



Universiteit
Leiden
The Netherlands

A new challenge: approaching Tetris Link with AI

Müller-Brockhausen, M.F.T.; Preuss, M.; Plaat, A.

Citation

Müller-Brockhausen, M. F. T., Preuss, M., & Plaat, A. (2021). A new challenge: approaching Tetris Link with AI. *2021 Ieee Conference On Games (Cog)*.
doi:10.1109/CoG52621.2021.9619044

Version: Publisher's Version

License: [Licensed under Article 25fa Copyright Act/Law \(Amendment Taverne\)](#)

Downloaded from: <https://hdl.handle.net/1887/3249627>

Note: To cite this publication please use the final published version (if applicable).

A New Challenge: Approaching Tetris Link with AI

Matthias Müller-Brockhausen, Mike Preuss, Aske Plaat

Leiden Institute of Advanced Computer Science

Leiden University

The Netherlands

m.f.t.muller-brockhausen@liacs.leidenuniv.nl

Abstract—Decades of research have been invested in making computer programs for playing games such as Chess and Go. This paper introduces a board game, Tetris Link, that is yet unexplored and appears to be highly challenging. Tetris Link has a large branching factor and lines of play that can be very deceptive, that search has a hard time uncovering. Finding good moves is very difficult for a computer player, our experiments show. We explore heuristic planning and two other approaches: Reinforcement Learning and Monte Carlo tree search. Curiously, a naive heuristic approach that is fueled by expert knowledge is still stronger than the planning and learning approaches. We, therefore, presume that Tetris Link is more difficult than expected. We offer our findings to the community as a challenge to improve upon.

Index Terms—Tetris Link, Heuristics, Monte Carlo tree search, Reinforcement Learning, RL Environment, OpenAI Gym

I. INTRODUCTION

Board games are favorite among AI researchers for experiments with intelligent decision-making and planning. For example, works that analyze the game of Chess date back centuries [1], [2]. Already in 1826, papers were published on machines that supposedly played Chess automatically [3], although it was unclear whether the machine was still operated somehow by humans. Nowadays, for some games, such as Chess [4] and Go [5], [6], we know for sure that there are algorithms that can, without the help of humans, automatically decide on a move and even outplay the best human players.

In this paper, we want to investigate a new challenge, Tetris Link, that has not yet received attention from researchers before, to the best of our knowledge (see Section II). Tetris Link is a manual, multiplayer version of the well-known video game Tetris. It is played on a vertical “board”, like Connect-4. The game has a large branching factor, and since it is not immediately obvious how a strong computer program should be designed, we put ourselves to this task in this paper. For that, we implement a digital version of the board game and take a brief look at the game’s theoretic aspects (Section III). Based on that theory, we develop a heuristic that we also test against human players (Section IV-A). Performance is limited, and we try other common AI approaches: Deep Reinforcement Learning (RL) [7] and Monte Carlo tree search (MCTS) [8]. In Subsection V-A, we look at MCTS agents, and in Subsection V-B, we look at RL agents and their performance in the game. In our design of the game environment for the RL agent, we

assess the impact of choices such as the reward on training success. Finally, we compare the performance of these agents after letting them compete against each other in a tournament (Section V-C). To our surprise, humans are stronger.

The main contribution of this paper is that we present to the community the challenge of implementing a well-playing computer program for Tetris Link. This challenge is much harder than expected, and we provide evidence (Section III-B) on why this might be the case, even for the deterministic 2-player version (without dice) of the game. The real Tetris Link can be played with four players using dice, which will presumably be even harder for an AI.

We document our approach, implementing three players based on the three main AI game-playing approaches of heuristic planning, Monte Carlo Tree Search, and Deep Reinforcement Learning. To our surprise and regret, all players were handily beaten by human players.¹ We respectfully offer our approach, code, and experience to the community to improve upon.

II. RELATED WORK

Few papers on Tetris Link exist in the literature. A single paper describes an experiment using Tetris Link [9]. This work is about teaching undergraduates “business decisions” using the game Tetris Link. To provide background on the game, we analyze the game in more depth in Section III. The authors are aware of a similar game called *Blokus* [10]. It is an interesting game as the branching factor is so large (≈ 32928 for *Blokus Duo* on a 14×14 board) that specialized FPGA’s have been applied to build good AI for it [11], [12]. Although visually similar, the gameplay is completely different, so *Blokus* strategies or heuristics do not transfer to Tetris Link.

The AI approaches that we try have been successfully applied to a variety of board games [13]. Heuristic planning has been the standard approach in many games such as Hex [14], Othello [15], Checkers [16], and Chess [4], [17]. MCTS has been used in a variety of applications such as Go and General Game Playing (GGP) [8], [18], [19]. Deep RL has seen great success in Backgammon [20], Go [5], and Othello [21]. Multi-agent MCTS has been presented in [22].

¹Humans only played against the Heuristic, not MCTS or DRL. The Heuristic is our strongest AI as can be seen in Figure 6.



Fig. 1: A photo of the original Tetris Link board game. The colored indicators on the side of the board help to keep track of the score.

III. TETRIS LINK

Tetris Link, depicted in Figure 1, is a turn-based board game for two to four players. Just as the original Tetris video game, Tetris Link features a ten-by-twenty grid in which shapes called tetrominoes² are placed on a board. This paper will refer to tetrominoes as blocks for brevity. The five available block shapes are referred to as: *I*, *O*, *T*, *S*, *L*.³ Every shape has a small white dot, also in the original physical board game variant, to make it easier to distinguish individual blocks from each other. Every player is assigned a color for distinction and gets twenty-five blocks: five of each shape. In every turn, a player must place precisely one block. A block fits if all of its parts are contained within the ten-by-twenty board. A player is skipped if they are unable to fit any of their remaining blocks. A player can never voluntarily skip if one of the available blocks fits somewhere in the board even if placing it is disadvantageous. The game ends when no block of any player fits into the board anymore.

