

全国高校地理信息科学实验教学系列丛书

时空大数据分析 实验教程

李晶 李柏延 刘宪锋 周自翔◎编著

清华大学出版社
北 京

内 容 简 介

本书是在新工科与新文科背景下地理信息科学及相关专业编写的实战型实验教程。全书以“数据密集型科学发现”为范式引领,紧扣时空大数据“获取—处理—分析—可视化—应用”的全生命周期主线,系统讲解时空地理数据处理、激光雷达(LiDAR)点云建模、空间统计学及 GeoAI 深度学习等前沿技术。

不同于传统教材对商业软件操作的按部就班,本书不仅注重代码层面的实现,更强调“地学逻辑”与“计算思维”的深度耦合。书中精选了黄河流域生态监测、古建筑数字化保护等源自真实科研一线的复杂工程案例,引导读者跳出“完美数据”的温室,在处理真实世界的噪声与缺失中掌握解决问题的核心能力。本书既可作为地理科学、地理信息科学、测绘工程等专业高年级本科生及研究生的实验教材,也可作为行业技术人员从传统 GIS 向时空智能转型的参考读物。

版权所有,侵权必究。举报:010-62782989, beiqinquan@tup.tsinghua.edu.cn。

图书在版编目(CIP)数据

时空大数据分析实验教程 / 李晶等编著. -- 北京:清华大学出版社, 2026. 9.
(全国高校地理信息科学实验教学系列丛书). -- ISBN 978-7-302-72769-9
I. P208.2
中国国家版本馆 CIP 数据核字第 2026S17T68 号

责任编辑:秦 娜 王 华

封面设计:谢晓翠

责任校对:王淑云

责任印制:宋 林

出版发行:清华大学出版社

网 址: <https://www.tup.com.cn>, <https://www.wqxuetang.com>

地 址:北京清华大学学研大厦 A 座 邮 编:100084

社 总 机:010-83470000 邮 购:010-62786544

投稿与读者服务:010-62776969, c-service@tup.tsinghua.edu.cn

质量反馈:010-62772015, zhiliang@tup.tsinghua.edu.cn

印 装 者:三河市天利华印刷装订有限公司

经 销:全国新华书店

开 本:185mm×260mm 印 张:15.75 字 数:379 千字

版 次:2026 年 9 月第 1 版 印 次:2026 年 9 月第 1 次印刷

定 价:49.00 元

产品编号:115988-01

全国高校地理信息科学实验教学系列丛书

编 委 会

顾 问 汤国安 宋关福

主 任 张书亮 刘耀林

副主任 熊礼阳 李 晶 程昌秀 吴 浩 王 春

洪 亮 任 福

编 委 (按姓氏拼音排序)

曹 凯 车 健 陈清祥 陈 颖 陈 勇

程 峰 邓 敏 杜世宏 段福洲 高照忠

谷 雷 郭一珂 何 宽 胡楚丽 黄丙湖

黄骁力 江 岭 黎 华 李柏延 李满春

李 苗 李思进 李 鑫 李 颖 林安琪

刘德儿 刘 涛 刘宪锋 罗国玮 梅 新

孟祥锐 明冬萍 年雁云 牛继强 潘梅娥

秦 昆 邵怀勇 宋立生 孙亚琴 万 幼

王 佳 王世东 位 宏 吴孟泉 吴田军

吴志峰 喜文飞 夏既胜 肖 佳 徐艳杰

杨建宇 姚晓军 张 虎 张韶华 张天媛

张王菲 张鲜化 张晓祥 赵建军 赵江洪

赵相伟 赵耀龙 郑光辉 郑江华 周 旭

周自翔



丛书序

FOREWORD

值此“全国高校地理信息科学实验教学系列丛书”首批成果付梓之际,作为地理信息系统(GIS)教育领域从业者、本套丛书编委会顾问,我深感欣慰与振奋。

地理信息科学是支撑国土空间规划、智慧城市建设、时空大数据应用等国家战略领域的核心基础学科,其学科属性决定了实验教学在人才培养中的核心地位。实验教学是衔接课堂理论与行业实操的关键纽带,是学生将抽象原理转化为解决问题能力的必经之路,直接决定着地理信息专业人才的培养质量。数十年来,我国 GIS 教育事业蓬勃发展,专业布点持续增多,人才培养规模稳步提升,但在实践教学环节,始终存在教材体系碎片化、教学标准不统一、实操内容与产业需求衔接不紧密、国产 GIS 技术融入不足等共性痛点,成为专业人才培养质量进一步提升的瓶颈。

近年来,随着云计算、大数据、人工智能与三维可视化等技术的深度融入,GIS 产业正加速向地理空间智能方向演进;与此同时,信创战略全面推进,国产 GIS 平台已成为行业应用的主流选择,行业对人才的实践能力、技术适配性与创新素养提出了全新要求。新时代的 GIS 教育,急需一套体系完整、标准统一、贴合产业、适配国产技术的实验教学资源,为全国高校提供规范化的实践教学参照,推动 GIS 实践教学从“各自探索”走向“规范优质”。

正是基于这样的行业共识与教育使命,在 2024 年 12 月举办的首届 GIS 教育大会上,丛书编委会正式成立,全面启动本套丛书的编撰工作。在 2025 年 10 月举办的第二届 GIS 教育大会上,编委会完成与教材配套软件 SuperMap GIS 的资源深度对接,同步完善组织架构,聘任多位全国高校名师组织分册的编写,为丛书编撰筑牢了专业根基。历经一年多的系统筹备与深耕打磨,依托全国数十所高校的优势教学资源与行业头部企业的技术力量,首批教材终于与广大读者见面。这既是全国 GIS 教育界协同攻关的成果,也是产学研深度融合助力人才培养的有益探索,更标志着我国系统化的地理空间智能实践教学体系落地迈出了关键一步。

本套丛书首批推出《地理信息系统原理实验教程》《空间数据库实验教程》《GIS 空间分析实验教程》《时空大数据分析实验教程》《三维 GIS 实验教程》《遥感地学分析与应用实验教程》,全面覆盖 GIS 基础原理、数据管理、空间分析、智能挖掘、三维建模、遥感应用等核心教学模块,构建了从入门基础到前沿应用、从方法原理到工程实践的完整实验教学体系。丛书的编写体例与逻辑均基于以国产 GIS 软件为核心的实验平台,精选真实行业案例与科研数据,既注重核心原理的拆解阐释,又强化实操能力的系统训练,着力破解“会操作而不明机理”的教学难题,兼具教学适用性与产业适配性。

作为一套面向全国高校的实验教学丛书,其出版填补了国内地理信息科学领域标准化实验教学系列丛书的空白,能够为各高校 GIS 专业的实验教学、课程改革提供权威规范的

教学支撑,更将有力推动全国 GIS 实验教学的均衡化、高质量发展,助力构建产学研用一体化的人才培养体系,为我国地理信息产业持续输送高素质、应用型、创新型专业人才。

丛书编撰过程中,各分册作者凝心聚力、精益求精,开展了近百次线上线下研讨打磨,对实验内容、操作步骤、案例数据反复进行论证。编委会联动出版社严格落实出版规范,完成全流程三审三校,统一全套丛书的目录体系与装帧排版,全方位保障了丛书的专业性、规范性与系统性。

地理信息科学技术迭代迅速、应用边界持续拓展,实验教学资源也须与时俱进、动态更新。未来,编委会将紧密跟踪学科前沿发展与全国高校教学实践反馈,稳步推出更多方向的实验教程,同时对已出版教材进行常态化修订升级,保障丛书始终贴合学科发展节奏与教学实际需求。

教育事业薪火相传,教材建设永无止境。希望这套凝聚了全国 GIS 教育人心血的丛书,能够成为广大师生教学与学习过程中的得力助手,也期待广大读者与同行多提宝贵意见,共同推动丛书不断完善,为我国地理信息科学教育事业的发展添砖加瓦。

汤国安

2026 年夏于南京师范大学



前言

PREFACE

“时空大数据挖掘与分析”是在新工科与新文科建设背景下,地理信息科学及相关专业的前沿核心课程,是推动学生从传统 GIS 应用迈向时空智能研究的关键进阶环节。课程以时空数据全生命周期处理与分析为主线,融合地理学、统计学与计算机科学多学科知识,兼具深厚的地学理论内涵与极强的工程实践属性。它既要求学习者掌握扎实的地理空间认知与地学逻辑,又需具备编程实现、算法设计与复杂问题求解的计算思维,是培养复合型地理信息人才的核心载体。

当前,大数据与人工智能技术深度融入地理学科发展,地理学研究范式正加速向“数据密集型科学发现”演进。从国土空间治理、生态环境监测到文化遗产数字化保护,各领域对时空大数据的深度挖掘与智能分析需求持续攀升,行业与科研端均迫切需要既懂地学规律又能驾驭算法与代码的复合型人才。与此同时,时空数据采集手段日益多元,激光雷达点云、移动轨迹、时序遥感影像等新型数据不断涌现,地理空间人工智能(GeoAI)、深度学习等前沿技术快速落地,行业发展现状对课程实验教学提出了更高要求——既要更新技术体系,更要回归地学本质,将理论方法与真实场景深度结合。

近年来,国内高校在时空大数据相关课程建设与教学改革中已取得积极进展,但在实验教学层面仍存在明显短板。一是多数传统实验教材偏重商业软件的流程化操作,对代码实现与算法原理拆解不足,地学逻辑与计算思维的耦合度不高;二是实验数据多为经过理想化处理的教学样本,与真实科研、工程场景中的噪声、缺失、异构数据存在差距,学生解决复杂实际问题的能力难以充分锻炼;三是前沿技术融入的系统性不足,缺乏覆盖数据全生命周期、兼顾基础方法与前沿方向的完整实验体系。

本书由多所高校 GIS 专业一线教师联合编写,依托团队十余年教学积累与科研实践经验,秉持“以地学逻辑为内核、以代码实现为路径、以真实案例为载体、以能力进阶为目标”的建设思路,旨在搭建地理科学原理与计算机技术之间的衔接桥梁,引导学生在代码实践中深化地学认知,在真实场景中掌握地学分析方法。

全书以“获取—处理—分析—可视化—应用”的时空大数据全生命周期为主线,系统涵盖时空地理数据处理、激光雷达点云建模、空间统计学分析及 GeoAI、深度学习等核心技术内容。教材突破传统实验的理想化数据模式,精选黄河流域生态监测、古建筑数字化保护等源自科研一线的真实工程案例,采用含噪声、缺失值的原始实测数据开展实验设计,引导学生在数据清洗、校正与分析的完整流程中,掌握真实场景下问题求解的核心能力。每个实验都兼顾代码实现与原理解读,注重强化地学逻辑与计算思维的深度融合,同时将家国情怀融入案例教学,在专业实践中传递生态保护、文化遗产的价值导向。

本书既可作为地理科学、地理信息科学、测绘工程等专业高年级本科生及研究生的实验教材,也可作为行业技术人员从传统 GIS 向时空智能转型的参考读物。编写过程中,教材

内容经多轮教学试用与迭代优化,吸纳了大量师生的反馈意见,多位研究生参与了文稿校正与实验验证工作,力求保障教材内容的科学性、可操作性与教学适用性。

希望本书的出版,能够为地理信息科学专业的前沿实验教学提供一套兼具深度与实战性的参考资源,助力学生构建完整的时空大数据分析能力体系,成长为适应时代需求的复合型地理信息人才。本书编写过程中,得到了陕西师范大学、西安科技大学、西南石油大学和西安云图信息技术有限公司广大师生的支持与协助,他们分别是:现任教师李晶、李柏延、李继园、梁伟、刘宪锋、党辉、高余婷、彭国强、杨喜平、曾莉、张学宝、周自翔,在读学生李晓梅、李悦宁、刘妍、刘子怡、沈齐靖、王位宁、吴馨婷、杨靖伦、尹艳红、岳伊文、钟赛英、张盼(按首字母排序)。同时编者参考了国内外相关领域的学术成果,在此一并表示衷心感谢。

受限于技术迭代速度与编者水平,书中难免存在疏漏与不足之处,恳请广大读者与同行专家批评指正,共同推动教学内容的持续完善。

作者

2026 年 4 月



目 录

CONTENTS

第 1 章	长时序植被覆盖时空演变分析	1
1.1	长时序植被覆盖时空演变理论与技术	1
1.1.1	长时序植被覆盖时空演变理论	1
1.1.2	长时序植被覆盖时空演变技术应用	3
1.2	长时序植被覆盖时空演变分析实验	5
1.2.1	植被覆盖的空间分布特征分析	5
1.2.2	植被覆盖的时间变化特征分析	12
1.2.3	植被覆盖的时间变化趋势分析	13
1.2.4	植被覆盖的可持续性分析	16
1.3	长时序植被覆盖时空演变分析实验思考与探索	21
第 2 章	生态系统服务	23
2.1	生态系统服务评估理论与技术	23
2.1.1	生态系统服务评估方法	23
2.1.2	生态系统服务评估技术应用	24
2.2	生态系统服务实验	26
2.2.1	遥感影像预处理	26
2.2.2	碳储量的计算	32
2.2.3	土地利用模拟	34
2.2.4	模拟精度验证	37
2.3	生态服务系统实验思考与探索	38
第 3 章	干旱事件识别和分析	40
3.1	干旱事件识别和分析理论与技术	40
3.1.1	干旱事件识别理论	40
3.1.2	干旱事件识别技术应用	45
3.2	干旱事件识别和分析实验	47
3.2.1	干旱事件的识别	47
3.2.2	干旱事件识别结果分析	51
3.2.3	投影转换、矢量化与面积统计	56

3.2.4 干旱事件时间趋势与热点分析	57
3.3 干旱事件识别和分析实验思考与探索	63
第4章 社区生活圈可步行性分析	65
4.1 社区生活圈可步行性分析理论与技术	65
4.1.1 社区生活圈可步行性分析理论	65
4.1.2 社区生活圈可步行性分析技术应用	68
4.2 社区生活圈可步行性分析实验	70
4.2.1 网络数据集构建	70
4.2.2 社区生活圈识别	74
4.2.3 最近设施查找	76
4.2.4 基础步行指数计算	79
4.2.5 步行指数修正	83
4.2.6 可视化结果	86
4.3 社区生活圈可步行性分析实验思考与探索	87
第5章 城市交通流和轨迹分析	89
5.1 城市交通流和轨迹分析理论与技术	89
5.1.1 城市交通流和轨迹分析理论	89
5.1.2 城市交通流和轨迹分析技术应用	90
5.2 城市交通流和轨迹分析实验	93
5.2.1 渔网划分	93
5.2.2 流量计算	103
5.2.3 空间自相关	105
5.2.4 OD通勤线的显示	106
5.2.5 社区探测	109
5.3 城市交通流和轨迹分析实验思考与探索	116
第6章 北斗/GNSS 卫星导航定位分析	118
6.1 北斗/GNSS 卫星导航定位理论与技术	118
6.1.1 北斗/GNSS 卫星导航定位理论	118
6.1.2 北斗/GNSS 卫星导航定位关键技术及模型	120
6.1.3 北斗/GNSS 数据评估指标	122
6.2 北斗/GNSS 卫星导航定位综合实验	124
6.2.1 基于 RTKLIB 的 SPP 伪距单点定位实验	124
6.2.2 北斗/GNSS 数据质量分析	133
6.2.3 实时 GNSS 数据导入 GIS 软件实验	141
6.3 北斗/GNSS 卫星导航实验思考与探索	144

第 7 章 激光点云建模历史文物·····	147
7.1 激光点云建模历史文物理论与技术 ·····	147
7.1.1 激光点云建模历史文物理论·····	147
7.1.2 激光点云建模历史文物技术应用·····	148
7.2 激光点云建模历史文物实验 ·····	150
7.2.1 文物单体人工建模与纹理贴图·····	150
7.2.2 多源数据集成展示·····	157
7.3 激光点云建模历史文物实验思考与探索 ·····	168
第 8 章 三维场景分析·····	170
8.1 三维场景分析理论与技术 ·····	170
8.1.1 三维场景分析理论·····	170
8.1.2 三维场景分析技术应用·····	171
8.2 三维场景分析实验 ·····	173
8.2.1 时序性三维淹没分析·····	173
8.2.2 洪水淹没分析·····	187
8.3 三维场景分析实验思考与探索 ·····	194
第 9 章 基于 GeoAI 融合遥感深度特征与 POI 的城市功能区识别 ·····	197
9.1 城市功能区识别理论与技术 ·····	197
9.1.1 城市功能区识别理论·····	197
9.1.2 城市功能区识别技术应用·····	199
9.2 城市功能区识别实验 ·····	202
9.2.1 城市功能区识别实验环境搭建·····	202
9.2.2 城市功能区识别数据获取与预处理·····	202
9.2.3 多源数据特征提取·····	205
9.2.4 空间增强 K -Means++ 聚类 ·····	211
9.2.5 无监督结果验证·····	213
9.3 城市功能区识别实验思考与探索 ·····	216
第 10 章 水保智脑平台 ·····	218
10.1 水保智脑平台理论基础 ·····	218
10.2 水保智脑平台实验 ·····	220
10.3 水保智脑平台实验思考与探索 ·····	231
参考文献·····	237



第1章

长时序植被覆盖时空演变分析

1.1 长时序植被覆盖时空演变理论与技术

1.1.1 长时序植被覆盖时空演变理论

归一化植被指数(normalized difference vegetation index, NDVI)是利用遥感技术反演植被覆盖状况的重要指标,基于植被在红光波段和近红外波段的光谱反射特性差异构建。该指数能够有效反映地表植被的生长状况、覆盖密度和生物量水平,是监测生态系统动态、评估植被生产力及研究全球变化生态学的核心参数。通过对 NDVI 时间序列数据的分析,可以揭示植被覆盖的时空演变规律,评估区域生态环境的变化趋势与可持续性。本章采用时序趋势分析方法,结合长程依赖性分析,构建完整的植被动态评估体系,为区域生态环境管理与决策提供科学依据。

1. NDVI 计算模型

NDVI 通过近红外波段(NIR)和红光波段(red)的反射率计算得到,计算公式为

$$NDVI = \frac{\rho_{NIR} - \rho_{red}}{\rho_{NIR} + \rho_{red}} \quad (1-1)$$

式中, ρ_{NIR} 为近红外波段反射率; ρ_{red} 为红光波段反射率。NDVI 数值范围为 $[-1, 1]$, 通常将 $NDVI > 0.1$ 的区域定义为植被覆盖区,用于后续分析。

