

NeuroLGP-SM: A Surrogate-assisted Neuroevolution Approach using Linear Genetic Programming

Fergal Stapleton¹[0000-0002-5347-1573], Brendan Cody-Kenny² and
Edgar Galván¹[0000-0001-8474-5234]

¹ Naturally Inspired Comp. Res. Group, CS Department, Hamilton Institute, Maynooth University, IE
fergal.stapleton.2020@mumail.ie, edgar.galvan@mu.ie
² Sema, Baltimore, MD, USA
brendan@semasoftware.com

Abstract. Evolutionary algorithms are increasingly recognised as a viable computational approach for the automated optimisation of deep neural networks (DNNs) within artificial intelligence. This method extends to the training of DNNs, an approach known as neuroevolution. However, neuroevolution is an inherently resource-intensive process, with certain studies reporting the consumption of thousands of GPU days for refining and training a single DNN network. To address the computational challenges associated with neuroevolution while still attaining good DNN accuracy, surrogate models emerge as a pragmatic solution. Despite their potential, the integration of surrogate models into neuroevolution is still in its early stages, hindered by factors such as the effective use of high-dimensional data and the representation employed in neuroevolution. In this context, we address these challenges by employing a suitable representation based on Linear Genetic Programming, denoted as NeuroLGP, and leveraging Kriging Partial Least Squares. The amalgamation of these two techniques culminates in our proposed methodology known as the NeuroLGP-Surrogate Model (NeuroLGP-SM). For comparison purposes, we also code and use a baseline approach incorporating a repair mechanism, a common practice in neuroevolution. Notably, the baseline approach surpasses the renowned VGG-16 model in accuracy. Given the computational intensity inherent in DNN operations, a singular run is typically the norm. To evaluate the efficacy of our proposed approach, we conducted 96 independent runs spanning a duration of 4 weeks. Significantly, our methodologies consistently outperform the baseline, with the SM model demonstrating superior accuracy or comparable results to the NeuroLGP approach. Noteworthy is the additional advantage that the SM approach exhibits a 25% reduction in computational requirements, further emphasising its efficiency for neuroevolution.

Keywords: Neuroevolution, Linear Genetic Programming, Surrogate-assisted Evolutionary Algorithms

1 Introduction

In the preceding decades, Evolutionary Computation (EC) [10] has attracted considerable attention in the burgeoning field of Neural Architecture Search (NAS) [12]. The combination of these two fields, commonly referred to as neuroevolution, is centred around automatically finding architectures of artificial neural networks (ANNs) by evolving network topologies, hyperparameters and/or weights. Typically, the architectures are evolved based on accuracy, but often consider other design characteristics, such as network latency [9]. Neuroevolution of deep neural networks (DNNs) [22], has seen growing interest in recent years as discussed in a recent survey [12], reviewing over 170 works, and has applications to state-of-the-art and emerging technologies such as in the field of autonomous vehicles [15,30].

When considering DNNs, in general, whether using random search [2], neuroevolution [12] or other approaches, there is a high computational overhead required in order to effectively search

for well-performing architectures. To tackle this, research has turned to the use of surrogate-assisted evolutionary algorithms (SAEAs) [19]. SAEAs can be used to estimate the fitness of expensive DNNs without requiring them to be fully trained. To build an effective modelling strategy for fitness estimation, we need to impute the accuracy of untrained or partially trained networks based on knowledge from previously fully-trained neural networks. However, comparisons between models often require ingenious encoding strategies [32] to compare genotypes as generally, it is not possible to create distance metrics for differing network topologies [31]. A more natural approach is to consider the behaviour of the DNN architectures rather than relying on genotypic information [18,31], however, to date, SAEAs based around phenotypic distance metrics have focused primarily on ANNs of only one to two layers of fully-connected layers [31]. Another limitation of this previous work is that Kriging, which is an interpolation technique widely used in SAEAs and is well suited for optimization problems [3], suffers computational overhead when using high-dimensional data which puts further restrictions on the size of phenotypic distance vectors that can be used.

The primary objective of this paper is to provide insight into the effective integration of surrogate models into neuroevolution, addressing a significant challenge posed by the commonly encountered high-dimensional data. The key contributions of this work are:

- Firstly, we implement and validate a baseline model employing a repair mechanism, a strategy frequently utilised in neuroevolutionary methods. This model surpasses the well-established VGG-16 model, setting a high-performance benchmark for comparison with our proposed approaches.
- Secondly, we introduce an innovative representation based on Linear Genetic Programming [5], termed NeuroLGP, facilitating the automatic discovery of well-performing DNNs that outperform the baseline model. The NeuroLGP approach is employed to compute the fitness of the entire population and serves as an excellent method to compare against a surrogate model (SM).
- Thirdly, to address the challenge of high-dimensional data and enable the use of a SM to reduce computational demands in neuroevolution, we employ Kriging Partial Least Squares [4]. The fusion of this technique with NeuroLGP results in our second proposed approach, referred to as NeuroLGP-SM. This approach consistently identifies DNNs that perform similarly to those discovered by NeuroLGP, simultaneously reducing fitness evaluations and training time required for these DNNs. In our previous work [28], we identified that it was not possible to use the original Kriging approach to this end.
- Fourthly, we demonstrate the effectiveness of our approach through a robust comparison involving 8 independent runs for each of the 2 aforementioned NeuroLGP approaches plus the baseline model, using 4 challenging image classification datasets. This extensive evaluation, deviates from the norm of a single run in the context of DNNs.
- Our fifth key contribution lies in the innovative management of the SM. This approach remains invariant to varying network topologies and robust to data augmentation techniques. Consequently, we can train our networks with a significantly reduced number of instances while maintaining the ability to generalise effectively to unseen data.

