

Principled Rewriting of ONNX Operators for Reluctant Solvers

Alban Grastien[✉], Guilhem Ardouin[✉], and Julien Girard-Satabin[✉]

Université Paris-Saclay, CEA, List, F-91120, Palaiseau, France
firstname.lastname@cea.fr
julien.girard2@cea.fr

Abstract. Neural networks (NN) are often represented using the Open Neural Network Exchange (ONNX) format that provides a rich set of operators. This richness implies that many tools support only a subset of ONNX operators and, consequently, a subset of NNs. However, some of these operators can be expressed equivalently as combinations of simpler, already supported, operators. In this work, we present an extension of CAISAR—a platform for verification of NN that uses external solvers—that automatically transforms these redundant operators into simpler ones before calling existing solvers, thus extending the range of these solvers. This work focuses on the ARGMAX operator, proposing a transformation, proving its correctness, and presenting experimental results.

Keywords: Formal Specification and Verification, Machine Learning

1 Introduction

As machine learning, and artificial intelligence in general, has become more and more prominent, there has been an increasing need for verification of such systems [6]. This has led to the introduction of standard representations, both of a dedicated exchange format for neural networks, ONNX [7], and of the problems to solve with the VNNLIB format [8].

While many solvers are being developed, the efforts are not always matched with appropriate resources. This is in particular true for contributions from the academic world. The ONNX format is comprehensive and includes no fewer than 196 operators. Some of these operators are redundant, in that they can be defined as a combination of simpler operators. While this is beneficial for expressivity—they make it easier to understand what their purpose is in

Solver	Nb op. (/196)
nnenum [5]	17
α - β -CROWN [12]	22
Marabou [9]	16
Maraboupy [13]	24
PyRAT [10]	54

Table 1: Number of operators supported by the most common solvers.

the architecture—they represent a burden for the developers of solvers who are more interested in testing new solving techniques than supporting an ever-growing list of operators. This is illustrated on Tab. 1 where we see that many solvers support a very narrow list of operators. Furthermore, a prover may support an operator that is unsupported by others. For instance, prover `nenum` supports vector concatenation which `Marabou` does not, while the latter supports vector transposition. We note that those operations are rather common in day-to-day machine learning practice. It is therefore unlikely that provers will be able to handle a standard neural network out-of-the-box, forcing the expert to do manual alterations to the architecture, which are error-prone and need to be done on a per-prover basis.

In this context, `CAISAR` [1] has been introduced as a verification platform that offers various tools to extend the capabilities of existing solvers to a larger array of problems. Of interest here, we introduce in this paper a new feature of `CAISAR` that allows it to reformulate existing neural networks into equivalent ones that use a “simpler” (or different) set of operators. Thanks to this reformulation, any solver that supports the restricted set of operators can be used to solve problems that include the more complicated operator, even if this operator is not supported natively by the solver. Similar reformulations, called *simplifications*, have been proposed in `DNNV` [11] that, as far as we know, is no longer maintained. Tensor compilers also often rewrite neural networks before sending them to the architecture where computation will take place to improve performance; there too, we see an effort to prove that these *rewrites* are valid [3], but notice that their objective is to improve computation time and not to increase support.

In particular, we focus on the `ARGMAX` operator, which computes the index of the elements in certain rows of a tensor that have maximal value. This choice is motivated by `ARGMAX` ubiquity in standard image classification settings, and hyper-properties specifications [2]. We propose a systematic transformation for any neural network that includes this operator, show the correctness of this transformation, and present an implementation of the transformation. We demonstrate its applicability on a test case where we simplify the expression of a neural network property by using an `ARGMAX` operator inside the model to be verified and then transform it to verify the property using existing solvers.

Section 2 introduces the problem and our contributions: the translation of a stripped-down version of `ARGMAX` into a different set of operators (support for the more advanced features is presented in Appendix A) and its integration in `CAISAR`. Section 3 provides a proof of the correctness of the translation. This extension is used experimentally in Section 4 before concluding in Section 5.

2 Preliminaries & Contributions

In this section, we provide the results in a slightly informal manner. The formal results are provided in Section 3. We first introduce the notions of tensors, neu-

ral networks, and ONNX. We then discuss how we propose to reformulate the ARGMAX operator, and finally discuss its integration in CAISAR.

2.1 Tensors, Neural Networks, and ONNX

We use the notation $[N]$ to represent the set $\{0, \dots, N-1\}$.

A *tensor* $\mathbf{t} \in \mathbf{T}$ is a multi-dimensional array of *primitive elements* such as integers or numbers in floating-point format. Tensor $\mathbf{t}$ has *shape* $sh(\mathbf{t}) = (\dim_1, \dots, \dim_r)$ where r is the *rank* of $\mathbf{t}$ and each $\dim_\ell \in \mathbf{N}$ is the *size* of the ℓ -th *dimension*. The elements of tensor $\mathbf{t}$ are accessed by $\mathbf{t}[\mathbf{i}]$ where $\mathbf{i} = [i_1, \dots, i_r] \in \text{dom}(sh(\mathbf{t})) \stackrel{\text{def}}{=} [\dim_1] \times \dots \times [\dim_r]$ is a vector of positions for each dimension.

A tensor $\mathbf{t}$ of shape $sh(\mathbf{t}) = (\dim_1, \dots, \dim_r)$ contains

$$|\text{dom}(sh(\mathbf{t}))| = \dim_1 \times \dots \times \dim_r$$

elements. Given a vector $\mathbf{i} = [i_1, \dots, i_r]$, we write $\mathbf{i}_{\ell \mapsto j}$ to refer to the same vector where the ℓ -th position is replaced with j , i.e., $[i_1, \dots, i_{\ell-1}, j, i_{\ell+1}, \dots, i_r]$.

