

IS UNIFORM EXPRESSIVITY TOO RESTRICTIVE? TOWARDS EFFICIENT EXPRESSIVITY OF GNNs

Sammy Khalife*
khalife.sammy@cornell.edu

Josué Tonelli-Cueto†
josue.tonelli.cueto@bizkaia.eu

ABSTRACT

Uniform expressivity guarantees that a Graph Neural Network (GNN) can express a query without the parameters depending on the size of the input graphs. This property is desirable in applications in order to have number of trainable parameters that is independent of the size of the input graphs. Uniform expressivity of the two variable guarded fragment (GC2) of first order logic is a well-celebrated result for Rectified Linear Unit (ReLU) GNNs Barceló et al. (2020). In this article, we prove that uniform expressivity of GC2 queries is not possible for GNNs with a wide class of Pfaffian activation functions (including the sigmoid and tanh), answering a question formulated by Grohe (2021). We also show that despite these limitations, many of those GNNs can still efficiently express GC2 queries in a way that the number of parameters remains logarithmic on the maximal degree of the input graphs. Furthermore, we demonstrate that a log-log dependency on the degree is achievable for a certain choice of activation function. This shows that uniform expressivity can be successfully relaxed by covering large graphs appearing in practical applications. Our experiments illustrates that our theoretical estimates hold in practice.

1 INTRODUCTION

Graph Neural Networks (GNNs) form a powerful computational framework for machine learning on graphs, and have proven to be very performant methods for various applications ranging from analysis of social networks, structure and functionality of molecules in chemistry and biological applications Duvinaud et al. (2015); Zitnik et al. (2018); Stokes et al. (2020); Khalife et al. (2021), computer vision Defferrard et al. (2016), simulations of physical systems Battaglia et al. (2016); Sanchez-Gonzalez et al. (2020), and techniques to enhance optimization algorithms Khalil et al. (2017); Cappart et al. (2021) to name a few. Significant progress has been made in recent years to understand their computational capabilities, in particular regarding functions one can compute or express via GNNs. This question is of fundamental importance as it precedes any learning related consideration, by asking a description the class of functions one *can hope* to learn with a GNN. A better understanding of the dependency on the architecture of the GNN (activation and aggregation functions, number of layers, inner dimensions, ...) can also help to select the most appropriate one, given some basic knowledge on the structure of the problem at hand. One common approach in this line of research is to compare standard algorithms on graphs to GNNs. For example, if a given algorithm is able to distinguish two graph structures, is there a GNN able to distinguish them (and conversely)? The canonical algorithm that serve as a basis of comparison is the *color refinement* algorithm (closely related to *Weisfeler-Lehman* algorithm), a heuristic that almost solve graph isomorphism Cai et al. (1992). It is for example well-known that the color refinement algorithm refine the standard GNNs, and that the activation function has an impact on the the ability of a GNN to refine color refinement Morris et al. (2019); Grohe (2021); Khalife & Basu (2023).

*School of Operations Research and Information Engineering, Cornell Tech, Cornell University, NY, USA

†Department of Applied Mathematics and Statistics, Johns Hopkins University, Baltimore, USA

The expressivity of GNNs can also be compared to queries of an appropriate logic. A first form of comparison relies on *uniform expressivity*, where the size of the GNN is not allowed to grow with the input graphs. Given this assumption, a first basic question is to determine the fragment of logic that can be expressed via GNNs, in order to describe queries that one can possibly learn via GNNs. For instance with first-order-logic, one can express the query that a vertex of a graph is part of a triangle. It turns out standard aggregation-combine GNNs cannot express that query, and more generally, a significant portion of the first order logic is removed. The seminal result of Barceló et al. (2020) states that aggregation-combine GNNs with sum aggregations (Σ -AC-GNNs) allow to express at best queries of a fragment of the first order logic (called graded model logic, or GC2), and can be achieved with Rectified Linear Unit (ReLU) activations. Conversely, ReLU Σ -AC-GNNs can express uniformly queries of GC2. In Grohe (2023), the author presents more general results about expressivity in the uniform setting, introducing a new logical guarded fragment (GFO+C) of the first order with counting (FO+C), in order to describe the expressivity of the general message passing GNNs, when no assumption is made on the activation or aggregation function. The author shows that $\text{GC2} \subseteq \text{GNN} \subseteq \text{GFO+C}$, where both inclusions are strict. See Figure 1 for a representation of the known result for uniform expressivity.

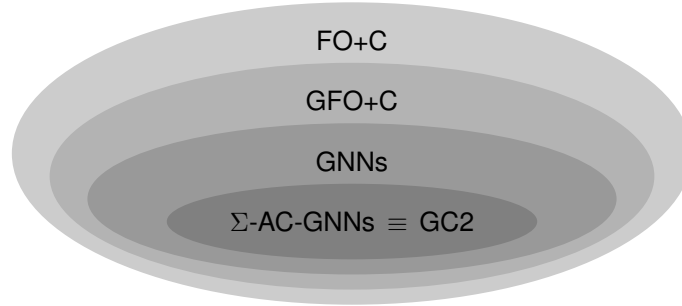


Figure 1: Main logical classes to describe the uniform expressivity of GNNs. If no restriction is made on the activation and aggregation functions, standard message passing GNNs can express slightly more than GC2 queries, but less than GFO+C queries. The sum-aggregation aggregation combine GNNs (Σ -AC-GNNs) exactly express GC2 queries.

Evidently, the logic of GNNs depends on the choice of architecture. Recent work include the impact of the activation function aggregation function on uniform expressivity Khalife (2023); Rosenbluth et al. (2023). In Grohe & Rosenbluth (2024), the authors also study the expressivity depending on the type of messages that is propagated (message that depends only on the source vs. dependency on source *and* target vertex). They show that in a non-uniform setting, the two types of GNNs have the same expressivity, but the second variant is more expressive uniformly. Additionally, Khalife (2023) also showed that GNNs with rational activations and aggregations do not allow to express such queries uniformly, even on trees of depth two. More generally, a complete characterization of the activation function that enables GC2 expressivity remains elusive.

By allowing the size of the GNN to grow with the input graphs, it is possible to relax the uniform notion of expressivity and ask analogous questions to the uniform case. For instance, what is the portion of first-order logic one can express over graphs of bounded order with a GNN? What is the required size of a GNN with a given activation to do so? In this setup, it becomes now possible that general GNNs express a *broader* class than the uniform one (GFO+C), and that Σ -AC-GNNs express a broader class than GC2. Such question has been explored in the work of Grohe (2023), that describes the new corresponding fragments of FO+C. In particular, the author proves an equivalence between a superclass of GFO+C, called GFO+C_{nu} and general GNNs with *rational piecewise linear* (rpl)-approximable functions. More precisely, a query can be expressed by a family of rpl-approximable GNNs whose number of parameters grow polynomially with respect to the input graphs, if and only if the query belongs to GFO+C_{nu} , a more expressive fragment than GFO+C. The GNNs used to prove that a GFC+C_{nu} this equivalence only requires linearized sigmoid activations ($x \mapsto \min(1, \max(0, x))$) and Σ -aggregation. Similarly to the uniform case, the set of activation functions that allows maximal expressivity still needs to be better understood. If one restricts to Σ -AC-GNNs and the well-known GC2 fragment, one can ask for example, given a family of activation functions, how well Σ -AC-GNNs with those activation functions can express GC2

queries. For example, given an activation function, and integer n , what is the number of parameters of a Σ -AC-GNNs needed to express GC2 queries of depth¹ d over graphs of order at most n ?

Main contributions. Our first contribution goes towards a more complete understanding of both uniform and non-uniform expressivity of GNNs, with a focus on Σ -AC-GNNs. We first show that similarly to rational GNNs, there exist a large family of GC2 that cannot be expressed with Σ -AC-GNNs with bounded Pfaffian activations including the sigmoidal activation $x \mapsto \frac{1}{1+\exp(-x)}$, function that was so far conjectured to be as expressive as the ReLU $x \mapsto \max(0, x)$ units in this setup Grohe (2021). More precisely, we prove that GNNs within that class cannot express all GC2 queries uniformly, in contrast with the same ones replaced with ReLU activations. Our second contribution adds on descriptive complexity results Grohe (2023) by providing new complexity upper-bounds in the non-uniform case for a class of functions containing the Pfaffian ones. Our constructive approach takes advantage of the composition power providing some simple assumptions on the activation functions, endowing deep neural networks to approximate the step function $x \mapsto \mathbf{1}_{\{x>0\}}$ efficiently. Our experiments confirm our theoretical statements on synthetic datasets.

The rest of this article is organized as follows. Section 2 presents the basic definitions of GNNs, Pfaffian functions and the background logic. In Section 3, we state our main theoretical results, in comparison with the existing ones. Section 4 presents an overview of the proof of our negative result that builds on these properties. Section 5 presents the core ideas to obtain our non-uniform expressivity upper-bounds. Technical details including additional definitions and lemmas are left in the appendix. In Section 6, we present the numerical experiments supporting our theoretical results. We conclude with some remarks, present the limitations of our work in Section 7.

2 PRELIMINARIES

Notations. In the following, for any positive integer n , $[n]$ refers to the set $\{1, \dots, n\}$. We assume the input graphs of GNNs to be finite, undirected, simple, and vertex-labeled: a graph is a tuple $G = (V(G), E(G), P_1(G), \dots, P_\ell(G))$ consisting of a finite vertex set $V(G)$, a binary edge relation $E(G) \subset V(G)^2$ that is symmetric and irreflexive, and unary relations $P_1(G), \dots, P_\ell(G) \subset V(G)$ representing $\ell > 0$ vertex labels. In the following, we suppose that the $P_i(G)$ ’s form a partition of the set of vertices of G , i.e. each vertex has a unique label. Also, the number ℓ of labels, which we will also call *colors*, is supposed to be fixed and does not grow with the size of the input graphs. In order to describe the logic of GNNs, we also take into account access to the color of the vertices into the definition of the logic considered. When there is no ambiguity about which graph G is being considered, $N(v)$ refers to the set of neighbors of v in G not including v . $|G|$ will denote the number of vertices of G . We use simple curly brackets for a set $X = \{x \in X\}$ and double curly brackets for a multiset $Y = \{\{y \in Y\}\}$.

