



Learning to Split: A Reinforcement-Learning-Guided Splitting Heuristic for Neural Network Verification (Invited Talk)

Maya Swisa and Guy Katz^(✉)

The Hebrew University of Jerusalem, Jerusalem, Israel
{maya.swisa,g.katz}@mail.huji.ac.il

Abstract. Modern neural network verifiers often encode neural network verification as constraint satisfaction problems. When dealing with standard piecewise-linear activation functions, such as ReLUs, verifiers typically employ branching heuristics that break a complex constraint satisfaction problem into multiple, simpler problems. The verifier’s performance depends heavily on the order in which this branching is performed: a poor selection may give rise to exponentially many sub-problems, hampering scalability. Here, we focus on the setting in which many related verification queries must be solved for the same neural network. The core idea is to use past experience to make good branching decisions, expediting verification. We present a reinforcement-learning-based branching heuristic that achieves this, by applying Deep Q-learning from Demonstrations (DQfD). Our experimental evaluation demonstrates a substantial reduction in average verification time and in the average number of iterations required, compared to modern splitting heuristics. These results highlight the great potential of reinforcement learning in the context of neural network verification.

1 Introduction

Deep neural networks (DNNs) have emerged as highly powerful tools, which are increasingly being integrated into complex systems in the fields of computer vision [8], natural language processing [18], and aerospace [15]. However, DNNs are also prone to mistakes: for example, they will often fail when presented with mildly perturbed (“adversarial”) inputs [3, 7, 22]. The growing reliance on DNNs, combined with their susceptibility to error, highlights the need for formal guarantees regarding their behavior; and indeed, great strides have been made on DNN verification in the past decade to try and meet this goal.

Despite these advancements, a substantial disparity remains between the scale of DNNs amenable to current verification techniques and the complexity of those required for real-world applications [1]. This is primarily due to the fact that DNN verification with piecewise-linear activations is NP-complete [13],

indicating that the computational cost can grow exponentially with the size of the network. This intractability necessitates the further development of sophisticated and efficient verification techniques.

Many modern verifiers operate by translating the verification problem into the Satisfiability Modulo Theories (SMT) setting (alongside other techniques, such as symbolic bound propagation [1, 24]). This approach provides sound and complete guarantees for a DNN’s correctness, or a counterexample demonstrating the undesirable behavior. While these methods guarantee a definitive answer, they rely on *case splitting*—the translation of piecewise-linear constraints to disjunctions of linear constraints—to handle the piecewise-linear activation functions of the DNN. This process transforms a complex problem into a long sequence of simpler problems, and is known to cause scalability issues [13]. It is known that the selection of neurons on which a solver splits can have a tremendous effect on the size of the search space that is eventually explored; and so multiple splitting heuristics have been proposed with the goal of minimizing this search space (e.g., Pseudo-Impact [26], polarity [25], and BaBSR [2]).

A fundamental limitation of existing approaches is that a single, static heuristic is typically chosen and applied throughout the entire verification process. This method overlooks three key points. First, it is very difficult to know, a priori, which heuristic will be the most useful for the problem at hand; second, it may be useful to change heuristics during the verification procedure, as different parts of the search space are traversed; and third, in many applications, a verification query is actually not stand-alone, but part of a set of related queries—and in that case, it can be useful to transfer knowledge from one run of the verifier to the next, potentially reducing the overall verification time.

To address these limitations, we propose a novel approach: we use Double DQN with demonstrations (DQfD) to learn a state-conditioned, adaptive splitting policy. Rather than fixing a single heuristic a priori, we train an agent that selects the next split according to an estimated action value at the current verifier state. This enables state-dependent branching decisions, allows the preferred split to change as the search evolves, and exploits experience from related verification queries through demonstrations.

To evaluate our approach, we implemented it on top of the Marabou solver [14], and trained and evaluated the learned policy on the ACAS Xu benchmark family [13]. We considered both the standard ACAS Xu safety properties [11] and local robustness queries, requiring the predicted advisory to remain unchanged within an ℓ_∞ ball around a given input. Compared with Marabou’s built-in splitting heuristics, our learned policy solved more instances and reduced average verification time, with improvements most evident on harder instances. Across the evaluated settings, the average time reduction ranged from 5.88% to 56.20%, highlighting the potential of learned branching for DNN verification.

