

[English]

In the history of infectious diseases, a single drug has been easily resisted by mutations in the pathogen. In contrast, when multiple drugs are used simultaneously, each helps prevent the emergence of resistance to the others. Ryfamate derives from Rifamate which is the combination of two medicines for tuberculosis to retard the development of drug resistance. As with this method, Ryfamate aims to combine an NNUE alpha-beta search with deep learning (DL) Monte Carlo tree search to compensate for each other's shortcomings.

Last year, Ryfamate used two NNUE functions and one ResNet function; this year, it will use one NNUE, one ResNet, and a new type of DL function, such as a deeper ResNet or Transformer. The new DL function currently has the problem of slow inference that leads to severe consequences in Shogi. However, it is considered to have a high recognition capacity. Therefore, Ryfamate will use the new function to compensate for the weak points of the previous Ryfamate with AdaBoost learning and ensemble inference.

[Library]

YaneuraOu : <https://github.com/yaneurao/YaneuraOu>

dlshogi : <https://github.com/TadaoYamaoka/DeepLearningShogi>

* Ryfamate also uses a large amount of data that Kano-san, Yamaoka-san, and Tayayan-san have published.

WCSC32 W@nderER アピール文書

Dated 2021/12/27

[昨年](#)に [引き続き](#)、入玉宣言による勝利を積極的に目指します。

チーム Barrel house のアピール文書@第 32 回世界コンピュータ将棋選手権

Barrel house とは岡山の駅前にあるビアバーです。元々チームメイトを求めさすらって行きついたところです。マスターに許可頂いたので名前をお借りしました。その後、メンバーが変わって当初の方向性とは全く違って来ましたが、まあ一度出した名前を変更するのもアレなのでそのままです。

プログラム名：白ビール

第 28 回の Hefeweizen が濁った白ビールでしたが、第 29 回では Kristallweizen とフィルターでろ過した透き通った白ビールでした。その後 Hefeweizen に戻すという形を取りました。しかしながら、欧州のチェスサイトでは命名由来から説明頂いているにも関わらず国内では白ビールとしか読んで頂けないので、昨年からもう白ビールでいいやってことです。

チームの特徴

フレッシュなチームを自認してますがどうやらおっさんチームのようです。本大会は体力勝負ではないので各方面に様々な知識や技術・勘などを働かせて今年も決勝に残ればいいかなあと思っています。諸事情でメンバーが多忙なため具体的な策はこれから詳細を詰めていくところです。（前回も前々回もそんなアピール文だった気もしますが）

評価関数

昨年、一昨年の計測でもトップチームと劣らないものであることを確認しました。（連盟モバイルでの採用など詳しくは blog にて）今後の調整等は未定です。

使用マシン

今年もノートパソコンを中継にクラウドの力をお借りする予定です。今年はクラウドの方も今年は若干様子が変わっているようでベンチマーク等も未だなので具体的なことは全く決まっていません。正直あまり予算がないのですが、マシンパワーだけで負けるのもアレなので予算内で一番速いマシンを借りようと考えています。

クラスタリングについて

第 5 回電王トーナメントで御披露目をした shotgun システムの進化版である Multi Ponder のクラスタリングを用います。制作メンバーが本年も外れておりますが、既に実装は文書化され多くの類似実装が選手権でも見られますので改良して用いるのは問題ないと考えます。

Sin-Daigorilla wcsc32 アピール文章

田中大吾

今回のアピール点として大きく 4 つある。

今回は NNUE 関数と dlshogi の合議、もしくはどちらか片方で出場する予定だ。

①モデル及び評価関数

NNUE に関しては、hkpe4 を採用した。Zero ベクトルから学習をスタートさせ、合議のため dl とのパラメーターを考えながら調整した。(例えば NNUE は中終盤型にし、dl は序盤重視みたいな)

やねうら王は version7.00

dl の教師データを利用し更に強化した。

Test 対局では Suisho5/6.50 に対し 100 戦 4 割程度である。

(条件は 80t 1 手 1.2 秒 投了値 500 定跡なし 初手から)

DL の方は 10 ブロック 約 35 億のデータ(psv と公開された hcpe 半分程度の割合)を学習させた。3epoch 程度が適切らしい。

やねうら教師との学習を交互に行うと train_loss と test_loss がめっちゃくちゃになるため、それらの調整も課題になってくる。

現時点では r4300 程度である。

合議 tool の方はメモリを大量に喰うため出来るだけ削減できる方法を探したい。Source については公開されているものを使用している。

②定跡

ほとんど手をつけていないが、後手の勝率が低下が進む中、後手の評価が高い局面から徹底的開始させ先手と後手の思考にノード差を付けることでバランスの取れた定跡を作成出来るのではないかと考えた。

ある程度 kif がとれたらテラショックコマンドで進めていき、ふかうら王で生成したスーパーテラショックや過去の定跡と merge してみた。

しかし効果は不明である。

③探索

やねうら王 7.00 を hkpe4 に対応させハイパーパラメーターの序盤重視率関連を変更した。

使用したソフト

Suisho5

やねうら王

Python dlshogi

[第32回世界コンピュータ将棋選手権]

大將軍

(たいしょうぐん)

アピール文書

横内 健一

横内 靖尚

大將軍の概要

- ◆ 評価関数に主眼を置いた将棋プログラム。
- ◆ 評価関数の特徴としては、盤上の3駒の位置関係を考慮。

過去には4駒の位置関係の評価関数で参加したこともありますが、結局は3駒の位置関係に落ち着きました。

大將軍の概要

- ◆ 過去の遺産をベースとし、評価関数の学習においては、
(いまさらながらではあるが・・・)
 - ◆ 駒の位置関係の相対位置による評価
 - ◆ 手番の学習
- を考慮。
- ◆ 今回は、3駒関係(kppt型)を基本とするが、学習対象とする局面は浅い探索の末端局面と現局面をミックスし学習するよう調整。教師データも3駒系とNNUE系の棋譜をミックスして利用。

使用ライブラリ

- ◆ やねうら王

- ◆ 最新のStockfishの探索ルーチン

- ◆ わかりやすいコード(過去および現在のアイデアを適用しやすい)

- ◆ 水匠2～4

- ◆ 評価関数学習用の棋譜作成(教師データ作成)に利用

やねうら王 PR 文書

— スーパーテラショック定跡生成手法について —

やねうら王チーム

要 旨

近年、Deep Learning 系の将棋ソフトはプロ棋士の研究にも取り入れられてきている。Deep Learning 系の将棋ソフトは従来型の将棋ソフトに比べると序盤戦術に長けていると言える。本論文手法は、その優れた序盤感覚のソフトで大規模な定跡を自動生成しようという試みである。

第 1 章 従来の定跡生成手法

第 1 節 人間の棋譜から

コンピュータ将棋も 2010 年ごろまでは人間が定跡を手で入力していた。特にプロの棋譜は神聖視されており、プロの棋譜に出現した 32 手目までの手順をそのまま定跡としていたソフトも多数存在した。GPS 将棋[2009-]¹などもプロの棋譜の手順をそのまま定跡化しているようであった。

しかしながら、プロの棋譜には 76 歩 34 歩 68 銀 88 角成(角のタダ取り)で投了というような棋譜も混じっており、プロの棋譜をそのまま定跡として採用していた Bonanza[2005-]²がこの手順を指すということで話題になったこともある。

そこで、その棋譜の指し手が指された回数を集計し、指された回数が多かった指し手を指すことでこのような悪手を指さないようにするような試みがあった。Bonanza で言うと narrow book というエンジンオプションである。

あるいはプロの棋譜に対して自作の将棋ソフトで 1 局面 0.1 秒や探索深さ 12 など短い時間(or 深さ)で思考させ、悪い評価値の指し手だけを取り除くというハイブリッドな手法が取られることもあった。習甦³[2014]がそのようなアプローチを取っていた。

プロの棋譜を用いる場合、インターネットで入手できる棋譜の数が限られている(4 万局程度)というのがあるし、トッププロより将棋ソフトの方が棋力が高くなった

現在においては、もはや人間の棋譜は定跡生成には直接は使えないと言っていいだろう。

第 2 節 floodgate の棋譜から

棋譜の数で言う floodgate(コンピュータ将棋ソフトの対局サーバー)⁴の棋譜を用いるというのは魅力的である。floodgate の棋譜の数は 50 万棋譜以上存在するし、その時代時代の最新のソフトが参加しているため、floodgate の棋譜を用いればお手軽かつそれなりの精度で生成できる。2015 年ごろまでわりと流行していた手法である。

floodgate の棋譜から定跡を生成する場合、Bonanza の narrow book の考え方を踏襲し、定跡上の各局面では、なるべく floodgate で採用された回数が多い指し手を選択する。採用回数が多い指し手ほどその信頼性が高いと考えられるからである。Apery[2015]⁵や、やねうら王[2014]⁶の standard_book.db などに見られるようにこの方法で作成された定跡を搭載する将棋ソフトはわりと多い。

しかし採用回数ベースで定跡を生成する場合、棋譜の数が少ないとその局面での採用確率が 100%であっても信用ならないという欠点がある。例えばその局面を含む棋譜が 1 つしかない場合に、その棋譜の指し手を定跡として登録して良いのかという問題がある。

また弱いソフトの指し手が混じっている場合、それがどれほど信用できるのかという問題もある。ひどい例で言うと、定跡を抜けた直後の局面で先手必敗(評価値 -600)ということもあった。

¹ Wikipedia 「GPS 将棋」:

<https://ja.wikipedia.org/wiki/GPS%E5%B0%86%E6%A3%8B>

² Wikipedia 「Bonanza」:

<https://ja.wikipedia.org/wiki/Bonanza>

³ Wikipedia 「習甦」:

<https://ja.wikipedia.org/wiki/%E7%BF%92%E7%94%A6>

⁴ floodgate :

<http://wdoor.c.u-tokyo.ac.jp/shogi/floodgate.html>

⁵ Wikipedia 「Apery」:

<https://ja.wikipedia.org/wiki/Apery>

⁶ Wikipedia 「やねうら王」:

<https://ja.wikipedia.org/wiki/%E3%82%84%E3%81%AD%E3%81%86%E3%82%89%E7%8E%8B>

あと、floodgate での対局は、序盤での持ち時間を節約するために何らかの定跡を用いて対局させることが多いようで、実際の勝率は低いのにわりと採用されている序盤の指し手というのが存在していることが dlshogi の山岡の研究からわかっている。⁷

そこで、floodgate の棋譜を用いる場合、勝率ベースで指し手を選択するという方法も考えられる。その指し手を指したあと、実際に勝っているのかという点を考慮しようというわけである。この方法は、わりとうまく行くのだが、強いソフトが採用している定跡を過度に信用してしまう問題がある。例えば、floodgate には初手 38 銀に固定したソフトや、必ず振り飛車にする定跡を搭載した将棋ソフトが参加していた。おそらくそのような序盤定跡を用いた時に、普通に指すのに比べてどれくらいレートが落ちるのかを見たい開発者がいたのであろう。このような棋譜が混じっている場合、勝率ベースで定跡を生成する場合でも信用ならない指し手が紛れ込む可能性はいくらでもある。

floodgate の棋譜に対しても自作の将棋ソフトで短い時間で思考させ、明らかに悪い評価値の指し手は取り除くというハイブリッドな手法が取られることもあった。習甦の方法に倣い、やねうら王[2015]もそのような方法を用いたことがある。

しかし、floodgate の棋譜上に数回しか出現しなかった局面の指し手は信用ならないので、30 回以上出現した局面についてのみに絞って定跡化しようなどすると、局面数はかなり減ってしまい、定跡として十分な局面数を確保できないという問題があった。

いずれにしても棋力にばらつきのある外部の棋譜に頼る以上、その信頼性を担保するのが難しく、データも無限に用意できるわけではないので定跡の局面数にも自ずと限界があった。

⁷ 山岡忠夫 「floodgate の序盤 3 手の統計」
<https://tadaoyamaoka.hatenablog.com/entry/2021/11/06/010026>

第 3 節 人間+ソフトのハイブリッド

ある程度の棋力を有する人間が、将棋ソフトを使いながら、将棋ソフトで上位にくる指し手の変化をすべて調べ、そしてその定跡の末端の局面から対局させて勝率を計測し、形勢が互角であっても勝率が悪い局面への指し手は取り除くという、気の遠くなるような作業を経て作られた定跡がある。

このような定跡としては suimon_fan 氏が公開されている s-book_black⁸が優秀なことで有名だ。WCSC31 で優勝した elmo[2021]も手作業による定跡であり、いまだ上位ソフトの開発者でもこのような手作業での定跡生成を行っている開発者は少なくはない。

第 4 節 自己対局型自動生成

そこで出てきたのが定跡の自動生成という技術である。2015 年ぐらいから流行りだす。ただ、当時の将棋ソフトは序盤はわりと精度が低かったので、実際に序盤で思考させて、良い評価値の枝を掘っていく(思考させて定跡局面として登録していく)のは、わりと危なかったところがあった。

そこで、自己対局をして勝率の良いところを掘り進めていくという手法が用いられることが多かった。とりわけそのなかでも優秀だったのは shotgun 定跡⁹と呼ばれるもので、SDT5(第 5 回将棋電王トーナメント)で shotgun というソフトが採用し(shotgun は SDT5 準優勝)、その後、WCSC28 で Hefeweizen がそれを受け継いだ。¹⁰

⁸ suimon_fan : https://twitter.com/mztn7_fan

⁹ 芝世式, コンピュータ将棋における定跡生成法の一提案 :
https://ipsj.ixsq.nii.ac.jp/ej/?action=repository_uri&item_id=192055&file_id=1&file_no=1

¹⁰ 芝世式, Hefeweizen WCSC28 PR 文書
<https://www.apply.computer-shogi.org/wcsc28/appeal/Hefeweizen/appeal.pdf>

shotgunの方法は簡単に言ってしまうえば異種ソフトをローカル環境で対局させ、勝率ベースで定跡化するというもので、相手が既知の将棋ソフトの場合、わりと高い確率でその定跡の指し手を指すようである。

Hefeweizenの手法はほぼノーメンテナン스로定跡が生成していけるということで、定跡自動生成系としてはわりと高い評価がなされているが、自己対局ベースなので生成にとても時間がかかることや既知のソフト以外のソフトだと定跡から外れやすいという欠点はあった。

生成のためにとても時間がかかるのでは、定跡に登録できる局面数が稼げないことになる。いくら精度が高く優秀であっても相手より先に定跡が尽きてしまったのでは、大会では相手の方が多くの持ち時間を残したまま終盤に突入してしまい、結果、不利になってしまう。

このように定跡DBはその登録されている指し手の質、精度はもちろんのこと、なるべく実戦で実現しやすい局面が数多く登録されていることも重要である。

第5節 Minimax探索型(テラショック化)

どんな方法で定跡を生成するにせよ、その定跡のleaf node(末端の局面)での評価値が得られているなら、定跡のゲームツリー上でMinimax探索を行えば、先後が定跡上の指し手で最善を尽くした時に到達するleaf nodeが確定する。

やねうら王では、これを**テラショック化**と呼んだ。¹¹やねうら王[2017]には与えた棋譜の局面で思考させ、定跡DB上の各局面に評価値をつけるコマンドをすでに用意していたので、このコマンドを用いて、まずfloodgateの棋譜などを与え、各局面に評価値をつける。その後、その定跡DBをテラショック化することができれば、定

跡ツリー上の各局面からのPV(最善応手列)が得られる。それをその各局面の最善手とした定跡DBファイルを別途出力する。これがテラショック化コマンドの概要である。

このテラショック化コマンドは、やねうら王[2019]で実装され、このテラショック化した定跡を用いてやねうら王はWCSC29で優勝した。

ところが、その大会後に判明したことだが、Minimax探索は末端の評価値がroot(定跡の探索開始局面 = 初期局面)まで伝播してくるので、ノイズに極めて弱いという性質がある。従来の将棋ソフト(非Deep Learning系の将棋ソフト)では、評価値にずいぶんノイズが乗っているようで、テラショック化には向かないという意味があった。実際、やねうら王がWCSC29で用いた定跡と前述のs-book_blackとを対局させてみたところ、定跡を抜けた局面でやねうら王の定跡側が不利な局面であった。

また、上の手法は、与えた棋譜ファイル上でMinimax探索のようなことをやっていることになるわけだが、この時、与えた棋譜の指し手しか先後ともに指さないという暗黙の仮定をしていることになる。これはじゃんけんで言うとグーが封じられているようなもので、その状態で(グー、チョキ、パーすべてを使える相手に)勝てるのかという意味もある。

テラショック化という発想自体はさほど悪くなかったのだと思うのだが、以上のように明らかな欠点も抱えており、従来型のソフトには荷が重いと言わざるを得なかった。

第6節 第一章まとめ

floodgateの棋譜から勝率の高かった指し手のみを選んで定跡化するのはさほど悪くないものの、floodgateの棋譜に出現した指し手しか指せないという意味で、グーしか出せないジャンケンにも似ている。今後、将棋ソ

¹¹ やねうら王ブログ 「テラショック定跡の生成手法」：
<https://yaneuraou.yaneu.com/2019/04/19/tera-shock-book-generation/>

フトの序盤がシビアになってくるとそんな方法は通用しなくなってくるだろう。

結局、現時点で定跡の自動生成手法で成功しているのは shotgun 定跡ぐらいのものだが、実際にローカル対局を行うので費やすリソースのわりに掘れる(登録できる)定跡の局面数に自ずと限りがある。

テラショック化は悪いアイデアではなかったが、従来の将棋ソフトには評価値にムラがある(ノイズが乗っている)ので、意図通りには機能しない。また PV(最善応手列)の周りを重点的に掘りたいのに、事前に与えた棋譜の局面しか掘れず、テラショック化してからその定跡を用いて自己対局させて、その対局棋譜を与えて思考させて、またテラショック化する、というようないイテレーション(反復)でしか定跡を掘っていけないので思考対象局面の選出に非常に手間と時間がかかる。