The goal of the game is to obtain the most points. One point is awarded for every block, provided that it is connected to a group of at least three blocks. Not every block has to touch every other block in the group, as shown in Figure 2a.

The *I* block only touches the *T* but not the *L* on the far right. Since they together form a chained group of three, it counts as three points. Blocks have to touch each other edge-to-edge. In Figure 2b, the red player receives no points as the *I* is only connected edge-to-edge to the blue *L*.

A player loses one point per empty square (or hole) below block that was placed, with a maximum of two minus points per turn. Figure 2c shows how one minus point for red would look like. Moreover, the Figure underlines a fundamental difference to video game Tetris. In video game Tetris, blocks slowly fall, and one could nudge the transparent *L* under the *S*

²A shape built from squares that touch each other edge-to-edge is called a polyomino [23]. Because they are made out of precisely four squares, these shapes are called tetromino [24].

³The *S* and *L* blocks may also be referred to as *Z* [25] and *J* [26].

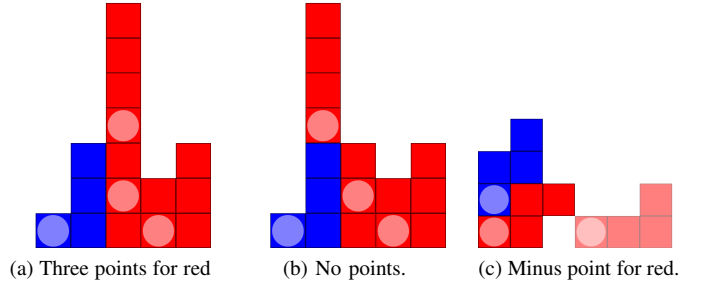


Fig. 2: Small examples to explain the game point system.

to fill the hole by precise timing of an action. In Tetris Link, one can only throw blocks into the top and let them fall straight to the bottom. In the original rules, a dice is rolled to determine which block is placed. If a player is out of a specific block, then the player gets skipped. Since every block could turn into one point, being skipped means potentially missing out on it. Although not a guaranteed disadvantage, as the opponent might also be skipped, the authors have abandoned the dice in their own matches as it resulted in too many games that felt unfairly lost.

The dice roll also makes Tetris Link a non-deterministic perfect information game. In this paper, there is no dice roll, so we analyze the deterministic version of Tetris Link. Note that we also focus on the two-player game only in this work. The three- and four-player versions are presumably even harder. In multiplayer games without teams, people might temporarily team up against the current leading player, which creates an unfair disadvantage [10]. Nevertheless, our web-based implementation for human test games⁴ can handle up to four players and can provide an impression of the Tetris Link gameplay.

A. Verification that all games can fill the board

Each game of Tetris Link can be played to the end, in the sense that there are enough blocks to fill the whole board without leaving any empty squares. This is easy to see by the following argument: The board is ten squares wide and twenty squares high, so it can accommodate 200 individual squares. Every player has twenty-five blocks, each consisting of four squares. There are always at least two players playing the game, so they are always able to fill the board.

$$\begin{aligned} & \text{playerBlocks} * \text{squaresPerBlock} * \text{playerAmount} \\ &= 25 * 4 * 2 = 200 \end{aligned} \quad (1)$$

B. Game complexity

An essential metric for search algorithms is the so-called branching factor, which specifies the average number of possible moves a player can perform in one state [27]. In order to compute this number, we look at the number of orientations for each block. The *I* block has two orientations as it can be

⁴<https://hizoul.github.io/contetro>

used either horizontally or vertically. The O block has only one orientation because it is a square. The T block has four different orientations for every side one can turn it to. The L and S are special cases as they can be mirrored. The L has four and the S two sides one can rotate it to. Mirroring them doubles the amount of possible orientations. Hence, in total, nineteen different shapes can be placed by rotating or mirroring the available five base blocks. Since the board is ten units wide, there are also ten different drop points per shape. In total, there can be up to 190 possible moves available in one turn. However, the game rules state that all placed squares have to be within the bounds of the game board. Twenty-eight of these moves are always impossible because they would necessitate some squares to exceed the bounds either on the left or right side of the board. Therefore, the exact number of maximum possible moves in one turn for Tetris Link is 162. Since the board gets fuller throughout the game, not all moves are always possible, and the branching factor decreases towards the end. In order to show this development throughout matches, we simulate 10,000 games.

We depict the average number of moves per turn in Figure 3. For the first eight to ten turns, all moves are available on average. Not depicted in the Figure but based on the data, this only holds true until turn 6. After that, the number of plays may already decrease. Tetris Link is a game of skill: random moves perform badly. A game consisting of random moves ends after only thirty turns. Many holes with many minus points are created, and the game ends quickly with a low score. The heuristic bars show that simple rules of thumb fill the board most of the time by taking more than forty turns. Furthermore, the branching factor in the midgame (turn 13-30) declines slower and hence offers more variety to the outcomes. Another thing that can be seen in Figure 3 is that only very few select plays can actually lead to positive points. Most of the turns will leave the score unchanged or even decrease it. A player would only consider taking minus points if it would cost the opponent even more points than the player loses, or there are no other options. To further underline this, we did an exhaustive tree search until depth 5, which is the first turn player one is able to achieve points. Of the over 111 Billion possibilities (111577100832), only $\approx 14.36\%$ allow for an increase in points. Moreover, only $\approx 3.06\%$ allows the optimal three-point gain, which the first player aims for. $\approx 69.08\%$ of the outcomes result in a point disadvantage, and in $\approx 16.55\%$ of the cases, no points have been gained nor lost due to blockage by the opponent or lack of connectivity.

