

(12) 특허협력조약에 의하여 공개된 국제출원

(19) 세계지식재산권기구
국제사무국

(43) 국제공개일

2025년 6월 19일 (19.06.2025)



(10) 국제공개번호

WO 2025/127262 A1

(51) 국제특허분류:

G06N 3/04 (2006.01)

G06N 3/084 (2023.01)

(21) 국제출원번호:

PCT/KR2024/005651

(22) 국제출원일:

2024년 4월 26일 (26.04.2024)

(25) 출원언어:

한국어

(26) 공개언어:

한국어

(30) 우선권정보:

10-2023-0178909 2023년 12월 11일 (11.12.2023) KR

(71) 출원인: 이화여자대학교 산학협력단 (EWha University - Industry Collaboration Foundation) [KR/KR]; 03760 서울특별시 서대문구 이화여대길 52 (KR).

(72) 발명자: 심재형 (SIM, Jae Hyeong); 04187 서울특별시 마포구 만리재로 74, 207동 801호 (KR). 임하영 (LIM, Ha Young); 03670 서울특별시 서대문구 증가로 150, 111동 1101호 (KR). 김주연 (KIM, Ju Yeon); 14407 경기도 부천시 오정구 역곡로504번길 52, 라동 401호 (KR). 장예서 (JANG, Ye Seo); 06284 서울특별시 강남구 삼성로 212, 19동 111호 (KR).

(74) 대리인: 김연권 (KIM, Youn Gwon); 05854 서울특별시 송파구 송파대로 201 B동 1416호 (KR).

(81) 지정국 (별도의 표시가 없는 한, 가능한 모든 종류의 국내 권리의 보호를 위하여): AE, AG, AL, AM, AO, AT, AU, AZ, BA, BB, BG, BH, BN, BR, BW, BY, BZ, CA, CH, CL, CN, CO, CR, CU, CV, CZ, DE, DJ, DK, DM, DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IQ, IR, IS, IT, JM, JO, JP, KE, KG, KH, KN, KP, KW, KZ, LA, LC, LK, LR, LS, LU, LY, MA, MD, MG, MK, MN, MU, 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, ST, SV, SY, TH, TJ, TM, TN, TR, TT, TZ, UA, UG, US, UZ, VC, VN, WS, ZA, ZM, ZW.

(84) 지정국 (별도의 표시가 없는 한, 가능한 모든 종류의 역내 권리의 보호를 위하여): ARIPO (BW, CV, GH, GM, KE, LR, LS, MW, MZ, NA, RW, SC, SD, SL, ST, SZ, TZ, UG, ZM, ZW), 유라시아 (AM, AZ, BY, KG, KZ, RU, TJ, TM), 유럽 (AL, AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI, FR, GB, GR, HR, HU, IE, IS, IT, LT, LU, LV, MC, ME, 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).

규칙 4.17에 의한 선언서:

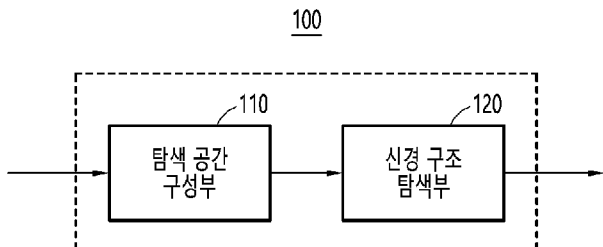
— 신규성을 해치지 아니하는 개시 또는 신규성 상실의 예외에 관한 선언 (규칙 4.17(v))

공개:

— 국제조사보고서와 함께 (조약 제21조(3))

(54) Title: NEURAL STRUCTURE SEARCH DEVICE BASED ON TEMPLATE AND METHOD THEREFOR

(54) 발명의 명칭: 템플릿에 기반하는 신경 구조 탐색장치 및 그 방법



110 ... Search space configuration unit
120 ... Neural structure search unit

(57) Abstract: The present invention relates to a neural structure search device and a method therefor. The neural structure search device according to one embodiment comprises: a search space configuration unit which configures a plurality of operator templates by matching, to each other, hardware blocks corresponding to the plurality of operators, and which configures a search space on the basis of the plurality of operator templates; and a neural structure search unit for deriving a target neural network through a neural search process based on the search space.

(57) 요약서: 본 발명은 신경 구조 탐색장치 및 그 방법에 관한 것으로서, 일 실시예에 따른 신경 구조 탐색장치는 복수의 연산자에 대응되는 하드웨어 블록을 각각 매칭하여 복수의 연산자 템플릿을 구성하고, 복수의 연산자 템플릿에 기초하여 탐색 공간을 구성하는 탐색 공간 구성부 및 탐색 공간에 기초하는 신경 탐색 과정을 통해 타겟 신경망을 도출하는 신경 구조 탐색부를 포함할 수 있다.

WO 2025/127262 A1

명세서

발명의 명칭: 템플릿에 기반하는 신경 구조 탐색장치 및 그 방법 기술분야

- [1] 본 발명은 신경 구조 탐색장치 및 그 방법에 관한 것으로, 보다 상세하게는 템플릿 구조로 구성된 탐색 공간을 이용하여 신경망 구조를 탐색하는 기술적 사상에 관한 것이다.
- [2] 본 출원은 2024년도 정부(과학기술정보통신부)의 재원으로 한국연구재단(No.RS-2023-00213548)과 정보통신기획평가원(No.RS-2022-00155966, 인공지능융합혁신인재양성(이화여자대학교), 국가과제고유번호: 1711179344, No.2021-0-02068, 인공지능 혁신 허브 연구 개발)의 지원을 받아 수행된 연구 결과이다.

배경기술

- [3] NAS(Neural Architecture Search)는 딥러닝 모델의 구조를 자동으로 찾아주는 AutoML(Automated Machine Learning)의 한 방식이다.
- [4] 기존 NAS는 모델의 정확도만을 고려하여 딥러닝 모델의 구조를 찾았기 때문에 하드웨어에서는 느리게 동작하는 경우가 많았으며, 하드웨어를 고려하더라도 특정 하드웨어에 맞춰서 동작하는 NAS 기법이 많았다.
- [5] 이에, 딥러닝 모델과 그 모델에 최적화된 하드웨어를 찾는 NAS 기법 (Hardware co-design NAS)에 대한 관심이 증가하고 있으나, 기존 하드웨어와 딥러닝 모델의 구조를 함께 찾는 NAS는 탐색 공간(Search Space)의 크기가 크다는 문제가 있다.
- [6] 구체적으로, 기존 기술들은 딥러닝 모델의 탐색 공간(일례로, 채널의 수, 층의 수 등)에 하드웨어 검색 공간(일례로, 프로세싱 요소(Processing Element)의 가로, 세로 크기, 버퍼 크기 등)이 추가되면서 탐색해야 하는 범위가 크게 증가했다.
- [7] 또한, 기존 기술들은 탐색 중에 후보 모델, 하드웨어 구조(즉, 하드웨어 블록)의 성능을 실시간으로 계산을 하는데, 이러한 방식은 시간이 오래 걸리고 부정확한 측정값을 가질 수 있으며, 이로 인해 하드웨어 후보들 중 최적의 구조가 아닌 구조를 도출하는 결과를 가져올 수 있다는 문제가 있다.

발명의 상세한 설명

기술적 과제

- [8] 본 발명은 각 연산에 대해 최적의 하드웨어 구조를 먼저 찾고 각 연산과 최적의 하드웨어 블록을 짝지은 연산자-하드웨어 쌍을 새로운 기본 요소를 정의하여 탐색 공간을 구성하는 신경 구조 탐색장치 및 그 방법을 제공하고자 한다.
- [9] 또한, 본 발명은 연산자-하드웨어 쌍에 기초하여 탐색 공간의 크기와, 신경망 탐색에 소요되는 시간 및 자원의 수를 최소화할 수 있는 신경 구조 탐색장치 및 그 방법을 제공하고자 한다.

