



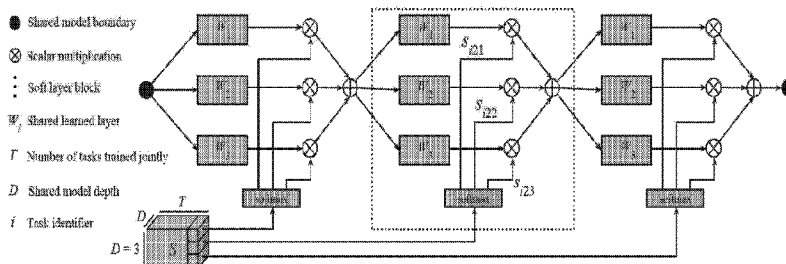
(12) **DEMANDE DE BREVET CANADIEN**
CANADIAN PATENT APPLICATION

(13) **A1**

(86) **Date de dépôt PCT/PCT Filing Date:** 2018/12/13
(87) **Date publication PCT/PCT Publication Date:** 2019/06/20
(85) **Entrée phase nationale/National Entry:** 2020/06/11
(86) **N° demande PCT/PCT Application No.:** US 2018/065472
(87) **N° publication PCT/PCT Publication No.:** 2019/118737
(30) **Priorités/Priorities:** 2017/12/13 (US62/598,409);
2018/02/06 (US62/627,166); 2018/04/24 (US62/662,082);
2018/12/07 (US16/212,830)

(51) **Cl.Int./Int.Cl.** **G06N 3/04** (2006.01),
G06N 3/08 (2006.01)
(71) **Demandeur/Applicant:**
COGNIZANT TECHNOLOGY SOLUTIONS U.S.
CORPORATION, US
(72) **Inventeurs/Inventors:**
LIANG, JASON ZHI, US;
MEYERSON, ELLIOT, US;
RISTO, MIKKULAINEN, US
(74) **Agent:** FASKEN MARTINEAU DUMOULIN LLP

(54) **Titre : EVOLUTION D'ARCHITECTURES POUR RESEAUX NEURONAUX MULTITACHES**
(54) **Title: EVOLUTION OF ARCHITECTURES FOR MULTITASK NEURAL NETWORKS**



(57) **Abrégé/Abstract:**

Evolution and coevolution of neural networks via multitask learning is described. The foundation is (1) the original soft ordering, which uses a fixed architecture for the modules and a fixed routing (i.e. network topology) that is shared among all tasks. This architecture is then extended in two ways with CoDeepNEAT: (2) by coevolving the module architectures (CM), and (3) by coevolving both the module architectures and a single shared routing for all tasks using (CMSR). An alternative evolutionary process (4) keeps the module architecture fixed, but evolves a separate routing for each task during training (CTR). Finally, approaches (2) and (4) are combined into (5), where both modules and task routing are coevolved (CMTR).

(12) INTERNATIONAL APPLICATION PUBLISHED UNDER THE PATENT COOPERATION TREATY (PCT)

(19) World Intellectual Property

Organization

International Bureau

(43) International Publication Date

20 June 2019 (20.06.2019)



(10) International Publication Number

WO 2019/118737 A1

(51) International Patent Classification:

G06F 17/30 (2006.01) G10L 15/06 (2013.01)

(21) International Application Number:

PCT/US2018/065472

(22) International Filing Date:

13 December 2018 (13.12.2018)

(25) Filing Language:

English

(26) Publication Language:

English

(30) Priority Data:

62/598,409 13 December 2017 (13.12.2017) US
 62/627,166 06 February 2018 (06.02.2018) US
 62/662,082 24 April 2018 (24.04.2018) US
 16/212,830 07 December 2018 (07.12.2018) US

(71) Applicant: **COGNIZANT TECHNOLOGY SOLUTIONS U.S. CORPORATION** [US/US]; 211 Quality Circle, College Station, TX 77845 (US).

(72) Inventors: **LIANG, Jason, Zhi**; 3509 Langdon Common, Fremont, CA 94538 (US). **MEYERSON, Elliot**; 550 Battery, #1905, San Francisco, CA 94111 (US). **RISTO, Mikku**; 339 Olmsted Road, Stanford, CA 94305 (US).

(74) Agent: **BEY, Dawn-Marie**; Bey & Cotropia PLLC, 213 Bayly Court, Richmond, VA 23229 (US).

(81) Designated States (unless otherwise indicated, for every kind of national protection available): AE, AG, AL, AM,

AO, AT, AU, AZ, BA, BB, BG, BH, BN, BR, BW, BY, BZ, CA, CH, CL, CN, CO, CR, CU, CZ, DE, DJ, DK, DM, DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IR, IS, JO, JP, KE, KG, KH, KN, KP, KR, KW, KZ, LA, LC, LK, LR, LS, LU, LY, MA, MD, ME, MG, MK, MN, MW, MX, MY, MZ, NA, NG, NI, NO, NZ, OM, PA, PE, PG, PH, PL, PT, QA, RO, RS, RU, RW, SA, SC, SD, SE, SG, SK, SL, SM, ST, SV, SY, TH, TJ, TM, TN, TR, TT, TZ, UA, UG, US, UZ, VC, VN, ZA, ZM, ZW.

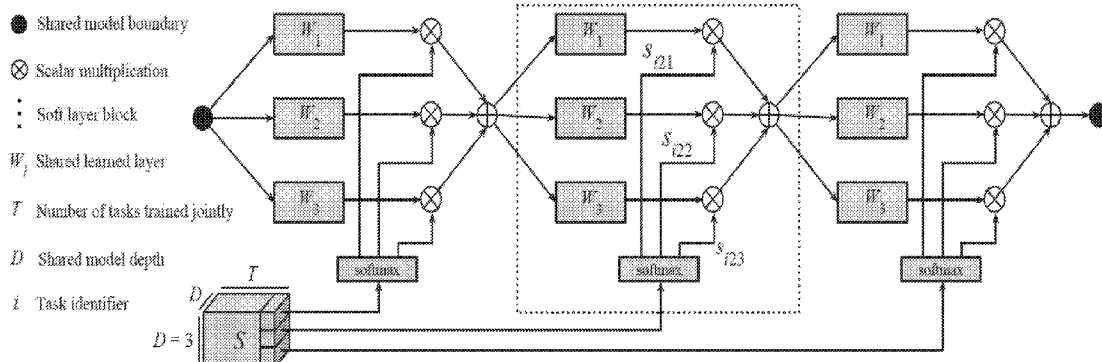
(84) Designated States (unless otherwise indicated, for every kind of regional protection available): ARIPO (BW, GH, GM, KE, LR, LS, MW, MZ, NA, RW, SD, SL, ST, SZ, TZ, UG, ZM, ZW), Eurasian (AM, AZ, BY, KG, KZ, RU, TJ, TM), European (AL, AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI, FR, GB, GR, HR, HU, IE, IS, IT, LT, LU, LV, MC, MK, MT, NL, NO, PL, PT, RO, RS, SE, SI, SK, SM, TR), OAPI (BF, BJ, CF, CG, CI, CM, GA, GN, GQ, GW, KM, ML, MR, NE, SN, TD, TG).

Published:

— with international search report (Art. 21(3))

(54) Title: EVOLUTION OF ARCHITECTURES FOR MULTITASK NEURAL NETWORKS

Figure 1



(57) **Abstract:** Evolution and coevolution of neural networks via multitask learning is described. The foundation is (1) the original soft ordering, which uses a fixed architecture for the modules and a fixed routing (i.e. network topology) that is shared among all tasks. This architecture is then extended in two ways with CoDeepNEAT: (2) by coevolving the module architectures (CM), and (3) by coevolving both the module architectures and a single shared routing for all tasks using (CMSR). An alternative evolutionary process (4) keeps the module architecture fixed, but evolves a separate routing for each task during training (CTR). Finally, approaches (2) and (4) are combined into (5), where both modules and task routing are coevolved (CMTR).



WO 2019/118737 A1

EVOLUTION OF ARCHITECTURES FOR MULTITASK NEURAL NETWORKS

Inventors: Jason Zhi LIANG, Elliot MEYERSON, Risto MIIKKULAINEN

CROSS-REFERENCE TO RELATED APPLICATIONS

[0001] The present application claims the benefit of priority to: U.S. Provisional Patent Application No. 62/598,409 for EVOLVING MULTITASK NEURAL NETWORK STRUCTURE filed on December 13, 2017, which is incorporated herein by reference in its entirety and is a continuation of U.S. Application Serial No. 16/212,830 entitled EVOLUTIONARY ARCHITECTURES FOR EVOLUTION OF DEEP NEURAL NETWORKS filed December 7, 2018 which claims benefit of priority to U.S. Provisional Patent Application No. 62/627,166 for EVOLUTIONARY ARCHITECTURE SEARCH FOR DEEP MULTITASK NETWORKS filed February 6, 2018; U.S. Provisional Patent Application No. 62/662,082 for EVOLUTIONARY ARCHITECTURES FOR EVOLUTION OF DEEP NEURAL NETWORKS filed April 24, 2018; and U.S. Provisional Patent Application No. 62/598,409 for EVOLVING MULTITASK NEURAL NETWORK STRUCTURE filed on December 13, 2017, each of which is incorporated herein by reference in their entireties.

INCORPORATIONS

[0002] In addition to the applications referenced above, the following materials are incorporated by reference as if fully set forth herein: J. Z. Liang, E. Meyerson, and R. Miikkulainen, "EVOLUTIONARY ARCHITECTURE SEARCH FOR DEEP MULTITASK NETWORKS," GECCO (2018) July 15–19, 2018, Kyoto, Japan and "FROM NODES TO NETWORKS: EVOLVING RECURRENT NEURAL NETWORKS," by Aditya Rawal and Risto Miikkulainen, GECCO, July 15-19, 2018, Kyoto, Japan; and "EVOLVING DEEP NEURAL NETWORKS" by Miikkulainen et al., CoRR, abs/1703.00548, March 2017.

FIELD OF THE TECHNOLOGY DISCLOSED

[0003] The technology disclosed is directed to artificial intelligence type computers and digital data processing systems and corresponding data processing methods and products for emulation of intelligence (i.e., knowledge based systems, reasoning systems, and knowledge

acquisition systems); and including systems for reasoning with uncertainty (e.g., fuzzy logic systems), adaptive systems, machine learning systems, and artificial neural networks. The technology disclosed generally relates to evolving neural networks architectures which are then used to solve complex problems.

BACKGROUND

[0004] The subject matter discussed in this section should not be assumed to be prior art merely as a result of its mention in this section. Similarly, a problem mentioned in this section or associated with the subject matter provided as background should not be assumed to have been previously recognized in the prior art. The subject matter in this section merely represents different approaches, which in and of themselves can also correspond to implementations of the claimed technology.

[0005] In multitask learning (MTL), a neural network is trained simultaneously to perform several different tasks at once. For instance, given an image as input, it can recognize the objects in it, identify the type of scene, and generate a verbal caption for it. Typically, the early parts of the network are shared between tasks, and the later parts, leading to the different tasks, are separate. The network is trained with gradient descent in all these tasks, and therefore the requirements of all tasks are combined in the shared parts of the network. The embeddings thus reflect the requirements of all tasks, making them more robust and general. Performance of a multitask network in each task can therefore exceed the performance of a network trained in only a single task.