The rest of the paper is organized as follows. Section 2 provides background on deep neural networks, their verification, and reinforcement learning. In Sect. 3, we present our methodology for learning adaptive splitting policies. Section 4 reports on implementation details and experimental evaluation on safety and

robustness benchmarks, followed by Sect. 5 where we discuss related work. Finally, Sect. 6 concludes and outlines directions for future work.

2 Preliminaries

2.1 Neural Networks and their Verification

Deep neural networks (DNNs) are composed of layers of interconnected neurons. The output of neuron j in layer l is $a_j^{(l)} = \sigma(W_j^{(l)} a^{(l-1)} + b_j^{(l)})$, where σ is the activation function. We focus on the ReLU activation function, $\text{ReLU}(x) = \max(0, x)$. For each ReLU neuron n_i , we write n_i^b and n_i^a for its pre- and post-activation values, respectively, so that $n_i^a = \max(0, n_i^b)$. This notation separates the network’s linear layers from its nonlinear parts.

Running Example. Consider the network in Fig. 1, with inputs $x_1 \in [-1, 1]$ and $x_2 \in [0, 1]$, two hidden ReLU units n_1, n_2 with pre-activations n_1^b, n_2^b and post-activations n_1^a, n_2^a , and a single output y . All bias values are zero.

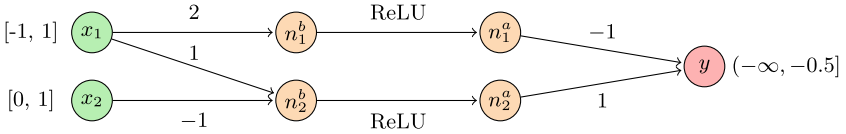


Fig. 1. A toy neural network with two inputs.

For this network, $n_1^b = 2x_1$, $n_2^b = x_1 - x_2$, $n_1^a = \text{ReLU}(n_1^b)$, $n_2^a = \text{ReLU}(n_2^b)$, and $f(x_1, x_2) = -n_1^a + n_2^a = -\text{ReLU}(2x_1) + \text{ReLU}(x_1 - x_2)$. For input $(x_1, x_2) = (0.3, 0.7)$, the neuron values are: $n_1^a = \text{ReLU}(0.6) = 0.6$, $n_2^a = \text{ReLU}(-0.4) = 0$, and $y = -0.6 \leq -0.5$.

Neural Network Verification. We formulate DNN verification as a satisfiability problem. Let $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ be a neural network, $X \subseteq \mathbb{R}^n$ an input domain, and ϕ a property over outputs that encodes an *undesirable* behavior (e.g., a violation of a safety constraint). The verification problem is to determine whether there exists $x \in X$ such that $\phi(f(x))$ holds. If so, the network exhibits the undesirable behavior, and the query is said to be satisfiable (SAT); otherwise, the network is safe with respect to ϕ , and the query is unsatisfiable (UNSAT) [4].

Example 1. Consider the network in Fig. 1, and let $X := (x_1, x_2) \in [-1, 1] \times [0, 1]$. Consider the property $\phi(f(x_1, x_2)) \equiv (f(x_1, x_2) \leq -0.5)$. The verification problem is to determine whether $\exists (x_1, x_2) \in [-1, 1] \times [0, 1] : f(x_1, x_2) \leq -0.5$ holds. For the specific input $(x_1, x_2) = (0.3, 0.7)$, the network evaluates to $f(0.3, 0.7) = -0.6 \leq -0.5$. Thus, $(0.3, 0.7)$ is a witness to satisfiability, and the query is SAT.

2.2 Verification via Branch-and-Bound (BaB)

DNNs with ReLU activations mix easy-to-solve, linear constraints (affine layers) with hard-to-solve non-linearities (ReLU) [13]. An *activation pattern* is an assignment of active/inactive status to all ReLUs. Once an activation pattern is fixed, the network reduces to a single affine transformation, and checking feasibility becomes easily solvable. However, a network with N ReLUs induces up to 2^N activation patterns, rendering exhaustive enumeration infeasible.