과제 해결 수단

- [10] 본 발명의 일실시예에 따른 신경 구조 탐색장치는 복수의 연산자에 대응되는 하드웨어 블록을 각각 매칭하여 복수의 연산자 템플릿을 구성하고, 복수의 연산자 템플릿에 기초하여 탐색 공간을 구성하는 탐색 공간 구성부 및 탐색 공간에 기초하는 신경 탐색 과정을 통해 타겟 신경망을 도출하는 신경 구조 탐색부를 포함할 수 있다.
- [11] 일측에 따르면, 복수의 연산자는 스킵 연결(skip connection) 연산자, 평균 풀링(average pooling) 연산자, 최대 풀링(max pooling) 연산자, 표준 컨볼루션(standard convolution) 연산자 및 깊이별 컨볼루션(depthwise convolution) 연산자 중 적어도 하나를 포함할 수 있다.
- [12] 일측에 따르면, 탐색 공간 구성부는 복수의 연산자에 대응되는 하드웨어 블록을 100%의 활용도(utilization)를 갖도록 설계하여 복수의 연산자 템플릿을 구성할 수 있다.
- [13] 일측에 따르면, 신경 구조 탐색부는 셀 기반의 탐색 및 경사 하강법에 기초하는 신경 탐색 과정을 통해 타겟 신경망을 도출할 수 있다.
- [14] 일측에 따르면, 신경 구조 탐색부는 타겟 신경망을 구성하는 복수의 노드와, 복수의 연산자 템플릿 중 적어도 하나의 연산자 템플릿으로 구성되며 복수의 노드를 서로 연결하는 타겟 연산자 템플릿을 각각 포함하는 복수의 셀을 추출하고, 적어도 하나의 연산자 템플릿 각각의 가중치를 설정하며, 가중치를 업데이트하는 학습 과정의 결과로 타겟 신경망에 대응되는 적어도 하나의 후보 신경망을 도출할 수 있다.
- [15] 일측에 따르면, 신경 구조 탐색부는 복수의 셀에 적어도 하나의 완전 연결 레이어(fully connected layer) 및 소프트맥스 활성화 함수(softmax activation function)를 적용하여 후보 신경망을 도출할 수 있다.
- [16] 일측에 따르면, 신경 구조 탐색부는 후보 신경망에 대한 평가 과정에 기초하여 타겟 신경망을 도출할 수 있다.
- [17] 일측에 따르면, 신경 구조 탐색부는 후보 신경망에 대한 모델 정확도 및 하드웨어 비용(hardware cost)을 평가하는 평가 과정에 기초하여 타겟 신경망을 도출할 수 있다.
- [18] 일측에 따르면, 신경 구조 탐색부는 교차 엔트로피 손실(cross-entropy loss) 및 하드웨어 비용에 기초하는 손실 함수에 대응하여 타겟 신경망을 도출할 수 있다.
- [19] 본 발명의 일실시예에 따른 신경 구조 탐색방법은 탐색 공간 구성부에서, 복수의 연산자에 대응되는 하드웨어 블록을 각각 매칭하여 복수의 연산자 템플릿을 구성하고, 복수의 연산자 템플릿에 기초하여 탐색 공간을 구성하는 단계 및 신경 구조 탐색부에서, 탐색 공간에 기초하는 신경 탐색 과정을 통해 타겟 신경망을 도출하는 단계를 포함할 수 있다.

발명의 효과

[20] 일실시에에 따르면, 본 발명은 각 연산에 대해 최적의 하드웨어 구조를 먼저 찾고 각 연산과 최적의 하드웨어 블록을 짝지은 연산자-하드웨어 쌍을 새로운 기본 요소로 정의하여 탐색 공간을 구성할 수 있다.

[21] 일실시에에 따르면, 본 발명은 연산자-하드웨어 쌍에 기초하여 탐색 공간의 크기와, 신경망 탐색에 소요되는 시간 및 자원의 수를 최소화할 수 있다.

도면의 간단한 설명

[22] 도 1은 일실시에에 따른 신경 구조 탐색장치를 설명하기 위한 도면이다.

[23] 도 2a 내지 도 2f는 일실시에에 따른 신경 구조 탐색장치에서 하드웨어 블록을 설계하는 예시를 구체적으로 설명하기 위한 도면이다.

[24] 도 3은 일실시에에 따른 신경 구조 탐색방법을 설명하는 도면이다.

발명의 실시를 위한 최선의 형태

[25] 본 명세서에 개시되어 있는 본 발명의 개념에 따른 실시예들에 대해서 특정한 구조적 또는 기능적 설명들은 단지 본 발명의 개념에 따른 실시예들을 설명하기 위한 목적으로 예시된 것으로서, 본 발명의 개념에 따른 실시예들은 다양한 형태로 실시될 수 있으며 본 명세서에 설명된 실시예들에 한정되지 않는다.

[26] 본 발명의 개념에 따른 실시예들은 다양한 변경들을 가할 수 있고 여러 가지 형태들을 가질 수 있으므로 실시예들을 도면에 예시하고 본 명세서에 상세하게 설명하고자 한다. 그러나, 이는 본 발명의 개념에 따른 실시예들을 특정한 개시 형태들에 대해 한정하려는 것이 아니며, 본 발명의 사상 및 기술 범위에 포함되는 변경, 균등물, 또는 대체물을 포함한다.

[27] 제1 또는 제2 등의 용어를 다양한 구성요소들을 설명하는데 사용될 수 있지만, 구성요소들은 용어들에 의해 한정되어서는 안 된다. 용어들은 하나의 구성요소를 다른 구성요소로부터 구별하는 목적으로만, 예를 들면 본 발명의 개념에 따른 권리 범위로부터 이탈되지 않은 채, 제1 구성요소는 제2 구성요소로 명명될 수 있고, 유사하게 제2 구성요소는 제1 구성요소로도 명명될 수 있다.

[28] 어떤 구성요소가 다른 구성요소에 "연결되어" 있다거나 "접속되어" 있다고 언급된 때에는, 그 다른 구성요소에 직접적으로 연결되어 있거나 또는 접속되어 있을 수도 있지만, 중간에 다른 구성요소가 존재할 수도 있다고 이해되어야 할 것이다. 반면에, 어떤 구성요소가 다른 구성요소에 "직접 연결되어" 있다거나 "직접 접속되어" 있다고 언급된 때에는, 중간에 다른 구성요소가 존재하지 않는 것으로 이해되어야 할 것이다. 구성요소들 간의 관계를 설명하는 표현들, 예를 들면 "~사이에"와 "바로~사이에" 또는 "~에 직접 이웃하는" 등도 마찬가지로 해석되어야 한다.

[29] 본 명세서에서 사용한 용어는 단지 특정한 실시예들을 설명하기 위해 사용된 것으로, 본 발명을 한정하려는 의도가 아니다. 단수의 표현은 문맥상 명백하게 다르게 뜻하지 않는 한, 복수의 표현을 포함한다. 본 명세서에서, "포함하다" 또는 "가지다" 등의 용어는 실시된 특징, 숫자, 단계, 동작, 구성요소, 부분품 또는

이들을 조합한 것이 존재함으로 지정하려는 것이지, 하나 또는 그 이상의 다른 특징들이나 숫자, 단계, 동작, 구성요소, 부분품 또는 이들을 조합한 것들의 존재 또는 부가 가능성을 미리 배제하지 않는 것으로 이해되어야 한다.

[30] 다르게 정의되지 않는 한, 기술적이거나 과학적인 용어를 포함해서 여기서 사용되는 모든 용어들은 본 발명이 속하는 기술 분야에서 통상의 지식을 가진 자에 의해 일반적으로 이해되는 것과 동일한 의미를 가진다. 일반적으로 사용되는 사전에 정의되어 있는 것과 같은 용어들은 관련 기술의 문맥상 가지는 의미와 일치하는 의미를 갖는 것으로 해석되어야 하며, 본 명세서에서 명백하게 정의하지 않는 한, 이상적이거나 과도하게 형식적인 의미로 해석되지 않는다.

[31]

[32] 이하, 실시예들을 첨부된 도면을 참조하여 상세하게 설명한다. 그러나, 특허출원의 범위가 이러한 실시예들에 의해 제한되거나 한정되는 것은 아니다. 각 도면에 제시된 동일한 참조 부호는 동일한 부재를 나타낸다.

[33]

[34] 도 1은 일실시예에 따른 신경 구조 탐색장치를 설명하기 위한 도면이다.

[35] 도 1을 참조하면, 신경 구조 탐색장치(100)는 각 연산에 대해 최적의 하드웨어 구조를 먼저 찾고 각 연산과 최적의 하드웨어 블록을 짝지은 연산자-하드웨어 쌍을 새로운 기본 요소로 정의하여 탐색 공간을 구성할 수 있다.

[36] 또한, 신경 구조 탐색장치(100)는 연산자-하드웨어 쌍에 기초하여 탐색 공간의 크기와, 신경망 탐색에 소요되는 시간 및 자원의 수를 최소화할 수 있다.

[37] 이를 위해, 신경 구조 탐색장치(100)는 탐색 공간 구성부(110) 및 신경 구조 탐색부(120)를 포함할 수 있다.

[38] 일실시예에 따른 탐색 공간 구성부(110)는 복수의 연산자에 대응되는 하드웨어 블록을 각각 매칭하여 복수의 연산자 템플릿(즉, 연산자-하드웨어 쌍)을 구성하고, 복수의 연산자 템플릿에 기초하여 탐색 공간을 구성할 수 있다.

[39] 다시 말해, 탐색 공간 구성부(110)는 사전에 복수의 연산자에 대응하여 최적의 구조를 갖는 하드웨어 블록 각각을 도출하고, 도출된 하드웨어 블록 각각을 대응되는 복수의 연산자 각각에 매칭하여 복수의 연산자 템플릿을 구성할 수 있다.

[40] 구체적으로, 탐색 공간 구성부(110)는 하기 수식1과 같은 복수의 연산자 템플릿($\mathcal{T}$)에 기초하여 하기 수식2와 같은 탐색 공간($\mathcal{S}$)을 구성할 수 있다.