[0006] Designing deep neural network architectures for multitask learning is a challenge: there are many ways to tie the tasks together, and the design choices matter. Much of the research in deep learning in recent years has focused on coming up with better architectures, and MTL is no exception. As a matter of fact, architecture plays possibly an even larger role in MTL because there are many ways to tie the multiple tasks together. The best network architectures are large and complex, and have become very hard for human designers to optimize.

[0007] MTL exploits relationships across problems to increase overall performance. The underlying idea is that if multiple tasks are related, the optimal models for those tasks will be related as well. In the convex optimization setting, this idea has been implemented via various

regularization penalties on shared parameter matrices. Evolutionary methods have also had success in MTL, especially in sequential decision-making domains.

[0008] Deep MTL has extended these ideas to domains where deep learning thrives, including vision, speech, natural language processing, and reinforcement learning. The key design decision in constructing a deep multitask network is deciding how parameters such as convolutional kernels or weight matrices are shared across tasks. Designing a deep neural network for a single task is already a high-dimensional open-ended optimization problem; having to design a network for multiple tasks and deciding how these networks share parameters grows this search space combinatorially. Most existing approaches draw from the deep learning perspective that each task has an underlying feature hierarchy, and tasks are related through an *a priori* alignment of their respective hierarchies. Another existing approach adapts network structure by learning task hierarchies, though it still assumes this strong hierarchical feature alignment.

[0009] Soft ordering is a recent approach that avoids such an alignment by allowing shared layers to be used across different depths. Through backpropagation, the joint model learns how to use each shared (potentially nonlinear) layer W_d at each depth d for the t -th task. This idea is implemented by learning a distinct scalar s_{td} for each such location, which then multiplies the layer's output. The final output at depth d for the task is then the sum of these weighted outputs across layers, i.e., a soft merge. **Figure 1** illustrates an exemplary soft ordering network with three shared layers. Soft ordering learns how to use the same layers in different locations by learning a tensor S of task-specific scaling parameters. S is learned jointly with the W_d , to allow flexible sharing across tasks and depths. This architecture enables the learning of layers that are used in different ways at different depths for different tasks.

[0010] Although soft ordering allows flexible sharing across depths, layers are still only applied in a fixed grid-like topology, which biases and restricts the type of sharing that can be learned.

[0011] As deep learning tasks and benchmarks become increasingly complex, finding the right architecture becomes more important. In fact, the performance of many state of the art networks depend mostly on novel and interesting architectural innovations. Unfortunately, discovering useful hyperparameters and architectures by hand is tedious and difficult; as a result, much research focuses on developing automated methods for doing it. Some promising methods for

hyperparameter search include deep Bayesian optimization and CMA-ES. One unique approach uses reinforcement learning to develop an LSTM policy for generating appropriate network topologies and hyperparameters for a given task.

[0012] One particularly promising area of research is the use of evolutionary algorithms (EAs) for performing architecture search. Evolutionary methods are well suited for these kinds of problems because they can be readily applied with no gradient information. Some of these approaches use a modified version of NEAT, an EA for neuron-level neuroevolution, for searching network topologies.

[0013] Others rely on genetic programming or hierarchical evolution. Along these lines, CoDeepNEAT combines the power of NEAT's neural topology search with hierarchical evolution to efficiently discover architectures within large search spaces. Networks evolved using CoDeepNEAT have achieved good results in image classification and image captioning domains, outperforming popular hand-designed architectures. The paper entitled "Evolving Deep Neural Networks" by Miikkulainen et al., CoRR, abs/1703.00548, March 2017, is descriptive of the CoDeepNeat method and is incorporated herein by reference.

[0014] CoDeepNEAT begins by initializing two populations, one of modules and one of blueprints, with minimal complexity. The blueprints and modules each contain at least one species and are evolved/complexified separately with a modified version of NEAT. An individual in the blueprint population is a directed acyclic graph (DAG) where each node contains a pointer to a particular module species. An individual in the module population is a DAG where each node represents a particular DNN layer and its corresponding hyperparameters (number of neurons, activation function, etc.). As shown in **Figure 2**, in a method of assembling networks for fitness evaluation in CoDeepNEAT, modules from species specified in the blueprint are inserted into locations specified in the blueprint, forming a temporary population of assembled networks. This approach allows evolving repetitive and deep structures seen in many successful DNNs. Each individual in this population is then evaluated by training it on a supervised learning task, and assigning its performance as fitness. The fitnesses of the individuals (networks) are attributed back to blueprints and modules as the average fitness of all the assembled networks containing that blueprint or module. One of the advantages of CoDeepNEAT is that it is capable of discovering modular, repetitive structures seen in state of the art networks such as Googlenet and Resnet.

[0015] Accordingly, there remains a need in the art for improved methods for designing deep neural network architectures for multitask learning.

SUMMARY OF THE EMBODIMENTS

[0016] The embodiments presented herein describe an automated, flexible approach for evolving architectures, i.e. hyperparameters, modules, and module routing topologies, of deep multitask networks. A recent deep MTL architecture called soft ordering is used as a starting point, in which a different soft sequence of modules is learned for each task. The present embodiments extend this architecture in several ways. First, a novel algorithm for evolving task specific routings that create a unique routing between modules for each task is employed. Second, more general modules with the same soft ordering architecture are evolved. Third, the general modules are evolved together with a blueprint, a shared routing for all tasks, that improves upon the soft ordering architecture. Fourth, as a capstone architecture, the task specific routing are evolved together with the general modules. Accordingly, the present embodiments generalize soft ordering layers to more general modules, and introduce evolutionary approaches to both design these modules and to discover how to assemble these modules into appropriate topologies for multitask learning.

[0017] In a first exemplary embodiment, a processor implemented method for evolving task-specific topologies in a multitask architecture includes: establishing a set of shared modules which are shared among each task-specific topology; initializing the shared modules

$\{\mathcal{M}_k\}_{k=1}^K$ with random weights; initializing a champion individual module routing scheme for each task (t), wherein the i th individual for the t th task is represented by a tuple (E_{ti}, G_{ti}, D_{ti}) , and further wherein E_{ti} is an encoder, G_{ti} is a DAG, which specifies the individual module routing scheme, and D_{ti} is a decoder, with E_{t1} and D_{t1} initialized with random weights; for each champion individual (E_{t1}, G_{t1}, D_{t1}) , generating a challenger (E_{t2}, G_{t2}, D_{t2}) by mutating the t th champion in accordance with a predetermined mutation subprocess; jointly training each champion and challenger for M iterations on a training set of data; evaluating each champion and challenger on a validation set of data to determine an accuracy fitness for each individual champion and challenger for its predetermined task; if a challenger has higher accuracy fitness than a corresponding champion, then the champion is replaced wherein $(E_{t1}, G_{t1}, D_{t1}) = (E_{t2}, G_{t2}, D_{t2})$.

D_{t2}); calculating an average accuracy fitness across all champions for tasks in the multitask architecture; and checkpointing the shared modules when the average accuracy is best achieved.

[0018] In a second exemplary embodiment, a processor implemented method for evolving task-specific topologies and shared modules in a multitask architecture includes: initializing a population of modules and randomly selecting modules (m) from each species in the population and grouping selected modules from each species (k) together into sets of modules M_k , providing the sets of modules M_k to a task-specific routing evolution subprocess, wherein the subprocess: establishes a set of shared modules which are shared among each task-specific topology; initializes a champion individual module routing scheme for each task (t), wherein the i th individual for the t th task is represented by a tuple (E_{ti}, G_{ti}, D_{ti}) , and further wherein E_{ti} is an encoder, G_{ti} is a DAG, which specifies the individual module routing scheme, and D_{ti} is a decoder, with E_{t1} and D_{t1} initialized with random weights; for each champion individual (E_{t1}, G_{t1}, D_{t1}) , generating a challenger (E_{t2}, G_{t2}, D_{t2}) by mutating the t th champion in accordance with a predetermined mutation subprocess; jointly training each champion and challenger for M iterations on a training set of data; evaluating each champion and challenger on a validation set of data to determine an accuracy fitness for each individual champion and challenger for its predetermined task; if a challenger has higher accuracy fitness than a corresponding champion, then the champion is replaced wherein $(E_{t1}, G_{t1}, D_{t1}) = (E_{t2}, G_{t2}, D_{t2})$; calculating an average accuracy fitness across all champions for tasks in the multitask architecture; checkpointing the shared modules when the average accuracy fitness is best achieved; and attributing the best achieved average accuracy fitness determined from the task-specific routing evolution subprocess to each module (m) as part of a module evolution subprocess which further includes applying evolutionary operators to evolve modules (m).

BRIEF DESCRIPTION OF THE DRAWINGS

[0019] In the drawings, like reference characters generally refer to like parts throughout the different views. Also, the drawings are not necessarily to scale, with an emphasis instead generally being placed upon illustrating the principles of the technology disclosed. In the following description, various implementations of the technology disclosed are described with reference to the following drawings, in which:

- [0020] **Figure 1** is an exemplary soft ordering network for use with one or more embodiments herein;
- [0021] **Figure 2** illustrates a method of assembling networks for fitness evaluation in CoDeepNEAT;
- [0022] **Figure 3** provides an overview of the methods employed in the embodiments herein in multitask learning;
- [0023] **Figure 4** illustrates an instance of coevolution of task routing with three tasks and four modules that are share across all tasks;
- [0024] **Figure 5** illustrates comparison of fitness over generations of single runs of CM (Coevolution of Modules), CMSR (Coevolution of Modules and Shared Routing), and CMTR in accordance with one or more embodiments herein;
- [0025] **Figure 6** illustrates comparison of fitness over number of meta-iterations of training for CTR (Coevolution of Task Routing) and CMTR (Coevolution of Modules and Task Routing) in accordance with one or more embodiments herein;
- [0026] **Figure 7** illustrates comparison of fitness over generations of CM with disabling, enabling, and evolving module weight sharing in accordance with one or more embodiments herein; and
- [0027] **Figures 8a, 8b, 8c, 8d** illustrates exemplary best performing modules from the CMTR experiment, and sample routing topologies evolved for different exemplary alphabets.

DETAILED DESCRIPTION

[0028] The following discussion is presented to enable any person skilled in the art to make and use the technology disclosed, and is provided in the context of a particular application and its requirements. Various modifications to the disclosed implementations will be readily apparent to those skilled in the art, and the general principles defined herein may be applied to other implementations and applications without departing from the spirit and scope of the technology disclosed. Thus, the technology disclosed is not intended to be limited to the implementations shown, but is to be accorded the widest scope consistent with the principles and features disclosed herein.

