

目 录

摘要.....	I
ABSTRACT	III
第一章 绪论.....	1
1.1 研究背景与意义	1
1.2 国内外研究现状	3
1.2.1 人工选择阶段	3
1.2.2 机器学习阶段	4
1.2.3 深度学习阶段	5
1.3 论文主要工作	6
1.4 论文结构安排	7
1.5 本章小结	8
第二章 基本理论.....	9
2.1 脉冲星候选体产生步骤	9
2.2 CRAFTS 巡天.....	13
2.3 相关深度学习理论和方法介绍	13
2.3.1 深度学习模型	13
2.3.2 ImageNet 数据集和预训练模型	14
2.3.3 PyTorch 介绍和 LogME 介绍	15
2.3.4 迁移学习	17
2.3.5 对比学习	18
2.3.6 自监督学习	18
2.4 本章小结	19
第三章 利用预训练-微调进行脉冲星候选体分类	20
3.1 数据集介绍	20
3.2 数据预处理	21

3.3 模型选择	22
3.4 模型微调	23
3.5 本章小结	26
第四章 利用自监督对比学习进行脉冲星候选体分类	27
4.1 数据集介绍	27
4.2 数据预处理	27
4.3 自监督对比方案设计	28
4.3.1 样本增强方式	29
4.3.2 整体模型架构和基础模型选择	30
4.3.3 损失函数	31
4.4 模型训练和微调	32
4.5 本章小结	34
第五章 实验结果	35
5.1 评价指标	35
5.2 结果分析	36
5.2.1 模型选择 LogME 结果分析	36
5.2.2 HTRU1 微调结果	37
5.2.3 自监督对比学习训练结果	39
5.2.4 自监督预训练模型迁移结果	40
5.3 本章小结	41
第六章 总结与展望	42
6.1 本文总结	42
6.2 研究展望	43
参考文献	45
附录	50
致谢	62
攻读硕士学位期间主要研究成果	63

贵州师范大学学位论文原创性声明	64
贵州师范大学学位论文使用授权书	64

摘要

候选体分类是脉冲星搜索流程中的重要一步，它承上启下，既是对搜索流程的阶段性的总结，又是对脉冲星性质研究的开始。随着中国“天眼”射电望远镜的正式投入使用，以及银道面脉冲星巡天、漂移扫描多科学目标同时巡天等项目的开展，脉冲星分类问题面临海量数据，单条赤纬三个小时的扫描数据即可产生百万级候选体，对分类工作提出了巨大的挑战。早期的人工筛选早已无法满足这一步骤的需求，尽管机器学习方法和监督深度学习方法在脉冲星候选体分类问题中表现良好，但是机器学习方法非常依赖专家对脉冲星样本特征的选取，而监督训练非常依赖样本的数量和质量，使得深度学习模型难以训练。为了更有效利用海量数据和各种优秀的视觉深度学习模型，受自监督学习和迁移学习启发，本文提出一种基于自监督对比学习和迁移学习的两阶段脉冲星候选体分类方法。

为了将预训练模型应用于 FAST 数据，本方法以自监督对比学习的方式对模型进行预训练得到预训练模型，再使用有标签数据集对模型进行微调。首先，利用 LogME 方法评估各个预训练模型的可迁移性，采用线性评估和微调评估两种方法，对预训练模型在 HTRU1 数据集上的表现进行了综合评估，衡量该方法在脉冲星数据上的有效性。然后，通过 LogME 评估得到在 FAST 数据集上可迁移性最好的预训练模型结构，利用 CRAFTS 巡天项目收集的大量无标签数据进行自监督对比学习预训练，使模型通过无标签脉冲星数据获得与识别脉冲星特征相关的能力。最后，将自监督对比学习得到的预训练模型中的

特征提取模块直接在 FAST 训练集上进行微调，并在测试集上进行效果测试。

结果表明，使用这种方式选择模型过程简单，训练过程无须进行大量数据标注，得到的模型分类效果好，给其它望远镜或其它巡天项目的周期性脉冲星采用计算机视觉领域的 SOTA 模型进行迁移学习提供了一定的参考。

关键词：脉冲星候选体分类；迁移学习；自监督学习；对比学习；可迁移性评估

ABSTRACT

Candidate classification is a crucial step in the pulsar search process. It bridges the gap between different stages, serving both as a summary of the search process and the beginning of pulsar property research. With the official commissioning of the Chinese FAST radio telescope and the initiation of projects such as the Galactic Plane Pulsar Snapshot survey and the Commensal Radio Astronomy FAST Survey, pulsar classification faces the challenge of massive amounts of data. A single three-hour scan at a given declination can generate millions of candidates, posing a significant challenge to the classification task. Early manual screening methods can no longer meet the demands of this step. Although machine learning and supervised deep learning methods perform well in pulsar candidate classification, machine learning heavily relies on experts selecting features of pulsar samples, and supervised training is highly dependent on the quantity and quality of samples, making deep learning models difficult to train. To more effectively utilize vast amounts of data and various excellent visual deep learning models, this paper proposes a two-stage pulsar candidate classification method inspired by self-supervised learning and transfer learning.

To apply pre-trained models to FAST data, this method first pre-trains the model using self-supervised contrastive learning to obtain a pre-trained model, and then fine-tunes the model using a labeled dataset. Initially, the

LogME method is employed to evaluate the transferability of various pre-trained models. Both linear evaluation and fine-tuning evaluation methods are used to comprehensively assess the performance of pre-trained models on the HTRU1 dataset, measuring the effectiveness of this method on pulsar data. Subsequently, the pre-trained model structure with the best transferability on the FAST dataset, as determined by LogME evaluation, is selected. A large amount of unlabeled data collected by the CRAFTS project is then used for self-supervised contrastive learning pre-training, enabling the model to acquire the ability to identify pulsar features from unlabeled pulsar data. Finally, the feature extraction module of the pre-trained model obtained through self-supervised contrastive learning is fine-tuned directly on the FAST training set, and its effectiveness is tested on the test set.

The results indicate that using this method for model selection is straightforward, the training process does not require extensive data annotation, and the resulting model achieves good classification performance. This provides a reference for applying state-of-the-art computer vision models for transfer learning in periodic pulsar detection in other telescopes or survey projects.

Keywords: pulsar candidate classification, transfer learning, self-supervised learning, contrastive learning, transferability evaluation

第一章 绪论

1.1 研究背景与意义

1934 年,天文学家 Walter Baade 和 Fritz Zwicky 首次提出了中子星^[1]的概念,他们认为恒星演化的终点可能会形成一种新类型的星体,它主要由中子组成,具有极小的半径和极高的密度。这一预言性质的概念在当时几乎是不可能通过仪器来直接观测的,因为从它的性质来看,中子星几乎不会向外发射光线。并且,在 Karl Jansky 初次记录到源自银河系中心的无线电波信号^[2]这一标志着射电天文学诞生的事件到这一猜测提出的这两年时间内,人们也并未通过射电方式探测到这类独特天体可能释放的特定信号。

约 30 年后的 1967 年,英国剑桥新建造的射电望远镜开始投入使用,用于观测星际介质不规则结构引起的闪烁,从而研究紧凑射电源的角结构。在这个背景下,研究生 Jocelyn Bell 对望远镜观测的数据进行处理时偶然发现了一系列奇怪的脉冲^[3]。这些脉冲是重复的,且具有 1.337 秒的稳定周期。最初,科学家们以为这些脉冲可能来自人造深空探测器、雷达或者地球反射的信号。然而,长达几个月的持续观测表明,这些脉冲的周期性和定期出现的性质排除了它是来自地球的干扰的可能。科学家们意识到这些信号应该是来自宇宙中某个天体,经过进一步的观测和研究,科学家们认为这些奇特的脉冲信号确实是由旋转的中子星产生的^[4-5]。这个发现证实了中子星的存在,并引发了人们对于脉冲星这类特殊天体的广泛探索。并且,脉冲星以其极高的密度、快速旋转和强大的磁场等物理特性,迅速成为了天文学和物理学研究的热门对象之一。

脉冲星是一种高速旋转的致密天体,其特征在于强大的磁场和特殊的电磁辐射特性^[6]。这些脉冲星所发出的电磁辐射可以被工作在不同波段的射电望远镜接收,方便研究人员从非可见光的角度进行研究。脉冲星的研究不仅局限于天文学领域,还涉及到物理学、宇宙学等多个学科领域,具有广泛的研究意义和应用价值。首先,脉冲星的发现和相关研究曾两次获得诺贝尔物理学奖,这充分表明了其在科学界的重要性和突破性。脉冲星的研究涉及到宇宙中极端条件下的物理现象和过程,对于理解宇宙的演化、结构和性质具有重要意义。其次,一些脉冲星位于超新星爆发的遗迹内部,这为研究超新星爆发过程和恒星演化机制提供了重