Definition 2.1 (Neural network). Fix an *activation function* $\sigma : \mathbb{R} \rightarrow \mathbb{R}$. For any number of hidden layers $k \in \mathbb{N}$, input and output dimensions $w_0, w_{k+1} \in \mathbb{N}$, a $\mathbb{R}^{w_0} \rightarrow \mathbb{R}^{w_{k+1}}$ *neural network with σ activation* is given by specifying a sequence of k natural numbers $w_1, w_2, \dots, w_k$ representing widths of the hidden layers and a set of $k+1$ affine transformations $T_i : \mathbb{R}^{w_{i-1}} \rightarrow \mathbb{R}^{w_i}$, $i = 1, \dots, k+1$. Such a NN is called a $(k+1)$ -layer NN, and is said to have k hidden layers. The function $f : \mathbb{R}^{w_0} \rightarrow \mathbb{R}^{w_{k+1}}$ computed or represented by this NN is:

$$f = T_{k+1} \circ \sigma \circ T_k \circ \dots \circ T_2 \circ \sigma \circ T_1.$$

The *size* of the neural network is defined as $w_1 + \dots + w_L$.

Definition 2.2 (Graph Neural Network (GNN)). A GNN² is a recursive embedding of vertices of a labeled graph by relying on the underlying adjacency information and node features using a neural network (cf. Definition 2.1). Each vertex v is attributed an indicator vector $\xi^0(v)$ of size ℓ , encoding the color of the node v : the colors being indexed by the palette $\{1, \dots, \ell\}$, $\xi^0(v) = e_i$ (the i -th canonical vector) if the color of the vertex v is i . A GNN is characterized by:

¹The notion of depth will be defined more precisely in Section 2. For the time being, one can think of the depth as a measure of complexity measure on the space of queries.

²This definition coincides with aggregation-combine GNNs (AC-GNNs) in the literature. In the following, for simplicity, unless state otherwise, we refer to AC-GNNs simply as GNNs.

◦ A positive integer T called the number of iterations, positive integers $(d_t)_{t \in \{1, \dots, T\}}$ and $(d'_t)_{t \in \{0, \dots, T\}}$ for inner dimensions. $d_0 = d'_0 = \ell$ is the input dimension of the GNN (number of colors) and d_T is the output dimension.

◦ A sequence of combination and aggregation functions $(\text{comb}_t, \text{agg}_t)_{t \in \{1, \dots, T\}}$. Each aggregation function agg_t maps each finite multiset of vectors of $\mathbb{R}^{d_{t-1}}$ to a vector in $\mathbb{R}^{d'_t}$. For any $t \in \{1, \dots, T\}$, each combination function $\text{comb}_t : \mathbb{R}^{d_{t-1} + d'_t} \rightarrow \mathbb{R}^{d_t}$ is a neural network with given activation function $\sigma : \mathbb{R} \rightarrow \mathbb{R}$.

The update rule of the GNN at iteration $t \in \{0, \dots, T-1\}$ for any labeled graph G and vertex $v \in V(G)$, is given by:

$$\xi^{t+1}(v) = \text{comb}(\xi^t(v), \text{agg}\{\{\xi^t(w) : w \in \mathcal{N}_G(v)\}\})$$

The *size* of a GNN is the maximum size of the combine neural network functions.

In the context of this paper, we will focus on the following activation functions:

- **ReLU:** The *Rectified Linear Unit* (ReLU) activation function $\text{ReLU} : \mathbb{R} \rightarrow \mathbb{R}_{\geq 0}$ is defined as $\text{ReLU}(x) = \max\{0, x\}$.
- **CReLU:** The *Clipped Rectified Linear Unit* (CReLU) activation function $\text{CReLU} : \mathbb{R} \rightarrow [0, 1]$ is defined as $\text{CReLU}(x) = \min\{\max\{0, x\}, 1\}$.
- **tanh:** The *Hyperbolic Tangent* activation function $\tanh : \mathbb{R} \rightarrow (-1, 1)$ is defined as $\tanh(x) = \frac{e^x - e^{-x}}{e^x + e^{-x}}$.
- **Sigmoid:** The *Sigmoid* activation function $\text{Sigmoid} : \mathbb{R} \rightarrow (0, 1)$ is defined as $\text{Sigmoid}(x) = \frac{1}{1 + e^{-x}}$.

Definition 2.3 (Guarded model logic (GC) and GC2 (Barceló et al., 2020; Grohe, 2021)). The fragment of guarded logic GC is formed using quantifiers that restrict to range over the neighbours of the current nodes. GC-formulas are formed from the atomic formulas by the Boolean connectives and quantification restricted to formulas of the form $\exists^{\geq p} y (E(x, y) \wedge \psi)$, where x and y are distinct variables and x appears in ψ as a free variable. Note that every formula of GC has at least one free variable. We refer to the 2-variable fragment of GC as GC2, also known as graded modal logic (with two variables).

A more constructive definition can be given as follows: a graded modal logic formula is either $\text{Col}(x)$ (returning 1 or 0 for one of the palette colors) or one of the following:

$$\neg \phi(x), \quad \phi(x) \wedge \psi(x), \text{ or } \quad \exists^{\geq N} y (E(x, y) \wedge \phi(y))$$

where ϕ and ψ are GC2 formulas and N is a positive integer:

Example 2.1 (Barceló et al. (2020)). All graded modal logic formulas are unary by definition, so all of them define unary queries. Suppose $\ell = 2$ (number of colors), and for illustration purposes $\text{Col}_1 = \text{Red}$, $\text{Col}_2 = \text{Blue}$. Let:

$$\gamma(x) := \text{Blue}(x) \wedge \exists y (E(x, y) \wedge \exists^{\geq 2} x (E(y, x) \wedge \text{Red}(x)))$$

γ queries if x has blue color and as at least one neighbor which has two blue neighbors. Then γ is in GC2. Now,

$$\delta(x) := \text{Blue}(x) \wedge \exists y (\neg E(x, y) \wedge \exists^{\geq 2} x E(y, x) \wedge \text{Red}(x))$$

is not in GC2 because the use of the guard $\neg E(x, y)$ is not allowed. However,

$$\eta(x) := \neg(\exists y (E(x, y) \wedge \exists^{\geq 2} x E(y, x) \wedge \text{Blue}(x)))$$

is in GC2 because the negation $\neg$ is applied to a formula in GC2.

Definition 2.4. (Khovanskii, 1991) Let $U \subseteq \mathbb{R}^n$ be an open set. A Pfaffian chain of order $r \geq 0$ and degree $\alpha \geq 1$ in U is a sequence of analytic real functions $f_1, \dots, f_r$ in U satisfying differential equations:

$$\frac{\partial f_i}{\partial x_j} = P_{i,j}(x, f_1(x), \dots, f_i(x))$$

where $P_{i,j} \in \mathbb{R}[x_1, \dots, x_n, y_1, \dots, y_i]$ are polynomials of $(n + i)$ variables of degree at most α .

Proposition 2.1. $\tanh$ and Sigmoid are bounded Pfaffian functions.

3 STATEMENT OF THE CONTRIBUTIONS

Our contributions focus on the expressivity of GC2 queries by GNNs.

Definition 3.1. Suppose that ξ is the vertex embedding computed by a GNN. We say that such GNN expresses a unary query Q over a set of graphs $\mathcal{X}$ if there is a real $\epsilon < \frac{1}{2}$ such that for all graphs $G \in \mathcal{X}$ and vertices $v \in V(G)$.

$$\begin{cases} \xi(G, v) \geq 1 - \epsilon & \text{if } v \in Q(G) \\ \xi(G, v) \leq \epsilon & \text{if } v \notin Q(G) \end{cases}.$$

We say that the GNN expresses *uniformly* Q if $\mathcal{X}$ is the set of all graphs.

The following result represent the state of the art regarding the expressivity of GC2 queries by GNN with activation function CReLU.

Theorem 3.1. (Barceló et al., 2020; Grohe, 2021) *For every GC2 query Q of size d , there is a GNN with activation function CReLU, d iterations and size at most $2d$ that express Q over all graphs. Conversely, any (aggregation-combine) GNN expresses with activation function CReLU expressed a GC2 query.*

Remark 3.1. ℓ being the number of colors of the vertices in the input graphs, the family of GNNs with $\text{agg} = \text{sum}$, and $\text{comb}(x, y) = \text{ReLU}(Ax + By + C)$ (where $A \in \mathbb{N}^{\ell \times \ell}$, $B \in \mathbb{N}^{\ell \times \ell}$ and $C \in \mathbb{N}^{\ell}$) is sufficient to uniformly express all queries of GC2 uniformly. Furthermore, for each query Q of depth q , there is a GNN of this type with at most q iterations that expresses uniformly Q .

Remark 3.2. Theorem 3.1 admits the a partial converse for beyond AC-GNNs Barceló et al. (2020): if a unary query is expressible by a GNN and also expressible in first-order logic, then it is expressible in GC2.

Until now, the question about uniform expressivity has been mainly open for other activation functions such as tanh and Sigmoid. Our contributions is two-fold. On the one hand, we show that uniform expressivity is not possible. On the other hand, we show that, despite this impossibility, a form of almost-uniform expressivity is possible.