[0029] **Figure 3** provides an overview of the methods employed in the embodiments herein in multitask learning. The foundation is (1) the original soft ordering, which uses a fixed

architecture for the modules and a fixed routing (i.e. network topology) that is shared among all tasks. This architecture is then extended in two ways with CoDeepNEAT: (2) by coevolving the module architectures (CM), and (3) by coevolving both the module architectures and a single shared routing for all tasks using (CMSR). This paper also introduces a novel approach (4) that keeps the module architecture fixed, but evolves a separate routing for each task during training (CTR). Finally, approaches (2) and (4) are combined into (5), where both modules and task routing are coevolved (CMTR). Table 1 gives high-level algorithmic descriptions of these methods, which are described in detail below.

TABLE 1

CM Algorithm	CMSR Algorithm	CTR Algorithm	CMTR Algorithm
Given fixed blueprint Initialize module population Each generation: • Assemble MTL networks with modules • Randomly initialize all weights • Train each MTL network with backprop • Assign fitnesses to modules	Initialize blueprint/module populations Each generation: • Assemble MTL networks with blueprints/modules • Randomly initialize all weights • Train each MTL network with backprop • Assign fitnesses to modules and blueprints • Update blueprint/module populations	Given set of modules Initialize topology population for each task Randomly initialize all weights Each meta-iteration: • Assemble networks • Jointly train all networks with backprop • Assign fitnesses to topologies • Update topology populations	Initialize module population Each generation: • Assemble sets of modules • Train sets of modules with CTR • Assign fitnesses to modules • Update module populations

• Update module populations			
--	--	--	--

Coevolution of Modules

[0030] In Coevolution of Modules (CM), CoDeepNEAT is used to search for promising module architectures, which are then inserted into appropriate positions to create an enhanced soft ordering network. The evolutionary process works as follows:

1. CoDeepNEAT initializes a population of modules MP. The blueprints are not used.
2. Modules are randomly chosen from each species in MP, grouped into sets M and are assembled into enhanced soft ordering networks.
3. Each assembled network is trained/evaluated on some task and its performance is returned as fitness.
4. Fitness is attributed to the modules, and NEAT evolutionary operators are applied to evolve the modules.
5. The process is repeated from step 1 until CoDeepNEAT terminates, i.e. no further progress is observed for a given number of generations.

Unlike in soft ordering, the number of modules and the depth of the network are not fixed but are evolved as global hyperparameters by CoDeepNEAT (however the layout is still a grid-like structure). Since the routing layout is fixed, the blueprint population of CoDeepNEAT, which determines how the modules are connected, is not used. Thus, one key operation in the original CoDeepNEAT, i.e. inserting modules into each node of the blueprint DAG, is skipped; only the module population is evolved.

[0031] To assemble a network for fitness evaluation, an individual is randomly chosen from each species in the module population to form an ordered set of distinct modules M. The hyperparameters evolved in each of the module's layers include the activation function, kernel size, number of filters, L2 regularization strength and output dropout rate. In addition, CoDeepNEAT also coevolves global hyperparameters that are relevant to the entire assembled network as a whole; these include learning rate, the number of filters of the final layer of each module, and the weight initialization method. Evolvable hyperparameters in each node include

the activation function, kernel size, number of filters, L2 regularization strength and output dropout rate. The modules are then transformed into actual neural networks by replacing each node in the DAG with the corresponding layer. To ensure compatibility between the inputs and outputs of each module, a linear 1×1 convolutional layer (number of filters determined by a global hyperparameter), followed by a max-pooling layer (provided that the feature map before pooling is at least 4×4) is included as the last layer in each module.

[0032] The modules are then inserted into the soft ordering network. The architecture of the network is interpreted as a grid of $K \times D$ slots, where d indicates the depth of the network and the slots with the same k value have the same module topology. For each available slot T_{kd} , the corresponding module M_k is inserted. If $k > |M|$, then $M_k \bmod |M|$ is inserted instead.

[0033] Finally, each module in a particular slot has the potential to share its weights with modules that have the same architecture and are located in other slots of the blueprint. Flag F_k in each module indicates whether or not the module's weights are shared. This flag is evolved as part of the module genotype in CoDeepNEAT. Also, there is also global flag F_d for each depth of the soft ordering network. If the M_k is placed in T_{kd} and both F_k and F_d are turned on, then the module is able to share its weights with any other M_k whose slot have both flags turned on as well. Such an arrangement allows each slot to have sharing enabled and disabled independently.

[0034] The assembled network is attached to separate encoders and decoders for each task and trained jointly using a gradient-based optimizer. Average performance over all tasks is returned as fitness back to CoDeepNEAT. That fitness is assigned to each of the modules in the assembled network. If a module is used in multiple assembled networks, their fitnesses are averaged into module fitness. After evaluation is complete, standard NEAT mutation, crossover, and speciation operators are applied to create the next generation of the module population.

Coevolution of Modules/Shared Routing

[0035] Coevolution of Modules and Shared Routing (CMSR) extends CM to include blueprint evolution. Thus, the routing between various modules no longer follows the fixed grid-like structure, but instead an arbitrary DAG. Each node in the blueprint genotype points to a particular module species. During assembly, the blueprints are converted into deep multitask networks as follows:

1. For each blueprint in the population, an individual module is randomly chosen from each species.
2. Each node in the blueprint is then replaced by the module from the appropriate species.
3. If a module has multiple inputs from previous nodes in the blueprint, the inputs are soft merged first.
4. The process is repeated from step 1 until reaching a target number of assembled networks.

As in CM, each node in the blueprint has a flag F_i that indicates whether node N_i should be shared or not. If two nodes are replaced by the same module and if both nodes have the sharing flag turned on, then the two modules will share weights. Such an arrangement allows each node to evolve independently whether they share weights or not. The training procedures for both CM and CMSR are otherwise identical. After fitness evaluation, the fitness is assigned to both blueprints and modules in the same manner as with CM. To accelerate evolution, the blueprint population is not initialized from minimally connected networks like the modules, but from randomly mutated networks that on average have five nodes.

Coevolution of Task Routing

[0036] The following multitask architecture search approach takes advantage of the dynamics of soft ordering by evolving task-specific topologies instead of a single blueprint. Like in soft ordering, in coevolution of task routing (CTR) there are K modules whose weights are shared everywhere they are used across all tasks. Like in blueprint evolution, CTR searches for the best ways to assemble modules into complete networks. However, unlike in blueprint evolution, CTR searches for a distinct module routing scheme for each task, and trains a single set of modules throughout evolution. Having a distinct routing scheme for each task makes sense if the shared modules are seen as a set of building blocks that are assembled to meet the differing demands of different problems. Training a single set of modules throughout evolution then makes sense as well: As modules are trained in different locations for different purposes during evolution, their functionality should become increasingly general, and it should thus become easier for them to adapt to the needs of a new location. Such training is efficient since the core structure of the network need not be retrained from scratch at every generation. In other words, CTR incurs no additional iterations of backpropagation over training a single fixed-topology

multitask model. Because of this feature, CTR is related to PathNet, which evolves pathways through modules as those modules are being trained. However, unlike in PathNet, in CTR distinct routing schemes are coevolved across tasks, modules can be applied in any location, and module usage is adapted via soft merges.

[0037] CTR operates a variant of a $(1 + 1)$ evolutionary strategy $((1 + 1)$ -ES) for each task. Separate ES for each task is possible because an evaluation of a multitask network yields a performance metric for each task. The $(1 + 1)$ -ES is chosen because it is efficient and sufficiently powerful in experiments, though it can potentially be replaced by any population-based method. To make it clear that a single set of modules is trained during evolution, and to disambiguate from the terminology of CoDeepNEAT, for CTR the term meta-iteration is used in place of generation.

[0038] Each individual constitutes a module routing scheme for a particular task. At any point in evolution, the i th individual for the t th task is represented by a tuple (E_{ti}, G_{ti}, D_{ti}) , where E_{ti} is an encoder, G_{ti} is a DAG, which specifies the module routing scheme, and D_{ti} is a decoder. The complete model for an individual is then given by

$$\mathbf{y}_t = (\mathcal{D}_{ti} \circ \mathcal{R}(G_{ti}, \{\mathcal{M}_k\}_{k=1}^K) \circ \mathcal{E}_{ti})(\mathbf{x}_t), \quad (1)$$

where $\mathcal{R}$ indicates the application of the shared modules $\mathcal{M}_k$ based on the DAG G_{ti} . Note that “ $\circ$ ” denotes function composition and $\mathcal{E}_{ti}$, and $\mathcal{D}_{ti}$ can be any neural network functions that are compatible with the set of shared modules. In the experiments discussed herein, each $\mathcal{E}_{ti}$ is an identity transformation layer, and each $\mathcal{D}_{ti}$ is a fully connected classification layer.

[0039] G_{ti} is a DAG, whose single source node represents the input layer for that task, and whose single sink node represents the output layer, e.g., a classification layer. All other nodes either point to a module $\mathcal{M}_k$ to be applied at that location, or a parameterless adapter layer that ensures adjacent modules are technically compatible. In the experiments discussed herein, all adapters are 2×2 max-pooling layers. Whenever a node of G_{ti} has multiple incoming edges, their contents are combined in a learned soft merge in accordance with the following learnable function:

$$\text{softmerge}(\text{in}_1, \dots, \text{in}_M) = \sum_{m=1..M} s_m \text{in}_m, \text{ with } \sum_{m=1..M} s_m = 1, \quad (2)$$

The algorithm (1) begins by initializing the shared modules $\{\mathcal{M}_k\}_{k=1}^K$ with random weights. Then, each champion (E_{t1}, G_{t1}, D_{t1}) is initialized, with E_{t1} and D_{t1} initialized with random weights, and G_{t1} according to some graph initialization policy. For example, the initialization of G_{t1} can be minimal or random. In the embodiments herein, G_{t1} is initialized to reflect the classical deep multitask learning approach, i.e.,

$$\mathcal{E}_{t1} \rightarrow \mathcal{M}_1 \rightarrow \mathcal{M}_2 \rightarrow \dots \rightarrow \mathcal{M}_K \rightarrow \mathcal{D}_{t1}, \quad (3)$$

with adapters added as needed.

[0040] At the start of each meta-iteration, a challenger (E_{t2}, G_{t2}, D_{t2}) is generated by mutating the t th champion as follows (the insertion of adapters is omitted for clarity):

1. The challenger starts as a copy of the champion, including learned weights, i.e., $(E_{t2}, G_{t2}, D_{t2}) := (E_{t1}, G_{t1}, D_{t1})$.
2. A pair of nodes (u, v) is randomly selected from G_{t2} such that v is an ancestor of u .
3. A module $\mathcal{M}_k$ is randomly selected from $\{\mathcal{M}_k\}_{k=1}^K$.
4. A new node w is added to G_{t2} with $\mathcal{M}_k$ as its function.
5. New edges (u, w) and (w, v) are added to G_{t2} .
6. The scalar weight of (w, v) is set such that its value after the softmax is some $\alpha \in (0, 1)$. To initially preserve champion behavior, α is set to be small. I.e., if $s_1, \dots, s_m$ are the scales of the existing inbound edges to v , s_{m+1} is the initial scale of the new edge, and $s_{\max} = \max(s_1, \dots, s_m)$ then

$$s_{m+1} = \ln \left(\frac{\alpha}{1-\alpha} \sum_{j=1..m} e^{s_j - s_{\max}} \right) + s_{\max} .$$

[0041] After challengers are generated, all champions and challengers are trained jointly for M iterations with a gradient-based optimizer. Note that the scales of G_{t1} and G_{t2} diverge during training, as do the weights of D_{t1} and D_{t2} . After training, all champions and challengers are evaluated on a validation set that is disjoint from the training data. The fitness for each individual is its performance for its task on the validation set and accuracy is the performance metric. If the challenger has higher fitness than the champion, then the champion is replaced, i.e., $(E_{t1}, G_{t1}, D_{t1}) = (E_{t2}, G_{t2}, D_{t2})$. After selection, if the average accuracy across all champions is the best achieved so far, the entire system is checkpointed, including the states of the modules. After evolution, the champions and modules from the last checkpoint constitute the final trained model, and are evaluated on a held out test set.