要的实验条件。通过观测和分析这些脉冲星，可以深入探讨恒星结构、恒星死亡和星际物质演化等课题，对于理解宇宙中恒星的演化和宇宙结构的形成具有重要意义。同时，脉冲星具有极高密度和强磁场的特性，因此研究其内部构造和电磁辐射机制可以帮助人们更好地理解物质在极端条件下的性质，不仅有助于推动核物理和高能物理的发展，还可以拓展人们对物质状态和行为的认识，对于推动科学技术的进步具有重要意义。另外，由于脉冲信号从发射源到达地球经过了宇宙空间，因此通过研究到达脉冲的性质可以了解信号途经的宇宙介质电子分布情况，有助于深入了解宇宙空间的物质分布和演化过程，对于理解宇宙结构和宇宙中的物质交互作用具有重要意义。此外，脉冲星的周期稳定性和极低的周期变化率使其成为时间标准的优良选择，可以用于测量宇宙学中的时间和距离，从而研究宇宙的结构、演化和加速膨胀等重要问题。脉冲星的应用还可以延伸到 X 射线脉冲星导航技术领域，该技术由于其抗干扰能力强、自主性强、导航误差小等优势^[7]，被视为新兴的航天器导航技术，对于太空探索和导航系统的发展具有重要意义。

到目前为止，已经发现了 3500 多颗脉冲星^[8]，其中大部分是使用射电望远镜开展的巡天项目观测得到的。被誉为“中国天眼”的 500 米口径球面射电望远镜(Five-hundred-meter Aperture Spherical radio Telescope, FAST)于 2016 年 9 月 25 日落成启用，是目前全球最大且最灵敏的射电望远镜^[9]。2020 年 1 月 11 日通过国家验收正式开放运行，到现在为止，已经开放运行 4 年多。基于超高灵敏度的明显优势，巡天时的实时的数据流量可达每秒 6GB，已成为中低频射电天文领域的观天利器。近几年来，FAST 已经取得非常不错的成果。到 2023 年 9 月，仅通过 FAST 银道面脉冲星(the Galactic Plane Pulsar Snapshot , GPPS)巡天¹就已经累计观测到 2628 颗脉冲星，含内银道面 1798 颗，外银道面 830 颗。到 2023 年 11 月，这一巡天项目累计新发现的脉冲星为 600 多颗^[10]。

对寻找脉冲星这个任务来说，由于其周期最低可达毫秒级别，为了采样这些数据，采样率通常需设置得很高，这也使得望远镜运行时产生的数据非常多。脉冲星候选体分类作为整个脉冲星搜索流程中的重要一环，为了从浩瀚的宇宙中找到通过星际介质传播时发生色散的周期性宽带脉冲信号，会为采集到的数据中满

1 <http://zmtt.bao.ac.cn/GPPS/>

足寻找标准的信号生成大量的候选体诊断图和相关的统计数据，以备后续分析。每个候选体都必须经过筛选和检查，以确定它是否为真的脉冲星信号。对于可能来自脉冲星的信号，会将其筛选出来以进行下一步验证，即需要申请望远镜时间对信号来源进行观测确认或利用其它望远镜进行交叉验证。如果分类效率太低，则采集的大量数据的积压将给存储资源和处理能力带来压力，如果分类效果太差，则会遗漏信号，导致错过重要的发现使得前功尽弃或者将干扰识别为脉冲星加重后续研究人员的负担^[11]。因此，对于脉冲星候选体分类方法的研究至关重要。脉冲星候选体分类方法可以根据其不同发展阶段主流方式分为人工选择阶段、人工特征选择辅助机器学习阶段、深度学习阶段。

1.2 国内外研究现状

1.2.1 人工选择阶段

人工选择指单纯依靠专家人工或仅依靠专家手动设置特征过滤条件的方式对脉冲星疑似信号进行筛选。在脉冲星发现伊始，射电望远镜数量较少，且受限于射电技术、计算机处理能力，在这个时期，脉冲星候选体数据量有限，专家们通常通过目视的方式检查每个疑似候选体的诊断图来进行筛选。如果诊断图显示出类似真实脉冲星具有的特征，则该候选体被保留进行进一步分析，否则被舍弃。比如在 1977 年进行的第二次 Molonglo 巡天中，总共发现了 2500 个疑似对象，其中大部分是重复检测到的强脉冲星信号和容易辨认的干扰^[12]。为了减少使用帕克斯(Parkes)望远镜进行再次观测确认的数量，专家们根据信噪比和脉冲轮廓进行人工评分，选出了 500 颗候选体。这些候选体中有 40%在多次观测中被相邻的扫描任务获取到，也因此获得了较高的评分。最终，这 500 颗候选体中共检测到了 224 颗脉冲星，且有 155 颗是新发现的，一举超过了之前 1967 年至 1976 年间总共发现的 149 颗。

随着技术进步和对短周期脉冲星关注的增加，望远镜的采样间隔下降到 2 毫秒以提高发现短周期脉冲星的灵敏度，这使得候选体数量大量增加。为了应对这种变化，一些巡天项目开始采用新的方法。例如，Joerell Survey B 使用计算机基于脉冲轮廓和信噪比对候选体进行过滤^[13]，得到一个合适范围内的疑似数目进行重新观测，以减少工作量；而 Arecibo Phase II 巡天直接将信噪比切割点设为 8σ ，将待分类的候选体数据个数缩减到 5405 个^[14]。然而，这些基于经验的方法

仅作用于搜索流程的最后阶段，效果十分有限。到了 20 世纪末 21 世纪初，随着候选体数量的增加，手动选择变得越来越不可行。为了进一步减少人工检查的候选体数量，一些新的图形工具如 runview、REAPER、JREAPER 和 PEACE 等应时而生。这些工具通过计算疑似对象的参数，根据专家的经验选择多个参数进行判断，并由程序自动打分并排名，从而去除大量噪声文件，提高了筛选的效率和准确性^[15-18]。

1.2.2 机器学习阶段

随着射电天文学领域的不断发展和技术的进步，脉冲星数据的数量和质量不断提高，这为研究者们开展脉冲星分类和识别提出了更多的要求和更高的挑战。在这个背景下，机器学习逐渐开始成为处理脉冲星数据的重要工具，因为它能够利用专家设计好的特征自动完成各个特征对应权重的学习，从而实现了对脉冲星进行分类的目标。以机器学习为基础的脉冲星候选体分类研究也取得了一系列重要进展。Eatough 等人的研究是其中的一次重要尝试，他们采用了多层感知机 (Multilayer Perceptron, MLP) 作为预测模型，利用 12 个受 JREAPER 评分系统启发得到的数值特征，成功发现了一颗新的脉冲星^[19]。此后，Bates 等人进一步扩充了特征集合，使得特征数目达到了 22 个，并利用基于多层感知机的神经网络分类器对 HTRU 巡天数据进行分类，取得了更好的性能^[20]。这些研究展示了机器学习在脉冲星候选体分类任务中的潜力和优势。然而，机器学习方法在脉冲星分类中也面临一些挑战和限制。例如，传统的特征设计往往受到研究者对数据特征认知的限制，难以充分挖掘数据的潜在信息。为了解决这一问题，研究者们提出了一系列新的特征选择和分类方法。Morello 等人^[21]通过引入计算机领域发展得到的最新成果，提出了 6 个经验特征并训练 SPINN 分类器，取得了较好的分类效果。Lyon 等人^[22]则设计了一种基于高斯海林格距离的快速决策树算法 (Gaussian Hellinger Very Fast Decision Tree, GH-VFDT)，并提出了 8 个新的特征，不过 Tan 等人^[23]指出 GH-VFDT 对于特定类型的脉冲星不敏感，如对具有宽脉冲星积分轮廓的脉冲星识别能力不够，因此又提出了一些新的特征并基于这些特征构建了集成分类器。由此可见，专家们针对不同巡天项目、不同类型的脉冲星，设计了各种特征和分类方法，以提高分类的准确性和鲁棒性，但也暴露出这一过程高度依赖于专家对脉冲星的认知。总的来说，机器学习在脉冲星分类中扮演着

重要的角色，通过不断优化特征选择和分类模型，可以实现对脉冲星的准确分类和识别。然而，未来的研究还需要更深入地理解脉冲星数据的特征和模式，进一步提升分类的精度和可靠性。

1.2.3 深度学习阶段

随着“数据上网”的发展，使得领域行业数据持续增加，深度学习也逐渐成为许多领域中挖掘数据特征的工具之一。相对于传统的机器学习算法，深度学习模型最大的优势之一在于其具有端到端的学习能力。端到端是指无需人工设计特征、进行特征提取、转换或其它中间步骤，只需为输入端提供合适的输入后，即可直接得到结果。传统机器学习算法通常需要手工设计特征，然后再将这些特征输入到模型中进行训练，这个过程中需要大量的人力和丰富的经验。深度学习模型则可以直接接受原始数据作为输入，在学习的过程中自动进行特征提取和权重分配，最终输出分类结果，无需人工干预，避免了人为因素对结果的影响，提高了模型的适用性。