2. Theil-Sen Median 趋势分析和曼-肯德尔检验

Theil-Sen Median 趋势分析和曼-肯德尔(Mann-Kendall)检验方法能够很好地结合起来,成为判断长时间序列数据趋势的重要方法,并且已经逐渐应用到植被长时间序列分析中。该方法的优点是不需要数据服从一定的分布,对数据误差具有较强的抵抗能力,对于显著性水平的检验具有较为坚实的统计学理论基础,使得结果较为科学和可信。其中,Theil-Sen Median 趋势分析是一种稳健的非参数统计的趋势计算方法,可以有效地减少数据异常值的影响。Theil-Sen Median 趋势计算 $\frac{n(n-1)}{2}$ 个数据组合的斜率的中位数,其计算公式为

$$S_{NDVI} = \text{Median}\left(\frac{NDVI_j - NDVI_i}{j - i}\right), \quad 2000 \leq i < j \leq 2025 \quad (1-2)$$

当 $S_{\text{NDVI}} > 0$ 时,反映 NDVI 呈现增长的趋势,当 $S_{\text{NDVI}} \leq 0$ 时,反映 NDVI 呈现退化的趋势。

Mann-Kendall 检验是一种非参数统计检验方法,用来判断趋势的显著性,它无须样本服从一定的分布,也不受少数异常值的干扰。计算公式如下:

设定 2000—2025 年 NDVI 时间序列为 $\{\text{NDVI } I_{t_k}\}$,其中序列索引 $k=1,2,\dots,n$, n 为时间序列样本总数; $k=1$ 对应 2000 年, $k=n$ 对应 2025 年。Mann-Kendall 检验标准化统计量 Z 定义为

$$Z = \begin{cases} \frac{S-1}{\sqrt{\text{Var}(S)}}, & S > 0 \\ 0, & S = 0 \\ \frac{S+1}{\sqrt{\text{Var}(S)}}, & S < 0 \end{cases} \quad (1-3)$$

其中,Mann-Kendall 秩序统计量 S 为

$$S = \sum_{j=1}^{n-1} \sum_{i=j+1}^n \text{sgn}(\text{NDVI}_{t_i} - \text{NDVI}_{t_j}) \quad (1-4)$$

符号函数 $\text{sgn}(\cdot)$ 定义为

$$\text{sgn}(\text{NDVI}_{t_i} - \text{NDVI}_{t_j}) = \begin{cases} 1, & \text{NDVI}_{t_i} - \text{NDVI}_{t_j} > 0 \\ 0, & \text{NDVI}_{t_i} - \text{NDVI}_{t_j} = 0 \\ -1, & \text{NDVI}_{t_i} - \text{NDVI}_{t_j} < 0 \end{cases} \quad (1-5)$$

秩序统计量 S 的方差 $\text{Var}(S)$ 计算如下:

$$\text{Var}(S) = \frac{n(n-1)(2n+5)}{18} \quad (1-6)$$

式中: Z 为 Mann-Kendall 检验标准化统计量; S 为 Mann-Kendall 秩序统计量; n 为 NDVI 时间序列样本总数,本研究中为 2000—2025 年总年份数; i, j 为时间序列数组索引,并非真实年份, j 代表时序靠前样本, i 代表时序靠后样本,满足 $i > j$,对应真实年份 $t_i > t_j$; t_i, t_j 分别为索引 i, j 对应的真实年份; $\text{sgn}(\cdot)$ 为符号函数; $\text{Var}(S)$ 为秩序统计量 S 的方差。统计量 Z 的取值范围为 $(-\infty, +\infty)$ 。在给定显著性水平 α 下,当 $|Z| > Z_{1-\alpha/2}$ 时,表明研究序列在 α 水平上存在显著变化趋势;当 $|Z| \leq Z_{1-\alpha/2}$ 时,表明不存在显著变化趋势。

3. Hurst 指数

Hurst 指数(Hurst exponent)是描述时间序列长期依赖性的重要方法,最早由英国水文学家 Hurst 提出。目前,该方法已广泛应用于水文学、经济学、气候学、地质学以及生态环境变化研究等领域。Hurst 指数主要通过重标度极差分析法(rescaled range analysis, R/S 分析)进行计算,其基本原理如下。

对于 NDVI 时间序列: $\{\text{NDVI}(t)\}$, $t=1,2,\dots,n$,其中, n 为时间序列样本总数。定义子序列长度为 τ ,其中: $2 \leq \tau \leq n$ 。

选取前 τ 个样本构成子序列,其平均值为

$$\overline{\text{NDVI}}_{(\tau)} = \frac{1}{\tau} \sum_{t=1}^{\tau} \text{NDVI}_{(t)} \quad (1-7)$$

累积离差定义为

$$X_{(t,\tau)} = \sum_{i=1}^t (\text{NDVI}_{(i)} - \overline{\text{NDVI}}_{(\tau)}), \quad 1 \leq t \leq \tau \quad (1-8)$$

子序列极差为

$$R_{(\tau)} = \max_{1 \leq t \leq \tau} X_{(t,\tau)} - \min_{1 \leq t \leq \tau} X_{(t,\tau)} \quad (1-9)$$

子序列标准差为

$$S_{(\tau)} = \left[\frac{1}{\tau} \sum_{t=1}^{\tau} (\text{NDVI}_{(t)} - \overline{\text{NDVI}}_{(\tau)})^2 \right]^{\frac{1}{2}} \quad (1-10)$$

若满足: $\frac{S(\tau)}{R(\tau)} \propto \tau^H$ 则说明该时间序列存在 Hurst 现象, 其中 H 即为 Hurst 指数。

对上述公式两边取常用对数, 可得到线性关系:

$$\lg \left[\frac{S(\tau)}{R(\tau)} \right] = \alpha + H \lg(\tau)$$

其中, H 可通过对 $\lg(\tau)$ 与 $\lg \left[\frac{S(\tau)}{R(\tau)} \right]$ 进行最小二乘线性拟合得到。

式中, n 为 NDVI 时间序列样本总数; t 为时间序列内部样本序号; τ 为子序列长度; $\text{NDVI}_{(t)}$ 为第 t 个样本对应的 NDVI 值; $\overline{\text{NDVI}}_{(\tau)}$: 长度为 τ 的子序列 NDVI 平均值; $X_{(t,\tau)}$ 为累积离差; $R(\tau)$ 为子序列极差; $S(\tau)$ 为子序列标准差; α 为线性拟合截距; H 为 Hurst 指数。

Hurst 指数取值范围为: $0 < H < 1$, 其含义如下: $0 < H < 0.5$ 时, 时间序列表现为反持续性, 未来变化趋势倾向于与过去变化趋势相反; $H = 0.5$ 时, 时间序列属于随机过程, 不存在明显的长期相关性; $0.5 < H < 1$ 时, 时间序列表现为持续性, 未来变化趋势倾向于延续过去变化规律, 具有明显的长期记忆特征。

1.1.2 长时序植被覆盖时空演变技术应用

1. 实验概述

本实验以植被覆盖时空变化分析为例, 介绍基于遥感数据的生态系统动态监测方法, 采用时序趋势分析方法, 结合长程依赖性分析, 构建完整的植被动态评估体系, 为区域生态环境管理与决策提供科学依据。

实验首先基于已有 NDVI 数据集提取研究区 NDVI 数据, 通过空间可视化展示植被分布格局; 随后构建年际 NDVI 变化序列, 采用 Theil-Sen Median 趋势分析和 Mann-Kendall 检验方法, 定量识别植被覆盖的变化趋势与空间分异特征; 在此基础上, 引入 Hurst 指数分析植被变化的长程依赖性, 评估未来变化趋势的持续性; 最后综合趋势分析与持续性评估结果, 对区域植被生态系统的可持续发展状态进行综合评价。

NDVI 时序分析可广泛应用于生态环境监测、植被动态评估、土地利用/覆被变化研究、生态系统服务功能评价、荒漠化监测与评估、农业生产力估算等领域, 是区域生态安全评估与可持续发展规划的重要决策支持工具。该方法体系同样适用于其他遥感指数(如 EVI 等)的时序分析, 为多维度生态环境监测提供通用技术框架。

2. 生态文明学思践行

本实验将思想教育有机融入专业教学全过程,通过 NDVI 植被动态分析实践,实现知识传授、能力培养与价值引领的统一。

首先,实验内容与国家生态文明建设战略紧密对接,引导学生树立“绿水青山就是金山银山”的绿色发展理念。通过对植被覆盖时空变化进行科学监测与分析,学生将深刻理解我国在生态环境保护、碳中和目标实现等方面的重大战略部署,增强对国家可持续发展道路的认同感和使命感,培养生态文明建设的责任意识。

其次,实验设计要注重系统思维与全局观念的培育。植被变化分析涉及自然、社会、经济等多重因素,要求学生运用系统思维理解生态环境问题的复杂性。通过分析植被变化与人类活动的关联,引导学生思考人与自然和谐共生的深刻内涵,培养辩证唯物主义世界观和整体性思维方法。

最后,在实验结果评价与反思环节,注重伦理意识与职业素养的塑造。学生需对分析结果进行真实性、可靠性评估,思考数据伦理在科学研究和决策支持中的重要性。通过反思技术应用的社会影响,培养学生以负责任的态度运用专业技术服务社会发展的职业伦理,强化科技为民的价值导向。

本实验通过“理论—方法—实践—反思”的全过程融入,使学生在掌握专业技术的同时,深化对国情社情的认识,增强服务国家生态安全战略的自觉性,实现专业技能提升与思想政治教育同向同行。

3. 实验目的

(1) 掌握基于时序遥感数据的 NDVI 动态分析基本原理与方法,包括 Theil-Sen Median 趋势分析、Mann-Kendall 检验及 Hurst 指数长程依赖性分析。

(2) 综合应用空间可视化、时间序列建模与趋势-持续性耦合分析方法,实现区域植被覆盖变化趋势的定量评估与未来变化方向的预测。

(3) 培养运用多源空间数据分析技术解决实际生态环境问题的能力,提升对区域生态系统变化过程的理解与科学评估水平。

4. 实验数据

本实验使用的数据存储于 data_expl 文件夹中,包含的数据内容及来源见表 1-1。



实验数据

表 1-1 实验数据内容及其来源

数据类型	数据来源	数据内容	数据说明
MODIS NDVI	USGS	研究区 2000—2025 年 NDVI 栅格数据	栅格数据,*.tif 格式
基础地理数据		研究区行政边界	矢量格式,线、面数据

5. 整体实验设计

本次实验旨在深入探讨研究区内植被覆盖区域的 NDVI 特性,具体将围绕以下 4 个核

心问题展开：①研究区植被覆盖区域 NDVI 的空间分布特征分析；②研究区植被覆盖区域 NDVI 的时间变化特征分析；③研究区植被覆盖区域 NDVI 的变化趋势分析；④研究区植被覆盖区域 NDVI 的可持续性分析。

本实验的整体流程如图 1-1 所示。

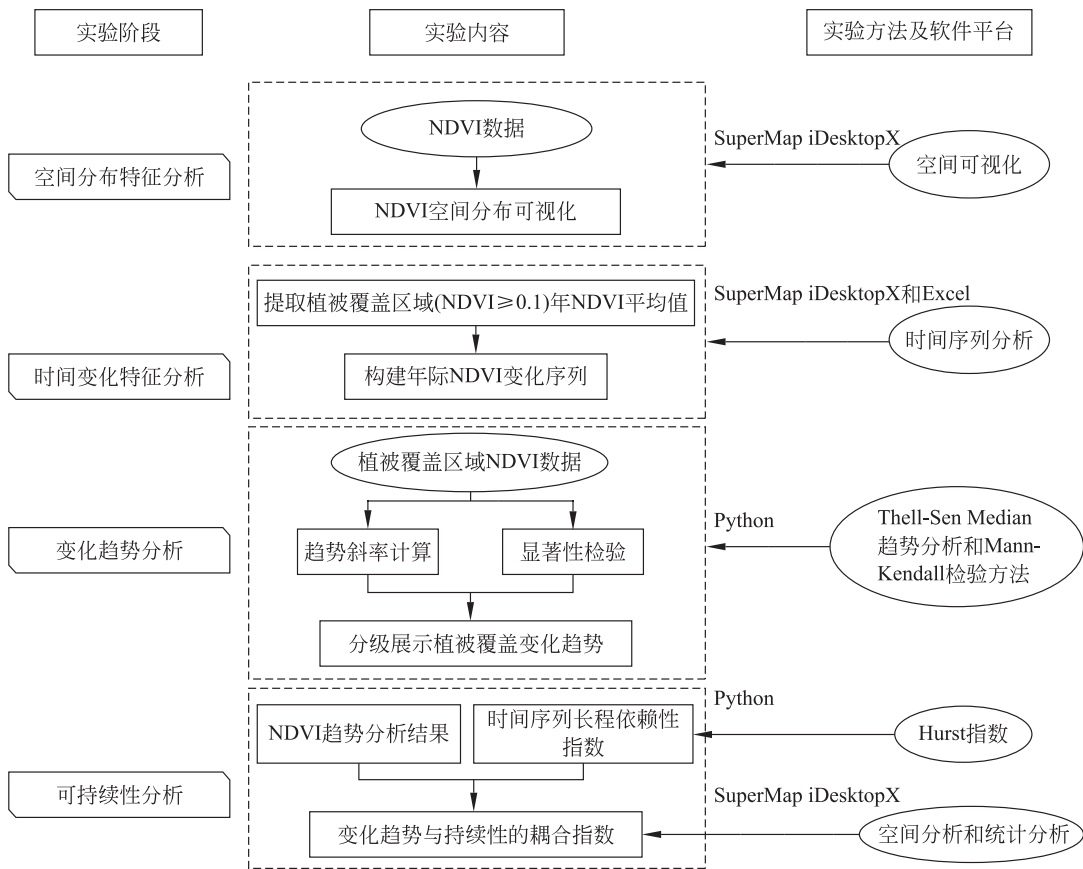


图 1-1 实验流程图

1.2 长时序植被覆盖时空演变分析实验

1.2.1 植被覆盖的空间分布特征分析

1. 栅格代数运算计算平均 NDVI

1) 导入数据集

运行 SuperMap iDesktopX, 右击【工作空间管理器】中的【数据源】, 选择【新建文件型数据源...】, 保存类型选择 SuperMap UDB 文件(*.udb)或 SuperMap UDBX 文件(*.udbx)。右击新创建的数据源, 选择【导入数据集...】, 在【数据导入】对话框中, 单击【添加文件】, 选择 2000—2025 年 NDVI 数据, 数据集类型选择栅格, 导入 NDVI 栅格数据

集(图 1-2)。

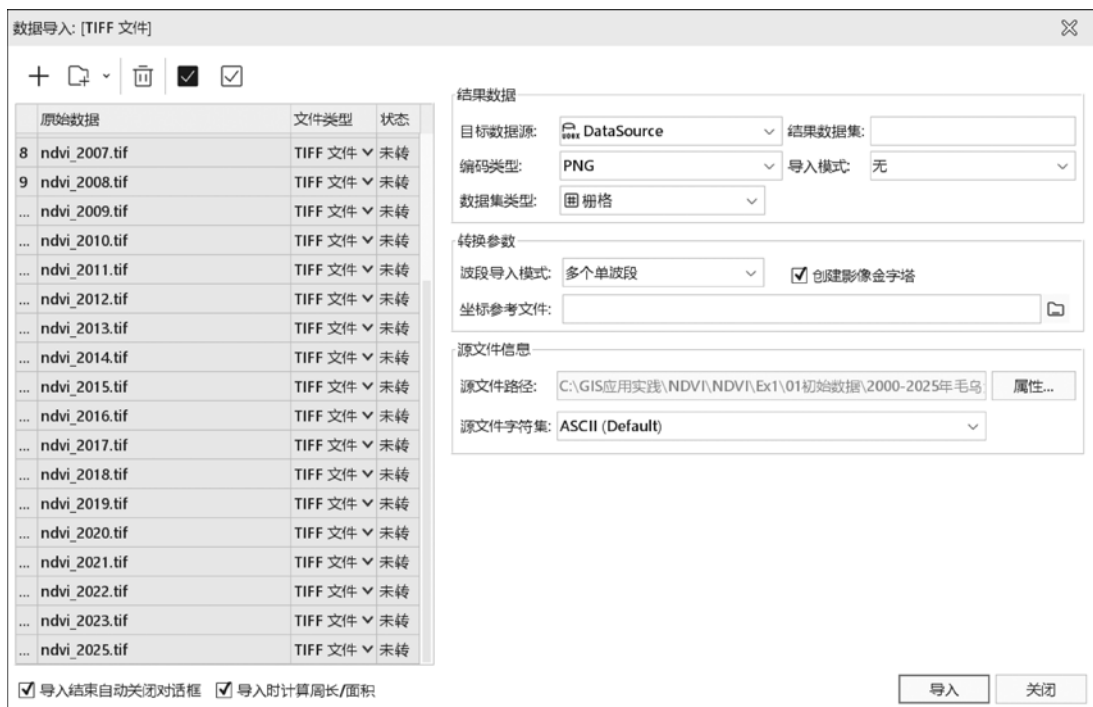


图 1-2 导入 NDVI 栅格数据集

2) 计算平均 NDVI

在 SuperMap iDesktopX 功能选项卡【数据】中定位到【数据处理】组,单击【栅格】栏目中的【代数运算】工具,在弹出的【代数运算】对话框中,找到代数运算表达式输入区域。输入表达式:

$$\begin{aligned}
 & ([DataSource.ndvi_{2000}] + [DataSource.ndvi_{2001}] + [DataSource.ndvi_{2002}] + \\
 & [DataSource.ndvi_{2003}] + [DataSource.ndvi_{2004}] + \\
 & [DataSource.ndvi_{2005}] + [DataSource.ndvi_{2006}] + \\
 & [DataSource.ndvi_{2007}] + [DataSource.ndvi_{2008}] + \\
 & [DataSource.ndvi_{2009}] + [DataSource.ndvi_{2010}] + \\
 & [DataSource.ndvi_{2011}] + [DataSource.ndvi_{2012}] + \\
 & [DataSource.ndvi_{2013}] + [DataSource.ndvi_{2014}] + \\
 & [DataSource.ndvi_{2015}] + [DataSource.ndvi_{2016}] + \\
 & [DataSource.ndvi_{2017}] + [DataSource.ndvi_{2018}] + \\
 & [DataSource.ndvi_{2019}] + [DataSource.ndvi_{2020}] + \\
 & [DataSource.ndvi_{2021}] + [DataSource.ndvi_{2022}] + \\
 & [DataSource.ndvi_{2023}] + [DataSource.ndvi_{2024}] + \\
 & [DataSource.ndvi_{2025}]) / 26
 \end{aligned} \quad (1-11)$$

在【参数设置】区域,设置像素格式为“双精度浮点型”(图 1-3)。在结果数据区域,选择目标数据源,并输入结果数据集名称,如“mean_NDVI”。确认表达式和参数设置无误后,单击运行按钮。运行成功后,结果数据集“mean_NDVI”将存储在指定的目标数据源中。



图 1-3 栅格代数运算参数设置界面

3) 制作 NDVI 空间分布图

(1) 绘制专题图:在工作空间管理器栏定位到存储结果数据“mean_NDVI”的数据源,右击“mean_NDVI”数据集,选择【添加到新地图】,目标数据就被添加到图层管理器里。

