

# 实时低功耗飞行器神经网络



张英<sup>1,2,3</sup> 陶磊岩<sup>4</sup> 曹健<sup>1</sup> 王世会<sup>2,3</sup> 赵茜<sup>2,3</sup> 张兴<sup>1</sup>

1 北京大学软件与微电子学院 北京 100871

2 北京航天自动控制研究所 北京 100854

3 宇航智能控制技术国家级重点实验室 北京 100854

4 北京遥感设备研究所 北京 100854

(zhangying\_@pku.edu.cn)

**摘要** 为了满足飞行器实时飞行过程中对大量异构输入数据的信息处理需求,文中提出了一种神经网络,其包括卷积定点滑动核、池化压缩量化核以及全连接压缩融合核,将飞行器异构传感器多路并行数据作为系统的输入,将辨识结果作为系统的输出。卷积滑动窗口核通过排除冗余数据的滑动窗快速实现数据特征的提取;池化压缩量化核使用压缩量化技术来提高系统的执行效率;全连接压缩融合核经删减量化后压缩融合并输出。该设计满足了飞行器对高可靠性、低功耗的在线智能集成需求。使用所提压缩量化方法,准确率最高可达 98.54%,压缩率为 77.8%,运行速度提升了 40 倍。

**关键词:** 低功耗;神经网络;实时在线;飞行器

**中图法分类号** TP311

## Real-time Low Power Consumption Aircraft Neural Network

ZHANG Ying<sup>1,2,3</sup>, TAO Lei-yan<sup>4</sup>, CAO Jian<sup>1</sup>, WANG Shi-hui<sup>2,3</sup>, ZHAO Qian<sup>2,3</sup> and ZHANG Xing<sup>1</sup>

1 School of Software and Microelectronics, Peking University, Beijing 100871, China

2 Beijing Aerospace Automatic Control Institution, Beijing 100854, China

3 National Key Laboratory of Science and Technology on Aerospace Intelligent Control, Beijing 100854, China

4 Beijing Institute of Remote Sensing Equipment, Beijing 100854, China

**Abstract** In order to meet the information processing requirements of a large amount of heterogeneous input data in the real-time flight of aircraft, this paper proposes a neural network, including convolution core with fixed-point sliding, pooling core with compression quantization and fully connected core with compression fusion. The input of the system is heterogeneous sensor data, and the output of the system is the identification results. Convolution core can extract data features quickly by eliminating redundant data sliding window. Pooling core improves system execution efficiency by using compression quantization technology. The design meets the on-line intelligent integration requirements of high reliability and low power consumption. With the proposed compression quantization method, the peak accuracy is 98.54%, the compression rate is 77.8%, and the running speed increases by 40 times.

**Keywords** Low power consumption, Neural network, Real-time online, Aircraft

## 1 引言

随着人工智能技术的不断发展<sup>[1-3]</sup>,基于人工智能技术的飞行器控制已经成为研究的热点<sup>[4-5]</sup>,其由于具有强大的处理非线性<sup>[6-7]</sup>、自学习以及并行运算的能力<sup>[8-10]</sup>,在复杂非线性系统的控制方面有很大的优势。其中,神经网络技术代表了一种新的方法体系,它以分布式的方式存储信息,利用网络的拓扑结构和权值分布实现非线性映射,并利用全局并行处理实现从输入空间到输出空间的非线性变换<sup>[11]</sup>,从而实现对飞

行器复杂系统的智能控制。

现代目标跟踪<sup>[12]</sup>和无人驾驶<sup>[13]</sup>等领域的智能集成需要丰富的硬件计算资源和较高的计算速度。清华大学的 Yin 等采用  $16 \times 16$  的可重构异构处理单元(Processing Unit, PU)来配置混合神经网络处理器以实现神经网络加速<sup>[14]</sup>。深鉴科技通过修剪和量化去除冗余的网络连接,在不降低识别精度的前提下微调剩余连接权重以减少参数和运算次数<sup>[15]</sup>。Han 等提出了负载均衡感知修剪方法,将长短时记忆网络(Long Short-Term Memory, LSTM)模型修剪为原来的

到稿日期:2019-12-23 返修日期:2020-04-23 本文已加入开放科学计划(OSID),请扫描上方二维码获取补充信息。

基金项目:国家自然科学基金(51877008)

This work was supported by the National Natural Science Foundation of China(51877008).