Zhu 等人在研究中设计了一种深度神经网络图像模式识别方式 PICS(Pulsar Image-based Classification System)，该系统集成了支持向量机(Support Vector Machine, SVM)、人工神经网络(Artificial Neural Network, ANN)、卷积神经网络(Convolutional Neural Network, CNN)和逻辑回归(Logistic Regression, LR)等多种技术，忽略以往机器学习方法研究中使用的人工设计的多种数值特征，直接将搜索流程中获得的四张子图直接输入到模型中，让模型自动学习提取特征并输出候选体评分^[24]。之后，Wang 等人对 PICS 进行了改进提出了 PICS-ResNet，相较于原来的 PICS，他引入了以残差连接为特点的 ResNet 网络作为二维特征提取网络取代原先的 CNN 网络，通过这一改进使 PICS 获得了更好的分类性能，同时利用 TensorFlow 框架使用 GPU 加速计算使得分类过程更加高效^[25]。然而，这些模型由于需要输入四张图像，其模型相对较为复杂，这可能会增加计算资源的需求和训练的时间成本。为了简化模型结构，Guo 等人提出仅使用两张子图（时间相位图和频率-相位图）即可满足脉冲星分类的需求，因为这两张图已经包含了丰富的信息用于脉冲星分类^[26]。Liu 和 Qian 等人在此基础上进行了研究，并取得了良好的效果^[27-28]。

相较于机器学习方法，首先，深度学习方法通常需要大量的数据进行训练，

特别是在模型复杂度较高的情况下,需要更多的数据来保证模型的泛化能力和分类稳定性。此外,模型的复杂度也会导致训练过程的计算资源和时间成本增加,需要权衡模型的复杂度和训练效率之间的关系。随着预训练模型和迁移学习技术的不断发展,越来越多的研究者开始将这些技术应用到脉冲星候选体分类中。Agarwal 等人^[29]使用迁移学习方法,基于 ImageNet 数据集的预训练图像分类模型,训练出的网络可以用于快速发现源的分类。Pang 等人^[30]也利用迁移学习进行单脉冲搜索,使得训练出的模型可以在多个巡天项目中进行迁移,提高了模型的通用性和适用范围。

本文尝试将迁移学习应用到周期性脉冲星候选体分类任务中来。在迁移学习的模型选择上,采用一种可迁移性评估方法对多种模型的可迁移性进行评估,在模型的输入特征选择上,仅将两张二维子图作为输入,在训练方式的选择上,采用自监督对比学习以进一步挖掘无标签数据的潜力,最终对训练得到的预训练模型进行微调并测试效果。

1.3 论文主要工作

迁移学习在近年来成为解决复杂下游任务的一种有效手段。脉冲星候选体分类任务面临着诸多挑战,如不同望远镜位置、数字后端设计的差异、脉冲星搜索流程的多样性以及周围环境的干扰变化等,使得设计深度学习模型并从头开始训练变得异常困难,因为这涉及到有标签数据少、数据极端不平衡、训练流程复杂、时效性差等问题。预训练模型是通过在大规模数据上进行训练得到的模型,通常具有丰富的语义信息和泛化能力。这使得它们能够应用于各种领域和任务,并且能够有效地减少从头训练的工作量和时间成本。通过使用预训练模型,可以利用已有的大规模数据集进行预训练,在脉冲星候选体分类等特定任务上进行微调。这种迁移学习的方法可以极大地提高模型的性能和泛化能力,同时也加快了模型训练的速度和效率。此外,预训练模型还可以帮助解决数据不平衡的问题,因为它们在预训练阶段已经学习到了各种特征和类别之间的关系,能够更好地处理脉冲星候选体分类所面临的极端不平衡数据带来的挑战。

本文主要工作包括:

- (1) 对多个基于 ImageNet 数据集进行预训练的模型在脉冲星数据集 HTRU1 上的表现进行对比分析,比较它们的 Precision、Recall 和 F1-

Score。首先，使用 LogME 方法对各个模型的在脉冲星数据集上的可迁移性进行评估；其次，使用后端融合的方式对频率-相位图和时间-相位图的两种特征进行融合；最后使用逐层解冻的方式对可迁移性评分高的 14 个模型进行微调 and 对比。

- (2) 整理 FAST 早期科学数据中心保存的部分 FAST 望远镜 19 波束 CRAFTS 数据。首先，使用 LogME 基于 FAST 公开数据集评估得到的可迁移性最好的模型，然后，利用自监督对比学习的方法，使用从 CRAFTS 无标签数据中提取的频率-相位数据和时间-相位数据进行模型预训练。最后，使用 FAST 验证集对模型效果进行验证。

1.4 论文结构安排

本文共有六个章节，各章节内容如下：

第一章为绪论，首先介绍了脉冲星的发现过程和研究意义、项目背景，国内外关于脉冲星候选体分类方法的发展历史和研究现状，以及本文的主要工作等内容。

第二章为相关工作，对使用 Presto 工具包搜索脉冲星的相关流程和本文涉及到的相关深度学习技术进行介绍。首先，介绍脉冲星候选体产生步骤和候选体诊断图，然后介绍各个深度学习相关的概念。

第三章利用公开数据集 HTRU1 对 ImageNet 预训练数据集利用 LogME 方法进行模型选择，并选择最后排名靠前的模型使用 HTRU1 数据集进行微调，初步测试 LogME 和预训练模型和微调方式在脉冲星候选体识别上的有效性。

第四章利用 FAST 早期科学数据中心的 CRAFTS 巡天产生的大量无标签数据进行训练，首先，使用 LogME 挑选出最适合 FAST 公开数据集迁移的模型。其次，利用自监督对比学习对挑选出来的最佳模型进行训练。最后，利用 FAST 公开数据集的训练集对模型进行微调迁移，并利用 FAST 公开数据集中的验证集进行测试。

第五章是实验和结果部分，LogME 在各个预训练模型上的评分和微调对比实验，自监督对比学习方式对比其它模型在 FAST 测试集上的表现。首先，介绍了实验的一些基本配置，接下来，详细分析了各个实验的结果。

第六章是总结和展望，主要分为两个部分，第一部分对本文的所有研究内容

进行总结，第二部分介绍了后续可以进行的工作。

1.5 本章小结

近几年来，计算机视觉领域发展迅速，涌现出来许多深度学习模型，但是这些深度模型由于脉冲星数据集本身的限制，无法利用现有的有限且分布于各个望远镜的标注数据集完成深度模型中大量参数的有效训练。迁移学习可以将已经学到的知识从一个任务（源任务）迁移到另一个任务（目标任务）来改善目标任务的学习效果，利用源任务中已有的知识，从而减少对大量训练数据的依赖，并且利用使用已经在源任务上训练好的模型或模型的部分作为目标任务的初始模型，也能一定程度上利用它的泛化能力和特征提取能力。因此，迁移学习的应用对于充分利用现有脉冲星数据有重要意义。基于此选题，本章介绍了脉冲星的发现过程和研究意义；随后介绍了脉冲星候选体分类的相关研究和研究现状，包括人工选择阶段、人工特征选择辅助机器学习阶段、深度学习阶段，并补充了迁移学习在脉冲星候选体分类的相关研究；最后介绍了本文的主要研究内容和本文的整体框架。

第二章 基本理论

2.1 脉冲星候选体产生步骤

脉冲星信号从信号源出发后，经历宇宙空间被位于世界各地的射电望远镜所接受，望远镜接收机有着提前设置好的中心观测频率，并有着固定的带宽，根据所观测对象的不同，望远镜有数字后端以合适的速率对望远镜接收到的信号进行采样和记录，最大化所搜索信号类型的灵敏度。采集下来的数据还会被分为多个频率通道，每个频率通道有

表 2-1 fits 文件部分信息展示

键	值
Telescope	FAST
Frontend	19BEAM
Backend	MB4K
Obs Date String	2020-04-09T16:12:35.889
MJD start time (DATE-OBS)	58948.67541538194444
MJD start time (STT_*)	58948.90380437646222
Sample time (us)	98.304
Central freq (MHz)	1250
Low channel (MHz)	1000
High channel (MHz)	1499.8779296875
Channel width (MHz)	0.1220703125
Number of channels	4096
Total Bandwidth (MHz)	500
DATA column	17
HDUs	primary, SUBINT

着固定的带宽，每个频率通道记录了以采样间隔取样的多个样本点。因此，对于固定长度内的脉冲星数据文件，可以简单的将它内部存储的数据理解为包含 N 个频率通道和 S 个采样点的一个 N 行 S 列二维矩阵 M 。对于寻找脉冲星这个任务来说，就是在多个固定时间长度的文件中找到最有可能是脉冲星候选体的信号。

原始文件的格式也有多种，FAST 数据主要使用的格式为 FITS(Flexible Image Transport System)格式，这是一种最初为天文学家设计的用于存储、传输和处理科学数据的文件格式，但现在也应用于地球科学、医学等领域。在 1982 年时 FITS 被国际天文学会确定为世界各天文台之间传输、交换数据的统一标准格式。它主要由一个二进制表格(Primary HDU)和零个或多个扩展的二进制表格(Extensions HDUs)组成。表 2-1 展示了一个文件 fits 文件存储的部分信息，包含望远镜相关的信息如名字、使用的接收机、数字后端、本次观测的开始日期和时间以及当前文件的开始时间、频率带宽信息以及 HDU 信息和数据存储的列位置信息。