在图层管理器中选中“mean_NDVI”图层,右击选择【制作专题图】,在弹出的【制作专题图】窗口左侧栏目中选择“栅格分段专题图”并单击【确定】,专题图创建后,会自动弹出【专题图】窗口,其中显示了当前专题图的默认设置。在【专题图】窗口的【属性】界面中,NDVI 制图的分段方法常用等距分段法,也可以按需要选择平方根分段法、对数分段法等。在【颜色方案】下拉列表中选择适用于 NDVI 的配色方案,在【适用 NDVI】分组中可直接选择适用于植被指数的颜色方案(图 1-4)。



图 1-4 年均 NDVI 空间分布显示



彩图 1-4

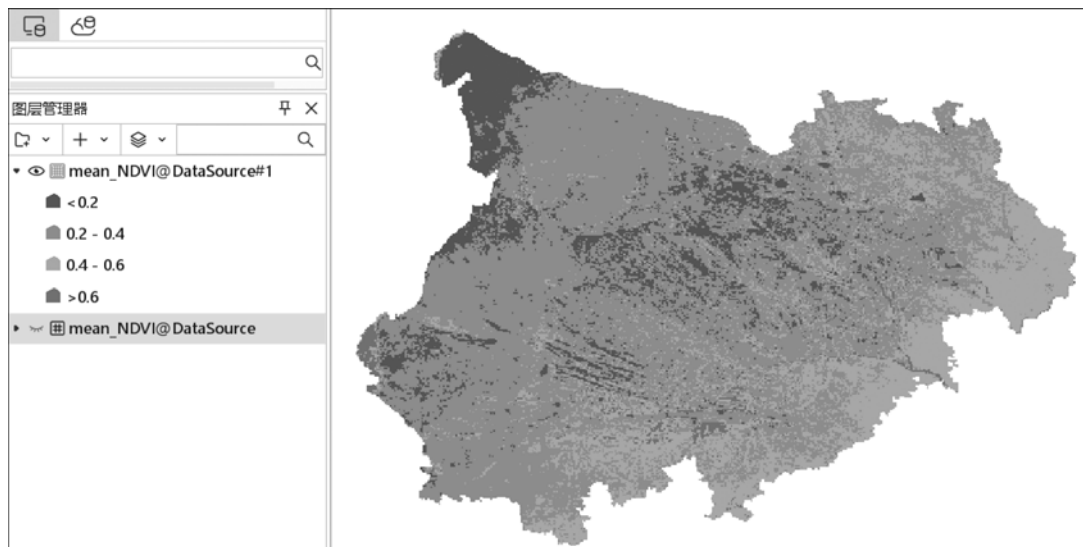


图 1-4(续)



图 1-5 新建布局

(2) 在地图窗口中确认地图渲染效果无误后,右击地图窗口选择**【保存地图】**。在弹出的**【保存地图】**对话框中,输入地图名称,例如“毛乌素 2000—2025 年平均 NDVI 空间分布”,单击**【确定】**保存。在工作空间管理器中,右击**【布局】**,选择**【新建布局】**(图 1-5);在**【新建布局】**对话框中,可以选择布局模板,或选择空白模板自行设计。在顶部功能选项卡**【布局】**中,根据地图形状选择合适方向;按出图需求选择 A4、A3 或自定义尺寸的纸张大小;设置页面上下边距的垂直和水平距离,确保地图在版面中位置适中。

确保当前活动窗口为布局窗口,单击顶部功能区**【对象操作】→【对象绘制】→【地图】**,鼠标指针变为绘制状态。在布局页面的合适位置,单击并拖曳鼠标,绘制一个矩形地图框。绘制完成后,在**【选择地图】**下拉列表中选择已保存的“毛乌素 2000—2025 年平均 NDVI 空间分布”地图,单击**【确定】**完成地图填充。在地图框上

右击,选择**【属性】**,在弹出的**【布局对象属性】**面板可设置地图框的边框样式(如边框类型、线条颜色、线条粗细等)。SuperMap iDesktopX 提供三种边框类型:无边框、单线边框、复杂边框。

若需调整地图框内的显示范围,选中地图对象后右击,选择**【锁定地图】**,此时可使用鼠标滚轮缩放、拖曳平移等操作,调整地图至合适的显示区域;调整完成后,在地图框内右击,再次选择**【锁定地图】**,返回布局编辑状态。

(3) 添加图例:在布局窗口中选中地图对象,单击顶部功能区**【对象操作】→【对象绘**

制】→【图例】，选择合适图例。在布局页面的合适位置，单击并拖曳鼠标绘制图例框，系统即根据地图内容自动生成图例。图例生成后，右击图例，选择【属性】进行以下设置：①图例标题：默认显示为“图例”，可修改为“年均 NDVI 值”等自定义标题；②图例列数：默认显示为 3 列，可根据图例项数量调整；③填充颜色：可为图例设置整体背景填充色；④字体设置：可分别设置图例标题、图例项、图例子项的字体、字号、颜色等风格。

添加指北针：在布局窗口中选中地图对象，单击顶部功能区【对象操作】→【对象绘制】→【指北针】，鼠标指针变为绘制状态。在地图框的合适位置单击即可添加默认样式的指北针。指北针添加后，右击指北针，选择【属性】进行以下设置：①指北针样式：选择不同的指北针样式模板；②大小与旋转：可调整指北针的宽度、高度及旋转角度。

添加比例尺：在布局窗口中选中地图对象，单击顶部功能区【对象操作】→【对象绘制】→【比例尺】。在布局页面中地图框的下方或适当位置，单击并拖曳鼠标绘制比例尺框，系统基于当前地图的显示比例自动生成比例尺。比例尺生成后，右击，选择【属性】进行以下设置：①比例尺类型：根据提供选项选择合适比例尺类型；②刻度设置：可调整比例尺的刻度数、刻度值、频数、高度等参数；③单位设置：可设置比例尺的显示单位；④字体与风格：可设置比例尺标签文字的字体、字号、颜色等风格。

添加地图标题：单击顶部功能区【对象操作】→【对象绘制】→【文本】，在布局页面的上方单击并拖曳鼠标绘制一个文本框，输入地图标题（如“毛乌素 2000—2025 年平均 NDVI 空间分布”），单击文本框外其他区域完成输入。标题添加后，右击选择【属性】进行以下设置：在【文本信息】选项卡中设置字体基本信息和字体效果；在【空间信息】选项卡中设置标题框的位置和大小。

完成上述所有布局元素的添加与设置后，执行输出操作：确保当前活动窗口为布局窗口，单击顶部功能区【布局】→【输出】→【输出图片】，在【输出图片】对话框中设置输出文件名、保存路径、输出分辨率（dpi，默认为 96，值越大图像越清晰，但文件大小也随之增加）。设置完成后单击【确定】，导出图片，如图 1-6 和图 1-7 所示。



图 1-6 输出布局



彩图 1-7

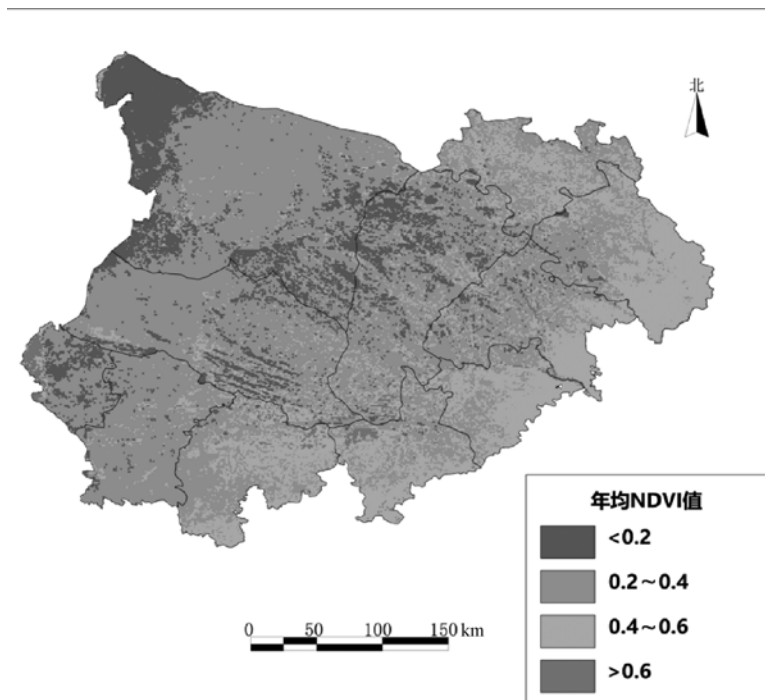


图 1-7 毛乌素 2000—2025 年平均 NDVI 空间分布图

2. NDVI 平均值重分级处理

1) 打开重分级工具

在 SuperMap iDesktopX 功能选项卡中单击【数据】→【数据处理】→【栅格】→【重分级】。

在弹出的【重分级】对话框的【源数据】区域,选择数据源为“DataSource”,数据集为“mean_NDVI”。在【参数设置】区域,选择“范围重分级”方式。单击【设置】,选择级数为 5,单击【确定】后分别设置段值和目标值(图 1-8)。

序号	段值下限	段值上限	目标值
1	0.0	0.1	1
2	0.1	0.4	2
3	0.4	0.5	3
4	0.5	0.6	4
5	0.6	0.8589769235025	

图 1-8 重分级参数设置界面

确认参数设置无误后,单击对话框底部的执行按钮,生成新的栅格数据集“result_Reclass”。

2) 查看重分级结果属性

在【工作空间管理器】中,右击新生成的“result_Reclass”数据集,选择“统计栅格值”。在统计表中可以查看重分级后的统计信息(图 1-9)。

序号	*SmID	SmUserID	GridValue	GridCount
1	1	0	1	49
2	2	0	2	57183
3	3	0	3	21146
4	4	0	4	10834
5	5	0	5	3132

图 1-9 重分级结果统计表

NDVI 平均值的分级统计结果表明: $NDVI < 0.1$ 的无植被覆盖区域占流域总面积的 0.05%,有植被覆盖区域占 99.95%。 $NDVI$ 为 $0.1 \sim 0.4$ 的低植被覆盖区域占 61.92%, $NDVI > 0.4$ 的高植被覆盖区域占 38.03%,其中 $NDVI$ 为 $0.4 \sim 0.5$ (含 0.5) 的区域占 22.90%, $0.5 \sim 0.6$ 的区域占 11.73%。 $NDVI > 0.6$ 的高植被覆盖区域占 3.39% (图 1-10)。对重分级后的 $NDVI$ 数据进行专题图制图显示(图 1-11)。

	A	B	C	D
1 ID	GridValue	GridCount	占比/%	
2	0	1	49	0.05
3	0	2	57183	61.92
4	0	3	21146	22.90
5	0	4	10834	11.73
6	0	5	3132	3.39
7				100.00

图 1-10 NDVI 平均值的分级统计结果

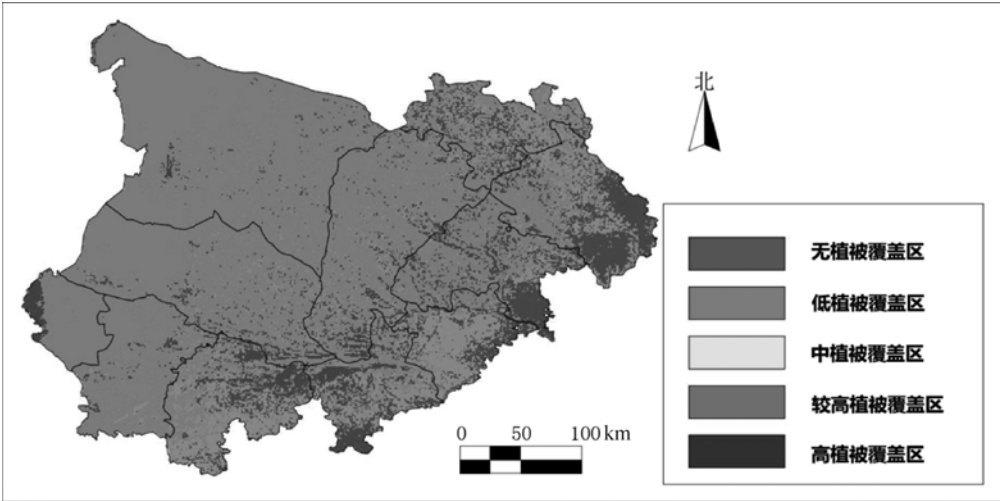


图 1-11 毛乌素 2000—2025 年植被覆盖空间分布图



彩图 1-11

1.2.2 植被覆盖的时间变化特征分析

1. 提取 NDVI 值 ≥ 0.1 区域

在 SuperMap iDesktopX 功能选项卡【数据】中定位到【数据处理】组,单击【栅格】栏目中的【代数运算】工具,在弹出的【代数运算】对话框中,找到代数运算表达式输入区域。输入条件表达式(以 2000 年为例):

$$\text{Con}([\text{DataSource.ndvi_2000}] \geq 0.1, [\text{DataSource.ndvi_2000}], 0) \quad (1-12)$$

设置结果数据集名称为“R_NDVI_2000”。单击执行按钮完成条件提取(图 1-12)。



图 1-12 栅格条件提取设置

2. 统计各年份 NDVI 平均值

在【工作空间管理器】中,右击每个筛选后的数据集(如“R_NDVI_2000”),选择【属性】,查看【统计信息】选项卡中的“平均值”,将各年份的平均值记录到 Excel 表格中(图 1-13)。

3. Excel 中制作年际 NDVI 变化图

1) 准备 Excel 数据表

打开 Microsoft Excel,创建新工作簿。在 Sheet1 中,按图 1-13 格式输入 2000—2025 年的年份和 NDVI 年平均值。

2) 创建年际 NDVI 变化图

选中 A1: B27 单元格区域(包括表头),单击【插入】选项卡,在【图表】组中选择“XY 散点图”。Excel 将自动生成年均 NDVI 变化散点图,设置图表属性,添加趋势线等,以直观展示 2000—2025 年毛乌素沙地年际 NDVI 变化(图 1-14)。

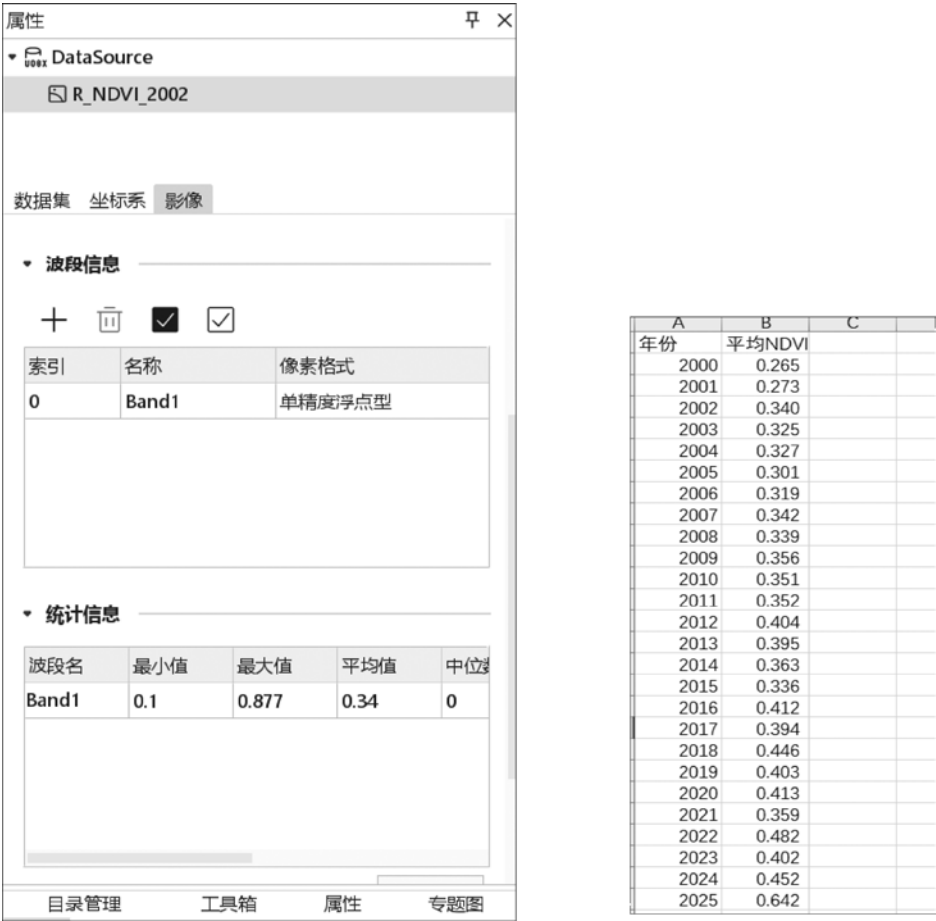


图 1-13 NDVI 年平均值统计

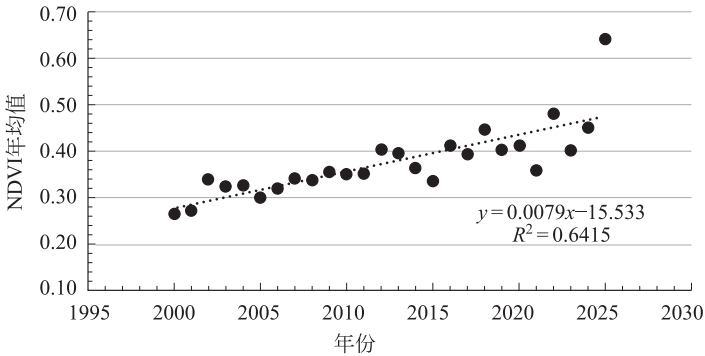


图 1-14 2000—2025 年毛乌素沙地年际 NDVI 变化

1.2.3 植被覆盖的时间变化趋势分析

1. 基于 Python 的 NDVI 时序趋势分析

1) NDVI 时序数据分析准备

确保已准备好 NDVI 时序数据文件，数据应存储在指定目录中，检查数据文件命名规

范：文件应包含年份信息，如 R_NDVI_2000.tif、R_NDVI_2001.tif 等；确认数据文件格式：所有文件均为 GeoTIFF 格式(.tif)，具有相同的空间参考系统和分辨率。

2) 配置代码参数

在 Python 中打开 sen.py, 在代码中找到主函数 main(), 修改关键参数, 如 NDVI 数据目录路径, 结果输出目录路径等(图 1-15)。

```
# 参数设置
tif_dir = "C:/NDVI/NDVI_RECLASS/R_ndvi0.1" # 修改为你的NDVI tif文件目录
output_dir = "C:/NDVI/sen&mk/sen/output" # 输出目录
start_year = 2000
end_year = 2025
fix_orientation = True

total_start_time = time.time()
```