An n -ary *ONNX operator* $\text{OP} : \mathbf{T}^n \mapsto \mathbf{T}$ of arity n is an operator that takes n tensors as input and returns a tensor. Throughout the paper, we distinguish the operator and its semantics: the operator is denoted with small caps (OP) and its mathematical semantics with italics (*op*). In this paper, we only consider the following operators:

1. $\text{CONST}_{\mathbf{t}}$ (which returns $\mathbf{t}$),
2. ADD (element-wise sum of two tensors),
3. SIGN (element-wise sign of a tensor, that returns 1, 0, -1 if the element is strictly positive, zero, or strictly negative, respectively),
4. RELU (element-wise application of the rectified linear unit activation function that, for any x , returns $\max(0, x)$),
5. TRANSPOSE_p (change of indexing of elements in the tensor following the permutation p),
6. MATMUL (standard matrix multiplication),
7. ARGMAX_ℓ (function that returns the element with highest value across dimension ℓ),
8. SOFTMAX (smooth approximation of a one-hot vector).

The semantics of each operator is formally defined in Section 3, but we clarify ARGMAX_ℓ as it is necessary to understand the translation. For a given vector $\mathbf{i} \in \text{dom}(sh(\mathbf{t}))$, for two indices $j, k \in \dim_\ell$, the element at position $\mathbf{i}_{\ell \mapsto j}$ (or, when clear from the context, simply *at position* j) is considered *higher* than the one at position k if the following property holds:

$$\begin{cases} \mathbf{t}[\mathbf{i}_{\ell \mapsto j}] > \mathbf{t}[\mathbf{i}_{\ell \mapsto k}] & \text{if } k < j, \\ \mathbf{t}[\mathbf{i}_{\ell \mapsto j}] \geq \mathbf{t}[\mathbf{i}_{\ell \mapsto k}] & \text{if } k > j, \end{cases} \quad (1)$$

i.e., the value $\mathbf{t}[\mathbf{i}_{\ell \mapsto j}]$ at position j is greater than that at position k , where tie breaking favours the smaller index. The *argmax_ℓ* semantics specifies that the size

of the output tensor at dimension ℓ would be one, and that the value at position $\mathbf{i}_{\ell \mapsto 0}$ should be the index $\hat{j}$ of highest value (amongst all the elements at position $\mathbf{i}_{\ell \mapsto k}$).

A *neural network* (NN) described in ONNX format is a dataflow (essentially a graph) of tensors where nodes are labelled with ONNX operators, as in Fig. 1 (here, limited to a non-branching graph). The input tensor is passed to the first node which performs its computation and passes it to the next node(s) until the end of the network is reached.

2.2 CAISAR

CAISAR is a platform that aims at helping practitioners verify NNs. For instance, it supports multiple solvers via a single interface, and makes it easy to switch between these solvers to build on the strengths of the different technologies. It also allows the practitioner to define high-level goals through the use of the rich specification language WhyML, as we illustrate in Section 4. Most solvers require a NN as input, which is provided by CAISAR in ONNX format.

In this paper, we introduce a new functionality to CAISAR, which is the ability to replace certain complex ONNX operators that are not supported by some solvers. In particular, we show in the next subsection that ARGMAX can be rewritten as a semantically equivalent network using simpler operators, so that any property that is satisfied by a NN that features this operator should be satisfied by a NN where this operator has been transformed (assuming all operators do satisfy their specification). This means that a solver that does not support ARGMAX can still be used to verify properties on NN that include this operator as long as it supports the operators of the translation.

CAISAR supports all the solvers mentioned in Table 1. Each solver can be run in different modes, based on the command arguments. For each of these *variants*, a *driver* needs to be provided that indicates what preprocessing needs to be performed. We specify the list of transformations in this driver file; before calling the solver, CAISAR applies these transformations to the ONNX.

In a future implementation, we hope to automatically determine which transformations are necessary. Notice however that there might be different ways to transform a complex operator with varying impact on performance; choosing the preferred transformations might be a delicate exercise.

2.3 Transforming the ARGMAX

We use Figure 1 as a reference. First, the ℓ -th dimension is swapped with the r -th one via a TRANSPOSE operation so that the MATMUL operation is applied on dimension r ; the TRANSPOSE operation is reversed at the end of the translation. The permutation is

$$p_{\ell,r} \stackrel{def}{=} [1, \dots, \ell - 1, r, \ell + 1, \dots, r - 1, \ell].$$

(If $\ell = r$, no transpose is performed.) Notice that TRANSPOSE can be a non-trivial operation when applied to a tensor; here however, the solver reasons about

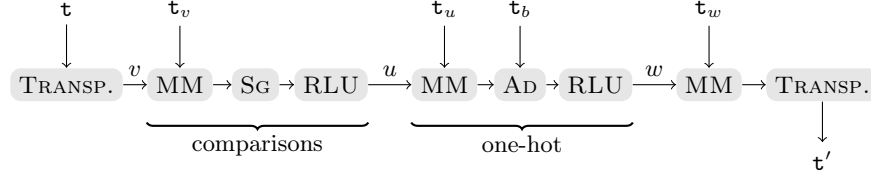


Fig. 1: Sub-network replacing the ARGMAX operator. The tensors and the semantics of v , u , and w are described in the text.

the effects of this operation but does not apply it in practice. At this point, the tensor is simply a collection of $|dom(sh(t))|/\dim_\ell$ vectors of dimension $\dim_\ell$; we simplify the discussion by considering one such vector that we denote v (i.e., $v[j] \stackrel{def}{=} t[i_{\ell \rightarrow j}]$ for the i of interest).

