



MIHAIL VACARCIUC

Licenciado em Ciências da Engenharia Eletrotécnica e de
Computadores

SEGMENTAÇÃO AVANÇADA DE FAIXAS DE GESTÃO DE COMBUSTÍVEL

DISSERTAÇÃO PARA A OBTENÇÃO DO GRAU DE MESTRE EM
ENGENHARIA ELETROTÉCNICA E DE COMPUTADORES

MESTRADO EM ENGENHARIA ELETROTÉCNICA E DE COMPUTADORES

Universidade NOVA de Lisboa
Setembro, 2022



SEGMENTAÇÃO AVANÇADA DE FAIXAS DE GESTÃO DE COMBUSTÍVEL

DISSERTAÇÃO PARA A OBTENÇÃO DO GRAU DE MESTRE EM
ENGENHARIA ELETROTÉCNICA E DE COMPUTADORES

MIHAIL VACARCIUC

Licenciado em Ciências da Engenharia Eletrotécnica e de Computadores

Orientador: André Damas Mora
Professor Auxiliar, NOVA University Lisbon

Segmentação avançada de Faixas de gestão de combustível

Copyright © Mihail Vacarciuc, Faculdade de Ciências e Tecnologia, Universidade NOVA de Lisboa.

A Faculdade de Ciências e Tecnologia e a Universidade NOVA de Lisboa têm o direito, perpétuo e sem limites geográficos, de arquivar e publicar esta dissertação através de exemplares impressos reproduzidos em papel ou de forma digital, ou por qualquer outro meio conhecido ou que venha a ser inventado, e de a divulgar através de repositórios científicos e de admitir a sua cópia e distribuição com objetivos educacionais ou de investigação, não comerciais, desde que seja dado crédito ao autor e editor.

A todos os que me apoiaram neste meu feito.

AGRADECIMENTOS

Em primeiro lugar, gostaria de agradecer ao professor André Damas Mora pelo seu esforço, apoio, disponibilidade, conselhos e orientação durante a realização desta dissertação. Agradeço-lhe também por me apresentar o tema da detecção remota e por me dar a oportunidade de realizar esta dissertação. Agradeço por me ter integrado no grupo de investigação CA3 e também lhes agradeço por me terem acolhido. Gostaria de agradecer especialmente ao João Pereira Pires pelo seu apoio incondicional, disponibilidade total e ajuda durante a dissertação.

Gostaria de agradecer também à Faculdade de Ciências e Tecnologia da Universidade Nova de Lisboa pela formação e oportunidade que me proporcionou.

Agradeço a todos os meus familiares, especialmente aos meus pais, Mariana e Alexandru, pelo apoio, confiança, força e por acreditarem sempre em mim. Agradeço o esforço e o sacrifício que fizeram para me proporcionar um bom bem-estar e oportunidades de escolha na minha carreira. À minha irmã, Alexandra, pela sua boa disposição e por se orgulhar de mim.

Por fim, agradeço a todos os meus colegas, especialmente aos meus grandes amigos, Guilherme, João Palma e Gonçalo, pelos momentos inesquecíveis dos últimos cinco anos e pela vossa grande amizade.

RESUMO

A constante evolução e investimento na aquisição de informação sobre o nosso planeta através de satélites resultou numa quantidade enorme de dados que é humanamente impossível de processar, analisar e extrair conclusões. Para solucionar este problema, surgiram novas tecnologias e métodos para realizar essas tarefas, tirando proveito da informação adquirida. Nesta dissertação serão estudadas algumas destas técnicas.

O Instituto de Conservação da Natureza e das Florestas (ICNF) definiu a implementação da rede primária de Faixas de Gestão de Combustível (FGC) como estratégia para prevenção e combate a incêndios florestais. Nessas faixas são realizadas intervenções para reduzir a vegetação nas mesmas, servindo como barreira à progressão do fogo. No entanto, devido ao elevado tamanho da rede, é difícil de a gerir e supervisionar. Para resolver este problema, foram criadas técnicas automáticas de deteção de manutenções nas faixas usando imagens do *Sentinel-2*.

Nesta dissertação foram estudadas duas metodologias para a deteção de tratamentos nas faixas. A primeira é uma metodologia baseada numa aprendizagem supervisionada e a segunda baseia-se no teste estatístico *Welch's t-Test*, metodologias de deteção com uma abordagem ao nível do píxel e ambas têm um pré-processamento comum, o co-registo, que consiste num alinhamento geométrico das imagens de satélite para garantir que o mesmo píxel represente o mesmo local nas várias imagens. Foi dado maior ênfase à segunda metodologia, identificando possíveis maneiras de melhorar o seu desempenho, através da variação de parâmetros e a comparação dos resultados obtidos com um conjunto de dados de validação. Por fim, para esta metodologia foi proposto e avaliado um pós-processamento para reduzir o número de classificações isoladas.

Palavras-chave: Deteção Remota, Faixas de Gestão de Combustível, *Sentinel-2*, Aprendizagem Supervisionada, *Welch's t-Test*, Teste de Normalidade, Co-Registo, Pós-Processamento

ABSTRACT

The continuous evolution and investment in acquiring information about our planet through satellite technology has resulted in an overwhelming amount of data that is beyond human capability to process, analyze, and derive meaningful conclusions from. To tackle this challenge, new methods and technologies have emerged to perform these tasks, leveraging the collected information. In this dissertation will be studied some of these techniques.

The Institute of Nature and Forest Conservation (ICNF) has established the implementation of the primary network of Fuel Breaks (FGC) as a strategy for preventing and fighting forest fires. Interventions are carried out in these zones to reduce vegetation, and their primary objective is to inhibit the spread of fire. However, due to the vast size of the network, it is challenging to manage and supervise effectively. To solve this problem, automatic techniques have been developed for maintenance detection in the Fuel Breaks using Sentinel-2 images.

In this dissertation, two methodologies were studied for detecting treatments in the Fuel Breaks. The first is a methodology based on supervised learning, while the second is based on the Welch's t-Test statistical test. They are detection methodologies with a pixel-level approach and both have a common pre-processing, co-registration, which consists of a geometric alignment of satellite images to ensure that the same pixel represents the same location across multiple images. Greater attention was given to the second methodology, with the aim of identifying potential ways to enhance its performance by varying parameters and comparing the results obtained with a validation dataset. Finally, a post-processing was proposed and evaluated for this methodology to minimize the number of isolated classifications.

Keywords: Remote Sensing, Fuel Break, Sentinel-2, Supervised Learning, Welch's t-Test, Normality Test, Co-Registration, Post-Processing

ÍNDICE

Índice de Figuras	x
Índice de Tabelas	xiii
Siglas	xiv
1 Introdução	1
1.1 Motivação	1
1.2 Faixas de Gestão de Combustível	2
1.3 Objetivos	3
1.4 Organização do Documento	4
2 Estado de Arte	5
2.1 Detecção Remota	5
2.1.1 Radiação Eletromagnética	6
2.1.2 Resoluções	8
2.1.3 Satélites de Observação da Terra	8
2.1.4 Sistemas de Informação Geográfica	11
2.2 Carta de Uso e Ocupação do Solo	13
2.3 Classificação por píxel contra classificação baseada em objeto	14
2.4 Segmentação	14
2.4.1 Segmentação Superpíxel	14
2.4.2 Segmentação Agrupamento Difuso	18
2.4.3 Algoritmo de segmentação multiresolução	19
2.5 Classificação	19
2.5.1 Algoritmos de Aprendizagem Automática	21
2.5.2 Algoritmos de Testes Estatísticos	22
3 Trabalho Desenvolvido	24
3.1 Áreas de Estudo	24

3.2	Metodologias para a detecção por píxel de tratamentos nas FGC	26
3.2.1	Pré-processamento das observações de satélite	26
3.2.2	Metodologia 1: Aprendizagem Supervisionada	28
3.2.3	Metodologia 2: Classificação com o teste estatístico <i>Welch's t-Test</i>	29
3.2.4	Primeira análise das metodologias	31
3.3	Testes à Metodologia 2: Classificação com o teste estatístico <i>Welch's t-Test</i>	33
3.3.1	Variação do nível de significância (α)	33
3.3.2	Teste de normalidade	34
3.3.3	Pós-processamento	35
4	Apresentação e Discussão dos Resultados	37
4.1	Áreas de Estudo	37
4.1.1	Fundão	38
4.1.2	Serra dos Candeeiros	39
4.1.3	Sertão	41
4.1.4	Seia	42
4.2	Resultados obtidos	42
4.2.1	Variação do nível de significância (α)	43
4.2.2	Teste de normalidade	46
4.2.3	Sumário	58
4.2.4	Pós-processamento	59
5	Conclusões e Trabalho Futuro	62
5.1	Conclusão	62
5.2	Trabalho Futuro	63
	Bibliografia	64

ÍNDICE DE FIGURAS

1.1	Valores Médios da Área ardida (ha) e o número de incêndios florestais entre 2010 a 2020 (fonte: PORDATA)	2
1.2	Secção transversal de uma Faixa de Gestão de Combustível (FGC) (DPFVAP, 2014).	3
2.1	Visão geral do processo de deteção remota. Adaptado de (J. B. Campbell e Wynne, 2011).	6
2.2	Espectro Eletromagnético e as respetivas regiões. Retirada de (“Espectro Eletromagnético”, 2017).	7
2.3	Espectro da radiação solar na Terra sem o efeito da atmosfera (amarelo) e com o efeito da atmosfera (vermelho). Retirada de (Rohde e Costa, 2020).	7
2.4	Primeira imagem televisiva da Terra captada pelo satélite TIROS-1. Retirada de (NASA, 2009).	9
2.5	Família de satélites <i>Sentinel</i> . Retirada de (ESA, s.d.-b).	10
2.6	Representação de dados em formato vetores, à esquerda, e em formato <i>raster</i> , à direita, de uma imagem de satélite, no centro. Retirada de (Najar et al., 2010).	12
2.7	Representação da (a) imagem da superfície terrestre e os (b) respetivos polígonos apresentados na COS2018v1.0. Retirada de (DGT, 2022).	13
2.8	Superpíxeis gerados por algoritmos baseados em bacias hidrografias, com 400 superpíxeis para a parte superior esquerda da imagem e 1200 superpíxeis para a parte inferior direita (Stutz et al., 2018).	15
2.9	Superpíxeis gerados por algoritmos baseados em grafos, com 400 superpíxeis para a parte superior esquerda da imagem e 1200 superpíxeis para a parte inferior direita (Stutz et al., 2018).	16
2.10	Superpíxeis gerados por algoritmos baseados em <i>clustering</i> , com 400 superpíxeis para a parte superior esquerda da imagem e 1200 superpíxeis para a parte inferior direita (Stutz et al., 2018).	17
2.11	Representação de uma Rede Neuronal Artificial com uma camada oculta. Retirado de (“Neural Network”, 2015).	22

3.1	Área de Estudo: Marisol, concelho Almada. A área a vermelho representa a faixa de intervenção e a área a verde representa uma zona de vegetação. . .	25
3.2	Área de estudo de FGC na Rede Primária de Faixas de Gestão de Combustível (RPFGC).	25
3.3	Fluxograma do pré-processamento.	27
3.4	Fluxograma da metodologia 1.	28
3.5	Fluxograma da metodologia 2.	29
3.6	Representação dos valores da imagem de <i>Scene Classification</i> (SCL). Retirado de (ESA, s.d.-a).	30
3.7	Legenda utilizada para a representação dos resultados obtidos.	32
3.8	Resultados obtidos por ambas as metodologias para a área do Marisol. . . .	32
3.9	Resultados obtidos para Marisol em 2018 para níveis de significância (α). .	33
3.10	Exemplo de uma aplicação de um filtro de mediana (3x3) numa imagem com os valores do <i>Welch's t-Test</i> para a série temporal dos píxeis interiores. . . .	35
3.11	Fluxograma da metodologia 2 com pós-processamento.	36
4.1	Legenda utilizada para a representação dos meses classificados.	38
4.2	Classificação anual quanto às intervenções numa faixa da zona do Fundão para o ano 2017.	39
4.3	Classificação anual quanto às intervenções numa faixa da zona da Serra dos Candeeiros para o ano 2017.	40
4.4	Classificação anual quanto às intervenções numa faixa da zona de Sertã para o ano 2017.	41
4.5	Classificação anual quanto às intervenções numa faixa da zona de Seia em formato vetorial para o ano 2017.	42
4.6	Resultados obtidos para a zona do Fundão para vários valores de α	43
4.7	Resultados obtidos para a zona da Serra dos Candeeiros para vários valores de α	44
4.8	Resultados obtidos para a zona da Sertã para vários valores de α	45
4.9	Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,05$. .	46
4.10	Resultados obtidos para a zona do Fundão para $\alpha = 0,05$ variando a janela temporal.	47
4.11	Resultados obtidos para a zona do Fundão para $\alpha = 0,005$ variando a janela temporal.	48
4.12	Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,005$. .	49
4.13	Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,05$. .	50
4.14	Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,005$. .	50
4.15	Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,0005$. .	50
4.16	Resultados obtidos para a zona do Fundão para $\alpha = 0,05$ variando a janela temporal.	51

4.17 Resultados obtidos para a zona do Fundão para $\alpha = 0,005$ variando a janela temporal.	52
4.18 Resultados obtidos para a zona do Fundão para $\alpha = 0,0005$ variando a janela temporal.	53
4.19 Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,05$	54
4.20 Resultados obtidos para a zona do Fundão para $\alpha = 0,05$ variando a janela temporal.	55
4.21 Resultados obtidos para a zona do Fundão para $\alpha = 0,005$ variando a janela temporal.	56
4.22 Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,005$	57
4.23 Resultados obtidos para a zona do Fundão aplicando um filtro de mediana.	59
4.24 Resultados obtidos para a zona da Serra dos Candeeiros aplicando um filtro de mediana.	60
4.25 Resultados obtidos para a zona da Sertã aplicando um filtro de mediana.	61

ÍNDICE DE TABELAS

2.1	Bandas Espectrais do <i>Sentinel 2</i> , dados retirados de (ESA, s.d.-d).	11
4.1	Número de observações usadas para cada região.	42
4.2	Exatidão para cada valor do α de cada região	45
4.3	Exatidão para a zona do Fundão variando a janela temporal.	49
4.4	Exatidão para a zona da Serra dos Candeeiros variando a janela temporal. .	54
4.5	Exatidão para a zona da Sertã variando a janela temporal.	57
4.6	Exatidão obtida variando a janela temporal.	58

SIGLAS

CIE	Comissão Internacional em Iluminação 17
COS	Carta de Uso e Ocupação do Solo 5, 13, 27, 30, 34, 46, 63
COSsim	Carta de Ocupação do Solo Simplificada 13
DGT	Direção-Geral do Território 13
ESA	Agência Espacial Europeia 8, 9, 26
ExG	Excess of Green Index 26, 28
ExR	Excess of Red Index 26, 28
FCM	Agrupamento difuso C-means 18
FGC	Faixa de Gestão de Combustível x, xi, 1, 2, 3, 24, 25, 26, 27, 28, 30
FIC	Faixa de interrupção de combustível 2
FRC	Faixa de redução de combustível 2
GEE	Google Earth Engine 26
ICNF	Instituto de Conservação da Natureza e das Florestas 1, 2, 24, 37
NASA	Administração Nacional da Aeronáutica e Espaço 8, 10
NDVI	Normalized Difference Vegetation Index 26, 30
OBIA	análise de imagens baseadas em objetos 19
QGIS	Quantum Geographic Information System 37, 62, 63
RNA	Redes Neurais Artificiais 21, 22, 28, 62
RPFGC	Rede Primária de Faixas de Gestão de Combustível xi, 2, 3, 24, 25, 27

SIG Sistema de Informação Geográfica [11](#), [12](#)

USGS Serviço Geológico dos Estados Unidos [8](#)

INTRODUÇÃO

1.1 Motivação

Nos últimos anos, tem havido um grande número de incêndios florestais de grande escala, tais como os ocorridos na Austrália e Califórnia em 2020, e aqueles que ocorrem recorrentemente nas florestas da Sibéria durante o verão. Como consequência, ficam milhões de hectares de área ardida, prejudicando toda a fauna e flora da zona, podendo levar à extinção de certas espécies. Além disso, há uma libertação de grande quantidade de gases de efeito estufa na atmosfera, como o dióxido de carbono (CO_2).

Em Portugal, todos os verões ocorrem incêndios florestais, sendo que 2017 foi o pior ano até hoje, com os incêndios de Pedrogão Grande, considerado o mais mortífero da história do país. Observando a figura 1.1, que representa os valores médios de área ardida (em hectares) e do número de incêndios florestais registados entre 2010 e 2020, constata-se que Portugal é o país com os valores mais elevados da União Europeia.

Para combater os incêndios florestais, o trabalho dos bombeiros e de outros membros competentes, como a Proteção Civil, é muito importante, mas apresenta elevado risco de vida. Assim, é essencial que existam um conjunto de medidas preventivas para aumentar a eficácia no combate aos incêndios. Em Portugal, o Instituto de Conservação da Natureza e das Florestas (ICNF) é responsável por criar e estabelecer essas medidas preventivas, sendo que a implementação das FGC faz parte dessas medidas.

O aumento do número de satélites de Observação da Terra e o acesso gratuito a essas observações permitiram desenvolver soluções que melhoram a prevenção de incêndios florestais. Com o passar do tempo, as FGC necessitam de manutenção, sendo essencial uma vigilância constante para avaliar o seu estado. Essa vigilância pode ser implementada através de observações periódicas realizadas por satélites.