获取到原始文件后，对于周期性脉冲星来说，使用深度学习进行脉冲星搜索的完整步骤如图 2-1 所示。候选体的生成主要采用 PRESTO(PulsaR Exploration and Search TOolkit)²中的流程，Presto 是一套由 Scott Ransom 开发的用于高效寻找脉冲星的搜索和分析软件工具包，主要由 C 和 Python 编写^[31]。流程中主要包含去干扰、消色散、快速傅里叶变换、加速搜索、折叠等主要步骤产生候选体。产生脉冲星候选体后，则可以使用计算机利用各种方法对候选体进行分类。在实际的搜索流程中，为了不漏过真实的信号，通常会将阈值会设置得比在方法研究过程中得出阈值更低，因此需要人工对模型分类的结果进行再次验证。人工确认后，即可对同一位置的疑似源使用同一望远镜或其它望远镜进行再次观测。接下来将对候选体特征提取前的各个步骤进行详细介绍。

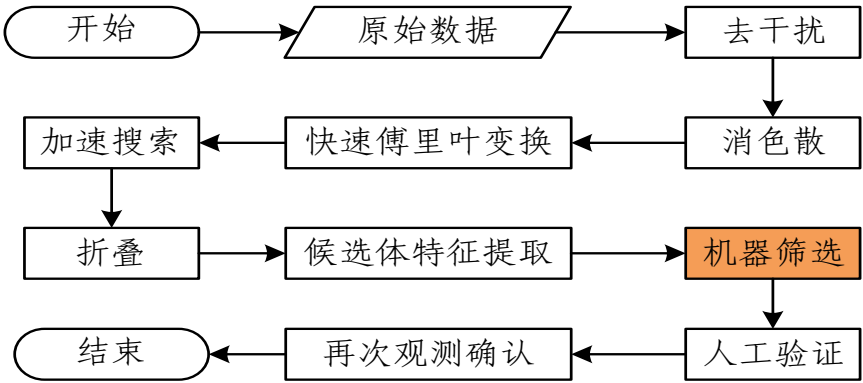


图 2-1 脉冲星搜索流程

其中，去干扰步骤的主要目的是去除能被望远镜接收到的但是属于非脉冲星发出的电磁辐射造成的干扰，它们可能来自于汽车启动、微波炉工作、飞机飞过、卫星通信等。许多脉冲星的信号本就微弱，因此如果不进行去干扰步骤则很可能

2 <https://github.com/scottransom/presto>

使得真实信号隐藏在地球上的这些人造干扰之下，导致遗漏微弱的脉冲星。

消色散步骤主要为了纠正射电信号在经过星际介质时产生的色散效应，即整个望远镜接收的带宽内的高低频信号到达望远镜时间不一致的问题。由于不同频率的射电信号到达时间不一致，会引起原本的脉冲星脉冲扩宽，影响到脉冲星的检测。纠正过程主要是通过一系列的信号处理技术，对接收到的信号进行色散导致的延迟进行补偿，使信号中的各频率通道的信号到达望远镜的时间基本一致。

快速傅里叶变换步骤通过对时域的射电信号进行 FFT 变换，将其转换为频域信号，从而能够清晰地展示信号在各个频率通道上的分布，如果原始时域数据中存在具有周期的信号，就能在周期所对应频率或者其谐波频率上出现峰。由于窄带脉冲的功率可能会分散在它的基波和谐波位置，出现基波功率太弱而达不到检测阈值的情况，通过谐波求和可以使其超过背景噪声或者是超出指定阈值，从而检测到脉冲信号以及确定它们的周期。

加速搜索的主要目的是检测和分析脉冲星运动中存在的加速度变化，由于脉冲星除了本身的自转外还可能伴随有其它的运动或者是伴星影响，使得脉冲星运动存在加速度变化从而产生多普勒效应，引起脉冲信号频率变化。加速度搜索通过检测和分析这些周期变化，能更有效地捕捉脉冲星。

折叠步骤是指将观测数据周期性地折叠，以增强脉冲星信号的信噪比，利用加速搜索给出的周期和色散，将位于同一相位的数据加和，这样就可以将多个周期内的信号叠加在一起，增强了信号的强度。对叠加后的数据进行平均化处理，以进一步减少噪声的影响，提升信号的检测信噪比，如果出现了一些信号显著超出了背景噪声水平，则可以产生认为它们可能是脉冲信号，需要做进一步筛选和判定。

在脉冲星研究中，折叠步骤是一个关键的过程，它利用先前得到的信息对数据进行折叠得到脉冲星的候选体。以 Presto 流程折叠后的结果为例，如图 2-2 所示，生成的脉冲星候选体图包含四个重要的子图，这些子图从 PICS^[24]开始一般作为各个端到端深度学习模型的输入，分别是平均脉冲轮廓图、时间-相位图、频率-相位图和 DM 曲线图。

- (1) 首先是平均脉冲轮廓图，它是根据脉冲星信号的周期将信号进行折叠后得到的脉冲轮廓图像。在该图中，横轴代表相位，纵轴代表辐射信号的

强度。通过观察平均脉冲轮廓，可以了解脉冲星在周期内辐射信号的变化情况，包括脉冲的形状、峰值和宽度等信息。

- (2) 接下来是时间-相位图，它展示了脉冲在一段时间内的信号持续情况。横轴表示相位，纵轴表示时间。通过时间-相位图，可以观察到脉冲在时间上的分布规律，了解脉冲信号的周期性和稳定性。
- (3) 第三个子图是频率-相位图，它显示了脉冲在不同频率通道上的信号分布情况。横轴代表相位，纵轴代表各个频率通道。频率-相位图可以帮助分析脉冲信号在频率上的特征变化，从而揭示脉冲星辐射信号的频率特性。
- (4) 最后是 DM 曲线图，DM 是色散测量值，表示数据在经过色散校正后，信号偏离白噪声的程度。DM 曲线图显示了在不同色散值下，信号的强度变化情况。通常情况下，DM 曲线的峰值对应的 DM 值就是为了校正该脉冲星色散效应而采用的色散测量值。在不考虑干扰的情况下（DM=0 处存在的峰值一般是干扰），数值越大的 DM 值表示信号与白噪声的偏离程度越大，也就越有可能是真实的信号。

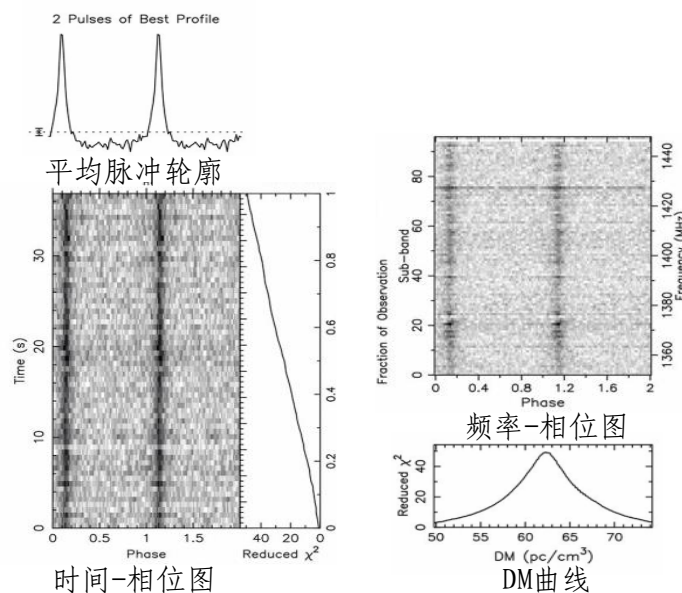


图 2-2 Presto 流程中产生的脉冲星候选体包含的主要图像特征（图片引用自 Presto 教程文件中的示例）

综合这四个子图的信息，可以对脉冲星的特征和性质有更深入的理解。平均脉冲轮廓图显示脉冲信号的形状和强度分布，时间-相位图和频率-相位图展示了脉冲在时间和频率上的变化规律，DM 曲线图则提供了关于信号色散校正的重要

信息。通过分析这些图像，研究人员可以更准确地确定脉冲星的性质，进行进一步的研究和分析。借助折叠步骤生成的脉冲星候选体图，可以对该信号是否为脉冲星信号作出初步判断。

2.2 CRAFTS 巡天

漂移扫描多科学目标同时巡天(The Commensal Radio Astronomy FAST Survey, CRAFTS)³是指利用 FAST 射电望远镜进行的多目标、多波段、多科学目标的巡天计划^[32]。该项目旨在充分利用 FAST 的强大观测能力，进行广泛的射电天文学研究，涵盖脉冲星搜寻、瞬态源监测、星际介质探测、快速射电暴(FRB)搜寻等多个前沿领域。它是控制 FAST 使用特定偏转角度漂移扫描，并产出满足寻找不同科学目标的数据，以实现超奈奎斯特采样大天区覆盖的一种工作模式^[33-35]，目前主要观测的区域为赤纬-14 度到+5.5 度内，这项巡天计划主要是为了充分利用 FAST 灵敏度，同时最大程度地降低系统操作的复杂性。利用 FAST 搭载的 19 波束 L 波段的接收机，将其对准特定角度并连续获取数据，通过多个后端同时处理，以获取脉冲星搜索、中性氢星系、中性氢成像和射电瞬变数据。截至 2022 年 8 月 1 号，已完成约 18%的天空覆盖。