图 1-15 参数设置代码

根据研究区域特点, 可调整分类阈值(图 1-16)。

```
def calculate_trend_classification(sndvi, z_value):
    """
    根据S_NDVI和Z值进行趋势分类

    分类规则:
    1. 明显改善: S_NDVI >= 0.0005 且 Z > 1.96
    2. 轻微改善: S_NDVI >= 0.0005 且 -1.96 < Z < 1.96
    3. 稳定不变: -0.0005 < S_NDVI <= 0.0005 且 -1.96 < Z < 1.96
    4. 轻微退化: S_NDVI < -0.0005 且 -1.96 < Z < 1.96
    5. 严重退化: S_NDVI < -0.0005 且 Z < -1.96

    """
    if np.isnan(sndvi) or np.isnan(z_value):
        return -9999 # 返回无效值标记

    # 明显改善
    if sndvi >= 0.0005 and z_value >= 1.96:
        return 1

    # 轻微改善
    elif sndvi >= 0.0005 and -1.96 < z_value < 1.96:
        return 2

    # 稳定不变
    elif -0.0005 <= sndvi <= 0.0005 and -1.96 <= z_value <= 1.96:
        return 3

    # 轻微退化
    elif sndvi < -0.0005 and -1.96 <= z_value <= 1.96:
        return 4

    # 严重退化
    elif sndvi < -0.0005 and z_value < -1.96:
        return 5
```

图 1-16 分类阈值调整代码

3) NDVI 趋势分析

执行 Python 脚本, 计算完成后, 控制台将显示详细的统计信息(图 1-17)。

```
=== NDVI趋势分类统计 ===
表1 NDVI趋势统计
```

S_NDVI	Z值	NDVI趋势变化	面积百分比/%	像元数
>0.0005	>1.96	明显改善	80.0	73,995
>0.0005	-1.96-1.96	轻微改善	17.5	16,220
-0.0005-0.0005	-1.96-1.96	稳定不变	1.3	1,187
<-0.0005	-1.96-1.96	轻微退化	1.0	959
<-0.0005	<-1.96	严重退化	0.2	145
总计			100.0	92,506

图 1-17 NDVI 趋势分类统计结果

4) 检查输出文件

导航到输出目录,检查生成的输出文件(图 1-18)。



图 1-18 代码输出文件

2. 年均 NDVI 变化趋势可视化

在【工作空间管理器】中,右击“DataSource”数据源,选择【导入数据集...】,在【数据导入】对话框中,单击【添加文件】按钮,找到并选择 Python 脚本生成的“ndvi_trend_classification.tif”文件,数据集类型选择“栅格”,导入 NDVI 栅格数据集。右击导入的数据集,选择【添加到新地图】,地图窗口中将显示 NDVI 趋势分类结果。

在图层管理器中,右击“ndvi_trend_classification”图层,选择【制作专题图】,在弹出的【制作专题图】窗口左侧栏目中选择“栅格单值专题图”并单击【确定】,专题图创建后,会自动弹出【专题图】窗口,根据 NDVI 分类结果设置显示参数和标题(图 1-19)。



图 1-19 专题图显示设置界面

将 Theil-Sen Median 趋势分析和 Mann-Kendall 检验结合起来,可以有效反映 2000—2025 年 NDVI 变化趋势的空间分布特征。由于基本上不存在 S_{NDVI} 严格等于 0 的区域,所以根据实际情况, S_{NDVI} 为 $-0.0005 \sim 0.0005$ 的划分为稳定不变区域, $S_{NDVI} \geq 0.0005$ 的划分为改善区域, $S_{NDVI} < -0.0005$ 的划分为退化区域。将 Mann-Kendall 检验在 0.05 置信水平

上的显著性检验结果划分为显著变化($Z > 1.96$ 或 $Z < -1.96$)和不显著变化($-1.96 \leq Z \leq 1.96$)。将 Theil-Sen Median 趋势分析的分级结果和 Mann-Kendall 检验的分级结果进行叠加,得到像元尺度上 NDVI 变化趋势数据,并将结果划分为 5 种变化类型(图 1-20)。



彩图 1-20

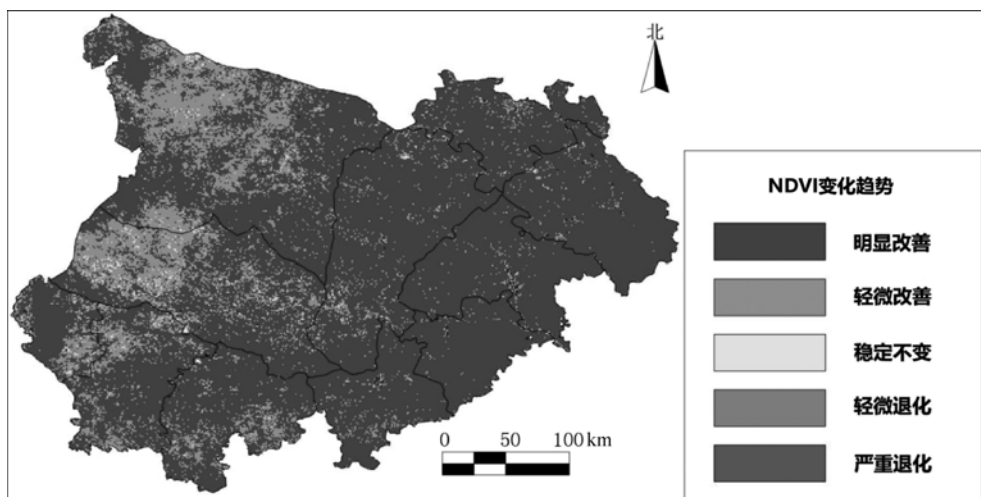


图 1-20 2000—2025 年毛乌素沙地 NDVI 变化趋势

1.2.4 植被覆盖的可持续性分析

1. 基于 Python 的 Hurst 指数计算

1) 准备 NDVI 时序数据文件

检查 NDVI 时序数据文件命名规范: 文件应包含年份信息,如 R_NDVI_2000.tif、R_NDVI_2001.tif 等; 确认数据文件格式: 文件应为 GeoTIFF 格式(.tif),具有相同的空间参考系统和分辨率; 检查数据时间序列的连续性: 建议至少有 10 年的连续数据。

2) 配置代码参数

在 Python 中打开 hurst.py,在代码中找到主函数 main()部分,修改参数如图 1-21 所示。

```
def main():
    print("NDVI Hurst指数计算器 ")
    print("="*60)

    # 获取文件列表
    input_path = 'C:/NDVI/NDVI_RECLASS/R_ndvi0.1/*.tif'
    ndvi_files = sorted(glob.glob(input_path))

    if not ndvi_files:
        print(f"错误: 在 {input_path} 未找到TIFF文件")
        return

    print(f"找到 {len(ndvi_files)} 个NDVI文件")

    # 预览文件
    for i, f in enumerate(ndvi_files[:5]):
        print(f" {os.path.basename(f)}")
    if len(ndvi_files) > 5:
        print(f"... 和 {len(ndvi_files)-5} 个其他文件")

    # 设置输出目录
    timestamp = time.strftime("%Y%m%d_%H%M%S")
    output_dir = f"./hurst_results_{timestamp}"
```

图 1-21 Hurst 指数参数设置代码

如果需要调整分类阈值,可以修改以下函数:

(1) 随机序列阈值调整: 在 analyze 函数中,调整随机序列的阈值区间,当前设置为 0.49~0.51。

(2) 分类阈值调整: 在 analyze 函数中,调整持续性强度分类阈值,当前设置为反持续性: $0 < H < 0.5$; 持续性: $0.5 < H < 1$ (图 1-22)。

```
print(f"\nHurst指数分类:")

# 反持续性 (0 < H < 0.5)
anti_persistent = np.sum((valid_hurst > 0) & (valid_hurst < 0.5))

# 随机序列 (H = 0.5, 实际使用一个小的区间)
random_series = np.sum((valid_hurst >= 0.49) & (valid_hurst <= 0.51))

# 持续性 (0.5 < H < 1)
persistent = np.sum((valid_hurst > 0.5) & (valid_hurst < 1))

# 额外统计: Hurst指数等于0.5的情况 (更严格的区间)
exact_half = np.sum((valid_hurst >= 0.499) & (valid_hurst <= 0.501))

categories = {
    '反持续性 (0 < H < 0.5)': anti_persistent,
    '随机序列 (H = 0.5, 实际区间0.49-0.51)': random_series,
    '持续性 (0.5 < H < 1)': persistent
}

for name, count in categories.items():
    percentage = count/valid_count*100
    print(f"{name}: {count:,} 像元 ({percentage:.1f}%)")
```

图 1-22 分类阈值调整代码

3) 运行 Hurst 指数计算

执行 Python 脚本,计算完成后控制台将显示详细的统计信息(图 1-23)。

```
结果统计:
有效Hurst指数数量: 84,906 (53.2%)
Hurst指数范围: [0.0000, 1.0000]
平均值: 0.7595
中位数: 0.7834
标准差: 0.1540

Hurst指数分类:
反持续性 (0 < H < 0.5): 5,450 像元 (6.4%)
随机序列 (H = 0.5, 实际区间0.49-0.51): 1,213 像元 (1.4%)
持续性 (0.5 < H < 1): 79,437 像元 (93.6%)
精确的H=0.5 (区间0.499-0.501): 131 像元 (0.2%)

Hurst指数强度分析:
反持续性平均H值: 0.4235 (越接近0, 反持续性越强)
持续性平均H值: 0.7826 (越接近1, 持续性越强)
```

图 1-23 Hurst 指数计算结果

4) 检查输出文件

导航到输出目录,检查生成的输出文件(图 1-24)。

```
hurst_results_YYYYMMDD_HHMMSS/    # 输出目录, 包含时间戳
├─ hurst_index.tif                  # Hurst指数栅格图 (浮点型)
├─ hurst_classification.tif         # Hurst分类栅格图 (整型)
└─ statistics.txt                   # 详细统计报告
```

图 1-24 Hurst 指数代码输出文件

2. Hurst 指数结果可视化

在 SuperMap iDesktopX 软件的【工作空间管理器】中,右击“DataSource”数据源,选择

【导入数据集...】，在【数据导入】对话框中，单击【添加文件】按钮，找到并选择 Python 脚本生成的“hurst_index.tif”文件，数据集类型选择“栅格”，导入 NDVI 栅格数据集。右击导入的数据集，选择【添加到新地图】，地图窗口中将显示 Hurst 指数计算结果。

右击“hurst_index”图层，使用专题图功能制图显示 2000—2025 年毛乌素沙地 Hurst 指数分布图(图 1-25)。



彩图 1-25

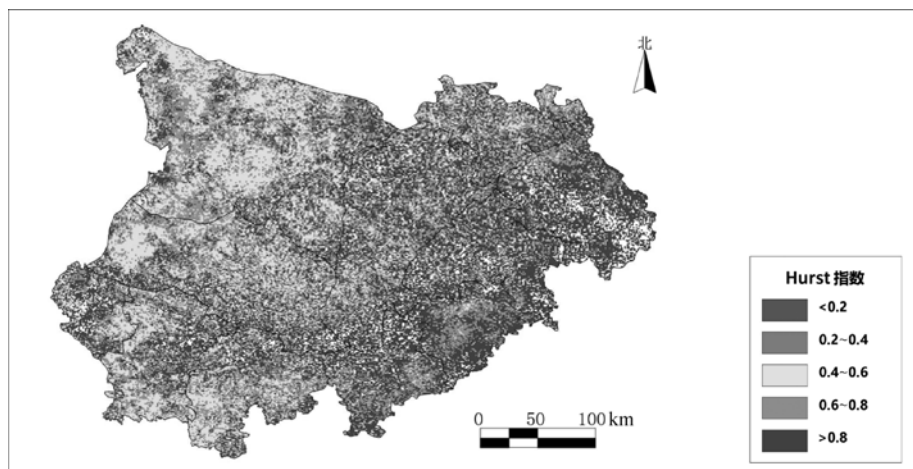


图 1-25 2000—2025 年毛乌素沙地 Hurst 指数分布图

3. 趋势-持续性耦合分析

1) 数据准备与加载

在 SuperMap iDesktopX 软件的【工作空间管理器】中，右击“DataSource”数据源，选择【导入数据集...】，在【数据导入】对话框中，单击【添加文件】按钮，找到并选择 Python 脚本生成的“hurst_classification.tif”文件，数据集类型选择“栅格”，导入 NDVI 栅格数据集。右击数据源中已导入的数据集“ndvi_trend_classification”和“hurst_classification”，选择【添加到新地图】，将数据添加到地图中进行后续操作。

2) 查看数据属性

在【工作空间管理器】中，分别右击“ndvi_trend_classification”和“hurst_classification”数据集，选择【统计栅格值】，查看数据集相关信息(图 1-26)。

ndvi_trend_classGridValue_1@DataSource				
序号	*SmID	SmUserID	GridValue	GridCount
1	1	0	1	73995
2	2	0	2	16220
3	3	0	3	1187
4	4	0	4	959
5	5	0	5	145

图 1-26 栅格值统计表

3) 数据预处理

为确保“ndvi_trend_classification”和“hurst_classification”数据集能够正确叠加分析，

需要统一两个数据集的无效值。在 SuperMap iDesktopX 功能选项卡【数据】中定位到【数据处理】组,单击【栅格】栏目中的【代数运算】工具,在弹出的【代数运算】对话框中,找到代数运算表达式输入区域。输入条件表达式:

$$\text{Con}(\text{IsNull}([\text{ndvi_trend_classification}])\text{OR}[\text{ndvi_trend_classification}] == -9999, 0, [\text{ndvi_trend_classification}]) \quad (1-13)$$

设置输出数据集名称为“ndvi_trend_processed”,单击【确定】执行运算。Hurst 分类数据中无效值已经是 0,无须特殊处理,确保两个数据集具有相同的空间参考系统和分辨率。

4) 数据重分级(可选)

如果两个数据集分类编码不一致或需要调整,可进行重分级。

在 SuperMap iDesktopX 功能选项卡【数据】中定位到【数据处理】组,单击【栅格】栏目中的【重分级】工具。在弹出的【重分级】对话框的【源数据】区域,选择数据源为“DataSource”,数据集为“ndvi_trend_processed”。

在【参数设置】区域,选择“单值重分级”方式。单击【设置】按钮,选择“级数”为 5,设置重分级规则:1→1(明显改善);2→2(轻微改善);3→3(稳定不变);4→4(轻微退化);5→5(严重退化);同时在工具对话框中勾选“未分级单元”参数选项,并设置其参数值为 0,输出数据集名称设为“ndvi_trend_classified”。

选择输入数据集“hurst_classification”,设置重分类规则:1→1(反持续性);2→2(随机序列);3→3(持续性);0→0(无效值),输出数据集名称设为“hurst_classified”。

5) 趋势-持续性耦合分析

在 SuperMap iDesktopX 功能选项卡【数据】中定位到【数据处理】组,单击【栅格】栏目中的【代数运算】工具,在弹出的【代数运算】对话框中,找到代数运算表达式输入区域,输入耦合分析表达式(图 1-27),设置输出数据集名称为“trend_persistence_coupling”,单击【确定】执行运算。

```
Con(
    [hurst_classified] == 3, // 如果是持续性
    [ndvi_trend_classified] + 3, // 将趋势分类(1-5)转换为耦合分类(4-8),后续再重映射
    Con(
        [hurst_classified] == 1 OR [hurst_classified] == 2, // 如果是反持续性或随机序列
        9, // 赋值为9(未来变化趋势不确定),后续再重映射为6
        0 // 无效值
    )
)
```

图 1-27 耦合分析表达式

6) 重分级为最终耦合分类

将计算结果重分级为 1~6 的最终分类。在 SuperMap iDesktopX 功能选项卡【数据】中定位到【数据处理】组,单击【栅格】栏目中的【重分级】工具。在弹出的【重分级】对话框的【源数据】区域,选择数据源为“DataSource”,数据集为“trend_persistence_coupling”。在【参数设置】区域,选择【单值重分级】方式。

单击【设置】按钮,选择“级数”为 6,设置重分级规则:4→1(持续显著改善→重映射为 1);5→2(持续轻微改善→重映射为 2);6→3(持续稳定不变→重映射为 3);7→4(持续轻微退化→重映射为 4);8→5(持续严重退化→重映射为 5);9→6(未来变化趋势不确

定→重映射为 6), 同时在工具对话框中勾选“未分级单元”参数选项, 并设置其参数值为 0, 输出数据集名称设为“trend_persistence_final”, 单击【确定】执行重分级(图 1-28)。



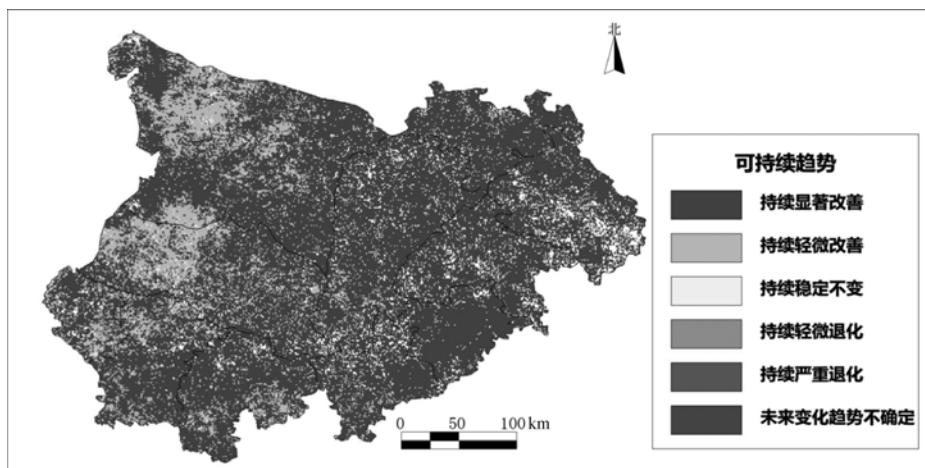
图 1-28 重分级参数设置界面

7) 结果可视化

在【工作空间管理器】窗口“DataSource”数据源中, 右击“trend_persistence_final”数据集, 选择【添加到新地图】。在图层管理器窗口中, 右击“trend_persistence_final”图层, 选择【制作专题图】, 在弹出的【制作专题图】窗口左侧栏目中选择【单值专题图】, 在弹出的【专题图】窗口中设置地图显示的相关属性(图 1-29)。对毛乌素沙地 NDVI 变化趋势持续性区域空间分布结果进行可视化(图 1-30)。



图 1-29 单值专题图参数设置窗口



彩图 1-30

图 1-30 毛乌素沙地 NDVI 变化趋势持续性区域空间分布

8) 结果统计与分析

统计各类别面积及占比,在【工作空间管理器】中,右击“trend_persistence_final”数据集,选择【统计栅格值】,查看数据集相关信息(图 1-31)。

根据像元数计算实际面积: 面积(km²)=像元数×(空间分辨率/1000)²

计算各类别百分比: 百分比(%)=(类别像元数/总有效像元数)×100

trend_persisten_1GridValue@DataSource			trend_persisten_2GridValue@I	
序号	*SmID	SmUserID	GridValue	GridCount
1	1	0	1	64023
2	2	0	2	13528
3	3	0	3	981
4	4	0	4	781
5	5	0	5	116
6	6	0	6	5465