Many verifiers adopt a branch-and-bound (BaB) search to tackle the verification problem [6]. BaB organizes the problem as a binary *splitting tree*: the root encodes the original verification problem; each node’s children correspond to a case split of one ReLU into its two phases; and each leaf corresponds to an activation pattern whose satisfiability has been decided. The process terminates when a SAT sub-problem is found (the original problem is also SAT); or when all paths in the search tree are determined to end in an UNSAT node.

In order to prune the search tree, BaB-based solvers usually apply *bound tightening*: they compute sound bounds on pre-activation values of the various ReLUs, and use them to identify ReLUs that are always active or always inactive for the current sub-problem. Whenever such a ReLU is discovered, there is no need to branch on it, and this expedites verification significantly.

Another important aspect in BaB-based techniques is the *prioritization of splits*: choosing which ReLU to split on, using the information the verifier has obtained so far on the current sub-problem (e.g., the currently known bounds). Effective prioritization directs the search quickly toward proofs or counter-examples, reducing the explored portion of the tree.

Running Example. We return to the toy network from Fig. 1 and the query

$$\exists(x_1, x_2) \in [-1, 1] \times [0, 1] : f(x_1, x_2) \leq -0.5.$$

Suppose the verifier first branches on n_1 , considering the case $n_1^b \leq 0$. Since $n_1^b = 2x_1$, this implies $x_1 \leq 0$. Because $x_2 \in [0, 1]$, it follows that $n_2^b = x_1 - x_2 \leq 0$, and so n_2 must also be inactive. Therefore $n_1^a = n_2^a = 0$, which yields $y = 0$. This contradicts the property $y \leq -0.5$, and thus the entire branch $n_1^b \leq 0$ is pruned without any further splitting.

Now, consider the case where the verifier first branches on n_2 . In the branch $n_2^b \leq 0$, we obtain only the relation $x_1 \leq x_2$; in the branch $n_2^b \geq 0$, we obtain only $x_1 \geq x_2$. In either case, the phase of n_1 is not determined. Thus, unlike in the previous case, the verifier must continue branching. This illustrates how a good branching decision can lead to a smaller search space.

2.3 Splitting Heuristics

A good heuristic prioritizes splits that are likely either to quickly lead to UNSAT by pruning large, infeasible subtrees, or to uncover a satisfying assignment early. Modern solvers employ several such heuristics, including Sum of Infeasibilities (SoI), which prioritizes nodes whose assignments in the current, relaxed solution

significantly violate the corresponding ReLU constraint [26]; the Pseudo-Impact heuristic, which tracks the historical impact of branching on each ReLU [26]; Polarity, which favors ReLUs whose pre-activation bound interval is more balanced around zero [25]; and BaBSR, which scores candidate splits by their estimated contribution to bound tightening [2].

2.4 Reinforcement Learning

Reinforcement Learning (RL) is a method for producing an autonomous agent that can make sequential decisions and interact with an environment, with the goal of maximizing a cumulative reward. The core components of an RL system are the agent, the environment, states, actions, and rewards [20]. In value-based RL, the objective is to learn an action-value function that estimates the long-term utility of taking a given action in a given state.

Double DQN [23] is a value-based method that improves training stability by decoupling action selection from action evaluation through a target network. Training typically follows an ϵ -greedy exploration strategy, whereby the agent selects the action with the highest estimated value with probability $1 - \epsilon$, and explores alternative actions with probability ϵ .

Deep Q-learning from Demonstrations (DQfD) [9] augments Q-learning with expert trajectories. The objective combines a value-learning term with a large-margin supervised term that encourages imitation of expert actions, enabling training to start from demonstration data rather than pure exploration.