[0042] More than most evolutionary methods, this algorithm reflects an artificial ecology. The shared modules can be viewed as a shared finite set of environmental resources that is constantly exploited and altered by the actions of different tasks, which can correspond to different species in an environment. Within each task, individuals compete and cooperate to develop mutualistic relationships with the other tasks via their interaction with this shared environment. A visualization of CTR under this perspective is shown in **Figure 4**. Importantly, even if a challenger does not outperform its champion, its developmental (learning) process still affects the shared resources. This perspective suggests a more optimistic view of evolution, in which individuals can have substantial positive effects on the future of the ecosystem even without reproducing.

Coevolution of Modules and Task Routing

[0043] Both CM and CTR improve upon the performance of the original soft ordering baseline. Interestingly, these improvements are largely orthogonal, and they can be combined to form an even more powerful algorithm called Coevolution of Modules and Task Routing (CMTR). Since evolution in CTR occurs during training and is highly computational efficient, it is feasible to use CoDeepNEAT as an outer evolutionary loop to evolve modules. To evaluate and assign fitness to the modules, they are passed on to CTR (the inner evolutionary loop) for

evolving and assembling the task specific routings. The performance of the final task-specific routings is returned to CoDeepNEAT and attributed to the modules in the same way as in CM: Each module is assigned the mean of the fitnesses of all the CTR runs that made use of that module. Another way to characterize CMTR is that it overcomes the weaknesses in both CM and CTR: CM's inability to create a customized routing for each task and CTR's inability to search for better module architectures. CMTR's evolutionary loop works as follows:

1. CoDeepNEAT initializes a population of modules MP. The blueprints are not used.
2. Modules are randomly chosen from each species in MP and grouped together into sets of modules M.
3. Each set of modules M_k is given to CTR, which assembles the modules by evolving task-specific routings. The performance of the evolved routings on a task is returned as fitness.
4. Fitness is attributed to the modules, and NEAT's evolutionary operators applied to evolve the modules.
5. The process repeats from step 2. until CoDeepNEAT terminates, i.e. no improvement for a given number of generations.

[0044] One difference between CMTR and CM is that each module's final convolutional layer has additional evolvable hyperparameters such as kernel size, activation function, and output dropout rate. Preliminary experiments suggested that the relatively complex routings in CMTR (when compared to CM and CMSR) require more complex final layers as well, thus evolving the complexity of the final layer is optimal. Like in CTR, the weights between modules are always shared in CMTR. If modules with completely new weights are added to the task routings, they have to be trained from scratch and may even hurt performance, whereas adding a module with already partially trained weights does not. In addition, as the routings evolved by CTR are much larger than those discovered by CM and CMSR, disabling or evolving weight sharing significantly bloats the total number of weight parameters and slows training significantly.

Experiments and Results

[0045] The Omniglot dataset consists of 50 alphabets of handwritten characters, each of which induces its own character recognition task. There are 20 instances of each character, each a 105×105 black and white image. Omniglot is a good fit for MTL, because there is clear

intuition that knowledge of several alphabets will make learning another one easier. Omniglot has been used in an array of settings: generative modeling, one-shot learning, and deep MTL. Previous deep MTL approaches used random training/testing splits for evaluation. However, with model search (i.e. when the model architecture is learned as well), a validation set separate from the training and testing sets is needed. Therefore, in the experiments utilizing the present embodiments, a fixed training/validation/testing split of 50%/20%/30% is introduced for each task. Because training is slow and increases linearly with the number of tasks, a subset of 20 tasks out of the 50 possible is used in the current experiments. These tasks are trained in a fixed random order. Soft ordering is the current state-of-the-art method in this domain. The experiments therefore use soft ordering as a starting point for designing further improvements.

[0046] For CoDeepNEAT fitness evaluations, all networks are trained using Adam for 3000 iterations over the 20 alphabets; for CTR, the network is trained for 120 meta-iterations (30,000 iterations). Each iteration is equivalent to one full forward and backward pass through the network with a single example image and label chosen randomly from each task. The fitness assigned to each network is the average validation accuracy across the 20 tasks after training.

[0047] For CM and CMSR, CoDeepNEAT is initialized with approximately 50 modules (in four species) and 20 blueprints (in one species). For CMTR, a smaller module population of around 25 (in two species) is found to be beneficial in reducing noise since each module is evaluated more often. During each generation, 100 networks are assembled from modules and/or blueprints for evaluation. With CoDeepNEAT, the evaluation of assembled networks is distributed over 100 separate EC2 instances with a K80 GPU in AWS. The average time for training is usually around 1-2 hours depending on the network size. With CTR, because it is a (1 + 1) evolutionary strategy with a small population size, it is sufficient to run the algorithm on a single GPU.

[0048] Because the fitness returned for each assembled network is noisy, to find the best assembled CoDeepNEAT network, the top 50 highest fitness networks from the entire history of the run are retrained for 30,000 iterations. For the CM and CMSR experiments, decaying the learning rate by a factor of 10 after 10 and 20 epochs of training gave a moderate boost to performance. Similar boost is not observed for CTR and CMTR and therefore learning rate is not decayed for them. To evaluate the performance of the best assembled network on the test set (which is not seen during evolution or training), the network is trained from scratch again for

30,000 iterations. For CTR and CMTR, this is equivalent to training for 120 meta-iterations. During training, a snapshot of the network is taken at the point of highest validation accuracy. This snapshot is then evaluated and the average test accuracy over all tasks returned.

[0049] **Figure 5** demonstrates how the best and mean fitness improves for CM, CMSR, and CMTR in the CoDeepNEAT outer loop where module/blueprint coevolution occurs. All three algorithms converge roughly to the same final fitness value, which is around 78% validation accuracy. CMTR converges the fastest, followed by CM, and lastly CMSR. This result is expected since the search space of CMTR is the smallest (only the modules are evolved with CoDeepNEAT), larger for CM (evolution of modules and weight sharing), and largest for CMSR (evolution of modules, blueprints, and weight sharing). Although CM, CMSR, and CMTR converge to the same fitness in evolution, CMTR achieves better final performance because training occurs via CTR. **Figure 6** compares how fitness (i.e. average validation accuracy) improves for CTR (using the default modules) and CMTR (using the best evolved modules discovered by CMTR) during training, averaged over 10 runs. Interestingly, while CTR improves faster in the first 10 meta-iterations, it is soon overtaken by CMTR, demonstrating how evolution discovers modules that leverage the available training better.

[0050] One open question is how much sharing of weights between modules affects the performance of the assembled network. Although disabling weight sharing is not optimal for CTR due to the complexity of the routing, both CM and CMSR may benefit since their routing topologies are much smaller (minimizing the effects of parameter bloat). **Figure 7** compares the effect of enabling, disabling, and evolving weight sharing with CM. Interestingly, disabling weight sharing leads to better performance than enabling it, but evolving it is best. Thus, the design choice of evolving sharing in CM and CMSR is vindicated. An analysis of the architecture of the best assembled networks shows that weight sharing in particular locations such as near the output decoders is a good strategy.

[0051] **Table 2** shows the validation and test accuracy for the best evolved network produced by each method, averaged over 10 runs. The best-performing methods are highlighted in bold and standard error for the 10 runs is shown in parenthesis. In addition, performance of the baseline methods are shown, namely (1) a hand designed single-task architecture, i.e. where each task is trained and evaluated separately, and (2) the soft ordering network architecture. Indeed, the methods improve upon the baseline according to increasing complexity: Evolving

modules and evolving topologies is significantly better than the baselines, and evolving both is significantly better than either alone. CMTR, the combination of CoDeepNEAT and routing evolution, combines the advantages of both and performs the best.

Table 2

Algorithm	Val Accuracy (%)	Test Accuracy (%)
1. Single Task [29]	63.59 (0.53)	60.81 (0.50)
2. Soft Ordering [29]	67.67 (0.74)	66.59 (0.71)
3. CM	80.38 (0.36)	81.33 (0.27)
4. CMSR	83.69 (0.21)	83.82 (0.18)
5. CTR	82.48 (0.21)	82.36 (0.19)
6. CMTR	88.20 (1.02)	87.82 (1.02)

[0052] This series of results confirms the intuition well: as a baseline, soft ordering performs significantly better in each task than single-task training (67% vs. 61% accuracy). Evolution of modules and topologies improves significantly upon soft ordering. Coevolution of modules and topologies together improves even more, and the capstone architecture turns out to be the best (at 88%). The results thus demonstrate three general points: evolutionary architecture search can make a large difference in performance of deep learning networks; MTL can improve performance of deep learning tasks; and putting these together results in a particularly powerful approach.

[0053] The best networks have approximately three million parameters. **Figures 8a, 8b, 8c** and **8d** visualize one of the best performing modules from the CMTR experiment (**8a, 8b**), and sample routing topologies evolved for the different alphabets (**8c, 8d**). Because the CoDeepNEAT outer loop is based on two species, the four modules passed to the CTR inner loop consist of two different designs (but still separate weights). Thus, evolution has discovered that a combination of simple and complex modules is beneficial. Similarly, while the routing topologies for some alphabets are simple, others are very complex. Moreover, similar topologies emerge for similar alphabets (such as those that contain prominent horizontal lines, like Gurmukhi and Manipuri). Also, when evolution is run multiple times, similar topologies for the same alphabet result. Such useful diversity in modules and routing topologies, i.e. structures that complement each other and work well together, would be remarkably difficult to develop by

hand. However, evolution discovers them consistently and effectively, demonstrating the power of the approach.

[0054] The experiments show that MTL can improve performance significantly across tasks, and that the architecture used for it matters significantly. First, modules used in the architecture can be optimized and do end up different in a systematic way. Unlike in the original soft ordering architecture, evolution in CM, CMSR, and CMTR results in discovery of a wide variety of simple and complex modules, and they are often repeated in the architecture. Evolution thus discovers a useful set of building blocks that are diverse in structure. Second, the routing of the modules matter as well. In CMSR, the shared but evolvable routing allows much more flexibility in how the modules can be reused and extends the principals that makes soft ordering useful. The power of CTR and CMTR is from evolving different topologies for different tasks, and tie the tasks together by sharing the modules in them. In addition, sharing components (including weight values) in CMTR is crucial to its performance. If indeed the power from multitasking comes from integrating requirements of multiple tasks, this integration will happen in the embeddings that the modules form, so it makes sense that sharing plays a central role. Third, compared to the CTR and CMTR, CM and CMSR have evolved away from sharing of module weights, despite the fact that module architectures are often reused in the network. This result makes sense as well: because the topology is shared in this approach, the differentiation between tasks comes from differentiated modules. Such an approach is an opposite way to solve the problem. Even though it is an effective approach as well, it is not quite as powerful as differentiated topologies and shared modules.