通信作者:曹健(caojian@ss.pku.edu.cn)

1/10<sup>[16]</sup>。ISSS 2017<sup>[17]</sup>采用多个 PE 进行卷积计算,实现低功耗的高性能加速设计。FPGA 2015<sup>[9]</sup>分析了计算吞吐量和存储器带宽的匹配,提出了 roofline 模型。FPGA 2017<sup>[18]</sup>设计了频域处理的加速器。

现有神经网络系统难以满足飞行器的实时在线需求。飞行器飞行速度快,要求在毫秒级时间内实现辨识输出,在完成大数据吞吐的同时,还要求系统运算的功耗低。现有技术虽然在网络规模较大时使用复杂的架构实现了卷积、池化和全连接层连接结构,但是由于未能对输入优先级最高的主惯导数据进行有效处理,因此无法满足实时在线飞行器 AI 神经网络系统的需求。

为了弥补现有技术的不足,本文提出了一种飞行器低功耗在线神经网络系统。该系统作为一种智能实时处理系统,有效地解决了飞行器实时飞行过程中对大量异构输入数据的高可靠性、低功耗智能集成处理的需求。

2 架构

如图 1 所示,本文提出的飞行器低功耗在线神经网络系统包括卷积定点滑动核、池化压缩量化核以及全连接压缩融合核,共  $i+1$  层,且每个子模块的卷积定点滑动核和池化压缩量化核的结构相同。其中,传感器信号 1 为对应飞行器优

先级最高的主惯导数据,它单独输入一个单元网络层后在第二次卷积时需控制后续 2 至  $i$  层的输入。将经过上述处理的飞行器异构传感器数据作为实时在线飞行器 AI 神经网络系统的输入,将辨识结果作为实时在线飞行器 AI 神经网络系统的输出。飞行器将实时采集的数据输入卷积定点滑动核,通过排除冗余数据的滑动窗快速实现数据特征的提取;随后将卷积滑动窗口核输出的结果发送给池化压缩量化核,使用数据采样和量化技术完成对数据的高效压缩,提高系统的执行效率;最后将池化压缩量化核的输出结果发送给全连接压缩融合核,对经过冗余删减和采样量化后的所有支路数据进行压缩融合处理,最终得到的输出结果可满足飞行器实时飞行过程中对大量异构输入数据的高可靠性、低功耗智能集成处理的需求。本文提出的实时在线飞行器 AI 神经网络系统包括  $i+1$  层神经网络单元,每层神经网络单元均包括若干交替进行的卷积定点滑动核和池化压缩量化核,卷积定点滑动核对输入的数据进行卷积处理,池化压缩量化核进行降低数据率的操作以及特征提取。外部输入的第一路传感器信号对应飞行器优先级最高的主惯导数据,被单独输入一层神经网络单元中;将经过该层神经网络单元中第一个卷积定点滑动核处理和第一个池化压缩量化核处理后的输出,作为控制所有  $i+1$  层神经网络单元中第二次卷积定点滑动核的输入。

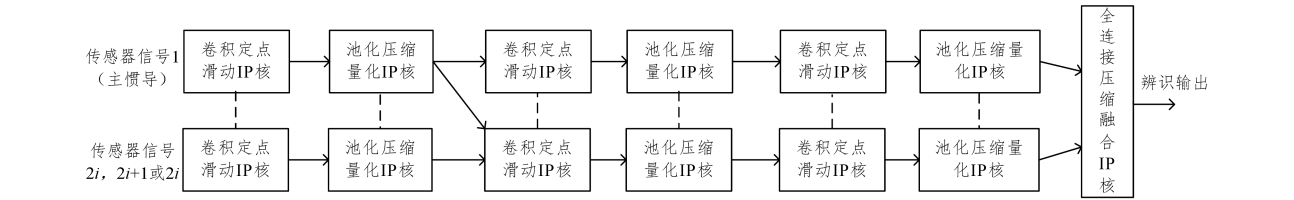


图 1 飞行器实时 AI 神经网络系统框图

Fig.1 Graph of aircraft real-time AI neural network system

所有  $i+1$  层神经网络单元的最终处理结果均被输入到全连接压缩融合核进行处理,以实现目标识别,而所述全连接压缩融合核的输出即为 AI 神经网络系统的输出。飞行器飞行时最主要的传感器信息是主惯导数据。传统的方法是将主