[41] [수식1]

$$[42] \mathcal{T}_i = (o_i, h_i)$$

[43] [수식2]

$$[44] \mathcal{S} = \{\mathcal{T}_0, \mathcal{T}_1, \dots, \mathcal{T}_N\}$$

[45] 여기서, o_i 는 i번째(여기서, i는 양의 정수) 연산자, h_i 는 연산자 o_i 를 구현하는 하드웨어 블록, N은 연산자 수를 의미한다.

- [46] 즉, 기존 기술은 연산을 먼저 실행하고 그 이후에 하드웨어를 탐색하는 방식을 사용하여 전체 탐색 공간이 연산의 탐색 공간과 하드웨어의 탐색공간의 곱으로 나타나 크기가 매우 컸으나, 본 발명의 복수의 연산자 템플릿에 기반하는 탐색 공간은 최적의 연산자-하드웨어 쌍을 템플릿으로 구성하여 연산의 수를 현저히 감소시킬 수 있으며, 이에 따라 신경 구조 탐색 프로세스를 연산을 실행하는 방식만 결정하면 되는 문제로 간소화하여, 탐색 시간과 필요한 자원의 수를 최소화할 수 있다.
- [47] 또한, 복수의 연산자 템플릿에 기반하는 탐색 공간은 신경 구조 탐색 전에 연산에 최적인 하드웨어의 구조를 찾아 놓기 때문에 정답 후보의 성능 측정을 계산을 통해 측정할 수 있으며, 이러한 계산을 사용하면 빠른 시간 내에 비교적 정확하게 하드웨어의 성능을 측정할 수 있어 가장 최적의 신경망(즉, 딥러닝 모델)과 이에 대응되는 하드웨어 구조를 찾을 수 있다.
- [48] 예를 들면, 복수의 연산자는 스킵 연결(skip connection) 연산자, 평균 풀링(average pooling) 연산자, 최대 풀링(max pooling) 연산자, 표준 컨볼루션(standard convolution) 연산자 및 깊이별 컨볼루션(depthwise convolution) 연산자 중 적어도 하나를 포함할 수 있다.
- [49] 보다 구체적인 예를 들면, 복수의 연산자는 하기 표1의 연산자를 포함할 수 있으나, 이에 한정되는 것은 아니다.
- [50] [표1]

연산자	설명
nop	No operation
skip_connect	스킵(잔차) 연결
avg_pool_3x3	3x3 평균 풀링
max_pool_3x3	3x3 최대 풀링
conv_1x1	1x1 표준 컨볼루션
conv_3x3	3x3 표준 컨볼루션
conv_5x5	5x5 표준 컨볼루션
dw_conv_3x3	3x3 깊이별 컨볼루션
dw_conv_5x5	5x5 깊이별 컨볼루션

- [51] 일측에 따르면, 탐색 공간 구성부(110)는 복수의 연산자에 대응되는 하드웨어 블록을 100%의 활용도(utilization)를 갖도록 설계하여 복수의 연산자 템플릿을 구성할 수 있다. 다시 말해, 탐색 공간 구성부(110)는 복수의 연산자 각각에 대응되는 하드웨어 블록의 각 구성요소들이 모두 활용될 수 있도록, 하드웨어 블록을 최적화 설계하여 효율을 극대화할 수 있다.
- [52] 일실시예에 따른 탐색 공간 구성부(110)에서 하드웨어 블록을 100%의 활용도를 갖도록 설계하는 방법은 이후 실시예 도 2a 내지 도 2f를 통해 보다 구체적으로 설명하기로 한다.
- [53] 일실시예에 따른 신경 구조 탐색부(120)는 탐색 공간에 기초하는 신경 탐색 과정을 통해 타겟 신경망(즉, 딥러닝 모델)을 도출할 수 있다.
- [54] 일측에 따르면, 신경 구조 탐색부(120)는 셀 기반의 탐색 및 경사 하강법(Gradient Descent)에 기초하는 신경 탐색 과정을 통해 타겟 신경망을 도출할 수 있으며, 이를 통해 탐색 비용을 절감하고 최적화 문제를 해결할 수 있다.
- [55] 여기서, 경사 하강법은 기계 학습과 최적화 문제에서 매개 변수를 조정하여 모델을 학습시키는데 사용되는 최적화 알고리즘 중 하나로, 주어진 함수의 기울기(경사)를 활용하여 함수의 최소값을 탐색하되, 기울기가 감소하는 방향으로 반복적으로 이동하여 최적화된 값을 도출할 수 있다.
- [56] 일측에 따르면, 신경 구조 탐색부(120)는 셀 추출 과정, 연산자 가중치 설정 과정, 가중치 갱신 과정, 연산자 선택 과정 및 연산자 평가 과정으로 구성되는 신경 탐색 과정을 통해 타겟 신경망을 도출할 수 있다.
- [57] 구체적으로, 신경 구조 탐색부(120)는 타겟 신경망을 구성하는 복수의 노드와, 복수의 연산자 템플릿 중 적어도 하나의 연산자 템플릿으로 구성되며 복수의 노드를 서로 연결하는 타겟 연산자 템플릿을 각각 포함하는 복수의 셀을 추출할 수 있다.
- [58] 예를 들면, 신경 구조 탐색부(120)는 노멀 셀(Normal Cell)과 리덕션 셀(Reduction Cell)을 포함하는 복수의 셀을 추출할 수 있으며, 여기서, 노멀 셀과 리덕션 셀은 기 설정된 순서대로 배치될 수 있다.
- [59] 보다 구체적인 예를 들면, 하나의 노멀 셀은 다른 2개의 노멀 셀 사이, 또는 다른 하나의 노멀 셀과 하나의 리덕션 셀 사이에 위치할 수 있고, 하나의 리덕션 셀은 2개의 노멀 셀 사이, 또는 M개(여기서, M은 양의 정수)의 노멀 셀마다 위치할 수도 있다.
- [60] 다음으로, 신경 구조 탐색부(120)는 적어도 하나의 연산자 템플릿 각각의 가중치를 설정할 수 있다.
- [61] 다시 말해, 신경 구조 탐색부(120)는 타겟 연산자 템플릿에 대한 후보 연산자 템플릿(즉, 탐색 공간에 구비된 복수의 연산자 템플릿 중 적어도 하나의 연산자 템플릿) 각각에 대한 가중치를 설정할 수 있다.
- [62] 예를 들면, 연산자 템플릿의 가중치는 너무 많은 후보 연산자가 사용되면 각 후보 연산자에 대응하는 가중치가 너무 작아서 학습이 제대로 수행되지 않으며, 검

색 공간이 커져서 학습 시간이 과도하게 오래 걸리는 문제가 발생할 수 있으므로 최적화된 값으로 설정될 수 있으며, 여기서 가중치의 최적화된 값은 관리자의 입력에 따라 설정될 수 있다.

[63] 다음으로, 신경 구조 탐색부(120)는 가중치를 업데이트하는 학습 과정(즉, 신경망 학습 과정)의 결과로 타겟 신경망에 대응되는 후보 신경망을 도출하고, 후보 신경망에 대한 평가 과정에 기초하여 타겟 신경망을 도출할 수 있다.

[64] 일측에 따르면, 신경 구조 탐색부(120)는 복수의 셀에 적어도 하나의 완전 연결 레이어(Fully Connected Layer) 및 소프트맥스 활성화 함수(Softmax Activation Function)를 적용하여 후보 신경망을 도출할 수 있다.

[65] 다시 말해, 신경 구조 탐색부(120)는 복수의 셀을 여러 번 쌓고 완전 연결 레이어와 소프트맥스 활성화 함수를 적용하여 완전한 네트워크를 구축할 수 있다.

[66] 예를 들면, 신경 구조 탐색부(120)는 타겟 신경망에 대응되는 학습 데이터를 수신하고, 수신한 학습 데이터에 기초하는 학습 과정을 통해 가중치를 업데이트할 수 있다.

[67] 또한, 신경 구조 탐색부(120)는 후보 신경망을 평가 세트(evaluation set)에 적용하는 평가 과정에 기초하여 타겟 신경망을 도출할 수 있으며, 여기서 평가 세트는 후보 신경망에 대한 모델 정확도 및 하드웨어 비용을 평가하는 수단을 포함할 수 있다.

[68] 보다 구체적인 예를 들면, 신경 구조 탐색부(120)는 교차 엔트로피 손실(cross-entropy loss) 및 하드웨어 비용에 기초하는 손실 함수에 대응하여 타겟 신경망을 도출할 수 있다.

[69] 구체적으로, 신경 구조 탐색부(120)는 하기 수식3에 기초하는 평가 과정을 통해 타겟 신경망을 도출할 수 있다.

[70] [수식3]

$$[71] \quad \min_{\alpha} \mathcal{L}_{val}(\mathbf{w}^*(\alpha), \alpha)$$

$$\text{s. t. } \mathbf{w}^*(\alpha) = \operatorname{argmin}_{\mathbf{w}} \mathcal{L}_{train}(\mathbf{w}, \alpha)$$

[72] 여기서, $\alpha \in \mathcal{S}$ 는 후보 연산자 템플릿의 선택에 따른 연산자 가중치, $\mathbf{w}$ 는 네트워크 가중치, $\mathbf{w}^*$ 는 최적의 가중치를 의미한다.