GridValue	变化趋势持续性	GridCount	面积/km2	占比/%
1	持续显著改善	64023	64023	75.42
2	持续轻微改善	13528	13528	15.94
3	持续稳定不变	981	981	1.16
4	持续轻微退化	781	781	0.92
5	持续严重退化	116	116	0.14
6	未来变化趋势不确定	5465	5465	6.44
				100.00

图 1-31 各类别面积及占比统计结果

1.3 长时序植被覆盖时空演变分析实验思考与探索

1. 思考题

1) 根据提供的多期 NDVI 数据,完成植被覆盖动态分析,编写关键实验步骤并保存相关结果图与数据。

2) 回答下列练习题,并填写关键结果。

(1) 研究区 2023 年 $NDVI \geq 0.1$ 的植被覆盖区域面积为_____ km²。

- (2) 2000—2025 年,研究区整体 NDVI 的 Theil-Sen Median 趋势斜率为_____。
- (3) Mann-Kendall 检验显示,整体变化趋势的 Z 统计量为_____,在 0.05 置信水平上的显著性检验是否显著_____(是/否)。
- (4) 研究区 NDVI 序列的 Hurst 指数均值为_____,表明未来变化整体具有_____(持续性/反持续性/随机性)。
- 3) 植被覆盖变化趋势可进行分级评价,不同等级对应不同的生态恢复或退化状态。请根据趋势分析结果,填写表 1-2,并列举出完成该表所使用的分析功能。

表 1-2 植被覆盖变化趋势等级面积统计

变化趋势等级	面积/km ²	占比/%
明显改善		
轻微改善		
基本稳定		
轻微退化		
严重退化		

2. 探索题

本实验的重点在于利用时序遥感数据与经典统计方法(Theil-Sen Median 趋势分析、Mann-Kendall 检验、Hurst 指数)对区域植被覆盖进行监测与评估,定量揭示植被变化的趋势、显著性及未来持续性,为生态系统管理与规划提供数据支持。尽管本实验建立了一套较为完整的分析框架,但在方法深化与应用拓展方面仍存在诸多值得探索的空间,例如,没有针对不同植被类型的变化进行分析;植被变化受气候、人文和政策的综合影响,因此需要进一步加强对 NDVI 变化的原因进行分析;植被变化的地域差异性十分显著,需要加强对不同地域植被变化差异的研究。

另外,除了本实验所用的方法,是否还有其他更简便或更稳健的趋势计算方法? Hurst 指数的计算方法除经典的 R/S 分析外,还有去趋势波动分析(DFA)、小波分析等方法。不同方法计算的 Hurst 指数结果是否存在差异? 如何评估并选择最适合植被时序数据长程依赖性分析的方法? 植被变化受气候、人文和政策的综合影响,如何将气象数据(降水、温度)、地形数据、土地利用、覆被变化数据,以及人类活动数据(如夜间灯光、人口密度)与 NDVI 趋势进行空间关联或耦合建模,以更深入地揭示植被变化的自然与人为驱动机制? 感兴趣的读者可以尝试对实验进行深化与拓展,从而更全面地掌握生态环境遥感时序分析的前沿方法与复杂应用。



第2章

生态系统服务

2.1 生态系统服务评估理论与技术

2.1.1 生态系统服务评估方法

生态系统服务(ecosystem services)是指人类从生态系统获得的所有惠益,包括供给服务(如提供食物和水)、调节服务(如控制洪水和疾病)、文化服务(如精神、娱乐和文化收益),以及支持服务(如维持地球生命生存环境的养分循环)。其中,固碳是生态系统服务的重要功能之一,能够通过碳储量和碳汇功能对生态系统进行评估,对分析研究区气候变化和可持续发展具有重要意义。

1. InVEST 模型

生态系统服务和权衡的综合评估(integrated valuation of ecosystem services and trade-offs, InVEST)模型是由美国斯坦福大学、大自然保护协会(TNC)与世界自然基金会(WWF)联合开发的生态系统服务评估工具。该模型通过模拟不同土地覆被情景下生态系统服务物质和价值量的变化,填补了生态系统服务功能定量评估领域的空白,为自然资源管理决策提供量化支撑,其评估结果通过空间化地图实现可视化表达。InVEST 模型中包含多个模块,如产水量模型、生物多样性模型、碳储量模型等,这些模块均具有重要的研究和利用价值。

2. 碳储量计算

InVEST 模型的碳储量模块(carbon storage and sequestration)可用于评估固碳服务,在导入与目标区域相关的土地利用数据和碳库表后,模型通过将四个碳库(地上生物量、地下生物量、土壤有机物、枯落物)的平均碳密度乘以不同土地利用/土地覆被类型的面积来计算生态系统碳储量,其具体计算公式如下:

$$C_{\text{above}} = \sum_i^n c_{i_{\text{above}}} \times A_i \quad (2-1)$$

$$C_{\text{below}} = \sum_i^n c_{i_{\text{below}}} \times A_i \quad (2-2)$$

$$C_{\text{soil}} = \sum_i^n c_{i_{\text{soil}}} \times A_i \quad (2-3)$$

$$C_{\text{dead}} = \sum_i^n c_{i_{\text{dead}}} \times A_i \quad (2-4)$$

$$C_{\text{total}} = C_{\text{above}} + C_{\text{below}} + C_{\text{soil}} + C_{\text{dead}} \quad (2-5)$$

式中, $c_{i_{\text{above}}}$ 、 $c_{i_{\text{below}}}$ 、 $c_{i_{\text{soil}}}$ 、 $c_{i_{\text{dead}}}$ 分别表示植被地上碳密度、地下碳密度、表层土壤有机碳密度及枯落物碳密度; A_i 表示第 i 类地类的面积; C_{total} 、 C_{above} 、 C_{below} 、 C_{soil} 、 C_{dead} 分别表示总碳储量,地上部分、地下部分、土壤以及枯落物碳储量(单位为 t)。

3. 元胞自动机

元胞自动机(cellular automaton, CA)是一种时间、空间与状态均离散的动力学模型,由元胞、状态、邻域与转换规则组成,其基本原理是元胞在 $t+1$ 时刻的状态是其在 t 时刻状态和邻域状态的函数。每个元胞代表研究区域内最小的空间单元,会依据自身当前状态、邻域元胞的状态以及预设的转换规则进行同步更新,通过大量简单元胞的局部交互,最终涌现出复杂的全局时空演化格局,适用于土地利用预测。

2.1.2 生态系统服务评估技术应用

1. 实验概述

实验以特定研究区域为对象,首先基于研究区遥感影像,对其进行辐射定标等预处理,为后续分析奠定基础;随后进行土地利用类型监督分类,明确不同地表覆被类型的空间分布格局;并使用 InVEST 模型的碳储量模块,结合土地利用数据、碳库参数等量化该区域固碳量,实现生态系统固碳服务的定量化评估;在此基础上,运用元胞自动机对土地利用变化进行模拟;最后通过 Kappa 系数等指标对模拟结果进行精度验证,确保结果的可靠性。

实验全过程贯穿“数据—方法—应用”的逻辑主线,既注重遥感与地理模型的技术融合,又强调从空间视角理解土地利用变化与生态系统服务的耦合关系,为后续生态系统服务优化调控、国土空间规划等实践应用提供科学支撑,同时培养学生运用跨学科技术解决生态环境问题的综合能力。

2. 生态系统文明学思践行

通过实验中对生态系统固碳服务的量化评估与土地利用变化的动态模拟,引导学生直观认识生态系统在应对全球气候变化、维持碳平衡中的核心作用,深刻理解“绿水青山就是金山银山”的生态价值内涵,树立尊重自然、顺应自然、保护自然的生态文明观,增强守护生态环境的责任感与使命感。

同时,实验各环节环环相扣,从遥感影像预处理的细节把控、土地利用分类的精度控制,到模型参数的校准、模拟结果的验证,每一步都直接影响最终结论的可靠性。借此培养学生严谨细致的实验习惯,引导其以实事求是的态度对待数据与结果,树立科学研究的诚信意识。

生态系统服务评估是一个典型的跨学科领域,涉及遥感、地理学、环境科学等多学科知识。实验过程中引导学生打破学科壁垒,认识专业知识在服务国家生态安全、推动绿色低碳发展中的重要作用,激发其将个人专业发展与国家战略需求相结合,树立“学用结合、报效祖国”的价值追求。

3. 实验目的

(1) 掌握遥感影像预处理(辐射定标、大气校正、监督分类等)的核心原理与操作方法,理解预处理对后续分析结果的影响机制。

(2) 熟悉 InVEST 模型的基本架构与固碳模块的核心原理,掌握模型参数的获取、校准

与输入方法,理解模型量化生态系统固碳服务的核心机制。

(3) 了解元胞自动机模型的基本原理与应用场景,掌握基于 CA 模型模拟土地利用变化的核心步骤,理解自然与社会经济驱动因素对土地利用变化的影响,提升空间模拟与综合分析能力。

4. 实验数据

本实验所使用的数据存储于 data_exp2 文件夹中,包含的数据内容及来源见表 2-1。



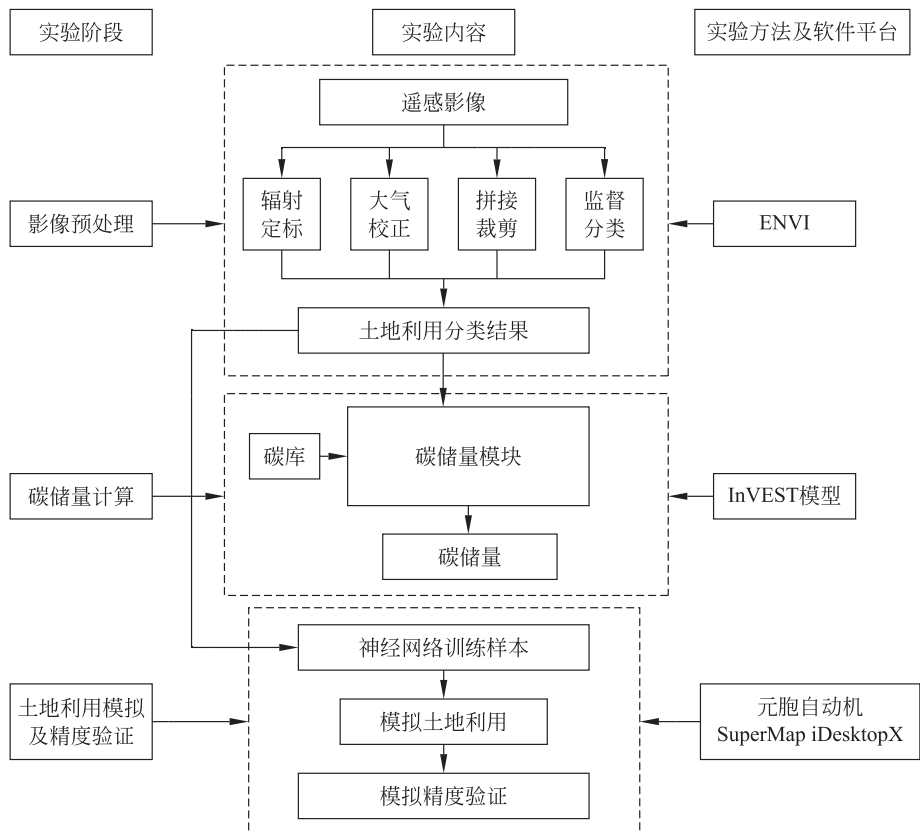
实验数据

表 2-1 实验数据内容及其来源

数据类型	数据来源	数据内容	数据说明
Landsat 遥感影像	地理空间数据云	研究区遥感影像	栅格数据,* .tif 格式
基础地理数据	—	研究区行政边界	矢量格式,线、面数据
GDP	参考文献[11]	国内生产总值	栅格数据,用于计算适宜性概率
逐年降水	国家青藏高原科学数据中心平台	—	栅格数据
逐年相对湿度	国家冰川冻土沙漠科学数据中心	—	栅格数据
逐年平均气温	国家青藏高原科学数据中心平台	—	栅格数据

5. 实验整体设计

本实验的整体流程如图 2-1 所示。



2.2 生态系统服务实验

2.2.1 遥感影像预处理

1. 辐射定标

打开 ENVI 5.6 软件,导入遥感影像,在右侧工具箱中搜索【Radiometric Calibration】,选中需要辐射定标的影像后,单击【Apply FLAASH Settings】并设置好输出影像的路径及文件名,单击【OK】即可完成辐射定标(图 2-2)。

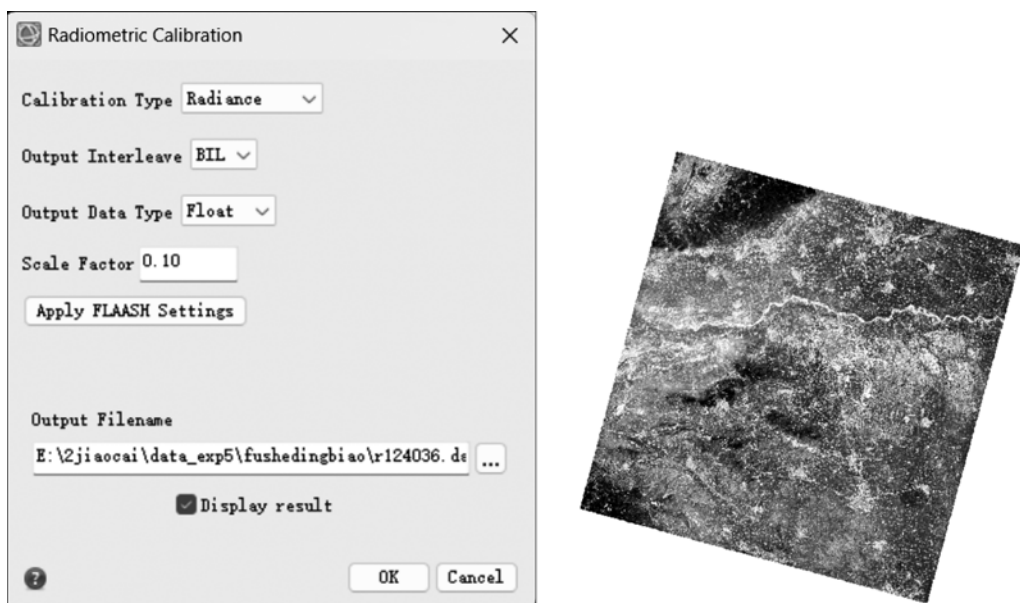


图 2-2 ENVI 辐射定标

2. 大气校正

在右侧工具箱处搜索大气校正工具【FLAASH Atmospheric Correction Module Input Parameters】,导入辐射定标结果影像,选择【Use single scale factor for all bands】,单击【OK】(图 2-3)。

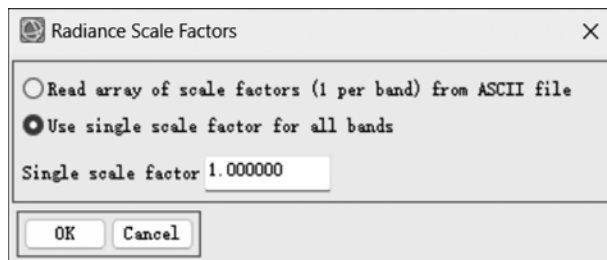


图 2-3 设置大气校正参数

设置文件输出路径和中间文件路径,中间文件路径需与输出路径一致,但不需要文件名。中心点经纬度由图像的地理坐标自动获取,传感器类型选择【Landsat-8 OLI】。设置好成像时间、大气模型与气溶胶模型等参数,如图 2-4 所示,大气校正结果如图 2-5 所示。

图 2-4 大气校正参数设置

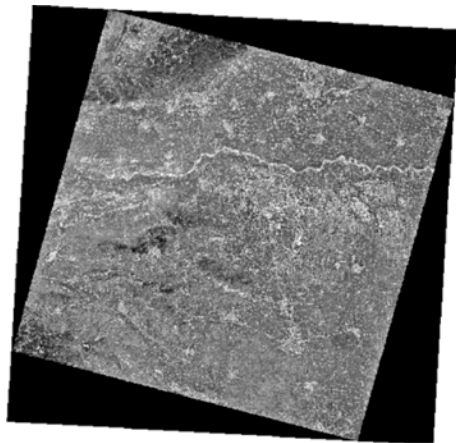
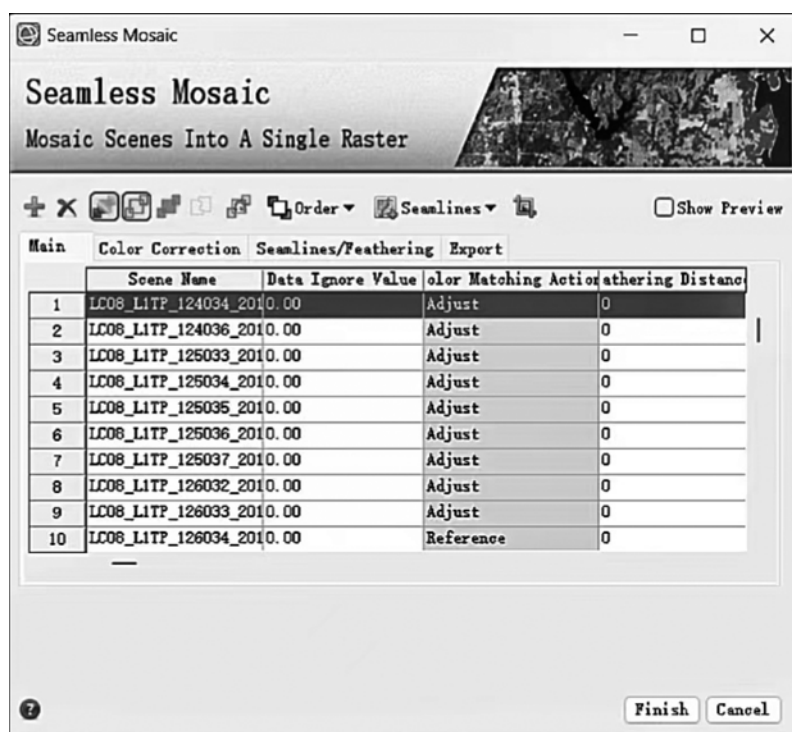


图 2-5 大气校正结果

3. 镶嵌裁剪

使用 ENVI 5.6 中的【Seamless Mosaic】工具对遥感影像进行镶嵌,镶嵌参数设置如图 2-6 所示,图(a)先导入需要镶嵌的所有遥感影像;图(b)对颜色校正进行参数选择;图(c)羽化镶嵌边缘,使接缝处过渡自然;图(d)设置输出路径;图(e)再使用【Subset Data from ROIs】导入研究区域的掩膜数据进行裁剪,ENVI 镶嵌裁剪结果如图 2-7 所示。



