seven things to remember from this briefing

1. **Tabular prediction is becoming a query, not a project.** TabFM predicts on unseen tables in a single forward pass; BigQuery AI.PREDICT will put that behind a SQL command.
2. **The value is the removed cost structure, not marginal accuracy.** Feature engineering, tuning, training, and retraining pipelines disappear from the per-use-case bill.
3. **This is a three-way platform contest, not a product launch.** SAP has committed over EUR 1B to Prior Labs, Amazon ships Mitra in AutoGluon, Google wires TabFM into the warehouse. The capability is infrastructure.
4. **Start in the long tail.** The fastest returns are the many small, high-maintenance models, and the small-data problems classical ML always served poorly.
5. **The infrastructure ask is modest.** Single mid-range GPUs on the self-hosted path, none on the managed path. Scrutinize AI.PREDICT pricing, not capex.
6. **Governance shifts, it does not shrink.** Version pinning, per-use-case validation evidence, explainability, and drift monitoring replace retraining audits as the control points.
7. **Redeploy, do not just reduce.** Context design, evaluation, and governance are the new high-value roles; the second dividend comes from moving specialists up the stack.

VERDICT: PILOT NOW

Run the champion-challenger benchmark on five to ten of your own tables this quarter. It costs days, not months. Adopt where the foundation model wins, keep the incumbent where it does not, and have governance in place before the first production score ships.

Five questions to ask your data leadership this month

- How many supervised tabular models do we maintain today, and what does each one cost per year to keep alive?
- Which ten use cases would we benchmark first, and do leakage-checked prediction views exist for them?
- If AI.PREDICT reached general availability tomorrow, which teams would be allowed to use it, under what controls?
- Do our model risk and validation frameworks cover a pretrained model that never retrains, including version pinning and drift monitoring?
- What is our redeployment plan for the data science hours that zero-shot prediction frees up?

CONCLUSION

Prediction becomes a question the business asks its data directly

For two decades, the answer to every tabular prediction problem was the same: assemble data, engineer features, train a model, tune it, deploy it, and repeat for the next dataset. TabFM is the clearest signal yet that this loop is being replaced by a different primitive: a single pretrained model that reads a table the way a language model reads a prompt, and answers immediately.

The competitive stakes are visible in the money: SAP committing over a billion euros to Prior Labs, Amazon embedding Mitra into AutoGluon, and Google wiring TabFM into BigQuery. When three ecosystems converge on the same capability within eighteen months, the capability is no longer a research curiosity. It is infrastructure.

The organizations that benefit first will not be the ones with the largest data science teams. They will be the ones that redesign their intake, governance, and platform choices, along the lines of the reference architecture above, so that a prediction becomes what it should always have been: a question the business asks its data directly, and an answer that arrives in seconds.

One model. Any table.
Answers in seconds.

THE NEW DEFAULT FOR ENTERPRISE PREDICTION · CXO INTELLIGENCE SERIES

STAY WITH THE SERIES

Read all articles: akhawat.com/writing

Subscribe to the CXO Intelligence Series newsletter on LinkedIn:

linkedin.com/newsletters/7475953227239419904

SOURCES

References

1. Google Research. Introducing TabFM: A zero-shot foundation model for tabular data. Google Research Blog, 30 June 2026. research.google/blog/introducing-tabfm-a-zero-shot-foundation-model-for-tabular-data

2. Google Research. TabFM repository, v1.0.0: code, TabArena reproduction, per-fold metrics, and head-to-head win rates. GitHub. github.com/google-research/tabfm
3. Google. TabFM v1.0.0 model weights, PyTorch and JAX backends. Hugging Face. huggingface.co/google/tabfm-1.0.0-pytorch
4. Hollmann, N., Muller, S., Purucker, L., et al. Accurate predictions on small data with a tabular foundation model. Nature, January 2025. doi:10.1038/s41586-024-08328-6
5. Prior Labs. TabPFN repository and documentation, including TabPFN-3 capacity limits and hardware guidance. GitHub. github.com/PriorLabs/TabPFN
6. SAP. Announcement of the acquisition of Prior Labs and investment of over EUR 1 billion across four years. SAP News, May 2026. news.sap.com
7. Amazon Science. Mitra: Mixed synthetic priors for enhancing tabular foundation models. amazon.science; model weights under Apache 2.0 within AutoGluon, auto.gluon.ai
8. TabArena living benchmark and Elo leaderboard, linked from the TabFM GitHub repository.
9. MarkTechPost. Google AI introduces TabFM: A hybrid-attention tabular foundation model for zero-shot classification and regression. 1 July 2026. marktechpost.com
10. Pandey, Y. R. I tried to break Google's new tabular foundation model. Then I fixed it. Independent evaluation, July 2026. yashrajpandey.com/writing/breaking-google-tabfm
11. Google Research. TimesFM: a decoder-only foundation model for time-series forecasting. github.com/google-research/timesfm
12. Prior Labs. Case study: TabPFN with Taktile for financial risk decisioning, including the account-opening fraud benchmark. priorlabs.ai/case-studies/taktile
13. Prior Labs. TabPFN-3 technical report and production application areas, including the Creditplus Bank AG endorsement. priorlabs.ai/technical-reports/tabpfn-3
14. Mission Lane Tech Blog. TabICL under the microscope: benchmarking tabular foundation models for enterprise credit risk. Medium, February 2026.

GPU sizing tiers in Table 3 are indicative estimates pending official TabFM hardware guidance; TabArena figures in Figure 5 are an illustrative representation of published rank order. All sources accessed July 2026.

This article is part of the CXO Intelligence Series. All articles: akhawat.com/writing · Newsletter: [CXO Intelligence Series on LinkedIn](#)

How to Cite

Akhawat, P. (2026). *Prediction Without Training: Google TabFM and the Rise of Tabular Foundation Models*. CXO Intelligence Series