We are now ready to calculate the approximate size of the game tree complexity for the deterministic variant of Tetris Link, in order to compare it to other games. On average, across all three agents shown in Figure 3, a game takes 37 turns and allows for 74 moves per turn ($74^{37} \approx 1.45 * 10^{69}$). The game tree complexity is similar to Chess (10^{123}) [28] but smaller than in Go (10^{360}) [29].

Agent	Random	Random-H	User-H	Tuned-H
Win Rate	48.16%	47%	71.65%	70%
Unique Games	10,000	10,000	7	50

TABLE I: First move advantage, over 10,000 games. All turns are compared for uniqueness to see whether the same games keep repeating. The -H in the agent name stands for Heuristic (see Section IV-A). Draws are not counted towards wins.

C. First move advantage

An important property of turn-based games is whether making the first move gives the player an advantage [30]. To put this into numbers, we let different strategies play against themselves 10,000 times to determine if the starting player has an advantage. All moves are recorded and checked for uniqueness.

As can be seen in Table I, the win rate for *random heuristic* as starting player is almost 50%. Although the win rate for the first player is higher for the *tuned heuristics*, these numbers are not as representative because the heuristic repeats the same tactics resulting in only seven or twenty-nine unique game starts. If we repeat the same few games, then we will not truly know whether the first player has a definite advantage. Especially considering that at least until turn six, all moves are always possible, there are around 10^{13} or 18 Trillion ($BranchingFactor^{Turns} = 162^6 = 18,075,490,334,784$) possible outcomes. Since the *random heuristic* has more deviation and plays properly as opposed to random moves, we believe that it is a good indicator of the actual first player advantage. Note that 47% is close to an equal opportunity. Different match history comparisons of Chess measure a difference of around two to five percent in win rate for the first player [30]. However, since neither Tetris Link nor Chess have been mathematically solved, one cannot be certain that there is a definite advantage.

IV. AI PLAYER DESIGN

In this section, we describe the three different types of AI players that we implemented, based on heuristics, MCTS, and RL, respectively. For the experiments (Section V), the game is coded in Rust and JavaScript (JS). The Rust version is written for faster experiments with MCTS and RL, and the JavaScript version is written to visually analyze games and also do a human play experiment. Both implementations share a common game log format using JSON in order to enable interoperability. To underline the importance of a performance-optimized version, we measured the average time it takes to simulate one single match where always the first legal move is made. The Rust implementation requires $590\mu s$ for that, whereas the JavaScript implementation needs 82ms.

A. Heuristic

We now describe the design of our heuristic player. A heuristic is a rule of thumb that works well most of the time [31]. For Tetris Link, we identify four heuristic measures:

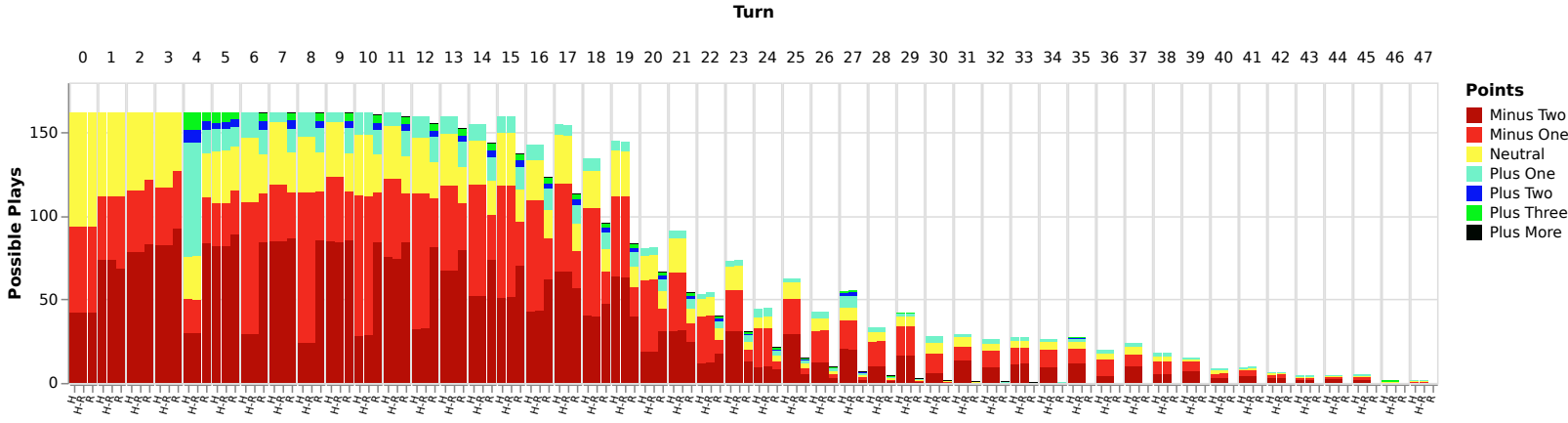


Fig. 3: This graph underlines the game’s difficulty, especially for search algorithms, by showing the average number of possible moves and their associated effect on the player’s score. The agent name abbreviations stand for: R = Random, H-R = Random Heuristic, and H = User Heuristic. See Section IV-A for an explanation of how the heuristic agents work.

the number of connectable edges, the size of groups, the player score, and the number of blocked edges. The number of blocked edges is the number of edges belonging to opponents that are blocked by the current players’ blocks. From our experience these heuristic values are positively related to the chance of winning.