[0055] One skilled in the art will recognize the variations to the embodiments herein that, though not explicitly described, are within the scope of the embodiments.

CLAIMS:

1. A processor implemented method for evolving task-specific topologies in a multitask architecture comprising:

establishing a set of shared modules which are shared among each task-specific topology;

initializing the shared modules $\{\mathcal{M}_k\}_{k=1}^K$ with random weights;

initializing a champion individual module routing scheme for each task (t), wherein the i th individual for the t th task is represented by a tuple (E_{ti}, G_{ti}, D_{ti}) , and further wherein E_{ti} is an encoder, G_{ti} is a DAG, which specifies the individual module routing scheme, and D_{ti} is a decoder, with E_{t1} and D_{t1} initialized with random weights;

for each champion individual (E_{t1}, G_{t1}, D_{t1}) , generating a challenger (E_{t2}, G_{t2}, D_{t2}) by mutating the t th champion in accordance with a predetermined mutation subprocess;

jointly training each champion and challenger for M iterations on a training set of data;

evaluating each champion and challenger on a validation set of data to determine an accuracy fitness for each individual champion and challenger for its predetermined task;

if a challenger has higher accuracy fitness than a corresponding champion, then the champion is replaced wherein $(E_{t1}, G_{t1}, D_{t1}) = (E_{t2}, G_{t2}, D_{t2})$;

calculating an average accuracy fitness across all champions for tasks in the multitask architecture; and

checkpointing the shared modules when the average accuracy is best achieved.

2. The process according to claim 1, wherein the predetermined mutation subprocess of includes:

- (i) start as a copy of the champion, including learned weights, wherein $(E_{t2}, G_{t2}, D_{t2}) := (E_{t1}, G_{t1}, D_{t1})$;
- (ii) randomly select a pair of nodes (u, v) from G_{t2} such that v is an ancestor of u ;
- (iii) randomly select a module M_k from the shared modules;
- (iv) add a new node w to G_{t2} with M_k as its function;
- (v) add new edges (u, w) and (w, v) to G_{t2} ;
- (vi) set the scalar weight of (w, v) such that its value after softmax is some $\alpha \in (0, 1)$.

3. The process according to claim 1, wherein the training set of data and the validation set of data are disjointed.

4. The process according to claim 1, wherein G_{ti} is initialized in accordance with a graph initialization policy.

5. The process according to claim 1, wherein a model for an individual is then given by

$$y_t = (\mathcal{D}_{ti} \circ \mathcal{R}(G_{ti}, \{\mathcal{M}_k\}_{k=1}^K) \circ \mathcal{E}_{ti})(x_t),$$

where $\mathcal{R}$ indicates application of the shared modules $\mathcal{M}_k$ based on the DAG G_{ti} .

6. The process according to claim 5, wherein $\mathcal{E}_{ti}$ and $\mathcal{D}_{ti}$ are selected from a grouping consisting of neural network functions that are compatible with the set of shared modules

7. The process according to claim 6, wherein each $\mathcal{E}_{ti}$ is an identity transformation layer, and $\mathcal{D}_{ti}$ is a fully connected classification layer.

8. The process according to claim 1, wherein G_{ti} is a DAG whose single source node represents the input layer for that task (t), and whose single sink node represents the output layer and further wherein all other nodes either point to a module $\mathcal{M}_k$ to be applied at that location, or to a parameterless adapter layer for ensuring adjacent modules are technically compatible.

9. A processor implemented method for evolving task-specific topologies and shared modules in a multitask architecture comprising:

initializing a population of modules and randomly selecting modules (m) from each species in the population and grouping selected modules from each species (k) together into sets of modules $\mathcal{M}_k$;

providing the sets of modules $\mathcal{M}_k$ to a task-specific routing evolution subprocess, wherein the subprocess:

establishes a set of shared modules which are shared among each task-specific topology;

initializes a champion individual module routing scheme for each task (t), wherein the i th individual for the t th task is represented by a tuple (E_{ti}, G_{ti}, D_{ti}) , and further wherein E_{ti} is an encoder, G_{ti} is a DAG, which specifies the individual module routing scheme, and D_{ti} is a decoder, with E_{t1} and D_{t1} initialized with random weights;

for each champion individual (E_{t1}, G_{t1}, D_{t1}) , generating a challenger (E_{t2}, G_{t2}, D_{t2}) by mutating the t th champion in accordance with a predetermined mutation subprocess;

jointly training each champion and challenger for M iterations on a training set of data;

evaluating each champion and challenger on a validation set of data to determine an accuracy fitness for each individual champion and challenger for its predetermined task;

if a challenger has higher accuracy fitness than a corresponding champion, then the champion is replaced wherein $(E_{t1}, G_{t1}, D_{t1}) = (E_{t2}, G_{t2}, D_{t2})$;

calculating an average accuracy fitness across all champions for tasks in the multitask architecture;

checkpointing the shared modules when the average accuracy fitness is best achieved;

attributing the best achieved average accuracy fitness determined from the task-specific routing evolution subprocess to each module (m) as part of a module evolution subprocess which further includes applying evolutionary operators to evolve modules (m).

10. The process according to claim 9, wherein the predetermined mutation subprocess of the task-specific routing evolution subprocess includes:

- (i) start as a copy of the champion, including learned weights, wherein $(E_{t2}, G_{t2}, D_{t2}) := (E_{t1}, G_{t1}, D_{t1})$;
- (ii) randomly select a pair of nodes (u, v) from G_{t2} such that v is an ancestor of u ;
- (iii) randomly select a module M_k from the shared modules;
- (iv) add a new node w to G_{t2} with M_k as its function;
- (v) add new edges (u, w) and (w, v) to G_{t2} ;
- (vi) set the scalar weight of (w, v) such that its value after softmax is some $\alpha \in (0, 1)$.