Section 2 details related work, Section 3 discusses background. Section 4 describes the NeuroLGP method along with the surrogate model in detail, Section 5 details the experimental setup, Section 6 offers analysis of our results and Section 7 offers concluding remarks.

2 Related Work

2.1 Neuroevolution using Surrogate Models

An old, but still highly relevant survey on surrogate models by Jin [19], covers some of the fundamental aspects of surrogate-assisted EAs, such as surrogate model management strategies and acquisition functions. A more recent survey by Khaldi and Draa [21], covers more state-of-the-art concepts, such as considering the computational complexity of a surrogate model when performing neuroevolution. Next, we will cover some of the most relevant works.

Performance prediction can be categorized into two branches: (i) approaches that infer the performance of unseen networks based on specific characteristics of previously evaluated networks and which typically have been fully trained, and (ii) approaches that instead used partially trained information to predict the future performance of a network.

A major work in recent years of the former approach is the End-to-End Performance Predictor (E2EPP) [32] proposed by Sun et al. This approach uses an offline surrogate model based on random forests to search for optimal Convolutional Neural Networks (CNN) architectures. The CNN is cleverly encoded such that it maps to a numerical decision variable, which can then be processed by a random forest surrogate. This approach alleviates restrictions found in other approaches such as assuming a smooth learning curve and not requiring large amounts of training. The authors approach was shown to speed up fitness evaluations while also achieving the best performance compared to other peer performance predictors.

Approaches falling into the second category include the Freeze-Thaw Bayesian Optimization (FBO) technique proposed by Swersky et al. [34]. It uses Bayesian optimization to decide if a partially trained network should be trained to completion. The fundamental idea is that a human expert is quickly able to assess if a network is likely to result in poor performance and as such can decide to halt training. Based on this notion the authors designed an approach that at its core is a form of performance prediction, where Bayesian Optimization is used to determine whether a particular partially trained network will yield preferable results compared to other networks. Unlike E2EPP, this approach does require a smooth learning curve which can be hampered if different learning rates are considered. Of interest, is that the FBO approach relies on the phenotypic behaviour of the network in order to decide which networks to evaluate.

Gaier et al. [11] devised a kernel-based surrogate model to use with NEAT [26] referred to as Surrogate-assisted Neat. To overcome the limitation of variable length genotypes, they exploited a distance metric inherent in the NEAT algorithm. This distance metric which is known as the ‘compatibility distance’ was originally designed to help promote diversity amongst the network topologies through speciation. Within this work they also use it as a distance metric to use for Gaussian Processes. They demonstrated that they were able to achieve similar performance to NEAT with much fewer evaluations.

Greenwood and McDonnell [17] proposed a grammar-based approach which generates a tensor representation of variable length DNN topologies. An advantage of using formal grammar as a representation is that they can be designed to ensure validity of models by describing the space of allowable topologies. Their approach modifies the DeepNeat algorithm first proposed by Miiikkulainen [23]. The modelling strategy of this work is designed around a two-phase approach. Firstly, during the initialization phase, the DeepNeat algorithm is used to evolve the population of neural networks and these models are used to formulate the surrogate. Secondly, the active learning phase is implemented where new networks are evaluated using the surrogate model, along with a subset of networks that are fully trained and evaluated to further inform and improve the surrogate model. They noted a five times improvement in compute time.

While NeuroLGP and NeurpLGP-SM are presented here for the first time, an analysis of the architectures discovered by these two approaches has not been presented in this work but

has subsequently been addressed in our upcoming paper [29]. Additionally, this work includes additional analysis of the surrogate model and energy consumption of the two approaches.

3 Background

3.1 Surrogate assisted models: Kriging and Kriging Partial Least Squares

The aim of a surrogate model, also referred to as a meta-model, is to sufficiently approximate an estimate of the fitness values of a potential solution, while reducing the run time of the evolutionary process, compared to the run time of using the real fitness function alone. This requires that the surrogate model is well-posed and requires a robust model management strategy, otherwise, the evolutionary algorithm may converge to a false optimum [19].

Kriging allows for interpolation by assuming spatial correlation exists between known data points. The process itself is informed by the distance and variation between these data points. More explicitly, a kernel function $K(\cdot)$ denotes the spatial correlation between two samples x_i and x_j as shown in Equation 1,

$$K(x, x') = \prod_{i=1}^m \exp(-\theta_i (x_i - x'_i)^2) \quad (1)$$

where the θ parameter determines the rate at which the correlation decays to zero and m is the dimension. The θ parameter is determined by the Maximum Likelihood Estimator (MLE). One drawback, however, is that for large m , the cost increases significantly since the MLE algorithm requires calculating the inverse of the correlation matrix many times.