We briefly overview these two contributions below, giving more details in Sections 4 and 5.

3.1 IMPOSSIBILITY OF UNIFORM EXPRESSIVITY

Definition 3.2 (Superpolynomial). Let $f : \mathbb{R} \rightarrow \mathbb{R}$ be a function having a finite limit $\ell \in \mathbb{R}$ in $+\infty$. We say that f *converges superpolynomially* to ℓ iff for every polynomial $P \in \mathbb{R}[X]$, $\lim_{x \rightarrow +\infty} P(x)(f(x) - \ell) \rightarrow 0$.

We say that a Pfaffian function is *superpolynomial* if it non constant and verifies this condition.

Our first result extends the negative result of rational GNNs by Khalife (2023) to the class of bounded superpolynomial Pfaffian GNNs.

Theorem 3.2. *Let $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ be a bounded superpolynomial Pfaffian activation function. Let ξ be the output of a GNN with activation function σ . Then there is a GC2 query Q , such that for any $\epsilon > 0$ there exist a pair of rooted trees (T, T') of depth 2 with roots (s, s') such that*

$$i) \quad Q(T, s) = 1 \quad \text{and} \quad Q(T', s') = 0$$

$$ii) \quad |\xi(s', T') - \xi(s, T)| < \epsilon$$

Corollary 3.3. *Let $\sigma \in \{\text{Sigmoid}, \text{tanh}\}$. Then there is a GC2 query Q such that for any GNN with activation σ , the GNN cannot express Q uniformly over all graphs.*

We want to emphasize that the above query does not depend on the choice of the GNN provided it has this type of activation function.

3.2 ALMOST-UNIFORM EXPRESSIVITY

Our second contribution shows that for a large class of activation functions, which we call step-like activation function (see Definition 5.1 in Section 5 below) we have almost-uniform expressivity of GC2 queries.

Theorem 3.4. *Let σ be a step-like activation function. Then, for every GC2 query Q of size d and Δ , there is a GNN with activation function σ , d iterations and size $O(d + \log \Delta)$ that express Q over all graphs with degree $\leq \Delta$. Moreover, if the parameter η of σ is zero, the size of the GNN can be further reduced to $O(d + \log \log \Delta)$.*

Corollary 3.5. *Let $\sigma \in \{\text{Sigmoid}, \tanh\}$. Then, for every GC2 query Q of size d and Δ , there is a GNN with activation function σ , d iterations and size $O(d + \log \log \Delta)$ that express Q over all graphs with degree $\leq \Delta$.*

We want to emphasize that the function $\log \log \Delta$ grows so slowly that for almost all theoretical and practical purposes we can consider it to be a constant function. Hence the term “almost-uniform expressivity”.

4 SUPERPOLYNOMIAL BOUNDED PFAFFIAN GNNs: IMPOSSIBILITY

Let us overview why Theorem 3.2 holds, expanding all details in Appendix B. Our first result makes use of specific families of input trees that allow us to *saturate* the output signal of a GNN by increasing appropriately the number of vertices in certain regions of these graphs. The input trees, pictured in Figure 2 are carefully chosen in order to return distinct outputs to query of GC2. The gist of our proof consists in showing that by increasing the order of the trees, the output signals of the GNNs of the considered class can be made arbitrarily close. This is achieved by proving a stronger property on all vertices on the trees stated below. To do so, we first need the intermediate definition:

$$S_r := \{\phi : \mathbb{R}^{1+(r-1)} \rightarrow \mathbb{R} : \forall P \in \mathbb{R}[X], \forall x \in \mathbb{R} \lim_{y \rightarrow +\infty} |P(y_{r-1})\phi(x, y_1, \dots, y_{r-1})| = 0\}$$

where for vectors $z \in \mathbb{R}^d$, and functions $F : \mathbb{R}^d \rightarrow \mathbb{R}$, $\lim_{z \rightarrow +\infty} F(z) = 0$ means that for every $\epsilon > 0$ there exists $\exists \beta \geq 0$, such that $\min(z_1, \dots, z_d) \geq \beta$ implies $|F(z)| < \epsilon$.

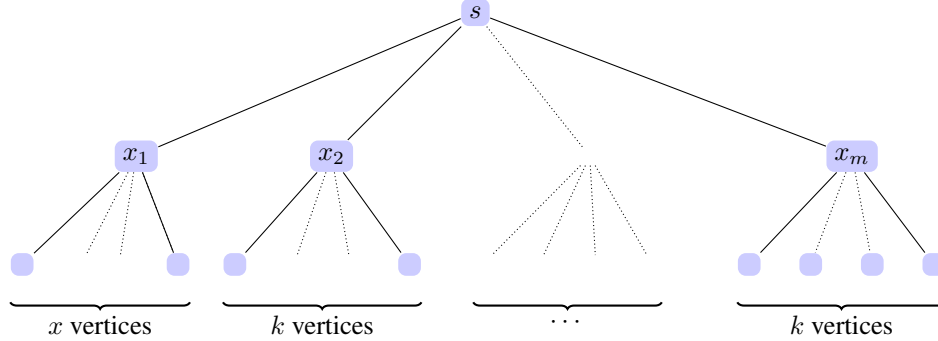


Figure 2: $T[x, m, k]$

For non-negative integers k, m , let $T[x, k, m]$ be the rooted tree of Figure 2. For ease of notation, if the context allows it, we will refer to $T[x, k, m]$ simply as T . For any non-negative integer t , $v \in V(T)$, let $\xi^t(v, T)$ be the embedding returned for node v by a GNN with activation function $\sigma \in H$ after t iterations. Let $M > 0$ be an integer. We will prove by induction on t that for any $t \in \{1, \dots, M\}$:

$$\begin{aligned} \xi^t(s, T) &= v_t + \eta_t(x, k, m) \\ \xi^t(x_{i \geq 2}, T) &= g_t(k) + \epsilon_t(x, k, m) \\ \xi^t(l_{i \geq 2}, T) &= h_t(k) + \nu_t(x, k, m) \\ \xi^t(x_1, T) &= g_t(k) + \epsilon_t^1(x, k, m) \\ \xi^t(l_1, T) &= h_t(k) + \nu_t^1(x, k, m) \end{aligned}$$

with the following properties:

- i) v_t is constant with respect to x, k and m .

- ii) g_t and h_t are bounded Pfaffian functions that depend only on the combination and aggregation functions and the iteration t .
- iii) For every $t \in \{1, \dots, M\}$, $\nu_t, \eta_t, \epsilon_t, \epsilon_t^1, \nu_t^1 \in S_3$.

A few comments are in order. Recall that S_3 is the set of functions $\mathbb{R}^3 \rightarrow \mathbb{R}$ such that, for any fixed first argument, the function of the second and third argument tend (collectively) *superpolynomially* to 0 with respect to the third argument. This asymptotic behavior will arise from our assumptions made on the Pfaffian activation functions, and will be used in our inductive argument. One immediate corollary is that we can decompose the signal at the root vertex into a main component v_t and an error term that tends to 0 when the order of the trees to $+\infty$. The main signal v_t will depend *only* on the number of iterations, not on x, m or k . In the other vertices, the main component will only be a function of k . This will prove for instance, that the two outputs of the GNN on the pair $T([0, k, m], T[1, k, m])_{(k,m) \in \mathbb{N}^2}$ converge to the same limit, when k and m tend to $+\infty$. However, for every $k > 0$ and $m > 0$, $T[0, k, m]$ and $T[1, k, m]$ have distinct output at the root for the following query:

$$Q(v) := \neg(\exists^{\geq 1} y (E(y, v) \wedge (\neg \exists^{\geq 2} v E(v, y)))) = \forall y (E(y, v) \wedge \exists^{\geq 2} z E(z, y))$$

expressing that vertex v 's neighbors have degree at least 2. Clearly, Q is a GC2 query of depth 4. We can easily generalize to any GC2 query for which the root vertex of both trees have different output. In conclusion, the uniform expressivity of GNNs with superpolynomial bounded activation functions is weaker than those of the ReLUs.

5 TOWARDS EFFICIENT NON-UNIFORM EXPRESSIVITY

Our second result focuses on *non-uniform* expressivity, where the size of the GNN is allowed to grow with the size of the input graphs. The result that we obtain holds for the wide class of activation functions defined below, and to whichever activation function can express them. This class (See Lemma C.1 in Appendix C) is characterized by a fast converge to the linear threshold activation function

$$\sigma_*(x) := \begin{cases} 0 & \text{if } x < 1/2 \\ 1/2 & \text{if } x = 1/2 \\ 1 & \text{if } x > 1/2 \end{cases}. \quad (1)$$

Definition 5.1 (Step-Like Activation Function). A *step-like activation function* is a C^2 -map $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ satisfying for $\eta \in [0, 1)$, $\varepsilon \in (0, 1/2)$, $N \in \mathbb{N}$ and $H > 0$ the following:

- (a) $\sigma(0) = 0$ and $\sigma(1) = 1$.
- (b) $|\sigma'(0)|, |\sigma'(1)| \leq \eta$.
- (c) for all $x \notin (\varepsilon, 1 - \varepsilon)$, $|\sigma^N(x) - \sigma_*(x)| \leq \min\{\varepsilon, (1 - \eta)/H\}$.
- (d) for all $x \in \sigma^N(\mathbb{R} \setminus (\varepsilon, 1 - \varepsilon))$, $|\sigma''(x)| \leq H$.

Our main contribution presented Theorem 3.4 can be deduced from its technical version stated in the theorem below. All proof details are relegated to Appendix C.