3 Methodology

3.1 Adaptive Splitting Heuristic

While existing splitting heuristics can significantly reduce the complexity of DNN verification queries, they operate under a fundamental limitation: they are static policies, in the sense that a single heuristic is typically chosen at the beginning of the verification procedure, and is then applied throughout the search process. The performance of these heuristics is highly sensitive to the specific network architecture and the property being verified. However, there is often no a-priori way of knowing which heuristic will be most effective for a given task, and a choice that works well in one scenario may fail miserably in another.

We argue that it is preferable to consider a splitting strategy that is not static, but which can instead adapt as the search progresses, choosing nodes and split variables based on the current state. Many practical settings involve repeated queries on the *same* network (often over related input regions or properties): e.g., a lane-detection model verified under different conditions [5, 11, 16]. Treating each query in isolation discards valuable information about the network’s structure and the efficacy of past branching decisions. By learning from this cumulative data, we can tap into information that was previously unavailable. An adaptive heuristic can observe the success or failure of past splitting choices and use this feedback to inform future decisions, effectively transferring knowledge from one

query to the next. To achieve these goals, we propose to leverage reinforcement learning to train an agent that implements a dynamic branching heuristic.

3.2 Formulation

MDP Formulation. We frame the ReLU splitting process within a Markov Decision Process (MDP), which allows an RL agent to learn an optimal policy for making sequential branching decisions. Intuitively, the agent’s task is to learn a policy that maps the current state to the most effective action, aiming to either quickly prune an infeasible branch or find a counter-example. The MDP in our problem is defined by the tuple $(\mathcal{S}, \mathcal{A}, P, R, \gamma)$:

- **State $\mathcal{S}$:** A state $s \in \mathcal{S}$ represents a specific node in the verification search tree, together with a numerical summary of the current subproblem. The state is given by a feature vector containing both local and global information. The local features, for each candidate ReLU, include its lower and upper pre-activation bounds, phase status, and heuristic scores such as SoI, polarity, and BaBSR. The global features summarize the search, including the total number of unfixed ReLUs, the current search tree depth, and the number of splits performed so far. These features capture information commonly used by splitting-based verification algorithms.
- **Action $\mathcal{A}$:** An action $a \in \mathcal{A}$ corresponds to selecting an unfixed ReLU and choosing one of its two branching phases (*active* or *inactive*). Thus, at each search node, the agent scores candidate (ReLU, phase) pairs, and the selected action determines the next branching decision.
- **Transition Function $P(s'|s, a)$:** The transition from state s to the next state s' after taking action a is determined by the underlying verifier. This includes bound propagation, relaxation tightening, and constraint feasibility/validity checks. The next state s' is the updated subproblem.
- **Reward (penalty) function $R(s, a)$:** We empirically observed that setting a uniform -1 reward per step, which is a common choice in RL, did not guide the search towards effective splits well enough—presumably due to the large search space. We therefore use a task-aligned reward function: each split action on neuron n incurs a delayed, normalized penalty proportional to the size of the subtree that action induces. Let $k(n)$ be the number of unfixed neurons before action a splits on neuron n . The potential subtree size, $\text{full}(n) = 2^{k(n)} - 1$, upper-bounds the number of internal nodes in a full binary tree with up to $2^{k(n)}$ leaves. Let $\text{actual}(n)$ denote the number of internal splits (actions) performed within the subtree rooted at n until that subtree’s search is complete (either pruned, or fully explored). We assign a single, delayed penalty when the subtree closes: $R(s, a) = -\frac{\text{actual}(n)}{\text{full}(n)}$. This normalization makes penalties comparable across neurons with different remaining depth: actions that trigger early pruning yield small (near-zero) penalties, whereas actions that force exhaustive exploration yield penalties near -1 .
- **Discount Factor $\gamma \in [0, 1]$:** Controls the agent’s sensitivity to future rewards. A value closer to 1 prioritizes long-term rewards, encouraging the

agent to consider each split’s future impact. We empirically set $\gamma = 0.9$, utilizing this effect while avoiding the instability associated with higher discount factors.