惯导数据与其他辅助的传感器信息进行卡尔曼滤波、预估计等操作以完成信息融合,而本文方法是基于输入数据驱动的。本方法利用第一路主惯导数据对其余所有路信息进行修正,大幅提升了输出结果的准确性。

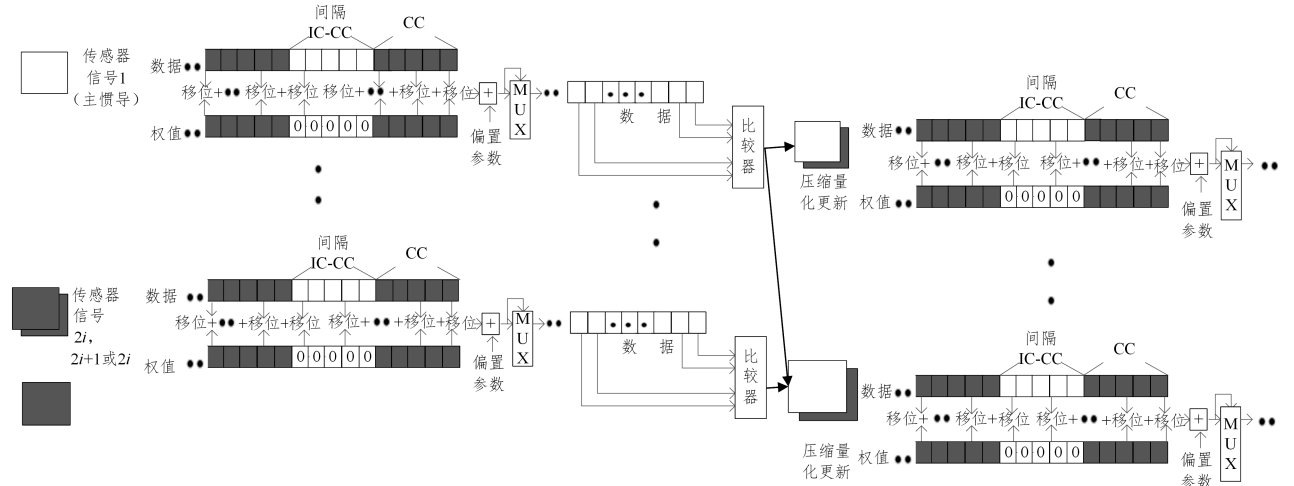


图 2 飞行器实时 AI 神经网络数据流图

Fig.2 Graph of aircraft real-time AI neural network dataflow

因为输入数据经过同样的卷积定点滑动窗口核和池化压 缩量化核,以及基于相同的数字变换,所以可以确保在第二次

卷积定点滑动核第一层神经网络处理过的主惯导数据与本层神经网络处理过的数据的输入信息数据具有相同的维数,确保了核的兼容性。

如图 2 所示,神经网络单元的层数具体为:1)当只有 1 个传感器输入时,层数为 1;2)当有 2 个及上传感器输入时,层数为  $i+1$ , $i$  为传感器数量的  $1/2$ ,并向下取整。

将外部输入的第一路传感器信号输入到第一层神经网络单元中;将第二路和第三路传感器信号均送入到第二层神经网络单元中;将第四路和第五路传感器信号均送入到第三层神经网络单元中。以此类推,当传感器数量为奇数时,第  $2i$  路和第  $2i+1$  路传感器信号均送入第  $i+1$  层神经网络单元中;当传感器数量为偶数时,第  $2i$  路传感器信号送入第  $i+1$  层神经网络单元中。