第2章 スーパーテラショック

定跡生成手法

第1節 Deep Learning 系思考エンジンの利用

従来型の非 Deep Learning 系の将棋ソフトの序盤には評価値にムラがある(ノイズが乗りやすい)、不正確であるという弱点があった。

そこで定跡の生成には Deep Learning 系の将棋ソフトを用いることにした。現在、ソースコードが公開されている Deep Learning 系の将棋ソフトのなかで最強である dlshogi¹²の互換エンジンとしてふかうら王という将棋ソフトを開発した。

今回の定跡生成では、このふかうら王で1局面1秒程度思考させる。Deep Learning 系の将棋ソフト、特にふかうら王では、思考させた局面のすべての合法手に対する評価値が得られる。なので、この得られた評価値をその指し手で1手進めた局面での評価値であると仮定して、有望そうな局面のみを効率的に展開していくことができる。

そのあと任意のタイミングで定跡 DB に対してテラショック化を行い、最終的な定跡 DB を得る。これがスーパーテラショック定跡生成手法の骨子である。

第2節 思考対象局面の選出

どの局面を思考対象として選ぶかという問題がある。外部から棋譜を与えてその局面を思考させるという方法もあるが、それだと与えた棋譜の局面しか思考しない

ので、「自分だけがグーの出せないじゃんけん」のようになってしまう。

テラショック化した定跡を用いて、実際に対局させ、その対局棋譜上に現れた局面について再度思考させるようなイテレーション(反復)も考えられる。やねうら王[2019]ではそうやって定跡を生成していたのだが、テラショック化自体に時間が相当かかるし、ローカル対局にも時間を要する。この方法だとあまり PV まわりを掘り進められない。

そこで、定跡ツリー上で、 $\alpha\beta$ 探索¹³を行い、PV(最善応手列)を取得し、その leaf node(末端の局面)を展開する(思考対象とする)ような仕組みを用意することにした。

これをここでは「(次に思考すべき局面の)選出」と呼ぶことにする。選出された局面をふかうら王の思考エンジンを用いて思考していく。これをここでは「思考」と呼ぶことにする。つまり、「選出」と「思考」。今回の定跡の自動生成はこの二つで成り立っている。

第3節 選出の高速化と思考との並列化

この「思考」中にも「選出」を行わなければ思考エンジンに遊び時間が出来てしまうから、「選出」と「思考」は並行して行える必要がある。また「選出」自体を高速化しないと、思考エンジンに遊び時間ができてしまう。

そのためには、

- 1) 思考エンジンが思考中にも次に有望そうな局面を何局面でも選出できること
 - 2) 選出が思考より遥かに短い時間で完了すること
- この二つが必須条件である。

¹² dlshogi GitHub : <https://github.com/TadaoYamaoka/DeepLearningShogi>

¹³ Wikipedia : Alpha beta pruning : https://en.wikipedia.org/wiki/Alpha%E2%80%93beta_pruning

最終的な局面数は 1500 万局面程度を想定している。1 局面 1 秒で間断なく思考しつづける時、1 日 86400 局面思考できるから、半年で 1500 万程度であるという想定である。

この 1500 万局面の定跡をメモリ上に読み込んでいる状態においても 1 回の選出が 1 秒未満で完了する必要があると考え、今回はそれを目標に開発を行った。

そこで今回、

1) のために ranged alpha beta search という技法を開発した。

2) のために 高速化のための探索上の工夫を複数実施した。

第 4 節 ranged alpha beta search

まず、典型的な alpha beta 探索の疑似コードを示す。これはゲーム木探索の教科書などに載っている。¹⁴

```
// pos   : 局面
// alpha : alpha beta 探索の alpha 値
// beta  : alpha beta 探索の beta 値
Value search(pos, alpha, beta):
  foreach m in 全合法手:
    if pos を m で進めた局面が定跡 DB 上にある:
      pos.do_move(m) // m で 1 手進める
      value = -search(pos, -beta, -alpha)
      pos.undo_move() // 1 手戻す
    else:
      value = pos の局面で m を指した時の評価値
      // 思考した時に全合法手に対応する評価値が得られている。
      if alpha < value:
        if beta <= value:
          return value # beta-cut
        alpha = value # update alpha
  return alpha
```

上記のプログラムの

foreach m in 全合法手:

の直後に、pos を m で進めた局面が思考中であるなら、この指し手 m はないもの (非合法手) として扱えば、次に良い評価値を持つ leaf node が探せるはずである。基本的な考えかたはそれだけなのだが、これを高速に行える必要がある。

やねうら王では、局面に対応する hash key として Zobrist Hashing という技法を用いている。¹⁵ また、やねうら王には、ある局面で指し手 m を指した局面に対応する hash key を求める関数も用意されている。

これを用いれば実際には盤面を動かさずして、指し手 m を指したあとの局面の hash key が得られる。この hash key が思考中の局面集合のなかにあれば、ないものとして扱う。つまり、さきほどの部分は、次のようなコードを書いてある。

foreach m in 全合法手:

```
// 探索中のノードの hash key 集合に m を指した局面の hash key が含まれるなら、m をないものとして扱う。
if searching_nodes.contains(pos.key_next(m)):
  continue
```

ranged alpha beta search の基本的なコードは以上である。

この疑似コードの searching_nodes は実際のプログラムでは C++ の std::unordered_set を用いている。そのためこの集合に含むかどうかの判定は極めて短い時間で行われる。

hash key だとたまたま衝突してしまうと困るので、本来は局面を一意に識別できるものを使うべきである。

¹⁴ Wikipedia 「アルファ・ベータ法」:
<https://ja.wikipedia.org/wiki/%E3%82%A2%E3%83%AB%E3%83%95%E3%82%A1%E3%83%BB%E3%83%99%E3%83%BC%E3%82%BF%E6%B3%95>

¹⁵ Zobrist hashing :
https://en.wikipedia.org/wiki/Zobrist_hashing

そのようなものとしては SFEN 文字列¹⁶やそれを 256bit に符号化した PackedSfen¹⁷などがあるが、その変換コストは馬鹿にならない。

やねうら王では、hash key を通常使う 64bit のみならず、128bit、256bit に変更する機能があるので、本手法では、128bit の hash key を用いることにした。この場合、天文学的な確率でしか hash 衝突は生じないので、定跡生成の用途で用いるならば問題ないと言えるだろう。

第 5 節 棋譜上の局面を掘る

本提案手法の目的は、外部棋譜を用いず、自分の思考エンジンのみでノーメンテナン스로低コスト(少ない計算量)で定跡を自動生成することにあるが、実際には、与えられた棋譜の局面を掘り進めたいこともある。

そのため、さきほどの ranged alpha beta search の疑似コードの以下の行

value = pos の局面で m を指した時の評価値
の続きに以下のコードを追加する。

```
// kif_nodes: 棋譜上に出現した局面の hash key の集合
if kif_nodes.contains(pos.key_next(m)):
    value += 100
```

棋譜上に出現した局面であれば、leaf node の評価値に加点している。これにより、ranged alpha beta search を行う時にこの leaf node は優遇される。

尤も、元があまりに悪い評価値であれば 100 点を加点したところで評価値としては悪いままであるから、この leaf node が選出されることはない。このように、与えた棋譜であっても展開するに値しないならばそこ以上は掘らないというのも自動的に上のコードでなされるので、有望な枝だけが延長される。このため大量の棋譜を与えてもそのすべての局面について思考するわけではなく、優れた局面の「選出」性能を発揮する。

第 6 節 千日手局面回避のテクニック

本手法では、千日手局面に到達した時のスコアは 0 として扱っている。通常の探索においては先手は、初期局面においてプラスの評価値がついている。後手は(先手が初手で最善手を指した場合)自分の局面でマイナスの評価値がついている。すなわち、ranged alpha beta search では先手は千日手を回避しながらプラスの評価値の局面に辿り着きたいし、後手はマイナスの評価値を避けるため千日手を目指すような探索をしていることになる。

この手法を実装する上で難しい点として、千日手まわりの処理が挙げられる。そこまでの経路(手順)と同じ局面に循環してしまう時、それをどのように取り扱うかという問題がある。

局面の合流を取り扱わなければ千日手に関してあまり難しい問題は発生しないのだが(Deep Learning 系で MCTS を用いる将棋ソフトや囲碁のソフトはそうなのが多い)、定跡の生成で合流を取り扱わない場合、将棋では特に角換わりで手待ちを繰り返して同じ局面に合流することが多く、そこで組み合わせ爆発を生じてしまい、そこ以降の定跡を掘れなくなってしまう。

このため、定跡生成においては局面の合流を適切に取り扱う必要があるが、合流を取り扱うと千日手の問題が出てくる。

¹⁶ SFEN 文字列 「USI プロトコルとは」

: <http://shogidokoro.starfree.jp/usi.html>

¹⁷ やねうら王では SFEN 文字列をハフマン符号化し、256bit にする機能が備わっている。これを PackedSfen と呼ぶ。

そのため、今回、千日手局面を ranged alpha beta search で回避するための手法を開発した。次のように行う。

「思考」済みの局面が再び「選出」された時、それは千日手手順でその局面に到達したものだと考えられる。これを banned nodes 集合と呼ぶ。PV が千日手である場合、循環局面に至る指し手 m とその局面の hash key をペアで保持しておき、その組み合わせの指し手 m を非合法手として扱えば良いのである。

この状態で「選出」を行えば、自動的にその次に有望な leaf node が得られるので、千日手局面が PV になり続けて「思考」できないという状況は回避できる。

これで「選出」時の千日手に関する問題は解決したのだが、テラショック化の時に千日手局面をどう解決するかと言う問題はある。これについては、第 9 節で詳しく述べる。

第 7 節 PV and nonPV

ranged alpha beta search は、通常の α β 探索と同様に再帰的に行われるが、ある局面の探索結果を記憶しておき、二度目にその局面に訪問した時は前回の結果を用いたい。

しかし、PV の leaf node は「選出」されたあと、「思考」が完了するまで無いものとして扱われるので、記憶していた各局面での探索結果の値が不正なものとなってしまう。

そこで、PV と NonPV という考え方を導入する。これは、チェスの Stockfish¹⁸ というチェスのソフトなど α

β 探索を行う現代の将棋・チェスのソフトで広く使われている手法である。

PV node : PV (最善応手列) に関係する node に関する探索。枝刈りはほとんどしない。

NonPV node : PV 以外の node に関する探索。枝刈りわりとする。

すなわち、探索する時に PV node 用の探索モードと、NonPV node 用の探索モードとの二つの探索モードがあると考えられる。PV はしっかり leaf node まで確認するが、NonPV node では枝刈りをできる限りする。PV (最善応手列) の読みが担保されていればまあいいだろうと言う考えかたである。

例えば、ranged alpha beta search で指し手 m で 1 手進めて再帰的に search 関数を呼び出すところ (次行) は、

value = -search(pos, -beta, -alpha)

以下のように、最初にまず NonPV モードで探索をし、 α 値を更新しそうな時だけ真面目に PV モードとして探索しなおすというコードになる。

```
value = -search<NonPV>(pos, -beta, -alpha)
if nodeType == PV && alpha < value:
    value = -search<PV>(pos, -beta, -alpha)
```

また、NonPV node では、前回その局面に訪問して探索した時の探索のスコア (best alpha) を記録しておき、それを用いて枝刈りを行う。

ただし、 α β 探索では、探索窓として区間 (α, β) の範囲に収まる値を探索している。探索結果がこの値の範囲内に収まれば、再度訪問した時に、その値をそのまま使えるが、そうでない場合はどうなるのだろうか？これについては次節で述べる。

¹⁸ Stockfish : <https://stockfishchess.org/>

第8節 fail low/fail high

alpha beta 探索は、その局面で区間(alpha, beta)の間に収まる値のなかで最大のものを探す探索である。beta 以上の値がひとつでも見つければ、その局面の探索を即座に中断し、search 関数はその値を返す。この状態は fail high と呼ばれ、そこで関数から抜ける枝刈りのことは beta cut と呼ばれる。

fail high が生じた時、この node のなかの指し手で beta を超えるものが一つ以上あることはわかるが、これを次回の探索の時に活かすにはどうすれば良いだろうか？

これを解決するのが、bound lower(下界)という考え方である。¹⁹ fail high が生じた時の value(探索値)をこの局面の探索の値 search_value として保存しておく。また、同時にこの値に関するラベルとして bound lower であるとラベルをつけておく。

実際は search_value 以上の値を記録する指し手 m があるかも知れないが、前回の訪問時にすべての指し手を調べたわけではないので(beta cut したので)、そこはわからない。ただ、少なくとも search_value の値を持つ指し手 m が 1 つは存在するとは言えている。

再度訪問したときに、bound lower の値として search_value が記録されていたとして、その時の探索窓の beta が

```
beta <= search_value
```

であるなら、beta cut して良いことになる。(下界による枝刈り)

同様に、前回の探索で alpha を上回る指し手 m がなかったとして、この場合 fail low と呼ばれる状態だが、

¹⁹ 英語では lower bound と呼ぶが、Stockfish などのソースコード上では BOUND_LOWER という定数ラベルが用いられているので、それに倣い、ここでは bound lower と書く)

この時の最大を記録した value を search_value に保管し、bound upper(上界)とラベルをつけておく。

次にこの node に訪問したとき、alpha が

```
alpha >= search_value
```

であれば、どの指し手 m でも alpha を上回らず fail low が起きる見込みが高いので即座に関数から抜け出すことができる。(上界による枝刈り)

また、alpha < search_value < beta であるような search_value の時は、search_value に bound_exact(正確にその範囲)のラベルをつけておく。

以上3つの枝刈りは NonPV でのみ行われる。そこで、その3つを合わせると以下のような枝刈りが search 関数の先頭で行えることがわかる。

```
// すでに探索済みでかつ千日手が絡まないスコアが記録されている時の枝刈り。(NonPV 限定)
if nodeType == NonPV and search_value != VALUE_NONE:
    if bound == BOUND_EXACT
        or (bound == BOUND_LOWER && search_value >= beta)
        or (bound == BOUND_UPPER && search_value <= alpha)
        return search_value
```

第9節 cyclic score

循環スコア(cyclic score)とは千日手が絡んだ探索の値のことである。いかに NonPV node であっても循環スコア(cyclic score)を信じない方が良いと思われる。

千日手は経路(そこまでの手順)に依存するので、経路が異なれば千日手にならないこともある。

そこで、千日手を検出した場合、循環スコアには cyclic フラグを立てて、子 node に cyclic フラグが立

っていれば、その node の探索値 `search_value` は `cyclic score` であるとみなし、それを信用しない、とすることは出来る。

ただ、こうしてしまうと組み合わせ爆発が起きるよう
でテラショック化が現実的な時間で終わらなくなっ
てしまう。主な原因として次のような局面がある。

60 手目付近で矢倉に組み上げて、先後ともに角を動
かして手待ちを繰り返す。角の動かし場所は 5 箇所程度
あり、先後で 5×5 の組み合わせがある。`cyclic score`
だからと言って、前回の探索値が丸々使えないのでは、
こういう部分で組み合わせ爆発を容易に起こすよう
である。

`df-pn` による詰将棋にも似た問題があつて、GHI 問題
という有名な問題がある。`df-pn` に出現する GHI 問題自
体は、千日手を検出した場合、そこまでの手順を `hash`
`key` 化したものを置換表に記憶しておくことで、GHI 問
題自体は回避できる。^{20 21}

しかし、テラショック化の時に GHI 問題の回避のため
にそのように経路を保存しておいても上に書いたよう
にその経路自体が容易に組み合わせ爆発を起こすので
無限に近い経路の数をその局面の探索結果の情報とし
て保持する必要が出てくる。

これを回避するためには、千日手を検出した時、そこ
までの経路で千日手に関与しない局面の `hash key` の集
合(`white list`)と、千日手に関与する局面の `hash key`
の集合(`black list`)とその時の探索の結果値をその局
面の情報として保存しておく必要がある。これは実装自

²⁰ Kishimoto, Akihiro and Martin Müller. “A
solution to the GHI problem for depth-first
proof-number search.” *Inf. Sci.* 175 (2005): 296-314.