Kriging Partial Least Squares (KPLS) seeks to address this by effectively reducing the number of parameters calculated [4]. It does so using partial least squares which project the high-dimensional data into a lower dimension using principal components. Equation 2 details the KPLS kernel,

$$K(x, x') = \prod_{k=1}^h \prod_{i=1}^m \exp(-\theta_k (w_{*i}^{(k)} x_i - w_{*i}^{(k)} x'_i)^2) \quad (2)$$

where w_* are rotated principal directions which maximize the covariance and are a measure of how important each principal component is. The number of principal components $h \ll m$ which allows for the substantial improvement in computational cost associated with the KPLS approach. For full details of the method, please refer to [4].

3.2 Phenotypic Distance

The rationale behind employing phenotypic distance draws inspiration not only from Stork et al.'s work [31] but also from our research in Genetic Programming (GP) [14,16], focusing on semantics. Specifically, our contributions in the realm of semantic distance metrics include advancements in underexplored domains, such as multi-objective optimisation [13,14,16,27]. In traditional GP, semantics refer to the behaviour of a program given a finite set of inputs.

In the context of neuroevolution, we can define a solution sample x as having the semantic or phenotypic behaviour of the i^{th} program such that $x_i = s(p_i)$, where the semantics $s(p)$ of a program p is the vector of values from the output. From Equation 1 we can then define our phenotypic distance d as

$$d(x_i, x_j) = d(s(p_i), s(p_j)) \quad (3)$$

As such the distance metric is dependent on the outputs of each network. In this work, we use the convention established by Stork et al. [31], where the x_i is a flattened vector containing the output of the nodes at the final layer for all data instances. As such, our phenotypic distance vector length is given as the number of images of the validation dataset times the number of classes. It is important to note that this approach can be extended to any deep learning model architecture that can have its output represented in vectorised form, such as transformers, however, further research would be required to determine the limitations and scalability of applying this approach to other deep learning approaches.

4 Methodology

4.1 NeuroLGP

With the original LGP paradigm [5], representation is inspired by the von Neumann architecture of modern CPUs which use registers for storing or modifying small amounts of data. The content of each register in this architecture can be changed using instruction operations. An instruction, in this context, has three main components: an *operand*, *registers* for which the operand operates on (one register for 2-register instruction and two registers for 3-register instruction) and the *destination register* which stores the computed results. In the context of LGP, genetic operations work by either modifying these registers or the instructions that operate on them. While the original approach of LGP uses registers to store small amounts of data or instructions, it is possible to abstract this form so that the registers are themselves pointers to much larger amounts of data stored in memory. Our proposed approach does as such, using these pointers to instead control the flow of the initial and intermediary data from each outputted layer of our evolvable DNN. The left of Figure 1, demonstrates psuedocode how the expected representation would look, where $r[i]$ denotes the i^{th} register. Each line of code is executed imperatively. This example contains both effective and non-effective lines of code. The non-effective lines of code are commented out and, as such, are not compiled. The register $r[0]$ is a specially designated register for the final output of the program. To the best of our knowledge, LGP has not yet been used for neuroevolution.

<pre>def neuroLGP (...) { r[0] := Conv(r[1]) // r[4] := BatchNorm(r[3]) r[5] := MaxPool(r[0]) r[11] := BatchNorm(r[5]) r[0] := Dense(r[11]) ... }</pre>	<pre>input = tf.keras.Input(shape=(input_dim)) x = tf.keras.layers.Conv2D(32, 3, ...)(input) # x = tf.keras.layers.BatchNormalization()(x) x = tf.keras.layers.MaxPooling2D(3)(x) x = tf.keras.layers.BatchNormalization()(x) output = tf.keras.layers.Dense(x) model = tf.keras.Model(output)</pre>
---	---

Fig. 1: *left*: NeuroLGP psuedocode for python. *right*: Functional API example in TensorFlow.

To understand how this representation can be useful for the case of neuroevolution, some example code in TensorFlow is given to demonstrate the imperative nature of defining models as seen in the right of Figure 1. This code demonstrates an example of a model definition using the functional API from TensorFlow. On Lines 3, 5 and 6, to the left of the assignment, a variable x will hold the outputs of each statement and for Lines 5 – 6, to the right of the assignment, x is passed to each layer. As such, the variable x holds transient data which updates as each line is executed. The aim is to replace x using abstract or virtual registers (i.e $r[0]$, $r[1]$, $r[2]$...etc),

where the x variables to the left of the assignment are represented using a *destination register* and values to the right represent *registers* to be operated on. Line 4 is not executed.

4.2 NeuroLGP with Surrogate Model (NeuroLGP-SM)

To determine which individuals should be fully evaluated, we use the expected improvement (EI) criteria [6,20]. EI is an acquisition function where the next candidate solution to be evaluated is based on the expected improvement over the current best solution so far. As such, EI allows us to evaluate solutions from regions of the search space likely to see the best improvement. To estimate the fitness using the surrogate model, first, we evaluate the fitness based on a limited number of epochs. From here, we split our population based on the EI criterion, whose parameters are informed by the surrogate model. Based on the EI criterion we fully evaluate a subset of the population, leaving the remaining proportion of the population to be evaluated using the surrogate model. Using the fully evaluated fitness we can now update the surrogate training data and re-train our surrogate model. In this sense, the use of the EI criterion is two-fold. Firstly, it allows us to evaluate individuals we expect to offer the greatest improvement, ensuring we expend our resources on the most promising networks. Secondly, it allows us to select iteratively, informing our surrogate model with better and better solutions. Figure 2 demonstrates our framework. On the left, we can see the workflows of the evolutionary process and the surrogate model. On the right, we can see in more detail how the expected EI criterion is used within our surrogate model management strategy.