卷积定点滑动核对数据进行卷积运算。卷积定点滑动核有 2 个缓存器,1 个为权值缓存器,用于存储权值;1 个为数据缓存器,用于存储传感器输入处理后的数据。

将权值缓存器存储量化为 2 的指数次幂的固定权值,数据缓存器按时钟周期移位缓存输入数据;权值缓存器的大小为  $CC \times CR$ ,数据缓存器的大小为  $(CR-1) \times IC + CC$ 。其中, $CC$  为卷积核列数, $CR$  为卷积核行数, $IC$  为输入特征数列数。

权值缓存器对应一组移位乘操作。按时钟完成数据缓存器移动后,将数据缓存器和权值缓存器的对应单元进行移位乘操作。

卷积定点滑动核的加法器用于累加移位乘操作结果和预设偏移参数。加法器个数为移位乘操作个数加 1,加法器包括与移位乘操作一一对应的部分,额外的 1 个加法器用于累加偏移参数。

卷积定点滑动核的多路选择器用于模拟激活函数。多路选择器的输入为所有加法器累加计算后的结果,输出为卷积定点滑动核的结果。

池化压缩量化核对输入数据进行降低数据率处理以及特征提取。首先池化压缩量化核的缓存器读入卷积定点滑动核,以处理输出的特征图数据,池化压缩量化核的缓存器为移位缓存器。然后,将缓存器中存储的数据输入到池化压缩量化核的比较器中进行数值比较,比较器的输出为数据压缩后的结果,该结果即为池化压缩量化核的处理结果,比较器的处理在一个时钟周期内完成。

当输入特征图数据持续流入缓存器时,实现了池化压缩量化核的滑动平移池化,池化压缩量化核的第一个输出结果为有效数据,后  $CC-1$  个结果为无效数据,第  $CC$  个输出结果为有效数据,后  $2CC-1$  个结果为无效数据,以此类推。然后,将最终得到的所有有效数据进行 2 的指数次幂定点缩减量化更新。定点量化是将大型矩阵进行乘法操作,转换为同或门进行真实的位操作。

举例说明 1:若量化为  $+1$  或  $-1$ ,则按如下可快速执行运行运算的简化公式。

$$X = \begin{cases} +1, & X \text{ 的概率分布为 } \sigma(x) \\ -1, & X \text{ 的概率分布为 } 1-\sigma(x) \end{cases}$$

其中, $\sigma(x)$ 为预设的分布函数。

举例说明 2:若量化为  $+1,0$  或  $-1$ ,则按如下可快速执行运算的简化公式。

$$X = \begin{cases} +1, & X > \emptyset \\ 0, & |X| \leq \emptyset \\ -1, & X < -\emptyset \end{cases}$$

其中, $\emptyset$  为预设的尺度因子。

全连接压缩融合核对提取结果进行处理,把提取到的特征综合起来,进行参数和输入特征图数据的运算,并将运算结果通过端口输出,结果为辨识输出。全连接压缩融合核包括缓存器,缓存器又包括存储预设参数的全连接权值缓存器、存储输入数据的全连接数据缓存器以及存储预设为定点量化的梯度的偏移量缓存器。全连接压缩融合核对提取结果进行处理的具体步骤如下。

首先对全连接权值缓存器中预存的  $n \times n$  个预设参数数据进行预压缩,量化为  $n$  类,并用每类量化值代表每类的权值,得到量化后的具有  $n \times n$  个索引的权值矩阵。

然后,将偏移量缓存器同样量化为  $n$  类,对每类进行梯度求和得到每类偏置,将每类偏置与量化值一起更新得到新的权值,并存储到全连接权值缓存器中。

最后,将全连接权值缓存器存储量化为 2 的指数次幂的固定权值,全连接数据缓存器按时钟周期移位缓存输入数据。其中,全连接权值缓存器对应一组移位乘操作,每组移位乘操作的个数与全连接权值缓存器的深度相同。按时钟完成数据缓存器移动后,将全连接数据缓存器和全连接权值缓存器的对应单元进行移位乘操作。