(a)



(b)

图 2-6 ENVI 镶嵌参数设置



(c)



(d)

图 2-6(续)



(c)

图 2-6(续)

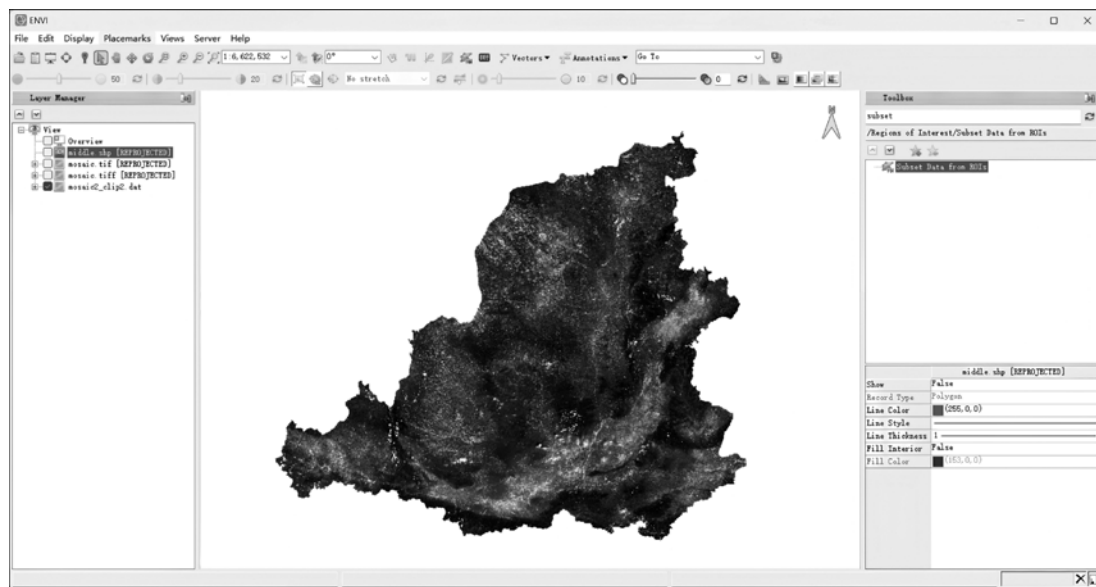


图 2-7 ENVI 镶嵌裁剪结果

4. 监督分类

在 ENVI 5.6 中导入裁剪好的影像数据右击, 选中该图层, 选择【New Region of Interest】以绘制训练样本, 这里按耕地、林地、草地、水体、建筑用地、裸土六个类型进行分



彩图 2-7

类(图 2-8)。

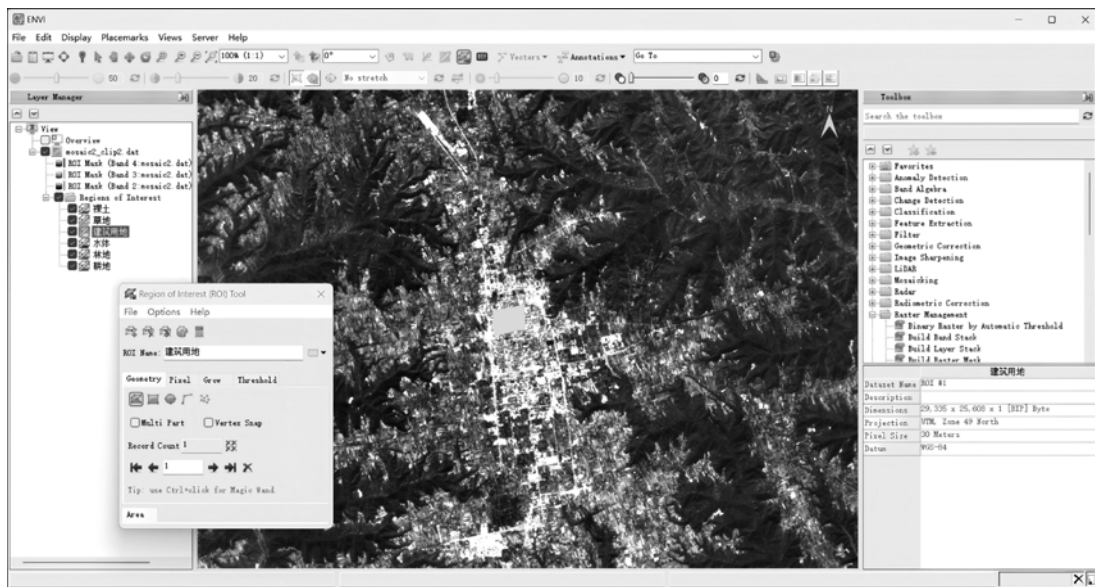


图 2-8 绘制训练样本

样本绘制完成后,单击【Options】→【Compute ROI Separability】对所绘制的样本进行分离度计算,样本类别之间分离度大于 1.8 为佳(图 2-9)。

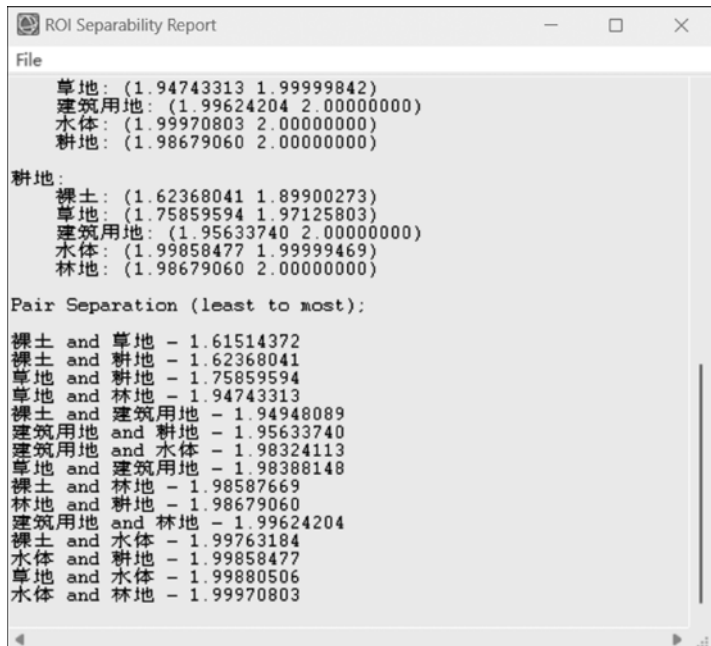


图 2-9 样本分离度

在右侧工具箱中单击选择【Classification】→【Supervised Classification】→【Maximum

Likelihood Classification】，对遥感影像进行监督分类处理，结果如图 2-10 所示。

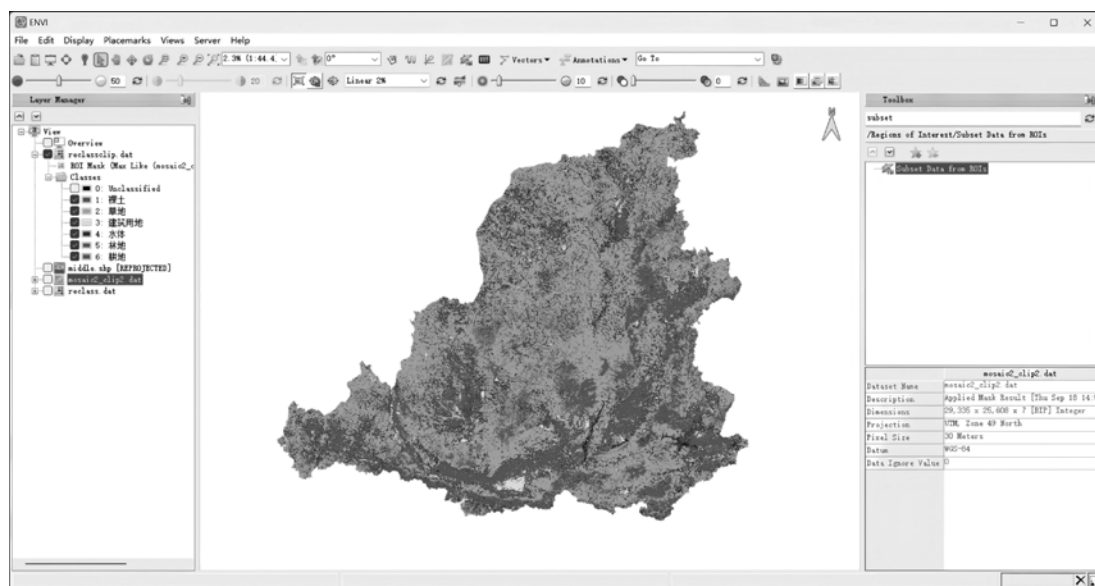


图 2-10 监督分类结果

2.2.2 碳储量的计算

打开 InVEST 软件，选择【Carbon Storage and Sequestration】模块对研究区域的碳储量情况进行计算(图 2-11、图 2-12)。

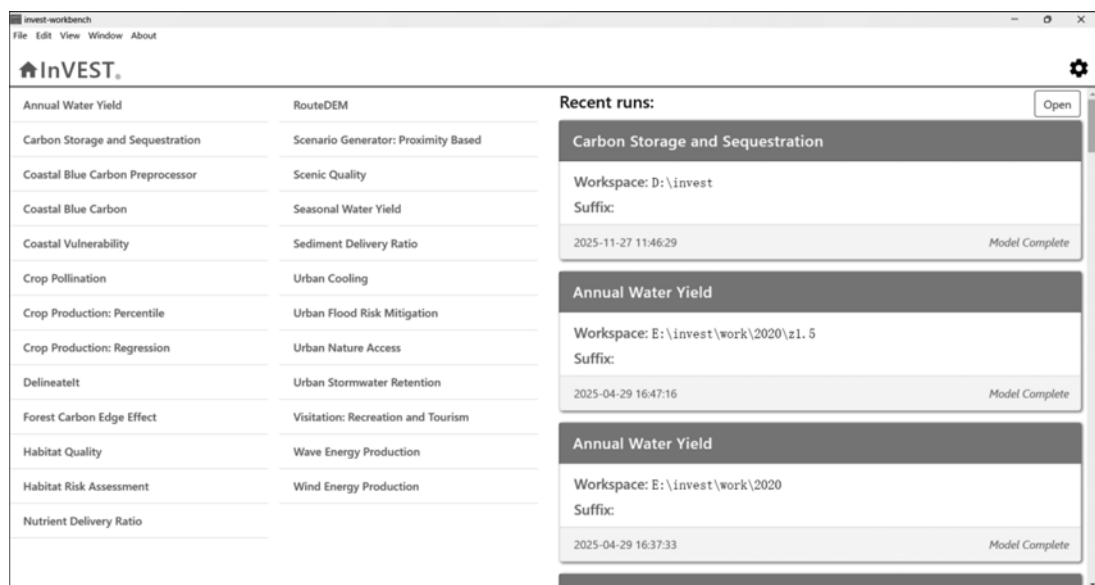


图 2-11 InVEST 主页

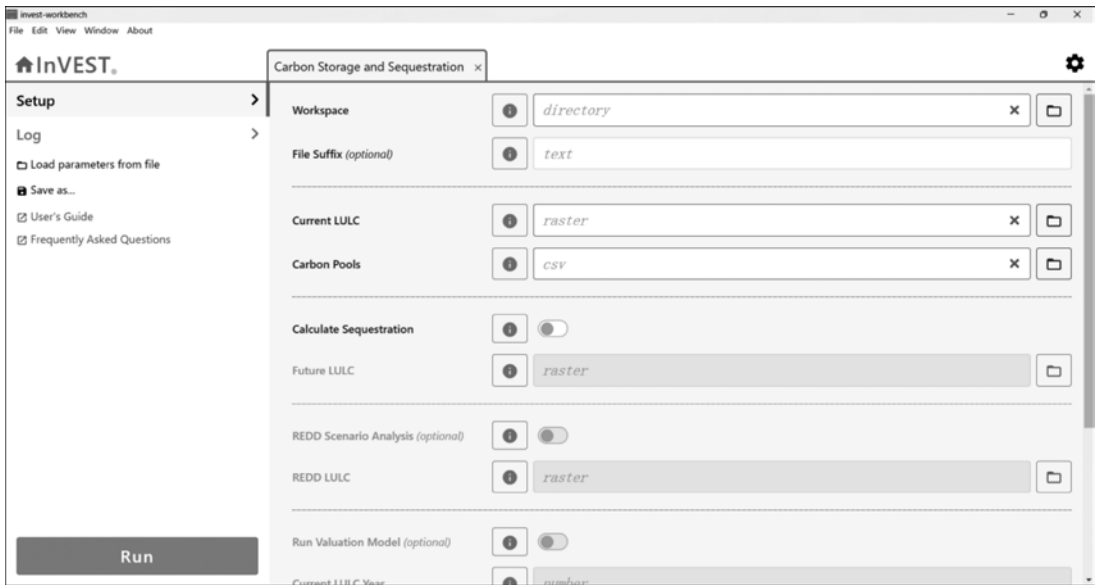


图 2-12 【Carbon Storage and Sequestration】模块

在相应位置设置好工作路径,并导入分类好的遥感影像数据和碳库 csv 表,单击【Run】即可成功运行该模型,结果如图 2-13 所示。

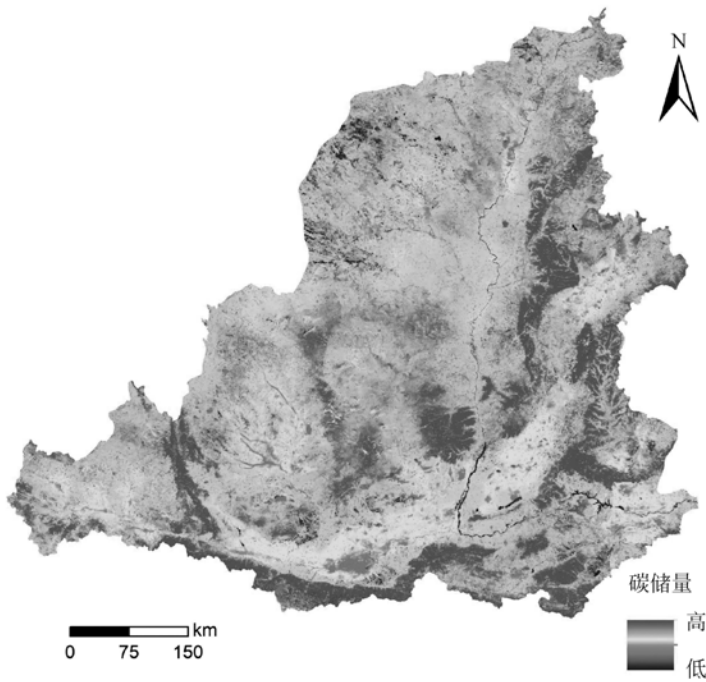


图 2-13 碳储量计算结果



彩图 2-13

2.2.3 土地利用模拟

1. 启动模块

选择主菜单的【FLUS Model】→【ANN-based Probability-of-occurrence Estimation】或者单击工具条中的工具按钮。打开【ANN-based Probability-of-occurrence Estimation】操作窗口。在【Land Use Data】中输入土地利用数据。单击【Set NoData Value】，依据第一列 Land Use Code 中的各类数值，在第二列 NoData Option 中选择对应的 Valid Data 或 NoData Value；第三列会同步显示该数值对应的标签说明，第四列则列出各土地利用类型对应的栅格像元总数。单击【Accept】按钮，完成土地利用数据的设置（图 2-14、图 2-15）。

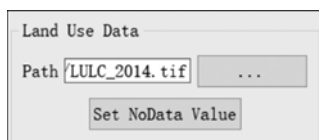


图 2-14 【Land Use Data】界面

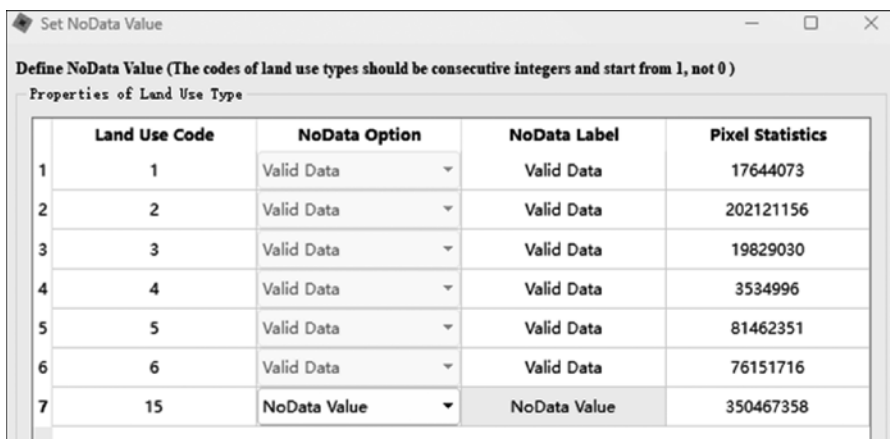


图 2-15 【Set NoData Value】界面

在【ANN Training】界面设置神经网络训练样本的采样模式，选择【Uniform Sampling】或【Random Sampling】。该示例中，选择【Uniform Sampling】，设置 Sampling Rate 为 20，Hidden Layer 为 12（图 2-16）。

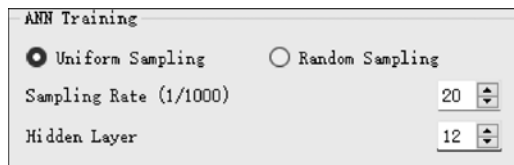


图 2-16 【ANN Training】界面

在【Save path】界面指定保存路径并输入要生成的出现概率数据的文件。勾选【Single

Accuracy】将生成 Float 类型影像,勾选【Double Accuracy】将生成 Double 类型影像。一般情况下,默认勾选【Single Accuracy】(图 2-17)。

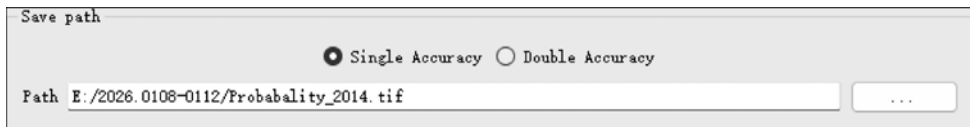


图 2-17 【Save path】界面

在【Driving Data】界面加载驱动力因子数据,单击【+】按钮添加新的驱动力因子数据;选中某个错误添加的驱动力因子,单击【-】按钮将其删除。

根据驱动力因子的实际情况,选择是否对驱动力因子进行归一化处理。系统默认选项为【Normalization】。若输入数据已预先完成归一化,则用户可选择【No Normalization】跳过该步骤(图 2-18)。