Theorem 5.1. *Let σ be a step-like activation function. Then, for every GC2 query Q of size d and Δ , there is a GNN with activation function σ , d iterations and size*

$$2d + N + \frac{2 + \log(\Delta + 2)}{1 - \log(1 - \eta)}$$

parameters per iteration that express Q over all graphs with degree $\leq \Delta$. Moreover, if $\eta = 0$, the size of the GNN can be further reduced to

$$2d + N + \log(2 + \log(\Delta + 2)).$$

Remark 5.1. We note that the fact that the size of the GNN depends on the degree Δ of the class of graphs we are working with only logarithmically guarantees that the GNN has a reasonable size for applications. Moreover, if $\eta = 0$, then we have a log-log dependency, which for whichever practical setting is an almost constant function.

Now, to give meaning to the above result, we will prove the following proposition giving us a way of constructing step-like activation functions out of tanh and Sigmoid.

Proposition 5.2. *There is a step-like activation function $\bar{\sigma}_{\arctan}$ with $\eta = 0.64$, $\varepsilon = 0.1$, $N = 0$ and $H = 1.52$ that can be expressed with a 2-layer NN with activation function $\arctan$ of size 2.*

Proposition 5.3. *There is a step-like activation function $\bar{\sigma}_{\tanh}$ with $\eta = 0$, $\varepsilon = 0.2$, $N = 0$ and $H = 2.2$ that can be expressed with a 3-layer NN with activation function $\tanh$ of size 5.*

Proposition 5.4. *There is a step-like activation function $\bar{\sigma}_{\text{Sigmoid}}$ with $\eta = 0$, $\varepsilon = 0.2$, $N = 0$ and $H = 2.2$ that can be expressed with a 5-layer NN with activation function Sigmoid of size 9.*

Using these propositions, we get the following two important corollaries of Theorem 5.1 to get the almost-uniform expressivity of GC2 queries for tanh and Sigmoid, for which uniform expressivity is not possible due to Theorem 3.2.

Corollary 5.5. *For every GC2 query Q of size d and Δ , there is a GNN with activation function $\tanh$, d iterations and size at most $2d + 1.4 + 0.7 \log(\Delta + 2)$ that express Q over all graphs with degree $\leq \Delta$.*

Corollary 5.6. *For every GC2 query Q of size d and Δ , there is a GNN with activation function $\tanh$, d iterations and size at most $2d + 5 \log(2 + \log(\Delta + 2))$ that express Q over all graphs with degree $\leq \Delta$.*

Corollary 5.7. *For every GC2 query Q of size d and Δ , there is a GNN with activation function $\tanh$, d iterations and size at most $2d + 9 \log(2 + \log(\Delta + 2))$ that express Q over all graphs with degree $\leq \Delta$.*

6 NUMERICAL EXPERIMENTS

6.1 QUERIES

In this section we demonstrate how the size of the neural networks impacts the GNN's ability to express GC2 queries via two experiments. For both of them, we consider the following queries:

$$Q_1(v) := \neg (\exists^{\geq 1} y (E(y, v) \wedge (\neg \exists^{\geq 2} v E(v, y)))) = \forall y (E(y, v) \wedge \exists^{\geq 2} z E(z, y))$$

$Q_1(v)$ is expressing that all neighbors of v have degree at least two. Note that Q is in GC2 and has depth 4³. In this case the input trees are unicolored.

$$Q_2(v) = \text{Red}(v) \wedge (\exists^{\geq 1} x E(x, v) \wedge (\exists^{\geq 1} v E(v, x) \wedge \text{Blue}(v)) \wedge (\exists^{\geq 1} v E(v, x) \wedge \text{Red}(v)))$$

$Q_2(v)$ is expressing that v is red, and has a neighbor that has a red neighbor and a blue neighbor. Q_2 is in GC2 as well and has depth 7. The vertices of the trees are colored as follows: the source and depth-one vertices are red, and only the leaves are blue.

6.2 EXPERIMENTAL SETUP

Computing and approximating queries. For the first experiment, our set of input graphs is composed of pairs of trees $(T[0, k], T[1, k])_{k \in \mathbb{N}}$ presented in Section 4. We compare for $i \in [2]$, the two following GNNs:

- i) The ReLU GNN that exactly computes the query Q_i .

For Q_1 , the ReLU GNN has sum-aggregation, 4 iterations (depth of Q). The input and output dimensions are set for convenience to 4 to have a constant dimension for all iterations. In particular, all features vectors of the input vertices are set to e_1 , the first canonical vector of $\mathbb{R}^4$. The weights and biases are set accordingly to the construction in Section 5 (specified in Appendix C). Their exact values are explicated in the Appendix D.

$$\text{For every } t \in [4], \text{ and every } x, y \in \mathbb{R}^4 \quad \text{comb}_t(x, y) := \text{ReLU}(Ax + By + c)$$

For Q_2 , the ReLU GNN has 7 iterations.

³Here, since trees are unicolored, we removed the color atomic queries for the sake of presentation. Otherwise, the depth of the query would be 5.

- ii) The Pfaffian GNN that expresses Q_i non uniformly, with activation function: $\sigma : x \mapsto \frac{4}{\pi} \arctan(x)$, and sum-aggregation, with the same construction as in the proof of Theorem 5.1. Each combination function is given by a NN with m layers, and one neuron per layer. The bias of the first layer is set to $-\frac{1}{2}$. In the last non-hidden layer (no activation), the weight is set to $\frac{1}{2}$, and the bias is set to 0.5 to the last layer.

Setting, $f_m := \sigma^m(\cdot)$, the neural network computes the function $g_m := \frac{1}{2} + \frac{1}{2}f_m(\cdot - 1/2)$. Note that g_m is indeed a function computed by a neural network with activation function σ by adding a second neuron in each layer, setting the weights to 0, and biases to 0 except in the last layer where the bias is set to 1 and the weight set to $\frac{\pi}{4 \arctan(1)}$ to add a bias of $\frac{1}{2}$.

$$\text{For every } t \in [4], \text{ and every } x, y \in \mathbb{R}^4 \quad \text{comb}_t(x, y) := g_m(Ax + By + c)$$

where A , B , and c are defined as in i) above, i.e., given in Appendix D below.

We report for different values of iterations m and orders of graph inputs n , how close the two outputs of the Pfaffian GNNs are, on tree inputs $T[0, k]$ and $T[1, k]$ are. Since the root of these trees have respectively 0 and 1 as output of the query Q , the ReLU GNN constructed will have a constant gap of 1. Theorem 3.2 states that these outputs will tend to be arbitrarily close when the order of the graph increases, for a fixed m . In addition, Theorem 3.4 tells us that the number of iterations should impact exponentially well how large graphs can be handled to be at a given level of precision of a gap 1.

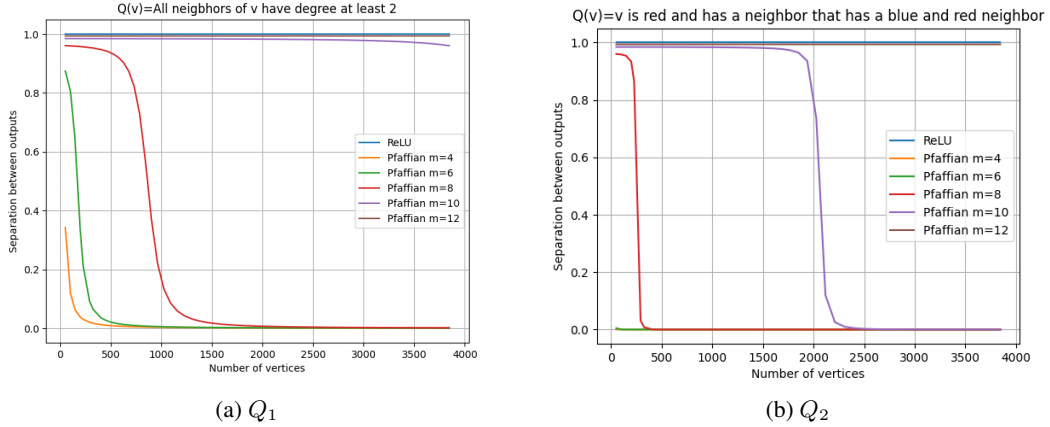


Figure 3: Separation power of CReLU GNN vs. step-like Pfaffian GNNs as a function of the size of the input.

Learning queries. For the second experiment, we compare the GNN’s ability to learn GC2 queries, depending on the activation considered (ReLU vs. Pfaffian). We consider the same two queries as above and train two distinct GNNs: (a) A first GNN with Sigmoid activation function, and (b) A second GNN with CReLU activation functions. Both GNNs have 4 and 7 iterations when trained to learn the first and second query respectively. Each iteration is attributed his own combine function, a feedforward neural network with one hidden layer. This choice is justified by our theoretical results that one can either compute exactly a GC2 query (with the ReLU one), or approximate it efficiently (with the Pfaffian one) with only one hidden layer for the combine function. Our training dataset is composed of 3750 graphs of varying size (between 50 and 2000 nodes) of average degree 5, and generated randomly. Similarly, our testing dataset is composed of 1250 graphs with varying size between 50 and 2000 vertices, of average degree 5.

The experiments were conducted on a Mac-OS laptop with an Apple M1 Pro chip. The neural networks were implemented using PyTorch 2.3.0 and Pytorch geometric 2.2.0. The details of the implementation are provided in the supplementary material.

6.3 EMPIRICAL RESULTS

Figure 3 reports our results for the first experiment. As our theory predicts, the CReLU GNN designed with the appropriate architecture and weights is able to distinguish the source vertices of both trees with constant precision 1, regardless of the number of vertices. On the other end, the Pfaffian GNNs struggle to do separation for values of $m \leq 8$ (recall that m is the number of composition, i.e. the number of layers in the combinations functions) for graphs of less than 1500 vertices. As the value of m increases slightly to 10 for the first query and 12 for the second, the capacity of these GNNs dramatically increases to reach a plateau close to 1 for graphs up to 4000 vertices. This also illustrates second our second theoretical result predicting that the step-like activation functions endow GNNs efficient approximation power of GC2 queries.