全连接压缩融合核的加法器用于累加移位乘操作结果和预设梯度的偏移量缓存器参数。加法器个数为移位乘操作个数加 1,加法器包括与移位乘操作一一对应的部分,额外的 1 个加法器用于累加偏移参数。

### 3 仿真实验

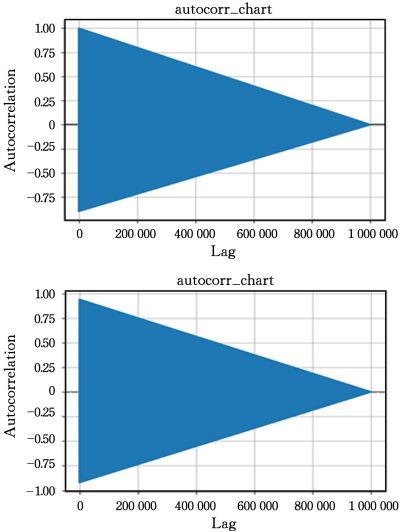
进行压缩量化后,本文使用下述方法进行数据相关度量,以检验序列是否平稳。若时间序列有有穷的二阶矩,且  $X_t$  满足如下两个条件,则称该序列为时间平稳序列。

$$\begin{aligned} (1) & \mu_t = EX_t = c, \\ (2) & \gamma(t,s) = E(X_t - c)(X_s - c) = \gamma(t-s,0). \end{aligned}$$

由分析结果可知,数据序列的自相关系数能够衰减至零,并且之后只在零附近进行小范围波动,因此可以认为数据具有平稳性。任取两组数据的自相关系数图进行展示,如图 3 所示。

本文进行了神经网络系统的训练和预测实验。实验按照上文的结构构建了神经网络并使用了 100 000 组数据进行测试,1 批 batchsize 包括 10 组训练数据,每 6 000 批 batchsize 为 1 个迭代周期(epoch),实验共进行了 22 个 epoch 的训练。每 1 个 batchsize 训练完成后,先求 batchsize 内数据的预测(inferences)和标签(labels)的损失函数,然后经过换算得到准确率。可以看出,在前 6 000 批数据内准确率的增长速度总体较快,在第 6 000 批之后准确率的增长速度总体趋于平缓,但仍保持波动增长的状态。当训练批次的值为 122 899 时,准确率达到峰值,为 98.54%。虽然压缩后的网络的准确率在后

续训练中趋于平缓且略微呈下降态势,但仍明显优于同等条件下传统方法的准确率,打破了准确率无法达到 90% 的壁垒。本次实验在 de2i-150, Intel® Xeon® CPU e3-1230 V2, GPU NVIDIA Quadro 4000 的硬件资源上运行。本文模型的压缩率可以达到 77.8%, 运行速度增加了 40 倍。



注:Lag 表示数据, Autocorrelation 表示自相关系数

图3 两组数据的自相关系数图

Fig. 3 Graph of autocorrelation coefficients of two sets of data

本文所提系统在 FPGA DE2i-150, Intel® Xeon® CPU E3-1230 V2 和 GPU NVIDIA Quadro 4000 的硬件资源上运行。实际运行模型可实现 77.8% 的压缩率, 识别性能提升了至少 20 倍, 速度最多提高了 40 倍。以 6 路传感器信号输入为例, 本文飞行器实时在线 AI 神经网络系统与其他文献提出的系统相比具有显著优势, 具体内容如表 1 所列。延迟时间指标如表 2 所列, 本文的 FPGA 实现方案较文献[22]和本文的 CPU、GPU 方案均有较大提升。

表1 性能及效率对比

Table 1 Performance and efficiency comparison

|        | 性能/GMACs | 效率/GOPs |
|--------|----------|---------|
| 文献[19] | 3.37     | 6.74    |
| 文献[20] | 2.6      | 5.25    |
| 文献[21] | 3.5      | 7.0     |
| 本文方案   | 9.86     | 19.72   |