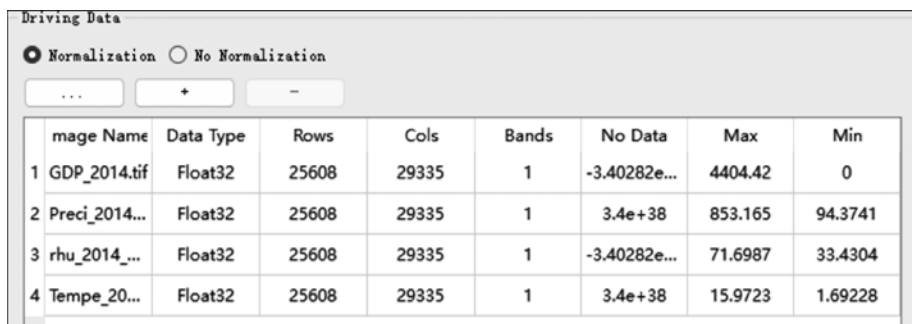


图 2-18 【Driving Data】界面

完成以上设置后,单击【Start Running】,开始进行神经网络模型的适宜性概率训练。训练结束后,将弹出【Finished training!】信息提示框(图 2-19)。

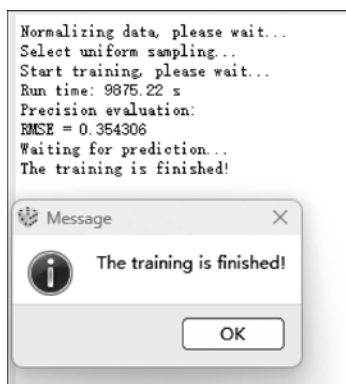


图 2-19 完成界面

2. 基于元胞自动机的预测模拟

单击主菜单的【FLUS Model】→【Self Adaptive Inertia and Competition Mechanism CA】或者单击工具条中的工具按钮,弹出【Land Use Data】界面。

单击【Set Land Use Type, Color Display and NoData Value】,列表的第一列 Land Use Code 为土地利用类型的代码值,在第二列 NoData Option 的下拉菜单中需要为每个代码选择对应的 Valid Data 或 NoData Value。用户可在 Land Use Type 一栏中输入对应的土地利用类型名称,如不输入将由默认名称依次填充(图 2-20、图 2-21)。

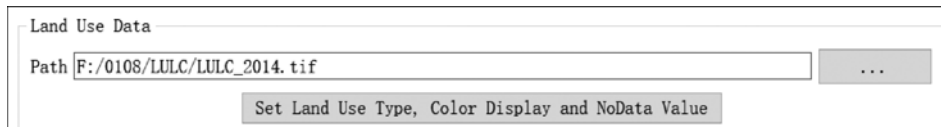


图 2-20 【Land Use Data】界面

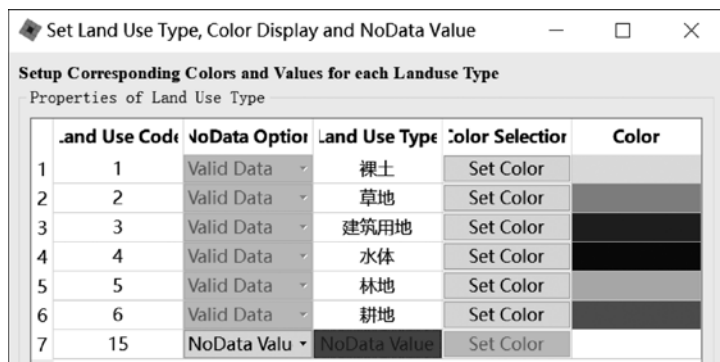


图 2-21 【Set Land Use Type, Color Display and NoData Value】界面

在【Probability Data】界面输入适宜性概率数据(图 2-22)。在【Save Simulation Result】界面设置模拟结果的保存路径(图 2-23)。

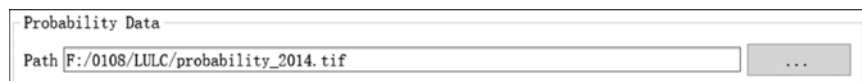


图 2-22 【Probability Data】界面

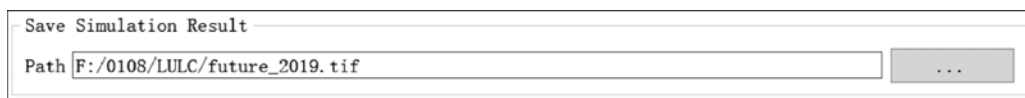


图 2-23 【Save Simulation Result】界面

在【Simulation Setting】界面设置模拟参数。Maximum Number of Iteration、Neighborhood、Accelerate、Tnread 分别默认为 300、3、0.1 和 1(图 2-24)。

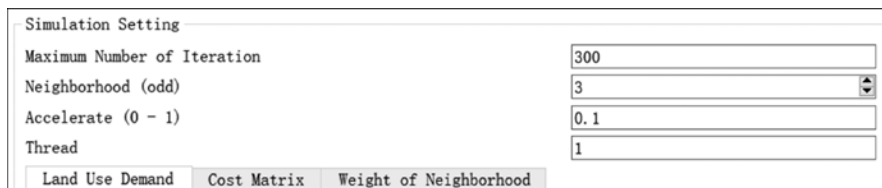


图 2-24 【Simulation Setting】界面

Initial Pixel Number 表示初始年份各类土地利用类型的像元数量,Future Pixel Number 表示未来年份各类土地利用类型的像元数量。用户需依据研究区域的实际发展情况,采用土地利用数量预测模型等方法,对未来的土地需求量进行预测,并在 Future Pixel Number 中输入各类土地利用类型变化后的目标数量(图 2-25)。

Land Use Demand	Cost Matrix		Weight of Neighborhood			
	裸土	草地	建筑用地	水体	林地	耕地
Initial Pixel Number	17644073	202121156	19829030	3534996	81462351	76151716
Future Pixel Number	12884212	196459628	25489184	3654884	89456865	72798549

图 2-25 【Land Use Demand】界面

在【Cost Matrix】界面需要设置各类土地利用类型在模拟转换中的成本矩阵。本示例中,基于对研究区域的实际认识,设定了转换限制规则。当某一用地类型不允许转换为另一类型时,需将矩阵中对应的值设为 0,允许转换时设为 1(图 2-26)。

Land Use Demand	Cost Matrix		Weight of Neighborhood			
	裸土	草地	建筑用地	水体	林地	耕地
裸土	1	0	1	1	0	0
草地	1	1	1	0	0	1
建筑用地	1	0	1	0	0	0
水体	1	0	0	1	0	0
林地	1	1	1	0	1	1
耕地	1	1	1	1	1	1

图 2-26 【Cost Matrix】界面

在【Weight of Neighborhood】界面设置各类土地利用类型的邻域因子,其取值范围为 0~1。数值越接近 1,表示该土地利用类型的空间扩张能力越强(图 2-27)。

Land Use Demand	Cost Matrix		Weight of Neighborhood			
	裸土	草地	建筑用地	水体	林地	耕地
Weight of neighborhood	0.2	0.5	1	0.3	0.5	0.6

图 2-27 【Weight of Neighborhood】界面

完成以上各项设置后,单击【Accept】确认。选择【Show】,进行土地利用变化的模拟。单击【Run】开始模拟,窗口左侧上方的图表面板将显示各土地利用类型在迭代过程中数量变化的曲线。当达到用户设置的迭代次数或未来土地类型的数量目标时,软件将自动停止迭代,模拟完成。

2.2.4 模拟精度验证

选择主菜单【Validation】→【Precision Validation】→【Kappa】,有两种采样模式:Random Sampling 或 Uniform Sampling。选择 Random Sampling 需在下方的输入框中设置随机采样点的比例,选择 Uniform Sampling 需输入每类用地的采样点个数。本示例中选择 Random Sampling,Sampling rate for total valid pixels 设为 10。单击【Start】进行计算,计算结果会在界面上显示,并保存在软件目录下的 Kappa.csv 文件中(图 2-28)。

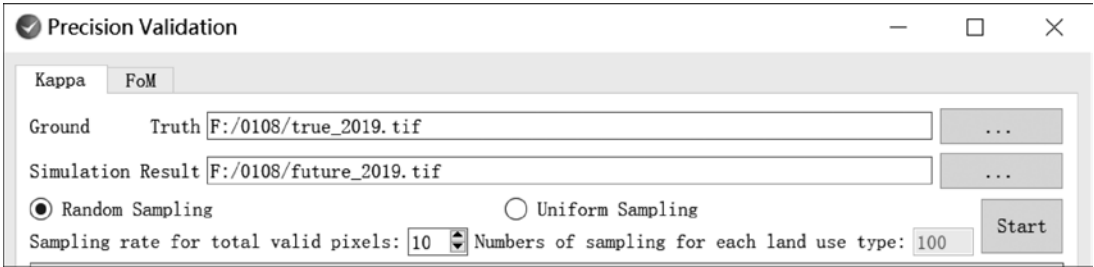


图 2-28 【Precision Validation】界面

2.3 生态服务系统实验思考与探索

1. 思考题

1) 结合研究区遥感影像与辅助数据,完成生态系统固碳服务评估及土地利用变化模拟,编写关键实验步骤并保存相关结果图与数据。

2) 完成下列练习题,并填写关键结果。

(1) 研究区遥感影像预处理后,大气校正后的地表反射率均值为_____ (保留 2 位小数)。

(2) 土地利用分类结果的混淆矩阵总体精度为_____% ,Kappa 系数为_____。

(3) InVEST 模型计算得出的研究区总固碳量为_____ t,其中_____ 的单位面积固碳量最高。

(4) 基于 CA 模型模拟的 2019 年研究区林地面积变化为_____ km²,对应的固碳量_____ (增加/减少)。

3) 请根据你的土地利用预测结果,填写表 2-2,并列举出完成该表所使用的分析功能。

表 2-2 土地利用预测结果

	2014 年利用面积/km ²	2019 年利用面积/km ²	变化率/%
耕地			
林地			
草地			
水体			
建筑用地			
裸土			

2. 探索题

(1) 本实验采用 InVEST 模型计算固碳量,除该模型外,还有 CASA(Carnegie-Ames-Stanford approach)模型可基于遥感数据估算植被净初级生产力(NPP)并推导固碳量。请分别用 InVEST 与 CASA 模型计算同一研究区的固碳量,对比两种模型的结果差异,并分析差异产生的核心原因(提示:关注模型原理、参数需求的区别)。

(2) 本实验中 CA 模型模拟土地利用变化时,主要考虑了社会经济驱动因素的“静态”影响。请收集研究区的气象数据(降水、温度)、人类活动数据(夜间灯光等),通过地理加权回归(GWR)分析不同驱动因素对土地利用变化的“空间异质性”影响,并将这种空间异质性关联结果融入 CA 模型的转换规则中,重新模拟土地利用变化,对比优化前后模型的模拟精度。



第3章

干旱事件识别和分析

3.1 干旱事件识别和分析理论与技术

3.1.1 干旱事件识别理论

干旱是指由降水和蒸发不平衡导致的水分亏缺现象。通常将干旱划分为气象干旱、农业干旱、水文干旱和社会经济干旱四种干旱类型,当降水减少时最先发生的是气象干旱,随着水分的持续蒸发和水资源的减少,逐渐引发农业干旱和水文干旱。当农业干旱和水文干旱发展到一定程度时就会引发社会经济干旱。因此可将不同干旱类型理解为:农业干旱即气象干旱对农业系统的影响,水文干旱即气象干旱对水文系统的影响,而社会经济干旱则是气象干旱对社会经济系统的影响。从干旱监测层面看,干旱作为一种复杂的现象,通常难以直接观测其发生时间、发展过程和影响范围,因而通常采用干旱指标对干旱事件进行描述。目前对干旱情况进行表征的干旱指数已发展百余种,本书以综合考虑降水和蒸散发的标准化降水蒸散指数(standardized precipitation evapotranspiration index, SPEI)为例开展干旱事件的识别与分析。

1. 标准化降水蒸散指数

标准化降水蒸散指数(SPEI)是由 Vicente-Serrano 等于 2010 年提出的一种多尺度干旱监测指数。SPEI 综合考虑了降水和潜在蒸散发两个关键气候要素,能够反映气候变暖背景下温度升高对干旱的影响,弥补了传统标准化降水指数(SPI)仅考虑降水的不足。

SPEI 指数具有以下特点:①多尺度性:可以计算不同时间尺度(如 1 个月、3 个月、6 个月、12 个月等)的干旱指数,适用于不同类型干旱的监测;②标准化特性: SPEI 值服从标准正态分布,均值为 0,标准差为 1,便于不同地区和时段的对比;③温度敏感性:考虑了潜在蒸散发,能够反映气温变化对干旱的影响;④时空可比性:标准化处理使得不同区域的干旱程度具有可比性。

SPEI 的计算主要包括以下步骤:

(1) 计算水分平衡量

水分平衡量(D)是降水量(P)与潜在蒸散发量(PET)的差值:

$$D_i = P_i - PET_i \quad (3-1)$$

式中, D_i 为第 i 月的水分平衡量(mm); P_i 为第 i 月的降水量(mm); PET_i 为第 i 月的潜在蒸散发量(mm)。

潜在蒸散发量 PET_i 可采用 Thornthwaite 方法和 Penman-Monteith 方法计算。其中 Thornthwaite 方法仅需温度数据, 计算公式为

$$PET = 16 \times \left(\frac{10T}{I} \right)^a \quad (3-2)$$

式中, T 为月平均气温($^{\circ}\text{C}$); I 为年热量指数, $I = \sum_{i=1}^{12} \left(\frac{T_i}{5} \right)^{1.514}$; a 为由年热量指数(I) 确定的无量纲经验指数, 用于表征潜在蒸散发对气温变化的非线性响应, 不具有独立、明确的物理意义。其数值由经验多项式计算: $a = 0.49239 + 1.792 \times 10^{-2} I - 7.71 \times 10^{-5} I^2 + 6.75 \times 10^{-7} I^3$ 。

本实验直接使用来自 ERA5-Land 的潜在蒸散发产品数据, 不用再次计算。

(2) 累积水分亏缺序列

对于不同时间尺度的 SPEI(如 SPEI-3 表示 3 个月尺度), 需要对水分平衡量进行累积:

$$D_n^k = \sum_{i=0}^{k-1} (P_{n-i} - PET_{n-i}), \quad n \geq k \quad (3-3)$$

式中, D_n^k 为第 n 月、时间尺度为 k 个月的累积水分平衡量; k 为时间尺度(月)。

(3) 概率分布拟合

将累积水分平衡量序列 D 采用三参数 log-logistic 概率分布进行拟合, 其概率密度函数为

$$f(x) = \frac{\beta}{\alpha} \left(\frac{x - \gamma}{\alpha} \right)^{\beta-1} \left[1 + \left(\frac{x - \gamma}{\alpha} \right)^{\beta} \right]^{-2} \quad (3-4)$$

累积概率分布函数为

$$F(x) = \left[1 + \left(\frac{\alpha}{x - \gamma} \right)^{\beta} \right]^{-1} \quad (3-5)$$

式中, α 为尺度参数($\alpha > 0$); β 为形状参数($\beta > 0$); γ 为位置参数。参数可采用线性矩法(L-moments)估计。

(4) 标准化处理

将累积概率 $F(x)$ 转换为标准正态分布, 得到 SPEI 值:

$$SPEI = W - \frac{C_0 + C_1 W + C_2 W^2}{1 + d_1 W + d_2 W^2 + d_3 W^3} \quad (3-6)$$

式中, $W = \sqrt{-2 \ln P}$, 当 $P \leq 0.5$ 时, $P = F(x)$; 当 $P > 0.5$ 时, $P = 1 - F(x)$ 。

系数取值为: $C_0 = 2.515517, C_1 = 0.802853, C_2 = 0.010328, d_1 = 1.432788, d_2 = 0.189269, d_3 = 0.001308$ 。

当 $P \leq 0.5$ 时, SPEI 符号为负; 当 $P > 0.5$ 时, SPEI 符号为正。

2. 干旱等级划分

根据国家标准《气象干旱等级》(GB/T 20481—2017) 和国际通用标准, 基于 SPEI 值的

干旱等级划分如表 3-1 所示。

表 3-1 干旱等级划分

SPEI 值范围	干旱等级	干旱类型
$\text{SPEI} > -0.5$	1	无旱
$-1.0 < \text{SPEI} \leq -0.5$	2	轻旱
$-1.5 < \text{SPEI} \leq -1.0$	3	中旱
$-2.0 < \text{SPEI} \leq -1.5$	4	重旱
$\text{SPEI} \leq -2.0$	5	特旱

在本研究中,当 $\text{SPEI} < -1.0$ 时,认为发生干旱事件。

3. 基于游程理论的干旱特征识别

1) 游程理论基本概念

游程理论(run theory)最早由 Yevjevich 于 1967 年提出,用于水文时间序列的分析,是目前干旱事件识别和特征提取的经典方法。游程理论的核心思想是:在一个时间序列中,当变量值连续低于(或高于)某一临界阈值时,这一连续时段被定义为一个“游程”。

在干旱研究中,游程理论通过设定干旱阈值(截断水平),将 SPEI 时间序列中连续低于该阈值的时段识别为一次独立的干旱事件,并可提取干旱事件的多个特征参数。

2) 干旱事件定义

基于游程理论,干旱事件可定义为:在 SPEI 时间序列中,当 SPEI 值连续低于预设阈值 x_0 (如 $\text{SPEI} = -1.0$) 时,从首次低于阈值到重新回到阈值以上的连续时段,构成一次完整的干旱事件。

需要注意的是,在实际应用中,为避免将短暂的 SPEI 值回升误判为干旱事件的结束,通常设置最小独立性准则。①最小干旱历时:一次干旱事件至少持续的时间(如 3 个月);②最小间隔期:两次干旱事件之间至少间隔的时间(如 1 个月),若两次干旱事件之间的间隔期小于设定值,则将其合并为一次干旱事件。

3) 干旱特征参数

基于游程理论识别出干旱事件后,可以提取以下特征参数定量描述干旱事件的时空特征:

(1) 干旱历时。干旱历时(drought duration, D)是指干旱事件持续的时间长度,通常以月为单位:

$$D = t_e - t_s + 1 \quad (3-7)$$

式中, D 为干旱历时(月); t_s 为干旱开始时间; t_e 为干旱结束时间。

(2) 干旱强度。干旱强度(drought severity, S)是指干旱期间水分亏缺的累积量,用 SPEI 负值的累加表示:

$$S = \sum_{i=t_s}^{t_e} |\text{SPEI}_i| \quad (3-8)$$

式中, S 为干旱强度(无量纲); SPEI_i 为第 i 月的 SPEI 值; $|\text{SPEI}_i|$ 表示 SPEI 绝对值。

(3) 干旱烈度/平均强度

干旱烈度(drought intensity, I)是干旱强度与历时的比值,表示单位时间内的平均干旱程度:

$$I = \frac{S}{D} = \frac{\sum_{i=t_s}^{t_e} |SPEI_i|}{t_e - t_s + 1} \quad (3-9)$$