2.3 相关深度学习理论和方法介绍

2.3.1 深度学习模型

深度学习模型是一类基于深层神经网络结构的机器学习模型，它通过多层次的非线性变换和特征抽取来实现对复杂数据的建模和学习。与传统的浅层模型相比，深度学习模型具有以下几个显著的区别和优势：

- (1) 深度层次：深度学习模型具有多个隐层，层次结构更加复杂，可以学习到数据的多层次抽象特征，从而提高了模型的表达能力和学习能力。
- (2) 非线性变换：深度学习模型中每一层都包含非线性激活函数，例如 ReLU、Sigmoid 等，这些函数能够引入非线性因素，使得模型可以处理更为复杂的数据模式和关系。
- (3) 端到端学习：深度学习模型可以进行端到端的学习，直接从原始数据中学习特征表示和任务映射，无需手工设计特征提取器，可以自动学习特

3 <http://groups.bao.ac.cn/ism/CRAFTS/CRAFTS/>

征表示，不需要人工提供特定领域的特征工程，能够更好地适应不同数据和任务，简化了模型设计和开发流程。

- (4) 大数据处理：深度学习模型对大规模数据的处理能力较强，可以利用大量数据来训练模型，从而提高模型的泛化能力和性能。

残差网络(ResNet)作为深度学习模型中的一种重要结构，在加深网络层次方面起到了至关重要的推动作用。它通过引入残差连接(Residual Connection)解决了深层网络训练中的梯度消失和梯度爆炸问题，从而可以训练非常深的网络，它的核心思想是通过跨层连接将输入信号与输出信号相加，使得模型可以学习残差映射而不是直接映射，从而更容易地训练深层网络^[36]。ResNet 的结构设计使得模型在增加深度的同时，保持了良好的性能和稳定的训练过程，成为了深度学习领域中的重要突破和进展之一。此外，近几年十分火热的 Attention 注意力机制也广泛应用于各种模型中，它模拟了人类在处理信息时的注意力分配过程。它的核心思想是在模型中引入一种机制，使得模型可以自动学习并集中注意力于对当前任务最有帮助的信息^[37]。残差连接和注意力机制的出现极大地推动了深度学习模型的发展，许多借鉴这种残差连接和注意力机制思想的模型被提出来，例如 SENet、Vit、Swin 等。这些模型在深度学习领域取得了显著的成果，不断加深网络结构，形成了更大规模、更高性能的大模型^[38-40]。大模型在许多领域都展现了出色的表现，例如在图像分类、目标检测、语音识别和自然语言处理等任务上取得了令人瞩目的成就。因此，不断提出新的结构和加深模型层次提升模型的表达能力、泛化能力和适应性，从而更好地应对复杂的数据任务和挑战。

然而，这类大模型通常需要庞大的训练数据和性能高且丰富的计算资源，对于脉冲星候选体分类这样的任务来说，既无法收集到庞大的高质量训练集，也由于搜索流程中的其它处理工作如消色散等步骤需要消耗大量计算资源，此外，由于实验室设备显存有限，无法满足从头大模型训练的需求。因此，如果要使用这些大模型就需要另辟蹊径使其能有效应用于脉冲星候选体分类任务。

2.3.2 ImageNet 数据集和预训练模型

为了解决这个问题，研究者基于模型的浅层特征包含更多与语义无关的、涉及图像的一些细节例如颜色、纹理、边缘和棱角等信息这一认知，认为一个深度模型中的浅层学习到特征表示可以迁移到其它任务中使用。为了得到一个适合迁

移的模型，需要有大型数据集，并基于这个大型训练集对模型进行充分预训练。

ImageNet 1k 数据集，也称 ISLVR2012，是较大的 ImageNet 数据集的一个子集，通常用于计算机视觉任务中来训练深度学习模型。它涵盖了 1000 个物体类别，训练集包含 1281167 张图像，验证集包含 50000 张图像^[41]。

预训练模型是一种已经在大型数据集进行充分训练以完成特定任务的深度学习模型，它可以直接使用，也可以根据实际的应用需求自由定制，一旦预训练完成，这些模型通常会成为各种下游任务的基础。它的意义在于无需从头开始构建一个全新的深度学习模型。而是直接使用经过预训练得到的表现最优模型，从中可以获取为了完成某个特定任务训练完成得到的深度学习模型中各个层次各个权重，从而节省时间、金钱和精力。这些权重可以进行修改，还可以向模型中添加更多的层以进一步定制或微调它，借助预训练模型，可以更快地搭建起用于脉冲星候选体分类任务的 AI 模型。微调是深度迁移学习中一种常见的技术，通过微调可以充分利用预训练期间学到的通用特征，并借助新任务上的数据对高层模型参数进行微调。这样的方法大大加速了模型的训练过程，并提高了模型在特定任务上的性能。在自然语言处理领域，预训练模型经常用于微调到情感分析、问答系统、文本生成等特定任务。通过在大量语言数据上训练，预训练模型可以学习到丰富的语义信息，为这些任务提供有力支持。同样，在计算机视觉领域，预训练模型也经常被微调到医学成像、自动驾驶、生物分析、药物发现等特定领域^[42]，为这些应用提供了高效的特征提取能力。

ImageNet 1K 数据集由于其庞大的类型数目和各个类型中数据的数据量以及高质量标注信息，使得它在计算机视觉领域被广泛使用，也出现了许多基于 ImageNet 数据集的预训练模型。通过在这个数据集上进行训练，可以获得泛化能力较强的预训练模型，为其他小型特定领域任务的发展提供了重要的基础，成为许多应用领域中的重要推动力量。

2.3.3 PyTorch 介绍和 LogME 介绍

PyTorch 是一个由 Meta 的人工智能研究团队开发和维护的开源深度学习开发框架⁴。它以其提供的动态计算图、张量操作、自动求导、模块化设计、GPU 加速以及与 numpy 包的高效转换接口而闻名。这些功能为研究人员和开发者提供

4 <https://pytorch.org/>

了便利，使他们能够更轻松构建和训练深度学习模型。动态计算图是 PyTorch 的一个重要特性，它允许用户按照需要创建计算图，而不需要预先定义。这使得模型的构建过程更加灵活，能够更快地实验不同的结构和算法。另外，PyTorch 提供了丰富的张量操作函数，使用户可以方便地进行各种数学运算和张量操作，进而加速模型的开发和优化过程。自动求导是 PyTorch 的另一个重要功能，它可以自动计算张量的梯度，并且支持动态图的梯度计算。这使得用户无需手动编写反向传播算法，大大简化了深度学习模型的训练过程。PyTorch 还采用了模块化设计，用户可以将模型按照组件进行构建，方便了模型的重用和调试。此外，PyTorch 提供了 GPU 加速，可以在 GPU 上高效运行深度学习计算，加快了模型训练的速度。最重要的是，PyTorch 拥有庞大的社区资源，提供了许多预训练模型和训练好的参数。用户可以直接通过代码加载这些模型，快速验证模型的可用性，节省了开发时间，加速了研究和应用的进程。因此，PyTorch 在深度学习领域中具有广泛的影响力和应用价值。

Pytorch 框架通过 torchvision 提供了大量预训练模型⁵，本文基于 0.16 版本的分类预训练模型如表 2.2 所示。

表 2-2 torchvision 中提供的分类预训练模型

• AlexNet	• ConvNeXt	• DenseNet	• EfficientNet
• EfficientNetV2	• GoogLeNet	• Inception V3	• MaxVit
• MNASNet	• MobileNet V2	• MobileNet V3	• RegNet
• ResNet	• ResNeXt	• ShuffleNet V2	• SqueezeNet
• SwinTransformer	• VGG	• VisionTransformer	• Wide ResNet

为了从众多的预训练模型中选择合适的预训练模型，本文采用 LogME(Log Maximum Evidence)方法，LogME 是一种用于加速预训练模型选择过程的方法，原理是将预训练模型视为特征提取器，与微调过程类似，最后增加一个线性层建立特征与标签的关系，并用统计学的证据来衡量效果的好坏^[43]。如今，对于深度学习框架来说，如果要为实际的应用选择一个合适的预训练模型，通常会使用一些简单粗暴的方法，比如查看类似的研究或应用中一般采用的基础模型，或者是直接选用预训练在原始任务上指标高的。如果要准确选择效果最好的模型，还需

5 <https://pytorch.org/vision/0.16/models>

要在大量候选模型中对每个模型进行微调查看效果，这将是一个非常耗时的操作，因为微调涉及到更改模型和训练模型，通常需要耗费大量的时间。为了对 torchvision 中提供的各个预训练模型在脉冲星上的效果作对比，显然需要一种高效的方式来评估各个模型的效果，LogME 相较以往的方法，将利用预训练模型在训练集上进行调参、微调再为模型打分的方式扩展到了除有监督分类迁移任务外的其它场景，覆盖了视觉、自然语言处理、分类、回归、无监督预训练模型等方向。它优越的性能来自于无须梯度计算、无须超参数优化以及算法实现上的优化，因此，本文将采用 LogME 用于预训练模型选择和自监督对比学习过程中的模型训练早停方法。