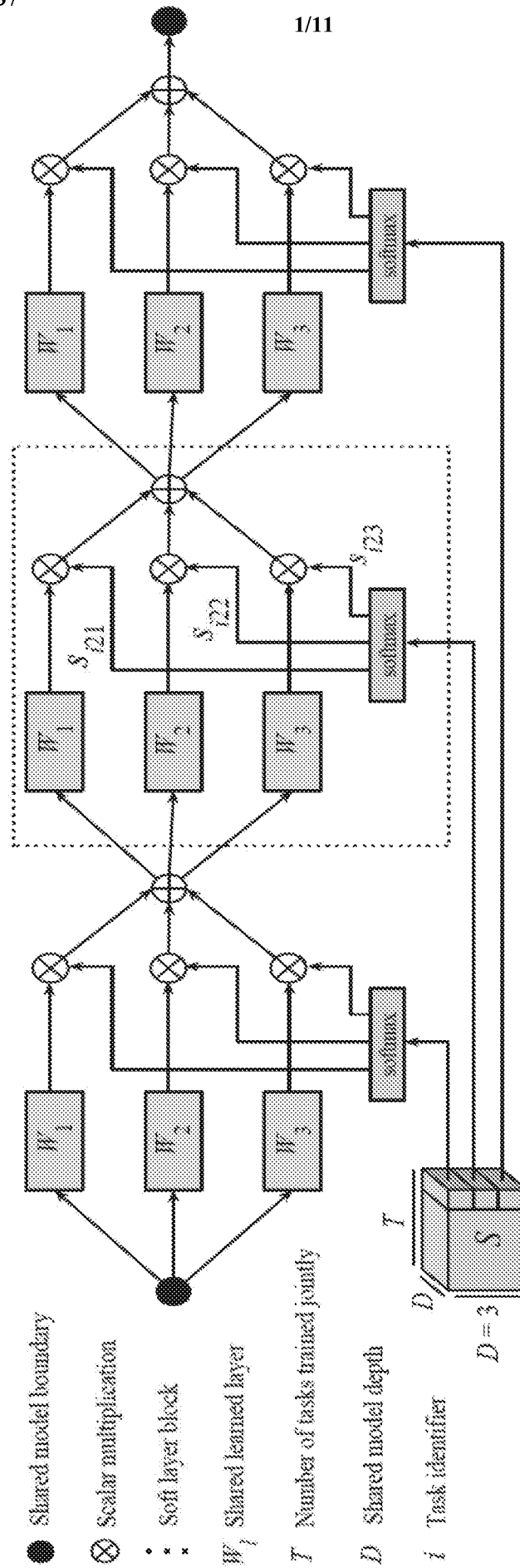


Figure 1

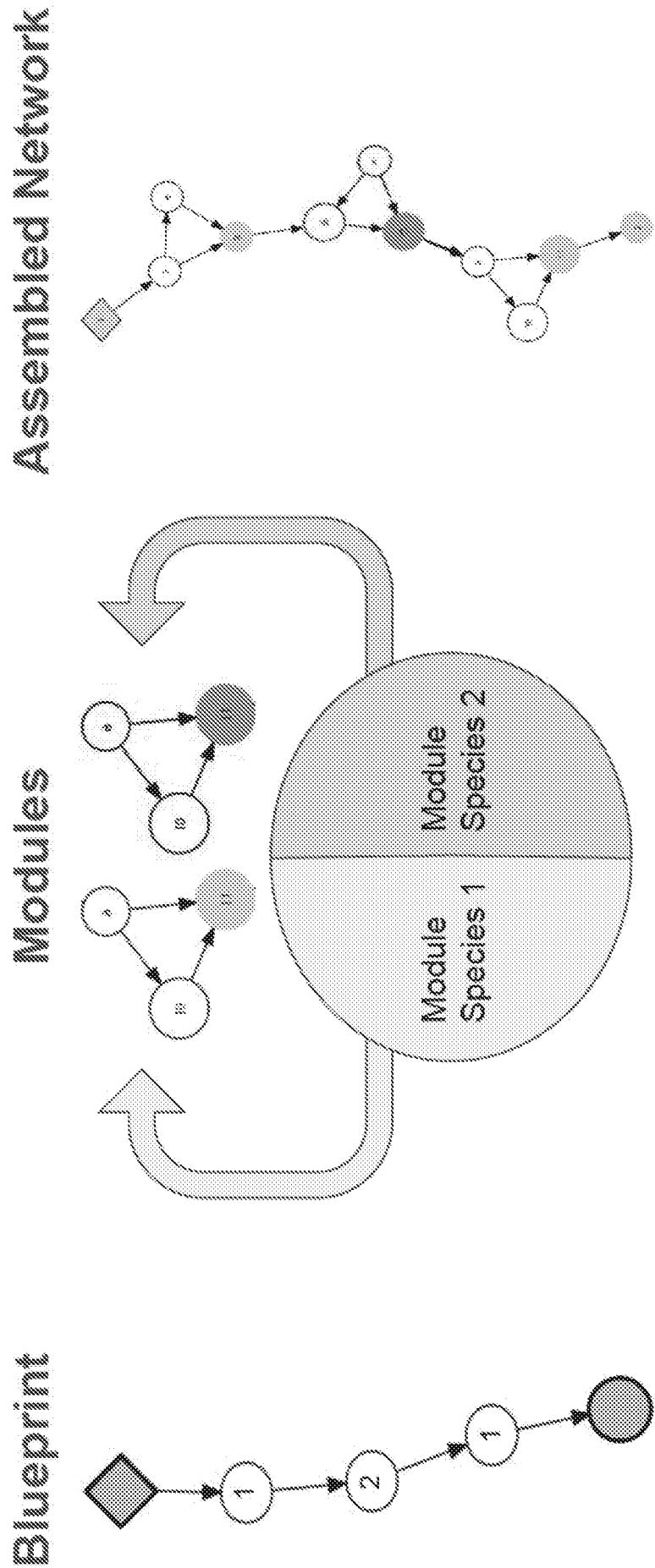


Figure 2

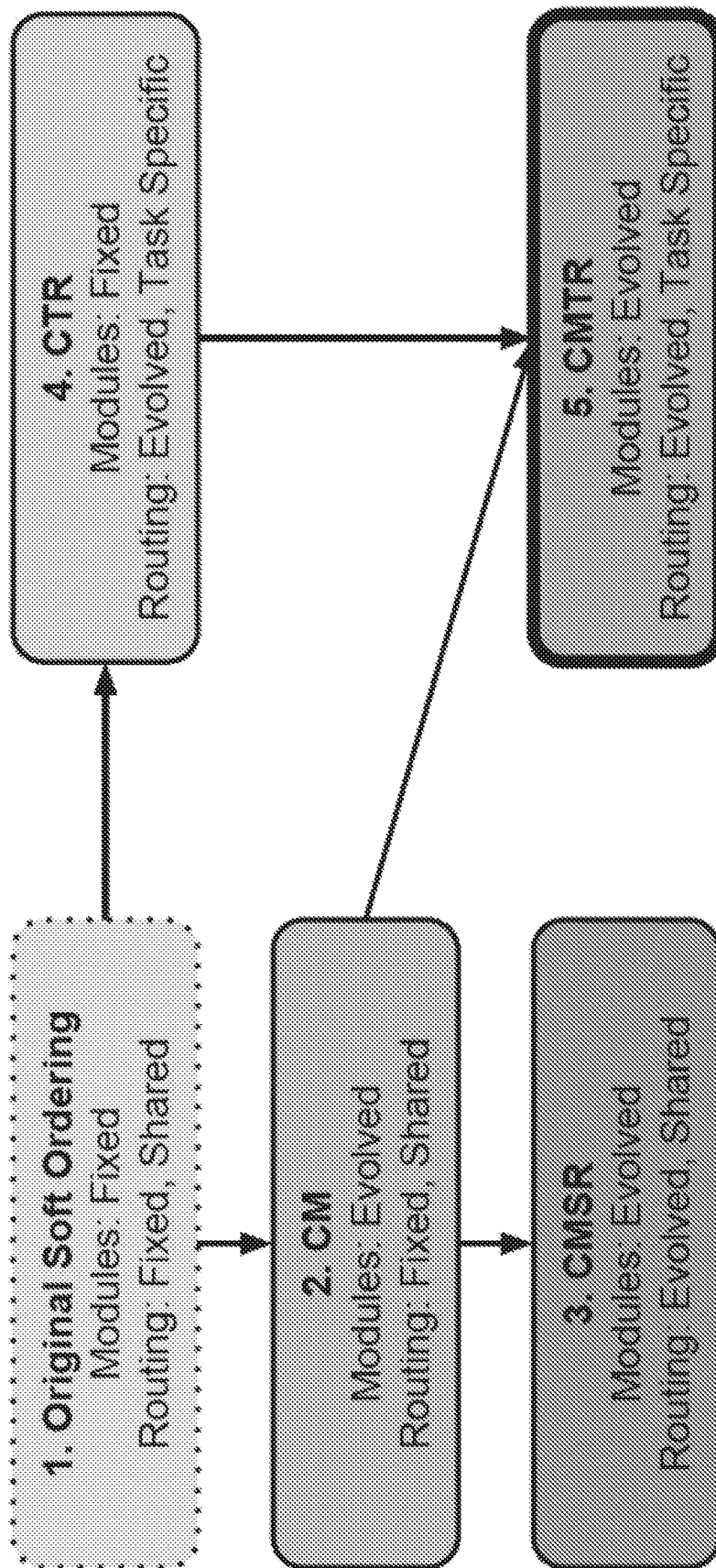


Figure 3

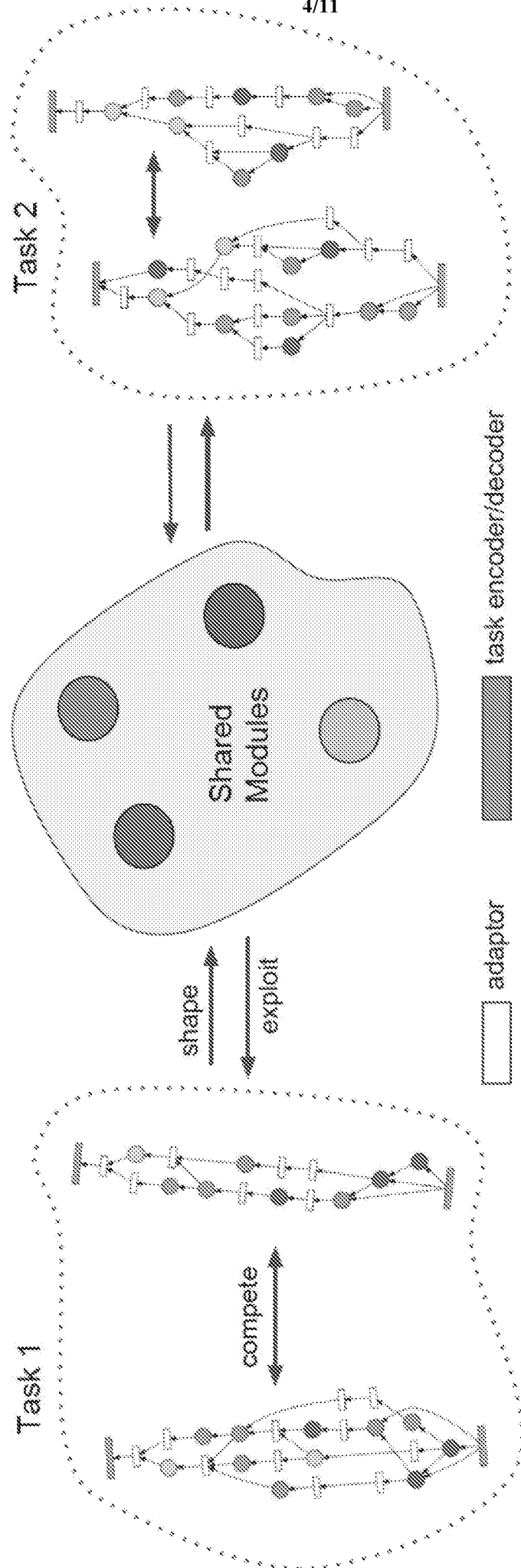


Figure 4