In a second step, the sub-network computes a vector u of size $|C_{\dim_\ell}^2| = \frac{\dim_\ell(\dim_\ell - 1)}{2}$ (where $C_{\dim_\ell}^2$ is “choose 2 from $[\dim_\ell]$ ”). Each element of vector u is associated with a pair $p \in C_{\dim_\ell}^2$ and we write $\xi(p)$ the index of this element. In other words:

1. $\forall p \in C_{\dim_\ell}^2. \xi(p) \in [|C_{\dim_\ell}^2|]$ and
2. $\forall \{p, p'\} \subseteq C_{\dim_\ell}^2. p \neq p' \Rightarrow \xi(p) \neq \xi(p')$.

For $j < k$, the sub-network ensures $u[\xi(\{j, k\})] = 1$ if $v[k]$ is higher than $v[j]$, and $u[\xi(\{j, k\})] = 0$ otherwise. This is done by computing $v[k] - v[j]$ through a MATMUL operation followed with a SIGN and a RELU operations.

In a third step, it computes a one-hot vector w where $w[\hat{j}]$ is 1 if $\hat{j}$ is the index with highest value in v , and all other $w[k]$ are 0. This is done via a MATMUL operation that adds up $u[\xi(\{j, k\})]$ for all $j > k$ and $(1 - u[\xi(\{j, k\})])$ for all $k > j$ (this should equate $\dim_\ell - 1$ for $\hat{j}$ and less for other indices) and subtracts $\dim_\ell - 2$ to the result, and followed with a RELU.

In the fourth step, the one-hot vector is multiplied (via a MATMUL) with vector $[0, 1, \dots, \dim_\ell - 1]$ so that the last dimension has size one and its value is the position of the 1 in w (starting from index 0).

Description of the intermediate tensors For completeness, we describe the tensors used in the translation.

Tensor t_v is a tensor of shape $(\dim_\ell, C_{\dim_\ell}^2)$ defined by

$$\forall \{j, k\} \in C_{\dim_\ell}^2. \forall i \in [\dim_\ell]. t_v[i, \xi(\{j, k\})] = \begin{cases} 1 & \text{if } i = \max(j, k) \\ -1 & \text{if } i = \min(j, k) \\ 0 & \text{otherwise.} \end{cases}$$

Tensor $\mathbf{t}_u$ is the transpose of $\mathbf{t}_v$, i.e., a tensor of shape $(C_{\dim_\ell}^2, \dim_\ell)$ such that

$$\forall \{j, k\} \in C_{\dim_\ell}^2. \forall i \in [\dim_\ell]. \mathbf{t}_u[\xi(\{j, k\}), i] = \begin{cases} 1 & \text{if } i = \max(j, k) \\ -1 & \text{if } i = \min(j, k) \\ 0 & \text{otherwise.} \end{cases}$$

Tensor $\mathbf{t}_b$ is a tensor of shape $(\dim_1, \dots, \dim_{r-1}, \dim_\ell)$ such that

$$\forall \mathbf{i} \in \text{dom}(sh(\mathbf{t}_b)). \forall j \in [\dim_\ell]. \quad \mathbf{t}_b(\mathbf{i}_{r \mapsto j}) = j - 1.$$

Tensor $\mathbf{t}_w$ is a tensor of shape $(\dim_\ell, 1)$ such that

$$\forall j \in [\dim_\ell]. \quad \mathbf{t}_w[j, 0] = j.$$

3 Proof of Correctness

In this proof, we demonstrate that the sub-NN described in Fig. 1 satisfies the *argmax* specification, i.e., that an ARGMAX operator can be replaced with the sequence of operators from the figure.

3.1 Specification

We provide the specification for each operator. Each ONNX operator OP is assumed to satisfy the corresponding specification *op*. While *softmax* is not used in the proof, we included it nonetheless as it is relevant for Section 4.

Notice that, for simplicity, we ignore or enforce certain aspects of these operators such as the broadcasting of binary operations. We also assume a simplified version of *argmax* which we extend in the Appendix. In what follows, a *shape-preserving* operator refers to an operator whose output tensor has the same shape as its input tensor(s).

- *const_t* is a 0-ary operator that outputs tensor $\mathbf{t}$: $\text{const}_t() = \mathbf{t}$.
- *add* is a shape-preserving binary operator s.t., for all $\mathbf{i} \in \text{dom}(sh(\mathbf{t}_1))$, $\text{add}(\mathbf{t}_1, \mathbf{t}_2)[\mathbf{i}] = \mathbf{t}_1[\mathbf{i}] + \mathbf{t}_2[\mathbf{i}]$ holds.
- *sign* and *relu* are shape-preserving unary operators s.t., for all $\mathbf{i} \in \text{dom}(sh(\mathbf{t}))$, hold

$$\text{sign}(\mathbf{t})[\mathbf{i}] = \begin{cases} 1 & \text{if } \mathbf{t}[\mathbf{i}] > 0 \\ -1 & \text{if } \mathbf{t}[\mathbf{i}] < 0 \\ 0 & \text{if } \mathbf{t}[\mathbf{i}] = 0, \end{cases} \quad \text{and} \quad \text{relu}(\mathbf{t})[\mathbf{i}] = \begin{cases} \mathbf{t}[\mathbf{i}] & \text{if } \mathbf{t}[\mathbf{i}] \geq 0 \\ 0 & \text{if } \mathbf{t}[\mathbf{i}] < 0. \end{cases}$$