2.3.4 迁移学习

迁移学习(Transfer Learning)最初起源于心理学和教育学领域，强调一种学习对另一个学习的影响。随着机器学习和人工智能的发展，迁移学习也逐渐成为了计算机科学领域的研究热点之一^[44]。迁移学习旨在解决数据稀缺或标注困难的问题，通过利用源领域的知识来改善目标领域的学习性能，为模型训练提供更多的先验知识和信息。在视觉领域，迁移学习发挥了重要作用，特别是在图像分类、目标检测和语义分割等任务中。通过迁移学习，可以将大规模数据集上训练得到的深度神经网络模型中的特征和知识迁移到相对较小的目标数据集上，从而加速模型的训练过程并提升模型的泛化能力。例如，使用预训练的卷积神经网络模型在大规模图像数据集上进行训练，然后通过微调或特征提取的方式将这些模型迁移到目标任务上，可以获得更好的性能表现。根据迁移学习的方式和目标任务的不同，可以将迁移学习方法分为几种主要类型：

- (1) 基于样本的迁移学习：在源领域和目标领域之间，对不同样本赋予不同的权重，从而关注哪些对目标更有帮助的样本，以便更好地适应目标任务。
- (2) 基于特征的迁移学习：通过对源域和目标域的特征进行变换，可以通过自编码器和生成对抗网络进行特征变换，使它们在同一特征空间中更相似。
- (3) 基于模型的迁移学习：使用源领域的预训练模型，然后在目标领域上进行微调，这适用于目标领域数据不足的情况；在一个模型中同时处理源

领域和目标领域的任务，适用共享的特征提取层促进知识迁移。

- (4) 基于关系的迁移学习：通过挖掘源领域和目标领域之间的关系，将已学到的知识迁移到目标任务中，比如“猫”、“狗”在具有相似的特征，因此可以将其中一类领域中学到的知识应用到目标领域的分类任务。

本文采用第三种方式来完成从 ImageNet 分类任务到脉冲星候选体分类任务的迁移学习，首先选择出在脉冲星数据上可迁移性表现好的预训练模型，再从模型的角度将其迁移到具体的任务上来。

2.3.5 对比学习

对比学习(Contrastive Learning)是一种无监督学习方法，旨在通过比较不同事物之间的差异和相似性来学习有效的表示^[45]。深度学习通常需要大量标记数据，这既耗时又昂贵，无监督学习通过自主发现数据中的结构来节省时间和资源。通过比较图像对、视频帧或图像和文本之间的差异和相似性，可以学习到更具有判别性和泛化性的特征表示。例如，对比学习网络和孪生网络等模型结构被广泛用于学习图像对之间的相似性和差异性，从而可以用于图像检索和人脸识别等应用场景。根据学习目标和方法的不同，对比学习可以分为几种主要类型：

- (1) 生成式学习：这类学习方式以自编码器为代表，通过生成数据使其在整体或高级语义上与训练数据相似，这类学习方式需要关注实例上的细节，因此模型较复杂。
- (2) 对比式学习：对比式学习着重于学习同类实例之间的共同特征，区分非同种实例之间的不同之处。

与生成式学习相比，对比式学习不需要关注实例上的细节，只需在抽象语义级别的特征空间上学会对数据进行区分，因此模型更简单且泛化能力更强，也是本文使用的方式。

2.3.6 自监督学习

自监督学习(Self-Supervised Learning)是一种无监督学习方法，其核心思想是利用无标签数据自身来构造监督信号，通过设计巧妙的代理任务（proxy task 或 pretext task），使模型能够在没有人工标注的情况下学习到有价值的数据表示^[46]。这种学习方式的关键在于能够从数据中自动产生“伪标签”，以此模拟有监督学习中的标签信息，驱动模型进行自我训练。

在视觉领域，自监督学习被广泛应用于图像表示学习、目标检测和图像生成等任务中。通过利用数据自动生成的标签或任务来学习数据的表示，可以在无需人工标注大量数据的情况下获得有效的模型。例如，图像颜色化、图像补全和图像旋转等任务可以被用作自监督学习的训练目标，从而学习到更具有判别性和泛化性的特征表示。

与监督学习依赖于大量的带有标注的数据来进行模型训练相比，自监督学习则通过数据自动生成的方式来学习数据的表示，无需大量的人工标注数据。与无监督学习主要依赖于数据本身的结构和分布来进行学习，无需任何标签信息相比，自监督学习则利用数据自动生成的标签或任务来学习数据的表示，相对于无监督学习更加灵活和可控。

自监督学习在图像识别、语义理解和推荐系统等领域有着广泛的应用，可以有效地利用大量无标注数据进行模型训练，降低了对于大量标注数据的依赖性。

2.4 本章小结

在本章中，对文章的相关技术进行了详细介绍，包括周期性脉冲星搜索流程，CRAFTS 巡天以及相关的深度学习理论。周期性脉冲星搜索流程部分主要详细介绍了到产生候选体之前的各个步骤，同时还介绍了本文数据集中单个数据所包含的图像信息，即脉冲星候选体折叠图。同时，CRAFTS 巡天数据作为本文数据处理的对象，也对其进行了简单介绍。随后，介绍了深度学习相关的理论，包括本文主要使用的深度学习框架，以及本文主要使用的预训练模型评估方法，随后介绍了本文涉及到的预训练微调、迁移学习、对比学习、自监督学习等相关理论，为后续内容的实验奠定基础。

第三章 利用预训练-微调进行脉冲星候选体分类

3.1 数据集介绍

本章节所采用的公开数据集为 HTRU1 Batched Dataset⁶，该数据集源自 High Time Resolution Universe Survey(HTRU)巡天项目^[21,47]中，针对银道纬度 $\pm 15^\circ$ 范围内的训练集部分，其初始目的是为了训练 SPINN(Straightforward Pulsar Identification using Neural Networks)分类器。HTRU 巡天项目是一项颇具影响力的天文研究计划，它主要依赖于澳大利亚联邦科学与工业研究组织(CSIRO)所运营的帕克斯射电望远镜进行观测工作。帕克斯射电望远镜，坐落在澳大利亚境内，以其直径达 64 米的庞大口径，在全球射电天文观测领域占据重要地位，堪称世界上最著名的射电天文观测设施之一。其卓越之处在于配备了多种类型接收器，使得望远镜能够对天空中多个区域多种类型的感兴趣对象进行观测，大幅度提升了巡天效率，对于广泛搜寻和精确识别脉冲星等宇宙现象具有重大意义。帕克斯多波束接收器在执行 HTRU 巡天时使用的 13 波束接收机，操作时设定在 1369MHz 附近的频率范围内，整个接收频率带宽约 400MHz。脉冲星因其强烈的射电辐射特性，在这一频段内易于被探测到，从而使得帕克斯望远镜能够在一次观测中捕获到大量的潜在脉冲星信号。鉴于 HTRU 巡天与 FAST 开展的多波束 CRAFTS 巡天在多波束技术和观测频段上存在相似性，故 HTRU 数据集是评估预训练效果的理想之选。

具体而言，HTRU1 Batched Dataset 包含了 60,000 张尺寸为 32×32 的三通道图像。这些图像的各个通道分别承载着不同的脉冲星特征信息：通道 1 记录了脉冲星的周期和色散特性，这两项数据对于理解脉冲星的自转规律及星际介质影响不可或缺；通道 2、3 则分别保存了频率-相位数据和时间-相位数据。

在全部 60,000 张图像中，有 50,000 张被划分为训练集，剩余 10,000 张用作测试集。数据集内共包含 1,194 个脉冲星信号样本，以及 58,806 个非脉冲星信号样本。这些正负样本按照一定比例均匀分布于训练集和测试集中，旨在确保模型训练过程中的泛化能力和测试结果的有效性。通过精心构造且富含脉冲星特性的

⁶ <https://as595.github.io/HTRU1/>

HTRU1 数据集，能够有效地评估各种预训练模型在脉冲星候选体分类任务上的性能。

3.2 数据预处理

公开数据集 HTRU1 Batched Dataset 受 CIFAR-10 数据集启发，以 `numpy` 数组的形式进行保存，这种格式提供了极大的便利，允许直接使用 `numpy` 库进行快速、高效的数据读取操作。为了能在 LogME 方法中对数据集基于 ImageNet 预训练模型上的可迁移性进行评估，需要对原始数据集进行特定的预处理调整，使之满足模型输入要求。

首先，需要从数据集中提取出与脉冲星识别密切相关的频率-相位信息和时间-相位信息。这些信息分别存储在数据集的第二、第三通道中，与对应的脉冲星或非脉冲星标签紧密关联。为了实现这一目标，设计了一个自定义数据集类，该类继承自 PyTorch 框架中的 `Dataset` 类。通过自定义数据集不仅能够根据实际需求灵活定义数据集的加载方式、预处理步骤以及数据增强策略，还能确保新数据集与 PyTorch 原生的数据加载工具——`DataLoader` 高效对接。这种高度定制化的数据集不仅简化了数据处理流程，还充分利用了 PyTorch 提供的丰富功能和工具，极大地提升了数据处理、模型训练和评估的效率。