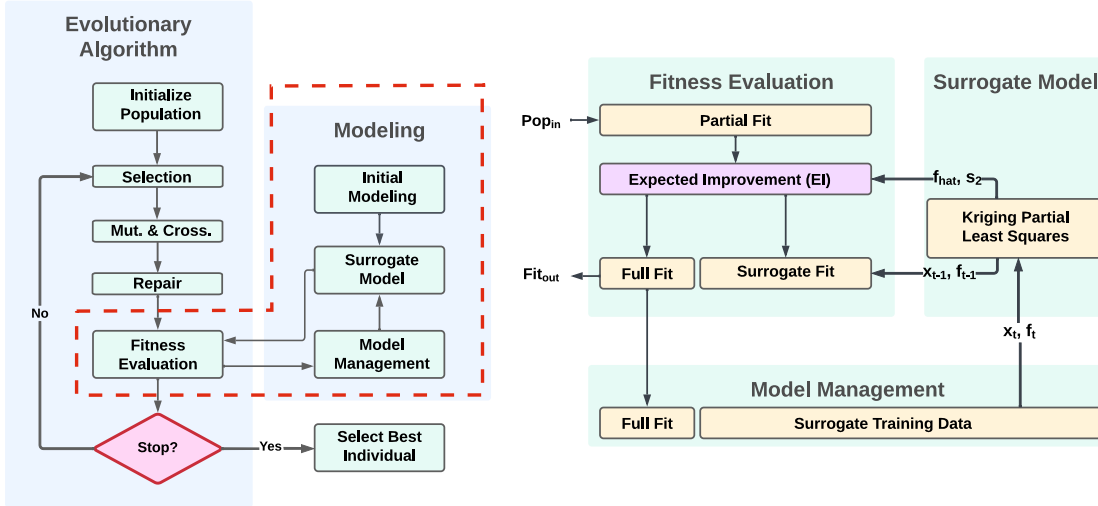


Fig. 2: *Left:* Diagram showing the interplay between a typical evolutionary algorithm and a surrogate model approach. *Right:* The surrogate model management strategy is shown on a more granular level.

4.3 Genetic Operations and Repair Mechanism

The mutation operator we use in this work mutates either a single input or output register or the operand. The mutation operator works on both effective code and non-effective code.

Additionally, we make use of a novel effective crossover operator. This operator selects two crossover points from the effective sections of the parent code and then transfers segments to create offspring. While the ends of each segment contain effective code the code within may be non-effective. A further point mutation repair is applied at either segment end to ensure input and output registers match after crossover has been applied. Ultimately, this results in always producing a valid network.

A repair mechanism is incorporated when performing the genotype-to-phenotype mapping to ensure models are compilable. There are several conditions for incorporating a repair mechanism, but it is important to note that in each case the repair is only performed on the effective code, by either inserting or deleting a specific layer.

- When there is no effective code in the genotype, we perform a single effective mutation inserting a single convolutional layer from our feature list.
- It is possible for the program to compile without a convolutional layer at the start (i.e., Max pooling, dropout, etc). We add an effective insertion of a convolutional layer before any layer that does not match this criteria.
- When the input data dimensions are too low or if the number of data-reducing operations exceeds the dimensions of the data, the neural network will fail to compile, we remove layers iteratively until the model is valid.
- If the input data dimensions are too high or if the number of data-reducing operations are too few, our fully-connected layer can be composed of an arbitrarily large number of nodes. Furthermore, adding further fully-connected layers can cause the models to become prohibitively expensive to evaluate, as such, we put additional conditions in to check and add a fully-connected layer before the final output layer. The size of the fully connected layer is dependent on the number of parameters from the CNN portion of the architecture.

5 Experimental Setup

The Breast Cancer Histopathological Image Classification (BreakHis) [25] dataset is a binary classification dataset consisting of microscopic images containing 2,480 benign and 5,429 malignant tumours, split across four types of magnification (40x, 100x, 200x and 400x). Each image consists of 700X460 pixels and 3-channel RGB with 8-bit depth in each channel. These images are re-scaled to a lower resolution of 64x64 pixels with 3-channels for each experiment. The overall data split for training, validation, test and test-2 is approximately (63.5%, 12.5%, 12.5%, 12.5%) or at the training level for each network, when only one test set is considered (70%, 15%, 15%). The first test set is used to evaluate the accuracy of each network on unseen data and is used to generate a fitness function for the NeuroLGP approach. The second data set is used to evaluate the NeuroLGP approach after evolution and as such is unseen to both the neural network and evolutionary processes. Since the dataset is imbalanced, synthetic minority oversampling technique (SMOTE) [8] is used to up-sample the minority class. In our work we use image level accuracy (ILA), which is the standard classification accuracy at the image level [1]. We performed a substantial number of independent runs, 96 in all, totalling approximately 80 GPU days (3 approaches x 4 datasets x 8 runs). A generation size of 15 and a population size of 50 were selected for all 3 approaches. Experiments were conducted using Kay supercomputer provided by the Irish Centre for High-End Computing (ICHEC). Runs were run in parallel, with each individual run assigned to a single Nvidia Tesla V100 GPU with 16GB Ram. A set of three experiments were designed to test three approaches:

1. A *baseline* approach which is a random search approach that initializes the network architecture layers randomly. These networks are not evolved but are repaired in line with the repair

mechanism as discussed in Section 4.3. All networks are fully trained to the full number of epochs (30 epochs),

2. The *expensive* approach where we evolve the structure of our networks using the NeuroLGP approach for the max number of epochs (Pop. size = 50, Gen. size = 15),
3. The *surrogate* approach where we employ the NeuroLGP-SM. The surrogate model is informed using partially trained networks (10 epochs) for 60% of the population size.

6 Results

6.1 Preliminary Analysis of the Baseline Model

The baseline in our approach makes use of random search over the feature extraction portion of our network architecture with a set of repair operations to ensure the architecture can be compiled. Typical network sizes from this approach ranged from $\sim 20\text{K}$ to $\sim 12.5\text{M}$ parameters with the baseline. For example of VGG-16 has $\sim 40\text{M}$ parameters [24]. Figure 3 shows the distribution of accuracy values for the baseline, for magnification x200, for all networks across 8 runs (750 individuals per run, 6000 individuals overall). It is important to point out that across all 8 runs, every model found by the baseline was valid and compilable. Slight peaks at 0.33 and 0.66 represent poor performing networks that that classify all the data as a single class. We can see that despite a tailed distribution to the left, the vast majority of individuals are centred around the median peak of 0.83. Additionally, the performance of VGG-16 is also shown in the figure with an accuracy of 0.87, which was verified in our own experiments and documented in the work by Cascianelli et al. [7]. The main takeaway is the baseline approach can find competitive architectures to compare against the surrogate and expensive approaches.

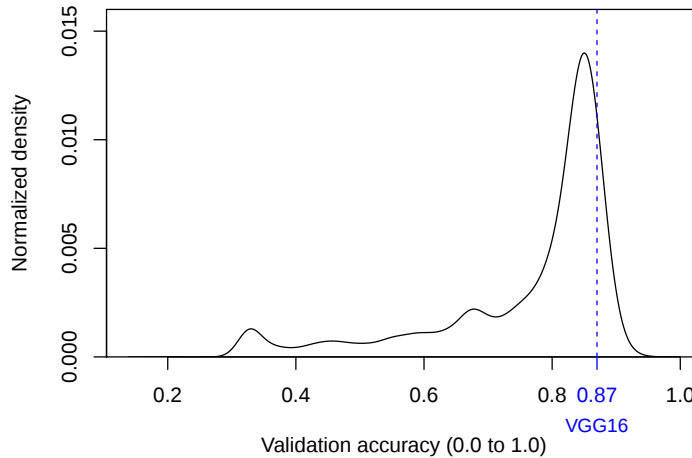


Fig. 3: Accuracy of 6000 individuals (read text).

6.2 Comparison of Models

A recent survey by Benhammou on the BreakHis dataset lists a comprehensive comparison of various machine learning approaches in terms of accuracy [1]. The results demonstrated here

match the state-of-the-art approaches in terms of performance and are particularly notable as the approach we use does not make use of transfer learning techniques.

A comparison is conducted between the three experiments as outlined in Section 5 for the BreakHis dataset. The violin plots in Figures 4(a-d) plot the accuracies of the best individual in terms of fitness, for each run, for the x40, x100, x200 and x400 magnifications, respectively. For reference, if we were to look at baseline as an example (left-hand side in each plot of Figure 4), then the individuals shown in these plots represent individuals found at the far right of the density plot in Figure 3. Initially, we ran the three models for 4 runs only, which took approximately 40 GPU days in total. It should be noted that in neuroevolution the norm is usually to do a single run. The initial results showed a tendency for the surrogate and expensive models to outperform the baseline model for the x40, x100 and x400 models however for x200 the surrogate and expensive underperformed. It was decided to extend this, to 8 runs totalling 80 GPU days in all. The updated results are represented here in the violin plots, where again we can see that in the case of x40 and x100 magnification, the surrogate (green) and expensive (blue) models tend to outperform the baseline (red) model. Results are more mixed when we look at x200, but is markedly improved over the initial set of 4 runs. The x400 magnification appears somewhat mixed, however, with more runs x400 may have more separation for surrogate and expensive models over the baseline given there are few better-performing individuals above the 92% mark. Comparing surrogate and expensive models in particular, an important observation can be made, in that for all magnifications they have similar accuracies. Given the distribution of the baseline is not fully separable from the other methods, compounded by the low run number, it was decided analysis of the best individuals across all runs would be more informative rather than relying purely on statistical testing.

The accuracies for the best-performing individuals across all runs for the baseline, surrogate and expensive models are listed respectively as such; for the x40 magnification the accuracies are 0.889, 0.913 and 0.930; for the x100 magnification are 0.869, 0.903 and 0.916; for the x200 magnification are 0.946, 0.970 and 0.960 and finally for the x400 magnification are 0.914, 0.925 and 0.925. In no instance did the baseline produce the best-performing individual. For the x40 and x100 magnifications, the expensive model produced the best-performing individual. For the x200 magnification the surrogate model produced the best-performing individual and the x400 magnification has a tie for the best-performing individual. In summary, while on average our surrogate and expensive models are as good as or better than the baseline, when we consider the best-performing individuals across all runs the surrogate and expensive models are better.

