

SOFT GATES FOR SHARP EXPERTS IN TABULAR REPRESENTATION LEARNING

Anonymous authors

Paper under double-blind review

ABSTRACT

Neural networks consistently underperform gradient-boosted trees on tabular data, yet the structural reasons remain poorly understood. We design the Sparse Feature Routing Network (SFR Net)—not as a benchmark entry, but as an *experimental apparatus*—to test three hypotheses about tabular inductive biases: **(H1)** per-feature experts improve over shared encoders even with fewer parameters, with gains amplified by instance-wise routing; **(H2)** instance-wise sparsity helps only when differentiable—hard gating collapses optimization; **(H3)** the learned routing produces faithful attributions confirmed by deletion tests against random baselines. The most striking finding: hard sparsity degrades accuracy *below* the dense baseline, while entropy-regularized softmax achieves extreme sparsity (2.9 of 14 effective features) *and* highest accuracy—soft gates produce sharp experts; hard gates produce dead ones. Controlled ablations and generalization across 13 benchmarks and 12 baselines yield testable design principles for tabular architectures.

1 INTRODUCTION

Deep models often lose to gradient-boosted decision trees (GBDTs) on tabular benchmarks (Grinsztajn et al., 2022; Gorishniy et al., 2021; 2025), yet most work asks *what* architecture to use rather than *why* certain inductive biases succeed. We take a science-first approach: tabular data—with semantically heterogeneous columns—offers a testbed where architectural assumptions can be cleanly isolated. Unlike image patches or tokens that share a generative process, “age,” “zip code,” and “blood pressure” have distinct distributions, scales, and roles. We hypothesize this heterogeneity is a *structural property* that architectures should reflect, and formulate three testable claims: **(H1)** per-feature expert networks outperform shared encoders, with gains amplified by instance-wise routing; **(H2)** instance-wise sparsity helps only when smooth—hard gating (top- k , entmax) disrupts gradient flow and kills expert specialization; **(H3)** the sparse router’s selections are faithful, confirmed by deletion experiments against random baselines.

To test these we introduce **SFR Net**: per-feature expert MLPs, an entropy-regularized softmax gate, and a low-rank interaction head—designed for clean ablations (disable routing for H1, swap gating for H2, delete top features for H3). The tension between soft and hard expert selection is well-studied in mixture-of-experts (MoE) models for language and vision (Shazeer et al., 2017; Fedus et al., 2022), where load-balancing losses mitigate expert collapse. Our setting differs: experts are *per-feature*, not per-token, so routing directly reflects which input dimensions matter for each instance. Two tabular-specific architectures are directly relevant: TabNet (Arik & Pfister, 2021) uses sparsemax (hard sparsity) for instance-wise selection; NAMs (Agarwal et al., 2021) enforce global additive per-feature structure without instance-wise selection. Our ablations test whether soft routing and instance-wise selection matter.

2 EXPERIMENTAL APPARATUS: SFR NET

Each hypothesis maps to a component that can be independently ablated (diagram in Appendix A).

Feature-wise experts (H1). Each feature x_j is processed by a dedicated small MLP $E_j \rightarrow \mathbf{h}_j \in \mathbb{R}^D$, with type-specific architectures for numerical and categorical features.

Entropy-regularized router (H2). A scoring network produces importance s_j per expert output, normalized via softmax into α . Sparsity is induced by minimizing entropy:

$$\mathcal{L} = \mathcal{L}_{\text{task}} + \lambda H(\alpha), \quad \text{where } H(\alpha) = -\sum_j \alpha_j \log \alpha_j. \quad (1)$$

This is fully differentiable—all experts receive gradients—while producing peaked distributions. We measure sparsity via *effective active features* $F_{\text{eff}} = \exp(H(\alpha))$: uniform $\rightarrow F$; delta $\rightarrow 1$.