²¹ `koumori-n` , コウモリのちょーおんば 「詰将棋ソ
ルバーにおける GHI 問題対策」
[https://komorinfo.com/blog/and-or-tree-ghi-proble
m/](https://komorinfo.com/blog/and-or-tree-ghi-problem/)

体が難しく、また、よほど工夫しないとメモリ空間的に
も現在の PC では厳しいので、今回の採用は見送る。

そこで本手法では、`cyclic score` に関してはきちん
と取り扱うことは諦め、NonPV では `cyclic score` であ
ろうと信用するような実装にしてある。

第 10 節 本手法の効果

本論文の手法で 80 万局面ほど思考し、テラショック
化を行ったスーパーテラショック定跡では、初手 76 歩
に対する 34 歩を明確に悪手(評価値-243)だと認定した。
²²

2 手目 34 歩はプロの将棋においても昔からよくある
オープニング(序盤)だが、本論文の手法は、この 2 手目
の 34 歩を咎めるところまでできている。実際、そのあと
ソフト同士の対局によると先手勝率が 6 割程度であり、
このオープニングは現在プロの将棋でも減りつつある。
²³

また、第一章で紹介した `s-book_black` とスーパーテ
ラショック定跡とで対局させた場合、29 手目まで定跡
で進行し、スーパーテラショック定跡側の定跡が先に尽
きたが、その局面から最新の将棋ソフト(水匠 4)で対局
を引き継ぎ、思考時間を変えて 100 戦やってみたところ
55 勝 42 敗 3 引き分けであった。`s_book_black` は先手専
用定跡のため、スーパーテラショック定跡側は後手番で
ある。後手番で勝ち越しているのはこれは十分な戦果で
あり、人間がソフトを活用し、長い時間と労力を掛けて
作り上げた定跡集を、外部棋譜を一切使わず、かつ、ノ
ーマンテナンスな自動定跡生成により打ち破ったと言
えるだろう。

²² やねうら王ブログ , スーパーテラショック定跡が
76 歩に 34 歩を全否定 :
[https://yaneuraou.yaneu.com/2021/11/05/super-te
ra-shock-book/](https://yaneuraou.yaneu.com/2021/11/05/super-tera-shock-book/)

²³ 山口 祐, ツイッター
[https://twitter.com/ymg_aq/status/1456283886439
129089](https://twitter.com/ymg_aq/status/1456283886439129089)

第 11 節 本手法のまとめ

現代にふさわしい定跡生成手法として、次のような手法の誕生が望まれていた。

- 完全な自動生成 (ノーメンテナンス)
- 少ない計算資源で掘っていける
- (floodgate の棋譜など) 外部の棋譜を用いない
- (shotgun 定跡のように) 対戦相手の思考エンジンを仮定しない
- 思考対象局面の選出が小さな計算コストで出来る
- テラショック化のような定跡ツリー上で Minimax を行い、評価値の高い leaf node を目指す定跡が生成される

本論文で提案するスーパーテラショック定跡生成手法では、これらの要求をすべて満たすことが出来た。

また、スーパーテラショック定跡生成手法は実際にやねうら王のソースコード上に実装し、公開もしている。

²⁴

参考にしていただければ幸いである。

²⁴ やねうら王 GitHub , makebook2021.cpp :
<https://github.com/yaneurao/Yaneura0u/blob/e620b83d4ecb110120cf2abe4f4c806e0634c9b4/source/book/makebook2021.cpp>

チームビール工房 HFT 支店@wcsc32 アピール文

昨年同様 Have Fun Tech 代表の曾根壮大と第 28 回優勝、第 29 回準優勝の芝世武のチームです。

昨年同様巨大な深層学習モデルで最高精度の局面評価を誇ります。

今の GPU では探索速度が文字通り桁違いですのでどこまで伸びるかギリギリまで強化学習を進めようと考えています。

また、モデルに関してはブロック数、チャンネル数のみならずポリシーヘッド、バリューヘッドの部分変更など複数パターンをテスト中ですので本戦で用いるのが 40 ブロックになるかどうかは乞うご期待と言った感じです。

恐らく選手権後にご報告致します。

ㄣ (Qha)のPR文章

Ryoto Sawada, Yuki Ito, Toshihiro Shirakawa, Keigo Nitadori (Quorax 党 将棋部)



DeepMind ! ?
破壊した
はずでは...

進捗を575でまとめると

MuZeroはAlphaZeroより弱かった

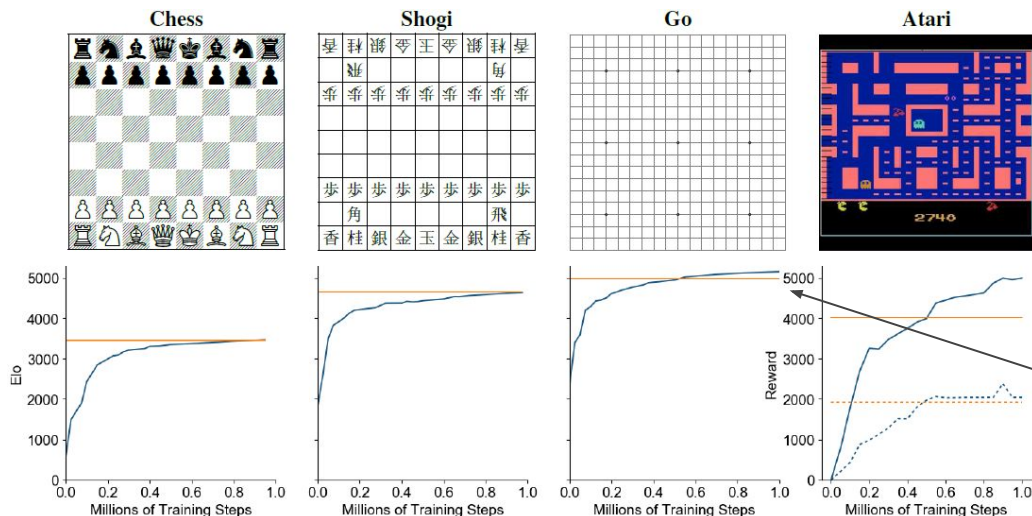


以下、MuZeroがどんなものであるか、どういう実験をしてどう
いう失敗をしたかについて解説します

大会で出すネタがねえどうしよう

そもそもMuZeroとは何者か

MuZeroはAlphaZeroの発展形のひとつ。探索時にゲームのシミュレータを必要としないことでより幅広い問題に適用することができる。また、原著論文によれば囲碁将棋などのゲームでもAlphaZeroと同等以上の性能を発揮している



次ページからAlphaZeroとMuZeroの違いを解説



囲碁ではAlphaZeroよりも強いと言ってる

AlphaZeroがどのようにして盤面を評価するか

AlphaZeroは局面を評価値に変換するAIとシミュレータ(図中のパソコン)からなる

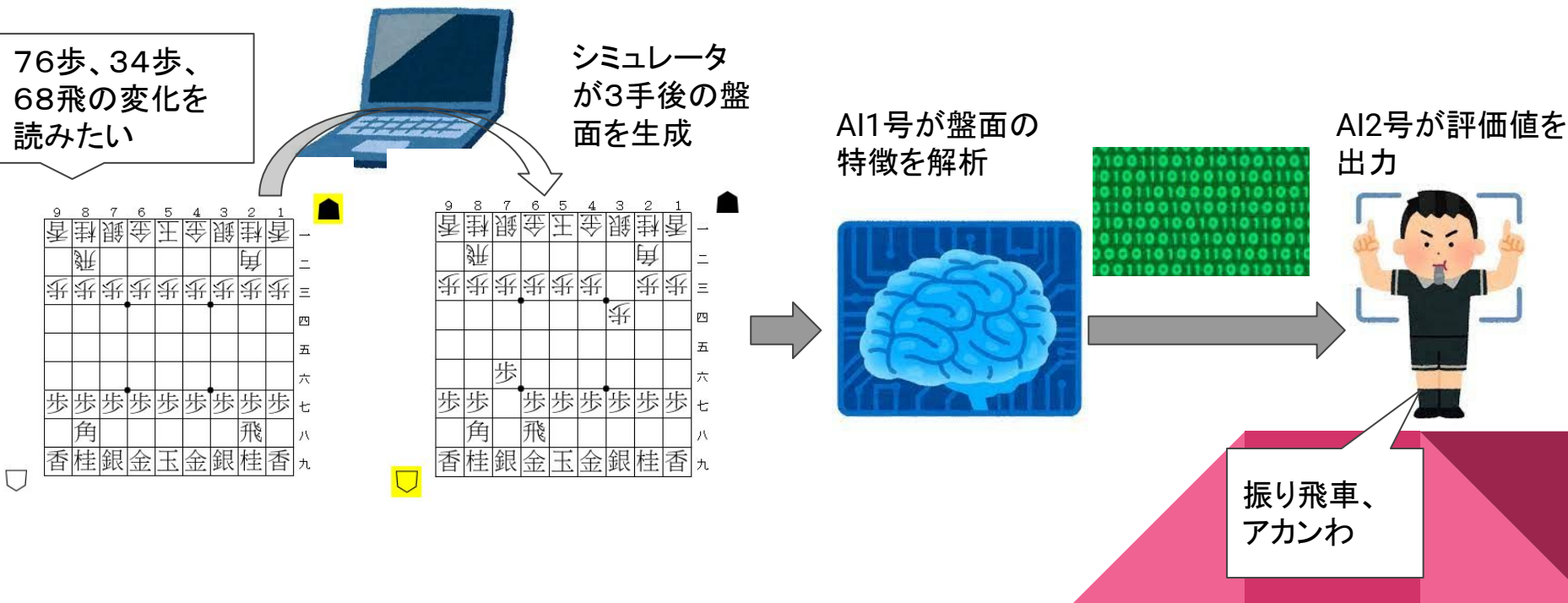
76歩、34歩、
68飛の変化を
読みたい

シミュレータ
が3手後の盤
面を生成

AI1号が盤面の
特徴を解析

AI2号が評価値を
出力

振り飛車、
アカンわ



MuZeroがどのようにして盤面を評価するか

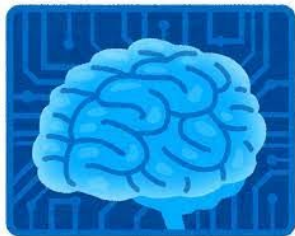
MuZeroは探索時にシミュレータを使わない

76歩、34歩、68
飛の変化を読み
たい



9	8	7	6	5	4	3	2	1	
香	桂	銀	金	玉	金	銀	桂	香	一
	飛							角	二
	歩							歩	三
									四
									五
									六
歩	歩	歩	歩	歩	歩	歩	歩	歩	七
	角							飛	八
香	桂	銀	金	玉	金	銀	桂	香	九

AI1号が盤面の
特徴を解析



76歩

34歩

68飛



AI3号が指し手の情報から数手
先の盤面の特徴を生成する



AI2号が評価値を
出力



振り飛車、
アカンわ

AlphaZeroとMuZeroの違い

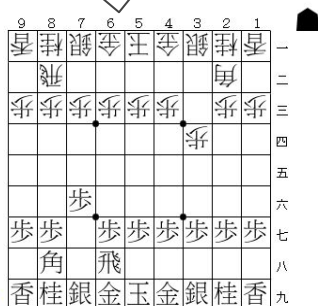
シミュレータが重いケース(テレビゲームなど)に対してMuZeroは有効

AlphaZeroでも対応可能

AlphaZeroには厳しい



シミュレータが3
手後の盤面を生
成(一瞬)



マ○オをジャンプさせた後の
盤面を生成(高コスト)



AlphaZeroとMuZeroの違い

囲碁、将棋などのシミュレータは極めて軽いため、MuZeroを使う必然性はない。
AlphaZeroとの優劣は数手先の局面をどれだけ正確に評価できるかで決まる

Q: 3手後の局面をどちらがより
正確に評価できるか



シミュレータが3
手後の盤面を生
成



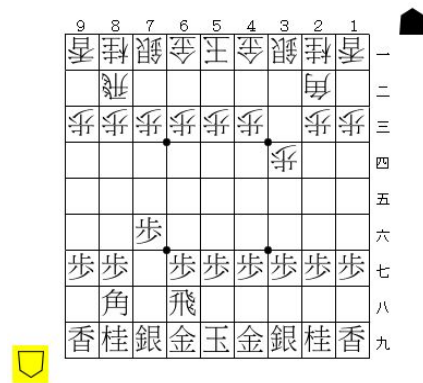
VS



AlphaZero vs MuZero

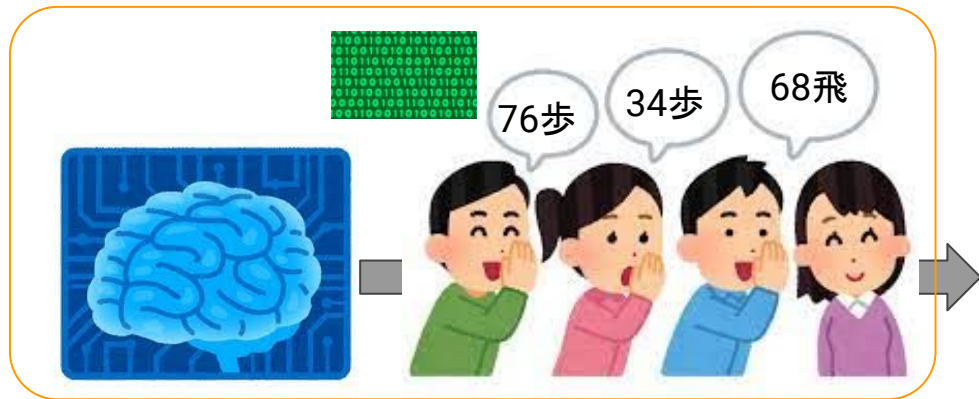
AlphaZero派の主張: AIを使って3手後の局面を完全に復元できるとは限らない。
シミュレータを使えば正確な結果が出るのだからシミュレータを使ったほうが良いに
決まってる

AIが将棋のルールを完全に理
解できるの？ どこかでエラーが
でたりしないの？

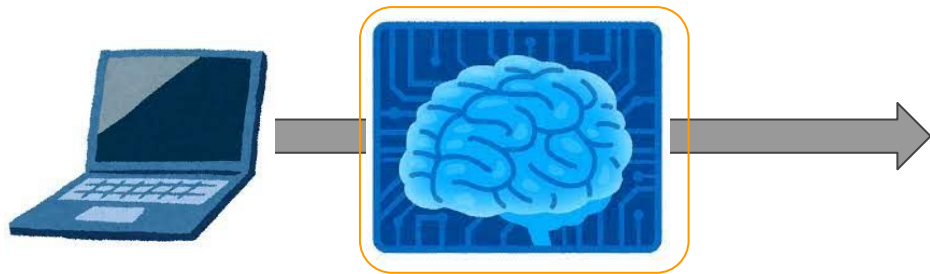


AlphaZero vs MuZero

MuZero派の主張: AIのモデルは基本的に大きいほど精度が高い。盤面の伝言ゲームを経ることにより高度な情報を抽出することができる



評価値を計算するまでに経由するAIの数が多い → より沢山の計算を行っている → 精度が高い



This suggests that MuZero may be caching its computation in the search tree and using each additional application of the dynamics model to gain a deeper understanding of the position.
(原著論文より引用)

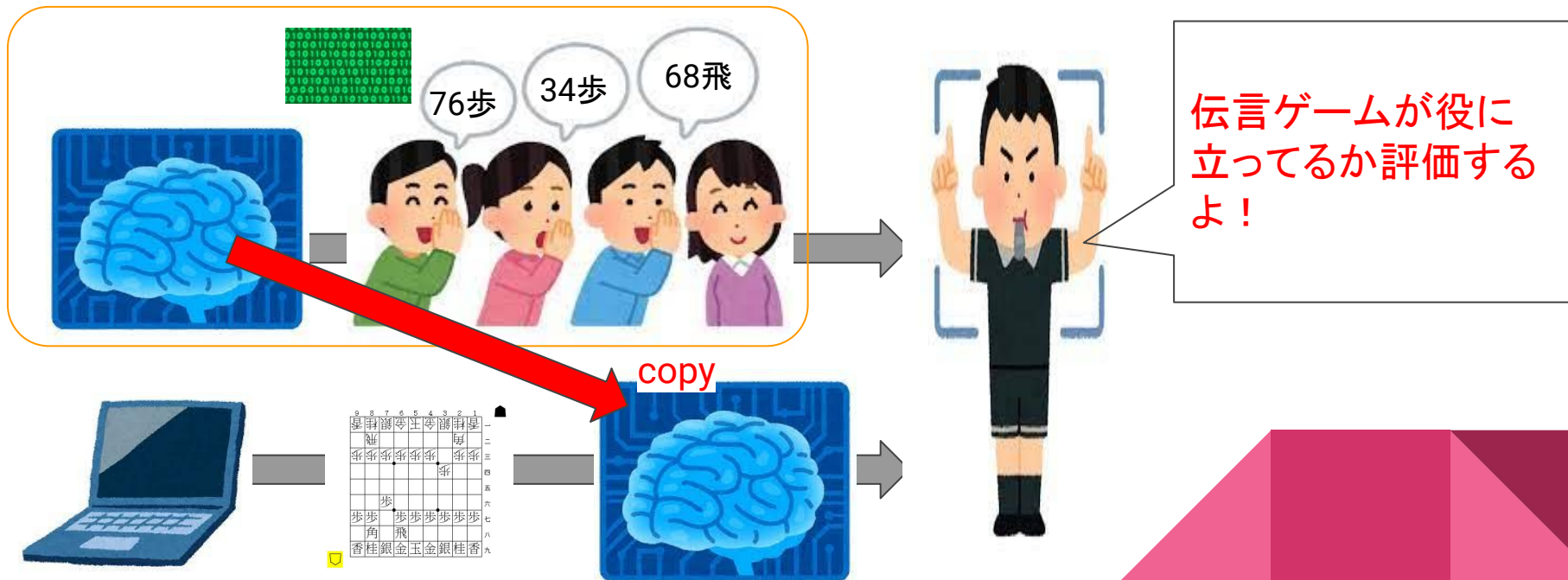
実験してみる

- python-dlshogi2を改造してMuZero形式に対応した → 全然勝てなかった
 - 1手1秒で20連敗したあたりで心が折れた
- 原著論文ではMuZero同士の対局では1手800 simulation(800nodesに相当？)、elmoなどとの対局では1手0.1秒で対局したらしい
 - スレッド数などの条件は不明。というか、1手0.1秒っておま、WCSC27のelmoってなどとツツコミどころは絶えない
 - ディープ系列同士の対局は同じ局面に偏りがちなので初期局面をどうするの問題もあまり触れられていないように見える



勝率以外の指標でも評価してみる

MuZeroが内包している盤面を特徴量に変換する部分を切り出し、AlphaZeroと同じ挙動をさせたうえで、各々のモデルの盤面評価制度を比較する



伝言ゲームの効果はあるのか

MuZero形式にすることで盤面評価精度が下がった

	MuZeroの一致率(伝言ゲーム5回後の一致率)	MuZeroから切り出したAlphaZero部分+シミュレータの一致率
公開モデルからの finetune	40.6%	53.1%
ゼロからの学習	35.3%	48.1%