- *transpose_p* where p is a permutation of the list of indices $[1, \dots, r]$ is a unary operator over tensors of rank r . Given a tensor $\mathbf{t}$ of shape $sh(\mathbf{t}) = (\dim_1, \dots, \dim_r)$, it generates a tensor of shape $(\dim_{p[1]}, \dots, \dim_{p[r]})$ such that, $\forall [i_1, \dots, i_r] \in \text{dom}(sh(\mathbf{t}))$, holds

$$\text{transpose}_p(\mathbf{t})[i_{p[1]}, \dots, i_{p[r]}] = \mathbf{t}[i_1, \dots, i_r].$$

- *matmul* is a binary operator that represents the standard matrix multiplication over tensors of same rank. Specifically, let $\mathbf{t}_1$ be a tensor of shape $(\dim_1^1, \dots, \dim_r^1)$. Tensor $\mathbf{t}_2$ must be a tensor of rank 2 of shape $(\dim_1^2, \dim_2^2)$ such that $\dim_r^1 = \dim_1^2$. Then, tensor $\mathbf{t}' = \text{matmul}(\mathbf{t}_1, \mathbf{t}_2)$ has rank r with shape $sh(\mathbf{t}') =$

$$(\dim_1^1, \dots, \dim_{r-2}^1, \dim_{r-1}^1, \dim_2^2)$$

that satisfies $\forall \mathbf{i} = [i_1, \dots, i_r] \in \text{dom}(sh(\mathbf{t}')), \text{matmul}(\mathbf{t}_1, \mathbf{t}_2)[\mathbf{i}] =$

$$\sum_{j \in [\dim_r^1]} \mathbf{t}_1[\mathbf{i}_{r \rightarrow j}] \times \mathbf{t}_2[j, i_r].$$

- *argmax_ℓ* is a unary operator that applies to tensors of rank ℓ or more and that indicates the index of the highest value across the ℓ -th dimension. In case of a tie, *argmax_ℓ* reports the lowest index. If tensor $\mathbf{t}$ has shape $(\dim_1, \dots, \dim_r)$, then tensor $\mathbf{t}' = \text{argmax}_\ell(\mathbf{t})$ has shape

$$(\dim_1, \dots, \dim_{\ell-1}, 1, \dim_{\ell+1}, \dots, \dim_r),$$

i.e., identical to that of $\mathbf{t}$ except that the ℓ -th dimension collapses to 1. Furthermore, $\forall \mathbf{i} \in \text{dom}(sh(\mathbf{t}')), \mathbf{t}'[\mathbf{i}_{\ell \rightarrow 0}]$ is the position of the highest element of $[\mathbf{t}[\mathbf{i}_{\ell \rightarrow 0}], \dots, \mathbf{t}[\mathbf{i}_{\ell \rightarrow \dim_\ell - 1}]]$.

- *softmax* is a shape-preserving unary operator that performs the soft-max overall the last layer of the tensor. Specifically, $\forall \mathbf{i} = [i_1, \dots, i_r]$, $\text{softmax}(\mathbf{t})[\mathbf{i}]$ equals

$$\frac{e^{\mathbf{t}[\mathbf{i}]}}{\sum_{k \in [\dim_r]} e^{\mathbf{t}[\mathbf{i}_{r \rightarrow k}]}}$$

3.2 Proof

Consider a vector $\mathbf{i} \in \text{dom}(sh(\text{ARGMAX}_\ell(\mathbf{t})))$ and index $\hat{j} \stackrel{\text{def}}{=} \text{ARGMAX}_\ell(\mathbf{t})[\mathbf{i}]$. With reference to Figure 1, our transformation replaces ARGMAX with a sub-NN that contains 9 ONNX operators. The proof consists in tracking properties of the tensor at different steps of this sub-NN in relation with $\mathbf{i}$ and $\hat{j}$. To this end, for each step $s \in \{0, \dots, 9\}$, we use $\mathbf{t}_s$ to refer to the tensor after the s -th operation. For instance, $\mathbf{t}$ is also called $\mathbf{t}_0$ and the tensor after the *sign* operator is $\mathbf{t}_3$. Our goal is to demonstrate that $\mathbf{t}_9[\mathbf{i}]$ evaluates to $\text{ARGMAX}_\ell(\mathbf{t})[\mathbf{i}]$, i.e., $\hat{j}$.

Properties of $\mathbf{t}_0$ By definition of $\hat{j}$, we have:

- $\forall k \in [\dim_\ell]. \mathbf{t}_0[\mathbf{i}_{\ell \rightarrow \hat{j}}] \geq \mathbf{t}_0[\mathbf{i}_{\ell \rightarrow k}]$ and
- $\forall k \in [\dim_\ell]. k < \hat{j} \implies \mathbf{t}_0[\mathbf{i}_{\ell \rightarrow \hat{j}}] > \mathbf{t}_0[\mathbf{i}_{\ell \rightarrow k}]$.

Properties of $\mathbf{t}_1 = \text{transpose}_{p_{\ell,r}}(\mathbf{t}_0)$ Let $\mathbf{i}'$ be the vector $\mathbf{i}$ after the permutation has been applied; vector $\mathbf{i}$ will not be used again until the end of the proof. We can rewrite the previous properties with the new vector:

- $\forall k \in [\dim_\ell]. \mathbf{t}_1[\mathbf{i}'_{r \rightarrow \hat{j}}] \geq \mathbf{t}_1[\mathbf{i}'_{r \rightarrow k}]$ and
- $\forall k \in [\dim_\ell]. k < \hat{j} \implies \mathbf{t}_1[\mathbf{i}'_{r \rightarrow \hat{j}}] > \mathbf{t}_1[\mathbf{i}'_{r \rightarrow k}]$.