6.3 Analysis of the Surrogate Model

We use three metrics to determine how well our surrogate model performs in terms of predicting the fitness of our partially trained models. Firstly, we use the Mean Squared Error (MSE) between the predicted fitness and actual fitness. Values closer to 0 indicate our surrogate model is accurate in predicting fitness. Secondly, Kendall’s Tau is used to measure the correlation between the predicted fitness and the actual fitness [33]. A coefficient value of -1 indicates a perfect negative correlation, while a coefficient of +1 indicates a perfect positive correlation. Values close to 0 would indicate no discernible correlation. When considering all runs, a mean value close to 0 would indicate our surrogate model is performing poorly while a mean value closer to 1 would indicate an ideal-performing surrogate model for Kendall’s Tau. Thirdly, the R^2 score [3], is used to measure how close the predicted value is to its true value and ranges from $-\infty$ to 1, and explains how much variability is in the prediction model. R^2 scores close to 1 would represent a model which perfectly fits the data.

Table 1 summarizes the effectiveness of the surrogate model for each of the datasets, as seen in the ‘Quality of fit NeuroLGP-SM’ header of this table. The low MSE values (second column),

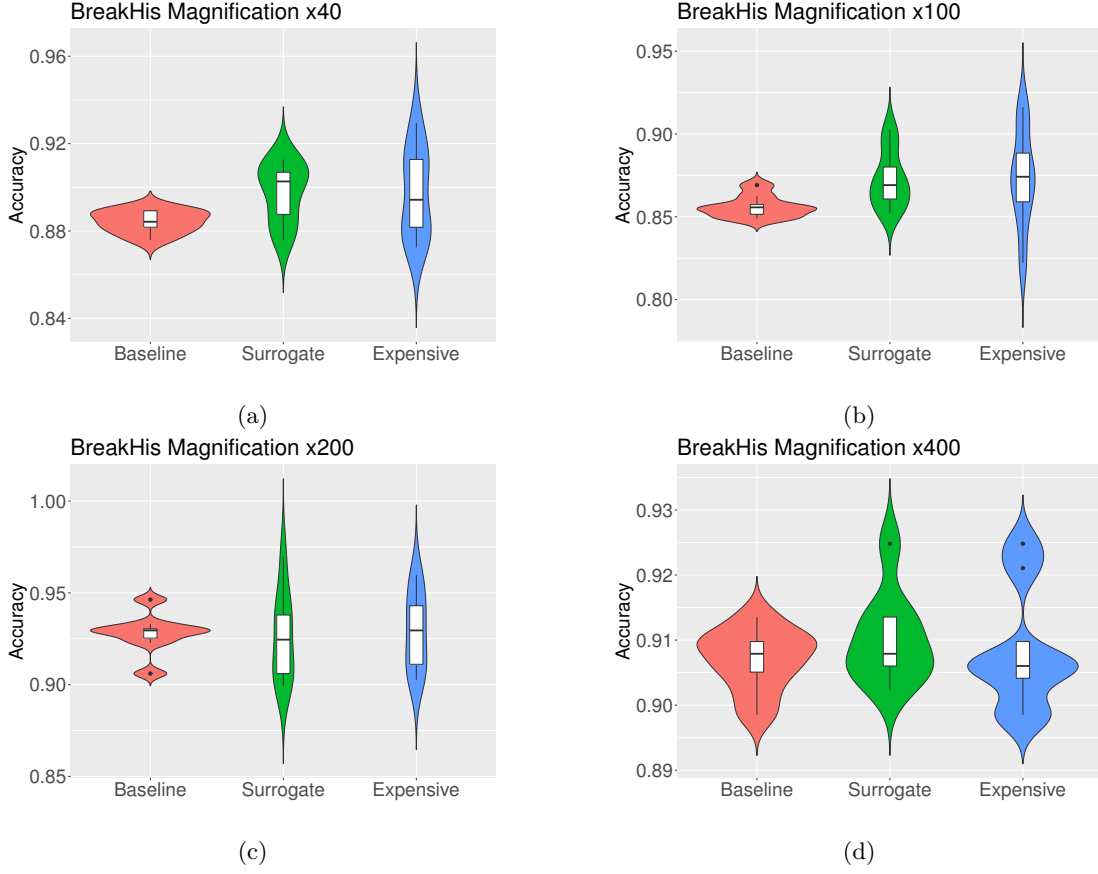


Fig. 4: Figures 4a through 4d plot violin plots, showing the distribution of accuracies of the architectures with the best fitness for each of the four data sets across 8 runs.

show a relatively low error between the predicted and actual fitness, however, the MSE value for x40 magnification was comparatively high at 0.0037, compared to x100, x200 and x400 at magnifications 0.0017, 0.0014 and 0.0009 respectively. Furthermore, Kendall’s Tau values range from 0.5647 to 0.6791 across the four datasets indicating a strong positive correlation between the actual and predicted fitness. We can see that for x40 we have R^2 of 0.5026 indicating a moderate level of fit being captured by the surrogate model while x100, x200 and x400 we have R^2 values of 0.6665, 0.7079 and 0.7786 indicating a strong level of fit. A deeper dive into why x40 had a comparatively lower performance revealed that for a particular run, the surrogate model estimated poorly on the lower fitness models but in general was better for higher fitness models and was better overall for the other seven runs.