表2 延迟时间对比

Table 2 Time delay comparison

|         | 所有用时                 |
|---------|----------------------|
| 文献[22]  | $7.3 \times 10^{-4}$ |
| 本文 CPU  | $2.1 \times 10^{-1}$ |
| 本文 GPU  | $8.3 \times 10^{-2}$ |
| 本文 FPGA | $7.6 \times 10^{-5}$ |

**结束语** 本文所提系统与现有技术相比具有如下优点。

(1)系统构建优点。本系统解决了飞行器在实时飞行过程中对大量异构输入数据进行并行处理时存在的问题。通过本系统中的卷积核, 将多维不同传感器产生的异构数据分别与图像及参数存储模块中的压缩量化数据对应。经过并行  $2i$  或  $2i+1$  个目标传感器的处理, 训练后的低功耗实时系统可

准确地完成信息识别和定位。其中, 传感器信号 1 为飞行器中优先级最高的主导数据, 它单独输入一个单元网络层, 在第二次卷积时需控制后续 1 至  $i+1$  层的输入。

(2)压缩量化优点。本系统解决了庞大复杂的神经网络在实时飞行过程中存在的问题。其使用了数据的定点压缩和压缩量化技术, 降低了模型参数及计算复杂度, 提高了系统的稳定性, 同时使用了共享权值技术, 有效减少了参数的数量。此外, 该系统还降低了数据存储要求并简化了计算过程, 其将相关数据量化为 2 的指数次幂, 乘法运算通过数据移位运算来解决。因为存储器和乘法器决定了神经网络的优先运算加速能力, 所以本文方法将这些核构建为一个完整的低存储无乘法器实时在线 AI 网络。在不损失精度的前提下, 仅通过使用与、或、非、异或门, 本文方法的乘法等效运算速度就提升了 10 倍, 除法等效运算速度提升了 40 倍。举例来说, 使用本文方法的移位乘操作, 仅需消耗 1 个单位, 但若使用其对应的 32 位乘法则将损耗 200 个单位。

(3)本文方法设计了适用于不同层的卷积定点滑动核、池化压缩量化核以及全连接压缩融合核, 避免了多余的输入数据的操作。该核在相同的神经网络结构下, 运用了压缩处理后得到的是非归零压缩编码, 有效降低了功耗, 且其计算能力和延迟时间与其他方法相比有所提升, 实现了不同输入数据的并行处理。

参 考 文 献

[1] LU L Q, LIANG Y, XIAO Q C, et al. Evaluating Fast Algorithms for Convolutional Neural Networks on FPGAs [C] // 2017 IEEE 25th Annual International Symposium on Field-Programmable Custom Computing Machines (FCCM). 2017:101-108.

[2] GIRSHICK R, DONAHUE J, DARREKK T, et al. Rich Feature Hierarchies for Accurate Object Detection and Semantic Segmentation [C] // 2014 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). Columbus, 2014:580-587.

[3] GIRSHICK R. Fast R-CNN [C] // IEEE International Conference on Computer Vision (ICCV). Santiago, Chile: IEEE, 2015: 1440-1448.

[4] MC CULLOCH W S, PITTS W H. A Logical Calculus of the Ideas Immanent in Nervous Activity [J]. Bulletin of Mathematical Biophysics, 1943, 5(5):115-133.

[5] MINSKY M, PAPERT S. Perceptrons: An Introduction to Computational Geometry [M]. USA, Massachusetts: The MIT Press, 1987:5-308.

[6] RUMELHART D E, HINTON G E, WILLIAMS R J. Learning Representations by Back-propagating Errors [J]. Nature 1998, 323(6088):533-536.

[7] KRIZHEVSKY A, SUTSKEVER I, HINTON G E. ImageNet Classification with Deep Convolutional Neural Networks [C] // International Conference on Neural Information Processing System, Kyoto, Japan. VLSI Secretariat Japan and Asia, 2012:1097-1105.

[8] SCHERER D, SCHULZ H, BEHNKE S. Accelerating Large-Scale Convolutional Neural Networks with Parallel Graphics



Multrocessors[C]//International Conference on Artificial Neural Networks(ICANN). 2010(6354):82-91.