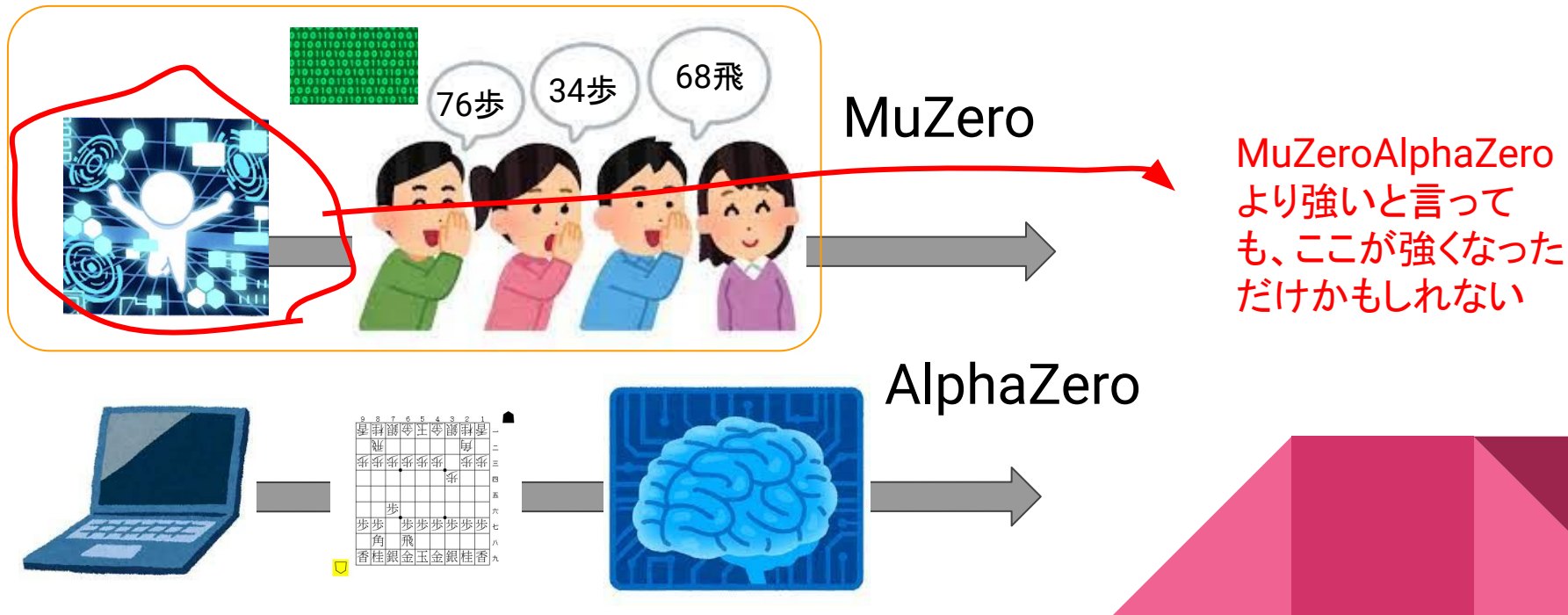


伝言ゲームをやるぐらいなら、シミュレータ回したほうがいいよ

- train/validationはfloodgateの棋譜から生成したデータを利用 (train:2000万、test:10万)
- MuZeroは「局面」ではなく「棋譜」のデータが必要で既存教師データを流用できない
- 教師の数としては正直頼りなくはある

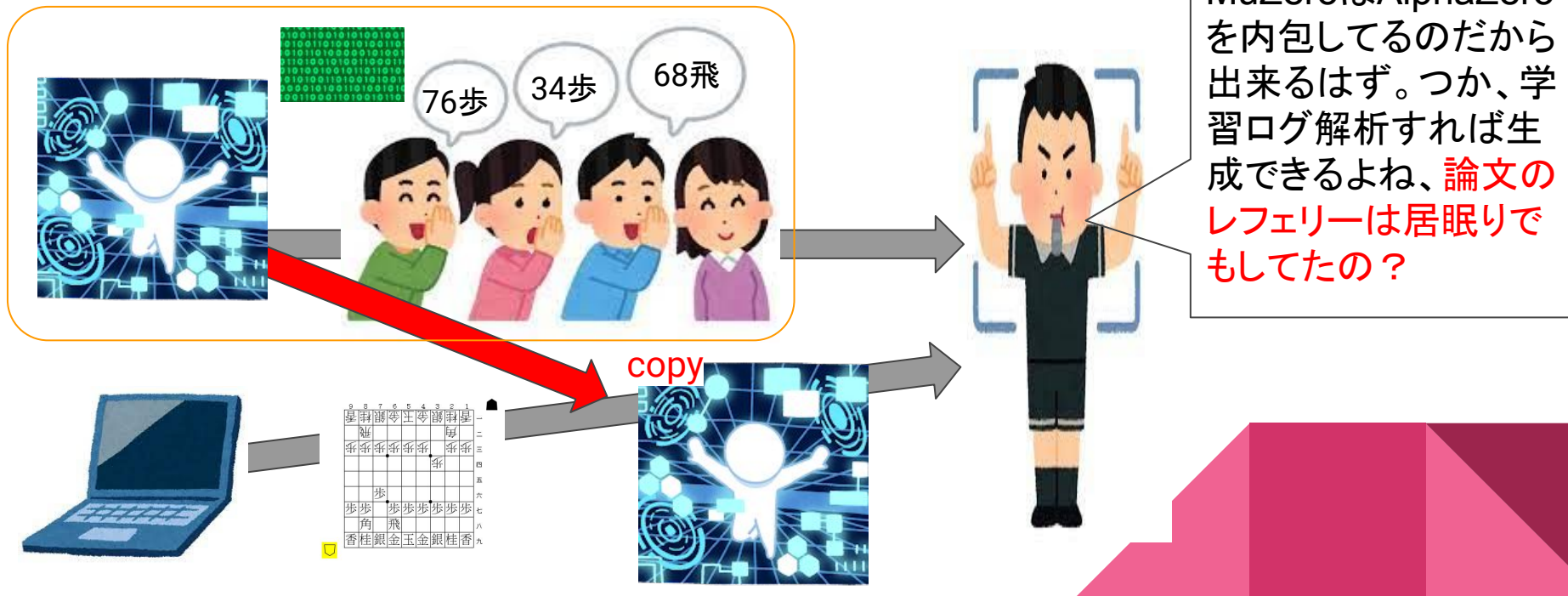
原著論文と実験結果の違いに関する考察

AlphaZeroとMuZeroとでは学習環境が厳密には一致していない。伝言ゲームが足を引っ張っていたとしても他の部分で強くなっていた可能性がある



(rsawadaの考える)原著論文の問題点

AI3号(原著論文でいうdynamics)が盤面評価の精度にどの程度影響を与えているかを評価するべきだった



原著論文と実験結果の違いに関する考察(論文を擁護)

- そもそもrsawadaの再現実装が失敗している
 - 学習の条件に敏感な手法であるとか、教師の数が少ないとか
 - とくに、伝言ゲームのネットワークが浅すぎた(伝言ゲームがシミュレータの質を超えるにはある程度のlayer数が必要)はありそう
- シミュレータ不要であることが売りであり、そもそも強さは売りにしてない
 - 原著論文では囲碁でAlphaZeroより強いことをあまり推してない(かも)
 - いや、abstractでもガッツリ触れてるな
- 大会が終わったらソースコードを公開するので遊んでみて欲しい
 - ~~コメント内に書き込んだDM社への悪口を消さないと~~

困ったぞ、出すものがない



この時期にコンテンツが
何もないのは数年ぶりだ

- 今から頑張って飛車を振る
 - ~~負けても振り飛車のせいにできる~~
- 実況を頑張る
 - 実況ツールの公開もやらないと(linuxはともかく、windowsで動かない)
- 別に負けてもいいやとMuZeroを放り込む
 - これ面白いかな.....？
 - 教師データの生成から学習ルーチン、探索部まで創っておいてお蔵入りするのも癪だけど
 - ~~もう全部、論文のレフェリーが悪い~~

マメット・ブンブク アピール文書

ザイオソフト コンピュータ将棋サークル
野田久順 岡部淳 鈴木崇啓
河野明男 伊苅久裕

目次

- マメット・ブンブク
- 改良点
- 使用ライブラリ

たぬきち

- 元ネタは『ファイナルファンタジーXIV』に登場するミニオンの名前です。
 - ミニオンとは、プレイヤーと一緒に連れて歩くことができる「ペット」のようなものです。
- ミニオンのように、いつも一緒にいてもらえるような将棋ソフトを目指しています。

改良点（1）

- 定跡生成手法に、たややん2020手法を使用しました。
 - floodgate の棋譜から定跡を作りました。
 - レーティング 3900 以上のソフト同士の対局の棋譜のみ使用しました。
 - 勝率が 33% 以上の指し手のみ使用しました。

改良点（2）

- 学習データの生成条件を変更しました。
 - 自己対局の対局開始局面を変更しました。
 - floodgate のレーティング 3900 以上のソフト同士の棋譜を使用しました。
 - 32 手目までからランダムに局面を選択しました。
 - ランダムムーブを入れないようにしました。
 - 探索深さ 9 で対局しました。

改良点（3）

- 局面数を数億から数十億に増やしました。
- 学習データの生成に「水匠 5」を使用させていただきました。
 - 貴重なソフトを公開してくださりありがとうございます。

改良点（4）

- 入力特徴量に HalfKP_vm を採用しました。
 - 玉が 6～9 筋にいるとき、盤面を左右反転させて入力します。
 - 左右対称次元下げと同じ効果が得られ、かつネットワークパラメーター数が減ります。
 - より少ない学習データで学習できると考えられます。

改良点（5）

- 機械学習のパラメーターを変更しました。
 - elmo 式学習法における、勝敗項の教師信号 (t) を $t=0.80$ としました。
 - 評価値のスケールが小さくなりました。
 - K・P・HalfKP 相対次元下げを無効化しました。

使用ライブラリ

- やねうら王
 - やねうら王を元に改造した思考部を使用している。
 - 独自の工夫を加えるにあたり、改造しやすく、レーティングも高いため。
- 水匠5
 - 学習データの生成に使用している。
 - レーティングが高く、学習データの生成速度が速いため。
- tanuki-
 - 学習データの生成に使用している。
 - 過去に開発した資産の再利用のため。

よろしくお願いします

第 32 回世界コンピュータ将棋選手権

dlshogi with HEROZ アピール文章

山岡忠夫
加納邦彦
山口祐
大森悠平
2021/3/29

※下線部分は、第 31 回世界コンピュータ将棋選手権からの差分を示す。

1 dlshogi のアピールポイント

dlshogi は、ディープラーニングを使用した将棋 AI である。

2017 年より、AlphaGo の手法を参考に開発を行っている。

ディープラーニング系の将棋 AI は、大局観に優れており、序中盤の形勢判断が従来型将棋 AI と比べて正確であるという特徴がある。一方、終盤の読みが重要になる局面では、従来型将棋 AI の方が正確な場合がある。

dlshogi は、終盤の課題に対処するために、独自の工夫を行っている。

具体的には、「MCTS の葉ノードでの短手数 of 読み探索」、「ルート局面で df-pn による長手数の読み探索」、「勝敗が確定したノードのゲーム木への伝播」、「PV 上の局面に対する長手数の読み探索」、「強化学習時に初期局面集を使用して局面の多様性を確保する」、「強化学習時に df-pn により読み探索を行い読みを報酬とする」という工夫を行っている。

これらのいくつかは、現在、dlshogi 以外のディープラーニング系の将棋 AI には取り入れられているが、dlshogi 以前にこれらを導入しているディープラーニング系の将棋 AI はなかった。

2 チーム参加について

今大会では、HEROZ チームとして参加する。

3 dlshogi の特徴

- ディープラーニングを使用
- 指し手を予測する Policy Network
- 局面の勝率を予測する Value Network
- 入力特徴にドメイン知識を積極的に活用
- モンテカルロ木探索
- 未探索のノードの価値に親ノードの価値を使用
- GPU によるバッチ処理に適した並列化

- 自己対局による強化学習
- 詰み探索結果を報酬とした強化学習
- 既存プログラムを加えたリーグ戦による強化学習
- 既存将棋プログラムの自己対局データを混ぜて学習
- 既存将棋プログラムの自己対局データを使った事前学習
- ブートストラップ法による Value Network の学習
- 引き分けも含めた学習
- 指し手の確率分布を学習
- 同一局面を平均化して学習
- 評価値の補正
- SWA(Stochastic Weight Averaging)
- 末端ノードでの短手順の詰み探索
- ルートノードでの df-pn による長手順の詰み探索
- 勝敗が確定したノードのゲーム木への確実な伝播
- PV 上の局面に対する長手数の詰み探索
- 序盤局面の事前探索（定跡化）
- 定跡作成時に floodgate の棋譜の統計を利用した確率分布を方策に利用
- マルチ GPU 対応 (NVIDIA A100×8 を使用予定)
- TensorRT を使用
- Optuna による探索パラメータの最適化
- 確率的な Ponder
- ノードのガベージコレクションとノード再利用処理の改良
- 飛車と角の利きのビット演算
- 2 値の入力特徴量を 1 ビットで転送することで推論のスループットを向上
- KL 情報量を利用した時間制御

4 使用ライブラリ

- Apery¹ (WCSC28)
→局面管理、合法手生成のために使用

4.1 ライブラリの選定理由

本プログラムは、将棋におけるディープラーニングの適用を検証することを目的としており、学習局面生成、局面管理、合法手生成については、使用可能なオープンソースがあれば使用する方針である。そのため、学習局面を圧縮形式(hcpe)で生成する機能を備えていて、合法手生成を高速に行える Apery を選定した。

¹ <https://github.com/HiraokaTakuya/apery>

5 各特長の具体的な詳細（独自性のアピール）

5.1 ディープラーニングを使用

DNN(Deep Neural Network)と MCTS を使用して指し手を生成する。
従来の探索アルゴリズム(α β 法)、評価関数(3 駒関係)は使用していない。

5.2 Policy Network

局面の遷移確率を Policy Network を使用して計算する。

Policy Network の構成には、Residual Network を使用した。

入力の畳み込み 1 層と、ResNet 15 ブロック(畳み込み 2 層で構成)と出力層の合計 32 の畳み込み層で構成した。フィルターサイズは 3 (入力層の持ち駒のチャンネルのみ 1)、フィルター数は 224 とした。

5.3 Value Network

局面の勝率を Value Network を使用して計算する。

Value Network は、Policy Network と出力層以外同じ構成で、出力層に全結合層をつなげ、シグモイド関数で勝率を出力する。

5.4 入力特徴にドメイン知識を積極的に活用

Alpha Zero では、入力特徴に呼吸点のような囲碁の知識を用いずに盤面の石の配置と履歴局面のみを入力特徴とすることで、ドメイン知識なしでも人間を上回ることが示された。しかし、その代償として、入力特徴にドメイン知識を活用した AlphaGo Lee/Master に比べて倍のネットワークの層数が必要になっている。AlphaGo Zero の論文の Figure 3 によると、ネットワーク層数が同一のバージョンでは Master を上回る前にレーティングが飽和している。

強い将棋ソフトを作るという目的であれば、積極的にドメイン知識を活用した方が計算リソースを省力化できると考えられる。

そのため、本ソフトでは、入力特徴に盤面の駒の配置の他に、利き数と王手がかかっているかという情報を加えている。それらの特徴量が学習時間を短縮する上で、有効であることは実験によって確かめている。

5.5 モンテカルロ木探索

対局時の指し手生成には、Policy Network と Value Network を活用したモンテカルロ木探索を使用する。

ノードを選択する方策に、Policy Network による遷移確率をボーナス項に使用した PUCT アルゴリズムを使用する。PUCT アルゴリズムは、AlphaZero の論文²の疑似コードに記述さ

² <http://science.sciencemag.org/content/362/6419/1140>

れた式を使用した。

また、末端ノードでの価値の評価に、Value Network で計算した勝率を使用する。

通常のモンテカルロ木探索では、末端ノードから複数回終局までプレイアウトを行った結果（勝率）を報酬とするが、将棋でランダムなプレイアウトは有効ではないため、プレイアウトを行わず Value Network の値を使用する。

5.6 未探索のノードの価値に親ノードの価値を使用

モンテカルロ木探索の UCB の計算時に、未探索の子ノードがある場合、そのノードの価値に何らかの初期値を与える必要がある。子ノードの価値は親ノードの価値に近いだろうという将棋のドメイン知識を利用し、それまでの探索で見積もった親のノードの価値を動的に初期値として使用する。ただし、ノードの訪問回数が増えるに従い、その価値の減衰を行い、幅より深さを優先した探索を行う (FPU reduction)。

5.7 GPU によるバッチ処理に適した並列化

複数回のシミュレーションを順番に実行した後、それぞれのシミュレーションの末端ノードの評価をまとめて GPU でバッチ処理する。その後、評価結果をそれぞれのシミュレーションが辿ったノードにバックアップする。以上を一つのスレッドで行うことで、マルチスレッドによる実装で課題となる GPU の計算後にスレッドが再開する際にリソース競合が起きる問題（大群の問題）を回避する。

GPU で計算中は、CPU が空くため、同じ処理を行うスレッドをもう一つ並列で実行する。2つのスレッドが相互に CPU と GPU を利用するため、利用効率が高い処理が可能となる。

5.8 自己対局による強化学習

事前学習を行ったモデルから開始して、AlphaZero³と同様の方式で強化学習を行う。自己対局により教師局面を生成し、その教師局面を学習したモデルで、再び教師局面を生成するというサイクルを繰り返すことでモデルを成長させる。

2018年の大会で使った elmo で生成した教師局面で収束するまで学習したモデルに比べて、自己対局による強化学習によって有意に強くすることができた。

5.9 詰み探索結果を報酬とした強化学習

自己対局時に終局まで対局を行うと、モンテカルロ木探索の特性上、詰むまでの手順が長くなる傾向がある。勝率予測に一定の閾値を設けることで、終局する前に勝敗を判定することで対局時間を短縮できるが、モデルの精度が低い場合は誤差が大きいため、学習精度に影響する。