Leveraging Expert Demonstrations. Training a DRL agent from scratch proved highly sample-inefficient: with sparse effective policies and a complex environment, purely exploratory learning from the available state features incurred prohibitively high computational cost and yielded no end-to-end gains. To mitigate this, we leveraged existing splitting heuristics as *expert demonstrations*. Although sub-optimal, these heuristics encode useful structure; using them to guide the policy via imitation pretraining followed by RL fine-tuning provides a viable path to faster, more reliable learning than uninformed exploration.

Q-Learning Integration. We solve the MDP using Double DQN with demonstrations (DQfD) [9]. Training proceeds in two phases. First, the agent is pre-trained on expert transitions generated by running the verifier with the heuristics described in Sect. 2 (although additional domain- or verifier-specific heuristics could be used as well). This learning-from-demonstrations phase optimizes a combined objective consisting of a Q-learning loss and a margin-based imitation loss, allowing the agent to learn from the outcomes of expert actions and to mimic them. Second, the agent is fine-tuned by interacting with the environment, while continuing to learn from both expert- and self-generated transitions. Over time, the weight of the imitation term is gradually reduced, allowing the policy to move beyond imitation and discover improved splitting strategies. During evaluation, the trained policy is deployed directly within the solver’s branching loop, where it replaces existing heuristics.

Hyperparameters and Trade-Offs. Training involves balancing learning from expert demonstrations with exploration driven by the ε -greedy strategy. Greater weight on demonstrations or the imitation loss yields a stronger initial policy, but may limit the discovery of improved splitting decisions. We use a standard ε -decay schedule, starting at 1.0 and multiplying by 0.95 each iteration until reaching a minimum of 0.05.

As in most reinforcement-learning settings, there is also a trade-off between training time and final performance; accordingly, all reported experiments use a fixed hyperparameter configuration selected in preliminary experiments and applied uniformly throughout the evaluation.

Initial Splits: Pseudo-impact. The Pseudo-Impact heuristic employs a specialized strategy for the initial splits. We found this strategy to be efficient for networks with relatively low input dimensions. We applied a common initial branching policy across all strategies to ensure a fair comparison and to isolate the contribution of our learned heuristic to the deeper, more complex portions

of the search tree. For all experiments presented in this paper, we applied the Pseudo-Impact heuristic for the first three splits of the search tree (depth ≤ 3), with the respective splitting policies taking over from that point onward.

4 Results

4.1 Implementation

To implement our learned branching policy, we integrated a DQfD agent into the SMT-based verifier *Marabou* [14]. As fixed heuristic baselines, we used *Polarity*, *Pseudo-Impact*, and *BaBSR* (with the first splits always performed using Pseudo-Impact, as previously discussed). All of our code and experiments are available online [21].

We evaluated our implementation on the ACAS-Xu networks for airborne collision avoidance [11], under two verification setups:

- (i) **Safety specifications.** The original specifications used in the ACAS-Xu benchmarks, denoted ϕ_1 - ϕ_4 , across all 45 ACAS-Xu networks [12], resulting in a total of 180 queries (4 properties $\times$ 45 networks).
- (ii) **Local robustness.** ℓ_∞ local-robustness verification queries, generated for randomly selected inputs on ACAS-Xu network $N_{1,1}$. For an input x_0 , a local robustness query verifies that the network’s predicted advisory remains unchanged throughout the ball $\|x - x_0\|_\infty \leq \delta$ around x_0 . We tested 1000 randomly selected inputs with three radius values: $\delta \in \{0.08, 0.09, 0.1\}$, producing a total of 3000 queries. The varying inputs and robustness radii produced a mix of satisfiable, unsatisfiable, and timed-out instances, allowing for a more informative comparison between strategies.

We trained a separate agent for each of the two setups. Training included a DQfD warm start followed by Double DQN and used approximately 5% of the queries available in the respective benchmark; and evaluation was then performed on the benchmark’s remaining queries.

The training budget was fixed a priori to 45,000 splitting steps: 5 epochs $\times$ 1,000 steps with demonstrations (DQfD), followed by 40 epochs $\times$ 1,000 steps with self-exploration/exploitation (DDQN). Each agent required about 2 h of offline training. All experiments were conducted on single-CPU machines with 2 GB of memory, running Debian 12, and with a 1-hour timeout.