式中, I 为干旱烈度(无量纲),干旱烈度可以反映干旱发展的速度和严重程度,烈度大表示干旱发展迅速或程度严重。

(4) 干旱峰值

干旱峰值(drought peak, P)是指干旱期间 SPEI 的最小值,代表干旱最严重的时刻:

$$P = \min(SPEI_{t_s}, SPEI_{t_s+1}, \dots, SPEI_{t_e}) \quad (3-10)$$

式中, P 为干旱峰值(无量纲),干旱峰值反映干旱事件的极端程度。

(5) 干旱频次

干旱频次(drought frequency, F)是指在研究时段内某地区发生干旱事件的总次数:

$$F = \sum_{j=1}^N n_j \quad (3-11)$$

式中, F 为干旱总次数; n_j 为第 j 个空间单元的干旱次数; N 为研究时段的总长度。

(6) 干旱累计强度

干旱累计强度(cumulative drought severity, CS)是指研究时段内某地区所有干旱事件强度的总和:

$$CS = \sum_{k=1}^F S_k \quad (3-12)$$

式中, CS 为累计干旱强度; S_k 为第 k 次干旱事件的强度; F 为干旱总次数。累计强度综合反映研究期内某地区干旱的总体受灾程度。

4) 干旱特征参数关系

干旱历时越长,通常干旱强度越大;干旱烈度反映了干旱强度与历时的综合效应;干旱频次高的地区,累计强度往往较大;干旱峰值体现了干旱的极端性,与强度不完全一致,这些参数从不同角度描述了干旱的时空特征,为全面评估区域干旱状况提供了定量依据。

4. 基于 Getis-Ord G_i^* 统计量的热点分析

空间自相关(spatial autocorrelation)是指空间中某一位置的属性值与其邻近位置属性值之间的相关程度,反映了地理现象在空间上的聚集或分散特征。Tobler 第一地理学定律指出:“任何事物都与其他事物相关,但近处的事物比远处的事物更相关”,这是空间自相关分析的理论基础。热点分析(hot spot analysis)是一种局部空间自相关分析方法,用于识别

在空间上显著聚集的高值区域(热点, hot spot)和低值区域(冷点, cold spot)。在干旱研究中, 热点分析可以识别出干旱高风险区域和低风险区域, 为防灾减灾提供科学依据。

1) Getis-Ord G_i^* 统计量原理

Getis-Ord G_i^* 统计量是由 Getis 和 Ord 于 1992 年提出的局部空间关联指标, 用于检测某一位置及其邻域范围内属性值的空间聚集程度。与其他局部空间自相关指标相比, G_i^* 统计量的特点是将目标要素本身纳入计算范围(即包括自身的影响), 更适合识别热点和冷点区域。

(1) Getis-Ord G_i^* 统计量的计算公式为

$$G_i^* = \frac{\sum_{j=1}^n w_{ij} x_j - \bar{X} \sum_{j=1}^n w_{ij}}{S \sqrt{\frac{n \sum_{j=1}^n w_{ij}^2 - \left(\sum_{j=1}^n w_{ij} \right)^2}{n-1}}} \quad (3-13)$$

式中, G_i^* 为位置 i 的 G_i 统计量; x_j 为位置 j 的属性值(本研究为干旱累计强度); w_{ij} 为空间权重矩阵, 表示位置 i 和位置 j 之间的空间关系; n 为研究区域内要素的总数; $\bar{X}$ 为所有要素属性值的平均值, $\bar{X} = \frac{\sum_{j=1}^n x_j}{n}$; S 为属性值的标准差, $S = \sqrt{\frac{\sum_{j=1}^n x_j^2}{n} - \bar{X}^2}$ 。

(2) G_i^* 统计量进一步标准化为 Z 得分:

Z 得分(Z -score)是 G_i^* 统计量的标准化形式, 服从标准正态分布, 其值即为上述公式计算的结果。 Z 得分的统计意义为

$Z > 0$: 表示位置 i 及其邻域内的属性值高于全局平均水平, 可能是高值聚集区(热点);

$Z < 0$: 表示位置 i 及其邻域内的属性值低于全局平均水平, 可能是低值聚集区(冷点);

$|Z|$ 越大表示空间聚集程度越显著。

2) 空间权重矩阵

空间权重矩阵 w_{ij} 定义了空间要素之间的邻近关系, 是热点分析的关键参数。常用的空间权重矩阵构建方法包括:

(1) 固定距离带(fixed distance band)

当位置 i 和位置 j 之间的距离小于设定的距离阈值 d 时, 认为二者为邻近关系:

$$w_{ij} = \begin{cases} 1, & d_{ij} \leq d \\ 0, & d_{ij} > d \end{cases} \quad (3-14)$$

式中, d_{ij} 为位置 i 和 j 之间的欧氏距离。

(2) 反距离权重(inverse distance)

权重值随距离增加而衰减:

$$w_{ij} = \frac{1}{d_{ij}^p} \quad (3-15)$$

式中, p 为距离衰减指数, 通常取 1 或 2。

(3) K 近邻(K nearest neighbors)

选择距离位置 i 最近的 K 个要素作为邻域:

$$w_{ij} = \begin{cases} 1, & j \in N_k(i) \\ 0, & j \notin N_k(i) \end{cases} \quad (3-16)$$

(4) 邻接关系(contiguity)

对于面状要素,可采用边界邻接或顶点邻接定义空间关系。在本实验中,根据研究区的空间尺度和数据特征选择合适的空间权重构建方法。

3.1.2 干旱事件识别技术应用

1. 实验概述

标准化降水蒸散指数(SPEI)是刻画区域气候水分平衡异常状况、识别干旱过程的重要指标,能够综合反映降水变化与潜在蒸散发对地表水分条件的影响。该指数通过描述气候水分收支相对于其长期平均状态的偏离程度,有效表征气象干旱及其向农业干旱发展的过程。基于 SPEI 的时序分析可识别干旱事件的发生、发展与缓解过程,揭示干旱在时间演变和空间分布上的差异特征,是干旱监测与评估研究中的重要技术手段。

本实验以长江中下游地区为研究对象,介绍基于 SPEI 时间序列数据的干旱事件识别、事件分布特征分析以及空间热点识别方法,构建区域尺度干旱时空动态监测的技术流程。

实验首先基于已有降水与潜在蒸散发数据计算研究区不同时间尺度的 SPEI 指数,并通过空间可视化手段展示长江中下游地区干旱状况的总体分布格局;随后构建 SPEI 年际变化序列,采用标准化异常分析方法识别干旱事件及其强度等级,定量刻画干旱发生的时间特征与空间分异规律;在此基础上,引入空间统计分析方法,对干旱事件进行热点识别与聚集特征分析,揭示干旱高发区域及其空间集聚特征;最后综合干旱强度、发生频率与空间热点分布结果,对长江中下游地区干旱的时空演变特征及其区域差异进行综合分析。

基于 SPEI 的干旱识别与热点分析方法可广泛应用于气象干旱监测、农业生产风险评估、水资源短缺分析、生态脆弱区识别及防灾减灾决策支持等领域,对于认识干旱灾害的形成机制及其空间格局具有重要意义。该分析框架同样适用于其他干旱指标(如 SPI、SSI 等)的时序分析与空间热点识别,为区域尺度干旱监测与风险评估提供通用的技术思路。

2. 水资源安全学思践行

本实验以长江中下游地区干旱事件识别与时空分析为教学载体,将素质培养贯穿于专业知识学习与实践操作的全过程,引导学生在掌握干旱监测与空间分析技术方法的同时,深化对区域生态安全与水资源治理问题的认识,实现专业能力的培养与价值引领的有机融合。

在教学内容设置上,实验紧密联系我国生态文明建设和水资源安全的现实需求,强调集约用水和系统治理的重要性。长江中下游地区是我国重要的农业生产和经济发展区域,同时也是旱涝灾害交替频发的典型区域,干旱问题对粮食安全、生态稳定和社会发展具有重要影响。通过对研究区干旱时空特征的定量分析,学生能够直观认识区域干旱风险的空间差异性和阶段性变化特征,理解国家在流域水资源管理、防灾减灾和区域协调发展方面所面临的现实挑战,从而增强服务国家重大战略需求的责任意识。

在实验方法与分析过程中,教师注重引导学生形成系统思维和综合分析能力。干旱的形成与演变并非单一气候因素作用的结果,而是气候变化、生态系统状态,以及人类活动共同影响的综合体现。通过多尺度干旱分布特征及热点集聚特征的分析,帮助学生从整体视角认识自然过程与人类活动之间的相互作用关系,理解区域生态风险的差异性,培养科学分析复杂地理环境问题的能力。

在结果解读与反思环节,强调科学规范、数据意识与职业责任。学生需对干旱识别结果进行合理解读,认识模型假设和数据不确定性对分析结论的影响,避免简单化和片面化判断。结合干旱监测成果在农业管理、水资源调配和灾害风险防控中的应用,引导学生树立以社会需求和公共利益为导向的专业价值观,增强将专业技术成果服务于国家和社会发展的责任感。

通过“理论认知—方法应用—结果分析—反思提升”的教学设计,本实验在提升学生干旱监测与时空分析专业能力的同时,引导其在实践中深化对我国生态环境现状与发展目标的理解,推动专业教育与思想政治教育协同推进、同向发力。

3. 实验目的

(1) 掌握基于降水与潜在蒸散发时间序列数据的干旱识别基本原理与方法,理解标准化降水蒸散指数(SPEI)的构建思想及其在干旱监测中的应用意义。

(2) 综合应用时间序列异常分析、空间可视化与空间统计分析方法,实现干旱事件的定量识别与空间分布特征分析,刻画干旱发生频次、强度及其时空差异特征。

(3) 培养运用多源时空数据与 GIS 空间分析技术开展干旱监测与评估的实践能力,提升对区域干旱演变过程及其对生态环境与水资源影响的综合认识与科学评估水平。

4. 实验数据

本实验所使用的数据存储于 EX3 文件夹中,包含的数据内容及来源见表 3-2 及二维码。



实验数据

表 3-2 实验数据内容及其来源

数据类型	数据内容	数据说明
ERA5-Land 再分析数据	研究区 2000—2024 年降水、潜在蒸散发栅格数据	栅格数据,*.nc 格式
基础地理数据	研究区行政边界	矢量格式,线、面数据

5. 整体实验设计

本次实验以长江中下游地区干旱事件识别与时空分析为研究目标,基于标准化降水蒸散指数(SPEI)构建区域干旱监测与分析框架,重点围绕以下 4 个核心问题展开:一是研究区干旱事件在空间上的总体分布格局及其区域差异特征;二是研究区干旱发生频次、强度及持续时间等特征在时间尺度上的变化规律;三是研究区干旱发生频次在不同阶段(5 年尺度)上的演变趋势及其变化特征;四是研究区干旱事件在空间上的集聚特征,干旱高发区域(热点)与低发区域(冷点)主要分布在哪里。

本实验的整体流程如图 3-1 所示。

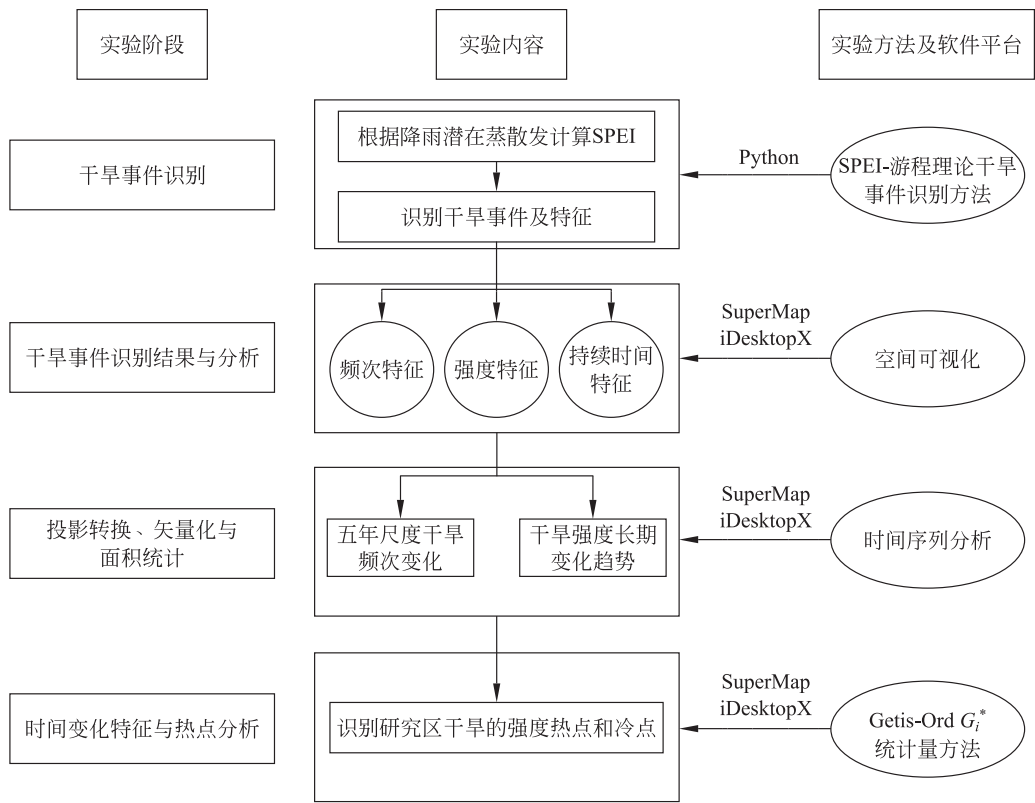


图 3-1 实验流程图

3.2 干旱事件识别和分析实验

3.2.1 干旱事件的识别

1. 计算标准化降水蒸散指数(SPEI)

(1) 数据准备：全球逐月降雨和潜在蒸散发数据，来自 ERA5-Land 再分析数据集（见二维码）。

先将其裁剪至研究范围并对研究区数据进行掩膜，只保留研究区域内的数据。确保已准备好数据文件，数据应存储在指定目录中；检查数据文件命名规范：文件应包含年份信息；确认数据文件格式：所有文件应为 nc 格式(.nc)，具有相同的空间参考系统和分辨率。

(2) 配置代码与参数。

计算 SPEI 及类似的气候指标时，调用 climate_indices 包可以大大减少代码工作量（图 3-2）。

运行以下代码，逐像元计算中国区域内的 SPEI 值，将其保存为.nc 文件，用于接下来的干旱事件识别（图 3-3）。

```

import xarray as xr
import numpy as np
import geopandas as gpd
from shapely.geometry import Point
import pandas as pd
from climate_indices import indices, compute
import os

```

图 3-2 Python 中调用 climate_indices 包计算 SPEI 值示例

```

# 逐像元计算
valid_count = 0
for i in range(n_lats):
    for j in range(n_lons):
        precip_series = precip[:, i, j].values
        pet_series = pet[:, i, j].values
        if np.all(np.isnan(precip_series)) or np.all(np.isnan(pet_series)):
            continue
        precip_series = np.maximum(precip_series, 0)
        pet_series = np.maximum(pet_series, 0)
        try:
            spei = indices.spei(
                precips_mm=precip_series,
                pet_mm=pet_series,
                scale=scale,
                distribution=indices.Distribution.gamma,
                periodicity=compute.Periodicity.monthly,
                data_start_year=start_year,
                calibration_year_initial=start_year,
                calibration_year_final=end_year
            )
            spei_array[:, i, j] = spei
            valid_count += 1
        except:
            pass
    if (i + 1) % 10 == 0:
        print(f"进度: {(i+1)/n_lats*100:.1f}%")
print(f"\n 计算完成: 有效像元 {valid_count}/{n_lats * n_lons}")

```

图 3-3 逐像元计算 SPEI 值代码示例

2. 识别干旱事件及其特征

执行脚本,设定干旱事件阈值为-1.0,最小持续时间为3个月(图3-4)。同样地,可以识别干旱事件的强度相关特征(图3-5)。

```
# 配置参数
spei_file = r'E:\supermap_drought_analyze\results\spei_3_monthly_China.nc'
output_folder = r'E:\supermap_drought_analyze\results\temporal'

# 干旱阈值
drought_threshold = -1.0    # 基础干旱阈值
min_duration = 3           # 最小持续时间(月)
```

图 3-4 设定干旱事件阈值及最小持续时间

```
# 逐像元识别干旱事件(在整个时间段上)
print("\n正在识别干旱事件...")
drought_events = []
for i in range(n_lats):
    for j in range(n_lons):
        series = spei[:, i, j]
        # 跳过无效像元
        if np.all(np.isnan(series)):
            continue
        # 识别干旱事件
        is_drought = series < drought_threshold
        in_drought = False
        start_time = None
        current_duration = 0
        current_spei_values = []
        for t in range(n_times):
            if np.isnan(series[t]):
                if in_drought and current_duration >= min_duration:
                    # 计算干旱强度特征
                    severity = np.sum(np.abs(current_spei_values)) # 累积强度
                    mean_intensity = np.mean(np.abs(current_spei_values)) # 平均强度
                    max_intensity = np.max(np.abs(current_spei_values)) # 最大强度

                    drought_events.append({
                        'lat_idx': i,
                        'lon_idx': j,
                        'start_time': start_time,
                        'duration': current_duration,
                        'severity': severity,
                        'mean_intensity': mean_intensity,
                        'max_intensity': max_intensity
                    })
                in_drought = False
                current_duration = 0
                current_spei_values = []
```

图 3-5 识别干旱事件发生次数、持续时间及强度等特征

```

        continue

    if is_drought[t]:
        if not in_drought:
            start_time = times[t]
            in_drought = True
            current_duration += 1
            current_spei_values.append(series[t])
        else:
            if in_drought and current_duration >= min_duration:
                severity = np.sum(np.abs(current_spei_values))
                mean_intensity = np.mean(np.abs(current_spei_values))
                max_intensity = np.max(np.abs(current_spei_values))

                drought_events.append({
                    'lat_idx': i,
                    'lon_idx': j,
                    'start_time': start_time,
                    'duration': current_duration,
                    'severity': severity,
                    'mean_intensity': mean_intensity,
                    'max_intensity': max_intensity
                })

            in_drought = False
            current_duration = 0
            current_spei_values = []
            # 处理最后一个干旱事件
    if in_drought and current_duration >= min_duration:
        severity = np.sum(np.abs(current_spei_values))
        mean_intensity = np.mean(np.abs(current_spei_values))
        max_intensity = np.max(np.abs(current_spei_values))
        drought_events.append({
            'lat_idx': i,
            'lon_idx': j,
            'start_time': start_time,
            'duration': current_duration,
            'severity': severity,
            'mean_intensity': mean_intensity,
            'max_intensity': max_intensity
        })

    if (i + 1) % 10 == 0:
        print(f"识别进度: {(i+1)/n_lats*100:.0f}%")

print(f"共识别干旱事件: {len(drought_events)} 个")

```

图 3-5(续)