この課題の対策として、df-pn による高速な長手数の詰み探索の結果を報酬とした。単純に

³ <https://arxiv.org/abs/1712.01815>

すべての局面で詰み探索を行うと、自己対局の実行速度が大幅に落ちてしまう。自己対局は複数エージェントに並列で対局を行わせ、各エージェントからの詰み探索の要求をキューに溜めて、詰み探索専用スレッドで処理するようにした。エージェントが GPU の計算待ちの間に詰み探索が完了する。エージェントが探索している局面は別々のため、時間のかかる詰み探索の要求が集中することは少ない。これにより自己対局の速度を大幅に落とすことなく長手数の詰み探索を行えるようになった。

5.10 既存プログラムを加えたリーグ戦による強化学習

自分自身のプログラムのみで強化学習を行うと戦略に弱点が生まれる可能性がある。弱点をふさぐには多様なプログラムによるリーグ戦が有効だが、複数のエージェントを学習するにはエージェント数の分だけ余分に計算資源が必要になる。

計算資源を省力化して、リーグ戦の効果を得るために、オープンソースで公開されている既存の将棋プログラムを 1/8 の割合でリーグに加えて強化学習を行うようにした。

5.11 既存将棋プログラムの自己対局データを使った事前学習

本プログラムを使用して、Alpha Zero と同様に、ランダムに初期化されたモデルから強化学習を行うことも可能だが、使用可能なマシンリソースが足りないため、スクラッチからの学習は行わず、既存将棋プログラムの自己対局データを教師データとして、教師あり学習でモデルの事前学習を行う。

教師データには、elmo で生成した自己対局データを使用した。

5.12 既存将棋プログラムの自己対局データを混ぜて学習

以前の dlshogi は、入玉宣言勝ちできる局面でなかなか入玉宣言勝ちを目指さないという課題があった。

自己対局では入玉宣言勝ちの棋譜が少ないため、それを補うため既存将棋プログラム(水匠)の自己対局で、入玉宣言勝ちの棋譜を生成し、dlshogi の自己対局のデータに混ぜて学習した。

5.13 ブートストラップ法による Value Network の学習

Value Network の学習の損失関数は、勝敗を教師データとした交差エントロピーと、探索結果の評価値を教師データとした交差エントロピーの和とした。

このように、本来の報酬（勝敗）とは別の推定量（探索結果の評価値）を用いてパラメータを更新する手法をブートストラップという。

経験的にブートストラップ手法は、非ブートストラップ手法より性能が良いことが知られている。

5.14 引き分けも含めた学習

将棋はルールに引き分けがあるゲームであるため、引き分けも正しく学習できる方が望ましい。そのため、自己対局で引き分けとなった対局も学習データに含めて学習した。

5.15 指し手の確率分布を学習

以前の dlshogi では、指し手のみを学習していたが、AlphaZero と同様に、自己対局時で MCTS で探索した際のルート局面の子ノードの訪問回数に従った確率分布を学習するように変更した。確率分布を学習することで、最善手と次善手の行動価値が近い場合に、次善手の行動価値を正しく学習できるようになる。

確率分布を学習することで、floodgate の棋譜に対する一致率が向上することが確認できたが、対局して強さを計測すると弱くなるという現象が確認できた。原因は、モデルの方策の出力の性質が変わるため、探索パラメータの調整が必要なためであった。Optuna を使用して探索パラメータを最適化(エラー! 参照元が見つかりません。参照)することで、指し手のみを学習したモデルよりも強くすることができた。

5.16 同一局面を平均化して学習

自己対局では、序盤の同一の局面の教師データが多く生成される。それらの重複した局面を別のサンプルとして学習すると、モデルの学習に偏りが起きる。

局面の偏りをなくするために、同一の局面を集約し、指し手の確率分布と勝敗を平均化し、1 サンプルとして学習した。

5.17 評価値の補正

自己対局で生成するデータには、MCTS で探索して得られた勝率(最善手の価値)を局面の評価値を記録し、学習時にブートストラップ項(エラー! 参照元が見つかりません。参照)として使用している。記録した評価値(勝率)が、実際の対局の結果から算出した勝率と一致しているか調べたところ、乖離しているという現象が確認できた。そのため、評価値を実際の自己対局での勝率に合うように、補正を行った。

5.18 SWA(Stochastic Weight Averaging)

画像認識の分野でエラー率の低減が報告されている手法である、SWA(Stochastic Weight Averaging)をニューラルネットワークの学習に取り入れた。一般的なアンサンブル手法では、推論結果の結果を平均化するが、SWA では学習時に一定間隔で重みを平均化することでアンサンブルの効果を実現する。

5.19 末端ノードでの短手順の詰み探索

モンテカルロ木探索の末端ノードで、5 手の詰み探索を行い、詰みの局面を正しく評価できるようにする。並列化の方式により、GPU で計算中の CPU が空いた時間に詰み探索を行う

ため、探索速度が落ちることはない。

5.20 ルートノードでの df-pn による長手数詰み探索

モンテカルロ木探索は最善手よりも安全な手を選ぶ傾向があるため詰みのある局面で駒得になるような手を選ぶことがある。

対策として、詰み探索を専用スレッドで行い、詰みが見つかった場合はその手を指すようにする。

詰み探索は、df-pn アルゴリズムを使って実装した。優越関係、証明駒、反証明駒、先端ノードでの 3 手詰めルーチンにより高速化を行っている。

5.21 勝敗が確定したノードのゲーム木への確実な伝播

モンテカルロ木探索で構築したゲーム木の末端ノードで詰みが見つかった場合、その結果をゲーム木に伝播して利用する。つまり、モンテカルロ木探索に、AND/OR 木の探索を組み合わせ、詰みの結果を確実にゲーム木に伝播するようにする。

5.22 PV 上の局面に対する長手数詰み探索

ディープラーニング系の将棋 AI は、選択的に探索を行うために、終盤の局面で読み抜けがあると、頓死することある。

頓死を防ぐため、PV 上の局面に対して、df-pn による長手数詰み探索を行い、詰みが見つかった場合、局面の価値を更新するようにする。

5.23 序盤局面の事前探索（定跡化）

出現頻度の高い序盤局面は、対局時に探索しなくても、事前に探索を行い定跡化しておくことができる。また、事前に探索することで、対局時よりも探索に時間をかけることができる。

ゲーム木は指数関数的に広がるため、固定の手数までの定跡を作成するよりも、有望な手順を選択的に定跡に追加する方が良い。自分が指す手は、1 つ局面につき最善手を 1 手（または数手）登録し、それに対する応手は、公開されている定跡や棋譜の統計情報を使って確率的に選択する。その手に対して、また最善手を 1 手（または数手）登録する。この手順により、頻度の高い局面については深い手順まで、頻度の低い局面については短い手順の定跡を作成することができる。

5.24 定跡作成時に floodgate の棋譜の統計を利用した確率分布を方策に利用

定跡を自分自身の探索のみで作成した場合、読み抜けがあった場合に定跡を抜けた後に不利な局面になる恐れがある。そのため、モンテカルロ木探索の PUCT の計算で、方策ネットワークの確率分布と floodgate の棋譜の統計を利用した確率分布を平均化した確率分布を利

用し、致命的な読み抜けを防止する。

5.25 マルチ GPU 対応

複数枚の GPU を使いニューラルネットワークの推論を分散処理する。

「5.7 GPU によるバッチ処理に適した並列化」の方式により、GPU ごとに 2 つの探索スレッドを割り当てることで、GPU を増やすことでスケールアウトすることができる。ノードの情報は、すべてのスレッドで共有する。

確認できている範囲で 4GPU までほぼ線形で探索速度を上げることができている。

5.26 TensorRT を使用

モデルの学習にはディープラーニングフレームワークとして PyTorch を使用しているが、対局プログラムには、推論用ライブラリの TensorRT を使用する。

TensorRT を使うことで、事前にレイヤー融合などのニューラルネットワークの最適化を行うことで、推論を高速化することができる。TensorCore に最適化されており、TensorCore を搭載した GPU では CUDA+cuDNN で推論を行う場合より、約 1.33 倍の高速化が可能になる⁴。

また、対局の実行環境にディープラーニングフレームワークの環境構築を不要とすることを目的とする。

5.27 Optuna による探索パラメータの最適化

PFNにより公開された Optuna⁵を使用して、モンテカルロ木探索の探索パラメータ (PUCT の定数、方策の温度パラメータ) を最適化した。

Optuna は、主にニューラルネットワークの学習のハイパーパラメータを最適化する目的で利用されるが、将棋エンジン同士の連続対局の勝率を目的関数として、探索パラメータの最適化に使えるようにするスクリプト⁶を開発した。Optuna の枝刈り機能により、少ない対局数で収束させることができる。

5.28 確率的な Ponder

モンテカルロ木探索は確率的にゲーム木を成長させる。その特性を活かして、相手が思考中に、相手局面からモンテカルロ木探索を行うことで、確率的に相手の手を予測して探索を行うことができる。予測手 1 手のみを Ponder の対象とするよりも、効率のよい Ponder が実現できる。

⁴ <https://tadaoyamaoka.hatenablog.com/entry/2020/04/19/120726>

⁵ <https://optuna.org/>

⁶

https://github.com/TadaoYamaoka/DeepLearningShogi/blob/master/utis/mcts_params_optimizer.py

5.29 ノードのガベージコレクションとノード再利用処理の改良

世界コンピュータ将棋オンライン大会でノード再利用に 10 秒以上かかる場合があることがわかったため、ノード再利用の方式の見直しを行った。

以前は、オープンアドレス法でハッシュ管理を行っており、ルートノードから辿ることができないノードをすべてのハッシュエントリに対して線形探索してノードの削除をおこなっていた。

これを、Leela Chess Zero のゲーム木の管理方法⁷を参考に、ゲーム木をツリーで管理を行うようにし、ルートの兄弟ノードをガベージコレクションする方式に変更した。ノードの合流の処理が行えなくなるという欠点があるが、ノード再利用を短い時間で行えるようになった。

5.30 飛車と角の利きのビット演算

第 31 回世界コンピュータ将棋選手権の Qugy のアピール文章⁸による、飛車、角の利きをビット演算により求める方法を実装した（実装はやねうら王のソースコードを参考にした）。ZEN2 の CPU で NPS が約 1% 向上した。

5.31 2 値の入力特徴量を 1 ビットで転送することで推論のスループットを向上

マルチ GPU を使用した場合、4GPU 以上では CPU と GPU 間の帯域がボトルネックになるため、2 値の入力特徴量を float の代わりに、1bit で表現し、GPU にビットで転送後、GPU 側で CUDA のプログラムでバッチ単位に並列に float に戻す処理を実装した。こうすることで、転送量が削減でき、NPS が 36.6%向上した。

5.32 KL 情報量を利用した時間制御

探索のルート局面の方策の分布と探索後の訪問回数の分布の KL 情報量が高い局面は、探索がより重要になる局面と考えられるため、KL 情報量が高い場合、より探索に時間を使うようにした。

KL 情報量を利用した時間制御を行うことで、同じ持ち時間でレーティングが 44.1 向上した。

⁷ <https://tadaoyamaoka.hatenablog.com/entry/2020/05/05/181849>

⁸ https://www.apply.computer-shogi.org/wcsc31/appeal/Qugiy/appeal_210518.pdf

「習甦」

- ・探索：1スレッドはdf-pnにより詰みを確認，他のスレッドはStockfishから取捨選択したシンプルな α - β 探索
- ・評価関数：玉の位置に対する各駒の位置と全升目の利き数（先後別3まで）を特徴量とするニューラルネットワーク
- ・学習：入玉した場合，宣言法に従い自駒の守備力が上がり敵駒の攻撃力が下がるよう設計した関数の値を報酬に加算

アピール文

<プログラムの名称>

いちびん

<プログラムの特徴>

(1) 概要

2年ほど前までは、完全自作路線をとっていましたが、能力の限界を感じ、その後、公開されている他人様のプログラム（いわゆる、やねうら王系のプログラム）を改造・利用させていただいています。

(2) 評価関数系の特徴

NNUE 系の評価関数を使っています。数年前のタヌキさんの関数を祖先として、自分なりに強化学習を繰り返しています。ことしは、周回遅れですが、FV-SCALE のところを、いじっています。

(3) 探索系の特徴

基本的には、JAVA を使って、やねうら王系のプログラムの動作をコントロールしています。JAVA 部分は、自作です。クラーデース・エスティマトス (Clades-Estimatus : CS) という概念で、評価点がマウス 800 点からマウス 1200 点 (この CS 閾値は、対戦相手の想定強さに依存) より下回った場合には、この先、どんなに頑張っても、勝つ見込みがほとんどない状態「推定負け」と判断して、(a) 対戦相手が読み落としている数十手先の相手方トン死を探索する、(b) 数十手先で、もしかしたら効果が出てきそうな升目に利きをきかせる、いわゆる勝負手を探索する。この 2 つのシミュレーションに、時間切れ負けを無視して、残り時間を投入するアイデアで挑みます。

(4) その他

以上のようなプログラムの特徴から、負けそうになると、異常な長考モードに突入します。往生際が悪いプログラムですが、お許しください。

第 32 回 世界コンピュータ将棋選手権

Novice アピール文章

中屋敷 太一 熊谷 啓孝 笹井 雄貴 矢内 洋祐

2022 年 2 月 15 日

1 概要

本稿では、コンピュータ将棋ソフトウェア Novice の用いているアルゴリズムについて記述する。筆者らは、Novice を AlphaZero [1] に倣って開発した。Novice は、AlphaZero 同様、探索（先読み）にモンテカルロ木探索 (Monte-Carlo Tree Search, MCTS)，そして局面の評価（大局観）にニューラルネットワークを用いている。自己対戦による棋譜の生成と、その棋譜からのニューラルネットワークの学習を繰り返すことにより、強い評価関数を学習することを目指している。

本稿では、Novice が AlphaZero のアルゴリズムに加えた工夫について記述する。

2 Novice の工夫

本節では、ニューラルネットワークの学習手法について記述する。ニューラルネットワークの学習は、自己対戦により教師データを生成し、その教師データを用いてニューラルネットワークのパラメータを更新して行う。初めに、自己対戦の用いている工夫を記述する。そしてその次に、ニューラルネットワークの学習に関する工夫について記述する。

2.1 自己対戦

AlphaZero の手法では、自己対戦による教師データから学習が行われているが、学習の際には膨大な計算資源が使われた [1]。そのため、この再現実験を行うことは容易ではなく、現実的な時間で強い将棋プログラムを作成するためには、いくつかの工夫が必要となる。

筆者らは、学習の予備実験の結果、3 つの次に記述する問題に遭遇した。1 つ目は、少ない教師データが、千日手による引き分けによって終局していること、2 つ目は、優勢な局面でも勝ち切るまでに多くの手数を要していること、そして 3 つ目は、戦型選択が偏り、いくつかの戦法の学習が行えていないことであった。本小節では、これらの問題について、筆者らが用いた解決策を記述する。

引き分け率の調整

予備実験では、教師データ内に、引き分けて終局している棋譜が少なくなかった。筆者らは、この原因は、攻める際に評価値が一瞬下がることが多いこと（将棋では攻める際に先に駒を取られることが多い）に由来すると考えた。特に (Novice の) 自己対戦では、MCTS のシミュレーション回数が少なく (具体的には 800 回程度)、探索が浅いため、この仮説は有力であると考えられる。

この問題を解決するために、生成している教

師データ内で引き分け率が2%を超えている場合には、その次に行う自己対戦の設定を、引き分けを先手勝ちまたは後手勝ちの設定で行うようにした。なお、教師データ内で先手勝率が50%を超えている場合には引き分けを後手勝ち、後手勝率が50%を超えている場合には引き分けを先手勝ちとした。

ニューラルネットワークがこれらの対局条件を区別して学習できるように、前述したようにニューラルネットワークに引き分けの際の点数（先手勝ちなら先手+1 かつ後手-1）を入力特徴量として与えた。

以上の工夫により、実際に生成される教師データ内の棋譜の引き分け率を常におおよそ2%に保つことに成功した。

様々な最大手数設定の対局

予備実験では、勝勢な局面でも、終局まで多くの手数を要している棋譜が多く発見された。筆者らは、この問題は、終局時に勝てる棋譜であれば、要する手数の長さを考慮していないことが原因であると考えた。

この問題を解決するために、様々な最大手数の設定で自己対戦を行うことで解決を図った。なお、最大手数に達した場合には、即座に引き分けとして扱った。この変更により、勝勢の局面でも、長く手数を要する勝ち手順には、引き分けとして終局する可能性が生じる。そのため、できるだけ短い手順で終局する手順が学習できることが期待される。また、副次的に、劣勢な局面では最大手数まで粘ろうとする手順を学習していることが期待される。

ニューラルネットワークが様々な最大手数に対応できるように、前述したようにニューラルネットワークに現在の手数および最大手数を入力特徴量として与えた。

初期局面集の使用

予備実験として、将棋の初期配置のみから学習を行うと、安定しない評価値が出力されたり、良い手を見落としていたりしていることが発見された。そこで、初期局面集を用意し、自己対戦による棋譜生成の際には、対局ごとにその中からランダムで1局面を選び、その局面を初期局面として使用した。初期局面集には、筆者らが選んだ約100局面を使用した。なお、初期局面集を用いるという工夫はやねうら王^{*1}やdlshogi^{*2}に既に用いられており [2, 3], 筆者らは大いに先行事例を参考にした。

^{*1} <https://github.com/yaneurao/YaneuraOu.git>

^{*2} <https://github.com/TadaoYamaoka/DeepLearningShogi>

表 1 ニューラルネットワークの入力特徴量

特徴量	チャンネル数
手番の駒配置	14
相手の駒配置	14
手番の持ち駒	24
相手の持ち駒	24
手番	2
王手	1
各列の歩の存在	2
現在の手数 / 最大手数	1
1.0 / 最大手数	1
引き分け点数	2
合計	85

2.2 ニューラルネットワークの構造と学習

本小節では, Novice の用いたニューラルネットワークの構造と学習手法について記述する.