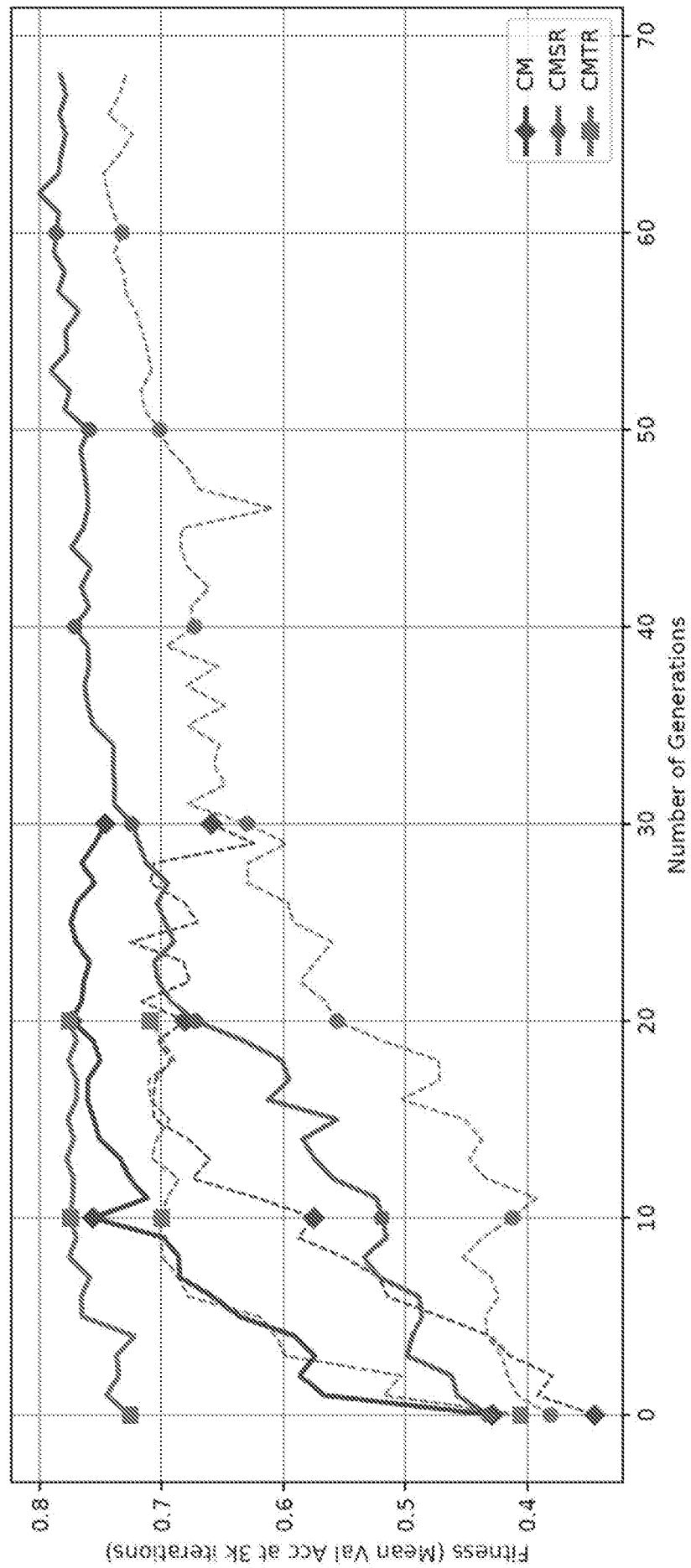


Figure 5

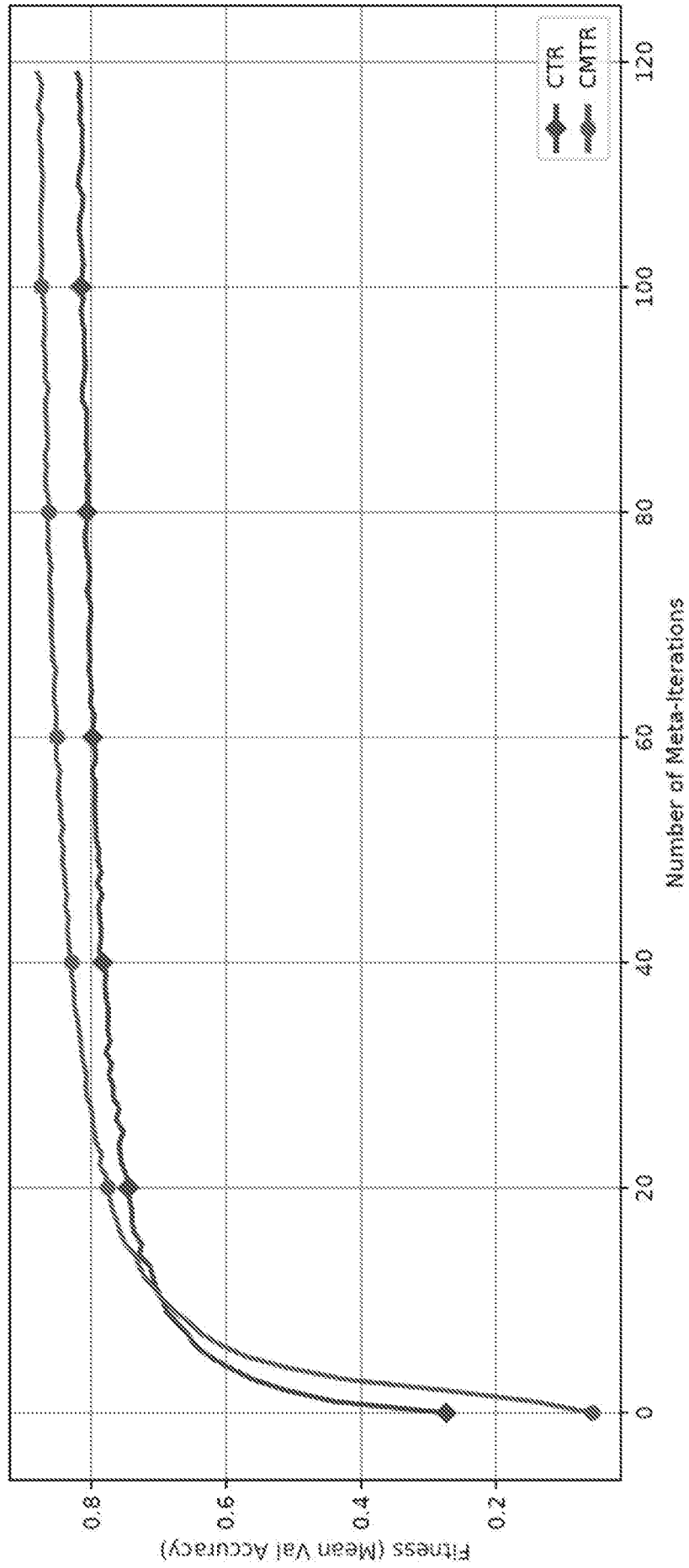


Figure 6

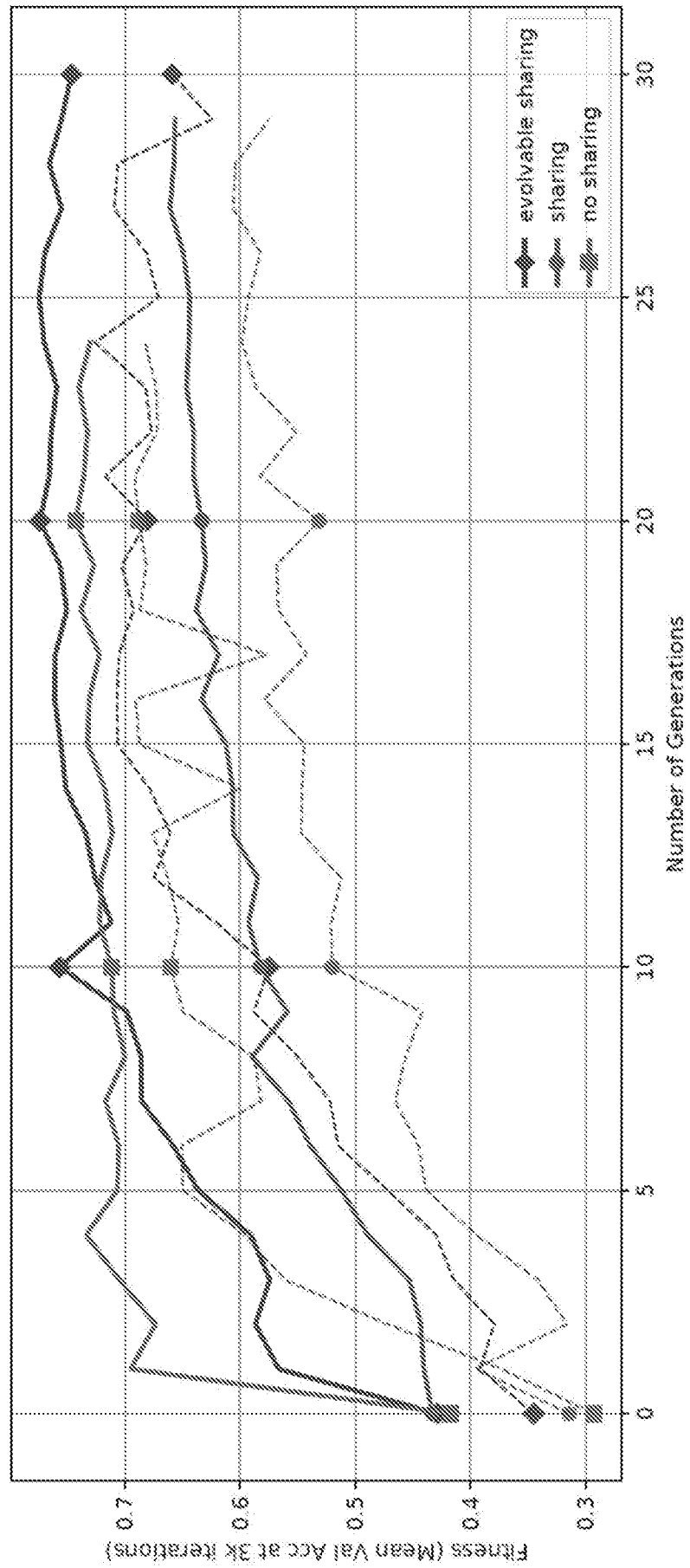


Figure 7

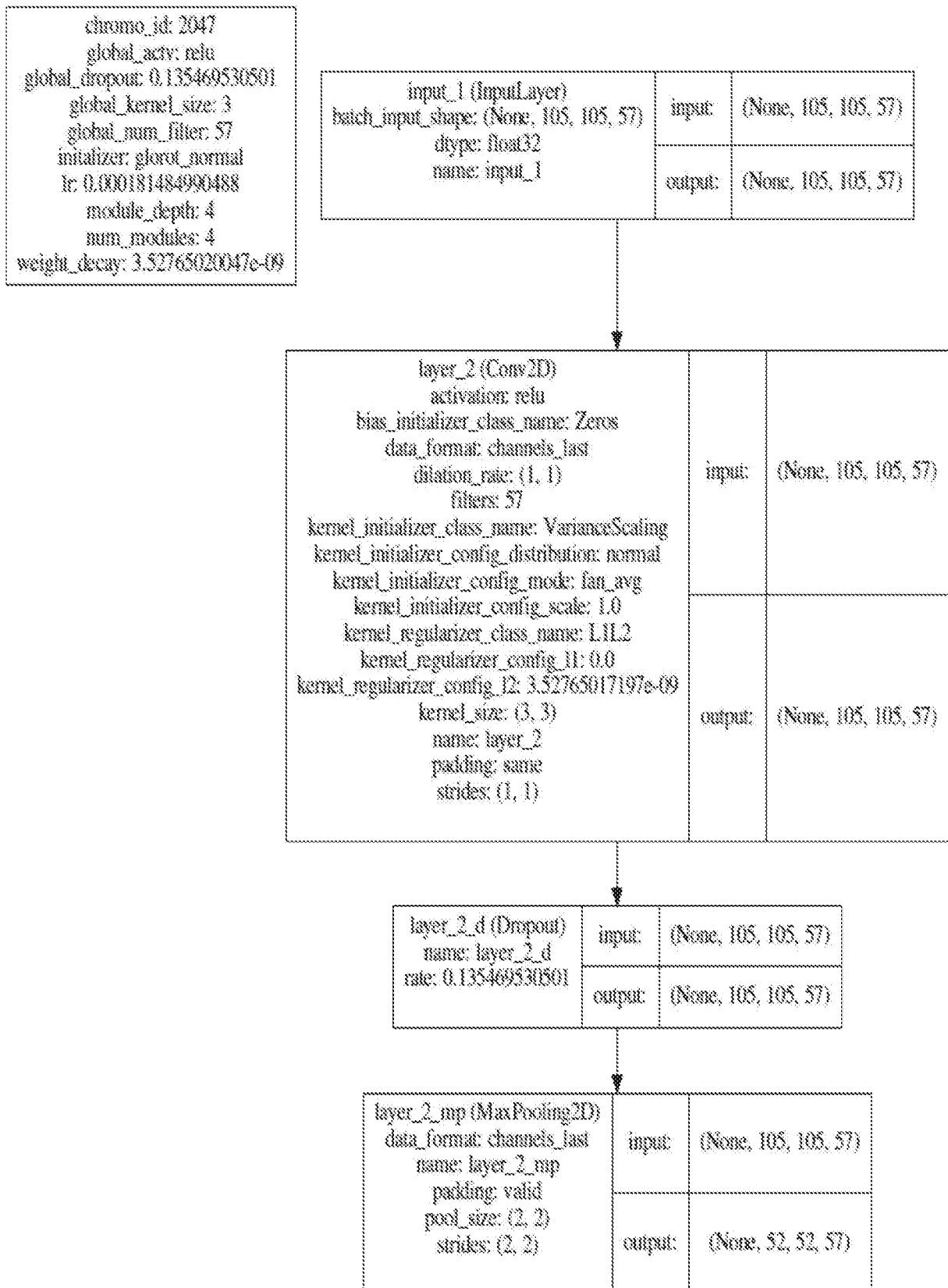