[73] 즉, 신경 구조 탐색부(120)는 최적의 가중치($\mathbf{w}^*$)를 수식3에 개시된 바와 같이 훈련 손실($\mathcal{L}_{train}$)을 최소화하여 도출할 수 있으며, 이때 훈련 손실($\mathcal{L}_{train}$)에 대응되는 손실 함수($\mathcal{L}$)는 하기 수식4를 통해 도출될 수 있다.

[74] [수식4]

$$[75] \quad \mathcal{L} = \mathcal{L}_{CE} + \lambda \mathcal{C}_{HW}$$

[76] 여기서, $\mathcal{L}_{CE}$ 는 교차 엔트로피 손실, $\mathcal{C}_{HW}$ 는 하드웨어 비용(일레로, 대기 시간, 면적, 에너지 소비), λ 는 하이퍼 파라미터를 의미한다.

- [77] 한편, 일실시예에 따른 신경 구조 탐색장치(즉, TD-NAAS)(100)는 하기 표2에 개시된 바와 같이, 기존 기술(즉, EDD, DANCE) 대비 정확도는 0.8% 및 0.2% 증가하고, 탐색 속도는 2.75배 빨라진 것을 확인할 수 있다.
- [78] 여기서, 표3의 실험 결과는 학습 데이터로 CIFAR-10 데이터셋을 사용하고, 연산자에 최적인 하드웨어 구조는 Synopsys Design Compiler를 사용하였으며, 합성 후 RTL 시뮬레이션 결과를 하드웨어 성능 평가 값으로 사용하고, 딥러닝 모델은 파이토치를 사용했고, NVIDIA RTX A6000 GPU 1개로 훈련한 결과를 나타낸다.
- [79] [표3]

Algorithm	Accuracy(%)	Latency(ms)
EDD	94.4	N/A
DANCE	95	11
TD-NAAS	95.2	4

- [80] 도 2a 내지 도 2f는 일실시예에 따른 신경 구조 탐색장치에서 하드웨어 블록을 설계하는 예시를 구체적으로 설명하기 위한 도면이다.
- [81] 도 2a 내지 도 2f를 참조하면, 도면부호 210은 일실시예에 따른 신경 구조 탐색장치에서 3x3 표준 컨볼루션 연산자에 대응하여 하드웨어 블록을 최적화 설계하는 방법을 설명하기 위한 도면이고, 도면부호 220 내지 240은 2x2 표준 컨볼루션 연산자에 대응하여 하드웨어 블록을 최적화 설계하는 방법을 설명하기 위한 도면이다.
- [82] 또한, 도면부호 250은 신경 구조 탐색장치에서 3x3 평균 풀링 연산자에 대응하여 하드웨어 블록을 최적화 설계하는 방법을 설명하기 위한 도면이며, 도면부호 260은 신경 구조 탐색장치에서 3x3 최대 풀링 연산자에 대응하여 하드웨어 블록을 최적화 설계하는 방법을 설명하기 위한 도면이다.
- [83] 구체적으로, 일실시예에 따른 신경 구조 탐색장치는 복수의 연산자에 대응되는 하드웨어 블록을 각각 매칭하여 복수의 연산자 템플릿을 구성하고, 복수의 연산자 템플릿에 기초하여 탐색 공간을 구성할 수 있다.
- [84] 다시 말해, 신경 구조 탐색장치는 연산자와 이에 대응되는 하드웨어 블록을 하나의 템플릿으로 묶어서 새로운 기본 요소로 정의하고, 이러한 기본 요소로 탐색 공간을 구성함으로써, 기존 기술 대비 탐색 공간의 크기를 줄여 효율적인 탐색이 가능해지며, 이로 인해 보다 좋은 성능의 신경망 구조가 도출되도록 지원할 수 있다.
- [85] 또한, 신경 구조 탐색장치는 연산자에 대응되는 하드웨어 블록이 항상 100% 활용도를 유지하도록 최적화 설계함으로써, 신경망 구조를 보다 효율적으로 처리하도록 지원할 수 있다.

- [86] 예를 들면, 신경 구조 탐색장치는 도면부호 210에 도시된 바와 같이, 연산자가 3x3 표준 컨볼루션 연산자이면, 9개의 활성값(I0 내지 I8)과 9개의 가중치(W0 내지 W8)에 대응되는 9개의 곱셈기와 이를 모두 더하는 덧셈-트리를 포함하는 하드웨어 블록을 구성하고, 3x3 컨볼루션 윈도우(convolution window)를 일렬로 펼쳐 하드웨어 블록에 공급할 수 있으며, 이를 통해 하드웨어 블록은 활용도가 항상 100%가 되도록 설계될 수 있다.
- [87] 보다 구체적으로, 도면부호 220 내지 240에 도시된 2x2 컨볼루션 연산자로 예시하면, 신경 구조 탐색장치는 4개의 곱셈기를 만들고 4:1 덧셈-트리를 구성하면(도면부호 220), 곱셈기(즉, 4개의 곱셈기)가 모두 동작하므로 하드웨어 블록의 활용도가 항상 100%가 되도록 설계될 수 있다.
- [88] 이 경우, 컨볼루션 윈도우가 (I0, I1, I5, I6)에서 (I1, I2, I6, I7)으로 슬라이딩(sliding) 되더라도(도면부호 230), 하드웨어 블록의 활용도가 여전히 100%를 유지하는 것을 확인할 수 있다.
- [89] 반면, 일실시에에 따른 신경 구조 탐색장치와 같이 2x2 컨볼루션 연산자에 대응되는 하드웨어 블록을 4개의 곱셈기로 최적화 설계하지 않고 5개의 곱셈기로 설계하여 5:1 덧셈-트리를 구성하면(도면부호 240), 80%의 곱셈기(즉, 5개 중 4개의 곱셈기)만 동작하므로 하드웨어 블록의 활용도가 80%가 되어 비효율이 발생하는 것을 확인할 수 있다.
- [90] 예를 들면, 신경 구조 탐색장치는 도면부호 250에 도시된 바와 같이, 연산자가 3x3 평균 풀링 연산자이면, 9개의 활성값(I0 내지 I8)에 대응되는 덧셈-트리 및 최종 곱셈기를 포함하는 하드웨어 블록을 구성하고, 3x3 풀링 윈도우(pooling window)를 일렬로 펼쳐 하드웨어 블록에 공급할 수 있으며, 이를 통해 하드웨어 블록은 활용도가 항상 100%가 되도록 설계될 수 있다.
- [91] 또한, 신경 구조 탐색장치는 도면부호 250에 도시된 바와 같이, 연산자가 3x3 최대 풀링 연산자이면, 9개의 활성값(I0 내지 I8)에 대응되는 9개의 비교기-트리를 포함하는 하드웨어 블록을 구성하고, 3x3 풀링 윈도우를 일렬로 펼쳐 하드웨어에 공급할 수 있으며, 이를 통해 하드웨어 블록은 활용도가 항상 100%가 되도록 설계될 수 있다.
- [92] 한편, 도 2a 내지 도 2f에서는 설명의 편의 상 복수의 연산자를 3x3 표준 컨볼루션 연산자, 2x2 표준 컨볼루션 연산자, 3x3 평균 풀링 연산자 및 3x3 최대 풀링 연산자로 예시하였으나, 일실시에에 따른 복수의 연산자는 이에 한정되지 않고 NxN(여기서, N은 양의 정수) 표준 컨볼루션 연산자, NxN 평균 풀링 연산자, NxN 최대 풀링 연산자를 포함할 수 있다.
- [93] 구체적으로, 신경 구조 탐색장치는 연산자가 NxN 표준 컨볼루션 연산자이면, N² 개의 곱셈기와, 곱셈기의 출력을 모두 더하는 덧셈-트리를 포함하는 하드웨어 블록을 구성하고, NxN 컨볼루션 윈도우를 일렬로 펼쳐 하드웨어 블록에 공급할 수 있으며, 이를 통해 하드웨어 블록은 활용도가 항상 100%가 되도록 설계될 수 있다.

