

[ICML 2026 Spotlight]

TabSwift: An Efficient Tabular Foundation Model with Row-Wise Attention

Presenter Kyungseon Lee

2026-06

Seoul National University

Background — Tabular Foundation Model

Tabular prediction.

Train dataset $\mathcal{D} = \{(\mathbf{x}_i, y_i)\}_{i=1}^N$, test dataset $\mathbf{x}_q \in \mathbb{R}^d$, y_i is a class label or a real value.

Tabular foundation model, TFM.

A Transformer-based model f_θ is pretrained once on synthetic tasks and then reused for every dataset. Prediction is done by **in-context learning**. Let $\mathbf{x}_q$ be the *test data (query)* and $\hat{y}_q$ the *estimated label*.

$$\hat{y}_q = f_\theta(\mathcal{D}, \mathbf{x}_q).$$

No gradient update at test time. The most well-known model is **TabPFN**.

Problem.

Recent TFMs use heavier and heavier architectures to gain accuracy.
Goal: improve the old lightweight TFM and reach high performance.

Background — TabPFN v1: Row-wise Attention

One row = one token.

N train rows and T test rows, so $n = N + T$ tokens in total.

- ❶ **Unify the input dimension.** Each dataset has a different number of features F , so we map all of them to a fixed size $F_{\max}$, using zero-padding or PCA. $\tilde{\mathbf{x}} \in \mathbb{R}^n \times \mathbb{R}^{F_{\max}}$.
- ❷ **Embedding.** Train and test rows are each embedded into the model dimension d_{model} . $\tilde{\mathbf{x}} \in \mathbb{R}^n \times \mathbb{R}^{d_{\text{model}}}$.
- ❸ **Inference (regression case)** $\hat{y} = f_{\text{TabPFNv1}}(\tilde{\mathbf{x}}) \in \mathbb{R}^n$

Computational cost. $\mathcal{O}(n^2 d_{\text{model}})$ per attention layer.

Background — TabPFN v2: Row & Column Attention

The difference from TabPFN version 1 is that **one cell = one token**. That is, the input is treated as an $n \times d$ grid of tokens.

- ① **Turn each scalar cell into a vector token.** With learnable weights $\mathbf{w}_{\text{val}}, \mathbf{b}_{\text{val}}, \mathbf{c}_j \in \mathbb{R}^{d_{\text{model}}}$,

$$\mathbf{h}_{ij}^{(0)} = x_{ij} \mathbf{w}_{\text{val}} + \mathbf{b}_{\text{val}} + \mathbf{c}_j \in \mathbb{R}^{d_{\text{model}}}, \quad x_{ij} \in \mathbb{R}$$

- ② **Attention in two directions.** Arranging $\mathbf{h}_{ij}^{(0)}$ into an $n \times d$ grid gives $\mathbf{H} \in \mathbb{R}^{n \times d \times d_{\text{model}}}$. Let $\mathbf{H}_{:,j} \in \mathbb{R}^{n \times d_{\text{model}}}$ be the column tokens of feature j , and $\mathbf{H}_{i,:} \in \mathbb{R}^{d \times d_{\text{model}}}$ the row tokens of instance i . Let Attn be self-attention over a token sequence, and assume a single attention layer.

$$\text{row attention: } \mathbf{H}_{:,j} \leftarrow \text{Attn}(\mathbf{H}_{:,j}) \quad \forall j, \quad \text{column: } \mathbf{H}_{i,:} \leftarrow \text{Attn}(\mathbf{H}_{i,:}) \quad \forall i$$

The cost is $\mathcal{O}(dn^2d_{\text{model}}) + \mathcal{O}(nd^2d_{\text{model}})$, more than d times that of v1.

TabSwift — Key Question

Can the lightweight row-wise backbone of v1 reach v2-level performance if we add modern techniques?

The authors add three things to the original TabPFN v1.

- Add missing model capacity - Learnable register tokens.
 - A technique proven in ViT. K learnable global tokens are appended to the input and used as memory slots.
- Stabilize the attention update - Gated attention.
 - A technique proven in LLMs. Multiply the attention output by an element-wise learnable gate.
- Make it more efficient - Adaptive early exit.
 - Easy test data finish their prediction at a shallow layer and stop computing. The early-exit decision is made token by token, not for the whole test set at once.

TabSwift — Register Tokens

Input. After preprocessing steps 1 and 2 from v1, we start from

$$\tilde{\mathbf{x}} \in \mathbb{R}^{n \times d_{\text{model}}}.$$

Register tokens. K learnable tokens $\mathbf{R}^{(0)} \in \mathbb{R}^{K \times d_{\text{model}}}$ are placed in front of the row tokens:

$$\mathbf{X} = \left[\mathbf{R}^{(0)}; \mathbf{H}_{\text{rows}}^{(0)} \right] \in \mathbb{R}^{S \times d_{\text{model}}}, \quad S = K + n.$$

- Learned independently of the data and shared across all tasks. They pass through every layer and get updated together.
- Role: a global memory that gathers information about the whole dataset through attention.
- Since $K \ll n$, the extra cost is negligible.