在明确所需提取的特征之后，发现每个单一特征图呈现出 32×32 像素的单通道灰度图像，这样的图像格式并不符合基于 ImageNet 预训练模型的输入要求。为此，借助 PyTorch 内置的 `transforms` 模块对这些特征图进行了以下一系列转换：

- (1) 通道扩展：将原始的单通道图像转化为三通道图像。具体操作是将单通道图像在红(R)、绿(G)、蓝(B)三个通道上进行重复赋值，即 R 通道、G 通道和 B 通道的像素值完全相同，均等于原单通道图像的像素值。这样就实现了从灰度图像到彩色图像的转变，满足了 ImageNet 预训练模型对多通道输入的要求。
- (2) 尺寸调整：将图像的分辨率由 32×32 像素调整至 224×224 像素。这一改动确保了预处理后的图像尺寸与基于 ImageNet 预训练模型的标准输入尺寸($3 \times 224 \times 224$)相匹配，确保模型能够顺利接收并处理这些数据。
- (3) 类型转换：将图像数据类型从原始格式转换为 PyTorch 张量类型。这种转换使得图像数据可以直接利用 PyTorch 提供的各类张量运算函数进行

计算，极大地增强了数据处理的灵活性和效率。

- (4) 标准化:对图像进行归一化处理,即将图像各通道的像素值减去 ImageNet 数据集相应的平均值,并除以相应标准差。这一标准化步骤有助于消除不同图像间的亮度和对比度差异,使预处理后的图像与 ImageNet 预训练模型在训练过程中所见数据保持一致的统计特性,从而提高模型对 HTRU1 数据集的适应性。

通过上述一系列定义好的数据预处理操作,原本的 HTRU1 数据集成功转化为符合 ImageNet 预训练模型输入要求的新数据集。

3.3 模型选择

本文选择使用清华大学 You 等人提出的 LogME 方法来评估预训练模型在脉冲星识别任务中的可迁移性。LogME 作为一种专门针对预训练模型评估的技术手段,旨在量化模型在特定下游任务中的潜在表现,无需进行实际的微调或训练,从而省大量计算资源和时间成本^[43]。

LogME 的核心思想主要是两点,第一点是最大证据估计,通过计算预训练模型提取特征后标签证据的最大值(边缘似然性)来评估模型与目标任务的兼容性。这种方法对过拟合具有免疫力,并适用于分类和回归任务。第二点是对数最大证据(LogME),将最大证据取对数,作为预训练模型在迁移学习中的评估指标。更高的 LogME 值表示该模型在数据集上有更好的迁移性能。

为了利用 LogME 完成可迁移性评估,首先,将待评估的预训练模型中的可训练参数冻结,将其视为一个固定的特征提取器,确保在整个 LogME 评估过程中不更新其内部参数。接着,定义钩子函数获取传递到模型分类头前的输出,将 HTRU1 数据集中的经过预处理的图像输入到该预训练模型中进行前向传播,将得到的特征与对应的脉冲星或非脉冲星标签进行对比。使用概率密度函数 $p(y|F)$ 来量化特征与标签之间的关系,其中 F 代表特征矩阵,包含了所有样本的特征向量, y 则表示对应的标签。通过对 $p(y|F)$ 的分析,可以从统计学角度理解预训练模型提取的特征对于区分脉冲星与非脉冲星的能力。在统计学中,证据(evidence)是一个衡量模型与数据集整体兼容性的量,反映了模型在给定数据生成假设下的合理性。在 LogME 方法中,通过计算边缘化似然

$$p(y|F) = \int p(w)p(y|F, w)dw \quad (3-1)$$

来衡量特征 F 与标签 y 之间的全局关联性，其中 w 代表线性层的权重参数， $\int$ 符号表示对所有可能权重 w 的积分操作。

最后，LogME 通过最大化证据来直接求解线性层的超参数 α 和 β ，这两个参数决定了特征与标签之间线性关系的强度和方向。在完成超参数求解后，即可得到对数最大证据值，该值直观反映了预训练模型在脉冲星识别任务上的潜在适用性。较高的 LogME 值意味着预训练模型所提取的特征在无需微调的情况下，与脉冲星识别任务具有较高的契合度，有望在实际微调后取得良好的识别效果。反之，较低的值则暗示模型可能需要较大的调整才能适应新任务，或者在该任务上表现不佳。

模型的选择过程非常简单，选择模型，输入预训练数据，得到 LogME 评分结果，利用从 torchvision 库共计 76 个涵盖了多种主流架构及其变种的预训练模型。这些模型包括但不限于表 2-2 中的模型以及变种，它们各自代表了深度学习领域在图像分类任务上的不同设计理念和技术演进。

对于每个预训练模型，LogME 主要关注其最后一层全连接层（分类头）之前的输出，即所谓的“瓶颈”特征。这是因为这些特征经过多层非线性变换，往往蕴含了丰富的语义信息，且更易于跨任务迁移。

选取得分排名前 10 的预训练模型作为验证模型。进行实际的微调过程。在预训练模型的基础上，添加一个任务特定的输出层（如全连接层），并使用 HTRU1 中的训练集进行有监督训练。微调过程中，仅更新新增加的输出层及从后往前解冻的指定层数的参数，保持预训练模型其余部分的权重不变。最后使用独立的测试集对微调后的模型进行评估，计算各项评价指标以衡量模型在脉冲星识别任务上的实际性能。对比这些模型在测试集上的表现，检验 LogME 方法应用到脉冲星候选体分类任务的可迁移性评估的有效性。

3.4 模型微调

在经过 LogME 方法的评估后，得到了分别针对时间-相位图和频率-相位图数据的可迁移性排名前十的预训练模型。尽管有些周期性候选体脉冲星分类方法倾向于使用同一模型提取两种图像特征，以节约资源，但在本研究中仍对时间-

相位图和频率-相位图采用两个独立的预训练模型进行特征提取，以充分挖掘两种图像各自特有的信息。

微调模型的整体框架如图 3-1 所示，处理时间-相位图和频率-相位图时，选用相同的预训练网络作为基础模型来提取特征。这种做法在脉冲星识别领域十分常见，因为两种特征存在一定的相似性，但又不完全一致，使用同种基础模型但分开训练的方式使得它能有效地从两种类型的图像中捕获关键的时空信息和模式，为后续的分类决策提供强有力的支持。

选定的预训练网络分别对输入的时间-相位图和频率-相位图进行特征提取。经过多层卷积和池化操作后，网络将原始图像转化为高维、低分辨率的特征张量。这些特征张量不仅浓缩了图像的视觉特征，还包含了对脉冲星信号频率特性和相位分布等物理属性的高度抽象表示。

为了保证模型能够综合多种特征，并使得深度模型取得最优的性能，通常需要选择一种将各个特征融合起来的方式。根据所采取方式的不同，可简单分为前端融合方式、后端融合方式和混合融合方式，分别对应原始数据拼接和特征拼接等方式、模型输出特征融合和决策层融合等以及综合两种方式。本文从效率和效果两方面考虑，将从时间-相位图和频率-相位图各自提取到的特征张量进行直接求和操作，实现特征融合。这种简单的融合策略旨在整合来自不同图像源的特征，使得模型能够同时考虑时间域和频率域的特征信息，形成对脉冲星信号更为全面、立体的理解。尽管此处采用了简单的求和操作，但在实际应用中，根据任务需求和数据特性，也可选择更为复杂的融合策略，如注意力机制、拼接融合等。

融合后的特征张量作为输入传递给逻辑回归模型。逻辑回归是一种简单而高效的二分类模型，它通过学习权重参数，将特征张量映射到一个介于 0 和 1 之间的连续数值，该数值可以理解为模型对输入图像属于脉冲星类别的置信度或概率。逻辑回归模型的优势在于模型结构清晰、易于解释，且在小规模数据集上训练速度快、泛化性能稳定。最后，将逻辑回归模型输出的置信度值与预先设定的阈值进行比较。

通常情况下，若置信度值大于阈值，则将该图像判定为脉冲星类别；反之，则认为其属于非脉冲星类别。需要注意的是，由于 HTRU1 数据的标签 0 表示脉冲星，而 1 表示非脉冲星，所以在 HTRU1 数据集最后的阈值比较上需要反向判

定。阈值的选择直接影响到模型的查准率和查全率，因此需要根据实际任务需求和数据分布进行调整和优化。例如，在对新发现脉冲星敏感度要求较高的情况下，可以选择较低的阈值以提高查全率，容忍一定的假阳性率；而在追求高精度识别的场景下，则应适当提高阈值，以降低假阳性误判。在实际的应用中，为了不放过任何可疑的信号，通常会设置阈值尽可能使得查全率达到 100%。