Properties of $\mathbf{t}_4 = \text{relu}(\text{sign}(\text{matmul}(\mathbf{t}_1, \mathbf{t}_v)))$ The following properties hold:

- The values in tensor $\mathbf{t}_4$ are all in $\{0, 1\}$. This is due to the application of *sign* followed with *relu*.
- $\forall k < \hat{j}. \mathbf{t}_4[\mathbf{i}'_{r \mapsto \xi(k, j)}] = 1$. This is because $\mathbf{t}_1[\mathbf{i}'_{r \mapsto \hat{j}}]$ is strictly greater than $\mathbf{t}_1[\mathbf{i}'_{r \mapsto k}]$, so that their difference is positive, the *sign* operation returns 1, and the *relu* operation returns 1.
- $\forall k > \hat{j}. \mathbf{t}_4[\mathbf{i}'_{r \mapsto \xi(k, j)}] = 0$. This is because $\mathbf{t}_1[\mathbf{i}'_{r \mapsto \hat{j}}]$ is greater than or equal to $\mathbf{t}_1[\mathbf{i}'_{r \mapsto k}]$, so that their difference is 0 or less, the *sign* operation returns 0 or -1 , and the *relu* operation returns 0.

We do not need to consider the values associated with pair $\{j, k\} \not\propto \hat{j}$ (except for the fact that they belong to $\{0, 1\}$, first property), because the comparison of k with $\hat{j}$ will suffice to infer relevant information about this index k .

Properties of $\mathbf{t}_5 = \text{matmul}_u(\mathbf{t}_4)$ The *matmul* operation between $\mathbf{t}_4$ and $\mathbf{t}_u$ implies that, for any $j \in [\text{dim}_\ell]$, the following holds:

$$\mathbf{t}_5[\mathbf{i}'_{r \mapsto j}] = \sum_{k < j} \mathbf{t}_4[\mathbf{i}'_{r \mapsto \xi(k, j)}] \times 1 + \sum_{k > j} \mathbf{t}_4[\mathbf{i}'_{r \mapsto \xi(k, j)}] \times (-1)$$

where k ranges over $[\text{dim}_\ell]$. Thus we can show:

- $\mathbf{t}_5[\mathbf{i}'_{r \mapsto \hat{j}}] = \hat{j}$. From the previous paragraph, we know that the value on $\mathbf{t}_4$ is 1 when k is lower than $\hat{j}$ and 0 otherwise. Thus, the sum adds up exactly to $\sum_{k < \hat{j}} 1 \times 1 + \sum_{k > \hat{j}} 0 \times (-1) = \hat{j}$.
- $\forall j < \hat{j}. \mathbf{t}_5[\mathbf{i}'_{r \mapsto j}] \leq j - 1$. Because the values of tensor $\mathbf{t}_4$ are between 0 and 1, an upper bound to this sum is j ; however, the index $k = \hat{j} > j$ yields a value $\mathbf{t}_4[\mathbf{i}'_{r \mapsto \xi(k, j)}] = 1$, so that the upper bound for the sum is $\sum_{k < j} 1 \times 1 + \sum_{k > j, k \neq \hat{j}} 0 \times (-1) + 1 \times (-1) = j - 1$.
- $\forall j > \hat{j}. \mathbf{t}_5[\mathbf{i}'_{r \mapsto j}] \leq j - 1$. A similar reasoning applies except that, this time, the known value is for index $k = \hat{j} < j$: $\mathbf{t}_4[\mathbf{i}'_{r \mapsto \xi(k, j)}] = 0$. This implies that the upper bound to the sum is $\sum_{k < j, k \neq \hat{j}} 1 \times 1 + 0 \times 1 + \sum_{k > j} 0 \times (-1) = j - 1$.

Properties of $\mathbf{t}_7 = \text{relu}(\text{add}(\mathbf{t}_5, \mathbf{t}_b))$ The operations between $\mathbf{t}_5$ and $\mathbf{t}_7$ are such that, for any $j \in [\text{dim}_\ell]$, the following holds:

$$\mathbf{t}_7[\mathbf{i}_{r \mapsto j}] = \text{relu}(\mathbf{t}_5[\mathbf{i}'_{r \mapsto j}] + 1 - j).$$

Thus, we can show:

- $\mathbf{t}_7[\mathbf{i}'_{r \mapsto \hat{j}}] = \text{relu}(\hat{j} + 1 - \hat{j}) = 1$.
- $\forall j \in [\text{dim}_\ell]. j \neq \hat{j} \implies \mathbf{t}_7[\mathbf{i}'_{r \mapsto j}] = 0$. This is because the value before *relu* is at most $(j - 1) + 1 - j = 0$.

Properties of $\mathbf{t}_8 = \text{matmul}(\mathbf{t}_7, \mathbf{t}_w)$ The *matmul* operation is such that

$$\mathbf{t}_8[\mathbf{i}'_{r \mapsto 0}] = \sum_{j \in \text{dom}(\text{dim}_\ell)} j \times \mathbf{t}_7[\mathbf{i}'_{r \mapsto j}].$$

Since all $\mathbf{t}_7[\mathbf{i}'_{r \mapsto j}]$ equal 0 except for $j = \hat{j}$ (for which the value is 1), the sum is $\mathbf{t}_8[\mathbf{i}'_{r \mapsto 0}] = \hat{j}$.

Properties of $\mathbf{t}' = \mathbf{t}_9 = \text{transpose}_{p_{\ell,r}}(\mathbf{t}_8)$ This is a simple transpose: $\mathbf{t}_9[\mathbf{i}_{\ell \rightarrow 0}] = \hat{j}$.

This concludes our proof of correctness for the transformation of the ARGMAX operator.

4 Application: Confidence-Based Properties for Provers

We demonstrate the benefits of our work on the verification of robustness of NN.