Each parameter is multiplied by a weight, and the overall heuristic score is the sum of all four weighted values. For every possible move in a given turn, the heuristic value is calculated, and the one with the highest value is chosen. If multiple moves have the same maximum value, a random one of these best moves is chosen, which is why in Table I, the heuristics play more than one single unique repeating game. The initial weights were manually set by letting the heuristic play against itself and detecting which combination would result in the most points gained for both players. We refer to this as *user heuristic*. We then use Optuna [32], a hyperparameter tuner, to tune a set of weights that reliably beat the *user heuristic*. This version is called *tuned heuristic*. To achieve a greater variety in playstyle, we also test a *random heuristic* which at every turn generates four new weights between zero to fifteen. To have an estimate on the performance of the heuristic, we let actual human players ($n=7$) familiar with the game play against the heuristic via the JavaScript implementation (Figure 4). The *random heuristic* achieved a win rate of 23.07% across 13 matches, and the *user heuristic* a win rate of 33.33% across six matches. The sample size is very small, but it still indicates that the heuristic is not particularly strong. This is supported by a qualitative analysis of the game played by the authors, based on our experience. We conclude that our heuristic variants play the game in a meaningful way but are not particularly strong.

B. MCTS

For applications in which no efficient heuristic can be found, MCTS is often used, as it constructs a value function by averaging random roll-outs [8]. Our MCTS implementation uses the standard UCT selection rule [33]. As further

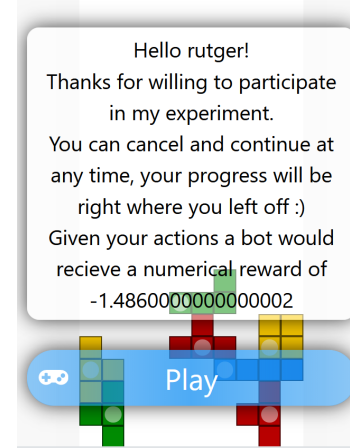


Fig. 4: The web interface to measure the effectiveness of our Heuristic agent against actual human players familiar with the game.

enhancements, we also use MCTS-RAVE [8] and MCTS-PoolRAVE [34] to see whether the modifications help in improving the quality of results. Furthermore, we experimented with improving the default (random) policy by replacing it with the heuristic. However, the heuristic calculation is so slow that it only manages to visit ten nodes per second. We want to stress that this is only the MCTS heuristic that is slow as it was not optimized for speed at all. The game simulation itself is fast processing a full match of random actions in on average $590\mu s$, so up to around 1694 matches per second.

MCTS is well-suited for parallelization, leading to more simulations per second and hence better play [8]. We implemented tree parallelization, a frequently used parallelization [35]. In tree parallel MCTS, many threads expand the same game tree simultaneously. Using 12 threads, we visit 16258 nodes per second on average with a random default policy. To put this into perspective, this is $1.63e^{-9}\%$ of all

10^{13} possibilities in the first six turns. Thus, only a small part of the game tree is explored by MCTS, even with parallel MCTS.

C. Reinforcement Learning Environment and Agent

A reinforcement learning environment requires an observation, actions, a reward [36], and an RL agent with an algorithm as well as a network structure. To prevent reinventing the wheel, we use existing code for RL, namely OpenAI gym [37] and the stable-baselines [38], which are written in Python. To connect Python to our Rust implementation, we compile a native shared library file and interact with it using Pythons *ctypes*. As RL Algorithm, we exclusively use the deep reinforcement learning algorithm PPO2 [39]. For the network structure, we increase the number of hidden layers from two layers of size 64 to three layers of size 128 because increasing the network size decreases the chances of getting stuck in local optima [40]. We do not use a two-headed output, so the network only returns the action probabilities but not the certainty of winning as in AlphaZero [6].

The observation portrays the current state of the game field. Inspired by AlphaGo, which includes as much information as possible (even the “komi”⁵), we add additional information such as the number of blocks left per player, the players’ current score, and which moves are currently legal. For the action space, we use a probability distribution over all moves. The probabilities of illegal moves are set to 0, so only valid moves are considered. For the reward, we have three different options.

- 1) Guided: $\frac{score + groupSize}{100} - scolding$
- 2) Score: $\frac{score}{100}$
- 3) Simple: ± 1 depending on win/loss

The *Guided* reward stands out because it is the only one that reduces the number of points via *scolding*. If the move with the highest probability is illegal, then the reward will be reduced by -0.004, and the first legal move with the highest probability is chosen. This should teach the agent to only make valid moves. This technique is called reward shaping, and its results may vary [41].

In order to detect which one of the three options is the most effective, we conduct an experiment. Per reward function, we collect the averages for the number of steps it took, the average reward achieved, and what the average score of the players was in the results. Our results, shown in Table II, indicate that the *Guided* reward function works best. It only takes around 3183 steps on average to reach a local optimum, and the average scores achieved in the matches are the highest. We define a local optimum as the same match repeating three times in a row. The *Score* reward function also lets the agent reach a local optimum, but it takes twice as long as the *Guided* function, and the score is slightly lower as well. The *simple* reward function seems unfit for training. It never reached a local optimum in the 10,000 steps we allowed it to run, and it got the lowest score in its games.