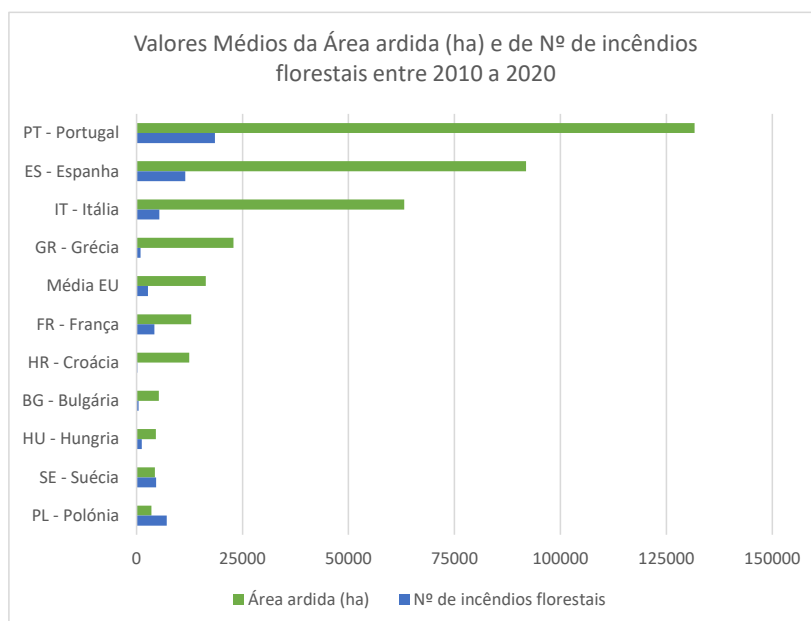


Figura 1.1: Valores Médios da Área ardida (ha) e o número de incêndios florestais entre 2010 a 2020 (fonte: PORDATA)

1.2 Faixas de Gestão de Combustível

Como medida de prevenção de incêndios florestais, o ICNF criou a Rede Primária de Faixa de Gestão de Combustível (RPFGC), cujas principais funções são reduzir a propagação de incêndios de grande dimensão, reduzir o efeito do alastrar dos mesmos, proteger as vias de comunicação e outras infraestruturas e isolar possíveis fontes de ignição. De uma forma geral, esta rede tem como propósito permitir um maior controlo dos fogos durante o seu combate.

Para implementar as FGC, são consideradas várias condições da região, como o elevado risco meteorológico, o histórico de incêndios, a predominância de eucaliptos e o declive médio.

No que diz respeito às especificações técnicas, a ICNF estabeleceu que as FGC devem ter uma largura mínima de 125 metros e delimitar compartimentos entre 500 e 10.000 hectares. As FGC são divididas em duas zonas: a Faixa de interrupção de combustível (FIC) e a Faixa de redução de combustível (FRC). Na FIC, a vegetação é totalmente removida e está localizada perto da rede viária da faixa, enquanto na FRC é removida apenas a vegetação superficial, como arbustos e plantas herbáceas, e as copas das árvores devem respeitar um distanciamento de 2 m na zona exterior da faixa e de 4 m para a zona mais próxima do centro da faixa, conforme ilustrado na figura 1.2 (DPFVAP, 2014).

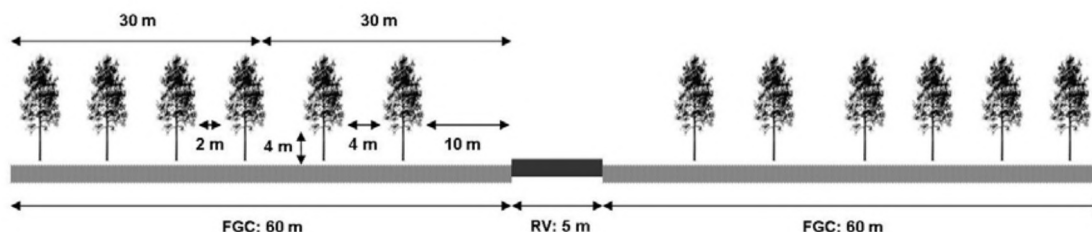


Figura 1.2: Secção transversal de uma FGC (DPFVAP, 2014).

1.3 Objetivos

Esta dissertação contribui para o estudo e análise de duas metodologias com uma abordagem ao nível do píxel para detetar tratamentos nas FGC. Ambas as metodologias geram imagens *raster* como resultado final, representando a classificação com o número do mês em que o tratamento foi identificado. Adicionalmente, foi implementado um pós-processamento para "limpar" os resultados obtidos, produzindo resultados finais nítidos sem classificações isoladas.

Numa primeira fase, foram seleccionadas as áreas de estudo a serem utilizadas nesta dissertação e foram recolhidas as imagens do *Sentinel-2* que abrangem essas zonas para um determinado ano. É importante que as áreas de estudo estejam incluídas na RPFGC e que se tenha conhecimento dos tratamentos que ocorreram lá durante esse ano, a fim de ser possível fazer uma comparação com os resultados obtidos.

Ambas as metodologias foram implementadas em *Python*. Como o pré-processamento é comum às duas metodologias, foi implementado um *script* individual para este efeito. Para a implementação da primeira metodologia, foram estudados métodos de aprendizagem supervisionada, com foco em Redes Neurais Artificiais. Para a implementação da segunda metodologia, foram estudados testes estatísticos, com foco no *Welch's t-Test*.

Após obter os primeiros resultados, foi realizada uma análise preliminar das metodologias, concluindo-se que a segunda metodologia obteve melhores resultados do que a primeira. Por isso, foi dado maior foco à segunda metodologia. Identificaram-se os principais parâmetros que afetam o desempenho da segunda metodologia e testou-se a variação dos mesmos. Para tirar conclusões acerca do desempenho da metodologia, os resultados obtidos foram comparados com um conjunto de dados de validação criado a partir das áreas de estudo.

Por fim, implementou-se um pós-processamento dos mapas gerados pelas duas metodologias e compararam-se os resultados com as áreas de estudo.

1.4 Organização do Documento

A dissertação está organizada em cinco capítulos, que incluem os seguintes conteúdos:

- **Capítulo 1:** : É apresentada a motivação, os objetivos desta dissertação e uma breve explicação sobre as Faixas de Gestão de Combustível;
- **Capítulo 2:** São apresentados os conceitos teóricos que servem de base a esta dissertação, explicando o conceito de detecção remota, onde são abordados os temas de radiação eletromagnética, resolução de imagem, as missões de observação da Terra e de sistemas de informação geográfica. Neste capítulo também se encontra o estado da arte de técnicas e métodos de segmentação e classificação usados em projetos de detecção remota;
- **Capítulo 3:** : É descrito o trabalho desenvolvido, apresentando as metodologias de classificação e como foram implementadas. São descritas as alterações e testes executados na metodologia de classificação pelo teste estatístico, *Welch's t-Test*, e o pós-processamento implementado;
- **Capítulo 4:** É apresentado o conjunto de dados de validação e os resultados obtidos das classificações;
- **Capítulo 5:** São apresentadas as conclusões obtidas nesta dissertação e quais são os possíveis trabalhos futuros

ESTADO DE ARTE

Neste capítulo serão abordados os aspectos teóricos relevantes para o desenvolvimento da dissertação. Será definido os conceitos de detecção remota e Carta de Uso e Ocupação do Solo (COS). Foi levantado o estado de arte de temas relevantes e adequados para a realização do objetivo proposto, para tal são apresentadas as principais diferenças das abordagens a píxel e ao objeto, os algoritmos de segmentação e por fim os tipos de classificadores apropriados para este caso.

2.1 Detecção Remota

O conceito de detecção remota está relacionado com o processo de aquisição de dados de um objeto ou fenómeno à distância, sem a necessidade de interagir fisicamente com o objeto de interesse ou de se deslocar para o local do mesmo. Inicialmente, este processo era realizado através da captação de imagens aéreas tiradas de aviões. Mais tarde, em 1957, foi lançado o primeiro satélite para o espaço, *Sputnik 1*. Com o forte investimento em tecnologias de satélite e de sensores, em 1972, foi lançado o satélite Landsat 1, considerado como o primeiro satélite para estudar e monitorizar a superfície terrestre. A partir dessa data, houve uma grande evolução em todas as tecnologias relacionadas com detecção remota. Atualmente, existem projetos como o *Sentinel* (ESA) e o *Landsat* (USGS/NASA), que disponibilizam regularmente dados da superfície terrestre gratuitamente.

Para o processo de detecção remota, podem-se considerar quatro etapas (Figura 2.1). A primeira consiste na definição do objeto físico que se pretende estudar/observar, podendo ser composto por edifícios, vegetação, solo, corpos de água e outros. Dessa forma, existem várias áreas científicas que podem ter interesse nesse estudo, como sejam o planeamento e gestão florestal, o urbanismo, a geologia e a cartografia. A segunda etapa consiste na captação de radiação eletromagnética emitida ou refletida pela superfície terrestre através de sensores. Na terceira etapa, o objetivo é analisar e interpretar os dados obtidos na segunda etapa para extrair informação e permitir a sua visualização. Por fim, a última etapa consiste em utilizar a informação recolhida para encontrar soluções para problemas práticos.

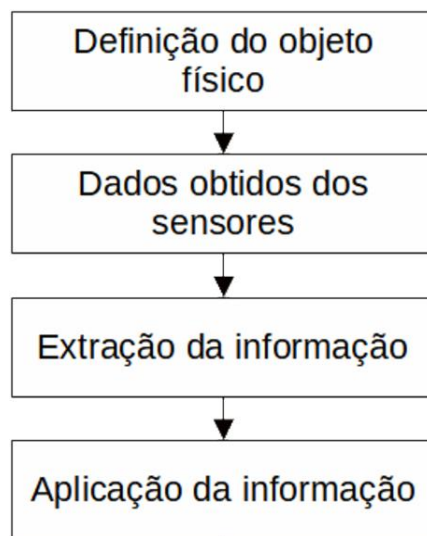


Figura 2.1: Visão geral do processo de detecção remota. Adaptado de (J. B. Campbell e Wynne, 2011).

2.1.1 Radiação Eletromagnética

Dado que os sensores de detecção remota são dispositivos que medem a radiação eletromagnética, é relevante abordar este tema com maior detalhe. Todos os objetos com uma temperatura superior a zero absoluto (0 Kelvin) emitem radiação eletromagnética, e também refletem radiação emitida por outros objetos. As duas principais fontes naturais de radiação eletromagnética utilizadas na detecção remota são o Sol e a Terra, pois os sensores captam a radiação presente na Terra, sendo que a maioria dela é radiação emitida pelo Sol e refletida pela superfície terrestre.

A radiação eletromagnética é constituída por um campo elétrico perpendicular a um campo magnético, que varia perpendicularmente com a direção da propagação, não necessitando de um meio para se propagar, podendo propagar-se no vácuo à velocidade da luz. A radiação eletromagnética pode ser caracterizada pelas seguintes propriedades: comprimento de onda, frequência, amplitude e a velocidade de propagação. A velocidade de propagação é obtida multiplicando o comprimento de onda pela frequência da mesma.

A luz visível é a radiação eletromagnética mais conhecida, mas representa apenas uma pequena porção do espectro eletromagnético. Por esta razão, é necessário analisar os outros intervalos do espectro, uma vez que estes podem apresentar comportamentos diferentes da luz visível. O espectro eletromagnético está dividido em subdivisões consoante as características de cada região. As subdivisões do espectro predefinidas são as seguintes (Figura 2.2): ondas de rádio, micro-ondas, infravermelhos, visível, ultravioleta, raios-X e raios gama.

Na detecção remota, as ondas eletromagnéticas utilizadas variam em comprimento de onda desde a região do visível (cerca de $0,4 \mu m$) até às micro-ondas (cerca de $1 mm$). No entanto, nem todas as frequências do espectro são utilizadas, uma vez que cada aplicação

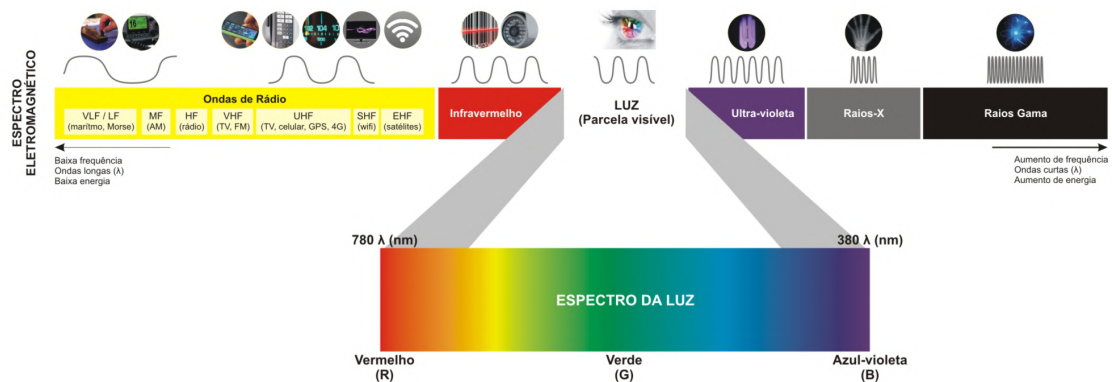


Figura 2.2: Espectro Eletromagnético e as respectivas regiões. Retirada de (“Espectro Eletromagnético”, 2017).

requer uma gama específica de comprimentos de onda, dependendo das características do alvo que se pretende detetar e das condições ambientais em que se encontra. Os sensores óticos trabalham na faixa do visível e do infravermelho próximo, enquanto os sensores de radar utilizam frequências na faixa das micro-ondas. Além disso, estes sensores estão limitados às janelas atmosféricas, as radiações que não são totalmente absorvidas pela atmosfera terrestre. A figura 2.3 ilustra a distribuição do espectro de radiação eletromagnética do Sol antes de entrar na atmosfera, representada a amarelo, e após entrar na atmosfera, representada a vermelho, e mostra também quais são os principais gases responsáveis pela absorção da radiação na atmosfera, como o vapor de água (H_2O), o dióxido de carbono (CO_2) e o ozono (O_3).

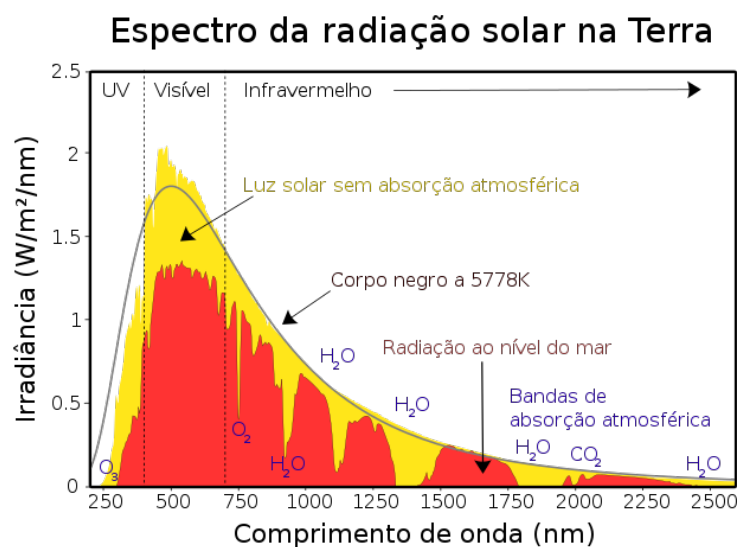


Figura 2.3: Espectro da radiação solar na Terra sem o efeito da atmosfera (amarelo) e com o efeito da atmosfera (vermelho). Retirada de (Rohde e Costa, 2020).

2.1.2 Resoluções

A resolução das imagens obtidas por detecção remota pode apresentar vários tipos diferentes, por isso é importante entender melhor cada tipo de resolução para garantir a qualidade dos dados captados. Os tipos de resolução em causa são a resolução espacial, espectral, temporal e radiométrica (Stevens et al., 2012).

- **Resolução espacial:** nesta resolução os píxeis das imagens representam o tamanho real. Por exemplo, uma resolução de 50 metros significa que cada píxel representa uma área de $50 \times 50 \text{ m}^2$ da superfície terrestre. Uma resolução alta significa uma representação mais precisa e detalhada, permitindo a detecção de objetos pequenos.
- **Resolução espectral:** existem três aspetos que caracterizam esta resolução, o número de bandas espectrais que o sensor tem, o tamanho do intervalo de comprimento de onda das bandas e a posição das bandas no espectro eletromagnético. Um sensor tem uma melhor resolução espectral quando apresenta mais bandas em diferentes posições do espectro e com intervalos menores de comprimento de onda.
- **Resolução temporal:** descreve o intervalo de tempo entre as aquisições de dados efetuadas pelo sensor para um determinado local. Quanto menor for esse intervalo, melhor será a resolução temporal. Esta resolução é importante para a detecção de alterações na superfície terrestre ao longo do tempo.
- **Resolução radiométrica:** está relacionada com o nível de sensibilidade do sensor em detetar variações na energia radiante para a mesma banda. A resolução é maior quando existem mais valores para representar as variações ocorridas. É comum os sensores terem 8, 10 ou 12 bits de resolução radiométrica.

2.1.3 Satélites de Observação da Terra

Hoje em dia, existem vários organismos e agências governamentais dedicados à exploração e operação do espaço. As principais agências espaciais que mais avanços alcançaram em satélites de observação da Terra são a Agência Espacial Europeia (ESA) com o programa *Copernicus*, e a Administração Nacional da Aeronáutica e Espaço (NASA) juntamente com Serviço Geológico dos Estados Unidos (USGS) com o programa *Landsat*. Esses programas disponibilizam gratuitamente dados de boa qualidade da superfície terrestre.

Em 1960 foi realizado o lançamento do satélite TIROS-1 do programa TIROS (*Television Infrared Observation Satellite*), o primeiro satélite meteorológico. Este programa foi o primeiro passo da NASA para estudar e avaliar a possibilidade de uso de satélites para o estudo e observação da Terra. Como tal, inicialmente a prioridade foi desenvolver um sistema capaz de captar informações meteorológicas (NASA, 2016). O TIROS-1 foi equipado por duas câmaras televisivas, uma de baixa e outra de alta resolução, e fita magnética para armazenar as imagens captadas. Este satélite só operou durante 78 dias, mas demonstrou que as tecnologias de satélite são importantes e úteis na observação da

Terra. Em 1972, foi lançado o satélite Landsat 1, com capacidade de observar e recolher dados da superfície terrestre.



Figura 2.4: Primeira imagem televisiva da Terra captada pelo satélite TIROS-1. Retirada de (NASA, 2009).

2.1.3.1 Programa *Copernicus*

Em 2014, o programa *Copernicus* foi aprovado pela União Europeia como o programa de observação da Terra da Europa, sucedendo ao programa anterior, *Global Monitoring for Environment and Security* (GMES). A agência responsável por este programa é a ESA, que é uma organização intergovernamental constituída por 22 Estados-membros, incluindo Portugal. O principal objetivo deste programa é fornecer informação útil para a gestão ambiental, perceber e mitigar os efeitos das alterações climáticas e preservar a segurança civil.