4.1 Confidence-based global robustness

Robustness is one of the most studied properties of NN. It specifies that a small change to the input should not modify dramatically the output of the network. When testing this property globally (i.e., over the complete range of possible inputs rather than on a few selected points), this property is never satisfied for classification NNs: in the region close to the frontier of decision, tiny changes to the input *will* modify the output. Thus, we are interested in verifying the property only when the decision is associated with a certain confidence.

The notion of *confidence-based safety properties* has been introduced by Athavale et al. [4]. In such a property, the *confidence* refers to the confidence level of the prediction made by the NN. The idea is to add a constraint to the original safety property of interest: robustness is only required when the confidence score of the input is above a given threshold. From this general concept, confidence-based properties can be derived to confidence-based global robustness, which allows some non-robust behaviour for low-confidence inputs. Its formal definition is given in Definition 1.

Definition 1 (Confidence-based global robustness). *Let $f : \mathbb{R}^m \rightarrow \mathbb{R}^n$ be a NN over tensors of shape (m) , $\mathcal{X} \subseteq \mathbb{R}^m$ be a subspace of input tensors, and $\varepsilon > 0$ and $\kappa \in (0, 1]$ be two constants. NN f is globally ε -robust for confidence κ over $\mathcal{X}$ if and only if it satisfies*

$$\begin{aligned} \forall \mathbf{x}, \mathbf{x}' \in \mathcal{X}. \quad & \bigwedge_{j \in [m]} |\mathbf{x}[j] - \mathbf{x}'[j]| \leq \varepsilon \wedge \text{CONF}(f(\mathbf{x})) > \kappa \\ & \Rightarrow \text{CLASS}(f(\mathbf{x})) = \text{CLASS}(f(\mathbf{x}')) \end{aligned}$$

where $\text{CONF}(f(\mathbf{x})) \stackrel{\text{def}}{=} \max_{j \in [n]} (\text{softmax}(f(\mathbf{x}))[j])$ and $\text{CLASS}(f(\mathbf{x})) \stackrel{\text{def}}{=} \text{argmax}(f(\mathbf{x}))$.

We note that the softmax operator in Def. 1 makes the problem very hard to verify. We now discuss two implementations of confidence-based global robustness based on the way the class equality is checked.

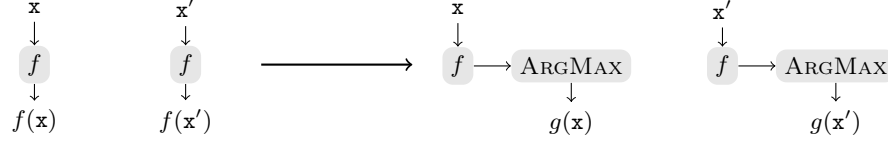


Fig. 2: g is obtained by adding an ARGMAX operator after f

Logical Implementation A practical implementation of confidence-based global robustness, LOGICAL, is available in CAISAR [2].

To check if the predicted class of two vectors $\mathbf{x}, \mathbf{x}' \in \mathbb{R}^m$ is the same with respect to f , one needs to verify if $\text{argmax}(f(\mathbf{x})) = \text{argmax}(f(\mathbf{x}'))$. Checking this condition using first-order logic requires the identification of the label of the class with the highest score. Given a label $l \in [n]$, $\text{CLASS}(f(\mathbf{x})) = l$ implies that for all $l' \in [n]$, we have $f(\mathbf{x})[l] \geq f(\mathbf{x})[l']$. Thus, the constraint $\text{CLASS}(f(\mathbf{x})) = l$ can be encoded using a predicate $\text{pred}_{f,\mathbf{x},l}$ defined such that

$$\text{pred}_{f,\mathbf{x},l} \stackrel{\text{def}}{=} \forall l' \in [n]. f(\mathbf{x})[l] \geq f(\mathbf{x})[l'].$$

We note that this implementation ignores the tie situations, which is a very minor issue. Therefore, checking $\text{CLASS}(f(\mathbf{x})) = \text{CLASS}(f(\mathbf{x}'))$ can be reduced to checking the satisfiability of the predicate

$$\text{eq}_{f,\mathbf{x},\mathbf{x}'} \stackrel{\text{def}}{=} \forall l \in [n]. \text{pred}_{f,\mathbf{x},l} \Leftrightarrow \text{pred}_{f,\mathbf{x}',l}.$$

The constraint $\neg \text{eq}_{f,\mathbf{x},\mathbf{x}'}$ is then directly passed to the solver (PyRAT). The logical implementation is therefore applicable because the *argmax* operation is not part of the input NN but is, instead, a constraint related to our specific *Robustness* problem.

ARGMAX operators It is also possible to verify the class equality of two executions of f by using the ARGMAX operator, ARGMAX. Figure 2 represents the two original executions that are compared and their reformulation using $g \stackrel{\text{def}}{=} \text{ARGMAX} \circ f$, obtained by adding an ARGMAX operator after f .

In this setting, checking the class equality between $f(\mathbf{x})$ and $f(\mathbf{x}')$ is reduced to checking $g(\mathbf{x}) = g(\mathbf{x}')$. This predicate is lighter than $\text{eq}_{f,\mathbf{x},\mathbf{x}'}$. In practice, to obtain the graph presented in Fig. 2, we use CAISAR. CAISAR edits the original computation graph by adding the ARGMAX operator at the end. Since the operator ARGMAX is unsupported by NN verification solvers, CAISAR transforms it with the reformulation given in Fig. 1.