Reward Type	Steps	Episode Reward	Score
Guided	3183.49	-0.17	-5.6
Simple	10000.0	-0.0	-12.25
Score	6214.45	-0.09	-6.88

TABLE II: Results of self-play with different reward types until either a local optimum or 10,000 steps have been reached. Step, Reward and Score show the average of all seeds.

V. AGENT TRAINING AND COMPARISON

For our experimental analysis, we first look at the performance of the MCTS agent (Section V-A) and the training process of the RL agents (Section V-B). Finally, we compare all previously introduced agents in a tournament to analyze their play quality and determine the currently best playing approach.

A. MCTS Effectiveness

1) *Setup*: Initial test matches of MCTS against the *user heuristic* resulted in a zero percent win rate, and a look at the game boards suggested near-random play. We use a basic version of MCTS with random playouts because using our heuristic as guidance was too slow. AlphaZero has shown that even games with high branching factors such as Go can be played well by MCTS when guided by a neural network [6]. However, without decision support from a learned model or a heuristic, we rely on simulations. In order to see if this guidance is the reason for bad MCTS performance, we abuse the fact that the *user heuristic* plays very predictably (Section III-C). We use the RAVE-MCTS variant (without the POOL addition), pre-fill the RAVE values with 100 games of the *user heuristic* playing against itself, and then let the MCTS play 100 matches against the *user heuristic*. We repeat this three times and use the average value across all three runs. We run this experiment with different RAVE- β parameter values. This parameter is responsible for the exploration/exploitation balancing and replaces the usual UCT C_p parameter. The closer the RAVE visits of a node reach RAVE- β , the smaller the exploration component becomes. Furthermore, we employ the slow heuristic default policy at every node in this experiment. We simulate one match per step because otherwise, the one-second thought time is not enough for the slow heuristic policy to finish the simulation step.

2) *Results*: Our MCTS implementation can play well with a decent win rate against the user-heuristic, as shown in Figure 5. This result underlines that in games with high branching factors, MCTS needs good guidance through the tree in order to perform well. Figure 3 also supports this as there are very few paths that will actually lead to good plays. The declining win rate with a higher RAVE- β value suggests that exploration on an already partially explored game tree worsens the result because the opponent does not deviate from its paths. The rise in win rate for a β value of 5000 after the large drop in 2500 underlines the effect of randomness involved in the search processes.

Even though the heuristic-supported playout policy works well, we will still use a random playout policy for the

⁵Komi refers to the first turn advantage points [5].

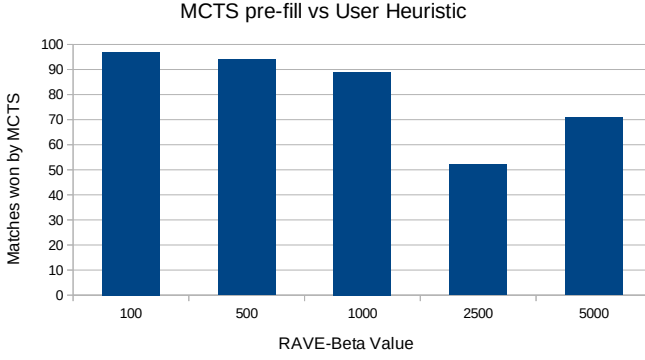


Fig. 5: Win rates of pre-filled MCTS playing against the *User-Heuristic* compared by the RAVE- β parameter.

tournament (Section V-C). Pre-filling the tree is very costly and would provide an unfair advantage to the MCTS method.

To underline that we believe the bad performance stems from the large branching factor and few paths that lead to positive points, we test our implementation on the game of Hex. In different board sizes (2x2 to 11x11) of Hex, our MCTS plays against a shortest path heuristic. The result is striking: as long as the branching factor stays below 49 (7x7), MCTS wins up to 90% of the matches. For larger branching factors, the win rate drops to 0% quickly.

B. RL Agents Training

1) *Agents*: We define an RL agent as the combination of reward type, algorithm, and training opponent. We use the guided reward function because it worked best in our experiment and call this agent *RL-Selfplay*. (This is a neural network only RL, without MCTS to improve training samples, that plays against itself [Selfplay].)

In addition to this rather simple agent, we introduce the *RL-Selfplay-Heuristic* agent. It builds on a trained *RL-Selfplay* agent where we continue training by playing against the heuristic. Observation and reward are the same as for *RL-Selfplay*.

From the first turn advantage experiment, we know that the heuristic plays well even with random weights. That is why we also introduce an agent called *RL-Heuristic*. This agent outputs four numbers that represent the heuristics weights (Section IV-A). We use a modified version of the guided reward function:

$$\frac{(\text{ownScore} - \text{opponentScore}) + \text{groupSize}}{100} \quad (2)$$

Group size stands for the total number of blocks that are connected with at least one other block. This is added because we want the algorithm to draw a connection between the number of points gained and the number of connected blocks. However, mainly the difference in points between itself and the opponent is used as a learning signal, so it aims to gain more points than the opponent. Scolding is not necessary anymore as we do not have to filter the output in any way.

2) *Setup*: In this section, we detail the training process of the RL agents. Each training is done four times, and only the best run is shown. Agents are trained with the default PPO2 hyperparameters, except for *RL-Heuristic*, which uses hand-tuned parameters.

When playing only against themselves, the networks still quickly reached a local optimum even with increased layer size. This optimum manifested in the same game being played on repeat and the reward per episode staying the same. This repetition is a known problem in self-play and can be called “chasing cycles” [42]. To prevent these local optima, we train five different agents against each other in random order. To be able to train against other agents, we modified the *stable-baselines* code.