Para cumprir os vários objetivos, foi desenvolvida uma nova família de missões chamada *Sentinel* (Figura 2.5), baseadas em constelações de satélites para fornecer dados de boa qualidade (ESA, s.d.-c). Cada missão foca-se em observações de aspetos específicos da Terra, como a atmosfera, oceanos, solo, entre outros. As finalidades de cada missão são as seguintes:

- ***Sentinel 1***: captação de imagens radar para serviços de monitorização da terra e oceanos durante o dia e noite e independentemente das condições meteorológicas;
- ***Sentinel 2***: captação de imagens óticas de alta resolução, permitindo a monitorização da vegetação, solo, áreas costeiras, entre outros objetos relevantes;
- ***Sentinel 3***: dedicado à observação marítima, utilizando imagens de radar e óticas para medição topográfica da superfície do mar, medições de temperatura da superfície terrestre e oceanos, e deteção da cor dos mesmos;

- **Sentinel 5P:** dedicado a medições atmosféricas, à qualidade do ar, às alterações climáticas, à camada do ozono e à radiação ultravioleta. Foi desenvolvido para preencher a falta de informação deste tipo após a paragem do satélite *Envisat* da NASA;
- **Sentinel 4 e 5:** ambos são dedicados à monitorização da qualidade do ar, observando a atmosfera terrestre, sendo que o *Sentinel 4* terá um maior foco na Europa e o *Sentinel 5* no mundo inteiro;
- **Sentinel 6:** dedicado à medição da altura do nível do mar, por meio de um radar altímetro.

Para o futuro, já estão em preparação os novos objetivos e prioridades para uma segunda geração do *Copernicus*, levando à conclusão das missões em falta (4 e 5), a continuação das existentes e a criação de novas missões para cumprir as metas estabelecidas pela União Europeia, com maior foco nas alterações climáticas e na monitorização do dióxido de carbono (CO_2).

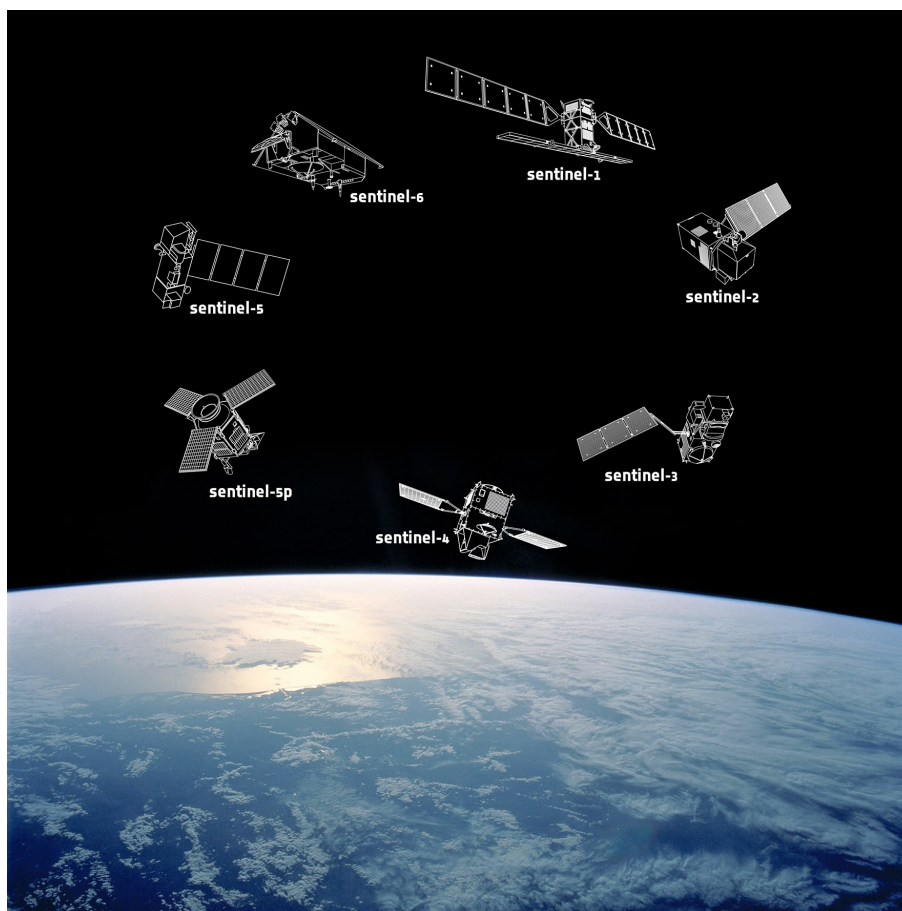


Figura 2.5: Família de satélites *Sentinel*. Retirada de (ESA, s.d.-b).

2.1.3.2 *Sentinel 2*

Para o tema abordado nesta dissertação, o satélite *Sentinel 2* é considerado o mais indicado. Como já foi mencionado anteriormente, as missões *Sentinel* são compostas por constelações de satélites, sendo neste caso uma constelação composta por dois satélites, o 2A e 2B, que se encontram em órbitas desfasadas de 180°, com uma resolução temporal de cerca de 5 dias. O sistema global cobre uma área que vai desde as latitudes 56° sul até 84° norte. Os satélites estão equipados com o *MultiSpectral Instrument* (MSI), que permite recolher informações de 13 bandas espectrais, com resoluções espaciais de 10, 20 e 60 m, como pode ser observado na tabela 2.1. A resolução radiométrica é de 12 bits, possibilitando a representação de valores entre 0 e 4095 (ESA, 2013).

O *Sentinel 2* disponibiliza dados que podem ser utilizados em diversas aplicações relacionadas com a superfície da Terra e zonas costeiras, permitindo a obtenção de informações relevantes para as atividades relacionadas com a agricultura, florestas, tudo relacionado com a vegetação, permitindo também análises de alterações, crescimento de plantas e em catástrofes naturais, tais como incêndios, cheias, entre outros.

Tabela 2.1: Bandas Espectrais do *Sentinel 2*, dados retirados de (ESA, s.d.-d).

Bandas <i>Sentinel-2</i>	Comprimento de onda (μm)	Largura de banda (μm)	Resolução (m)
Banda 1 — Aerossol Costeiro	0.443	0.02	60
Banda 2 — Azul	0.490	0.065	10
Banda 3 — Verde	0.560	0.035	10
Banda 4 — Vermelho	0.665	0.03	10
Banda 5 — Borda Vermelha de Vegetação	0.705	0.015	20
Banda 6 — Borda Vermelha de Vegetação	0.704	0.015	20
Banda 7 — Borda Vermelha de Vegetação	0.783	0.02	20
Banda 8 — NIR	0.842	0.115	10
Banda 8a — Narrow NIR	0.865	0.02	20
Banda 9 — Vapor de água	0.945	0.02	60
Banda 10 — SWIR - Cirrus	1.375	0.03	60
Banda 11 — SWIR	1.610	0.09	20
Banda 12 — SWIR	2.190	0.180	20

2.1.4 Sistemas de Informação Geográfica

Um Sistema de Informação Geográfica (SIG) é uma ferramenta computacional utilizada para capturar, armazenar, manipular, analisar e visualizar dados geográficos, dados que descrevem a localização e as características dos locais. Distinguindo-se assim de outros sistemas de informação e estabelecendo-se como uma tecnologia importante em diversas áreas, uma vez que é mais prático observar os dados geográficos num mapa do

que em tabelas ou outras formas (J. Campbell e Shin, 2011; Chang, 2006). Exemplos concretos de aplicações usando SIG são no planejamento do uso de terrenos, na gestão de recursos naturais, avaliações de risco de catástrofes naturais, gestão e planejamento florestal (Chang, 2006).

Para a representação de espaços geográficos, existem principalmente dois modelos de dados: dados *raster* e vetoriais (Figura 2.6). Os dados *raster* representam a informação através de píxeis e é um modelo de dados conhecido e usado, para além do SIG, na fotografia digital e em vários formatos de imagens digitais, como JPEG, BMP e TIFF, que são baseados no modelo *raster*. Os dados vetoriais utilizam pontos com as respetivas coordenadas para representar vértices de atributos espaciais, de forma a que os mapas se assemelham a polígonos. Os tipos de vetores que existem em SIG são pontos, linhas e polígonos (J. Campbell e Shin, 2011). É possível converter de um modelo para outro, mas normalmente esta conversão envolve a perda de precisão e informação. Por isso, é melhor utilizar os dados no formato que for mais apropriado, sendo que os SIG suportam os dois tipos de dados.

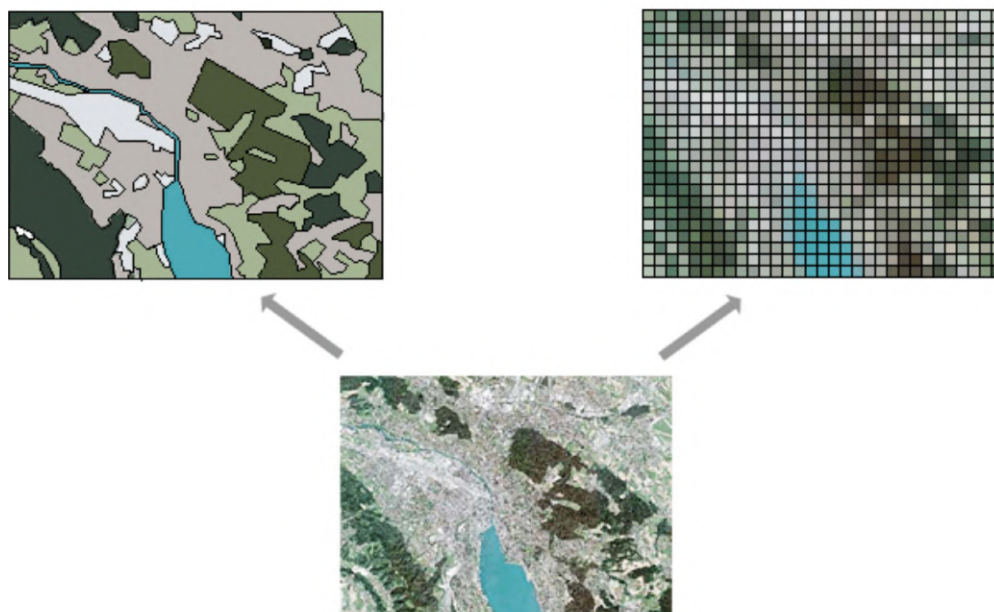


Figura 2.6: Representação de dados em formato vetores, à esquerda, e em formato *raster*, à direita, de uma imagem de satélite, no centro. Retirada de (Najar et al., 2010).

Atualmente, existem vários fornecedores de *software* de SIG, sendo um dos mais importantes a *Environmental Systems Research Institute* (ESRI), que desenvolve e distribui o programa ArcGIS, e a *Precisely*, com o programa MapInfo. Existem também *softwares* de código aberto, como o QGIS, GRASS, entre outros.

2.2 Carta de Uso e Ocupação do Solo

A carta de uso e ocupação do solo (COS) é uma cartografia de polígonos, cuja função é representar de uma forma homogênea unidades de uso e ocupação do solo. Uma unidade de uso e ocupação do solo é definida como qualquer área de terreno superior ou igual a 1 hectare, com uma distância entre linhas superior ou igual a 20 metros e com uma percentagem superior ou igual a 75% de uma determinada classe de uso e ocupação do solo (Caetano et al., 2019).

A primeira COS foi criada em 1990 e desde então tem sido atualizada, como o foi nos anos de 1995, 2007, 2010, 2015 e 2018, sendo desta forma possível fazer análises da evolução temporal da ocupação dos terrenos. A COS é importante para as decisões de ordenamento nacional do território, para a monitorização ambiental e como ferramenta de auxílio no planeamento político, económico e social. O mapa mais recente, COS2018, apresenta 83 classes de ocupação, mais 35 do que a versão anterior, COS2015. A nova versão foi produzida pela atualização da versão anterior, atualizando apenas as áreas que sofreram alterações de território e corrigindo os erros identificados da versão anterior.

A Direção-Geral do Território (DGT) é a entidade responsável pelo desenvolvimento da COS. Para além da COS, a DGT desenvolveu um novo produto, Carta de Ocupação do Solo Simplificada (COSsim), com o intuito de fornecer informações complementares à COS. De momento, já foram produzidas três COSsim para os anos 2018, 2020 e 2021, e a intenção é a de produzir uma COSsim anualmente. Esses novos mapas são produzidos utilizando métodos baseados em inteligência artificial com base em séries temporais de imagens de satélite *Sentinel-2*, e são compostos por três níveis de detalhe temático crescente, com 6, 9 e 15 classes por nível. O principal destaque desses novos mapas é que são produzidos em formato *raster* com uma resolução espacial de 10 m.

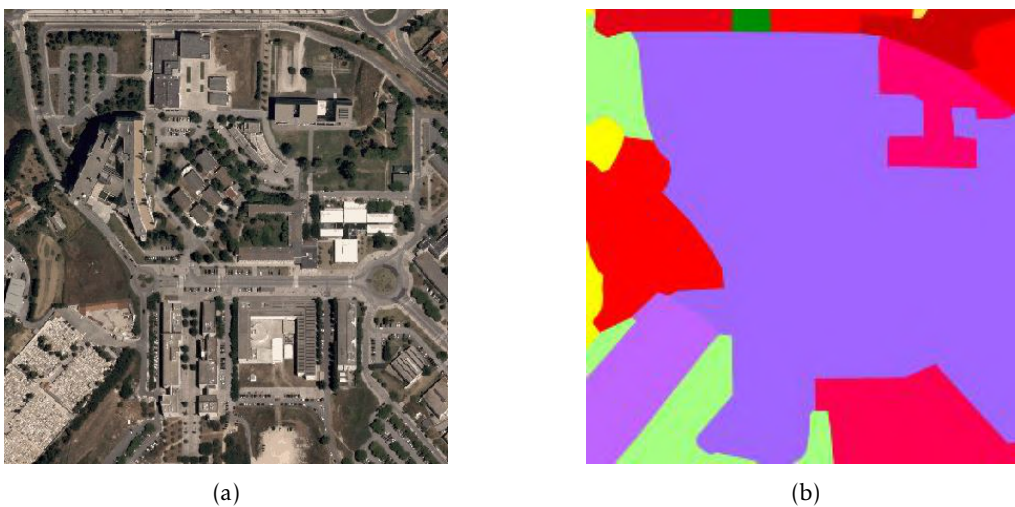


Figura 2.7: Representação da (a) imagem da superfície terrestre e os (b) respetivos polígonos apresentados na COS2018v1.0. Retirada de (DGT, 2022).

2.3 Classificação por píxel contra classificação baseada em objeto

A abordagem por píxel e a abordagem por objeto são duas metodologias utilizadas em processamento de imagens. Na abordagem por píxel, os métodos de classificação operam diretamente com os píxeis, enquanto que na abordagem por objeto, primeiro é feita uma separação dos píxeis em objetos homogêneos através de algoritmos de segmentação e, posteriormente, são usados os métodos de classificação. Uma das principais diferenças entre ambas as abordagens é que na abordagem por objeto cada objeto consegue ter um conjunto de atributos, para além das propriedades espectrais que o caracterizam, nomeadamente a nível espacial, ao nível da textura e contextual. Por outro lado, na abordagem por píxel apenas as propriedades espectrais de cada píxel são analisadas. Outra diferença importante é que a abordagem por píxel costuma produzir píxeis isolados em imagens de alta resolução, o que pode afetar a precisão do classificador. Na abordagem por objeto, esse efeito não costuma ocorrer visto que para a mesma classe a variação espectral é menor.

Apesar de uma maior aderência, a abordagem por objeto tem as suas limitações e problemas. Os dois erros mais comuns na segmentação de imagens são o excesso ou a falta de segmentação (Liu e Xia, 2010). No caso do segundo erro, os objetos criados podem incluir mais do que uma classe, o que pode levar a erros na classificação. Em ambos os erros, os atributos dos objetos criados podem não ser representativos dos objetos reais, como a área e forma, o que pode diminuir a precisão do classificador.

Em estudos de comparação entre as duas abordagens para a deteção de alterações temporais em imagens de satélite, ambas as abordagens mostram serem eficazes na deteção de alterações. Na abordagem por píxel, a informação extraída é mais detalhada, mas a abordagem por objeto permite a utilização dos objetos nas operações posteriores, tornando todo o processo mais eficiente (Pereira-Pires, Aubard, Silva et al., 2020).

2.4 Segmentação

A segmentação de imagem consiste em agrupar um conjunto de píxeis de uma imagem para formar um objeto. O objetivo principal da segmentação é identificar objetos reais nas imagens, facilitando e simplificando a análise de imagens. Existem várias técnicas e formas de usar e aplicar a segmentação, dependendo da aplicação. Para imagens de deteção remota, é comum utilizar métodos derivados da segmentação baseada em regiões.

2.4.1 Segmentação Superpíxel

Esta segmentação consiste em agrupar píxeis semelhantes ao nível da cor, textura e da forma. Uma das principais vantagens de segmentação superpíxel é a diminuição da complexidade das próximas operações de processamento de imagens. Em vários estudos,

os algoritmos de segmentação superpíxel foram comparados com uma segmentação em excesso, sendo considerados quase equivalentes. No entanto, posteriormente, surgiu a principal diferença entre ambas, os algoritmos de segmentação superpíxel permitem um controle sobre os números de objetos criados, enquanto que no outro método isso não é possível.

Algumas propriedades desejadas para os superpíxeis são: os objetos criados devem ser disjuntos e devem caracterizar todos os píxeis, é esperado que os objetos criados preservem os contornos na imagem, devem ser objetos compactos, regulares, com contornos lisos, os objetos criados não podem prejudicar o desempenho das próximas operações e, idealmente, devem melhorar esse desempenho. Por fim, deve criar-se o menor número possível de objetos para cumprir os requisitos mencionados e obter uma segmentação de boa qualidade (Stutz et al., 2018).

Existem vários algoritmos de segmentação superpíxel que se baseiam em diferentes métodos de segmentação. Estes são alguns exemplos:

- **Segmentação baseada em bacias hidrográficas:** nesta abordagem, é usada a transformada de *watershed*, que consiste numa segmentação a partir do mapa topográfico da imagem, onde cada píxel corresponde a uma posição e os níveis de cinza representam as altitudes. Os objetos são gerados agrupando os píxeis mais próximos de um ponto mínimo de altitude, de modo que cada objeto tenha um ponto mínimo local (Kornilov e Safonov, 2018). Quanto aos algoritmos de segmentação superpíxel baseados nesta abordagem, existe o algoritmo tradicional, *Watershed*, o *Compact Watershed*, entre outros. Ambos permitem controlar o número de superpíxeis e, no último, é possível controlar a compactidade dos mesmos (Stutz et al., 2018).

(a) *Watershed*(b) *Compact Watershed*

Figura 2.8: Superpíxeis gerados por algoritmos baseados em bacias hidrografias, com 400 superpíxeis para a parte superior esquerda da imagem e 1200 superpíxeis para a parte inferior direita (Stutz et al., 2018).