4.2 Experimental evaluation

To assess the relevance of this reformulation, we experimentally compared the two implementations, LOGICAL and ARGMAX, with PyRAT.¹ We note that this is not a direct comparison between using a solver that supports ARGMAX and using our reformulation. The goal of this section is to check whether the use of transformations is a viable method to handle the lack of support for some operators. To this end, we use several instances of neural network verification benchmarks.

COMPAS is a neural network that, based on a selection of features, assesses the probability that criminal defendants will commit future crimes; here, we use a feedforward neural network with 53 parameters. *Adult* is a neural network that estimates, using a person’s information, whether this person has an income above \$50,000 per year; here, we use a feedforward neural network with 34 parameters. *Industrial* is a neural network used in an undisclosed industrial setting; here, we use a feedforward neural network with 1630 parameters.

For these experiments, all instances are set with $\varepsilon = 0.001$, only the confidence threshold varies. The time taken for each instance is the mean across all executions of the same instance (10 executions) with the prover PyRAT. Note that in this study we do not focus on the efficient verification of confidence-based global robustness, but rather on performance differences between two implementations for class equality.

From Table 2, we observe that the ARGMAX implementation is significantly faster on both *COMPAS* and *Industrial*, but it is not for *Adult*, where the verification is slightly slower. It can be argued that, for a neural network with fewer than three output classes, the ARGMAX encoding is heavier to verify than the logical formula that is small (its size grows with the number of output classes). But with more output classes, performance difference is observed in favour of the ARGMAX implementation. Therefore, using the ARGMAX approach yields a faster verification for neural networks with more than two output classes.

This experimentation highlights two positive aspects of using the ARGMAX representation. First, it speeds up verification for neural networks with more than two output classes. Second, it helps to have an easier representation of the verification problem: the first-order logic encoding is replaced by an equality test on the output of two neural networks, as shown in Figure 2.

5 Conclusion

In this paper, we propose a transformation of the ARGMAX ONNX operator into an equivalent sequence of simpler ones. While this could be used in any context, we are particularly interested in using this for verification of neural networks where solvers generally support a restricted set of operators that excludes ARGMAX. We implemented this transformation in the CAISAR platform and

¹ Attempts to use other solvers ‘out-of-the-box’ failed, meaning that these solvers would require careful parameter-tuning to solve these problems in reasonable time.

		Validity	LOGICAL	ARGMAX	Time difference	
<i>COMPAS</i>	($\kappa = 0.9$)	✓	17.4s	12.3s	−5.1s	(−29.3%)
<i>COMPAS</i>	($\kappa = 0.8$)	✓	18.8s	13.5s	−5.3s	(−28.2%)
<i>COMPAS</i>	($\kappa = 0.7$)	✓	38.3s	30.0s	−8.3s	(−21.7%)
<i>COMPAS</i>	($\kappa = 0.6$)	✓	564.6s	493.7s	−70.9s	(−12.6%)
<i>COMPAS</i>	($\kappa = 0.2$)	✗	63.8s	34.2s	−29.6s	(−46.4%)
<i>COMPAS</i>	($\kappa = 0.1$)	✗	64.0s	34.4s	−29.6s	(−46.2%)
<i>Adult</i>	($\kappa = 0.9$)	✓	13.4s	13.5s	+0.1s	(+0.7%)
<i>Adult</i>	($\kappa = 0.8$)	✓	21.8s	23.0s	+1.2s	(+5.5%)
<i>Adult</i>	($\kappa = 0.7$)	✓	123.2s	138.4s	+15.2s	(+12.3%)
<i>Adult</i>	($\kappa = 0.2$)	✗	78.2s	86.7s	+8.5s	(+10.9%)
<i>Adult</i>	($\kappa = 0.1$)	✗	78.6s	86.8s	+8.2s	(+10.4%)
<i>Industrial</i>	($\kappa = 0.999$)	✓	87.8s	61.9s	−25.9s	(−29.5%)
<i>Industrial</i>	($\kappa = 0.99$)	✓	274.6s	188.1s	−86.5s	(−31.5%)
<i>Industrial</i>	($\kappa = 0.9$)	✓	2058.9s	1389.5s	−669.4s	(−32.5%)
<i>Industrial</i>	($\kappa = 0.2$)	✗	115.9s	68.8s	−47.1s	(−40.6%)
<i>Industrial</i>	($\kappa = 0.1$)	✗	116.2s	68.6s	−47.6s	(−41.0%)

Table 2: Runtime comparison on three benchmarks for different confidence thresholds
✓ stands for “Valid”, ✗ for “Invalid”

showed that it not only allows the use of solvers that do not support this operator but also yields better performance than an ad-hoc reformulation for the problem of confidence-based global robustness.

This type of contribution means that developers of dedicated solvers can now concentrate on the interesting aspects of their technology. Having access to this transformation means that they do not need to explicitly support ARGMAX, which implies less work and simpler, less error-prone, and easier-to-maintain code on their side. We recognise however that our implementation, being a generic one, could lead to a performance drop depending on the inner workings of the solver. For instance, an ARGMAX operator allows a solver to immediately infer the range $[\text{dim}_\ell]$ of values of the output network; inferring the same range for the network from Figure 1 requires non-trivial reasoning.

Our goal is to expand this work and to provide additional transformations. For instance, ARGMAX automatically gives us ARGMIN for free. The SOFTMAX operator can be easily derived from the EXP one. We want to provide the largest possible array of transformations to broaden the class of NN verification problems that can be solved by existing solvers. Ideally, we would like to define a minimal subset of supported operators and a partially-ordered set from “higher level” operators to the minimally-supported subset.