ニューラルネットワークの構造

ニューラルネットワークの入力層は, 1 局面あたり, 85 チャンネル, 9×9 ピクセルで構成される. 各チャンネルが表現する特徴量を表 1 に示す. 盤上の位置に依存する情報は, 9×9 のうち, 対応する場所のみが指定される値となる. それ以外のものは, 9×9 ピクセル全て同じ値が指定される. 駒配置については, 先手後手を含めて駒種を区別した各チャンネルを用意し, 各駒の存在するマスに対応したピクセルに 1 を設定した. 持ち駒については, 上限を 4 枚とし, dlshogi と同じ手法 [4] で設定した. 手番には 2 チャンネル割り当てられており, 与えられた局面の手番が先手番の時はこの 1 チャンネル目のみが 9×9 全て 1, 後手番の時はこの 2 チャンネル目のみが 9×9 全て 1 と設定される. なお, AlphaZero では, 手番を表すためには 1 チャンネルのみが使用され, 後手番の時にそのチャネ

ルを 9×9 全て 1 にしていた. しかし, Leela Zero^{*3}によって, 先手番での局面を入力とするニューラルネットワークの実行の際に, ニューラルネットワークが盤面を把握しづらくなる可能性が指摘されている [5].

leela-zero/README.md より抜粋

There are 18 inputs to the first layer, instead of 17 as in the paper. The original AlphaGo Zero design has a slight imbalance in that it is easier for the black player to see the board edge (due to how padding works in neural networks). This has been fixed in Leela Zero.

そのため, Novice も Leela Zero と同様の手法で 2 チャンネル使用した. Novice では, 学習の効率向上を期待し, 現在の局面が王手されているかどうかと, それぞれの升の列に歩があるかどうかを追加の情報として与えた. また, 後述する理由のため, 現在の手数及び最大手数に関する情報と, 引き分けで対局が終了した際に先手後手がそれぞれ得られる点数を入力に与えた.

中間層は SE block [6] 有りの Residual Network [7] で構成した. なお, 12 blocks $\times$ 256 filters の構成を用いた.

最終的に, ニューラルネットワーク $f(\cdot; \theta)$ は 4 つの値を出力する (式 1).

$$\mathbf{p}(s), \mathbf{v}(s), \mathbf{d}(s), \mathbf{z}(s) = f(s; \theta) \quad (1)$$

ここで $\mathbf{p}(s)$ は局面 s での次の一手の確率分布, $\mathbf{v}(s)$ は { 勝ち, 引き分け, 負け } の確率分布, $\mathbf{d}$ は { 先手の宣言勝ち, 後手の宣言勝ち, それ以外 } の確率分布, $\mathbf{z}(s)$ は 9×9 の各マスごとに先手後手それぞれ利きがあるかどうかの予測である.

^{*3} <https://github.com/leela-zero/leela-zero>

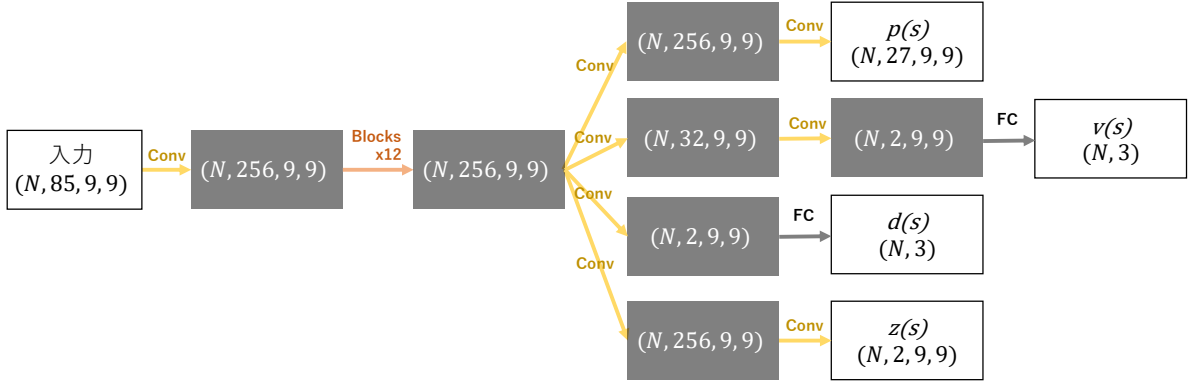


図1 ニューラルネットワークの概略図

ニューラルネットワークの学習

ニューラルネットワークの重み θ の学習は、自己対戦の棋譜を教師データとし、Stochastic Gradient Descent を用いて次のロス関数 $\mathcal{L}$ を最小化をすることで行った。

$$\begin{aligned}\mathcal{L}(s; \theta) = & -\mathbf{p}'(s) \cdot \log \mathbf{p}(s; \theta) \\ & -v'(s) \cdot \log v(s; \theta) \\ & -0.1 \cdot d'(s) \cdot \log d(s; \theta) \\ & -0.1 \cdot \mathbf{z}'(s) \cdot \log \mathbf{z}(s; \theta) \\ & -0.01 \cdot \mathcal{H}(\mathbf{p}(s; \theta))\end{aligned}$$

ここで $\mathbf{p}'(s), v'(s), d'(s), \mathbf{z}'(s)$ はそれぞれ、教師データの、MCTS の訪問回数の分布、対局結果、入玉結果、盤面の利きである。また $\mathcal{H}(\cdot)$ は方策のエントロピーであり、方策が確定的になるのを防ぐ効果を期待できる [8]。

2.3 詰み探索による強化

モンテカルロ木探索と合わせて、PV Mate [9, 10] と末端局面での5手詰め探索を用いている。PV Mate では、現在の探索の中で最も現れそうな進行に現れる各局面について、証明数探索 (Proof Number Search, PNS) [11] を用いて詰みの有無を探索している。なお、実装を簡単にするため df-pn [12] は未使用である。また、モンテカルロ木探索に現れる全ての局面で深さ5の深さ優先探索 (Depth First Search,

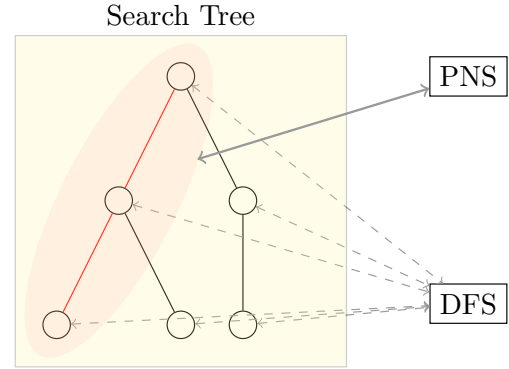


図2 詰み探索の概略図。赤い線はPVを表す。PV上の局面ではPNSによる詰み探索を行う。加えて、全ての局面でDFSによる詰み探索を行う。

DFS) を行っている。なお、PNSおよびDFSは、それぞれ専用のタスクのキューを持ち、モンテカルロ木探索とは独立したスレッドで動作させることで高速化をしている。

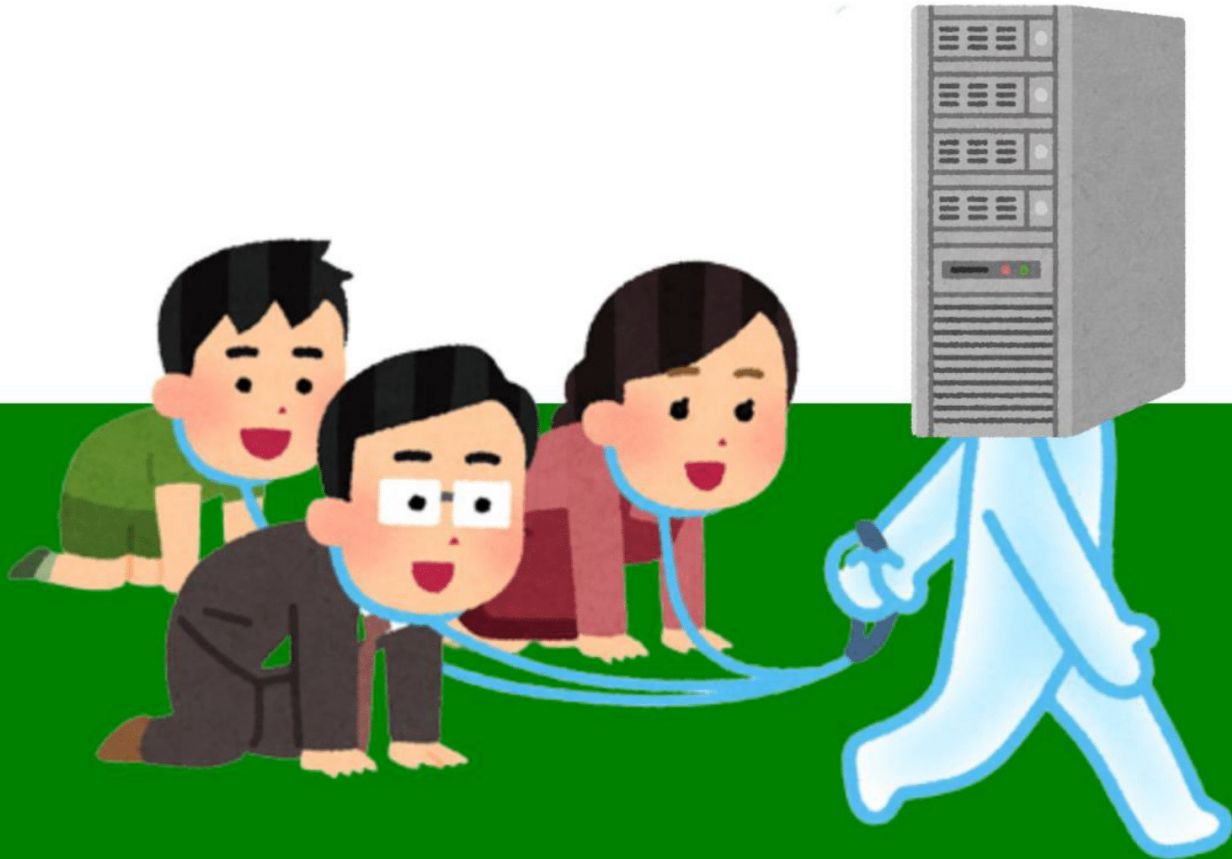
参考文献

- [1] David Silver, Thomas Hubert, Julian Schrittwieser, Ioannis Antonoglou, Matthew Lai, Arthur Guez, Marc Lanctot, Laurent Sifre, Dhharshan Kumaran, Thore Graepel, Timothy Lillicrap, Karen Simonyan, and Demis Hassabis. A general reinforcement learning algorithm that masters chess, shogi, and go through self-play. *Science*, Vol. 362, No. 6419, pp. 1140–1144, 2018.
- [2] Tadao Yamaoka. 将棋 AI の進捗 その 19 (初期局面集). <https://tadaoyamaoka.hatenablog.com/entry/2018/04/06/001745>. (2022/01/19 閲覧).
- [3] やねうらお. Yaneuraou/source/learn/learner.cpp. <https://github.com/yaneuraou/Yaneura0u/blob/cc30a4323fa323be29d8513bb44269aaa4b4d39a/source/learn/learner.cpp#L3421>. (2022/02/09 閲覧).
- [4] Tadao Yamaoka. 将棋 AI の実験ノート (入力特徴量の数値の表現方法). <https://tadaoyamaoka.hatenablog.com/entry/2019/06/12/230710>. (2022/01/19 閲覧).
- [5] leela-zero. README.md. <https://github.com/leela-zero/leela-zero/blob/e3ed6310d33d75078ba74c3adf887d18439fc2e3/README.md>. (2022/01/19 閲覧).
- [6] Jie Hu, Li Shen, Samuel Albanie, Gang Sun, and Enhua Wu. Squeeze-and-excitation networks, 2019.
- [7] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition, 2015.
- [8] Tadao Yamaoka. 強化学習におけるバッチサイズとエントロピー正則化. <https://tadaoyamaoka.hatenablog.com/entry/2019/05/10/234328>. (2022/02/09 閲覧).
- [9] Tadao Yamaoka. 第 2 回世界将棋 AI 電竜戦バージョン. <https://github.com/TadaoYamaoka/DeepLearningShogi/releases/tag/denryu2021>. (2022/01/19 閲覧).
- [10] Tadao Yamaoka. Merge feature/pv_mate. <https://github.com/TadaoYamaoka/DeepLearningShogi/commit/4f4d7eed209e7183384ee13425d4e3e42e1dcec5>. (2022/01/19 閲覧).
- [11] L.V. Allis. *Searching for solutions in games and artificial intelligence*. PhD thesis, Maastricht University, January 1994.
- [12] 長井歩, 今井浩. df-pn アルゴリズムの詰将棋を解くプログラムへの応用. 情報処理学会論文誌, Vol. 43, No. 6, pp. 1769–1777, Jun 2002.

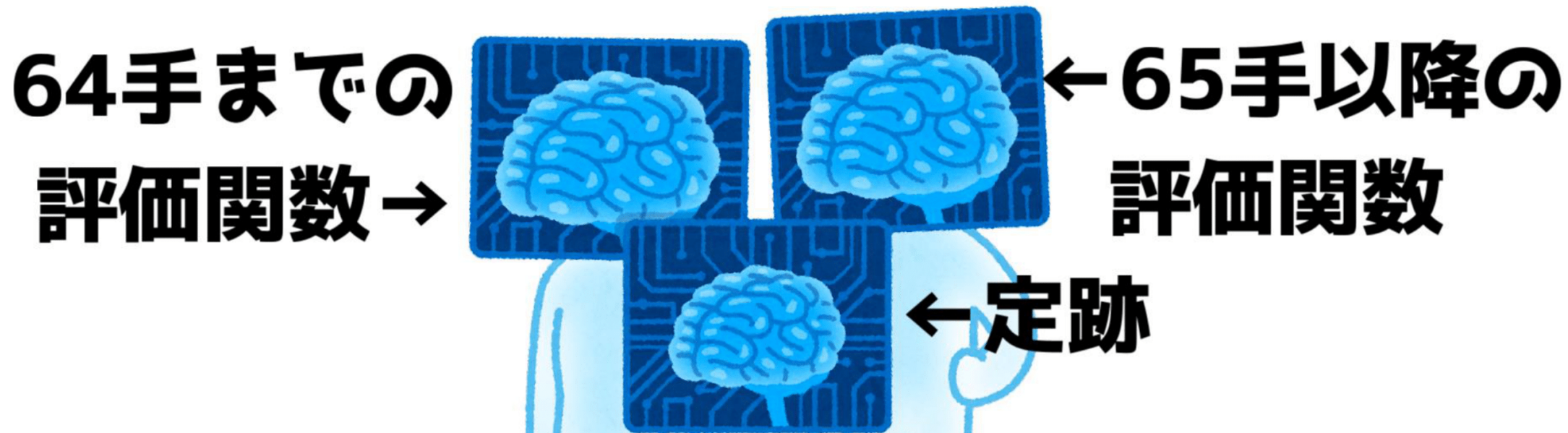
WCSC32 koronアピール文章

koron

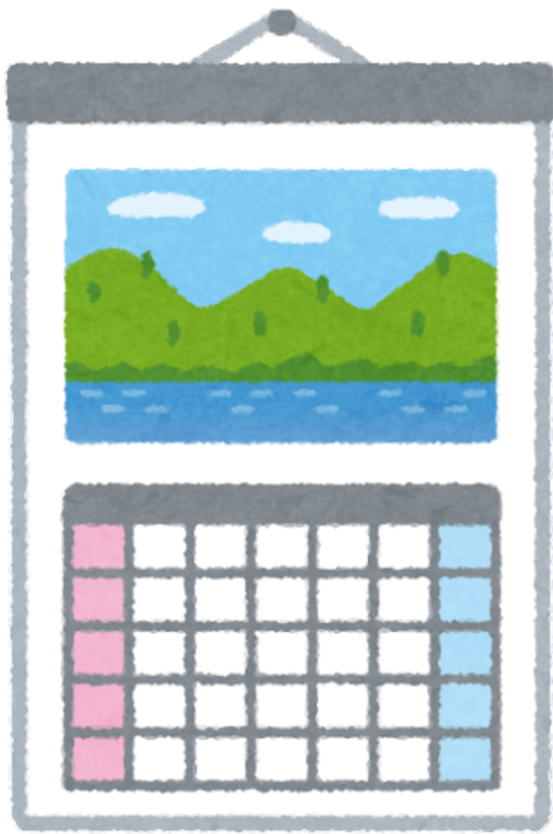
shogi AI



Q koronはWCSC32で何をするの？



A 64手で評価関数を入れ替えて
NNUEの表現力を増やせないかの検証



**GCP借りて会場参加の
予定でしたが受験生で
多忙なので
Ryzen 9 3950x使用の
オンライン参加に
変更しました**

名人コブラ アピール文書

—

松山洋章

概要

ディープラーニング評価関数とNNUE
評価関数をベースとしたアンサンブル
評価関数を使用します

評価関数

MultiPVで出力したdlshogiとやねうら
王の評価値や、局面情報等を入力特
徴としたスタッキング評価関数を作成し
ます。

使用ライブラリ

- dlshogi

局面評価のベースとして。

序中盤の局面評価に優れるため。

- やねうら王

局面評価のベースとして。

読みの速さに優れるため。

ソフト名の由来

劇団鋼鉄村松

「二手目8七飛車成り戦法」

登場人物より

参考:

https://stage.corich.jp/stage_main/23036

第32回世界コンピュータ将棋選手権

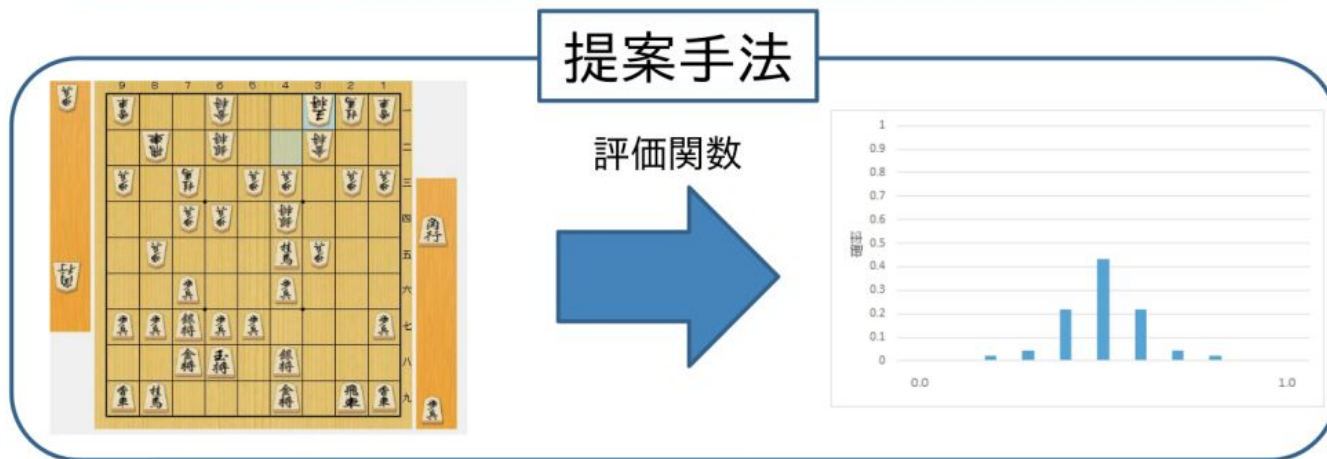
Miacisアピール文書

迫田真太郎

2022/03/29

基本構成と特徴

- 基本的にはよくある「DNN + MCTS構成」
 - [GitHub](#)
- 特徴



今年のテーマ

- とにかくネットワークを巨大化する
 - 理由
 - できるだけ探索に対する依存を脱したい
 - NPS増加による棋力の伸びはやや鈍化する傾向
(cf. [MuZero論文](#) Fig.3)
→低NPSで強いことが重要