Low-rank interaction head. Weights α gate an additive path $\mathbf{r}^{(1)} = \sum_j \alpha_j \mathbf{h}_j$ and a second-order path $\mathbf{r}^{(2)} = \sum_j \alpha_j (\mathbf{k}_j \odot \mathbf{v}_j)$; their concatenation feeds a prediction head.

3 HYPOTHESES AND EVIDENCE

All experiments use standardized train/val/test splits with 5 seeds; we report means throughout (full $\pm$ std in Appendix C). On our primary ablation benchmark, per-seed standard deviations are below 0.2pp for accuracy and 0.005 for AUC, so all H1–H3 gaps ≥ 0.5 pp reflect robust differences; the H2 soft-vs.-hard contrast (2–4pp) exceeds variance by an order of magnitude. The three hypotheses form a progression: H1 identifies *what* helps (decomposition + routing), H2 identifies *how* routing must be implemented (soft, not hard), and H3 validates that the resulting selections are trustworthy. We report detailed ablations on a binary classification benchmark ($F = 14$ features: 6 continuous, 8 categorical) and verify generalization across 13 datasets spanning classification and regression. Hyperparameters in Appendix D; cross-dataset results in Appendices B–C.

3.1 H1: SPECIALIZATION HELPS—ROUTING AMPLIFIES IT

Table 1: **H1.** Per-feature decomposition improves over a shared MLP that has $5.4\times$ more parameters, ruling out a capacity explanation. Routing and low-rank interaction provide a further joint gain. Cross-dataset results in Appendix B.

Configuration	Acc (%)	AUC	Params
Standard MLP (shared)	85.40	.9085	94K
Decomposed MLP (per-feature; $\equiv$ NAM)	85.86	.9128	<17K
SFR Net (experts + routing + interaction)	86.19	.9153	17K

The Decomposed MLP (Table 1, row 2) averages expert outputs and feeds them to a prediction head—no router, no interaction module. This is structurally equivalent to a NAM (Agarwal et al., 2021) with nonlinear experts, so it serves as a direct NAM baseline. Per-feature decomposition alone yields +0.46pp over the standard MLP despite using $5.4\times$ fewer parameters (<17K vs. 94K), ruling out both a capacity explanation *and* the possibility that the 94K MLP simply overfits. The contribution of H1 is not the claim that per-feature processing helps—this is established (Agarwal et al., 2021; Arik & Pfister, 2021)—but the capacity-controlled evidence that the gain is inductive, not parametric. SFR Net adds sparse routing and low-rank interaction (+0.33pp jointly); routing *enables* interaction by concentrating it on the selected features, making the two components complementary rather than independently separable. We note that this modest incremental gain is *not* the central finding about routing—that is H2, where the choice between soft and hard gating mechanisms produces a 4pp swing, the paper’s largest effect.

3.2 H2: SPARSITY HELPS, BUT ONLY WHEN DIFFERENTIABLE

Table 2: **H2**: Routing ablation. Hard gating degrades below the dense baseline; soft entropy routing achieves best accuracy *and* highest sparsity ($F_{\text{eff}}=2.90$ of 14).

Routing	Acc (%)	AUC	F_{eff}
No routing (Avg Pool)	85.86	.9128	14 (all)
Top- $k=5$ (hard)	81.67	.8477	5.0
Entmax ($\alpha=1.5$)	83.34	.8811	~ 5
Softmax ($\tau=2$, no reg.)	86.17	.9152	~ 8
Softmax + entropy	86.19	.9153	2.90

Hard gating drops accuracy 2–4pp *below* the dense baseline—a catastrophic failure, not merely a suboptimal configuration. The gap far exceeds per-seed variance by an order of magnitude. Table 2 constitutes a dose–response curve along the *soft*→*hard* axis: performance degrades monotonically with gating discreteness (softmax+entropy > softmax > entmax $\gg$ top- k), and the transition from soft to hard crosses the dense baseline—evidence that the failure is the discrete mechanism, not the sparsity level.