Another interesting extension is to propose multiple transformations, when possible, and study which one is the most efficient for which solvers. For instance,

the currently proposed transformation requires to apply a SIGN operator over a tensor whose size is quadratic in the input (this is because the transformation computes the difference between every pair of elements in the vector). A different transformation could require only a linear number of SIGN operations. This can be done by implementing a single-elimination procedure comprised of a log number of rounds in each of which half the elements are discarded. It is unclear whether this implementation would lead to better performance but there is an argument to be made that the SIGN operation is a critical path in the search.

A Variants of *argmax*

The specification of *argmax* includes two options that we have not yet discussed.

The first option indicates whether the dimension upon which the *argmax* is performed is removed altogether. This means that the output shape is not

$$(\dim_1, \dots, \dim_{\ell-1}, 1, \dim_{\ell+1}, \dim_r)$$

but

$$(\dim_1, \dots, \dim_{\ell-1}, \dim_{\ell+1}, \dim_r).$$

This can be implemented adding a RESHAPE operator that allows one to access the list of atomic elements in the tensor differently. For instance, a tensor of shape $(2, 3, 4)$ with 24 elements could be interpreted as a tensor of shape $(1, 6, 1, 4)$. In this case, because the reshape occurs on a dimension of size 1, the reshaping is trivial.

The second option pertains to the tie breaking rule between elements of identical value. When the *select last index* option is set, the tie breaking rule favours the element at highest index instead of the lowest index.

This interpretation can be achieved by modifying the constant tensors involved in the transformation. For instance, tensor $\mathbf{t}_v$ is defined so that the difference $v[k] - v[j]$ is computed, where $k > j$. The reason for this order is this operation is followed by a *sign* and a *relu*: the case $v[k] = v[j]$ is treated similarly as the case $v[k] < v[j]$. When the tie breaking rule is reversed, we want to compute $v[j] - v[k]$ instead. Thus, the values -1 and 1 are swapped in $\mathbf{t}_u$ (and in $\mathbf{t}_v$).

This implies that the values of $\mathbf{t}_b$ also need updating. The value associated with $\hat{j}$ in $\mathbf{t}_5$ is now $(\dim_\ell - \hat{j} - 1)$. The bias associated with j is therefore $2 - \dim_\ell + j$, so that applying this bias associates $\hat{j}$ in $\mathbf{t}_6$ to 1.

References

1. Alberti, M., Bobot, F., Chihani, Z., Grastien, A., Girard-Satabin, J.: The CAISAR platform: extending the reach of Machine Learning specification and verification. In: International Conference on Integrated Formal Methods (IFM). pp. 290–309 (2025)

2. Ardouin, G., Alberti, M., Girard-Satabin, J.: La confiance avec le contrôle : spécification et vérification d'hyperpropriétés sur réseaux de neurones. In: Journées francophones des langages applicatifs (JFLA). pp. 1–21 (2026)
3. Arora, J., Lu, S., Jain, D., Xu, T., Houshmand, F., Phothilimthana, P.M., Lesani, M., Narayanan, P., Murthy, K.S., Bodik, R., Sabne, A., Mendis, C.: TensorRight: automated verification of tensor graph rewrites. In: Proceedings of the ACM on Programming Languages. vol. 9, pp. 832–863 (2025)
4. Athavale, A., Bartocci, E., Christakis, M., Maffei, M., Nickovic, D., Weissenbacher, G.: Verifying global two-safety properties in neural networks with confidence. In: International Conference on Computer Aided Verification (CAV). pp. 329–351 (2024)
5. Bak, S.: nenum: verification of ReLU neural networks with optimized abstraction refinement. In: NASA Formal Methods Symposium (NFM). pp. 19–36
6. Brix, C., Bak, S., Johnson, T.T., Wu, H.: The fifth international verification of neural networks competition (VNN-COMP 2024): summary and results (2024)
7. Community: ONNX: Open Neural Network eXchange Format standard for machine learning interoperability (2022)
8. Demarchi, S., Guidotti, D., Pulina, L., Tacchella, A.: Supporting standardization of neural networks verification with VNN-LIB and CoCoNet. In: Workshop on Formal Methods for ML-Enabled Autonomous Systems (FoMLAS). pp. 47–58 (2023)
9. Katz, G., Huang, D.A., Ibeling, D., Julian, K., Lazarus, C., Lim, R., Shah, P., Thakoor, S., Wu, H., Zeljić, A., Dill, D.L., Kochenderfer, M.J., Barrett, C.: The Marabou framework for verification and analysis of deep neural networks. In: International Conference on Computer Aided Verification (CAV). pp. 443–452 (2019)
10. Lemesle, A., Lehmann, J., Gall, T.L., Chihani, Z.: Neural Network Verification with PyRAT. In: International Static Analysis Symposium (SAS). pp. 11–33 (2025)
11. Shriver, D., Elbaum, S., Dwyer, M.B.: DNNV: a framework for deep neural network verification. In: Computer Aided Verification. pp. 137–150. Lecture Notes in Computer Science (2021)
12. Wang, S., Zhang, H., Xu, K., Lin, X., Jana, S., Hsieh, C.J., Kolter, J.Z.: Beta-CROWN: efficient bound propagation with per-neuron split constraints for complete and incomplete neural network robustness verification. In: Advances in Neural Information Processing Systems (NeurIPS). pp. 29909–29921 (2021)
13. Wu, H., Isac, O., Zeljić, A., Tagomori, T., Daggitt, M., Kokke, W., Refaeli, I., Amir, G., Julian, K., Bassan, S., Huang, P., Lahav, O., Wu, M., Zhang, M., Komen-dantskaya, E., Katz, G., Barrett, C.: Marabou 2.0: a versatile formal analyzer of neural networks. In: International Conference on Computer Aided Verification (CAV). pp. 249–264 (2024)