Figure 4 reports the mean square error on our test set for learning the same two queries. We observe that the error is similar for both GNNs, with a slight advantage for Pfaffian GNNs. This suggests that better uniform capacities does not imply that learning the queries is getting easier, given the same architecture and learning method. In particular, we conjecture that the set of weights such that the ReLU GNN computes the query is challenging to reach with standard optimization learning techniques based on variants of stochastic gradient descent.

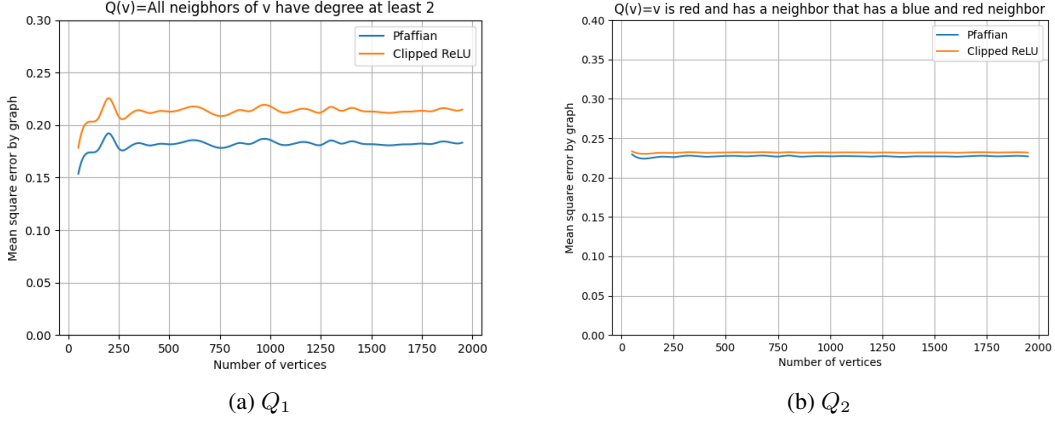


Figure 4: Learning queries with Pfaffian GNN vs ReLU GNN. The mean square error per graph on the test set is displayed as a function of the order of the graph.

7 DISCUSSION AND OPEN QUESTIONS

Uniform expressivity is a desirable property when some information on the function or query to learn is available. Structural information may allow to design the GNN appropriately in a hope to express the query by selecting the correct weights, for example after learning from samples. In this study, we showed that many activation functions are not as powerful as CReLU GNNs from a uniform standpoint. This limitation is counterbalanced by the efficiency of many activations (including the Pfaffian ones that are weaker than CReLU) for non-uniform expressivity, where one allows the number of parameters to be non-constant with respect to the input graphs. Therefore, we argue that despite being desirable, uniform expressivity has to be complemented by at different measures for a thorough study on the expressivity of GNNs.

Moreover, it is frequent that one does not know in advance the structure of the query to be learned. Even when knowing the appropriate architecture for which approximate expressivity is guaranteed, learning the appropriate coefficients can be challenging as illustrated in our numerical experiments. However, these experiments are only preliminary and constitute by no means an exhaustive exploration of all the architectures and hyperparameters to learn those queries. As such, we hope this study will constitute a first step towards provably efficient and expressive GNNs, and suggest new directions to bridge the gap between learning and expressivity.

ACKNOWLEDGMENTS AND DISCLOSURE OF FUNDING

Sammy Khalife gratefully acknowledges support from Air Force Office of Scientific Research (AFOSR) grant FA95502010341, and from the Office of Naval Research (ONR) grant N00014-24-1-2645. J.T.-C. thanks Jazz G. Suchen for useful discussions during the write-up of this paper, particularly regarding Proposition C.3.

REFERENCES

- Pablo Barceló, Egor V Kostylev, Mikael Monet, Jorge Pérez, Juan Reutter, and Juan-Pablo Silva. The logical expressiveness of graph neural networks. In *8th International Conference on Learning Representations (ICLR 2020)*, 2020.
- Peter Battaglia, Razvan Pascanu, Matthew Lai, Danilo Jimenez Rezende, et al. Interaction networks for learning about objects, relations and physics. *Advances in neural information processing systems*, 29, 2016.
- Jin-Yi Cai, Martin Fürer, and Neil Immerman. An optimal lower bound on the number of variables for graph identifications. *Combinatorica*, 12(4):389–410, 1992.
- Quentin Cappart, Didier Chételat, Elias Khalil, Andrea Lodi, Christopher Morris, and Petar Veličković. Combinatorial optimization and reasoning with graph neural networks. *arXiv preprint arXiv:2102.09544*, 2021.
- Michaël Defferrard, Xavier Bresson, and Pierre Vandergheynst. Convolutional neural networks on graphs with fast localized spectral filtering. *Advances in neural information processing systems*, 29, 2016.
- David K Duvenaud, Dougal Maclaurin, Jorge Iparraguirre, Rafael Bombarell, Timothy Hirzel, Alán Aspuru-Guzik, and Ryan P Adams. Convolutional networks on graphs for learning molecular fingerprints. *Advances in neural information processing systems*, 28, 2015.
- Martin Grohe. The logic of graph neural networks. In *2021 36th Annual ACM/IEEE Symposium on Logic in Computer Science (LICS)*, pp. 1–17. IEEE, 2021.
- Martin Grohe. The descriptive complexity of graph neural networks. *arXiv preprint arXiv:2303.04613*, 2023.
- Martin Grohe and Eran Rosenbluth. Are targeted messages more effective? *arXiv preprint arXiv:2403.06817*, 2024.
- Sammy Khalife. The logic of rational graph neural networks. *arXiv preprint arXiv:2310.13139*, 2023.
- Sammy Khalife and Amitabh Basu. On the power of graph neural networks and the role of the activation function. *arXiv preprint arXiv:2307.04661*, 2023.
- Sammy Khalife, Thérèse Malliavin, and Leo Liberti. Secondary structure assignment of proteins in the absence of sequence information. *Bioinformatics Advances*, 1(1):vbab038, 2021.
- Elias Khalil, Hanjun Dai, Yuyu Zhang, Bistra Dilkina, and Le Song. Learning combinatorial optimization algorithms over graphs. *Advances in neural information processing systems*, 30, 2017.
- Askold G Khovanskiĭ. *Fewnomials*, volume 88. American Mathematical Soc., 1991.
- Christopher Morris, Martin Ritzert, Matthias Fey, William L Hamilton, Jan Eric Lenssen, Gaurav Rattan, and Martin Grohe. Weisfeiler and leman go neural: Higher-order graph neural networks. In *Proceedings of the AAAI conference on artificial intelligence*, volume 33, pp. 4602–4609, 2019.
- Eran Rosenbluth, Jan Toenshoff, and Martin Grohe. Some might say all you need is sum. In *Proceedings of the Thirty-Second International Joint Conference on Artificial Intelligence*, pp. 4172–4179, 2023.

Alvaro Sanchez-Gonzalez, Jonathan Godwin, Tobias Pfaff, Rex Ying, Jure Leskovec, and Peter Battaglia. Learning to simulate complex physics with graph networks. In *International conference on machine learning*, pp. 8459–8468. PMLR, 2020.

Jonathan M Stokes, Kevin Yang, Kyle Swanson, Wengong Jin, Andres Cubillos-Ruiz, Nina M Donghia, Craig R MacNair, Shawn French, Lindsey A Carfrae, Zohar Bloom-Ackermann, et al. A deep learning approach to antibiotic discovery. *Cell*, 180(4):688–702, 2020.

A. J. Wilkie. A theorem of the complement and some new o-minimal structures. *Selecta Mathematica, New Series*, 5:397–421, 1999. doi: 10.1007/s000290050052. URL <https://doi.org/10.1007/s000290050052>.

Marinka Zitnik, Monica Agrawal, and Jure Leskovec. Modeling polypharmacy side effects with graph convolutional networks. *Bioinformatics*, 34(13):i457–i466, 2018.

A FACTS ABOUT PFAFFIAN FUNCTIONS

The following three results are folklore in the theory of Pfaffian functions. We finish with a proof of Proposition 2.1.

The first proposition is just a direct consequence of the so-called *o-minimality* of Pfaffian functions Wilkie (1999), which guarantees that Pfaffian systems share with polynomial systems their finiteness properties—among which we have the finiteness of zero-dimensional zero sets. The first lemma is an easy consequence of the chain rule and the definition of Pfaffian functions. The second lemma follows easily from the proposition, but we prove it since the proof is not straightforward.

Proposition A.1. *Let $f : \mathbb{R} \rightarrow \mathbb{R}$ be a Pfaffian function. Then f is either constant or has finitely many zeroes in $\mathbb{R}$.* $\square$

Lemma A.2. *The set of Pfaffian functions is stable by composition.* $\square$

Lemma A.3. *A bounded Pfaffian function $f : \mathbb{R} \rightarrow \mathbb{R}$ has a bounded derivative.*

Proof. Since f , f' and f'' are Pfaffian, by Proposition A.1, they have a finite number of zeros—if any of them are constant, then f cannot be bounded. Hence there is $R > 0$ such that all the zeros of f , f' and f'' are inside $(-R, R)$. In this way, since f' is bounded on any bounded interval, we only need to prove that f' is bounded on $[R, \infty)$ and on $(-\infty, -R]$. Without loss of generality, we will focus on the former case.