4.2 Performance

For each splitting policy, we report the number of queries where a counterexample was found (*SATs*), unsatisfiable queries (*UNSATs*), and runs that exceeded the time limit (*Timeouts*). We also report the average time per instance (*Avg Time*) and the average number of solver main-loop iterations (*Avg Iterations*). Both averages are computed over all queries. For runs that exceeded the time limit, we use the timeout threshold as the runtime.

Setup (i): ACAS Xu Safety Properties. On the original ACAS Xu properties, our learned splitting heuristic outperforms all other heuristics in the number of solved instances and average verification time. The aggregate results appear in Table 1, demonstrating the effectiveness of our method.

Table 1. ACAS-Xu runs across all networks: averages over all instances, with timeouts counted as 1 h.

Heuristic	SAT	UNSAT	Timeout	Avg. Time (ms)	Avg. Iter.
DQfD Agent	42.00	105.00	23.00	692691.04	838986.76
Polarity	40.00	103.00	27.00	771287.04	795486.79
Pseudo-Impact	23.00	84.00	63.00	1347021.98	1183586.92
BaBSR	41.00	104.00	25.00	775690.24	779218.49

Setup (ii): Local-Robustness Properties. Table 2 summarizes the results for setup (ii). The learned splitting heuristic achieves the best average verification time, although it requires a slightly greater number of iterations compared to the Polarity heuristic. Moreover, it solves a significantly higher number of instances compared to the baselines.

Table 2. Summary of robustness runs for common (ExampleID, Epsilon).

Heuristic	SAT	UNSAT	Timeout	Avg. Time (ms)	Avg. Iter.
DQfD Agent	253.00	2595.00	122.00	328457.31	235690.92
Polarity	252.00	2574.00	144.00	348988.84	222903.82
Pseudo-Impact	88.00	2319.00	563.00	749976.29	732681.33
BaBSR	244.00	2550.00	176.00	376507.73	265502.52

Figure 2 demonstrates that the learned splitting heuristic closely tracks the strongest baseline on easier instances—indicating an effective imitation of expert behavior. On harder instances, however, it surpasses the baseline, suggesting more effective splitting choices.

5 Related Work

Exact Verification via SMT and BaB. Complete solvers repeatedly perform case splitting on piecewise-linear activations and propagate bounds to prune the search space [13, 14]. Recent work strengthens this paradigm through tighter

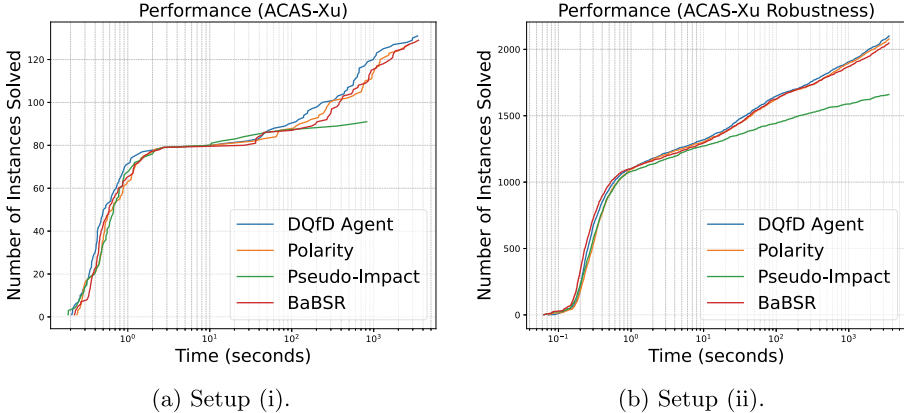


Fig. 2. Cumulative instances solved vs. time (log scale) on the two setups.

relaxations and improved branching scores [2]. Our contribution is complementary: we learn a splitting policy that can be inserted into a BaB framework or an SMT-based solver, replacing hand-crafted splitting heuristics.