- **Segmentação baseada em grafos:** esta abordagem é fundamentada na teoria dos grafos, sendo um ramo da matemática que estuda as relações entre os objetos num conjunto. Um grafo é representado por $G(V, E)$, onde V é o conjunto de vértices (ou nós) e E é o subconjunto de pares de vértices (ou arestas) que os ligam. Na segmentação de imagens com esta abordagem, os vértices, V , correspondem a cada píxel e as arestas, E , aos pares de píxeis vizinhos. As arestas têm pesos associados para medir as diferenças entre dois píxeis ligados por essa aresta (Felzenszwalb e Huttenlocher, 2004). Alguns exemplos de algoritmos de segmentação superpíxel baseados nesta abordagem são o *Entropy Rate Superpíxels (ERS)*, o *Normalised Cuts*, entre outros. Nestes, é possível controlar apenas o número de superpíxeis gerados (Stutz et al., 2018).

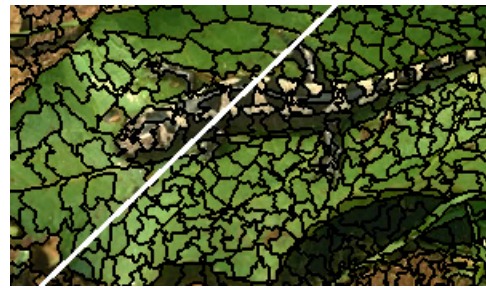
(a) *Normalised Cuts*(b) *Entropy Rate Superpíxels (ERS)*

Figura 2.9: Superpíxeis gerados por algoritmos baseados em grafos, com 400 superpíxeis para a parte superior esquerda da imagem e 1200 superpíxeis para a parte inferior direita (Stutz et al., 2018).

- **Segmentação baseada em *clustering*:** esta abordagem consiste na agregação de objetos ou píxeis com características semelhantes. Geralmente, é utilizada quando as classes já estão identificadas no conjunto de dados. Este método pode ser dividido em duas fases: na primeira fase, é necessário escolher o número de agrupamentos pretendidos e calcular os respetivos centros; na segunda fase, cada objeto ou píxel é movido para o agrupamento que tenha a menor distância relativamente ao centro (Juneja e Kashyap, 2016). Um dos algoritmos de segmentação superpíxel mais comum é o *Simple Linear Iterative Clustering (SLIC)*, embora existam outras variantes e melhorias deste algoritmo, como o *Preemptive SLIC (preSLIC)* (Stutz et al., 2018).



Figura 2.10: Superpíxeis gerados por algoritmos baseados em *clustering*, com 400 superpíxeis para a parte superior esquerda da imagem e 1200 superpíxeis para a parte inferior direita (Stutz et al., 2018).

2.4.1.1 Simple linear iterative clustering

O *Simple linear iterative clustering* (SLIC) é um dos algoritmos mais comuns de segmentação superpíxel. É considerado um algoritmo simples, contendo apenas um parâmetro, o número desejado de superpíxeis a serem gerados. O algoritmo também permite ajustar opcionalmente a compactidade (Achanta et al., 2012). O processo é realizado num espaço de 5 dimensões (5D), o espaço $[L, a, b, x, y]$, onde $[L, a, b]$ é o vetor da cor do píxel no espaço de cores CIELAB e $[x, y]$ representam as coordenadas do píxel. CIELAB, também conhecido como CIE $L^*a^*b^*$, é um espaço de cores definido pela Comissão Internacional em Iluminação (CIE) que descreve a cor com três componentes numéricas, onde L especifica a luminosidade da cor (varia de 0 — preto a 100 — branco) e a e b representam informações da cor, podendo variar entre valores positivos e negativos, representando vermelho/verde e azul/amarelo, respetivamente. Para agrupar os píxeis no espaço 5D, é utilizada uma função de distância (Equação 2.3) (Achanta et al., 2010). O algoritmo começa com a escolha do número de superpíxeis a serem criados (K) e a definição dos centros iniciais $C_k = [l_k, a_k, b_k, x_k, y_k]^T$, onde $k \in [1, K]$, amostrados numa grelha regular com um espaçamento de S píxeis, com $S = \sqrt{N/k}$ (N é o número de píxeis da imagem) para obter superpíxeis com tamanhos semelhantes. Em seguida, numa vizinhança de 3×3 do centro, o mesmo é movido para o píxel com o menor gradiente para evitar centrar os superpíxeis nas bordas (Achanta et al., 2012). Como é esperado que cada superpíxel tenha um tamanho de $S \times S$, a procura de píxeis semelhantes é realizada numa área de $2S \times 2S$ em torno do centro, tornando este algoritmo mais rápido em comparação com outros algoritmos de *clustering*, como o *k-means*, em que a comparação é feita com todos os centros. A distância (D_s) é obtida pela equação 2.3, em que i representa o píxel a comparar e m representa a compactidade desejada:

$$d_{lab} = \sqrt{(l_k - l_i)^2 + (a_k - a_i)^2 + (b_k - b_i)^2} \quad (2.1)$$

$$d_{xy} = \sqrt{(x_k - x_i)^2 + (y_k - y_i)^2} \quad (2.2)$$

$$D_s = d_{lab} + \frac{m}{S} d_{xy} \quad (2.3)$$

2.4.2 Segmentação Agrupamento Difuso

Como já foi referido em 2.4.1, o *clustering* consiste em formar grupos cujos elementos sejam o mais semelhantes possível entre si, utilizando medidas de distância, conectividade e/ou intensidade para apurar a similaridade entre os elementos. O agrupamento difuso (*fuzzy clustering*) é uma forma de agrupar elementos que podem pertencer a mais de um grupo. Por exemplo, neste método um pêssego pode conter as cores vermelha, laranja e amarelo, e nos outros métodos só poderia ser uma das cores.

O Agrupamento difuso C-means (FCM) é um dos algoritmos de agrupamento difuso mais conhecidos. Foi criado por J.C. Dunn em 1973 e mais tarde melhorado por J.C. Bezdek em 1981. Este algoritmo é semelhante ao algoritmo de agrupamento k-means, com a diferença de que no FCM um elemento pode pertencer a vários conjuntos e é atribuída uma probabilidade que representa o quão provável é esse elemento pertencer a um determinado conjunto. O objetivo principal do algoritmo FCM é minimizar a função objetivo, $J_{FCM}(U, V, X)$ (Jinlin et al., 2018):

$$J_{FCM}(U, V, X) = \sum_{i=1}^N \sum_{j=1}^c u_{ij}^m d^2(x_i, v_j) \quad (2.4)$$

Onde $m \geq 1$, o vetor $X = \{x_1, x_2, x_3, \dots, x_N\}$ representa os N píxeis da imagem, o vetor $V = \{v_1, v_2, \dots, v_c\}$ representa os centros dos *clusters*, $d^2(x_i, v_j)$ é o quadrado de uma expressão de similaridade entre o elemento x_i e o centro de um *cluster*, v_j , que pode ser a distância euclidiana ou outra, e u_{ij} representa o grau de pertença do elemento x_i ao *cluster* c_j . Para que o algoritmo execute corretamente, é necessário que sejam satisfeitas as seguintes condições:

$$0 \leq u_{ij} \leq 1 \quad (2.5)$$

$$\sum_{j=1}^c u_{ij} = 1 \quad (2.6)$$

$$\sum_{i=1}^N u_{ij} \leq N \quad (2.7)$$

No processo iterativo, os centros e os graus de similaridade são constantemente atualizados com o objetivo de minimizar o valor da função objetivo. As atualizações da matriz de similaridade, U e dos centros v_j são feitas por meio das seguintes fórmulas:

$$u_{ij} = \frac{1}{\sum_{k=1}^c \left(\frac{d_{ij}}{d_{ik}}\right)^{\frac{2}{m-1}}} \quad (2.8)$$

$$v_j = \frac{\sum_{i=1}^N u_{ij}^m x_i}{\sum_{i=1}^N u_{ij}^m} \quad (2.9)$$

2.4.3 Algoritmo de segmentação multiresolução

A segmentação multiresolução é um método comum de segmentação utilizado em processamento de imagens. Este algoritmo está integrado na aplicação da *eCognition*, que permite criar aplicações de análise de imagens baseadas em objetos (OBIA). O método é baseado na segmentação por união de regiões usando uma abordagem de *bottom-up*. Este processo é iniciado com objetos de um píxel e são iterativamente criados objetos maiores, derivados da junção de objetos adjacentes que atendem a critérios pré-estabelecidos, como um critério de heterogeneidade entre os objetos adjacentes (Baatz e Schape, 2000) e/ou outro critério, como o tamanho dos objetos (Benz et al., 2004). Esse processo iterativo é finalizado quando o valor da função de heterogeneidade entre objetos excede um limite pré-estabelecido.

A segmentação multiresolução da versão *eCognition* é mais completa, pois permite controlar a segmentação escolhendo certos atributos para os objetos, como o tamanho, a cor, a forma, a suavidade e a compactidade. Ao utilizar diferentes escalas de segmentação, é possível detetar objetos de diferentes tamanhos, permitindo utilizar este algoritmo para problemas de várias escalas, desde as menores até as maiores. É importante notar que a escolha de uma escala de segmentação que controle o tamanho médio dos objetos é crucial para obter uma boa segmentação, pois é importante que a escala escolhida seja semelhante ao tamanho real dos objetos a serem segmentados. Visto que a escolha de uma boa escala pode ser um processo exaustivo, algumas estratégias e métodos foram adaptados para automatizar a seleção da escala. Uma dessas ferramentas é a *Estimation of Scale Parameter* (ESP-2) (Drăguț et al., 2010).

2.5 Classificação

Após a recolha de dados, é necessário extrair informação dos mesmos. O passo seguinte é a utilização de classificadores para a extração de conhecimento. O processo de classificação consiste em estimar a classe a que algo pertence com base nos atributos extraídos dos dados recolhidos. Uma das formas mais conhecidas de classificação é a utilização de aprendizagem automática (*machine learning*). Existem dois tipos de aprendizagem, a supervisionada, onde o classificador utiliza dados exemplo previamente classificados, para calcular a estimativa da classe dos dados a classificar. Alguns exemplos de algoritmos de aprendizagem supervisionada são, por exemplo, as regressões lineares, florestas

aleatórias, redes neuronais, entre outras. A aprendizagem não supervisionada é utilizada para encontrar padrões no conjunto de dados, sem o conhecimento prévio das classes dos exemplos reais.

Além da utilização de aprendizagem automática, existem outras técnicas de classificação, como a utilização de testes estatísticos. O teste de hipóteses é uma técnica que coloca uma hipótese acerca dos dados em estudo e aceita ou rejeita a hipótese com base num resultado estatístico. Estes permitem a realização de testes de comparação entre dois ou mais grupos. Para a classificação por hipóteses, é necessário ter em conta cinco passos importantes:

- **1.º Passo:** Definir as Hipótese Nula (H_0) e Hipótese Alternativa (H_1). Por exemplo, na comparação de um certo valor entre duas amostras, querendo saber se as amostras são parecidas ou iguais, em que $H_0 : \mu_1 = \mu_2$, representa que a média dos dois grupos é igual, enquanto H_1 representa a negação de H_0 , sendo que $H_1 : \mu_1 \neq \mu_2$;
- **2.º Passo:** Escolha de um valor de significância/confiança, representado pela letra grega α . Este valor é um limite que determina se os resultados obtidos são estatisticamente significativos, sendo geralmente menor ou igual a 5%. Esta percentagem representa a probabilidade de rejeitar a Hipótese Nula (H_0) quando ela é verdadeira;
- **3.º Passo:** Obtenção e preparação dos dados consoante o teste estatístico apropriado para o estudo que se pretende realizar, respeitando as suas condições e regras;
- **4.º Passo:** Execução do teste estatístico, obtendo o valor- p . Este valor é uma probabilidade que permite verificar a significância da Hipótese Nula (H_0), permitindo aceitar ou rejeitar com base no valor de significância (α). Se o valor- p for igual ou menor a α , H_0 é rejeitado;
- **5.º Passo:** Realização da classificação consoante a aceitação ou rejeição das hipóteses.

Em estatística, para ser possível retirar conclusões válidas sobre uma população, é necessário ter em consideração alguns pressupostos por detrás dos dados. Esses pressupostos devem ser cuidadosamente escolhidos, pois podem levar a conclusões imprecisas. Alguns dos pressupostos mais comuns em estatística são:

- Normalidade: assume-se que os conjuntos de dados utilizados na análise têm uma distribuição normal;
- Linearidade: refere-se à relação entre as variáveis independentes e dependentes, utilizada em regressões lineares;
- Homogeneidade de variâncias: refere-se a uma variância semelhante ao longo dos vários grupos de dados.

Consoante a análise que se pretende realizar, existem testes apropriados para cada caso. Alguns dos mais comuns são:

- Regressão: é um teste que permite estimar a causalidade/relação entre variáveis dependentes e independentes. O teste estatístico mais comum para este caso é a Regressão Linear;
- Testes comparativos: estes testes verificam as diferenças de médias entre grupos. Os mais conhecidos são o *t-Test* e o ANOVA;
- Correlação: este teste procura associações entre variáveis e quantifica a relação entre elas. Alguns testes são: o coeficiente de correlação de Pearson e o teste de *Chi-Square* (χ^2).

2.5.1 Algoritmos de Aprendizagem Automática

Nesta secção serão apresentados e descritos alguns dos algoritmos mais comuns de aprendizagem automática que foram e são utilizados na deteção remota (Pereira-Pires, Aubard, Ribeiro et al., 2020).

2.5.1.1 Floresta Aleatória

O classificador Floresta Aleatória (*Random Forest*) introduzido por Breiman em 2001, é um algoritmo comum de aprendizagem de máquina que tem sido aplicado com sucesso em problemas de classificação, regressão e seleção de atributos. A Floresta Aleatória é um método de classificação por *ensemble* que consiste em utilizar múltiplas árvores de decisão para obter uma melhor eficácia de predição do que usar as árvores individualmente. Para construir uma Floresta Aleatória, são geradas várias árvores de decisão com dados retirados aleatoriamente de um conjunto de dados. Normalmente, dois terços dos dados de treino são usados para criar as árvores de decisão, enquanto um terço dos dados é usado para realizar uma validação cruzada e calcular o poder de predição das árvores de decisão geradas (Tonbul, 2018). Na classificação da Floresta Aleatória, é necessário especificar o número de árvores e o número de atributos a serem usados. Quanto maior o número de árvores, melhor a eficácia de predição, mas isso também aumenta o tempo de treino e o tempo de previsão (Breiman, 2001). O número de atributos é importante porque ajuda a reduzir a correlação entre as árvores e a aumentar a diversidade na Floresta Aleatória.

2.5.1.2 Redes Neurais Artificiais

As Redes Neurais Artificiais (RNA) são algoritmos de aprendizagem automática inspirados na estrutura do cérebro humano. A estrutura de uma RNA é composta por neurónios artificiais, representados por u_j na figura 2.11, onde cada neurónio tem valores de entrada e produz um resultado singular que por sua vez pode ser utilizado como valor de entrada nos neurónios das camadas seguintes. Esses neurónios podem ser organizados

em diferentes camadas, sendo que duas delas são para a entrada de valores do exterior, a camada de entrada, e a outra para a produção do resultado final, a camada de saída, representadas respetivamente por x_i e z_k na figura 2.11. Entre essas duas camadas podem existir zero ou mais camadas escondidas. Entre estas camadas podem existir vários padrões de conexões, onde cada neurónio de uma camada está conectado a todos os neurónios da camada seguinte, como está representado por w_{ij} e w'_{jk} da figura 2.11. Em certas arquiteturas, é possível ter um neurónio ligado não só à camada seguinte, mas também às posteriores. As RNA conseguem realizar aprendizagem automática e reconhecimento de padrões, tornando-as amplamente utilizadas em várias aplicações de deteção remota, como classificação de imagens, deteção de alterações e análise de séries temporais.

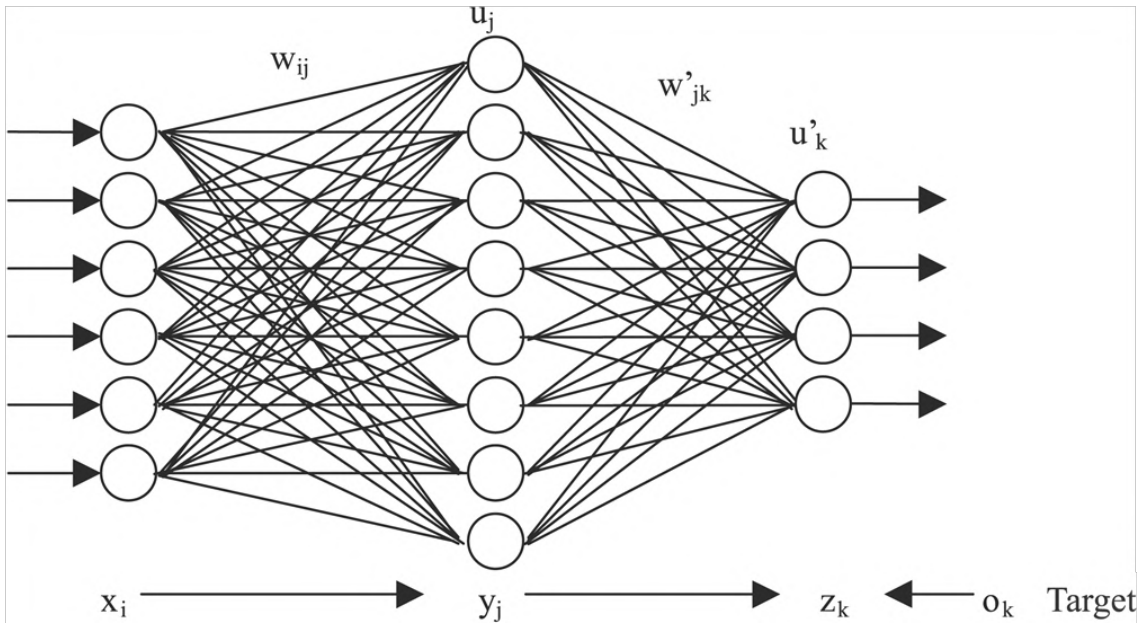


Figura 2.11: Representação de uma Rede Neuronal Artificial com uma camada oculta. Retirado de (“Neural Network”, 2015).