Transformer(Vision Transformer)

- 大量ステップ数の学習を行うことで損失は改善
 - 低NPSを補えるほどではない
 - FP16でしか動かないのもあり、難航中

	Policy損失	Value損失	NPS	基準プログラムに対する勝率
ResNet (20block 256ch)	1.760	0.462	11742.3	41.8%
ViT (12block, 384ch)	1.740	0.460	5936.2	37.0%

ResNetの巨大化(ブロック2倍)

- 40ブロックモデルは上手く学習ができなかった

	Policy損失	Value損失
ResNet (20block 256ch)	1.760	0.462
ResNet (40block 256ch)	1.925	0.471

ResNetの巨大化(ch数2倍)

- ch数を2倍にしたモデルでは途中までのところValueだけは改善

※まだ学習が終わっていない

※NPSが $\frac{1}{2}$ ~ $\frac{1}{3}$ 程度になることは確認済み

1280000 ステップ時点	Policy損失	Value損失
ResNet (20block 256ch)	1.751	0.475
ResNet (40block 256ch)	1.799	0.471

所感

- 単純な巨大化でそれほど上手くいっている感じはしない
 - 性能自体は微改善
 - NPSの低下の方が痛い
- ブロック数の長大化での学習失敗を改善したい
 - [DropPath](#)
 - [DeepNet](#)

第 32 回世界コンピュータ将棋選手権 アピール文書
プログラム名「あやめ」
2022 年 4 月 27 日

昨年の大会からほとんど変わっていません。前回大会の[アピール文書](#)をご確認ください。

第 32 回世界コンピュータ将棋選手権 参加ソフト

ねね将棋 アピール文書

日高雅俊

2022/05/01

方針

iPad 上で深層学習ベースの将棋 AI を実装する。

世界コンピュータ将棋選手権において、Windows 以外のタブレット端末(iOS / iPadOS / Android)が用いられるのは初※。

概要

近年のハイエンドスマートフォン・タブレットには深層学習モデルを動作させることを目的とした専用の回路 (iPad では Neural Engine) が組み込まれており、この処理能力によりどの程度強力な将棋 AI を構築できるかを検証する。ソフトウェア面だけでなく、冷却等モバイル端末特有の課題についても考察する。

開発言語には iPad アプリのデファクトスタンダードである Swift を用いてスクラッチで実装予定。こちらも世界コンピュータ将棋選手権においては初使用。

Apple 製ライブラリ・評価関数以外は独自に実装する。

評価関数には書籍「強い将棋ソフトの創りかた」(山岡忠夫、加納邦彦著、マイナビ出版刊)で解説されている深層学習モデルを用い、iPad 向けに実行速度の最適化を試みる。

棋力は、floodgate で約 3300。おおむね人間のプロ並みとなった。レートの基準となる gikou2_1c に勝ち越すが、下位のソフトに頓死させられることも頻繁にみられる。

なお、前年までのねね将棋と実装上の連続性はない。

評価関数

「強い将棋ソフトの創りかた」のサンプルコードを用いて、畳み込み層 15 ブロック(30 層)、224 チャンネルのモデルを学習させた。

探索コアはシングルスレッドで、バッチサイズ 1 で評価している。約 330nps で動作する。バッチサイズ 2 以上にすることで数十%の NPS 向上が可能だが、読めるノード数が少ない状態でバッチサイズを上げると逆に棋力が低下する場合もあるため採用しなかった。

探索

MCTS 探索。思考開始局面のみ、1 手詰め探索を行う。探索したノードでの詰め探索は未実装。

思考時間管理

固定で 10 秒思考（サーバ上では 11 秒と計測される）。相手の手番中に思考する Ponder 機能あり。

定跡

ソフトウェア実装は別となるものの、評価関数・探索アルゴリズムは dlshogi と類似しており、棋風としては dlshogi やそのライブラリを用いたソフトの探索が浅い状態のものとなることが予想される。他のソフトと異なる棋譜を残すため、初手では 2 六歩、7 六歩および悪手を除外した 20 通りの手からランダムに指すこととした。

「第 2 期電王戦バージョンの ponanza が指さない初手は

▲8 六歩、▲9 八香、▲6 八金、▲5 八金左、▲4 八飛、▲3 八飛、▲1 八飛、▲1 八香の 8 通り。

それ以外の 22 通りの初手を指す。」

参考：https://twitter.com/floodgate_fan/status/851019316522762240

後手番になった際は 2 手目をランダムに選択する。

1 手目、2 手目の組み合わせによる、ねね将棋の評価値（3 手目を思考する先手番からみたもの）を示す。

先手	1四歩	2四歩	3四歩	4四歩	5四歩	6四歩	7四歩	8四歩	9四歩	1二香	3二香	3二銀	4二銀	6二銀	7二銀	8二銀	4二馬	5二馬	6二馬	7二馬	9二馬	3二金	4二金	5二金左	6二金左	6二金右	7二金	4二玉	5二玉	6二玉
1六歩	121	706	55	116	224	234	233	56	107	340	431	352	227	188	181	229	208	266	341	417	383	153	275	323	146	281	372	204	234	333
2六歩	234	1399	145	635	539	424	478	101	204	557	594	343	302	281	305	996	605	549	604	526	601	113	372	862	278	381	413	249	334	477
3六歩	-37	658	-59	133	231	149	159	-68	-66	301	344	188	128	105	39	161	240	310	371	388	354	-34	159	364	96	176	302	151	63	247
4六歩	53	649	-96	119	260	137	196	-84	35	315	380	284	99	114	67	251	164	267	318	331	344	-36	187	374	121	197	242	108	120	264
5六歩	-59	589	-62	46	214	1	11	-101	-33	224	241	86	27	-46	-44	258	260	302	373	384	380	-15	70	338	-31	57	106	-15	95	275
6六歩	20	650	43	50	103	-1	37	34	11	242	299	60	72	31	32	100	160	259	7	269	195	71	145	188	43	136	195	65	165	292
7六歩	230	683	144	450	368	404	316	114	189	523	559	339	297	301	327	276	318	358	466	475	480	127	317	649	314	408	440	297	361	479
8六歩	-388	414	-436	-70	-306	-337	-255	-336	-254	-31	2	-133	-199	-314	-383	60	-35	21	47	59	19	-301	-264	-26	-337	-245	-192	-437	-308	-53
9六歩	110	671	55	143	338	238	273	65	129	390	472	318	144	172	141	401	351	382	417	421	412	71	263	373	156	274	320	200	206	345
1八香	-200	359	-275	-126	2	-100	-93	-281	-185	84	126	10	-171	-134	-185	21	4	59	126	125	143	-240	-38	89	-144	-73	-16	-144	-142	41
9八香	-163	394	-162	-102	9	9	-44	-184	-185	90	157	87	-101	-130	-144	70	22	65	139	175	171	-134	-1	73	-134	-48	32	-116	-72	107
3八銀	101	747	10	102	306	221	243	6	75	348	441	315	145	169	124	218	154	287	351	387	359	32	281	396	162	263	361	211	188	321
4八銀	84	709	-22	120	314	184	216	-26	63	368	422	312	126	159	107	251	239	324	401	419	406	1	259	379	148	238	350	142	165	319
6八銀	26	758	15	46	192	137	71	66	-26	297	334	315	207	74	84	146	123	146	217	283	229	64	266	278	127	194	274	212	205	278
7八銀	59	800	-148	60	130	74	69	64	14	242	365	313	237	80	80	128	141	182	249	259	249	57	259	265	124	207	258	200	203	312
1八飛	-231	114	-198	-62	-154	-134	-136	-237	-188	12	56	-68	-79	-189	-164	-15	21	33	31	24	29	-96	-139	-10	-191	-99	-77	-208	-63	81
3八飛	-207	109	-228	48	-153	-144	-150	-258	-193	63	106	-73	-121	-186	-172	-25	71	141	88	90	93	-120	-157	49	-179	-77	-29	-204	-71	120
4八飛	-211	161	-187	146	-122	-131	-141	-211	-178	30	67	-49	-77	-190	-156	45	109	29	34	74	54	-69	-134	3	-199	-96	-70	-220	-65	93
5八飛	-143	114	-91	-58	-68	-81	-136	-135	-55	129	141	-7	-18	-122	-70	49	86	38	131	38	116	-30	-87	38	-135	22	49	-110	47	193
6八飛	-80	153	-60	-41	-69	3	-66	-63	-102	120	154	19	8	-62	16	56	48	57	62	64	56	38	-19	82	-65	69	120	-55	76	152
7八飛	-136	172	-90	63	-64	-61	18	-124	-60	109	176	54	7	-72	-51	64	95	147	109	152	121	2	-36	93	-56	38	63	-68	51	199
3八金	-31	543	-167	17	182	43	76	-165	-53	228	285	218	-6	-3	-26	109	63	162	239	237	234	-59	123	245	10	61	143	61	-14	138
4八金	8	602	-114	16	197	119	177	-130	3	253	360	230	54	79	47	161	112	206	285	283	260	-45	164	245	64	172	280	98	84	218
5八金右	88	693	34	123	317	171	227	30	81	361	433	318	124	154	135	257	231	327	382	409	376	38	263	367	143	259	335	135	176	340
6八金左	-148	455	-130	-14	34	-81	-119	-137	-137	128	137	112	-146	-141	-122	114	89	134	181	179	195	-180	-4	166	-126	-45	30	-123	-116	138
6八金	-36	679	-117	46	265	168	182	-115	18	286	395	172	122	119	29	215	171	268	334	378	343	-120	169	298	77	155	298	95	115	240
7八金	131	784	96	58	222	188	239	81	96	298	421	343	273	213	266	178	114	199	240	305	272	126	288	300	176	335	436	257	334	331
4八玉	-115	534	-161	-177	-68	-52	-52	-190	-157	55	161	85	8	-98	-46	1	-59	7	50	63	59	-169	1	32	-118	17	137	-19	-19	85
5八玉	52	728	-17	13	121	158	145	-46	16	246	341	246	142	110	113	171	95	160	251	269	235	14	186	239	62	233	326	165	168	244
6八玉	109	774	4	123	325	259	172	12	127	391	500	264	141	118	106	278	243	337	407	441	421	8	297	410	205	210	208	169	195	337

これをもとに、3 四歩、8 四歩以外の上位 3 手からランダムに選択する。

3 手目以降は定跡なし。

ハードウェア

iPad (第 9 世代、2021 年発売)を用いる。OS は iPadOS 15 である。iPadOS はほぼ iPhone 向けの iOS と同等である。

対局に必要な有線 LAN への接続は、Lightning-USB アダプタおよび USB・Ethernet アダプタを組み合わせて行う。

冷却

評価関数単独での連続動作、将棋エンジンとしての動作いずれの場合も明らかな発熱はなく、ファンによる強制空冷は不要と判断した。

GUI

一般に将棋エンジンを開発する場合、GUI として将棋所や ShogiGUI が用いられる。しかし iPad 上でこれらの GUI は動作しないため、独自に実装を行った。

スクリーンショットを以下に示す。



画面右中央で MCTS の探索木から複数の読み筋（候補手 3 手、それぞれに対し応手 3 手）を表示する点が特徴的である。勝率の%の横にその標準偏差を表示し、局面の複雑さを示す。棒グラフ状の背景色は、その指し手を検討した回数を示す。

なお、指し手を手入力する機能は実装していない。このような特殊な状況での対局が必要な場合は、開発機となる Mac 上で将棋所 Mac を動作させ、TCP 経由で USI プロトコルを用いて連携を行う。

※ CSA Web サイトの過去の出場チームリストにて確認。オンライン大会である電竜戦については、第 2 回世界将棋 AI 電竜戦本戦（2021 年 11 月）においてチームシャイニーのマシン環境として「Android スマホアプリ」との記載がある。

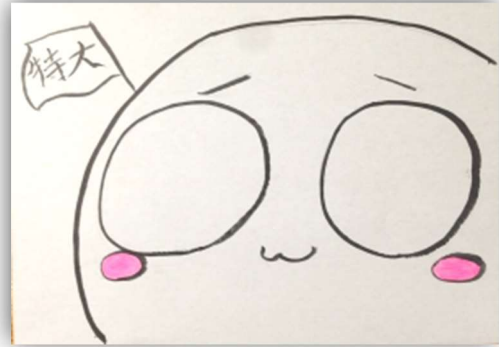
TMOQ アピール文書

2022 年 03 月 13 日 作成

2022 年 05 月 01 日 改訂

【ソフト名】TMOQ (特大もっきゅ)

“TMOQ”と書いて「特大もっきゅ」と呼びます、
愛娘が命名しデザインしたものです



【コンピュータ将棋大会実績】

2016 年の WCSC26 以降、ほぼ全ての大会に出場、
中位の成績をキープ

【TMOQ の特徴】

1. dlshogi ベース (WCSC29 より継続)
2. ネットワークに ResNet 9b を使用 (当初予定の GhostNet 15b は取下げ)
3. 学習データに「GCT の学習に使用したデータセット」+過去の TMOQ 棋譜を利用
4. 2000 万手を超える定跡 (Floodgate、電竜戦等の棋譜を寄せ集めたもの)
5. 『定跡チェッカー』別の USI エンジンが定跡手を評価・ダメ出し
6. 『入玉および駒得モード』MCTS の報酬に玉の縦位置および自分の駒数を考慮
7. 『入玉後ボタンタッチ』自玉が敵陣に入ったら、別の USI エンジンが宣言勝ちを目指す
8. Note PC を使用、莫大な計算資源がなくてもコンピュータ将棋は楽しめる！

【使用ライブラリ】

- ベースに「DeepLearningShogi」(Commit 790e2f4 on 3 Feb 2022) (GPL ライセンス)
- 『定跡チェッカー』および『入玉後ボタンタッチ』に「やねうら王 V7.00」(GPL ライセンス)を改悪して使用させていただきます

【御礼】

今回も山岡氏、加納氏&やねうらお氏を中心に、多くの方の公開情報により参加できました。
この場を借りて御礼申し上げます

第32回世界コンピュータ将棋選手権
なのは アピール文書 V0
2022年3月31日 川端一之

去年から なんの進歩も ありません

https://www.apply.computer-shogi.org/wcsc31/appeal/Nanoha/Nanoha_appeal2021v0.pdf

AobaZero のアピール文書

山下 宏

yss@bd.mbn.or.jp

1 AlphaZero の追試が最初の目的

AobaZero は Bonanza、LeelaZero のコードをベースに AlphaZero の追試をするべく MCTS + ディープラーニング で実装されてます。ネットワークは 3x3 のフィルタが 256 個の 20 block の ResNet でパラメータの個数は 2340 万個。棋譜生成をユーザの皆様と協力して行う分散強化学習です。オープンソースです*1。

2 AlphaZero の追試は 2021 年 4 月に終了

AlphaZero の将棋の追試は、2019 年 3 月から開始し、2021 年 4 月に 3900 万棋譜を作成して終了しました*2。その後、表 1 の変更を行っています。2022 年 3 月 30 日現在、5360 万棋譜を作成しています。

表 1 追試終了後の改良

日付	
2021 年 4 月	40 block に移行
2021 年 9 月	20 block に戻し温度を 1.0 から 1.3 に変えて序盤の変化を増やす
2021 年 12 月	Value の学習を「実際の勝敗」から「実際の勝敗」+「探索勝率」の平均に
2022 年 1 月	1 手 800payout 固定から 100~3200 と可変に
2022 年 2 月	ネットワークの構造を dlshogi 風に変更 利きの情報あり、ReLU を Swish に 30 手目までのランダム性を変更