Learning for Branching and Verification. Learning-guided branching has been explored in several related settings. In neural network verification, Jeong et al. [10] use a graph neural network to guide *node selection*. Our work is distinct: we use DQfD to learn the *branching rule* itself, directly affecting how subproblems are created and, consequently, the effectiveness of bound tightening and pruning. Qu et al. [19] apply reinforcement learning to branching in general MILP queries; however, DNN verification is structurally different, because branching must interact with specialized nonlinear relaxations and sound bound-tightening procedures. Liang et al. [17] propose an online adaptive branching heuristic for CDCL SAT solvers; while similarly motivated, their setting involves Boolean branching with online feedback, whereas our policy is learned offline from demonstrations and transferred across related verification queries.

6 Conclusion and Future Work

We introduced a novel reinforcement-learning-guided splitting heuristic for DNN verifiers. Our method integrates learning from demonstrations with Double DQN inside an SMT-based verifier, and improves verification performance speed.

Several directions remain for future work, aimed at making learned branching a more reusable and scalable. First, the current framework can be improved by enriching the MDP representation. Additional state features and alternative reward signals may better inform the agent’s decisions and help optimize overall verification performance. Second, our evaluation focuses on ACAS Xu and feed-forward networks with ReLU activations. Assessing the broader applicability of the approach will require experiments on a wider range of benchmarks and

network architectures. Finally, an important goal is to improve the portability of learned splitting policies across networks and properties. This includes studying transfer across networks with different input sizes and architectural variations. Progress in this direction could significantly reduce the cost of applying learned branching to new verification tasks.

Acknowledgments. This work was partially funded by the European Union (ERC, VeriDeL, 101112713). Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Research Council Executive Agency. Neither the European Union nor the granting authority can be held responsible for them. This research was additionally supported by a grant from the Israeli Science Foundation (grant number 558/24).

Data Availability Statement. The code and experiments reported in this paper are available online [21].

Disclosure of Interests. The authors have no competing interests to declare that are relevant to the content of this article.

References

1. Brix, C., Bak, S., Johnson, T., Wu, H.: The Fifth International Verification of Neural Networks Competition (VNN-COMP 2024): Summary and Results. Technical report (2024). <https://arxiv.org/abs/2412.19985>. <https://doi.org/10.48550/arXiv.2412.19985>
2. Bunel, R., Lu, J., Turkaslan, I., Torr, P.H.S., Kohli, P., Kumar, M.P.: Branch and bound for piecewise linear neural network verification. *J. Mach. Learn. Res.* **21**(42), 1–39 (2020)
3. Carlini, N., Wagner, D.: Towards evaluating the robustness of neural networks. In: Proceedings of IEEE Symposium on Security and Privacy (S&P), pp. 39–57 (2017)
4. Casadio, M., et al.: Neural network robustness as a verification property: a principled case study. In: Proceedings of 34th International Conference on Computer Aided Verification (CAV), pp. 219–231 (2022). https://doi.org/10.1007/978-3-031-13185-1_11
5. Elboher, Y., et al.: Robustness assessment of a runway object classifier for safe aircraft taxiing. In: Proceedings of 43rd Digital Avionics Systems Conference (DASC), pp. 1–6 (2024)
6. Ferrari, C., Müller, M.N., Jovanovic, N., Vechev, M.: Complete Verification via Multi-Neuron Relaxation Guided Branch-and-Bound. Technical report (2022). <https://arxiv.org/abs/2205.00263>. <https://doi.org/10.48550/arXiv.2205.00263>
7. Goodfellow, I., Shlens, J., Szegedy, C.: Explaining and harnessing adversarial examples. In: Proceedings of International Conference on Learning Representations (ICLR) (2015)
8. He, K., Zhang, X., Ren, S., Sun, J.: Deep residual learning for image recognition. In: Proceedings of IEEE Conference on Computer Vision and Pattern Recognition (CVPR), pp. 770–778 (2016)
9. Hester, T., et al.: Deep Q-learning from demonstrations. In: Proceedings of AAAI Conference on Artificial Intelligence, vol. 32, no. 1 (2018)