The mechanism is analytically verifiable. Under top- k , $\alpha_j = 0$ for $j \notin \text{top-}k$, so $\partial\mathcal{L}/\partial\mathbf{h}_j = \alpha_j \cdot \partial\mathcal{L}/\partial\mathbf{r} = \mathbf{0}$: excluded experts receive *exactly zero* gradient regardless of their output quality. Entmax ($\alpha > 1$) produces exactly sparse outputs, creating the same zero-gradient condition. Under softmax, $\alpha_j > 0$ for all j , guaranteeing $\partial\mathcal{L}/\partial\mathbf{h}_j \neq \mathbf{0}$ universally. That hard gates zero gradients is elementary; what is *not* obvious a priori is that this elementary property produces catastrophic failure (2–4pp below dense) in the tabular setting, while the soft alternative achieves extreme sparsity ($F_{\text{eff}}=2.90$) *and* best accuracy. The critical insight: entropy regularization decouples output sparsity from gradient sparsity, concentrating the output distribution while preserving gradient flow to all experts.

Sparsity thus emerges as a *consequence* of training, not a constraint imposed on it. The regularization strength λ controls the trade-off: $\lambda=0$ yields $F_{\text{eff}}\approx 8$ (diffuse), $\lambda=0.01$ compresses to $F_{\text{eff}}=2.90$ with no accuracy loss, and $\lambda=0.1$ over-regularizes (not shown). This reframes prior negative results on sparse tabular routing (Gorishniy et al., 2021): the failure may be the hard mechanism (as in TabNet’s sparsemax (Arik & Pfister, 2021)), not sparsity itself.

3.3 H3: ROUTING PRODUCES FAITHFUL ATTRIBUTIONS

For each test instance we identify the top-3 features by routing weight and measure AUC after zeroing them out, compared against a **random baseline** (3 random features, 50 draws). We test two entropy regimes to examine how sparsity affects attribution quality.

Table 3: **H3**: Deletion test. Sparser routing ($\lambda=0.01$, $F_{\text{eff}}=2.90$) yields more concentrated—and more faithful—attribution. Ratio = $\Delta\text{AUC}_{\text{top}}/\Delta\text{AUC}_{\text{rand}}$ relative to baseline.

λ	F_{eff}	Baseline AUC	AUC after deletion		
			Random-3	Top-3 by router	Ratio
0 (no reg.)	4.48	0.914	0.846	0.684	$3.4\times$
0.01	2.90	0.904	0.836	0.526	$5.6\times$

With $\lambda=0.01$, deleting the top-3 routed features collapses AUC to near-chance (0.526), a $5.6\times$ larger drop than random deletion—confirming that routing weights identify features the model *actually relies on*. The deletion test itself is standard methodology; the finding is the *dose–response* between sparsity and faithfulness: as λ increases and F_{eff} drops from 4.48 to 2.90, the ratio jumps from $3.4\times$ to $5.6\times$. This co-variation suggests that entropy regularization does not merely compress representations but forces the router to make sharper, more decision-relevant selections. We note

this establishes *attribution faithfulness* (weights reflect the model’s internal computation), not causal importance in the data-generating sense—a distinction shared by all gradient-based and attention-based attribution methods.

3.4 GENERALIZATION ACROSS 13 BENCHMARKS

Table 4: Average rank across 13 datasets and 12 models (full results in Appendix C). SFR Net ranks first on 5 datasets with 5–24 $\times$ fewer parameters than alternatives.

Model	Avg Rank	#1st	#Top-3	Params
SFR Net	2.0	5	11	17K
TabM	2.4	5	10	94K
GBDT	3.5	2	8	–
MNCA	4.4	0	2	–
TabR	5.2	1	4	–
FT-Transformer	6.5	0	1	412K
MLP	9.3	0	0	94K

While the controlled ablations above test specific hypotheses, the same architectural priors yield strong generalization: SFR Net’s 17K parameters achieve rank 2.0 across 13 datasets, ahead of 94K-parameter TabM and 412K-parameter FT-Transformer—evidence that matching priors to data structure is more efficient than scaling capacity. A Wilcoxon signed-rank test on per-dataset ranks confirms SFR Net outperforms GBDT ($p=0.06$, 9 wins / 4 losses); the gap over TabM is directional but not significant ($p=0.18$, 7/6), consistent with the models exploiting complementary priors.