The ‘GPU hours per run’ header of Table 1 shows a comparison between the expensive and surrogate model in average runtime per GPU hour, in the fifth and sixth columns respectively. Of note is that the surrogate model typically is an order of magnitude higher for standard deviation, meaning there is greater variance associated with runtime for the surrogate model approach. The last column shows the percentage in terms of time saved using the surrogate model over the expensive model. In general, we can see that given the parameters selected, we roughly save 25% in terms of GPU hours. Overall, if we consider the total time it took to run

Table 1: Average quality of fit and number of GPU hours per run.

Mag.	Quality of fit NeuroLGP-SM			GPU hours per run				
	MSE	Kendall's Tau	R^2	Expensive (hrs)		Surrogate (hrs)		Reduction Time
				Mean	Std	Mean	Std	
x40	0.0037	0.6019	0.5026	21.3	± 0.2	15.9	± 0.7	25.3%
x100	0.0017	0.6791	0.6665	22.7	± 0.1	16.6	± 0.7	26.7%
x200	0.0014	0.6225	0.7079	22.4	± 0.1	16.8	± 0.8	24.9%
x400	0.0009	0.5647	0.7786	20.0	± 0.1	15.3	± 0.6	23.8%

8 runs for all 4 datasets, the expensive models took ~ 28 GPU days and the surrogate models ~ 21 GPU days, saving approximately ~ 7 GPU days in all.

7 Conclusion and Future Work

In our work, we adapt the use of phenotypic distance vectors for surrogate-assisted neuroevolution of Deep Neural Networks (DNNs) using a variation of Kriging, entitled Kriging Partial Least Squares (KPLS). While phenotypic distance vectors have previously been used for traditional Artificial Neural Networks (ANNs) of only a few layers, their use in surrogate modelling for DNNs marks a significant leap forward in tackling high-dimensionality in neuroevolution. Our surrogate model management strategy makes use of a novel neuroevolutionary approach inspired by Linear Genetic Programming, entitled NeuroLGP which allows us to evolve compact, robust and variable-length architectures. This new approach was shown to outperform a competitive baseline for both the expensive (NeuroLGP) and surrogate-based variant (NeuroLGP-SM) using four magnification subsets of the BreakHis dataset when considering best-individuals across all runs. While we conducted an extensive set of experiments totalling 80 GPU days, in future, we would like to extend the number of runs from 8 to 16 to conduct a full statistical analysis. Furthermore, the runtime associated with NeuroLGP-SM was on average 25% faster compared to the expensive variant.

Acknowledgements

This publication has emanated from research conducted with the financial support of Science Foundation Ireland under Grant number 18/CRT/6049. For the purpose of Open Access, the author has applied a CC BY public copyright licence to any Author Accepted Manuscript version arising from this submission.

References

1. Y. Benhammou, B. Achchab, F. Herrera, and S. Tabik. Breakhis based breast cancer automatic diagnosis using deep learning: Taxonomy, survey and insights. *Neurocomputing*, 375:9–24, 2020.
2. J. Bergstra and Y. Bengio. Random search for hyper-parameter optimization. *Journal of machine learning research*, 13(2), 2012.
3. A. Bhosekar and M. Ierapetritou. Advances in surrogate based modeling, feasibility analysis, and optimization: A review. *Computers & Chemical Engineering*, 108:250–267, 2018.
4. M. A. Bouhlel, N. Bartoli, A. Otsmane, and J. Morlier. Improving kriging surrogates of high-dimensional design models by partial least squares dimension reduction. *Structural and Multidisciplinary Optimization*, 53:935–952, 2016.
5. M. Brameier and W. Banzhaf. *Linear genetic programming*, volume 1. Springer, 2007.