- [94] 또한, 신경 구조 탐색장치는 연산자가 $N \times N$ 평균 풀링 연산자이면, N^2 개의 활성 값 덧셈-트리와, 최종 곱셈기를 포함하는 하드웨어 블록을 구성하고, $N \times N$ 풀링 윈도우를 일렬로 펼쳐 하드웨어 블록에 공급(곱셈기에는 $1/N^2$ 을 공급)하고 할 수 있으며, 이를 통해 하드웨어 블록은 활용도가 항상 100%가 되도록 설계될 수 있다.
- [95] 또한, 신경 구조 탐색장치는 연산자가 $N \times N$ 최대 풀링 연산자이면, N^2 개의 비교기-트리를 포함하는 하드웨어 블록을 구성하고, $N \times N$ 풀링 윈도우를 일렬로 펼쳐 하드웨어 블록에 공급할 수 있으며, 이를 통해 하드웨어 블록은 활용도가 항상 100%가 되도록 설계될 수 있다.
- [96]
- [97] 도 3은 일실시예에 따른 신경 구조 탐색방법을 설명하는 도면이다.
- [98] 다시 말해, 도 3은 도 1 내지 도 2f를 통해 설명한 일실시예에 따른 신경 구조 탐색장치의 동작방법을 설명하는 도면일 수 있다.
- [99] 310 단계에서 신경 구조 탐색방법은 탐색 공간 구성부에서, 복수의 연산자에 대응되는 하드웨어 블록을 각각 매칭하여 복수의 연산자 템플릿을 구성하고, 복수의 연산자 템플릿에 기초하여 탐색 공간을 구성할 수 있다.
- [100] 다시 말해, 신경 구조 탐색방법은 탐색 공간 구성부에서, 복수의 연산자에 대응하여 최적의 구조를 갖는 하드웨어 블록 각각을 사전에 매칭하여 복수의 연산자 템플릿을 구성할 수 있다.
- [101] 예를 들면, 복수의 연산자는 스킵 연결(skip connection) 연산자, 평균 풀링(average pooling) 연산자, 최대 풀링(max pooling) 연산자, 표준 컨볼루션(standard convolution) 연산자 및 깊이별 컨볼루션(depthwise convolution) 연산자 중 적어도 하나를 포함할 수 있다.
- [102] 일측에 따르면, 310 단계에서 신경 구조 탐색방법은 탐색 공간 구성부에서, 복수의 연산자에 대응되는 하드웨어 블록을 100%의 활용도(utilization)를 갖도록 설계하여 복수의 연산자 템플릿을 구성할 수 있다.
- [103] 320 단계에서 신경 구조 탐색방법은 신경 구조 탐색부에서, 탐색 공간에 기초하는 신경 탐색 과정을 통해 타겟 신경망(즉, 딥러닝 모델)을 도출할 수 있다.
- [104] 일측에 따르면, 320 단계에서 신경 구조 탐색방법은 신경 구조 탐색부에서, 셀 기반의 탐색 및 경사 하강법(Gradient Descent)에 기초하는 신경 탐색 과정을 통해 타겟 신경망을 도출할 수 있으며, 이를 통해 탐색 비용을 절감하고 최적화 문제를 해결할 수 있다.
- [105] 예를 들면, 320 단계에서 신경 구조 탐색방법은 신경 구조 탐색부에서, 셀 추출 과정, 연산자 가중치 설정 과정, 가중치 갱신 과정, 연산자 선택 과정 및 연산자 평가 과정으로 구성되는 신경 탐색 과정을 통해 타겟 신경망을 도출할 수 있다.
- [106] 구체적으로, 320 단계에서 신경 구조 탐색방법은 신경 구조 탐색부에서, 타겟 신경망을 구성하는 복수의 노드와, 복수의 연산자 템플릿 중 적어도 하나의 연산

자 템플릿으로 구성되며 복수의 노드를 서로 연결하는 타겟 연산자 템플릿을 각각 포함하는 복수의 셀을 추출할 수 있다.

[107] 예를 들면, 320 단계에서 신경 구조 탐색방법은 신경 구조 탐색부에서, 노멀 셀(Normal Cell)과 리덕션 셀(Reduction Cell)을 포함하는 복수의 셀을 추출할 수 있다.

[108] 다음으로, 320 단계에서 신경 구조 탐색방법은 신경 구조 탐색부에서, 적어도 하나의 연산자 템플릿 각각의 가중치를 설정할 수 있다.

[109] 다시 말해, 320 단계에서 신경 구조 탐색방법은 신경 구조 탐색부에서, 타겟 연산자 템플릿에 대한 후보 연산자 템플릿(즉, 탐색 공간에 구비된 복수의 연산자 템플릿 중 적어도 하나의 연산자 템플릿) 각각에 대한 가중치를 설정할 수 있다.

[110] 다음으로, 320 단계에서 신경 구조 탐색방법은 신경 구조 탐색부에서, 가중치를 업데이트하는 학습 과정(즉, 신경망 학습 과정)의 결과로 타겟 신경망에 대응되는 후보 신경망을 도출하고, 후보 신경망에 대한 평가 과정에 기초하여 타겟 신경망을 도출할 수 있다.

[111] 예를 들면, 320 단계에서 신경 구조 탐색방법은 신경 구조 탐색부에서, 복수의 셀에 적어도 하나의 완전 연결 레이어(Fully Connected Layer) 및 소프트맥스 활성화 함수(Softmax Activation Function)를 적용하여 후보 신경망을 도출할 수 있다.

[112] 또한, 320 단계에서 신경 구조 탐색방법은 신경 구조 탐색부에서, 타겟 신경망에 대응되는 학습 데이터를 수신하고, 수신한 학습 데이터에 기초하는 학습 과정을 통해 가중치를 업데이트할 수 있다.

[113] 또한, 320 단계에서 신경 구조 탐색방법은 신경 구조 탐색부에서, 후보 신경망을 평가 세트(evaluation set)에 적용하는 평가 과정에 기초하여 타겟 신경망을 도출할 수 있으며, 여기서 평가 세트는 후보 신경망에 대한 모델 정확도 및 하드웨어 비용을 평가하는 수단을 포함할 수 있다.

[114] 보다 구체적인 예를 들면, 320 단계에서 신경 구조 탐색방법은 신경 구조 탐색부에서, 교차 엔트로피 손실(cross-entropy loss) 및 하드웨어 비용에 기초하는 손실 함수에 대응하여 타겟 신경망을 도출할 수 있다.

[115]

[116] 결국, 본 발명을 이용하면, 각 연산에 대해 최적의 하드웨어 구조를 먼저 찾고 각 연산과 최적의 하드웨어 블록을 짝지은 연산자-하드웨어 쌍을 새로운 기본 요소로 정의하여 탐색 공간을 구성할 수 있다.

[117] 또한, 본 발명을 이용하면, 연산자-하드웨어 쌍에 기초하여 탐색 공간의 크기와, 신경망 탐색에 소요되는 시간 및 자원의 수를 최소화할 수 있다.

[118]

[119] 이상과 같이 실시예들이 비록 한정된 도면에 의해 설명되었으나, 해당 기술분야에서 통상의 지식을 가진 자라면 상기의 기재로부터 다양한 수정 및 변형이 가능하다. 예를 들면, 설명된 기술들이 설명된 방법과 다른 순서로 수행되거나, 및/또는 설명된 장치, 구조, 장치, 회로 등의 구성요소들이 설명된 방법과 다른 형태

로 결합 또는 조합되거나, 다른 구성요소 또는 균등물에 의하여 대치되거나 치환되더라도 적절한 결과가 달성될 수 있다.

[120] 그러므로, 다른 구현들, 다른 실시예들 및 특허청구범위와 균등한 것들도 후술하는 특허청구범위의 범위에 속한다.

청구범위

- [청구항 1] 복수의 연산자에 대응되는 하드웨어 블록을 각각 매칭하여 복수의 연산자 템플릿을 구성하고, 상기 복수의 연산자 템플릿에 기초하여 탐색 공간을 구성하는 탐색 공간 구성부 및
상기 탐색 공간에 기초하는 신경 탐색 과정을 통해 타겟 신경망을 도출하는 신경 구조 탐색부
를 포함하는 신경 구조 탐색장치.
- [청구항 2] 제1항에 있어서,
상기 복수의 연산자는,
스킵 연결(skip connection) 연산자, 평균 풀링(average pooling) 연산자, 최대 풀링(max pooling) 연산자, 표준 컨볼루션(standard convolution) 연산자 및 깊이별 컨볼루션(depthwise convolution) 연산자 중 적어도 하나를 포함하는
신경 구조 탐색장치.
- [청구항 3] 제1항에 있어서,
상기 탐색 공간 구성부는,
상기 복수의 연산자에 대응되는 하드웨어 블록을 100%의 활용도(utilization)를 갖도록 설계하여 상기 복수의 연산자 템플릿을 구성하는
신경 구조 탐색장치.
- [청구항 4] 제1항에 있어서,
상기 신경 구조 탐색부는,
셀 기반의 탐색 및 경사 하강법에 기초하는 상기 신경 탐색 과정을 통해
상기 타겟 신경망을 도출하는
신경 구조 탐색장치.
- [청구항 5] 제1항에 있어서,
상기 신경 구조 탐색부는,
상기 타겟 신경망을 구성하는 복수의 노드와, 상기 복수의 연산자 템플릿 중 적어도 하나의 연산자 템플릿으로 구성되며 상기 복수의 노드를 서로 연결하는 타겟 연산자 템플릿을 각각 포함하는 복수의 셀을 추출하고, 상기 적어도 하나의 연산자 템플릿 각각의 가중치를 설정하며, 상기 가중치를 업데이트하는 학습 과정의 결과로 상기 타겟 신경망에 대응되는 적어도 하나의 후보 신경망을 도출하는
신경 구조 탐색장치.
- [청구항 6] 제5항에 있어서,
상기 신경 구조 탐색부는,

상기 복수의 셀에 적어도 하나의 완전 연결 레이어(fully connected layer) 및 소프트맥스 활성화 함수(softmax activation function)를 적용하여 상기 후보 신경망을 도출하는
신경 구조 탐색장치.