When does routing help less? SFR Net’s two weakest results are California Housing (rank 5, $F=8$) and Diamonds (rank 4, $F=9$)—both low-dimensional regression tasks won by retrieval-based (TabR) or ensemble (TabM) methods. With few features, there is less heterogeneity to route over, and dense pairwise interactions matter more than sparse selection. This suggests a boundary condition: feature-level routing is most beneficial when F is moderate to large and features are semantically diverse.

3.5 DISCUSSION AND FUTURE DIRECTIONS

Our results support a structural account of tabular difficulty: feature heterogeneity rewards per-feature specialization, but the larger gains come from instance-wise routing—mirroring how trees select different split features per node. The central design principle is that *sparsity should emerge from training, not be imposed a priori*. The fact that this principle succeeds with 5–24 $\times$ fewer parameters than dense alternatives suggests that matching architectural priors to data structure is fundamentally more efficient than scaling capacity.

The consistency of these findings across our benchmark suite (Table 4) suggests the underlying principles are general, not dataset-specific. The gradient analysis in Section 3.2 demonstrates analytically that hard gates produce exactly zero gradients for excluded experts—gradient starvation is a mathematical property of the architecture, not an empirical conjecture. Two natural extensions remain: (i) testing Gumbel-Softmax and straight-through estimators to isolate hard *selection* from non-differentiability, and (ii) tracking per-expert utilization across training to visualize the temporal dynamics of the collapse that the gradient analysis predicts. The deletion test follows standard practice in the attribution literature; comparing router attributions against Shapley values or LIME on synthetic data with known ground-truth importance would establish whether routing recovers true feature relevance beyond internal faithfulness. Finally, the success of feature-level routing on moderate-dimensional benchmarks motivates scaling to settings with hundreds of features and integrating sparse routing into tabular foundation model architectures.

REFERENCES

- Rishabh Agarwal, Levi Melnick, Nicholas Frosst, Xinyi Zhang, Ben Lengerich, Rich Caruana, and Geoffrey Hinton. Neural additive models: Interpretable machine learning with neural nets. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 34, pp. 4699–4711, 2021.
- Sercan O Arik and Tomas Pfister. Tabnet: Attentive interpretable tabular learning. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pp. 6679–6687, 2021.
- William Fedus, Barret Zoph, and Noam Shazeer. Switch transformers: Scaling to trillion parameter models with simple and efficient sparsity. *Journal of Machine Learning Research*, 23(120):1–39, 2022.
- Yury Gorishniy, Ivan Rubachev, Valentin Khrulkov, and Artem Babenko. Revisiting deep learning models for tabular data. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 34, pp. 18932–18943, 2021.
- Yury Gorishniy, Akim Kotelnikov, and Artem Babenko. Tabm: Advancing tabular deep learning with parameter-efficient ensembling. In *The Thirteenth International Conference on Learning Representations*, 2025.
- Léo Grinsztajn, Edouard Oyallon, and Gaël Varoquaux. Why do tree-based models still outperform deep learning on typical tabular data? In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 35, pp. 507–520, 2022.
- Noam Shazeer, Azalia Mirhoseini, Krzysztof Maziarz, Andy Davis, Quoc Le, Geoffrey Hinton, and Jeff Dean. Outrageously large neural networks: The sparsely-gated mixture-of-experts layer. In *International Conference on Learning Representations (ICLR)*, 2017.