10. Jaeckle, F., Lu, J., Kumar, M.P.: Neural Network Branch-and-Bound for Neural Network Verification. Technical report (2021). <https://arxiv.org/abs/2107.12855>. <https://doi.org/10.48550/arXiv.2107.12855>
11. Julian, K., Kochenderfer, M., Owen, M.: Deep neural network compression for aircraft collision avoidance systems. *J. Guid. Control. Dyn.* **42**(3), 598–608 (2019)
12. Katz, G., Barrett, C., Dill, D., Julian, K., Kochenderfer, M.: Reluplex: a calculus for reasoning about deep neural networks. *Formal Methods Syst. Des. (FMSD)* (2021). <https://doi.org/10.1007/s10703-021-00363-7>
13. Katz, G., Barrett, C., Dill, D.L., Julian, K., Kochenderfer, M.J.: Reluplex: an efficient SMT solver for verifying deep neural networks. In: Majumdar, R., Kunčák, V. (eds.) *CAV 2017. LNCS*, vol. 10426, pp. 97–117. Springer, Cham (2017). https://doi.org/10.1007/978-3-319-63387-9_5
14. Katz, G., et al.: The marabou framework for verification and analysis of deep neural networks. In: Dillig, I., Tasiran, S. (eds.) *CAV 2019. LNCS*, vol. 11561, pp. 443–452. Springer, Cham (2019). https://doi.org/10.1007/978-3-030-25540-4_26
15. Katz, G., Levy, N., Refaeli, I., Yerushalmi, R.: DEM: a method for certifying deep neural network classifier outputs in aerospace. In: *Proceedings of 43rd Digital Avionics Systems Conference (DASC)*, pp. 1–8 (2024)
16. Kessler, C., et al.: Neural network verification for gliding drone control: a case study. In: *Proceedings of International Symposium on AI Verification (SAIV)* (2025)
17. Liang, J., Ganesh, V., Poupart, P., Czarnecki, K.: Learning rate based branching heuristic for SAT solvers. In: *Proceedings of 19th International Conference on Theory and Applications of Satisfiability Testing (SAT)*, pp. 123–140 (2016)
18. OpenAI. ChatGPT (2022). <https://chatgpt.com>
19. Qu, Q., et al.: An Improved Reinforcement Learning Algorithm for Learning to Branch. Technical report (2022). <https://arxiv.org/abs/2201.06213>. <https://doi.org/10.48550/arXiv.2201.06213>
20. Sutton, R.S., Barto, A.G.: *Reinforcement Learning: An Introduction*. MIT Press (2018)
21. Swisa, M., Katz, G.: Learning to Split: A Reinforcement-Learning-Guided Splitting Heuristic for Neural Network Verification (Code) (2025). <https://github.com/mayaswissa/LearningToSplit>
22. Szegedy, C., et al.: Intriguing Properties of Neural Networks. Technical report (2013). <https://arxiv.org/abs/1312.6199>. <https://doi.org/10.48550/arXiv.1312.6199>
23. van Hasselt, H., Guez, A., Silver, D.: Deep reinforcement learning with double Q-learning. In: *Proceedings of AAAI Conference on Artificial Intelligence*, vol. 30, no. 1, pp. 2094–2100 (2016)
24. Wang, S., et al.: Beta-CROWN: efficient bound propagation with per-neuron split constraints for neural network robustness verification. In: *Proceedings of 35th Conference on Neural Information Processing Systems (NeurIPS)* (2021)
25. Wu, H., et al.: Parallelization techniques for verifying neural networks. In: *Proceedings of Formal Methods in Computer-Aided Design (FMCAD)*, pp. 128–137 (2020). <https://doi.org/10.34727/2020/isbn.978-3-85448-042-6.20>
26. Wu, H., Zeljić, A., Katz, G., Barrett, C.: Efficient neural network analysis with sum-of-infeasibilities. In: *TACAS 2022. LNCS*, vol. 13243, pp. 143–163. Springer, Cham (2022). https://doi.org/10.1007/978-3-030-99524-9_8

Open Access This chapter is licensed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