3) *Results*: The training process for *RL-Selfplay* peaked around 1.5 million steps. For *RL-Selfplay-Heuristic* we use the two best candidates from *RL-Selfplay*, namely #3 after one million steps with a reward of 0.04 and #1 after 1.5 million steps with a reward of 0.034. The training of *RL-Selfplay#1-Heuristic* reaches its peak after 3.44 and *RL-Selfplay#3-Heuristic* after 3.64 Million steps with a reward of 0.032 and 0.024. These are our first RL agents that can achieve a positive reward while playing against the heuristic.

The *RL-Heuristic* training worked well, achieving mostly a positive reward. However, looking at the output values, we realize the reward function design was unfortunate. It sets all weights to zero, except for the enemy block value to fifteen and the number of open edges between four and seven. So by negating the player’s score with the opponent’s score, we have unwillingly forced the heuristic to focus on blocking the opponent over everything else. Needless to say, with these weights, the *RL-Heuristic* rarely wins. Although it manages to keep the opponent’s score low, it does not focus on gaining points which leaves it with a point disadvantage.

C. Tournament

1) *Setup*: In the tournament, all presented AI approaches play against each other. Every bot will face every other bot in 100 matches. We have five different RL bots, three MCTS bots, and three heuristic bots. Each agent gets at most one second of time to decide on their move, and they are not allowed to think during the opponent’s turn. Every bot will play half of its matches as first and the other half as second player. The bot’s skill will be compared via a Bayesian Bradley Terry (BBT) skill rating [43]. The original BBT [43] ratings range from 0 to 50. By adjusting the BBT- β parameter, we change this to represent the ELO range (0 to 3000) [44].

2) *Results*: The final skill rating is portrayed in Figure 6. The three heuristic agents take the top 3, followed by RL and MCTS. Remarkably, the *tuned heuristic* performed best, even though it is only optimized to play well against the *user heuristic*, but yet it performs best across all agents.

Seeing *RL-Heuristic* as the best RL approach shows that the other RL agents are far from playing well. Yet all RL agents consistently beating MCTS with random playouts proves that the agents definitely learned to play reasonably.

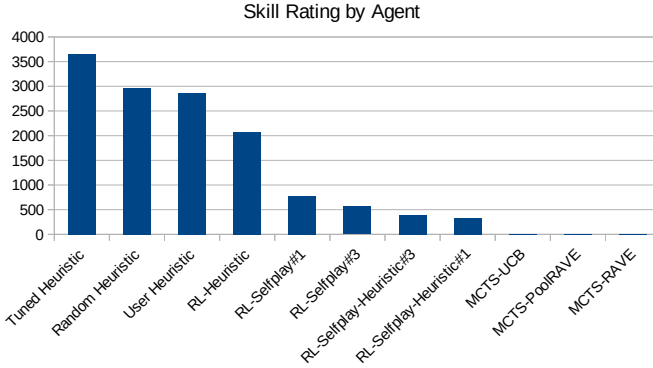


Fig. 6: The skill rating of the agents that participated in the tournament. MCTS uses a non pre-filled tree, resulting in bad performance (Section V-A2).



Fig. 7: Visualisation of the scores that agents achieved in the tournament. Agents are sorted by the skill rating in Fig. 6.

It is interesting to see that the MCTS-UCB (14% win rate) variant performed best because the other two variants [RAVE (0.02%), PoolRAVE (0.04%)] were conceived in order to improve the performance of UCB via slight modifications [34]. Please note that this poor performance in Figures 6 and 7 is caused by the tree not being pre-filled as it was in Figure 5 (Section V-A2). This shows the importance of pre-filling MCTS in Tetris Link due to the few paths that result in positive points (see Figure 5).

The skill rating omits information about the quality of the individual moves. To gain further insight into that, we provide Figure 7. Here, we can see that every agent manages at least once to gain 8 points or more. This means that every agent had at least one match it played well. Looking at the lowest achieved scores and average scores, we find that every agent, except for the pure heuristic ones, plays badly, considering that they only make ± 3 points on average.

VI. CONCLUSION AND FUTURE WORK

Board game strategy analysis has been done for decades, and especially games like Chess and Go have seen countless papers analyzing the game, patterns, and more to find the best

play strategies [6]. We contributed to that field by taking a close look at the board game Tetris Link. While the strategy is key to winning, some games, such as Hex, give the first player a definite advantage. We have experimentally shown that there is no clear advantage for the starting player in Tetris Link (Section III-C).

We have implemented three game-playing programs based on common approaches in AI: heuristic search, MCTS, and reinforcement learning. Despite some effort, none of our rule-based agents was able to beat human players.

In doing so, we have obtained an understanding of why it may be hard to design a good AI for Tetris Link:

- Especially at the beginning, the branching factor is large, staying at 162 for at least the first six turns.
- The majority of possible moves result in minus points (see Figure 3 and Section III-B).
- In our experience, mistakes/minus points can hardly be recovered from. The unforgivingness for these moves may make it harder to come up with a decent strategy, as generally postulated by [45].
- Many rewards in the game stack — they come delayed after multiple appropriate moves because groups of blocks count and not single blocks.

All this holds true for the simplified version we treat here: no dice, only two players. Adding up to two more players and dice will also make the game harder.

With a solid understanding of the game itself, we investigated different approaches for AI agents to play the game, namely heuristic, RL and MCTS. We have shown that all tested approaches can perform well against certain opponents. The best currently known algorithmic approach is the tuned heuristic, although it can not consistently beat human players.