TabSwift — Element-wise Gated Attention

Layer input $\mathbf{X} \in \mathbb{R}^{S \times d_{\text{model}}}$.

Assume the number of heads in multi-head attention is fixed to one.

Attention output. $\mathbf{W}_Q, \mathbf{W}_K, \mathbf{W}_V \in \mathbb{R}^{d_{\text{model}}}$:

$$\mathbf{O} = \text{softmax}\left(\frac{\mathbf{Q}(\mathbf{K})^\top}{\sqrt{d_{\text{model}}}}\right) \mathbf{V} \in \mathbb{R}^{S \times d_{\text{model}}}, \quad \mathbf{Q} = \mathbf{X}\mathbf{W}_Q$$

Gating. Build a gate from the input and multiply it into the output element-wise. $\mathbf{W}_g \in \mathbb{R}^{d_{\text{model}} \times d_{\text{model}}}$, $\mathbf{b}_g \in \mathbb{R}^{d_{\text{model}}}$, σ is the sigmoid:

$$\mathbf{G} = \sigma(\mathbf{X}\mathbf{W}_g + \mathbf{b}_g) \in (0, 1)^{S \times d_{\text{model}}}, \quad \tilde{\mathbf{O}} = \mathbf{G} \odot \mathbf{O}$$

The extra cost of the gate is linear, $\mathcal{O}(S d_{\text{model}})$. Negligible compared to the $\mathcal{O}(S^2 d_{\text{model}})$ of attention.

TabSwift — Early Exit (1): Architecture and Training

At every layer, compute the prediction we would get if we stopped now, and how much we can trust it.

Notation.

Let $\mathbf{H}^{(e)} \in \mathbb{R}^{S \times d_{\text{model}}}$ be the hidden state of layer e . Let $\mathbf{z}^{(e)} \in \mathbb{R}^{T \times d_{\text{model}}}$ be the rows of $\mathbf{H}^{(e)}$ at the test-data positions, and $\mathbf{r}^{(e)} \in \mathbb{R}^{K \times d_{\text{model}}}$ the rows at the register-token positions. There are T test rows and K register-token rows.

Prediction head $h_{\text{pred}}^{(e)}$.

The intermediate prediction using only the layers up to e

$$\hat{y}^{(e)} = h_{\text{pred}}^{(e)}(\mathbf{z}^{(e)})$$

Exit head $h_{\text{exit}}^{(e)}$.

Concatenate with the register mean $\mathbf{r}^{(e)} = \frac{1}{K} \sum_k \mathbf{R}_k^{(e)} \in \mathbb{R}^{d_{\text{model}}}$ and output a logit

$$s^{(e)} = h_{\text{exit}}^{(e)}\left(\left[\mathbf{z}^{(e)}; \mathbf{r}^{(e)}\right]\right) \in \mathbb{R}$$

TabSwift — Early Exit (2): Stopping Rule and Choosing τ

Train. The exit head learns a binary label: “is the intermediate prediction at layer e already correct?”

$$t^{(e)} = \mathbb{I}\{\hat{y}^{(e)} = y\} \in \{0, 1\}$$

The loss is binary cross-entropy:

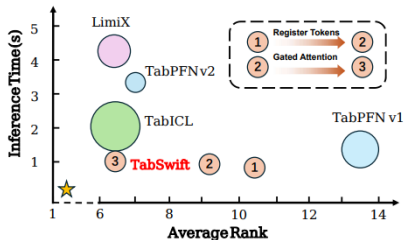
$$\mathcal{L}_{\text{exit}}^{(e)} = \text{BCE}(\sigma(s^{(e)}), t^{(e)})$$

Inference. Compute layers in order starting from 1, check $\sigma(s^{(e)})$ at every layer, and stop at the first layer where it passes the threshold τ :

$$e^* = \min \left\{ e : \sigma(s^{(e)}) \geq \tau \right\}, \quad \text{output is } \hat{y}^{(e^*)}.$$

If no layer passes, use the prediction from the full depth L .

Results



Backbone	# tokens	attention cost
Row-wise only, v1	n	$\mathcal{O}(n^2 d_{\text{model}})$
TabSwift	$S = K + n$	$\mathcal{O}(S^2 d_{\text{model}}) \approx \mathcal{O}(n^2 d_{\text{model}})$
Row & column, v2	$n \times d$	$\mathcal{O}(dn^2 d_{\text{model}} + nd^2 d_{\text{model}})$