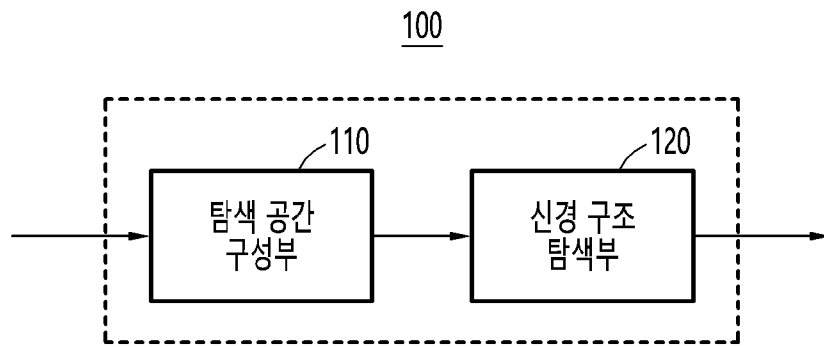
[청구항 7] 제5항에 있어서,
상기 신경 구조 탐색부는,
상기 후보 신경망에 대한 평가 과정에 기초하여 상기 타겟 신경망을 도출하는
신경 구조 탐색장치.

[청구항 8] 제7항에 있어서,
상기 신경 구조 탐색부는,
상기 후보 신경망에 대한 모델 정확도 및 하드웨어 비용(hardware cost)을 평가하는 상기 평가 과정에 기초하여 상기 타겟 신경망을 도출하는
신경 구조 탐색장치.

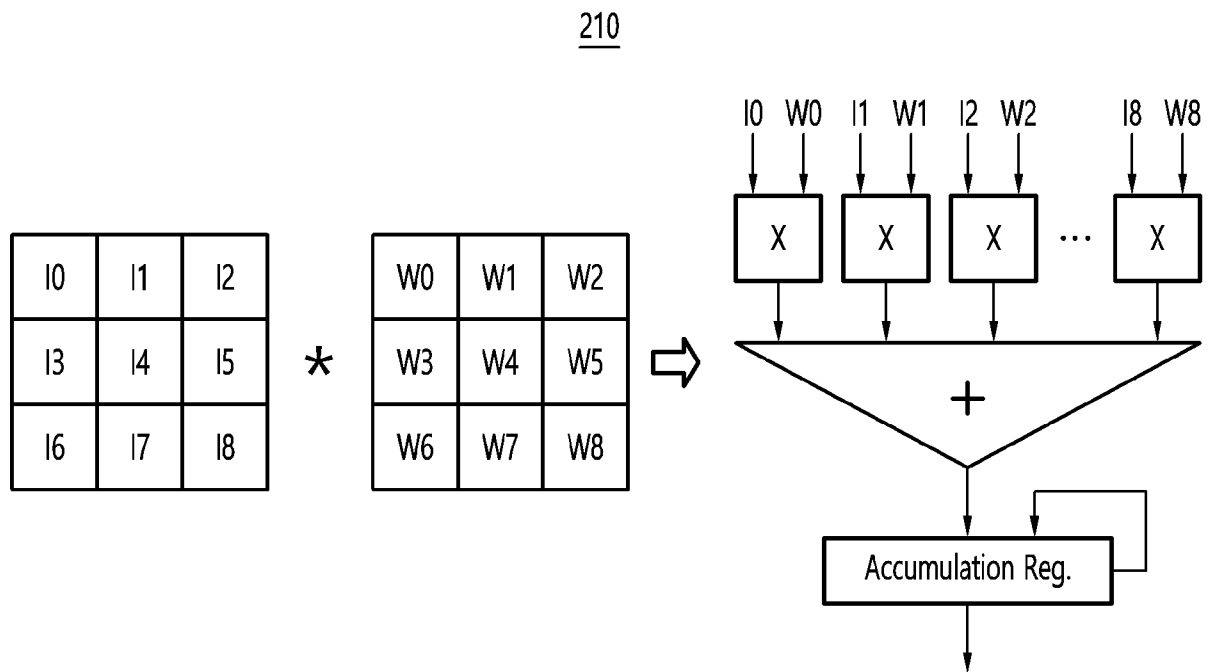
[청구항 9] 제7항에 있어서,
상기 신경 구조 탐색부는,
교차 엔트로피 손실(cross-entropy loss) 및 하드웨어 비용에 기초하는 손실 함수에 대응하여 상기 타겟 신경망을 도출하는
신경 구조 탐색장치.

[청구항 10] 탐색 공간 구성부에서, 복수의 연산자에 대응되는 하드웨어 블록을 각각 매칭하여 복수의 연산자 템플릿을 구성하고, 상기 복수의 연산자 템플릿에 기초하여 탐색 공간을 구성하는 단계 및
신경 구조 탐색부에서, 상기 탐색 공간에 기초하는 신경 탐색 과정을 통해 타겟 신경망을 도출하는 단계
를 포함하는 신경 구조 탐색방법.

[도1]

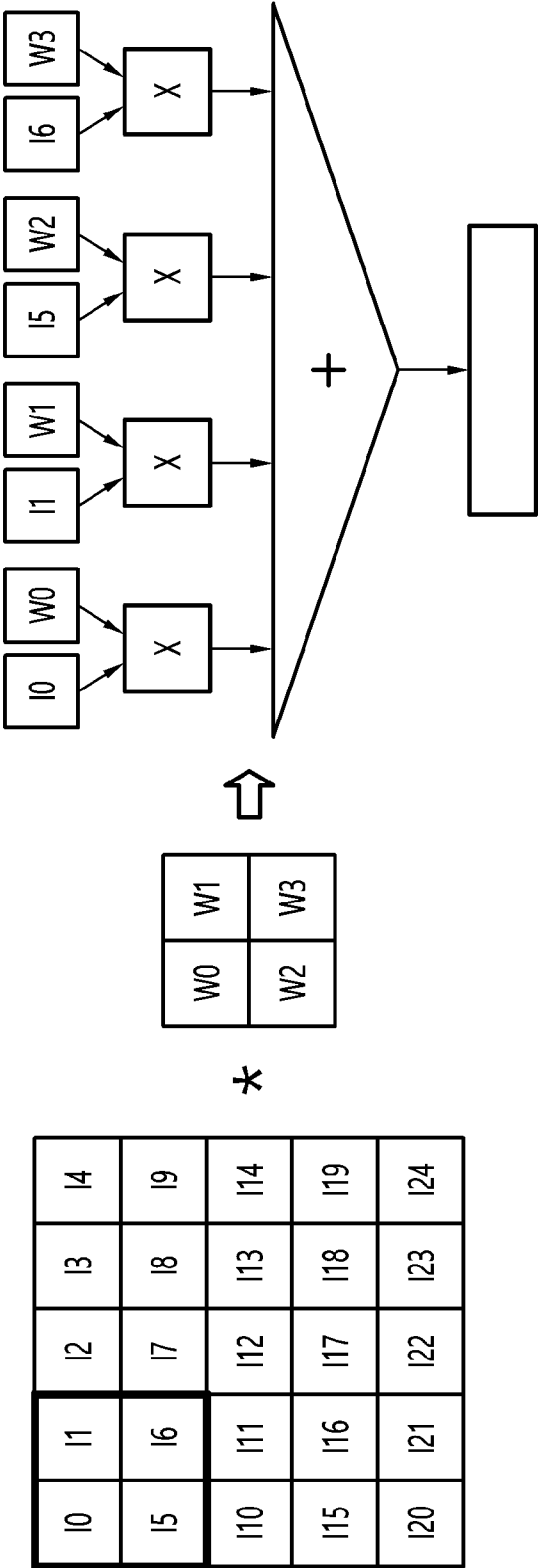


[도2a]