2.5.2 Algoritmos de Testes Estatísticos

Para o tema abordado nesta dissertação, é relevante apresentar algumas técnicas estatísticas frequentemente usadas na deteção de alterações em imagens de satélite, principalmente para abordagens por píxel.

2.5.2.1 Welch's *t*-Test

Este teste estatístico é uma adaptação ou melhoria do *Student's t-Test*, sendo que o primeiro consegue realizar um teste *t* para conjuntos de dados com variâncias e tamanhos diferentes, enquanto o segundo é indicado para conjuntos de dados com variância e tamanhos iguais. Comparando ambas as técnicas, conclui-se que o *Welch's t-Test* é mais

robusto e recomendável do que o *Student's t-Test*, pois no caso de utilização de conjuntos de dados com variância e tamanhos iguais o *Welch's t-Test* apresenta um resultado muito próximo do *Student's t-Test* e o contrário já não se verifica.

Este teste é geralmente utilizado para realizar um teste de hipótese que verifique a igualdade das médias de duas populações ou grupos de dados. Antes de realizar o teste, é necessário garantir que os conjuntos de dados possuem uma normalidade, o mesmo é necessário garantir para o *Student's t-Test*. Em seguida, são realizados os seguintes cálculos:

$$t = \frac{\bar{X}_1 - \bar{X}_2}{\sqrt{\frac{s_1^2}{N_1} + \frac{s_2^2}{N_2}}} \quad (2.10)$$

Onde $\bar{X}_i$ é a média, s_i representa o desvio padrão e o N_i o tamanho do conjunto de dados i . Os graus de liberdade, GL , são calculados utilizando a aproximação de *Satterthwaite*:

$$GL \approx \frac{\left(\frac{s_1^2}{N_1} + \frac{s_2^2}{N_2} \right)^2}{\frac{(s_1^2/N_1)^2}{N_1-1} + \frac{(s_2^2/N_2)^2}{N_2-1}} \quad (2.11)$$

Após obter o valor- t e os graus de liberdade (GL), é possível obter o valor- p através das tabelas de distribuição- t e tirar conclusões acerca das hipóteses. No caso de um teste bicaudal, testa-se se as médias de ambos os grupos são iguais, e no caso de um teste unicaudal, testa-se se a média de um grupo é maior (ou menor) ou igual a outro grupo.

TRABALHO DESENVOLVIDO

Neste capítulo é apresentado o trabalho desenvolvido para a detecção de tratamentos nas FGC com uma abordagem por píxel. Para tal, são propostas duas metodologias: a primeira, usando uma aprendizagem supervisionada, e a segunda, usando o teste estatístico *Welch's t-test*. Em seguida, é feita uma breve descrição da implementação das duas metodologias em *Python*. Para o segundo método, foram realizadas e testadas modificações do algoritmo com o intuito de melhorar a sua eficiência. Por fim, foi implementado um método de pós-processamento para reduzir o ruído na classificação final.

3.1 Áreas de Estudo

Para esta dissertação, foram definidos dois conjuntos de áreas de estudo. O primeiro conjunto foi utilizado para realizar uma primeira validação das metodologias descritas abaixo e é composto por uma área próxima de Marisol, no concelho de Almada, uma zona constituída principalmente por eucaliptos e que não faz parte da RPFGC. No entanto, existem intervenções para reduzir a biomassa por baixo de uma linha de alta tensão, criando uma faixa com cerca de 40 metros de largura. A localização desta zona pode ser visualizada na figura 3.1. Este primeiro conjunto é utilizado para uma primeira validação dos resultados obtidos das metodologias descritas abaixo, por ser um exemplo de pequena dimensão que não exige grande esforço computacional.

O segundo conjunto de áreas de estudo é composto por quatro zonas diferentes pertencentes à RPFGC, situadas nos seguintes locais: Fundão, Seia, Serra dos Candeeiros (situada nos concelhos do Rio Maior, Alcobaça e Porto de Mós) e Sertã (Figura 3.2). Essas zonas são utilizadas principalmente para a validação e análise de resultados, pois existe informação disponibilizada pelo ICNF quanto às intervenções ocorridas nas faixas para o ano de 2017 e 2018.

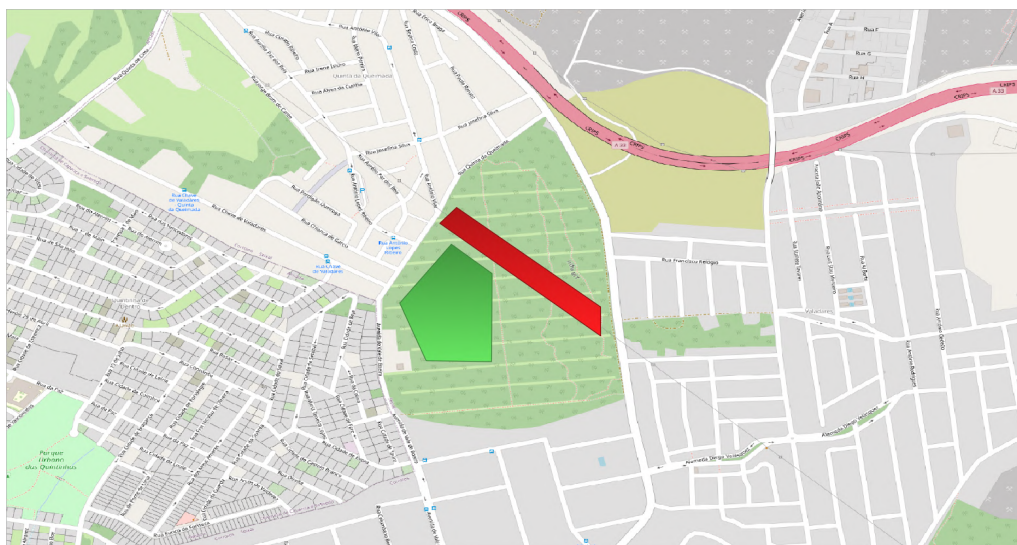


Figura 3.1: Área de Estudo: Marisol, concelho Almada. A área a vermelho representa a faixa de intervenção e a área a verde representa uma zona de vegetação.

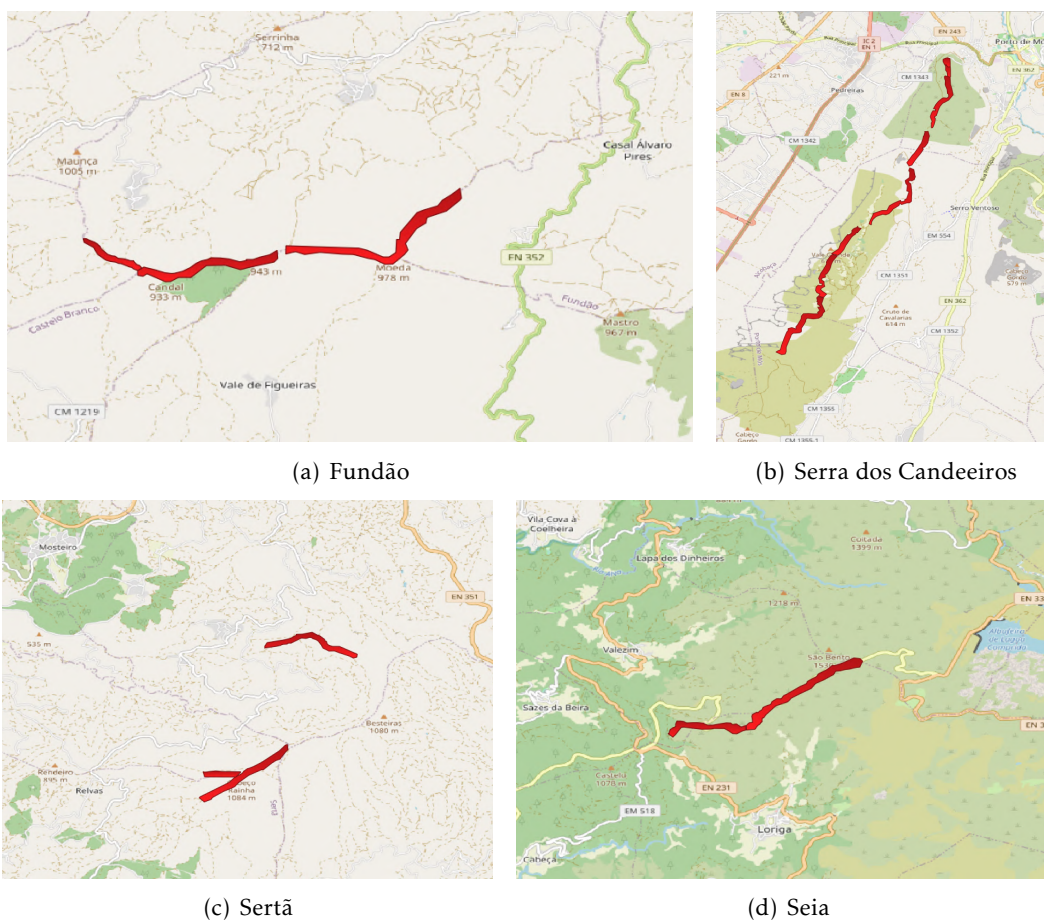


Figura 3.2: Área de estudo de FGC na RPFGC.

3.2 Metodologias para a detecção por píxel de tratamentos nas FGC

Foram implementadas e testadas duas metodologias para a detecção de tratamento das FGC por píxel. As metodologias em questão foram apresentadas em (Aubard et al., 2020) e implementadas na plataforma baseada em nuvem *Google Earth Engine (GEE)*. Nesta dissertação, foram implementadas em *Python*, não estando sujeitas aos termos de utilização da plataforma GEE.

O primeiro método utiliza um compósito de médias mensais dos atributos necessários para o classificador, sendo principalmente índices de vegetação, como o *Normalized Difference Vegetation Index (NDVI)*, o *Excess of Green Index (ExG)*, o *Excess of Red Index (ExR)*. O segundo método consiste na detecção de diferenças estatisticamente significativas em cada píxel das FGC de cada observação de satélite numa janela temporal. Como o segundo método é baseado em testes estatísticos, uma das suas vantagens é que não requer um conjunto de dados de treino para criar o classificador, mas necessita de um número elevado de observações para uma boa classificação. No entanto, como o teste estatístico é aplicado a cada píxel, é mais exigente computacionalmente do que o primeiro método.

Ambas as metodologias têm como principal limitação a existência de efeitos atmosféricos, como nuvens, sombras e outros, nas observações de satélite que impossibilitam a leitura do valor real dos píxeis das FGC, diminuindo o número de dados disponíveis para a classificação final. Para ambos os métodos é necessária a realização de um pré-processamento das observações de satélite antes de realizar a classificação.

3.2.1 Pré-processamento das observações de satélite

A primeira fase consiste na aquisição, armazenamento e preparação dos dados provenientes do satélite *Sentinel-2* para a fase seguinte, conforme ilustrado no esquema de execução apresentado na figura 3.3. Como as observações disponíveis das áreas de estudo selecionadas são a partir de abril de 2017, o estudo é iniciado a partir deste mês. Para a classificação das imagens, são utilizadas imagens do *Sentinel-2* do nível 2A em vez de 1C, pois as primeiras já possuem correções atmosféricas realizadas pelo algoritmo *Sen2Cor* da ESA e uma imagem de *Scene Classification (SCL)*, que fornece uma classificação de 0 a 11 para resoluções de 60 e 20 metros, indicando quais píxeis são nuvens, sombras, nevoeiro e outros, possibilitando a filtragem de píxeis (Figura 3.6). As imagens são armazenadas localmente, organizadas em pastas pelo *tile* que representam e, em seguida, em cada uma dessas pastas, são organizadas por ano de captura das imagens, o que simplifica e facilita a análise e/ou tratamento para anos e *tiles* específicos.

Um dos principais aspetos a ser tratado nesta fase é a correção de pequenos desvios dos píxeis, até 1,5 píxeis (15 metros), verificados entre observações captadas em datas diferentes por satélites diferentes. Esses desvios devem ser corrigidos para garantir que, ao longo de um ano de observações, um determinado píxel represente sempre o mesmo

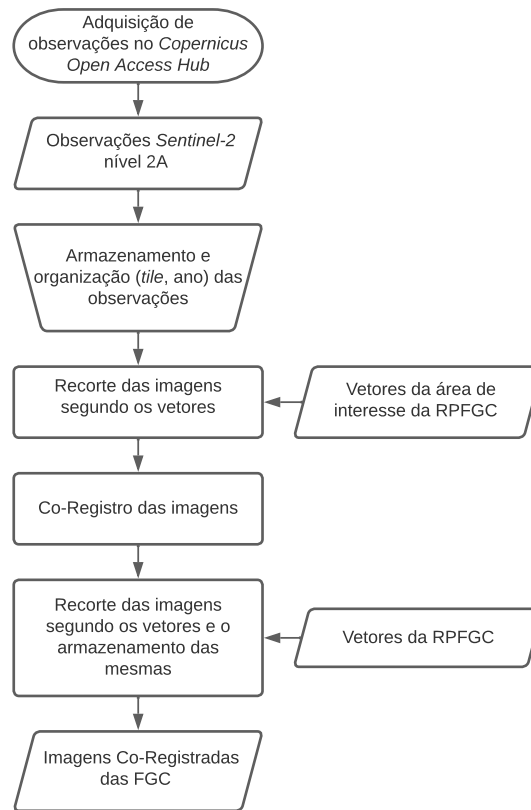


Figura 3.3: Fluxograma do pré-processamento.

espaço físico. Normalmente, esses desvios são mais evidentes e prejudiciais nos píxeis de fronteiras. Por exemplo, numa estrada próxima a uma área de vegetação, se o desvio não for corrigido, um píxel da vegetação dessa área pode alternar entre vegetação e estrada ao longo das observações, aumentando o número de falsos positivos. Para realizar a correção dos desvios, é necessário definir uma imagem de referência, escolhendo de preferência, uma imagem sem nuvens ou com poucas nuvens que preencha todo o *tile*, e utilizada a banda espectral 4 (vermelho) para o cálculo dos desvios. Para a implementação, é utilizado um pacote de *Python* que executa o co-registro de imagens de satélite, o AROSICS (Scheffler, 2017).

Na implementação em *Python*, antes de realizar o co-registro das imagens, são realizados cortes para as áreas de interesse que cobrem grandes partes da RPFGC, sendo realizado o co-registro posteriormente nas imagens cortadas. As imagens geradas são armazenadas localmente numa pasta diferente das originais e organizadas da mesma forma, adicionando-se mais um nível para organizar por áreas de interesse. Depois do co-registro, são geradas imagens que representam as FGC através de cortes nas imagens co-registradas. Dependendo da metodologia, são preparadas as imagens das bandas espectrais necessárias e, no caso da metodologia 2, é feita uma preparação da COS, que é necessária para a classificação.

3.2.2 Metodologia 1: Aprendizagem Supervisionada

Neste caso, a implementação em *Python* não foi totalmente idêntica ao algoritmo proposto em (Aubard et al., 2020). Com esta implementação, pretendeu-se testar uma RNA treinada com uma abordagem orientada a objetos para detetar tratamentos de FGC e perceber o seu funcionamento para uma abordagem a píxel, usando compósitos médios mensais. Na figura 3.4, é apresentado o esquema de execução desta metodologia.

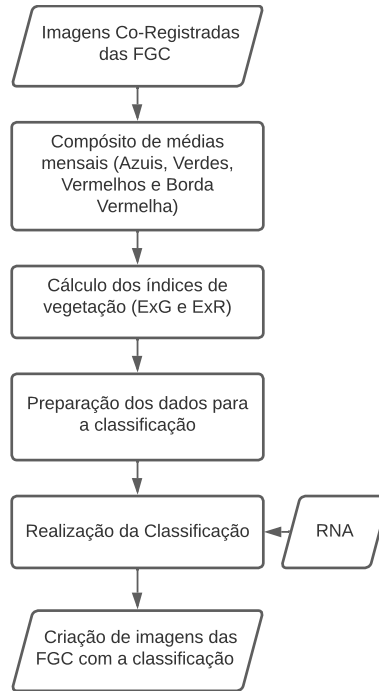


Figura 3.4: Fluxograma da metodologia 1.

Após o pré-processamento, é realizado um compósito de médias mensais para as bandas Azul (banda 2), Verde (banda 3), Vermelha (banda 4) e Vermelho Limite (banda 5) para os píxeis das FGC. Esses valores médios são posteriormente usados para obter os índices de vegetação, mais precisamente o ExG e ExR, que são calculados respetivamente com as fórmulas 3.1 e 3.2, onde G , R e B representam o verde, vermelho e o azul.

$$ExG = 2 \times G - R - B \quad (3.1)$$

$$ExR = 1.3 \times R - G \quad (3.2)$$

Após obter os índices de vegetação, os dados são organizados de forma a serem usados como entradas no classificador RNA. Para cada píxel de cada mês, os dados são organizados da seguinte forma num *array*, $[b5, ExG, ExR, b5', ExG', ExR']$, em que cada campo representa:

- **b5**: valor médio mensal do Vermelho Limite (banda 5);
- **b5'**: valor médio mensal do Vermelho Limite (banda 5) do mês anterior;
- **ExG**: ExG do mês atual;
- **ExG'**: ExG do mês anterior;
- **ExR**: ExR do mês atual;
- **ExR'**: ExR do mês anterior.

Isso resulta numa classificação binária, em que 1 indica que existiu tratamento naquela zona do píxel e 0 representa o oposto, a ausência de tratamento. Após se obter a classificação para cada píxel, são criadas imagens binárias mensais e anuais, onde é possível visualizar e analisar os resultados obtidos.

3.2.3 Metodologia 2: Classificação com o teste estatístico *Welch's t-Test*

Esta metodologia baseia-se na utilização do teste estatístico, *Welch's t-Test*, para a classificação de cada píxel quanto ao tratamento ou não dessa zona. Na figura 3.11, está representado o esquema de execução desta metodologia.