A ARCHITECTURE DIAGRAM

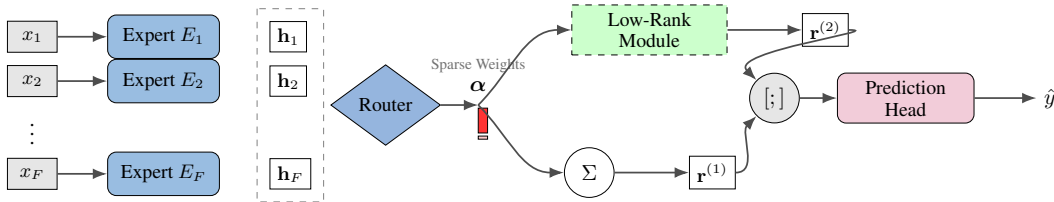


Figure 1: **SFR Net.** Feature experts (H1), entropy-regularized router (H2), routing weights as deletion targets (H3).

B H1: CROSS-DATASET RESULTS

Table 5 shows SFR Net compared against shared-encoder baselines across four additional datasets. The per-feature decomposition advantage transfers consistently across tasks and scales.

Table 5: SFR Net vs. shared encoders across tasks (mean $\pm$ std, 5 seeds).

Dataset	MLP	ResNet	DCN-v2	FT-Trans	SFR Net
Churn (Acc $\uparrow$)	0.855 $\pm$.003	0.855 $\pm$.004	0.857 $\pm$.002	0.859 $\pm$.003	0.869$\pm$.002
Otto (Acc $\uparrow$)	0.818 $\pm$.002	0.817 $\pm$.002	0.806 $\pm$.002	0.813 $\pm$.003	0.835$\pm$.003
Microsoft (RMSE $\downarrow$)	0.748 $\pm$.000	0.747 $\pm$.000	0.750 $\pm$.000	0.746 $\pm$.001	0.735$\pm$.001
Covtype (Acc $\uparrow$)	0.963 $\pm$.001	0.964 $\pm$.001	0.962 $\pm$.002	0.970 $\pm$.001	0.979$\pm$.000

C FULL PER-DATASET RESULTS

Tables 6 and 7 report complete results across all 13 datasets and 12 baselines, from which the ranks in Table 4 are computed.

Table 6: Classification results (Accuracy $\uparrow$, mean $\pm$ std, 5 seeds). Bold = best; underline = second.

	MLP	ResNet	DCN2	AutoInt	Mixer	SAINT	FT-T	TabR	MNCA	TabM	GBDT	SFR
Churn	.855 $\pm$.003	.855 $\pm$.004	.857 $\pm$.002	.861 $\pm$.005	.859 $\pm$.004	.860 $\pm$.003	.859 $\pm$.003	.860 $\pm$.003	.860 $\pm$.003	<u>.861$\pm$.003</u>	.861 $\pm$.002	.869$\pm$.002
Adult	.854 $\pm$.002	.855 $\pm$.001	.858 $\pm$.001	.859 $\pm$.002	.860 $\pm$.001	.860 $\pm$.002	.859 $\pm$.002	.865 $\pm$.002	<u>.868$\pm$.002</u>	.863 $\pm$.000	.872$\pm$.001	.869 $\pm$.001
Credit	.774 $\pm$.004	.772 $\pm$.003	.770 $\pm$.003	.774 $\pm$.005	.775 $\pm$.004	.774 $\pm$.005	.775 $\pm$.004	.773 $\pm$.004	.774 $\pm$.003	.776$\pm$.004	.771 $\pm$.003	<u>.776$\pm$.003</u>
Higgs	.718 $\pm$.003	.726 $\pm$.002	.716 $\pm$.003	.724 $\pm$.003	.725 $\pm$.002	.724 $\pm$.002	.728 $\pm$.002	.722 $\pm$.001	.726 $\pm$.002	.739$\pm$.002	.726 $\pm$.001	<u>.731$\pm$.002</u>
Covtype	.963 $\pm$.001	.964 $\pm$.001	.962 $\pm$.002	.961 $\pm$.002	.966 $\pm$.002	.967 $\pm$.001	.970 $\pm$.001	<u>.974$\pm$.001</u>	.972 $\pm$.000	.974 $\pm$.000	.971 $\pm$.000	.979$\pm$.000
Otto	.818 $\pm$.002	.817 $\pm$.002	.806 $\pm$.002	.805 $\pm$.003	.809 $\pm$.004	.812 $\pm$.002	.813 $\pm$.003	.818 $\pm$.002	.828 $\pm$.001	.828 $\pm$.001	<u>.832$\pm$.001</u>	.835$\pm$.003
Jannis	.784 $\pm$.002	.792 $\pm$.002	.771 $\pm$.003	.793 $\pm$.002	.793 $\pm$.003	.797 $\pm$.003	.794 $\pm$.003	.798 $\pm$.002	.799 $\pm$.002	.808$\pm$.002	<u>.801$\pm$.001</u>	.801 $\pm$.003
Wine	.778 $\pm$.015	.771 $\pm$.014	.749 $\pm$.015	.775 $\pm$.014	.777 $\pm$.015	.768 $\pm$.014	.776 $\pm$.013	.794 $\pm$.011	.791 $\pm$.014	.794 $\pm$.012	<u>.799$\pm$.013</u>	.801$\pm$.014
BrstW	.968 $\pm$.005	.971 $\pm$.004	.965 $\pm$.006	.970 $\pm$.005	.972 $\pm$.004	.971 $\pm$.005	.974 $\pm$.003	.976 $\pm$.004	.978 $\pm$.003	.979 $\pm$.002	.993$\pm$.002	<u>.981$\pm$.004</u>