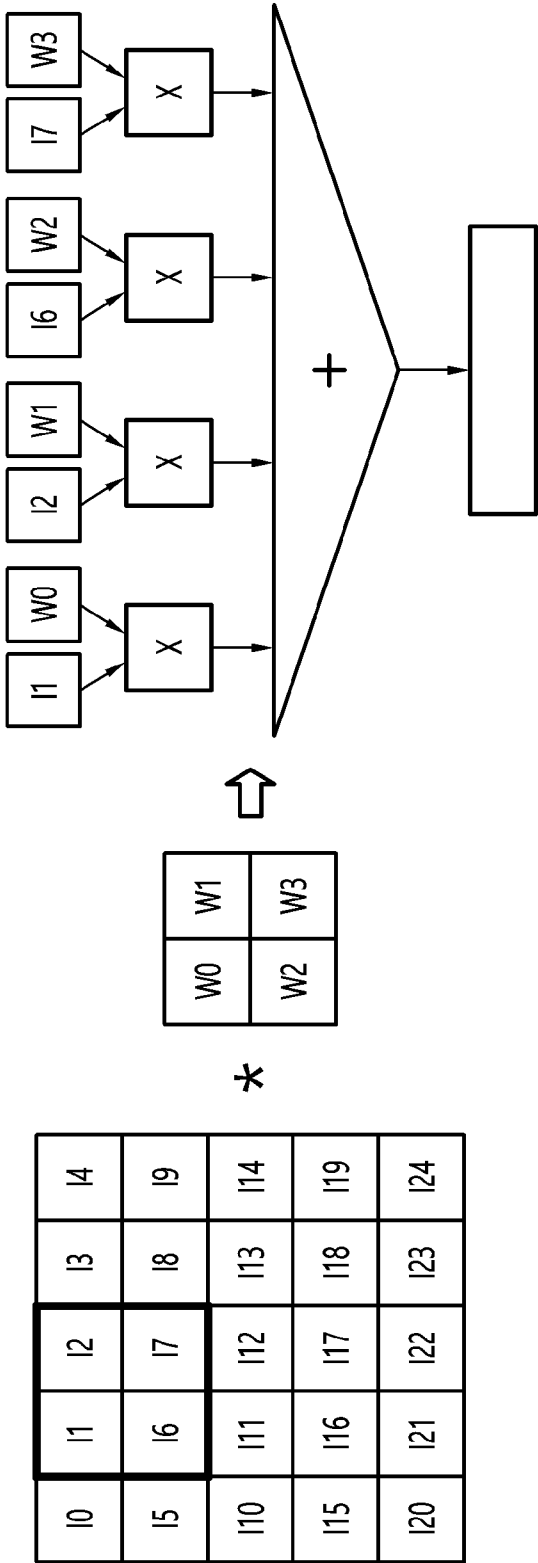
[도2b]

220



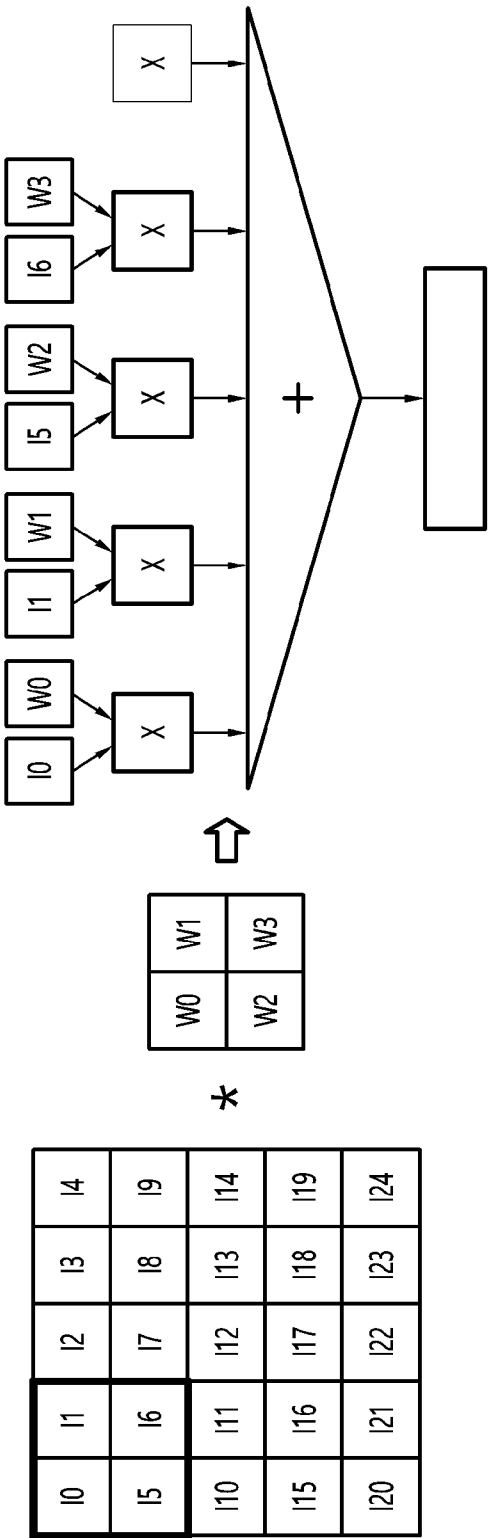
[도2c]

230

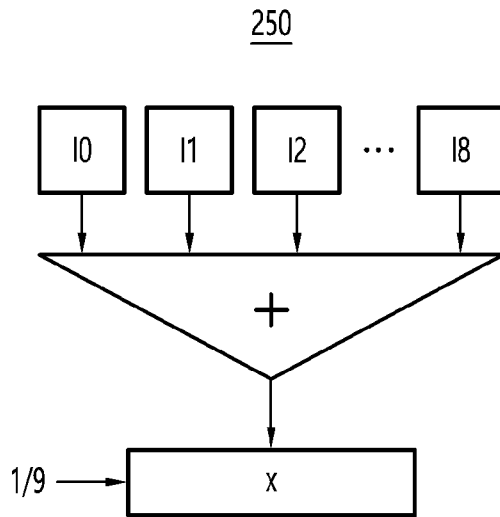


[도2d]

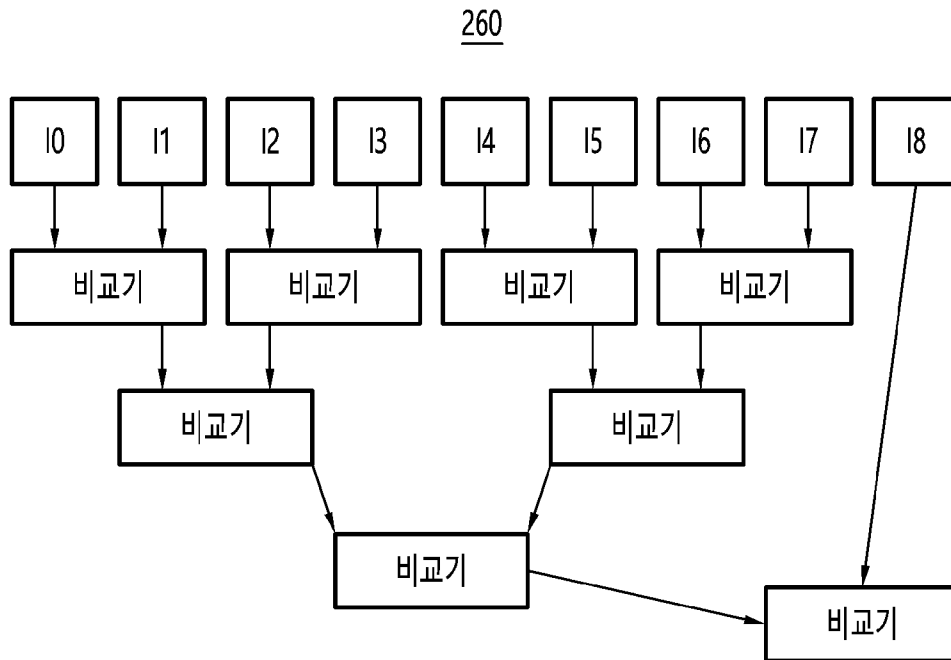
240



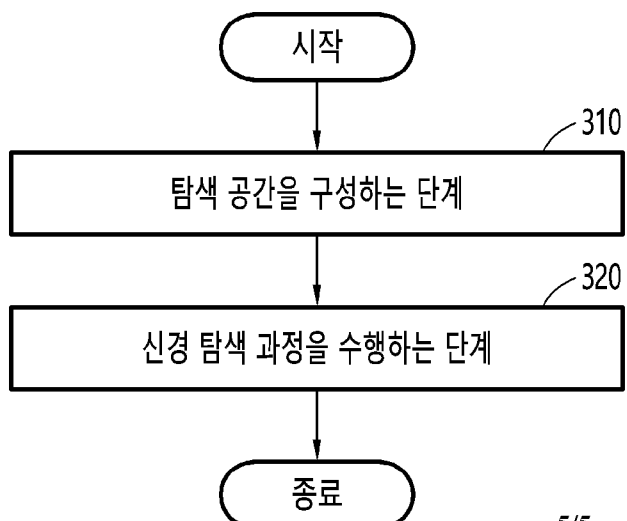
[도2e]



[도2f]



[도3]



INTERNATIONAL SEARCH REPORT

International application No.

PCT/KR2024/005651

A. CLASSIFICATION OF SUBJECT MATTER**G06N 3/04**(2006.01)i; **G06N 3/084**(2023.01)i

According to International Patent Classification (IPC) or to both national classification and IPC

B. FIELDS SEARCHED

Minimum documentation searched (classification system followed by classification symbols)

G06N 3/04(2006.01); G06F 15/18(2006.01); G06F 8/41(2018.01); G06F 9/48(2006.01); G06F 9/50(2006.01);
G06F 9/52(2006.01); G06N 3/065(2023.01); G06N 3/08(2006.01)

Documentation searched other than minimum documentation to the extent that such documents are included in the fields searched

Korean utility models and applications for utility models: IPC as above
Japanese utility models and applications for utility models: IPC as above

Electronic data base consulted during the international search (name of data base and, where practicable, search terms used)

eKOMPASS (KIPO internal) & keywords: 신경망(neural network), 연산자(operator), 하드웨어 블록(hardware block), 매칭(matching), 탐색 공간(search space), 도출(extracting)

C. DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
X	CN 116306856 A (ZHEJIANG LAB) 23 June 2023 (2023-06-23) See paragraph [0039]; and claim 1.	1-4,10
Y		5-9
Y	CN 103617146 B (BEIJING QIHOO TECHNOLOGY CO., LTD. et al.) 13 October 2017 (2017-10-13) See claim 1.	5-9
A	US 2019-0332441 A1 (HEWLETT PACKARD ENTERPRISE DEVELOPMENT LP) 31 October 2019 (2019-10-31) See paragraphs [0046]-[0051]; and figures 8A-8B.	1-10
A	KR 10-2021-0113004 A (SAMSUNG ELECTRONICS CO., LTD. et al.) 15 September 2021 (2021-09-15) See paragraphs [0177]-[0188]; and figure 13.	1-10



Further documents are listed in the continuation of Box C.