Training an RL agent (Section V-B) for Tetris Link has proven to be complicated. Just getting the network to produce positive rewards required much trial and error, and in the end, the agent did not perform well even when consistently achieving a positive reward. We believe the learning difficulty in Tetris Link comes from the many opportunities to make minus points in the game. One turn offers usually offers one plus to three points, or six at most if two groups are connected, but that means that multiple previous turns that were well planned and gave zero points if not even more minus points had to be made. Hence recovering from minus points is difficult, meaning small mistakes have graver consequences.

Although MCTS performed poorly in our tournament, we have shown that with proper guidance through the tree, MCTS can perform nicely in Tetris Link and Hex (Section V-A). This suggests that MCTS extensions like progressive bias and progressive unpruning are good candidates for future experiments [46]. Alternatively, a combination where RL guides an MCTS through the tree might work well, e.g., AlphaZero [6] or MoHex v3 [47].

As computer game researchers, we found ourselves challenged to create a good agent to play the game Tetris Link. We tried a large variety of classic approaches and were not

able to achieve the results we hoped for. We invite the research community to use our code and improve upon our approaches.⁶

REFERENCES

- [1] F. D. Philidor, *Analysis of the Game of Chess*. P. Elmsly, 1790.
- [2] T. Sprague, "On the different possible non-linear arrangements of eight men on a Chess-board," *Proceedings of the Edinburgh Mathematical Society*, vol. 8, pp. 30–43, 1889.
- [3] G. Bradford, *The History and Analysis of the Supposed Automation Chess Player of M. de Kempelen: Now Exhibiting in this Country, by Mr. Maelzel*. Hilliard, Gray & Company, 1826.
- [4] F.-h. Hsu, "Computer chess, then and now: The Deep Blue saga," in *Proceedings of Technical Papers. International Symposium on VLSI Technology, Systems, and Applications*. IEEE, 1997, pp. 153–156.
- [5] D. Silver, A. Huang, C. J. Maddison, A. Guez, L. Sifre, G. Van Den Driessche, J. Schrittwieser, I. Antonoglou, V. Panneershelvam, M. Lanctot *et al.*, "Mastering the game of Go with deep neural networks and tree search," *Nature*, vol. 529, no. 7587, pp. 484–489, 2016.
- [6] D. Silver, J. Schrittwieser, K. Simonyan, I. Antonoglou, A. Huang, A. Guez, T. Hubert, L. Baker, M. Lai, A. Bolton *et al.*, "Mastering the game of go without human knowledge," *Nature*, vol. 550, no. 7676, pp. 354–359, 2017.
- [7] R. S. Sutton and A. G. Barto, *Reinforcement learning: An introduction. Second Edition*. MIT press, 2018.
- [8] C. B. Browne, E. Powley, D. Whitehouse, S. M. Lucas, P. I. Cowling, P. Rohlfshagen, S. Tavener, D. Perez, S. Samothrakis, and S. Colton, "A survey of monte carlo tree search methods," *IEEE Transactions on Computational Intelligence and AI in games*, vol. 4, no. 1, pp. 1–43, 2012.
- [9] P. Orenstein, "Does experiential learning improve learning outcomes in an undergraduate course in game theory—a preliminary analysis," *Northeast Decision Sciences Institute*, 2014.
- [10] C. Chao, "Blokus Game Solver," in *Computational Engineering Commons*. California Polytechnic State University, 2018.
- [11] A. Jahanshahi, M. K. Taram, and N. Eskandari, "Blokus Duo game on FPGA," in *The 17th CSI International Symposium on Computer Architecture & Digital Systems (CADSD 2013)*. IEEE, 2013, pp. 149–152.
- [12] E. Qasemi, A. Samadi, M. H. Shadmehr, B. Azizian, S. Mozaffari, A. Shirian, and B. Alizadeh, "Highly scalable, shared-memory, Monte-Carlo tree search based Blokus Duo Solver on FPGA," in *2014 International Conference on Field-Programmable Technology (FPT)*. IEEE, 2014, pp. 370–373.
- [13] A. Plaat, *Learning to Play: Reinforcement Learning and Games*. Springer Nature, 2020.
- [14] J. Van Rijswijk, "Search and evaluation in Hex," *Master of science, University of Alberta*, 2002.
- [15] J. Schaeffer, J. van den Herik, and T.-s. Hsu, "Games, computers and artificial intelligence," *Chips Challenging Champions: games, computer and Artificial Intelligence*, pp. 3–9, 2002.
- [16] A. Plaat, J. Schaeffer, W. Pijls, and A. De Bruin, "Best-first fixed-depth minimax algorithms," *Artificial Intelligence*, vol. 87, no. 1-2, pp. 255–293, 1996.
- [17] S. J. Russell and P. Norvig, *Artificial intelligence: a modern approach*. Malaysia; Pearson Education Limited, 2016.
- [18] R. Coulom, "Efficient selectivity and backup operators in Monte-Carlo tree search," in *International conference on computers and games*. Springer, 2006, pp. 72–83.
- [19] B. Ruijl, J. Vermaseren, A. Plaat, and J. v. d. Herik, "Combining simulated annealing and Monte Carlo tree search for expression simplification," *arXiv preprint arXiv:1312.0841*, 2013.