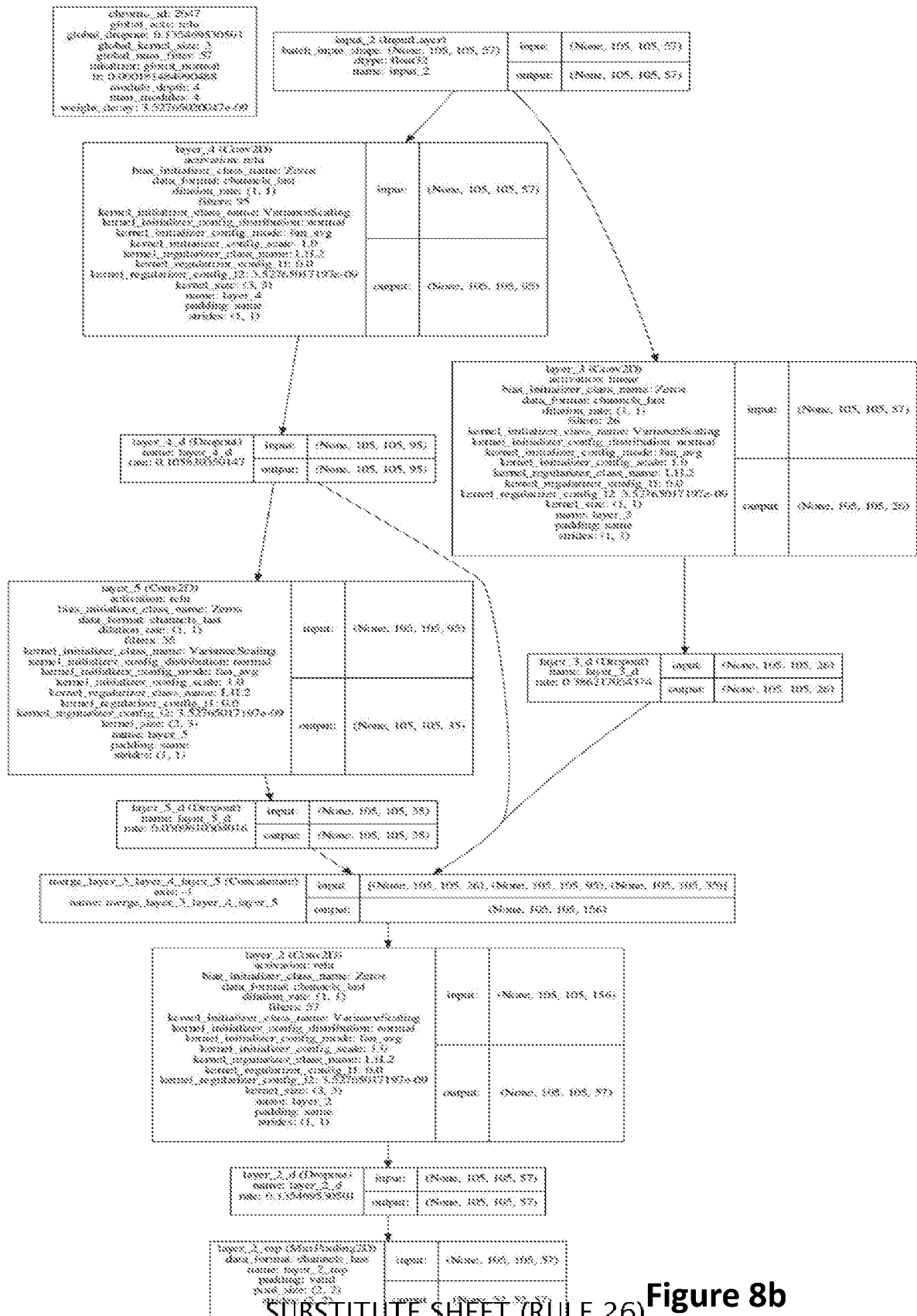
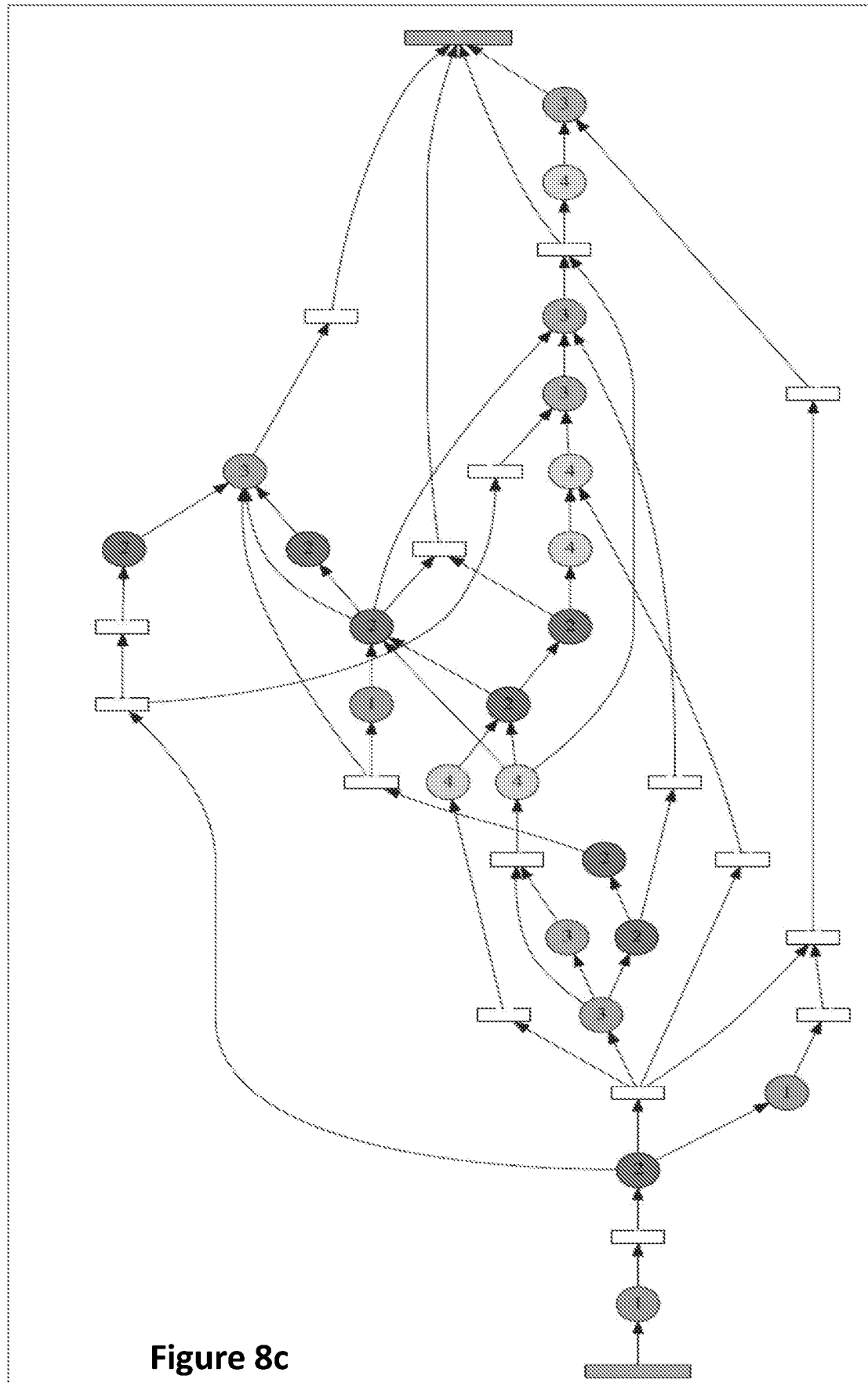


Figure 8a
 SUBSTITUTE SHEET (RULE 26)

SUBSTITUTE SHEET (RULE 26) **Figure 8b**



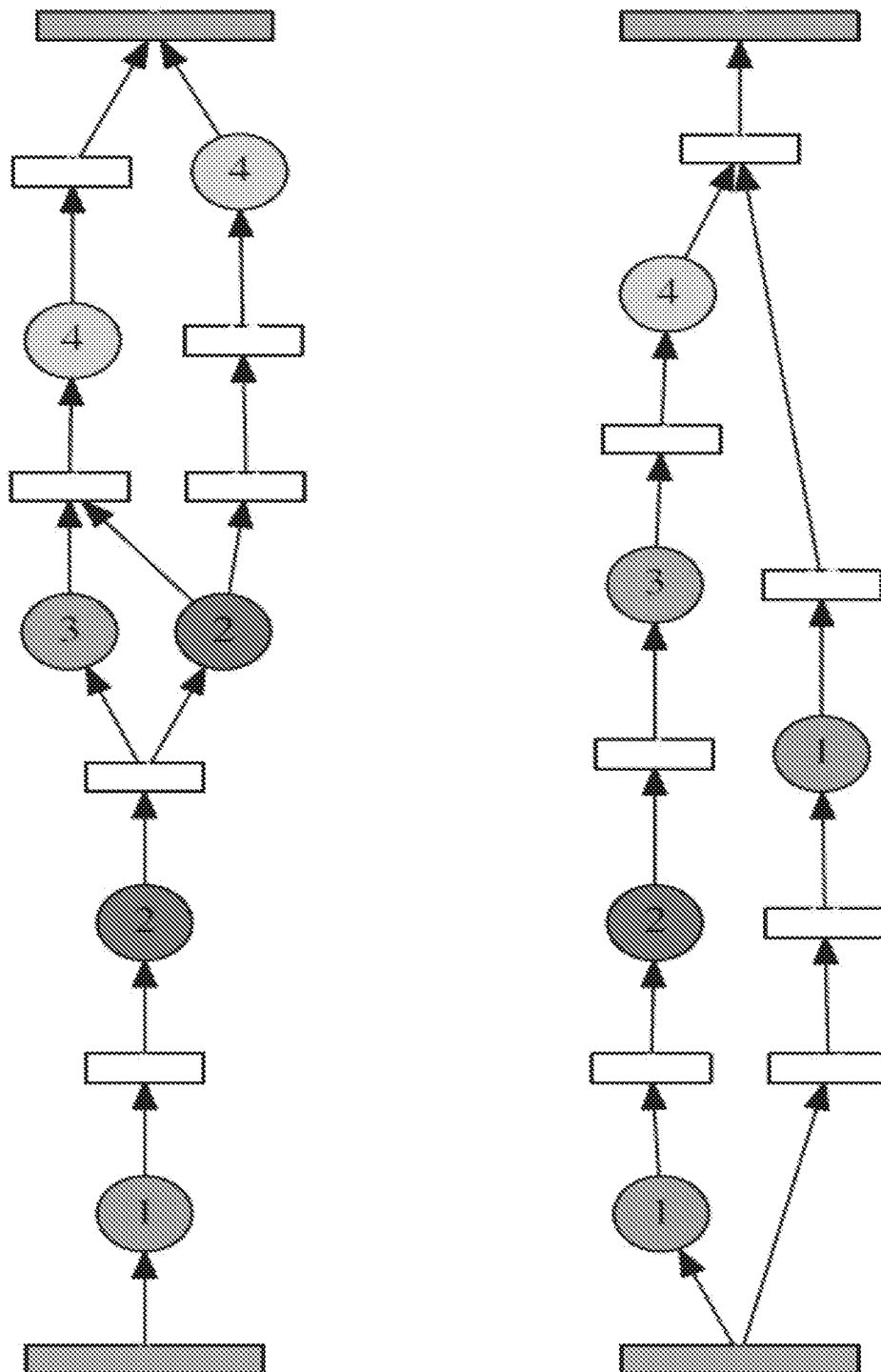


Figure 8d

Figure 1