6. E. Brochu, V. M. Cora, and N. De Freitas. A tutorial on bayesian optimization of expensive cost functions, with application to active user modeling and hierarchical reinforcement learning. *arXiv preprint arXiv:1012.2599*, 2010.
7. S. Cascianelli, R. Bello-Cerezo, F. Bianconi, M. L. Fravolini, M. Belal, B. Palumbo, and J. N. Kather. Dimensionality reduction strategies for cnn-based classification of histopathological images. In *Intelligent Interactive Multimedia Systems and Services 2017 10*, pages 21–30. Springer, 2018.
8. N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer. Smote: synthetic minority over-sampling technique. *Journal of artificial intelligence research*, 16:321–357, 2002.
9. J. Chen, S.-h. Kao, H. He, W. Zhuo, S. Wen, C.-H. Lee, and S.-H. G. Chan. Run, don’t walk: Chasing higher flops for faster neural networks. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 12021–12031, 2023.
10. A. E. Eiben and J. E. Smith. *Introduction to Evolutionary Computing*. Springer Verlag, 2003.
11. A. Gaier, A. Asteroth, and J.-B. Mouret. Data-efficient neuroevolution with kernel-based surrogate models. In *Proceedings of the genetic and evolutionary computation conference*, pages 85–92, 2018.
12. E. Galván and P. Mooney. Neuroevolution in deep neural networks: Current trends and future challenges. *IEEE Transactions on Artificial Intelligence*, 2:476–493, 2021.
13. E. Galván and M. Schoenauer. Promoting semantic diversity in multi-objective genetic programming. In A. Auger and T. Stützle, editors, *Proceedings of the Genetic and Evolutionary Computation Conference, GECCO 2019, Prague, Czech Republic, July 13-17, 2019*, pages 1021–1029. ACM, 2019.
14. E. Galván and F. Stapleton. Semantic-based distance approaches in multi-objective genetic programming. In *2020 IEEE Symposium Series on Computational Intelligence (SSCI)*, pages 149–156. IEEE, 2020.
15. E. Galván and F. Stapleton. Evolutionary multi-objective optimisation in neurotrajectory prediction. *Applied Soft Computing*, 146:110693, 2023.
16. E. Galván, L. Trujillo, and F. Stapleton. Semantics in multi-objective genetic programming. *Applied Soft Computing*, 115:108143, 2022.
17. B. Greenwood and T. McDonnell. Surrogate-assisted neuroevolution. In *Proceedings of the Genetic and Evolutionary Computation Conference*, pages 1048–1056, 2022.
18. A. Hagg, M. Zaefferer, J. Stork, and A. Gaier. Prediction of neural network performance by phenotypic modeling. In *Proceedings of the Genetic and Evolutionary Computation Conference Companion*, pages 1576–1582, 2019.
19. Y. Jin. Surrogate-assisted evolutionary computation: Recent advances and future challenges. *Swarm and Evolutionary Computation*, 1(2):61–70, 2011.
20. D. R. Jones, M. Schonlau, and W. J. Welch. Efficient global optimization of expensive black-box functions. *Journal of Global optimization*, 13:455–492, 1998.
21. M. I. E. Khaldi and A. Draa. Surrogate-assisted evolutionary optimisation: a novel blueprint and a state of the art survey. *Evolutionary Intelligence*, pages 1–31, 2023.
22. Y. LeCun, Y. Bengio, and G. Hinton. Deep learning. *Nature*, 521(7553):436–444, 2015.
23. R. Miikkulainen, J. Liang, E. Meyerson, A. Rawal, D. Fink, O. Francon, B. Raju, H. Shahrzad, A. Navruzyan, N. Duffy, et al. Evolving deep neural networks. In *Artificial intelligence in the age of neural networks and brain computing*, pages 293–312. Elsevier, 2019.
24. K. Simonyan and A. Zisserman. Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*, 2014.
25. F. A. Spanhol, L. S. Oliveira, C. Petitjean, and L. Heutte. Breast cancer histopathological image classification using convolutional neural networks. In *2016 international joint conference on neural networks (IJCNN)*, pages 2560–2567. IEEE, 2016.
26. K. O. Stanley and R. Miikkulainen. Evolving neural networks through augmenting topologies. *Evol. Comput.*, 10(2):99–127, June 2002.
27. F. Stapleton and E. Galván. Semantic neighborhood ordering in multi-objective genetic programming based on decomposition. In *2021 IEEE Congress on Evolutionary Computation (CEC)*, pages 580–587. IEEE, 2021.
28. F. Stapleton and E. Galván. Initial steps towards tackling high-dimensional surrogate modeling for neuroevolution using kriging partial least squares. In *Proceedings of the Companion Conference on Genetic and Evolutionary Computation, GECCO ’23 Companion*, page 83–84, New York, NY, USA, 2023. Association for Computing Machinery.
29. F. Stapleton and E. Galván. NeuroLGP-SM: scalable Surrogate-Assisted neuroevolution for deep neural networks. In *2024 IEEE Congress on Evolutionary Computation (CEC) (CEC 2024)*, page 8.81, Yokohama, Japan, June 2024.
30. F. Stapleton, E. Galván, G. Sistu, and S. Yogamani. Neuroevolutionary multi-objective approaches to trajectory prediction in autonomous vehicles. In *Proceedings of the Genetic and Evolutionary Computation Conference Companion, GECCO ’22*, page 675–678, New York, NY, USA, 2022. Association for Computing Machinery.

31. J. Stork, M. Zaefferer, and T. Bartz-Beielstein. Improving neuroevolution efficiency by surrogate model-based optimization with phenotypic distance kernels. In *International Conference on the Applications of Evolutionary Computation (Part of EvoStar)*, pages 504–519. Springer, 2019.
32. Y. Sun, H. Wang, B. Xue, Y. Jin, G. G. Yen, and M. Zhang. Surrogate-assisted evolutionary deep learning using an end-to-end random forest-based performance predictor. *IEEE Transactions on Evolutionary Computation*, 24(2):350–364, 2020.
33. Y. Sun, B. Xue, M. Zhang, and G. G. Yen. Evolving deep convolutional neural networks for image classification. *IEEE Transactions on Evolutionary Computation*, 24(2):394–407, 2019.
34. K. Swersky, J. Snoek, and R. P. Adams. Freeze-thaw bayesian optimization. *arXiv preprint arXiv:1406.3896*, 2014.