- [20] G. Tesaro, "Neurogammon: A neural network backgammon learning program," *Heuristic Programming in Artificial Intelligence: The First Computer Olympiad, Chichester, England*, 1989.
- [21] N. J. van Eck and M. van Wezel, "Application of reinforcement learning to the game of Othello," *Computers & Operations Research*, vol. 35, no. 6, pp. 1999–2017, 2008.
- [22] E. Galván-López, R. Li, C. Patsakis, S. Clarke, and V. Cahill, "Heuristic-Based Multi-Agent Monte Carlo Tree Search," in *IISA 2014, The 5th International Conference on Information, Intelligence, Systems and Applications*. IEEE, 2014, pp. 177–182.
- [23] E. D. Demaine and M. L. Demaine, "Jigsaw puzzles, edge matching, and polyomino packing: Connections and complexity," *Graphs and Combinatorics*, vol. 23, no. 1, pp. 195–208, 2007.
- [24] Wikipedia. (2019-09-09) Tetromino. [Online]. Available: <https://en.wikipedia.org/wiki/Tetromino>
- [25] H. Burgiel, "How to lose at Tetris," *The Mathematical Gazette*, vol. 81, no. 491, pp. 194–200, 1997.
- [26] D. Carr, "Applying reinforcement learning to Tetris," *Department of Computer Science Rhodes University*, 2005.
- [27] S. Edelkamp and R. E. Korf, "The branching factor of regular search spaces," in *AAAI/IAAI*, 1998, pp. 299–304.
- [28] C. E. Shannon, "Programming a computer for playing chess," in *first presented at the National IRE Convention, March 9, 1949, and also in Claude Elwood Shannon Collected Papers*. IEEE Press, 1949, pp. 637–656.
- [29] H. Matsubara, H. Iida, and R. Grimbergen, "Chess, Shogi, Go, natural developments in game research," *ICCA journal*, vol. 19, no. 2, pp. 103–112, 1996.
- [30] W. Streeter, "Is the first move an advantage," *Chess Review*, p. 16, 1946.
- [31] M. H. Romanycia and F. J. Pelletier, "What is a heuristic?" *Computational Intelligence*, vol. 1, no. 1, pp. 47–58, 1985.
- [32] T. Akiba, S. Sano, T. Yanase, T. Ohta, and M. Koyama, "Optuna: A Next-generation Hyperparameter Optimization Framework," in *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*. ACM, 2019, pp. 2623–2631.
- [33] L. Kocsis and C. Szepesvári, "Bandit based monte-carlo planning," in *European conference on machine learning*. Springer, 2006, pp. 282–293.
- [34] A. Rimmel, F. Teytaud, and O. Teytaud, "Biasing Monte-Carlo simulations through RAVE values," in *International Conference on Computers and Games*. Springer, 2010, pp. 59–68.
- [35] M. Enzenberger and M. Müller, "A lock-free multithreaded Monte-Carlo tree search algorithm," in *Advances in Computer Games*. Springer, 2009, pp. 14–20.
- [36] L. P. Kaelbling, M. L. Littman, and A. W. Moore, "Reinforcement learning: A survey," *Journal of artificial intelligence research*, vol. 4, pp. 237–285, 1996.
- [37] G. Brockman, V. Cheung, L. Pettersson, J. Schneider, J. Schulman, J. Tang, and W. Zaremba, "OpenAI Gym," *arXiv preprint arXiv:1606.01540*, 2016.
- [38] A. Hill, A. Raffin, M. Ernestus, A. Gleave, A. Kanervisto, R. Traore, P. Dhariwal, C. Hesse, O. Klimov, A. Nichol, M. Plappert, A. Radford, J. Schulman, S. Sidor, and Y. Wu, "Stable Baselines," <https://github.com/hill-a/stable-baselines>, 2018.
- [39] J. Schulman, F. Wolski, P. Dhariwal, A. Radford, and O. Klimov, "Proximal policy optimization algorithms," *arXiv preprint arXiv:1707.06347*, 2017.
- [40] S. Lawrence, C. L. Giles, and A. C. Tsoi, "Lessons in neural network training: Overfitting may be harder than expected," in *AAAI/IAAI*. Citeseer, 1997, pp. 540–545.
- [41] M. Grzes and D. Kudenko, "Plan-based reward shaping for reinforcement learning," in *2008 4th International IEEE Conference Intelligent Systems*, vol. 2. IEEE, 2008, pp. 10–22.
- [42] O. Vinyals, I. Babuschkin, W. M. Czarnecki, M. Mathieu, A. Dudzik, J. Chung, D. H. Choi, R. Powell, T. Ewalds, P. Georgiev *et al.*, "Grandmaster level in StarCraft II using multi-agent reinforcement learning," *Nature*, pp. 1–5, 2019.
- [43] R. C. Weng and C.-J. Lin, "A Bayesian Approximation Method for Online Ranking," *Journal of Machine Learning Research*, vol. 12, no. Jan, pp. 267–300, 2011.
- [44] M. E. Glickman and A. C. Jones, "Rating the chess rating system," *CHANCE-BERLIN THEN NEW YORK*, vol. 12, pp. 21–28, 1999.
- [45] M. Cook and A. Raad, "Hyperstate space graphs for automated game analysis," in *2019 IEEE Conference on Games (CoG)*. IEEE, 2019.
- [46] G. M. J. Chaslot, M. H. Winands, H. J. V. D. HERIK, J. W. Uiterwijk, and B. Bouzy, "Progressive strategies for monte-carlo tree search," *New Mathematics and Natural Computation*, vol. 4, no. 03, pp. 343–357, 2008.
- [47] C. Gao, R. Hayward, and M. Müller, "Move prediction using deep convolutional neural networks in Hex," *IEEE Transactions on Games*, vol. 10, no. 4, pp. 336–343, 2017.

⁶<https://github.com/Hizoul/tetris-link-research>