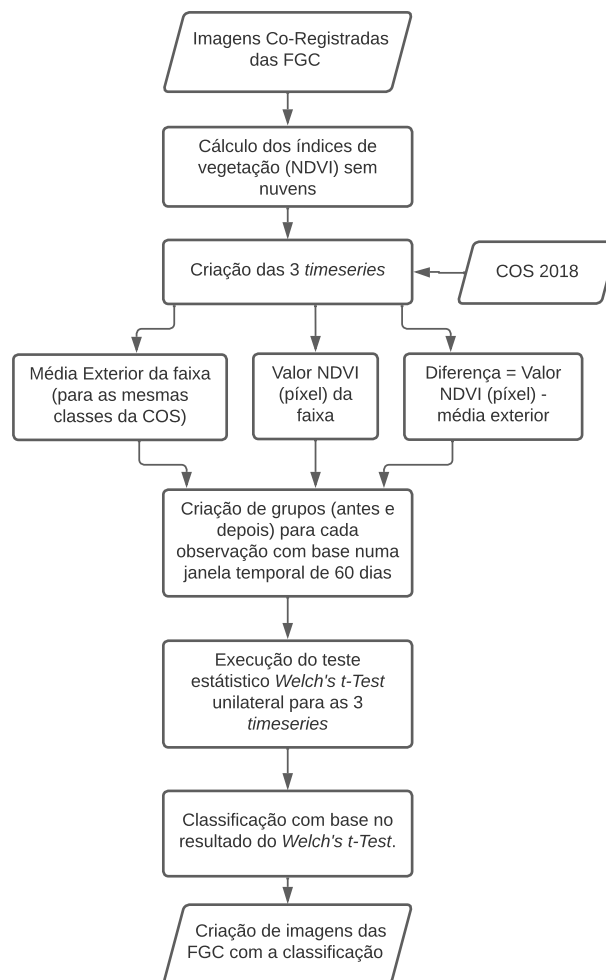


Figura 3.5: Fluxograma da metodologia 2.

$$NDVI = \frac{NIR - R}{NIR + R} \quad (3.3)$$

Após o pré-processamento, é calculado o índice de vegetação NDVI para cada píxel através da equação 3.3, onde *NIR* e *R* representam os infravermelhos próximos e o vermelho. Foi utilizada a imagem de *Scene Classification* (SCL) para realizar uma filtragem dos píxeis a utilizar nos processos seguintes. Apenas foram utilizados píxeis com uma classificação superior a 3 e inferior a 8, sendo os restantes excluídos. Na figura 3.6, está representada a interpretação dos valores da imagem de *Scene Classification* (SCL).

Label	Classification
0	NO_DATA
1	SATURATED_OR_DEFECTIVE
2	CAST_SHADOWS
3	CLOUD_SHADOWS
4	VEGETATION
5	NOT_VEGETATED
6	WATER
7	UNCLASSIFIED
8	CLOUD_MEDIUM_PROBABILITY
9	CLOUD_HIGH_PROBABILITY
10	THIN_CIRRUS
11	SNOW or ICE

Figura 3.6: Representação dos valores da imagem de *Scene Classification* (SCL). Retirado de (ESA, s.d.-a).

Em seguida, foram criadas três séries temporais para cada píxel, de cada observação:

- **Píxeis interiores:** é utilizado o valor do índice de vegetação, NDVI, de cada píxel do interior da FGC;
- **Média exterior:** é calculada uma média para o índice de vegetação, NDVI, dos píxeis do exterior da FGC, até uma distância de 500 metros para a mesma classe da COS de 2018, guardando esse valor nos píxeis interiores pertencentes à mesma classe da COS para facilitar processos futuros;
- **Diferença:** é calculado fazendo a diferença entre os **píxeis do interior** (NDVI) com a **média exterior** para cada píxel.

Ao longo do ano, ocorrem mudanças de estações e com isso verificam-se alterações fenológicas em certos meses. Para que estas alterações não influenciem a classificação, são utilizadas três séries temporais. A primeira serve para a detecção de uma variação da vegetação (possível intervenção ou variação fenológica), a segunda para perceber se existe variação fenológica e a última para revalidar/reforçar que existe uma variação da vegetação tendo em conta a variação fenológica. Explicado qual o objetivo de cada série temporal, a classificação final para a detecção de intervenções é feita da seguinte forma:

- Considera-se que um píxel sofreu intervenções se verificar uma variação significativa dos **píxeis do interior** e dos píxeis da **diferença** e não se verificar uma variação significativa da **média exterior** para a mesma data de observação.

Antes de executar o teste estatístico, são criados dois conjuntos de dados para cada data observada, um para agrupar os valores anteriores a essa data e outro para agrupar os valores posteriores à mesma data, incluindo o valor da data a observar no primeiro conjunto. Nesta metodologia, é sugerido utilizar uma janela temporal de 60 dias para criar os conjuntos e que cada conjunto deve ter no mínimo dois valores e um máximo de oito. Esta organização dos dados é feita para as três séries temporais.

Após a realização do teste *Welch's t-Test* para cada data observada, foram executadas as funções estatísticas do pacote *SciPy*, mais precisamente a função `scipy.stats.ttest_ind_from_stats()`. Foi realizado um teste unicaudal para cada série temporal, com o objetivo de testar se a média do primeiro grupo é superior ou igual à do segundo grupo. Esta função retorna dois valores, o valor-*t* e o valor-*p*, sendo que o valor-*p* é o mais relevante e é utilizado posteriormente na classificação. Foram testados três níveis de significância (α): igual a 0,05, 0,005 e 0,0005, tendo-se verificado que, de modo geral, os melhores resultados foram obtidos com o último valor.

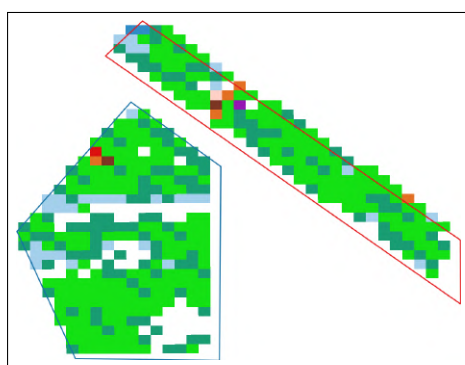
Por fim, a classificação foi realizada utilizando os critérios mencionados mais acima, produzindo imagens mensais e anuais que representam os píxeis que sofreram intervenções. A imagem anual é criada juntando todas as imagens mensais de um ano, e no caso de píxeis com múltiplas classificações mensais, é dada prioridade ao mês em que a intervenção foi detetada primeiro.

3.2.4 Primeira análise das metodologias

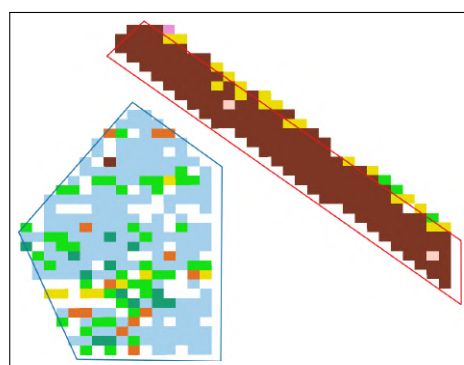
Como já foi referido, utilizou-se a área de estudo de Marisol para realizar uma análise inicial dos resultados obtidos. Sabe-se que na faixa desta zona foram realizados cortes em março de 2017 e em fevereiro de 2018. O primeiro tratamento não pode ser detetado porque não existem imagens do *Sentinel-2* do nível 2A para essas datas que cubram essa zona, no entanto, o segundo tratamento pode ser detetado. A figura 3.8 apresenta os resultados obtidos para cada metodologia nos anos de 2017 e 2018.

Value	Color	Label
1		1 - Janeiro
2		2 - Fevereiro
3		3 - Março
4		4 - Abril
5		5 - Maio
6		6 - Junho
7		7 - Julho
8		8 - Agosto
9		9- Setembro
10		10 - Outubro
11		11 - Novembro
12		12 - Dezembro

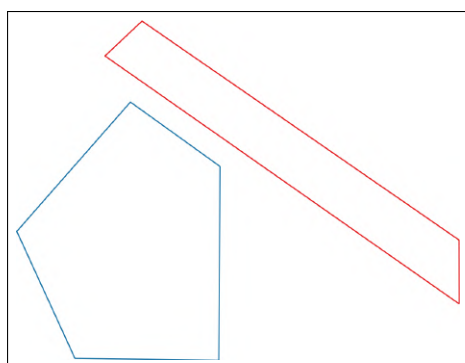
Figura 3.7: Legenda utilizada para a representação dos resultados obtidos.



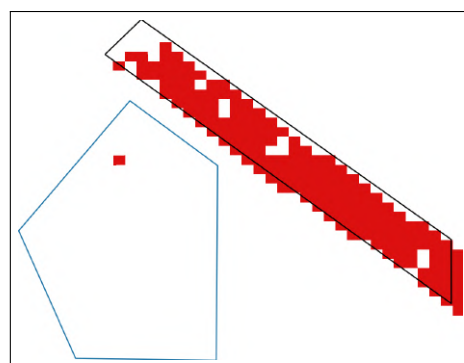
(a) **Metodologia 1** para o ano 2017.



(b) **Metodologia 1** para o ano 2018.



(c) **Metodologia 2** para o ano 2017.



(d) **Metodologia 2** para o ano 2018.

Figura 3.8: Resultados obtidos por ambas as metodologias para a área do Marisol.

Algumas das conclusões possíveis de retirar dos resultados obtidos são que a primeira metodologia não é adequada para este estudo e a segunda apresentou bons resultados, demonstrando ser uma boa metodologia para este tipo de análises. Na área de estudo de Marisol, foram delimitadas duas zonas para o estudo, uma que é a faixa (polígono retangular) e uma zona de vegetação (polígono pentagonal) onde nunca ocorreram tratamentos, logo as metodologias não deveriam detetar alterações nesta zona. No entanto, verificou-se que na primeira metodologia quase toda a zona de vegetação foi detetada como tendo sofrido alterações, enquanto na segunda metodologia, não se registaram alterações. Supõe-se que a primeira metodologia deteta alterações devido à fenologia, ou seja, esta metodologia não foi treinada para lidar com esses casos. Outro detalhe observado é que, na primeira metodologia, a deteção é feita com um mês de atraso, ou seja, um tratamento realizado em fevereiro é detetado em março, devido à utilização de compósitos de médias mensais. No caso da utilização de séries temporais (metodologia 2), a deteção coincide com o mês da intervenção.

3.3 Testes à Metodologia 2: Classificação com o teste estatístico Welch's t-Test

Com base nos resultados apresentados anteriormente, decidiu-se optar pela segunda metodologia para dar continuidade à dissertação. Foram realizados vários testes para aprimorar a metodologia e desenvolver um método de pós-processamento para reduzir o número de falsos positivos, especialmente os píxeis isolados.

3.3.1 Variação do nível de significância (α)

Foi testada a utilização de diferentes níveis de significância (α) para o teste estatístico realizado. Foram testados os valores de 0,05, 0,005 e 0,0005, e na figura 3.9 estão representados os resultados obtidos para vários valores de α na área de estudo de Marisol para o ano 2018. Observou-se que, de uma forma geral, o melhor nível de significância (α) é o 0,0005, pois além de identificar quase na totalidade as áreas que sofreram intervenção, também apresenta menos falsos positivos devido a variações fenológicas.

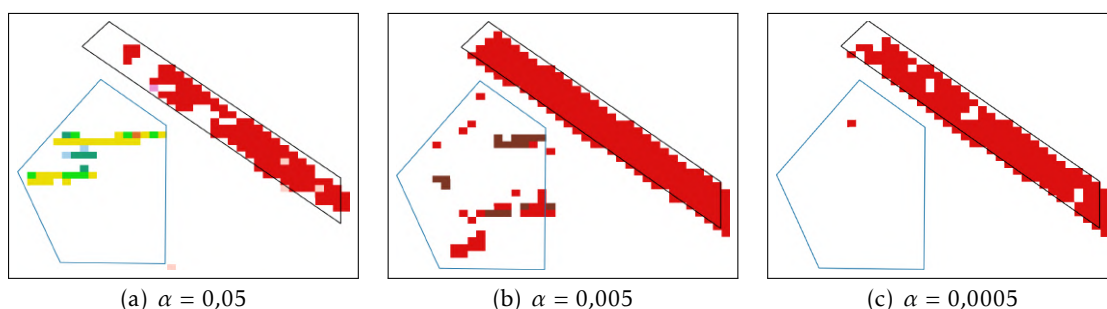


Figura 3.9: Resultados obtidos para Marisol em 2018 para níveis de significância (α).

Em vez de se testar com um único nível de significância para toda a faixa, realizou-se um pequeno estudo da variação do nível de significância para classes específicas da COS, com o objetivo de verificar se para algumas classes poderiam existir melhores detecções para valores diferentes de α . Os resultados desta análise serão apresentados no próximo capítulo.

3.3.2 Teste de normalidade

O teste estatístico *Welch's t-Test* pressupõe que os dados estejam distribuídos normalmente. Nesta metodologia, não se realiza qualquer teste aos dados quanto à normalidade, por isso, realizou-se o teste de *Shapiro-Wilk* para apurar se esse pressuposto é cumprido. Foram testados três níveis de significância (α), o 0,05, 0,005 e 0,0005, os mesmos utilizados no teste anterior. No caso de estudo de Marisol, verificou-se que para $\alpha = 0,0005$ e $\alpha = 0,005$ foram registradas menos ocorrências onde a normalidade não se verifica, podendo inferir-se que existe uma correlação entre uma boa classificação e um menor número de casos de falha de normalidade. No entanto, verificou-se que, mesmo para esses valores de α , o número de casos é razoavelmente elevado em alguns meses.

No que diz respeito a falhas de normalidade, existem três operações estatísticas para lidar com esses casos:

- É possível realizar o teste estatístico mesmo quando a amostra de dados não é normal, desde que seja adequada para o teste em questão, geralmente é necessária uma amostra de pelo menos 20 elementos;
- É possível executar alguma transformação ou operação nos dados de forma a ajustá-los a uma distribuição normal;
- É possível realizar um teste estatístico não paramétrico para a distribuição que os dados possam ter, como uma distribuição de *Poisson*, por exemplo.

Neste contexto, pensa-se que as primeiras duas opções são as mais adequadas. Na primeira opção, caso a classificação seja considerada boa, mesmo que os dados não cumpram a normalidade, aceitam-se os resultados. Na segunda opção, foi realizado um teste para verificar se é possível reduzir as ocorrências de falha da normalidade e, consequentemente, melhorar a qualidade da classificação. Neste caso, foram testados diferentes intervalos de tempo (usados na criação dos conjuntos de dados), nomeadamente 60, 50 e 40 dias. Observou-se, numa análise preliminar, que para um intervalo temporal de 30 dias, o conjunto de dados era demasiado pequeno. Por isso, a metodologia não foi testada para intervalos inferiores a 40 dias, pois a diminuição dos conjuntos de dados é desfavorável para a metodologia. A terceira opção não foi testada porque não existe ou não foi encontrado um teste estatístico não paramétrico semelhante ao *Welch's t-Test*. O teste U

de *Mann-Whitney* seria o mais próximo, mas é considerado o teste estatístico não paramétrico para o *Student's t-Test*. A análise dos diferentes intervalos de tempo é apresentada no capítulo 4.

3.3.3 Pós-processamento

Nesta metodologia, foi observado o aparecimento de píxeis isolados. Para resolver este problema optou-se por criar um pós-processamento que reduzisse a ocorrência desses casos. Para tal, utilizaram-se filtros de mediana. Foi testada a aplicação de um filtro de mediana (3x3) nas imagens com os resultados do teste *Welch's t-Test* para a série temporal dos píxeis internos, seguido da classificação. Na figura 3.10 é apresentado um exemplo de aplicação desse filtro, no qual se pode observar que o resultado obtido apresenta formas melhor delimitadas e mais homogêneas.

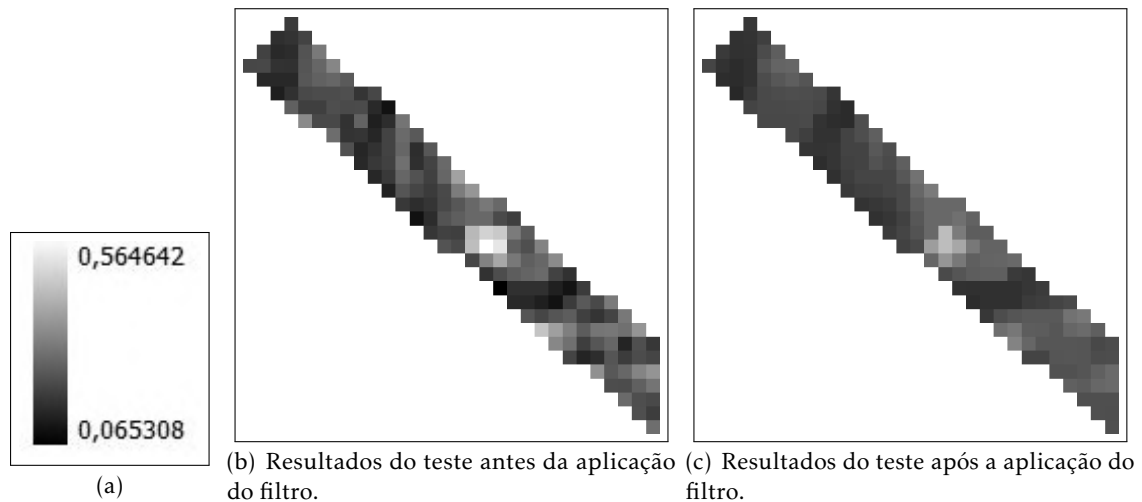


Figura 3.10: Exemplo de uma aplicação de um filtro de mediana (3x3) numa imagem com os valores do *Welch's t-Test* para a série temporal dos píxeis interiores.