Assume, without loss of generality, after multiplying f and the variable X by -1 if necessary, that $f(R) > 0$ and $f'(R) > 0$. Now, we either have that f'' is positive or negative on the whole $[R, \infty)$, and so for all $x \in [R, \infty)$, $f'(x) \geq f'(R) > 0$, if f'' positive, or $0 < f'(x) \leq f'(R)$, if f'' negative. In the first case, for $x \geq R$,

$$f(x) = f(R) + \int_R^x f'(s) ds \geq f(R) + (x - R)f'(R)$$

goes to infinity as $x \rightarrow \infty$, which contradicts our assumption on f . Thus, we are in the second case, in which f' is bounded as desired. $\square$

Now, we end with a proof of Proposition 2.1.

Proof of Proposition 2.1. Since

$$\tanh(x) = 2 \text{Sigmoid}(2x) - 1, \quad (2)$$

we only need to show that Sigmoid is Pfaffian, by Lemma A.2. Now, $\text{Sigmoid}'(x) = -e^{-x} \text{Sigmoid}(x)^2$. Therefore

$$\text{Sigmoid}'(x) = P(e^{-x}, \text{Sigmoid}(x)), \quad e^{-x} = Q(e^{-x})$$

with $P(U, V) = -UV^2$ and $Q(Z) = -Z$, and so Sigmoid is Pfaffian. $\square$

B PROOF OF THEOREM 3.2 AND COROLLARY 3.3

Proof of Theorem 3.2. We suppose here that $\text{agg}_t = \text{SUM}$, but the proof can be adapted to any aggregation function that admits a limit in $\{-\infty, +\infty\}$ when the size of the multiset with a repeated entry grows to $+\infty$. We consider the set of activation functions

$$H := \{\sigma : \mathbb{R} \rightarrow \mathbb{R} : \sigma \text{ is Pfaffian, bounded and } \sigma \text{ converges fast in } +\infty\}$$

Let S_r be the set of $\mathbb{R}^r \rightarrow \mathbb{R}$ functions of fast decay in the last variable:

$$S_r := \{\phi : \mathbb{R}^{1+(r-1)} \rightarrow \mathbb{R} : \forall P \in \mathbb{R}[X], \forall x \in \mathbb{R} \lim_{y \rightarrow +\infty} |P(y_{r-1})\phi(x, y_1, \dots, y_{r-1})| = 0\}$$

where for functions $F : \mathbb{R}^d \rightarrow \mathbb{R}$, $\lim_{z \rightarrow +\infty} F(z) = 0$ means that for every $\epsilon > 0$ there exists $\exists \beta \geq 0$, such that $\min(z_1, \dots, z_d) \geq \beta \implies |F(z)| < \epsilon$.

For non-negative integers k, m , let $T[x, k, m]$ be the rooted tree of Figure 2. For ease of notation, if the context allows it, we will refer to $T[x, k, m]$ simply as T . For any non-negative integer t , $v \in V(T)$, let $\xi^t(v, T)$ be the embedding returned for node v by a GNN with activation function $\sigma \in H$ after t iterations. Let $M > 0$ be an integer. We will prove by induction on t that for any $t \in \{1, \dots, M\}$:

$$\begin{aligned} \xi^t(s, T) &= v_t + \eta_t(x, k, m) \\ \xi^t(x_{i \geq 2}, T) &= g_t(k) + \epsilon_t(x, k, m) \\ \xi^t(l_{i \geq 2}, T) &= h_t(k) + \nu_t(x, k, m) \\ \xi^t(x_1, T) &= g_t(k) + \epsilon_t^1(x, k, m) \\ \xi^t(l_1, T) &= h_t(k) + \nu_t^1(x, k, m) \end{aligned}$$

with the following properties: a) v_t is constant with respect to x and m . b) g_t and h_t are bounded Pfaffian functions that depend only on the combination and aggregation functions and the iteration t , and c) For every $t \in \{1, \dots, M\}$, $\nu_t, \eta_t, \epsilon_t, \epsilon_t^1, \nu_t^1 \in S_3$.

Base case: Obvious as all embeddings are initialized to 1.

Induction hypothesis: Suppose that for some integer $t \geq 0$, for every $u \leq t$:

- i) the above system of equation is verified
- ii) g_u, h_u are Pfaffian and bounded functions
- iii) $\nu_u, \eta_u, \epsilon_u, \epsilon_u^1$ and $\nu_u^1 \in S_3$

Induction step: First note that comb has bounded derivatives as the derivative of a bounded Pfaffian function is also bounded. Let $K_{\text{comb}} := \sup_{x \in \mathbb{R}^2} \|\partial_{x_1} \text{comb}(x), \partial_{x_2} \text{comb}(x)\|_2$.

The update rule for the leaves writes

$$\begin{aligned} \xi^{t+1}(l_{i \geq 2}, T) &= \text{comb}(\xi^t(l_i, T), \xi^t(x_i, T)) \\ &= \text{comb}(h_t(k) + \nu_{t,m}(k), g_t(k) + \epsilon_{t,m}(k)) \\ &:= \text{comb}(h_t(k), g_t(k)) + \nu_{t+1,m}(k) \\ &:= h_{t+1}(k) + \nu_{t+1,m}(k) \end{aligned}$$

where

$$\begin{aligned} |\nu_{t+1}(x, k, m)| &= |\text{comb}(h_t(k) + \nu_t(x, k, m), g_t(k) + \epsilon_t(x, k, m)) - \text{comb}(h_t(k), g_t(k))| \\ &\leq K_{\text{comb}}(\nu_t(x, k, m), \epsilon_t(x, k, m))\|_2 \\ &\leq K_{\text{comb}}\sqrt{2}\|(\nu_t(x, k, m), \epsilon_t(x, k, m))\|_\infty \end{aligned}$$

similarly, $\xi^{t+1}(l_1, T) = \text{comb}(h_t(k), g_t(k)) + \nu_{t+1}^1(x, k, m)$ where

$$\nu_{t+1}^1(x, k, m) = |\text{comb}(h_t(k) + \nu_t^1(x, k, m), g_t(k) + \epsilon_t^1(x, k, m)) - \text{comb}(h_t(k), g_t(k))|$$

For the intermediary vertices, we have

$$\begin{aligned}
\xi^{t+1}(x_{i \geq 2}, T) &= \text{comb}(\xi^t(x_i, T), \xi^t(s, T) + k\xi^t(l_i, T)) \\
&= \text{comb}(g_t(k) + \epsilon_t(x, k, m), v_t + \eta_t(x, k, m) + k(h_t(k) + \nu_t(x, k, m))) \\
&= \text{comb}(g_t(k), v_t + kh_t(k)) + \epsilon_{t+1}(x, k, m) \\
&:= g_{t+1}(k) + \epsilon_{t+1}(x, k, m)
\end{aligned}$$

Let $\Delta := \eta_t(x, k, m) + k\nu_t(x, k, m)$

$$\begin{aligned}
|\epsilon_{t+1, m}(k)| &= |\text{comb}(g_t(k) + \epsilon_t(x, k, m), v_t + \Delta + kh_t(k)) - \text{comb}(g_t(k), v_t + kh_t(k))| \\
&\leq K_{\text{comb}}\sqrt{2}\|(\epsilon_t(x, k, m), \Delta)\|_{\infty}
\end{aligned}$$

Similar inequalities can be derived for $\xi^t(x_1, T)$. Finally, for the root we have

$$\begin{aligned}
\xi^{t+1}(s, T) &= \text{comb}(\xi^t(s, T), \sum_{i=1}^m \xi^t(x_i, T)) \\
&= \text{comb}(v_t + \eta_t(x, k, m), g_t(x) + (m-1)g_t(k) + \epsilon_t^1(x, k, m) + (m-1)\epsilon_t(x, k, m)) \\
&:= \text{comb}(v_t, g_t(x) + (m-1)g_t(k)) + \eta_{t+1}(x, k, m)
\end{aligned}$$

Using Corollary A.1, for every $u \in \{1, \dots, t+1\}$, g_u is either constant equal to 0 on $\mathbb{R}$ or has no zeroes after some rank. Therefore, there exists an integer k^* such that for every $k_i \geq k^*$, for every $u \in \{1, \dots, t+1\}$:

- Case a) $g_u(k_i) > 0$
- Case b) or $g_u(k_i) < 0$
- Case c) $g_u = \tilde{0}$ on $\mathbb{R}$

(i.e. we can select a common $\beta = k^*$ for each iteration). Now, suppose that $|k| \geq \beta$.