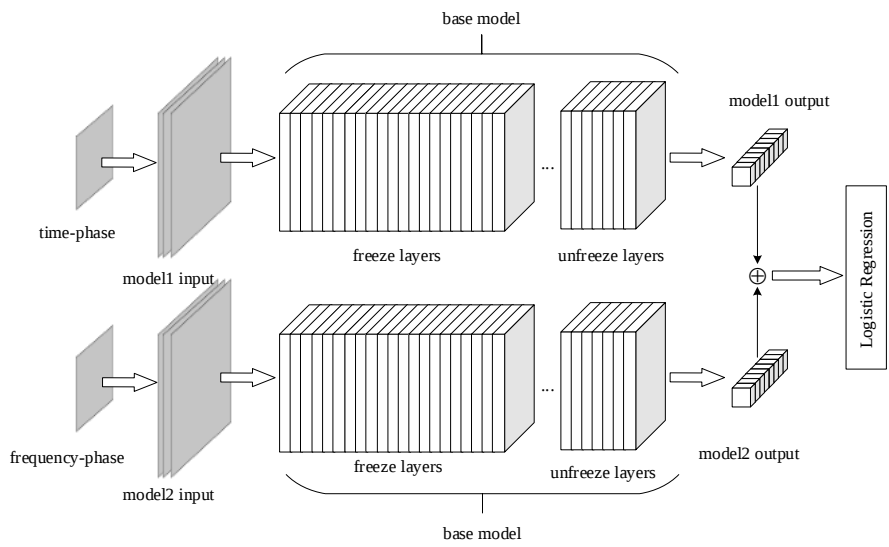


图 3-1 微调模型架构图

在模型微调过程中，本文采取了验证集来防止过拟合现象的发生，以确保模型在训练过程中能够保持良好的泛化能力。将原始训练集按照 80%：20%的比例随机划分为训练集和验证集。这一划分策略确保了大部分数据用于模型训练，而保留一部分数据即验证集用于实时监控模型在未见过数据上的表现。通过观察验证集上的性能指标，可以及时调整训练策略，避免模型过度适应训练数据，从而在新数据上出现性能下滑。

接下来，详细阐述微调过程中的参数设置：

- (1) 批次大小：将训练批次大小设置为 64。这意味着在每次迭代中，模型会一次性处理 64 张图像，并基于这些图像计算损失和梯度，进行参数更新。
- (2) 提前终止：为了避免过度训练，采用早停策略来动态控制训练迭代次数。具体来说，参考 fetch^[29]中迁移学习中微调实验设置，当验证集上的损失连续三次上升时，即表明模型开始出现过拟合迹象，此时停止训练。早停策略能够在保证模型性能的同时，有效地节约计算资源和时间。
- (3) 损失函数：在本章脉冲星候选体分类任务中，选用二分类交叉熵损失函

数(Binary Cross-Entropy Loss, BCELoss)作为模型的损失函数。BCELoss 通常用于二分类问题中，它的计算方式如式 3-2 所示，其中 N 是样本数量， y_i 是真实标签（本文中为 0 或 1）， $\hat{y}_i$ 是模型预测的概率值（0~1 之间的数）。BCELoss 能够有效衡量模型预测概率与真实标签之间的差距，并通过反向传播过程指导模型参数更新。由于其在二分类问题中具有良好的梯度传递特性，模型通常能够较快地收敛至最优状态。

$$BCELoss = -\frac{1}{N} \sum_{i=1}^N [y_i \cdot \log(\hat{y}_i) + (1 - y_i) \cdot \log(1 - \hat{y}_i)] \quad (3-2)$$

- (4) 优化器与学习率：为了优化模型参数，选择了随机梯度下降(stochastic gradient descent, SGD)算法作为优化器。SGD 以其简单高效的特点，在深度学习领域广泛应用。设定学习率为 10^{-4} ，这是一个较小的学习率值，有助于模型在微调时进行精细化调整，防止参数振荡。此外，设置动量参数为 0.99，动量机制能够加速模型在梯度方向上的移动，同时平滑训练过程，减少震荡。
- (5) 脉冲星分类阈值：最后，设定脉冲星分类的阈值为 0.5。在模型预测阶段，对于任一输入图像，若模型输出的脉冲星概率大于 0.5，则将其判定为非脉冲星；否则，判定为脉冲星（HTRU1 数据集标签与常用设置相反）。0.5 是二分类问题中常用的默认阈值，实际应用中需根据具体需求和数据分布特点进行调整。

3.5 本章小结

本章首先介绍了 HTRU1 数据集，紧接着介绍了针对 HTRU1 数据集的预处理和工程上的一些实现细节，随后，详细介绍了 LogME 评估预训练模型的原理，最后，主要是介绍了微调相关的参数设置和一些实验设计的细节。由于仅通过微调的方式来对脉冲星数据进行分类的效果一般，所以本文基于大量已有无标签数据，利用自监督对比学习对模型进行预训练。

第四章 利用自监督对比学习进行脉冲星候选体分类

4.1 数据集介绍

本章所使用的公开数据集是 FAST 数据集，这是 Wang 于 2019 年公开的数据集⁷，为方便起见，后续称其为 FAST-2019 数据集。该数据集分为训练集和测试集两个部分。为了平衡样本的数量，训练集采用了近似 1:1 的比例，总共包含 1,835 个样本，其中包括 837 个正样本和 998 个负样本。而测试集采用了接近真实线上情况的比例，没有进行平衡处理，其中包含 326 个脉冲星样本和 13,321 个 RFI 样本。

FAST-2019 数据集中所包含的数据为脉冲星折叠结果 pfd 文件，需要从 pfd 提取出所需要的频率-相位数据和时间-相位数据，并经过进一步的处理才能输入到模型中进行训练和验证。除了这个数据集之外，本文还利用了位于 FAST 早期科学数据中心的 CRAFTS 巡天数据。这批数据是 2020 年 4 月 12 日由 19 波束接收机中 M01 波束采集，后经 FAST 早期科学数据中心处理留档于 NFS 存储中的累计 1,092,783 个 pfd 数据文件，用于进行自监督对比学习的训练。

通过使用自有的 CRAFTS 无标签巡天数据和公开数据集 FAST-2019，本文得以构建起一个全面且具有代表性的训练和验证数据集，用于评估模型在脉冲星分类任务上的性能。与真实情况接近，通常拥有一批由望远镜接收并按照规定搜索流程处理完的无标签数据，和一批在调试阶段接收到的已知脉冲星的真实样本和干扰样本，通过这种方式，能快速地为每一个望远镜训练一个针对自身数据的深度模型用于脉冲星候选体分类任务。

4.2 数据预处理

由于 FAST-2019 所提供的数据格式与无标签自建数据集相同，都采用了 pfd 格式，其中包含了脉冲星的脉冲轮廓数据，即脉冲星信号在不同脉冲相位上的强度信息。此外，元数据中还可能包含有关观测参数、仪器设置和数据处理的信息。然而，由于 pfd 数据文件中包含了一些多余的信息，为了更好的满足模型输入的要求并且节约存储空间，只提取时间-相位图和频率-相位图的对应的那部分数据。

⁷ https://github.com/dzuwhf/FAST_label_data

为了实现数据的提取和处理，本文采用了 PICS 中使用的的数据提取模块。该模块可以从 pfd 文件中提取出需要的输入图片数据，并将其保存为 numpy 数据，格式为 npy，以方便在自定义数据集类中使用 numpy 模块加载。最终对应的数据大小范围两部分，时间-相位数据和频率-相位数据，分别占用 17GB 存储空间。在自定义数据集类中，定义了多种形式的获取方法，以满足多种模型输入数据集的需求。其中包括：

- (1) 单独获取时间-相位图和对对应标签的方法。这种方法适用于需要对时间-相位图进行独立处理的模型，例如对脉冲星信号的时间特征进行分析的模型。
- (2) 单独获取频率-相位图和对对应标签的方法。这种方法适用于需要对频率-相位图进行独立处理的模型，例如对脉冲星信号的频率特征进行分析的模型。
- (3) 获取匹配的时间-相位图、频率-相位图和对对应标签的方法。这种方法适用于需要同时考虑时间和频率特征的模型，例如对脉冲星信号的整体特征进行分析的模型。

这些不同形式的获取方法可以方便地应用于不同类型的模型中。可以使用单独获取时间-相位图或频率-相位图的方式进行各自可迁移性评估，而对于已经搭建好的模型，使用获取匹配的时间-相位图、频率-相位图和对对应标签的方法进行训练和测试。

4.3 自监督对比方案设计

尽管脉冲星领域可以称得上大数据领域，各个望远镜收集到的数据非常多，但是有标签的数据总是有限的，而且对于脉冲星领域来说，除了目前收集的有限的脉冲星外，其余的都是干扰，有用的数据也无法形成高质量的数据集。因此，自监督学习便可以利用这些大规模的无标签数据学习通用的特征表达，然后，将学习到的预训练模型迁移到后续的脉冲星候选体分类任务中，本文主要采用自监督对比学习的方式，即通过衡量正样本和负样本之间的距离进行自监督学习以获得通用的特征表达。

由于数据是无标签的数据，这里的正样本和负样本与传统的正负样本定义不一致，对于任意一个样本，称它为锚点样本，利用适用于脉冲星的样本增强方法

以上内容仅为本文档的试下载部分，为可阅读页数的一半内容。如要下载或阅读全文，请访问：<https://d.book118.com/455303304324012101>