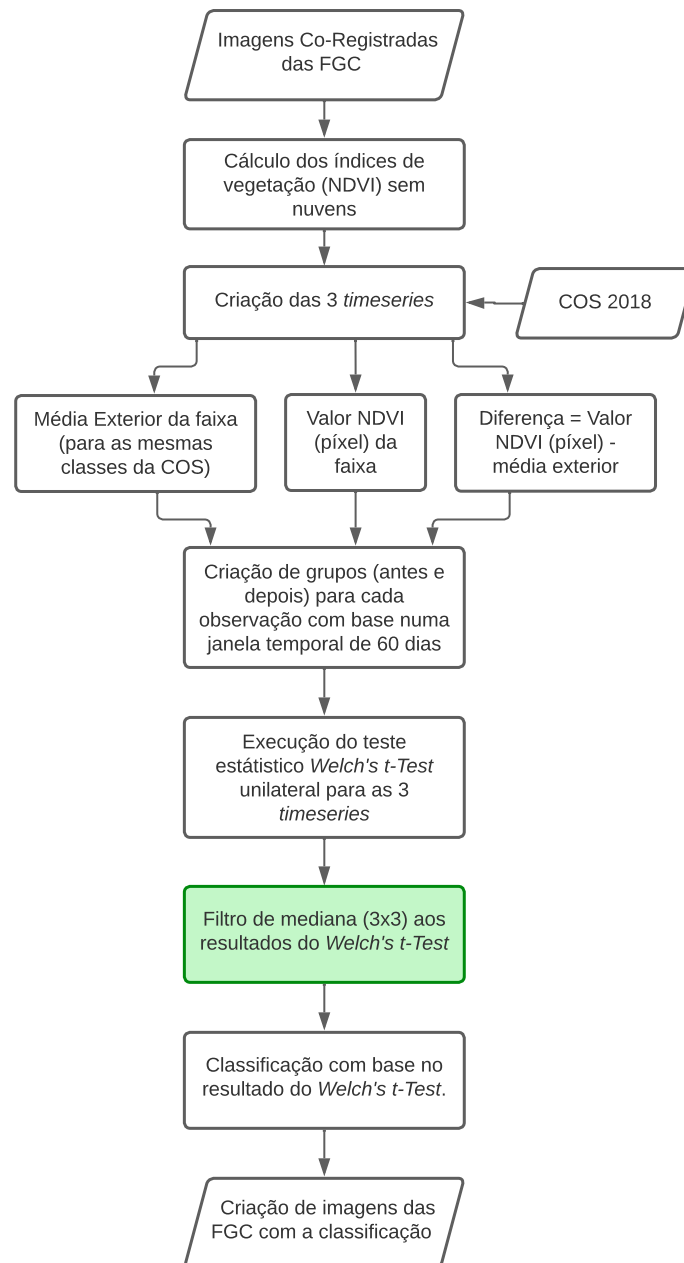


Figura 3.11: Fluxograma da metodologia 2 com pós-processamento.

APRESENTAÇÃO E DISCUSSÃO DOS RESULTADOS

Neste capítulo, apresentam-se os resultados do trabalho desenvolvido e realiza-se uma análise crítica dos mesmos. Para isso, são inicialmente apresentados os casos de estudo, seguidos dos resultados dos vários testes realizados com a metodologia de classificação com o teste estatístico *Welch's t-Test*, a variação do nível de significância (α), o teste de normalidade e os resultados do pós-processamento.

4.1 Áreas de Estudo

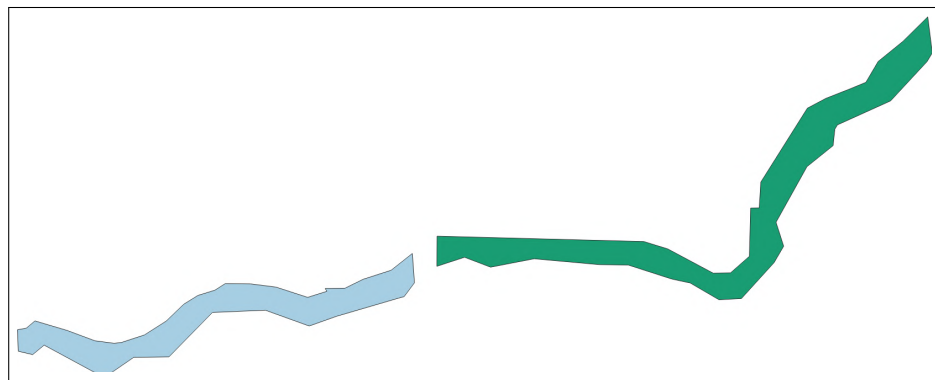
Serão apresentados os resultados das classificações realizadas pelo ICNF com uma abordagem por objeto, bem como as imagens *raster* correspondentes. Como os dados para a validação estão em formato vetorial e os resultados obtidos nesta dissertação estão em formato *raster*, não seria adequado nem possível avaliar os resultados obtidos. Assim, foram criadas imagens *raster* com o auxílio do *software* Quantum Geographic Information System (QGIS) para avaliar os resultados obtidos. Essas imagens foram criadas através de um processo manual, com base nos resultados em formato vetorial e na inspeção visual das observações de satélite, para apurar com melhor detalhe (a nível do píxel, com uma resolução de 10 metros) as áreas que sofreram tratamentos e o mês em que ocorreu. A legenda utilizada para identificação dos meses é a mesma que foi apresentada no capítulo anterior, e também está na figura 4.1.

Value	Color	Label
1		1 - Janeiro
2		2 - Fevereiro
3		3 - Março
4		4 - Abril
5		5 - Maio
6		6 - Junho
7		7 - Julho
8		8 - Agosto
9		9- Setembro
10		10 - Outubro
11		11 - Novembro
12		12 - Dezembro

Figura 4.1: Legenda utilizada para a representação dos meses classificados.

4.1.1 Fundão

Nesta zona, foi observado um tratamento quase total da faixa ao longo de três meses, de julho a agosto de 2017. No entanto, não houve tratamento total da faixa, devido à presença de estradas ou caminhos contidos na faixa que não sofreram intervenções. Ao analisar as observações de satélite, verificou-se que a maioria das intervenções ocorreu nos meses apresentados no formato vetorial, mas foram observadas algumas intervenções noutros meses (Figura 4.2).



(a) Formato vetorial.

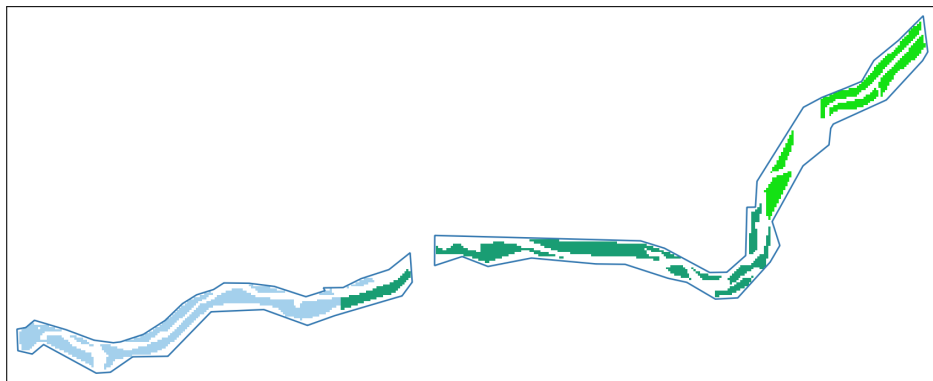
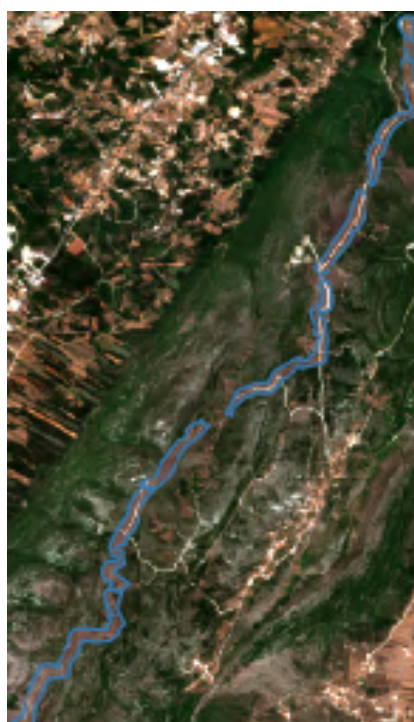
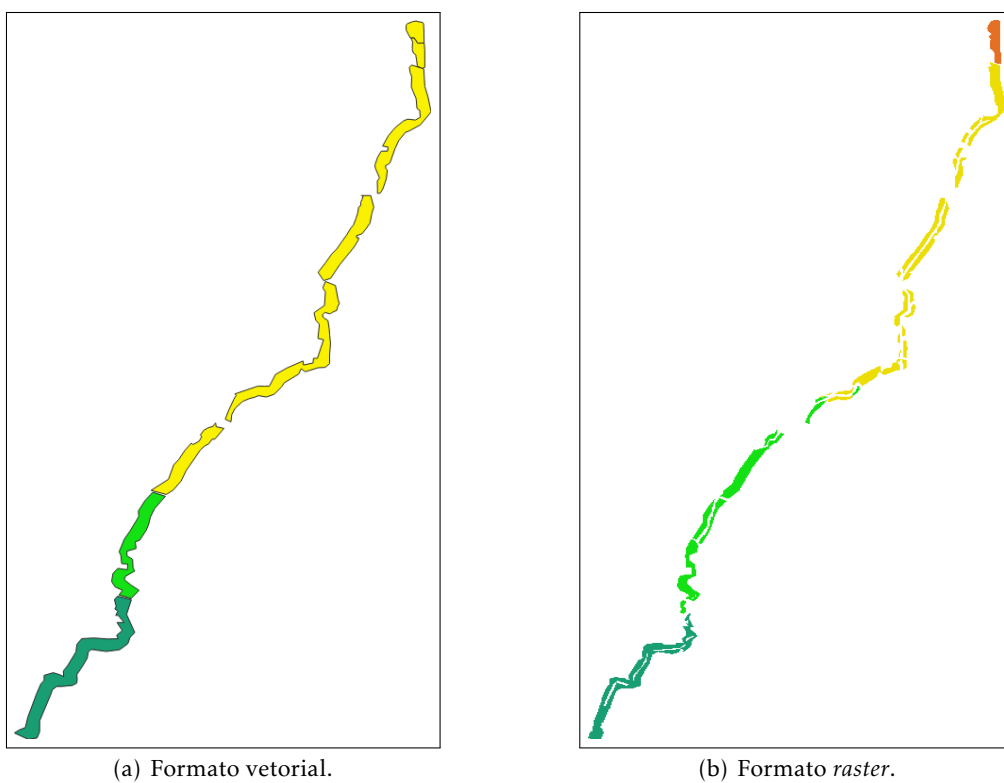
(b) Formato *raster*.(c) *True Color Image* do mês de setembro de 2017

Figura 4.2: Classificação anual quanto às intervenções numa faixa da zona do Fundão para o ano 2017.

4.1.2 Serra dos Candeeiros

Já para esta zona, foi verificado um tratamento quase total da faixa ao longo de quatro meses, de abril a julho de 2017. Novamente, a presença de estradas ou caminhos que não sofreram intervenções impediu que o tratamento fosse total. Ao inspecionar as observações de satélite, foram observadas algumas diferenças relativamente aos resultados em formato vetorial, principalmente na parte superior e numa parte do meio da faixa, em que se notou que houve tratamento em abril e julho, respetivamente (Figura 4.3).



(c) *True Color Image* do mês de agosto de 2017

Figura 4.3: Classificação anual quanto às intervenções numa faixa da zona da Serra dos Candeeiros para o ano 2017.

4.1.3 Sertã

Para esta zona, foi observado um tratamento quase total da faixa superior em maio de 2017, enquanto que na faixa da direita foram observados alguns pequenos tratamentos em junho de 2017. Não foi possível realizar uma análise sobre a faixa da esquerda, devido à falta de imagens do *Sentinel-2* do nível 2A para essa área em 2017 (Figura 4.4).

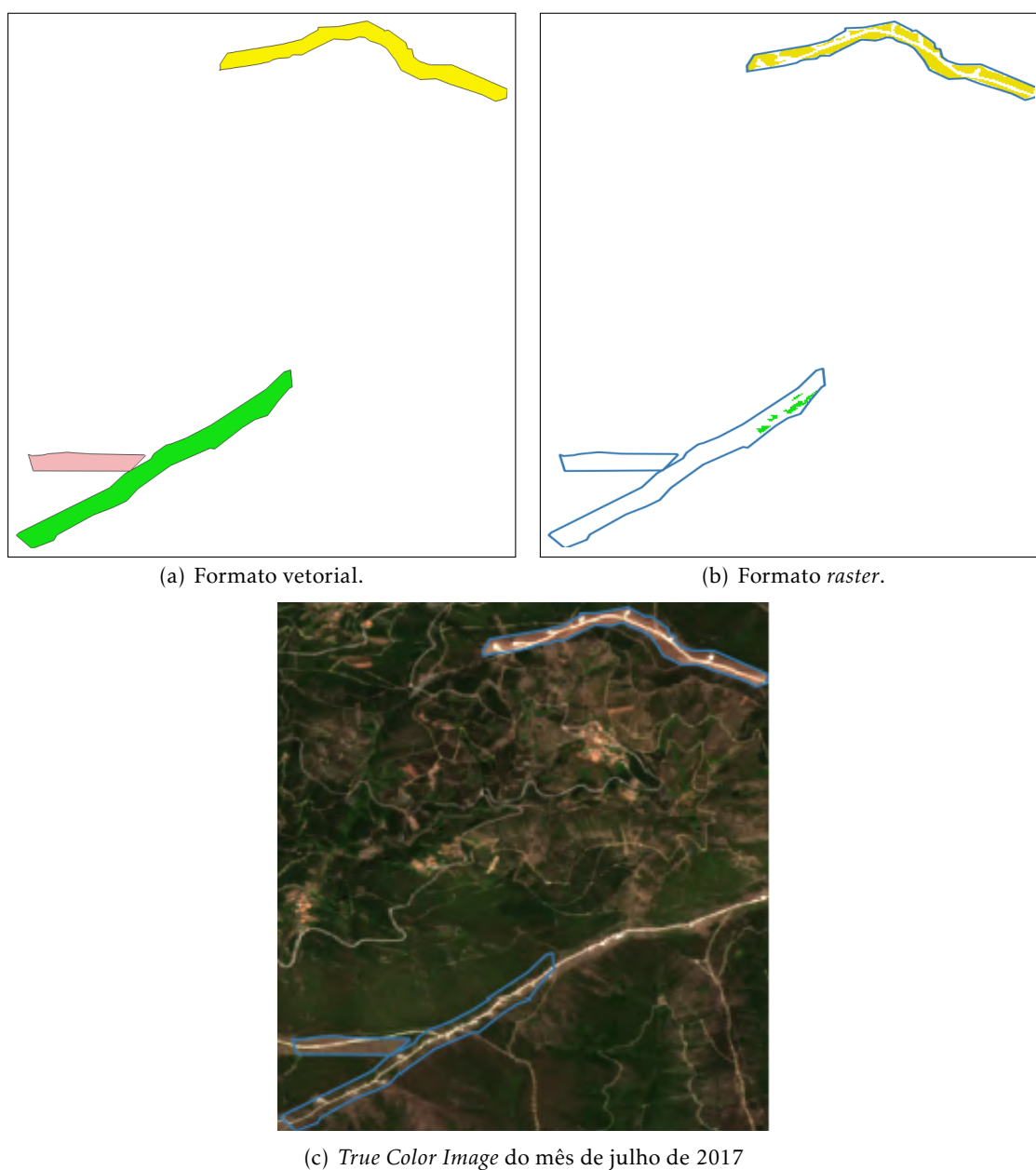


Figura 4.4: Classificação anual quanto às intervenções numa faixa da zona de Sertão para o ano 2017.

4.1.4 Seia

Finalmente, para esta zona, não foi possível criar a imagem *raster* através da inspeção das observações de satélite devido à falta de imagens perceptíveis, sem nuvens e/ou outros efeitos atmosféricos para os meses de maio, junho e julho de 2017. Por isso, esta zona foi considerada uma má escolha para área de estudo e não será analisada posteriormente (Figura 4.5).

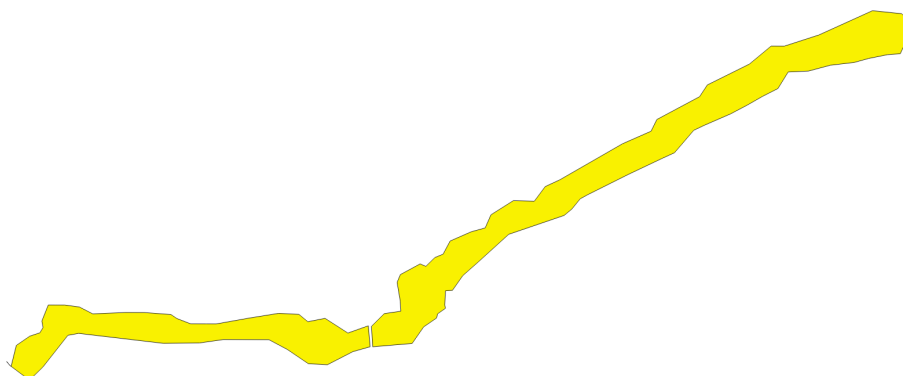


Figura 4.5: Classificação anual quanto às intervenções numa faixa da zona de Seia em formato vetorial para o ano 2017.

4.2 Resultados obtidos

Nesta secção, serão apresentados os resultados obtidos nesta dissertação para a realização dos vários testes mencionados no capítulo anterior, bem como uma avaliação dos mesmos, de modo a perceber quais são as melhores especificações para a metodologia em questão.

Na tabela 4.1, é indicado o número de observações dos anos de 2017 e 2018 usadas para cada região e o *tile* em que cada região se enquadra. Foram utilizadas observações a partir do mês de abril de 2017 até dezembro de 2017 e observações de janeiro de 2018.

Tabela 4.1: Número de observações usadas para cada região.

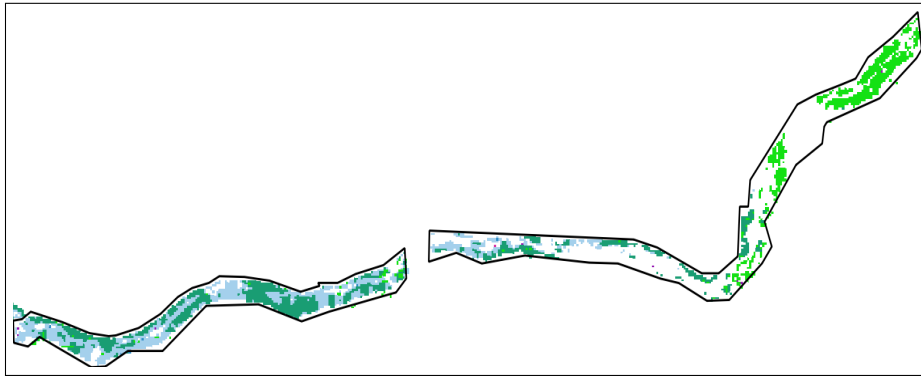
	2017	2018	<i>Tile</i>
Fundão	21	2	T29TPE
Serra dos Candeeiros	34	2	T29SND
Sertã	19	2	T29TNE

Para avaliar os resultados, foi utilizada a equação abaixo, que permite apurar a exatidão dos vários testes da metodologia, em percentagem. Esta equação ajuda a perceber qual a percentagem de semelhança entre os resultados obtidos com os da validação. Para o cálculo deste valor são usados o número de positivos verdadeiros (PV), número de negativos verdadeiros (NV), o número de positivos falsos (PF) e o número de negativos falsos (NF).