By continuity of comb , g_t having a limit in $+\infty$,

$$\forall \epsilon > 0, \exists M, \forall k \geq \beta, \forall m \geq M, |\text{comb}(v_t, g_t(x) + (m-1)g_t(k)) - l| \leq \epsilon$$

where $l = \lim_{y \rightarrow +\infty} \text{comb}(v_t, y)$ if case a) for $u = t+1$ and $l = \lim_{y \rightarrow -\infty} \text{comb}(v_t, y)$ if case b) for $u = t+1$, and $l = \text{comb}(v_t, 0)$ otherwise (case c) for $u = t+1$). Hence

$$\begin{aligned}
\xi^{t+1}(s, T) &= \text{comb}(v_t, g_t(x) + (m-1)g_t(k)) + \eta_{t+1}(x, k, m) \\
&= \begin{cases} \lim_{y \rightarrow -\infty} \text{comb}(v_t, y) & \text{if } g_t(k^*) < 0 \\ \lim_{y \rightarrow +\infty} \text{comb}(v_t, y) & \text{if } g_t(k^*) > 0 \\ \text{comb}(v_t, 0) & \text{otherwise} \end{cases} + \eta_{t+1}(x, k, m)
\end{aligned}$$

where $\eta_{t+1} = \eta_{t+1}^1 + \eta_{t+1}^2$ (in the following, we suppose without loss of generality that $g_t(k^*) > 0$) with

$$\begin{aligned}
|\eta_{t+1}^1(x, k, m)| &:= |\text{comb}(v_t + \eta_t(x, k, m), g_t(x) + (m-1)g_t(k) + \epsilon_t^1(x, k, m) \\
&\quad + (m-1)\epsilon_t(x, k, m)) - \text{comb}(v_t, g_t(x) + (m-1)g_t(k))|
\end{aligned}$$

and

$$|\eta_{t+1}^{(2)}(x, k, m)| := |\text{comb}(v_t, g_t(x) + (m-1)g_t(k)) - \lim_{y \rightarrow +\infty} \text{comb}(v_t, y)|$$

In this way,

$$|\eta_{t+1}^1(x, k, m)| \leq K_{\text{comb}}\sqrt{2}\|(\eta_t(x, k, m), \epsilon_t^1(x, k, m) + (m-1)\epsilon_t(x, k, m))\|_{\infty}$$

The error term gets multiplied by m . We know by the induction hypothesis that for every polynomial P

$$\forall \epsilon > 0, \exists \beta, \forall m \geq \beta, \beta \leq |k| \implies |P(m)\epsilon_t(x, k, m)| \leq \epsilon$$

In these conditions $|(m-1)\epsilon_t(x, k, m)| \leq \epsilon$. The property is still verified at rank $t+1$ as $|P(m)(m-1)\epsilon_t(x, k, m)| \leq \underbrace{|P(m)(m-1)|}_{Q(m)}|\epsilon_t(x, k, m)|$ where $Q(X) := (X-1)P(X)$ is a polynomial.

The triangular inequality $|\eta_{t+1, m}(k)| \leq |\eta_{t+1, m}^{(1)}(k)| + |\eta_{t+1, m}^{(2)}(k)|$, this gives the desired property (after splitting the ϵ in half) and ends the induction. $\square$

Proof of Corollary 3.3. This follows from Theorem 3.2 and Proposition 2.1. $\square$

C PROOFS FOR SECTION 5

C.1 PROOF OF THEOREM 5.1

Let us overview the proof of Theorem 5.1. First, we will show that the constructive proof of (Barceló et al., 2020, Proposition 4.2) for CReLU holds for the linear threshold activation function σ_* defined in equation 1. After this, we will exploit that for a fast convergence of step-like function (see Definition 5.1) σ , σ^m converges fast to σ_* as the following lemma shows.

Lemma C.1. *Let σ be a step-like activation function. Then for all $x \notin (\varepsilon, 1 - \varepsilon)$ and all $m \geq N$,*

$$|\sigma^m(x) - \sigma_*(x)| \leq \varepsilon \left(\frac{1 + \eta}{2} \right)^{m-N}. \quad (3)$$

Moreover, if further $\eta = 0$, then for all $x \notin (\varepsilon, 1 - \varepsilon)$ and all $m \geq N$,

$$|\sigma^m(x) - \sigma_*(x)| \leq \frac{\varepsilon}{2^{2^{m-N}-1}}. \quad (4)$$

Proof of Theorem 5.1. Let Q be a query of GC2 of depth d , and let $(Q_1, \dots, Q_d)$ be an enumeration of its subformulas. First, we extend the constructive proof of (Barceló et al., 2020, Proposition 4.2) to GNNs with Linear Threshold activations σ_* . In this extension, we have to make sure we can control the error made efficiently⁴.

We consider the combination functions for $t \in [d]$, $\text{comb}_t : (x, y) \mapsto \sigma_*(A_t x + B_t y + C_t)$, with the same choices of $d \times d$ matrices A_t , B_t and C_t initially set to 0 and filled as described below. Each row of A_t , B_t and C_t corresponds to a subformula Q_j with $j \leq t$ as follows:

- 1) (Queries of depth 1) The $\ell' \leq d$ color functions used by Q_j are encoded in the first ℓ' rows of A_1 . To do so, the first ℓ' rows are the transpose of the first i -th canonical vectors of $\mathbb{R}^d$ for $i \in [\ell']$.
- 2) For $j > 1$, if $Q_j(x) = \exists^{>K} y E(x, y) \wedge Q_i(y)$ for some i with $j > i$. Set $C_j = -(K - 1)$, $A_{j,k} = 0$ for every $k \in [j]$, $B_{j,i} = 1$ and $B_{j,k} = 0$ for every $k \in [j]$.
- 3) For $j > 1$, if $Q_j = \neg Q_i$ for some i with $j > i$. Set $C_j = +1$, $A_{j,i} = -1$ and $A_{j,k} = 0$ for $k \in [j] - \{i\}$, and $B_{j,k} = 0$ for every $k \in [j]$.
- 4) For $j > 2$, if $Q_j = Q_{i_1} \wedge Q_{i_2}$ for some i_1, i_2 in $[j - 1]$. Set $C_j = -1$, $A_{j,i_1} = A_{j,i_2} = 1$, $A_{j,k} = 0$ for $k \in [j] - \{i_1, i_2\}$, and $B_{j,k} = 0$ for every $k \in [j]$.

For any integer $t < d$ and for every graph G and each vertex $v \in V(G)$ of G , consider

$$q^{t+1}(G, v) := \text{comb}_t \left(q^t(G, v), \sum_{w \in \mathcal{N}_G(v)} q^t(G, w) \right). \quad (5)$$

Observe that for each $t \in [d]$, $q^t(G, v) \in \{0, 1\}^d$. With the exact same inductive argument as in (Barceló et al., 2020, Proposition 4.2), we have that

$$q^t(G, v) = (Q_1(G, v), \dots, Q_d(G, v))^T.$$

We will now consider the GNN where we have substituted σ_* by its approximation σ^m . The new GNN will have combination functions of the form $\text{comb}_t^m : (x, y) \mapsto \sigma^m(A_t x + B_t y + C_t)$, where $t \in [d]$, and $\xi_{t+1}^m(G, v) = \text{comb}_t^m \left(\xi_t^m(G, v), \sum_{w \in \mathcal{N}_G(v)} \xi_t^m(G, w) \right)$, where $t \in [d - 1]$. Observe that comb_t^m is a neural network of width 1 and depth m with activation function σ .

Let Δ be a positive integer and fix a graph G with maximal degree bounded by Δ . For $t \in [d]$, we define

$$\epsilon_t := \max_{v \in V(G)} \|\xi^t(G, v) - q^t(G, v)\|_\infty.$$

We shall prove that, under the assumptions of the theorem, the following claim:

⁴The original proof to exactly compute the query can actually be extended for any function $g : \mathbb{R} \rightarrow \mathbb{R}$ verifying: there exists real $a < b$ such that i) for every real $x \leq a$, $g(x) = 0$, ii) for every real $x \geq b$, $g(x) = 1$ and iii) g is non-decreasing on $[a, b]$.

Claim. For any $t \in [d-1]$,

$$\text{if } 2(\Delta+2)\epsilon_t < 1, \text{ then } \epsilon_{t+1} \leq \begin{cases} \left(\frac{1+\eta}{2}\right)^{m-N} \varepsilon & \text{if } \eta > 0 \\ \frac{\varepsilon}{2^{2m-N-1}} & \text{if } \eta = 0 \end{cases}$$

Under this claim, using induction, we can easily show that if either

$$\eta > 0 \text{ and } m \geq N + \frac{2 + \log(\Delta+2)}{1 - \log(1-\eta)}$$

or

$$\eta = 0 \text{ and } m \geq N + \log(2 + \log(\Delta+2)),$$

then $\epsilon_d \leq 1/3 < 1/2$. Effectively, under any of the above assumptions, we have that

$$2(\Delta+2)\epsilon_t \implies 2(\Delta+2)\epsilon_{t+1}$$

by the claim and the fact that $\varepsilon < 1/2$. Since $\epsilon_0 = 0$, we conclude, by induction, that

$$\epsilon_d \leq \begin{cases} \left(\frac{1+\eta}{2}\right)^{m-N} \varepsilon & \text{if } \eta > 0 \\ \frac{\varepsilon}{2^{2m-N-1}} & \text{if } \eta = 0 \end{cases} \leq 1/3 < 1/2,$$

providing the expressivity. Note that the size of the above GNN is $2d + m$, and so the statement of the theorem follows.

Now, we prove the claim. For $t \in [d-1]$ we have by the triangular inequality

$$\begin{aligned} \epsilon_{t+1} &= \max_{v \in V(G)} \|\xi^{t+1}(G, v) - q^{t+1}(G, v)\|_\infty \\ &= \max_{v \in V(G)} \left\| \text{comb}_t^m \left(\xi^t(G, v), \sum_{w \in \mathcal{N}_G(v)} \xi^t(G, w) \right) - \text{comb}_t \left(q^t(G, v), \sum_{w \in \mathcal{N}_G(v)} q^t(G, w) \right) \right\|_\infty \\ &\leq \max_{v \in V(G)} \left\| \text{comb}_t^m \left(\xi^t(G, v), \sum_{w \in \mathcal{N}_G(v)} \xi^t(G, w) \right) - \text{comb}_t \left(\xi^t(G, v), \sum_{w \in \mathcal{N}_G(v)} \xi^t(G, w) \right) \right\|_\infty \\ &\quad + \max_{v \in V(G)} \left\| \text{comb}_t \left(\xi^t(G, v), \sum_{w \in \mathcal{N}_G(v)} \xi^t(G, w) \right) - \text{comb}_t \left(q^t(G, v), \sum_{w \in \mathcal{N}_G(v)} q^t(G, w) \right) \right\|_\infty. \end{aligned}$$

The first maximum of the last inequality is bounded above by

$$\sup \{ \|\text{comb}_t^m(x, y) - \text{comb}_t(x, y)\|_\infty : \text{for all } i, x_i \notin (\epsilon_t, 1 - \epsilon_t), y_i \notin (\Delta\epsilon_t, 1 - \Delta\epsilon_t) \},$$

since each component of $\xi^t(G, v)$ is at distance at most ϵ_t from 0 or 1 by definition of ϵ_t . Now, $\text{comb}_t^m(x, y) = \sigma^m(A_t x + B_t y + C_t)$ and $\text{comb}_t(x, y) = \sigma_*(A_t x + B_t y + C_t)$ with $A_t, B_t \in \{-1, 0, 1\}^{d \times d}$ and $C_t \in \mathbb{Z}^d$. Moreover, the matrix A_t has at most two non-zero entries per row, and B_t at most one non-zero entry. Since the 1-norm of each row of A_t is at most 2, and the one of B_t is at most 1, we can improve the upper-bound with

$$\sup \{ \|\sigma^m(z) - \sigma_*(z)\| : z \notin ((\Delta+2)\epsilon_t, 1 - (\Delta+2)\epsilon_t) \},$$

By assumption, this quantity is bounded from above by the claimed by Lemma C.1.