See patent family annex.

* Special categories of cited documents:

“A” document defining the general state of the art which is not considered to be of particular relevance

“D” document cited by the applicant in the international application

“E” earlier application or patent but published on or after the international filing date

“L” document which may throw doubts on priority claim(s) or which is cited to establish the publication date of another citation or other special reason (as specified)

“O” document referring to an oral disclosure, use, exhibition or other means

“P” document published prior to the international filing date but later than the priority date claimed

“T” later document published after the international filing date or priority date and not in conflict with the application but cited to understand the principle or theory underlying the invention

“X” document of particular relevance; the claimed invention cannot be considered novel or cannot be considered to involve an inventive step when the document is taken alone

“Y” document of particular relevance; the claimed invention cannot be considered to involve an inventive step when the document is combined with one or more other such documents, such combination being obvious to a person skilled in the art

“&” document member of the same patent family

Date of the actual completion of the international search

03 September 2024

Date of mailing of the international search report

03 September 2024

Name and mailing address of the ISA/KR

**Korean Intellectual Property Office
Government Complex-Daejeon Building 4, 189 Cheongsaro, Seo-gu, Daejeon 35208**

Facsimile No. +82-42-481-8578

Authorized officer

Telephone No.

INTERNATIONAL SEARCH REPORT

International application No.

PCT/KR2024/005651

C. DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
A	KR 10-2022-0041926 A (GOOGLE LLC) 01 April 2022 (2022-04-01) See claims 1-10.	1-10
X	LIM, Ha Young et al. TD-NAAS: Template-Based Differentiable Neural Architecture Accelerator Search. 2023 20th International SoC Design Conference (ISOCC). 26 October 2023. See sections 1-4. ※ This document is a known document declaring exceptions to lack of novelty by the applicant.	1-10

INTERNATIONAL SEARCH REPORT
Information on patent family members

International application No.

PCT/KR2024/005651

Patent document cited in search report			Publication date (day/month/year)	Patent family member(s)		Publication date (day/month/year)	
CN	116306856	A	23 June 2023	CN	116306856	B	05 September 2023
CN	103617146	B	13 October 2017	CN	103617146	A	05 March 2014
US	2019-0332441	A1	31 October 2019	US	10698737	B2	30 June 2020
KR	10-2021-0113004	A	15 September 2021	CN	113361704	A	07 September 2021
				US	2021-0279587	A1	09 September 2021
KR	10-2022-0041926	A	01 April 2022	CN	114258538	A	29 March 2022
				CN	114258538	B	12 April 2024
				EP	3980888	A1	13 April 2022
				JP	2022-544796	A	21 October 2022
				JP	7342247	B2	11 September 2023
				TW	202113596	A	01 April 2021
				TW	I764236	B	11 May 2022
				US	2022-0326988	A1	13 October 2022
				WO	2021-034675	A1	25 February 2021

A. 발명이 속하는 기술분류(국제특허분류(IPC)) G06N 3/04(2006.01)i; G06N 3/084(2023.01)i		
B. 조사된 분야 조사된 최소문헌(국제특허분류를 기재) G06N 3/04(2006.01); G06F 15/18(2006.01); G06F 8/41(2018.01); G06F 9/48(2006.01); G06F 9/50(2006.01); G06F 9/52(2006.01); G06N 3/065(2023.01); G06N 3/08(2006.01) 조사된 기술분야에 속하는 최소문헌 이외의 문헌 한국등록실용신안공보 및 한국공개실용신안공보: 조사된 최소문헌란에 기재된 IPC 일본등록실용신안공보 및 일본공개실용신안공보: 조사된 최소문헌란에 기재된 IPC 국제조사에 이용된 전산 데이터베이스(데이터베이스의 명칭 및 검색어(해당하는 경우)) eKOMPASS(특허청 내부 검색시스템) & 키워드: 신경망(neural network), 연산자(operator), 하드웨어 블록(hardware block), 매칭(matching), 탐색 공간(search space), 도출(extracting)		
C. 관련 문헌		
카테고리*	인용문헌명 및 관련 구절(해당하는 경우)의 기재	관련 청구항
X	CN 116306856 A (ZHEJIANG LAB) 2023.06.23 단락 [0039]; 및 청구항 1	1-4,10
Y		5-9
Y	CN 103617146 B (BEIJING QIHOO TECHNOLOGY CO., LTD. 등) 2017.10.13 청구항 1	5-9
A	US 2019-0332441 A1 (HEWLETT PACKARD ENTERPRISE DEVELOPMENT LP) 2019.10.31 단락 [0046]-[0051]; 및 도면 8A-8B	1-10
A	KR 10-2021-0113004 A (삼성전자주식회사 등) 2021.09.15 단락 [0177]-[0188]; 및 도면 13	1-10
A	KR 10-2022-0041926 A (구글 엘엘씨) 2022.04.01 청구항 1-10	1-10
☒ 추가 문헌이 C(계속)에 기재되어 있습니다. ☒ 대응특허에 관한 별지를 참조하십시오.		
* 인용된 문헌의 특별 카테고리: “A” 특별히 관련이 없는 것으로 보이는 일반적인 기술수준을 정의한 문헌 “D” 본 국제출원에서 출원인이 인용한 문헌 “E” 국제출원일보다 빠른 출원일 또는 우선일을 가지나 국제출원일 이후에 공개된 선출원 또는 특허 문헌 “L” 우선권 주장에 의문을 제기하는 문헌 또는 다른 인용문헌의 공개일 또는 다른 특별한 이유(이유를 명시)를 밝히기 위하여 인용된 문헌 “O” 구두 개시, 사용, 전시 또는 기타 수단을 언급하고 있는 문헌 “P” 우선일 이후에 공개되었으나 국제출원일 이전에 공개된 문헌 “T” 국제출원일 또는 우선일 후에 공개된 문헌으로, 출원과 상충하지 않으며 발명의 기초가 되는 원리나 이론을 이해하기 위해 인용된 문헌 “X” 특별한 관련이 있는 문헌. 해당 문헌 하나만으로 청구된 발명의 신규성 또는 진보성이 없는 것으로 본다. “Y” 특별한 관련이 있는 문헌. 해당 문헌이 하나 이상의 다른 문헌과 조합하는 경우로 그 조합이 당업자에게 자명한 경우 청구된 발명은 진보성이 없는 것으로 본다. “&” 동일한 대응특허문헌에 속하는 문헌		
국제조사의 실제 완료일 2024년09월03일(03.09.2024)	국제조사보고서 발송일 2024년09월03일(03.09.2024)	
ISA/KR의 명칭 및 우편주소 대한민국 특허청 (35208) 대전광역시 서구 청사로 189, 4동 (둔산동, 정부대전청사) 팩스 번호 +82-42-481-8578	심사관 양정록 전화번호 +82-42-481-5709	

C. 관련 문헌

카테고리*	인용문헌명 및 관련 구절(해당하는 경우)의 기재	관련 청구항
X	HAYOUNG LIM 등, 'TD-NAAS: Template-Based Differentiable Neural Architecture Accelerator Search', 2023 20th International SoC Design Conference (ISODC), 2023.10.26 섹션 1-4 ※ 위 문헌은 출원인이 신규성 상실의 예외로서 선언한 공지문헌임.	1-10

국제조사보고서에서 인용된 특허문헌	공개일	대응특허문헌	공개일
CN 116306856 A	2023/06/23	CN 116306856 B	2023/09/05
CN 103617146 B	2017/10/13	CN 103617146 A	2014/03/05
US 2019-0332441 A1	2019/10/31	US 10698737 B2	2020/06/30
KR 10-2021-0113004 A	2021/09/15	CN 113361704 A	2021/09/07
		US 2021-0279587 A1	2021/09/09
KR 10-2022-0041926 A	2022/04/01	CN 114258538 A	2022/03/29
		CN 114258538 B	2024/04/12
		EP 3980888 A1	2022/04/13
		JP 2022-544796 A	2022/10/21
		JP 7342247 B2	2023/09/11
		TW 202113596 A	2021/04/01
		TW I764236 B	2022/05/11
		US 2022-0326988 A1	2022/10/13
		WO 2021-034675 A1	2021/02/25