[9] ZHANG C,LI P,SUN G Y,et al. Optimizing FPGA-based Accelerator Design for Deep Convolutional Neural Networks[C]//Proceedings of the 2015 ACM/SIGDA International Symposium on Field-Programmable Gate Arrays. New York, USA: ACM, 2015:161-170.

[10] LECUN Y,BOTTOU L,BENGIO Y,et al. Gradient-based Learning Applied to Document Recognition[J]. Proceedings of the IEEE,1998,86(11):2278-2324.

[11] SIMONYAN K,ZISSERMAN A. Very Deep Convolutional Networks for Large-scale Image Recognition[C]//Computer Vision and Pattern Recognition(CVPR). Columbus OH USA, IEEE Computer Society,2014(v1):1409-1556.

[12] REN S Q,HE K M,GIRSHICK R,et al. Faster R-CNN: Towards Real-Time Object Detection with Region Proposal Networks[J]. IEEE Transactions on Pattern Analysis and Machine Intelligence,2017(39):1137-1149.

[13] KIM H S,HONG S W,SON H R,et al. High Speed Road Boundary Detection on the Images for Autonomous Vehicle with the Multi-layer[C]//IEEE International Symposium on Circuits and Systems(ISCAS). Bangkok, Thailand, IEEE 2003:769-772.

[14] YIN S,OUYANG P,TANG S,et al. A 1.06-to-5.09 TOPS/W Reconfigurable Hybrid-Neural-Network Processor for Deep Learning Applications[C]//2017 Symposia on VLSI Technology and Circuits (VLSI). Kyoto, Japan. VLSI Secretariat Japan and Asia,2017:26-27.

[15] HAN S,POOL J,TRAN J,et al. 2015. Learning both Weights and Connections for Efficient Neural Networks[C]//Neural Information Processing System. Montreal, Quebec, Canada: MIT Press,2015:1506-1526.

[16] HAN S,KANG J,MAO H,et al. ESE:Efficient Speech Recognition Engine with Sparse LSTM on FPGA[C]//Proceedings of the 2017 ACM/SIGDA International Symposium on Field-Programmable Gate Arrays. Oulu, Finland: IEEE/ACM, 2017: 75-84.

[17] GONG L,WANG C,LI X,et al. Work-in-Progress:A Power-Efficient and High Performance FPGA Accelerator for Convolutional Neural Networks[C]//Proceedings of the 12th IEEE/ACM/IF International Conference on Hardware/Software Code-sign and System Synthesis Companion. Oulu, Finland: IEEE/ACM,2017:1-6.

[18] ZHANG C,PRASANNA V. Frequency Domain Acceleration of Convolutional Neural Networks on CPU-FPGA Shared Memory System[C]//Proceedings of the 2017 ACM/SIGDA International Symposium on Field-Programmable Gate Arrays. Oulu, Finland:IEEE/ACM,2017:35-44.

[19] SANKARADAS M,JAKKULA V,CADAMBI S,et al. A Massively Parallel Coprocessor for Convolutional Neural Networks [C]//2009 20th IEEE International Conference on Application-specific Systems, Architectures and Processors (ASAP). Montreal:IEEE,2009:53-60.

[20] FARABET C,POULET C,HAN J Y,et al. CNP:An Fpga-based Processor for Convolutional Networks[C]//Field Programmable Logic and Applications(FPL). Prague:IEEE,2009: 32-37.

[21] CADAMBI S,MAJUMDAR A,BECCHI M,et al. A Programmable Parallel Accelerator for Learning and Classification[C]//Proceedings of the 19th International Conference on Parallel Architectures and Compilation Techniques (PACT). Vienna, Austria:IEEE,2010:273-284.

[22] FANG R,LIU J H,XUE Z H,et al. FPGA-based design for convolution neural network[J]. Computer Engineering and Applications,2015,51(8):32-36.



**ZHANG Ying**, born in 1982, Ph.D, senior engineer. Her main research interest is intelligent control.



**CAO Jian**, born in 1980, Ph.D, associate professor, is a member of China Computer Federation. His main research interests include edge computing, intelligent hardware and system design.