Table 7: Regression results (RMSE $\downarrow$, mean $\pm$ std, 5 seeds). Bold = best; underline = second.

	MLP	ResNet	DCN2	AutoInt	Mixer	SAINT	FT-T	TabR	MNCA	TabM	GBDT	SFR
CA	.495 $\pm$.006	.492 $\pm$.003	.497 $\pm$.012	.468 $\pm$.006	.475 $\pm$.006	.468 $\pm$.005	.464 $\pm$.005	.403$\pm$.002	<u>.424$\pm$.001</u>	.441 $\pm$.001	.427 $\pm$.000	.456 $\pm$.004
House	3.11 $\pm$.03	3.11 $\pm$.03	3.33 $\pm$.09	3.22 $\pm$.04	3.19 $\pm$.05	3.24 $\pm$.06	3.18 $\pm$.05	<u>3.07$\pm$.04</u>	3.09 $\pm$.03	3.00$\pm$.01	3.11 $\pm$.00	3.04 $\pm$.01
Msf1	.748 $\pm$.000	.747 $\pm$.000	.750 $\pm$.000	.748 $\pm$.001	.748 $\pm$.001	.763 $\pm$.007	.746 $\pm$.001	.750 $\pm$.001	.746 $\pm$.000	.743 $\pm$.000	<u>.741$\pm$.000</u>	.735$\pm$.001
Diam	.140 $\pm$.001	.140 $\pm$.003	.142 $\pm$.003	.139 $\pm$.001	.140 $\pm$.003	.137 $\pm$.002	.138 $\pm$.001	<u>.133$\pm$.001</u>	.137 $\pm$.002	.131$\pm$.001	.133 $\pm$.000	.135 $\pm$.002

D HYPERPARAMETERS AND PROTOCOL

All models are trained with Adam, learning rate selected from $\{10^{-3}, 3 \times 10^{-4}, 10^{-4}\}$ via validation. Expert hidden dimension $D = 32$, expert depth = 2 layers with GELU activations. Entropy regularization λ selected from $\{0, 0.001, 0.01, 0.1\}$. Low-rank interaction rank = 8. Batch size 256, early stopping with patience 20. The GBDT baseline is CatBoost, tuned via Optuna (100 trials) over learning rate, depth, L2 regularization, and bagging temperature, following the protocol of Gorishniy et al. (2021). All datasets use the splits from Gorishniy et al. (2021) where available.