For the second maximum, define

$$x := \xi^t(G, v) \text{ and } y := \sum_{w \in \mathcal{N}_G(v)} \xi^t(G, w)$$

and

$$x' := q^t(G, v) \text{ and } y' := \sum_{w \in \mathcal{N}_G(v)} q^t(G, w)$$

to simplify notation. By definition of ϵ_t , we have

$$\|x - x'\|_\infty \leq \epsilon_t \text{ and } \|y - y'\|_\infty \leq \Delta\epsilon_t.$$

Now, arguing as above regarding A_t and B_t , we have that

$$\|(A_t x + B_t y + C_t) - (A_t x' + B_t y' + C_t)\|_\infty = \|A_t(x - x') + B_t(y - y')\|_\infty \leq (2 + \Delta)\epsilon_t.$$

In this way, $A_t x + B_t y + C_t$ is a real vector with ∞ -norm $< 1/2$ to the integer vector $A_t x' + B_t y' + C_t$ and so we conclude that

$$\text{comb}_t(x, y) = \text{comb}_t(x', y'),$$

because σ_* takes the same value on an integer and on a real number at distance less than $\frac{1}{2}$ from it. Hence the value of the difference under the second maximum is zero, proving the claim. $\square$

To prove Lemma C.1 we will need the following well-known lemma.

Lemma C.2. *Let $f : \mathbb{R} \rightarrow \mathbb{R}$ be a C^2 -function, $x_0 \in \mathbb{R}$ a fixed point of f (i.e., $f(x_0) = x_0$) and $r > 0$. If $|f'(x_0)| = \eta < 1$, $\sup_{t \in [x_0 - r, x_0 + r]} |f''(t)| = H > 0$ and*

$$Hr \leq 1 - \eta,$$

then for every integer $n \geq 0$ and $x \in [x_0 - r, x_0 + r]$,

$$|f^n(x) - x_0| \leq \left(\frac{1 + \eta}{2}\right)^n |x - x_0| \leq \left(\frac{1 + \eta}{2}\right)^n r.$$

Moreover, if $\eta = 0$, then for all n and $x \in [x_0 - r, x_0 + r]$,

$$|f^n(x) - x_0| \leq \left(\frac{H}{2}\right)^{2^n - 1} |x - x_0|^{2^n} \leq \frac{r}{2^{2^n - 1}}.$$

Proof of Lemma C.1. Fix $x \notin (\varepsilon, 1 - \varepsilon)$ and let $r = \min\{\varepsilon, (1 - \eta)/H\}$. By (c), we have that $|\sigma^N(x) - \sigma_*(x)| \leq r$ with $Hr \leq 1 - \eta$. Hence, since $\sigma_*(x)$ is a fixed point of σ , by (a), we have, by Lemma C.2, (b) and (e), that

$$|\sigma^m(\sigma^N(x)) - \sigma_*(x)|$$

converges to zero with the desired speed. $\square$

Proof of Lemma C.2. Let $H = \sup_{t \in [x_0 - r, x_0 + r]} |f''(t)|$. Then we have that $r \leq \frac{1 - \eta}{H}$ and so for $x \in [x_0 - r, x_0 + r]$, we will have, by Taylor's theorem, for some s between x and x_0 ,

$$f(x) = x_0 \pm \eta(x - x_0) + \frac{1}{2}f''(s)(x - x_0)^2,$$

where we write x_0 since $x_0 = f(x_0)$. Rearranging the expression and taking absolute values,

$$|f(x) - x_0| \leq \eta|x - x_0| + \frac{1}{2}H|x - x_0|^2 \leq \left(\eta + \frac{Hr}{2}\right)|x - x_0|,$$

since $|f''(s)| \leq H$ and $|x - x_0| \leq r$. Now, by assumption, $Hr/2 \leq \frac{1 - \eta}{2}$, and so

$$|f(x) - x_0| \leq \left(\frac{1 + \eta}{2}\right)|x - x_0|.$$

For the case $\eta = 0$, we have instead

$$|f(x) - x_0| \leq \frac{H}{2}|x - x_0|^2.$$

Hence in both cases, induction on n ends the proof. $\square$

C.2 PROOFS OF PROPOSITIONS 5.2, 5.3 AND 5.4

We now focus on Proposition 5.2, 5.3 and 5.4. First, we prove Proposition 5.2. Second, we show that Proposition 5.4 follows from Proposition 5.3. Third and last, we prove Proposition 5.3 by proving the more precise Proposition C.3.

Proof of Proposition 5.2. We just need to consider $\sigma_{\arctan}(x) = \frac{1}{2} + \frac{2}{\pi} \arctan(2x - 1)$. Then a simple computation shows that the given claim is true. $\square$

Proof of Proposition 5.4. By equation 2, we can express $\tanh$ as a NN with activation function Sigmoid. Hence, we can transfor $\bar{\sigma}_{\tanh}$ into σ_{Sigmoid} by substituting each $x \mapsto \tanh(x)$ by $x \mapsto 2 \text{Sigmoid}(2x) - 1$. A Simple counting finishes the proof. $\square$

Observe that Proposition 5.3 follows from the following precise proposition that we will prove instead.

Proposition C.3. *Let*

$$\sigma_{\tanh}(x) := \frac{1}{2} + \frac{\tanh(x - 1/2)}{2 \tanh(1/2)}$$

and

$$\bar{\sigma}_{\tanh}(x) := \sigma_{\tanh} \left(\frac{\tanh(2ax - a) - \alpha \tanh(4ax - 2a) + \tanh(6ax - 3a)}{2(\tanh(a) - \alpha \tanh(2a) + \tanh(3a))} + \frac{1}{2} \right)$$

where $a \in (0.45, 0.46)$ and $\alpha \in (3.14, 3.15)$ are such that

$$\min \left\{ \frac{\text{sech}^2(x) + 3 \text{sech}^2(3x)}{\text{sech}^2(2x)} \mid x \in \mathbb{R} \right\} = \frac{\text{sech}^2(a) + 3 \text{sech}^2(3a)}{\text{sech}^2(2a)} = \alpha.$$

Then $\sigma_{\tanh}$ and $\bar{\sigma}_{\tanh}$ are step-like activation functions with, respectively, $\eta = 0.86$, $\varepsilon = 0.16$, $N = 0$ and $H = 0.84$, and $\eta = 0$, $\varepsilon = 0.2$, $N = 0$ and $H = 2.2$.

Proof. The claim about $\sigma_{\tanh}(x)$ follows after a simple numerical computation. For the claim about $\bar{\sigma}_{\tanh}(x)$ observe that

$$\bar{\sigma}_{\tanh}(x) = \sigma_{\tanh}(\tau(x))$$

where, by the choices above, we have that:

- (0) $\tau(0) = 0$, $\tau(1) = 1$.
- (1) τ is strictly increasing.
- (2) $\tau'(x) = 0$ if and only if $x \in \{0, 1\}$.

Hence $\bar{\sigma}_{\tanh}$ is fixes 0 and 1, is strictly increasing and, by the chain rule, satisfies $\eta = 0$. The rest of the constants are found by numerical computation. $\square$

C.3 PROOFS OF COROLLARIES 5.5, 5.6 AND 5.7

Proof of Corollary 5.5. This is a combination of Theorem 5.1 with Proposition 5.2. $\square$

Proof of Corollary 5.6. This is a combination of Theorem 5.1 with Proposition 5.3. $\square$

Proof of Corollary 5.7. This is a combination of Theorem 5.1 with Proposition 5.4. $\square$

D NUMERICAL EXPERIMENTS

Computing and approximating queries. For Q_1 , the ReLU GNN has sum-aggregation, 4 iterations (depth of Q). The input and output dimensions are set for convenience to 4 to have a constant dimension for all iterations. In particular, all features vectors of the input vertices are set to e_1 , the first canonical vector of $\mathbb{R}^4$. The weights and biases are set as follows .

$$A = \begin{pmatrix} 0 & 0 & 0 & 0 \\ -1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & -1 & 0 \end{pmatrix}, B = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix}, c = \begin{pmatrix} -1 \\ 1 \\ 0 \\ 1 \end{pmatrix}$$

For every $t \in [4]$, and every $x, y \in \mathbb{R}^4$ $\text{comb}_t(x, y) := \text{ReLU}(Ax + By + c)$

.
For Q_2 , the ReLU GNN has 7 iterations. The weights and biases are set as:

$$A = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 1 & 0 \end{pmatrix}, B = \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}, c = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ -1 \\ 0 \\ -1 \end{pmatrix}$$