3 40 block に移行

まずネットワークのサイズを 20 block から倍の 40 block に変更しました。囲碁の KataGo では +150 Elo ほどの効果があり、この程度の上昇を期待していたのですが実際は +40 Elo 程度でした。原因はいくつか考えられます。

- 学習が収束した状態、の棋譜から再学習したもので開始したせい。最初から 40 block だと違う？

- 学習が収束した状態で、30 手後のユニークな局面は 40% 程度 (100 万棋譜で)。同じような (相掛かりの) 棋譜ばかりで多様性がない。30 手後の細かい定跡の変化で強くなっている？
- そもそも将棋では 20 block 以上にしても効果が少ない？ 20 block で作った棋譜を学習させた 10 block は -120 ELO 弱い。
40 block +40 ELO 強い
20 block
10 block -120 ELO 弱い
- 40 block の再学習で学習率の下げ方が早すぎた (Cosine Annealing を 1 回だけ、の方が良かった？)。
- バグ

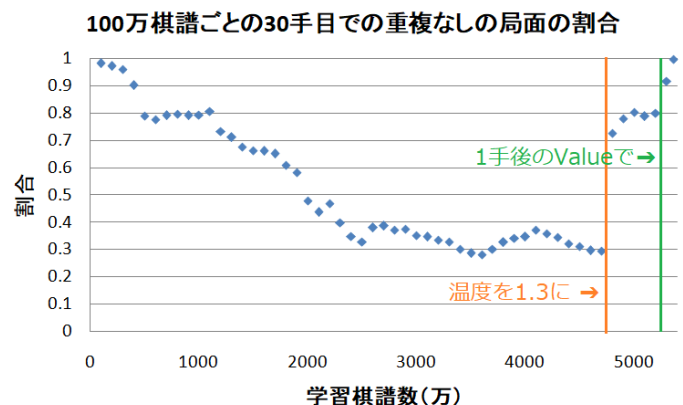


図 1 重複なしの (ユニークな) 局面の割合

4 重複局面を減らすために温度を 1.3 に

図 1 は学習棋譜の 30 手目で重複していない (ユニークな) 局面の割合です (100 万棋譜ごと)。学習開始時はほぼすべての棋譜がバラバラですが、徐々に同じ棋譜を生成するようになり、AlphaZero が学習を終了した 2400 万棋譜では 35% 程度まで下がります。具体的には初手から▲ 26 歩△ 84 歩の相掛かりの将棋ばかり指すようになります。重複を減らそうと 4700 万棋譜の時点で温度*3を 1.0 から 1.3 に変更して、初手から 30 手までは訪問回数が少ない手も選ばれやすいよ

*1 <https://github.com/kobanium/aobazero>

*2 <https://github.com/kobanium/aobazero/issues/54>

*3 温度→0 で訪問回数最大の手を、温度 1 で訪問回数の割合で、温度→∞ですべての手を均等に選ぶ

*10 最善の Value との差を diff とすると、 $1/\exp(\text{diff} \cdot 70)$

30 手までの手順が決まれば、少なくともその手を 1 回は探索するようにして、800playout 後に強制的にその手を選んでます。

これらによって棋譜の重複なしの割合は 99.1% とほぼ完全にバラバラになってます。他にすべてのノードで 3 手詰を調べるようにしました。これは +20 ELO 程度の向上でした。そもそも 1 手詰の形や 3 手詰の形は簡単なようで、NN が「覚えて」しまっています。これらの変更でも現状、棋力に変化はなく、正しい方向なのかは不明です。

9 人間の知識は使っていない、をおそらく継続

利きの情報の追加や 3 手詰などで AlphaZero からは離れてきましたが、まだ全体としては「人間の知識は使っていない」を継続していると考えています。

10 1 手 1playout で将棋クエストで 6 段に

AobaZero の w3880^{*11}が 1playout で将棋クエストの :FuriJirouBot というアカウントで長考 (持ち時間 10 分) で 6 段 (2250 点) になっています。名前から推測できるように振飛車しか指さない設定です^{*12}。序盤はユーザ棋譜から作った定跡で振飛車を選ぶようになってるそうです。現在の AobaZero は振飛車を指さないのですが (初期の 400 万棋譜時点では後手のみ四間飛車を指していました) それでも学習棋譜では時々出てくるのでそれなりに指せるようです。

floodgate で 1 手 1playout は 2150 点 (w3392) でした。将棋クエストの長考、のレートはほぼ一致するようです。floodgate は対戦相手が偏っているのと勝率が低い (9 勝 83 敗、勝率 0.10)、アンカー (3300 点) から離れてる、でやや不確かですが。

余談ですが囲碁の KataGo も 1playout で囲碁クエストで :Katago1pBot で動いており下の段位になっています。

- 9 路で 2470 点 (7 段)
- 13 路で 2820 点 (9 段)
- 19 路で 2750 点 (9 段)

11 NN は内部で探索してる？

まったく探索なしで局面を NN に与えて返ってくる最善手を指すだけでこれだけの強さがあるのは驚きです。ただ探索してないとはいえ、NN の内部ではおそらく複雑な if 文の

組み合わせで疑似的な探索をしている (例えば、この形は角を切って同玉なら頭金で詰む形なので角を切る手の確率が高い、みたいな) ので、まったく読んでいない、というのはやや語弊があるのかもしれませんが。

12 詰を読むと水平線効果の無駄な王手をするように

また棚瀬さんから指摘されたのですが、3 手詰を読むようにして学習させた NN の重みは 1 手 1playout でも負けの局面で無意味な王手をするようになりました。これは探索部が負けの局面では王手以外の手を指さなくなり、NN もそれを学習したためです。きれいな形作りには投了直前は人間の棋譜からのみ学習させる、などが必要なのかもしれません。

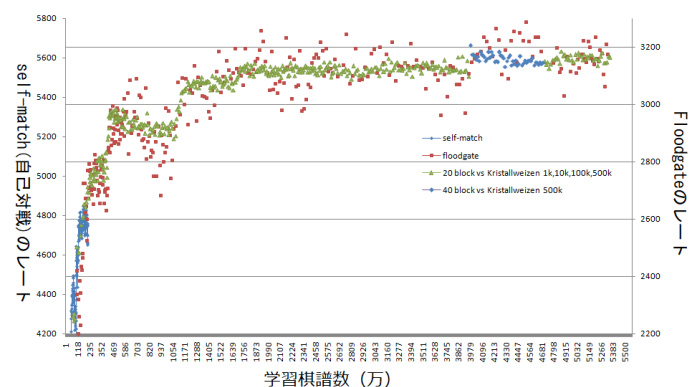


図3 AobaZero の棋力の推移。右軸が floodgate レート

13 で、強くなったの？

去年からほとんど強くなっていません。強くするのは大変ですね……。図3がELOの推移です。

14 3 年で 5300 万棋譜

5300 万棋譜、という膨大な棋譜を 3 年間で生成してきました。棋譜生成に協力していただいてる皆様に感謝いたします。

^{*11} w3880 で 3880 番目に作成された重み、を意味します。w3880 は利き情報を使ってない最後の重みでもあります。

^{*12} 将棋クエストはトライルール、AobaZero は宣言なので入玉になりにくいように振飛車に、とのことでした。

手抜きについて

手抜きチーム*

2022 年 3 月 31 日

手抜きは CSA プロトコルで対局を行うコンピュータ将棋プログラムです。開発者らが将棋プログラムの仕組みを理解するために作っています。

リポジトリ：<https://github.com/hikaen2/tenuki-d>

- 初段から二段くらいの棋力
- NNUE
- $\alpha\beta$ 探索
- D 言語

使用ライブラリ

- 『どうたぬき』(tanuki- 第 1 回世界将棋 AI 電竜戦バージョン) の評価関数ファイル nn.bin^{*1}

TODO

- 差分評価
- 千日手チェック
- 詰将棋
- ponder

* 鈴木太朗 (Twitter: @hikaen2), 玉川直樹 (Twitter: @Neakih_kick)

^{*1} <https://github.com/nodchip/tanuki-/releases/tag/tanuki-denryu1>



十六式いろは煌(きらめき)



自己紹介

末吉 竜介「学生の皆さんと、いろいろと学び試し楽しみながら作っていければと思います。」

岩田 大夢「AIを学んで早一年...まだまだ分からないことだらけですが、この大会を通じて成長していきたいと思っています！」

紺野 誠「将棋の知識を活かしていきたい」

若月 翔威「この大会を通して成長できるようにがんばりたいです」

小林 優輝「手取り足取りみんなで頑張っていきたいと思います！」

上野 勇樹「精一杯頑張りたいと思います！」

村越 小友梨「AIと将棋どちらも上達できるように頑張ります！」

長谷部 太一「大会を通して何か得られるものが有ればと思っています！」

畑中 慎吾「わからないことばかりですが、大会を通して経験を積み成長に繋がればと思います」

平沢 蒼「将棋もAIも未熟者ですが、精一杯頑張ります。」

李 愚昶「将棋AIは初めてでまだまだですが、上達できるようがんばります！」

若松 萌生「ひとつでも多くのことを学び、成長出来たらと思います」

山内 伊織「...ノーコメント」

風間 俊介「...ノーコメント」



「十六式いろは 煌」の由来

様々な名前の候補が上がり最終的に決まったのが考え始めてからなんと **1か月**かかりました！。

皇(すめらぎ)、煌(きらめき)、日本工学院、かまトウ(学校のマスコットキャラ) ...などなど。

「日本工学院の名前があった方がよいのではないか」や「ローマ字で書いた方がカッコいい！」

などかなりの意見などがありました但最终的には末吉先生の「十六式いろは」と生徒達で考えた「煌(きらめき)」を組み合わせで決定しました！



ソフトの概要

採用予定

- dlshogi
- やねうら王系
- 詰将棋(※)

(※)以下のいずれか

- 脊尾詰
- なのは詰め
- KomorningHeights

ソフトの説明

- 2～3台のマシンのクラスタ化(予定)。
- 合議制を採用。
- 持ち時間の管理を工夫(予定)。(詳細:序盤や深読み無しで指せる手はすぐに指して、熟慮が必要な場面に時間を多く割けるようにする。)
- floodgate、AobaZeroの棋譜に加えて、2022年3月現在で公開されている将棋ソフト同士の対戦の棋譜を用いて、dlshogiの機械学習をさせる。
- Optimizerの再選定。

第3 2回世界コンピュータ将棋選手権「水匠」アピール文書

令和4年3月31日

杉 村 達 也

第1 使う他者作成プログラム

やねうら王→探索部、学習部を使用させていただきます。

dlshogi→探索部、学習部を使用させていただきます。

第2 評価関数の学習

1 NNUE評価関数の学習

NNUEにはFV_SCALEという値があり、この値を変更しながら学習させると、評価関数の評価値スケールや勝敗項に関するlossの減り方、勝率項に関するlossの減り方が変わります。水匠評価関数の追加学習は、NNUE学習パラメータのlambdaを0、すなわち勝率項を見ずに勝敗項だけで学習させているので、FV_SCALEをいい感じに変更して学習させることが大切です（なぜ学習が上手くいっているのかはよくわかりません。）。

教師局面は水匠/やねうら王で1手約200万ノード探索させて対局させた棋譜を使用しています（現在約3億局面）。

2 dlshogi評価関数の学習

dlshogi評価関数は15ブロック224フィルタのモデルを学習させています。先ほどの水匠で作成した教師局面と、dlshogiで作成した教師局面をepoch毎に交互に学習させると徐々に強くなるので、半年前くらいから学習を継続しています（なぜ学習が上手くいっているのかはよくわかりません。）。

第3 その他

1 探索部の改善

最新のStockfishのコードで強くなる部分をマージします。

2 定跡の作成

- ① 任意の局面を与え、その局面から連続対局
- ② 特定の局面からの対局数が一定数(例:>100)を超えたら枝局面を開始局面として以下同様に繰り返す。
- ③ 最終的には勝率を用いてmin-max的な手法で、リーフ局面からルート局面までを伝播させる。

という手法で定跡を作成します。

3 楽観合議

学習させたNNUE評価関数とdlshogi型評価関数の両者とも使いたいので、評価値のスケールをいい感じに調整した上で、楽観合議をさせる予定です。

以上

prelude アピール文書

谷合廣紀, 中村朋生

2022 年 3 月 28 日

1 独自の工夫

基本は dlshogi/ふかうら王などのいわゆる DL 系で採用されている、policy+value Network + MCST のアプローチを取っています。dlshogi/ふかうら王と大きく異なるのは、モデル構造とその入出力です。

1.1 モデル入力のエンコード

モデルへの入力を画像的エンコード (9x9xfeatures) ではなく、文字列的エンコード (95xfeatures) としています。具体的には 1 一、1 二... 9 九の駒と先後の持ち駒を並べた 95 字を、自然言語処理と同様の枠組みで処理しています。基本的なアイデア・モチベーションは拙作の bert-mcts と同じで、広域的な相関を浅い層のうちから取ることがひとつの狙いです。

1.2 モデル構造

扱うデータが画像的な 2 次元ではなく 1 次元のため、MLP のみで構成した独自のモデルを構成しています。

1.3 モデル出力

dlshogi/ふかうら王で採用されている policy の出力と比べて、構造的かつ無駄のないラベルエンコーディングを導入して、学習が進みやすいように工夫しています。value 側では、損失関数に互角局面に近いほど重みを大きくするなどの工夫を施しています。

2 使用ライブラリ・使用データ

- ふかうら王
- 「強い将棋ソフトの創り方」公開データ

2022.3.31

神田 剛志

■開発コンセプト

名前の通り軽く速く。（末尾のEFはDNNのモデル(EfficientNet)が由来です）
ローカルPC環境でもCPU系/GPU系に限らず、つよつよインスタンス勢や
高級スリッパと戦うことがLightweightシリーズの目的です。

■アピールポイント

・ DNNモデルの改良

本家のResNetをEfficientNetで再構築し、1 から学習しなおしました。
第2回電竜戦時のモデルから6層追加した改良型に当たります。

・ USIエンジンのパラメータ設定変更によるNPS向上

GPUに局面を渡す際のバッチサイズを1024に上げています。
これと軽量なモデルと組み合わせにより、ローカルPC環境での平均NPSを向上
させています。

・ GCT学習データによる教師あり学習とLightweight-EF自身による強化学習

dlshogiチームが公開してくださっている学習データとLightweight-EFの
自己対局データを用いた強化学習を実施しています。
またこの際、学習時のバッチサイズをあげる事で学習の安定化を図っています。
（学習データを無償で公開されている山岡さん、加納さんには感謝申し上げます）

・ UCB選択アルゴリズムの一部変更（NPS向上のための簡易枝刈りの実施）

PUCTアルゴリズムに従って探索木を降りていく際、各子ノードの着手確率を
利用した簡易的な枝刈りを実施することで、最大UCB値の子ノード選択処理に
かかる時間を短縮しています。

・ 定跡の使用

初期局面の事前探索結果を利用することで、持ち時間の消費を抑えます。

■使用ライブラリ

dlshogi：自己対局データ生成・探索部・モデル学習・定跡作成に利用
Gikou2 ：検証時のテスト対局に使用
Suisho3：検証時のテスト対局に使用
Suisho4：検証時のテスト対局に使用
elmo for learn：学習データ作成に利用

NNUE評価関数と定跡の極北を目指しています。

- ・従来の強化学習手法に加え独自のNNUE評価関数の強化
- ・昨年優勝のelmoとほぼ同じ手作業による定跡の生成。特に後手番の定跡の強化に主眼を置いています
- ・探索部はやねうら王、いわゆるやねうら王チルドレンです。やねうら王+最新のStockfishのキャッチアップを予定

よろしくお願いいたします。

・役割論理とは

コンピュータ将棋においてついさっき確立された論理ですなwww

定跡とNNUE型評価関数と高級スリッパの圧倒的火力によって評価値ダメージレースを制する必勝の戦術ですぞwww

・定跡について

現状はほぼ先手必勝ですなwwwもちろんS振りではありませんぞwww

千日手模様になりやすい角換わりより青野流の方が深堀りされているようですなwww

独自の千日手回避手段を用いる場合もあるようですなwww

まずは後手番でどこまでダメージを最小限にするかが重要ですぞwww

互角の分かれならあとはNNUE型評価関数の終盤の強さと高級スリッパの超火力で踏み潰すだけですなwww

現状は後手番ではs-book_blackに少し分が悪いですぞwww

もちろん定跡を整備しないとDL系や水〇、や〇うら王などの強豪には手も足も出ませんなwww

相手が強豪で後手番ならもう千日手で逃げてしまう手もありますぞwww千日手専用定跡wwwありえないwww

そもそもプロ棋士でもないのに何でこんなことをする必要があるんですかなwww

・評価関数について

最低限水匠5より強くないと話になりませんなwwwただし評価関数の相性もあるので難しいですぞwww

もちろん振り飛車評価関数は総合的にロジックするまでもなくありえないwww

・必然力とは

論者を圧倒的勝利に押し上げる力ですなwww

絶対に先手を引く、スイス式の当たりが良い、裏街道最高wwwなどは必然力とされていますなwww

ヤーティ神への信仰によって得られる加護とされていますぞwww

現状これが最も重要なファクターとなっていますなwww

・ムックとは

ロジカル語法を使うもののヤーティを使わず、役割論理という戦術を理解・実践しない者を指しますなwww

つまりこのアピール文そのものですなwww完全に異教徒ですなwww

・さいごに

結論としては「評価関数の強化」と「定跡の強化」というごく当たり前の結論に至る訳ですなwww

最低限s-book_black対策くらいはしておきたいところですが難しいですなwww