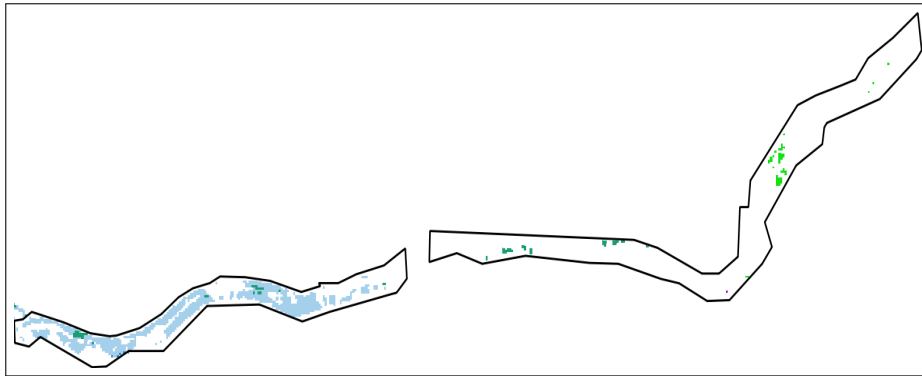
$$\text{Exatidão} = \frac{\text{Número de previsões corretas}}{\text{Número total de previsões efetuadas}} = \frac{PV + NV}{PV + NV + PF + NF} \quad (4.1)$$

4.2.1 Variação do nível de significância (α)

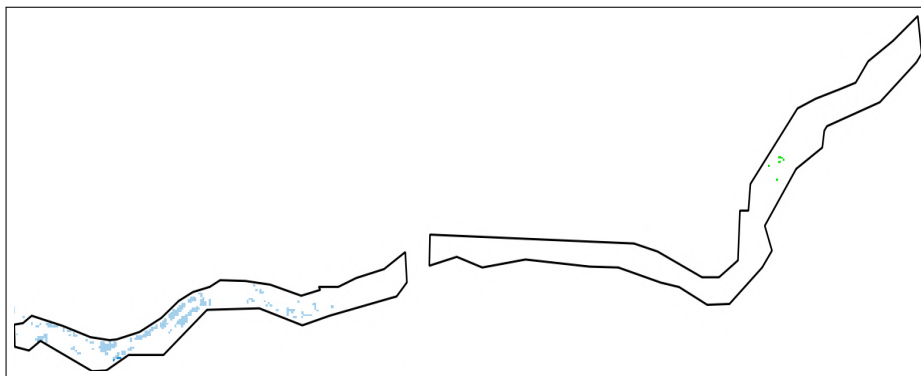
4.2.1.1 Fundão



(a) $\alpha = 0,05$



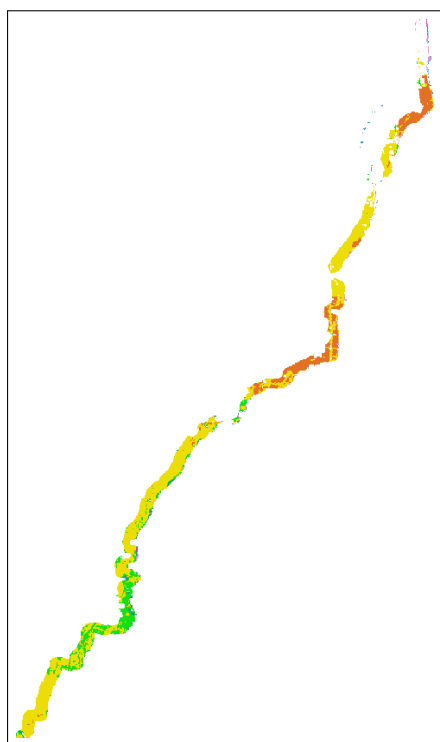
(b) $\alpha = 0,005$



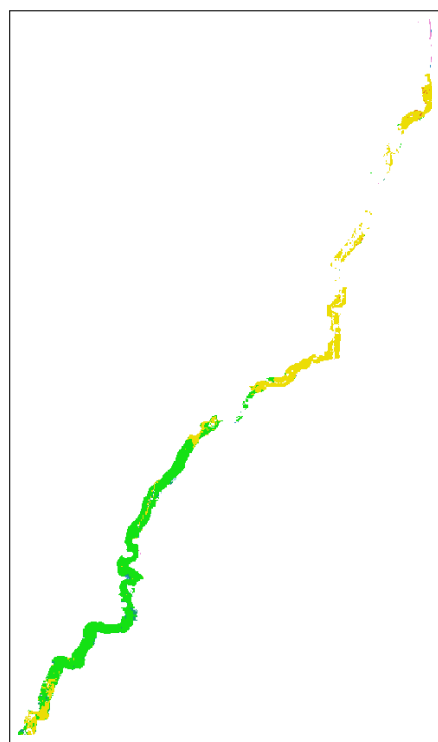
(c) $\alpha = 0,0005$

Figura 4.6: Resultados obtidos para a zona do Fundão para vários valores de α .

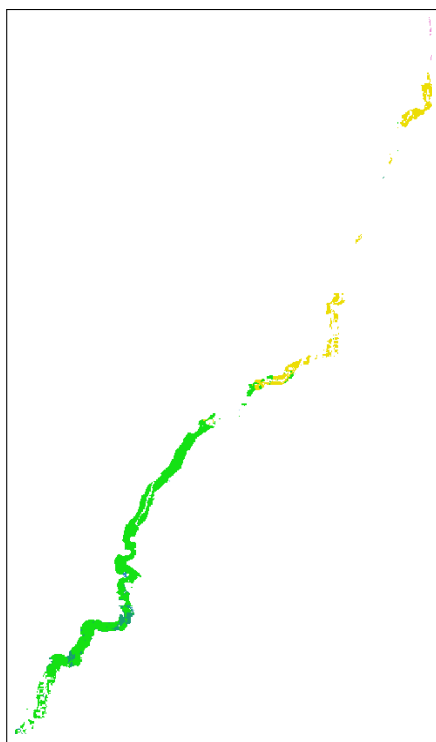
4.2.1.2 Serra dos Candeeiros



(a) $\alpha = 0,05$



(b) $\alpha = 0,005$



(c) $\alpha = 0,0005$

Figura 4.7: Resultados obtidos para a zona da Serra dos Candeeiros para vários valores de α .

4.2.1.3 Sertã

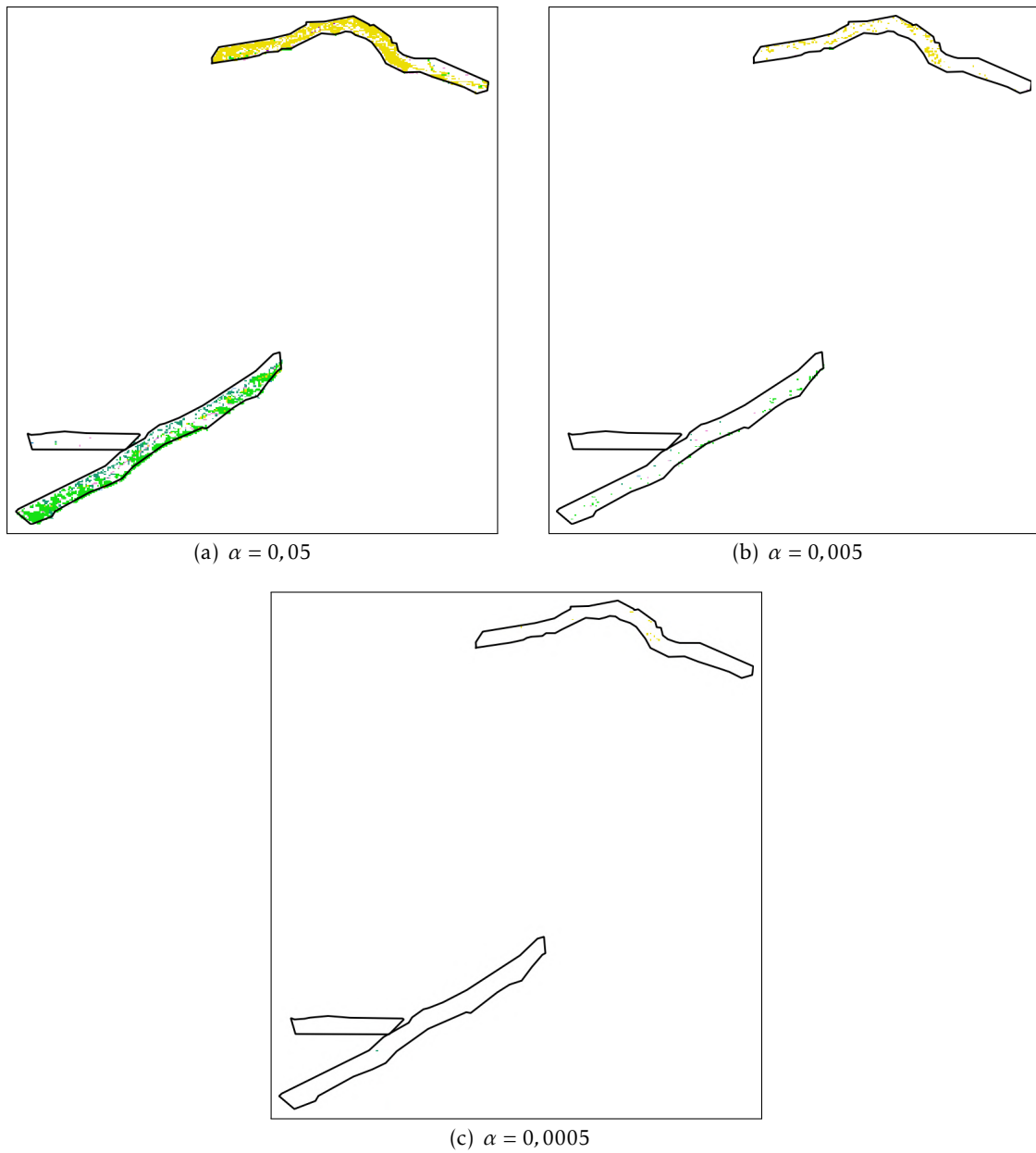


Figura 4.8: Resultados obtidos para a zona da Sertã para vários valores de α .

Tabela 4.2: Exatidão para cada valor do α de cada região

	0,05	0,005	0,0005
Fundão	59.4%	62.8%	59.0%
Serra dos Candeeiros	21.54%	47.3%	49.6%
Sertã	57.0%	74.0%	72.2%

Na tabela 4.2, estão apresentados os valores de exatidão de cada zona para vários valores de α . De modo geral, observou-se que, para valores de α altos (ex: 0,05), a metodologia deteta mais píxeis como tratados, não necessariamente bem classificados, e diminuindo o valor de α , diminui-se o número de píxeis mal classificados, tornando a metodologia menos sensível. Para valores baixos de α (ex: 0,0005), pode ocorrer a não detecção de zonas que sofreram tratamentos, como se observa nas imagens acima apresentadas.

Outro fator importante que influencia a qualidade dos resultados é o número de observações disponíveis, como se observa nos casos de Fundão e Sertã, em que a classificação é quase inexistente para o valor de $\alpha = 0,0005$, enquanto no caso da Serra dos Candeeiros, observam-se troços da faixa completamente classificados.

Quanto aos testes de vários valores α consoante a classe da COS, analisou-se a classe floresta, com o valor 5 na COS, e a classe mato, com o valor 6 na COS. A faixa da zona do Fundão é, na maioria, classificada como floresta e algumas pequenas áreas de mato. As outras duas zonas são, na maioria, classificadas como matos e com pequenas áreas de floresta. Concluiu-se que, para a classe mato, o número de detecções diminui mais do que na classe floresta com a diminuição do valor de α .

4.2.2 Teste de normalidade

4.2.2.1 Fundão

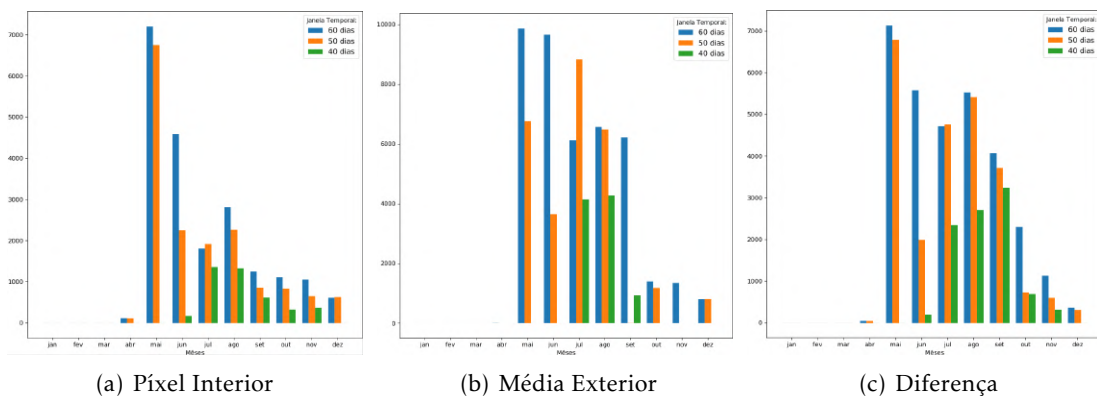
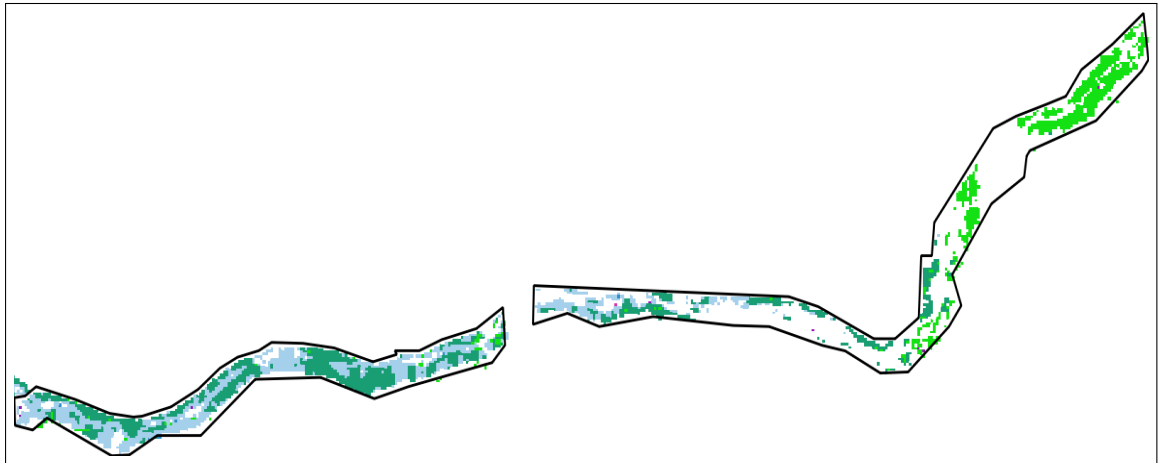
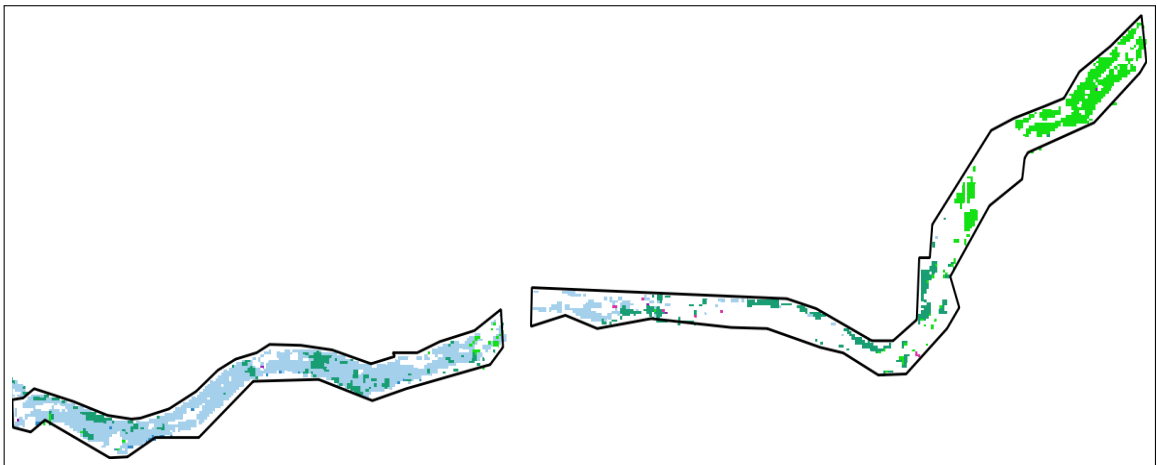


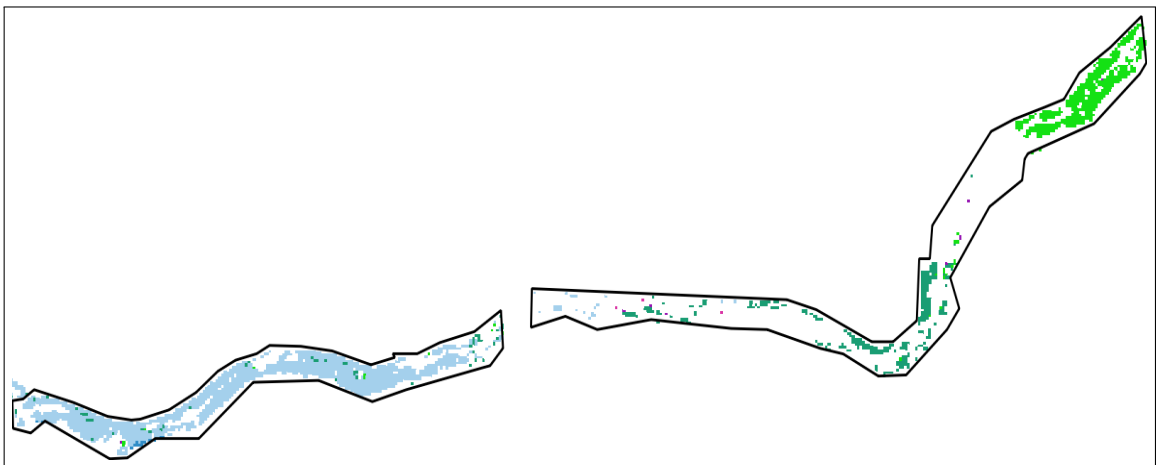
Figura 4.9: Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,05$.



(a) Janela Temporal = 60 dias.

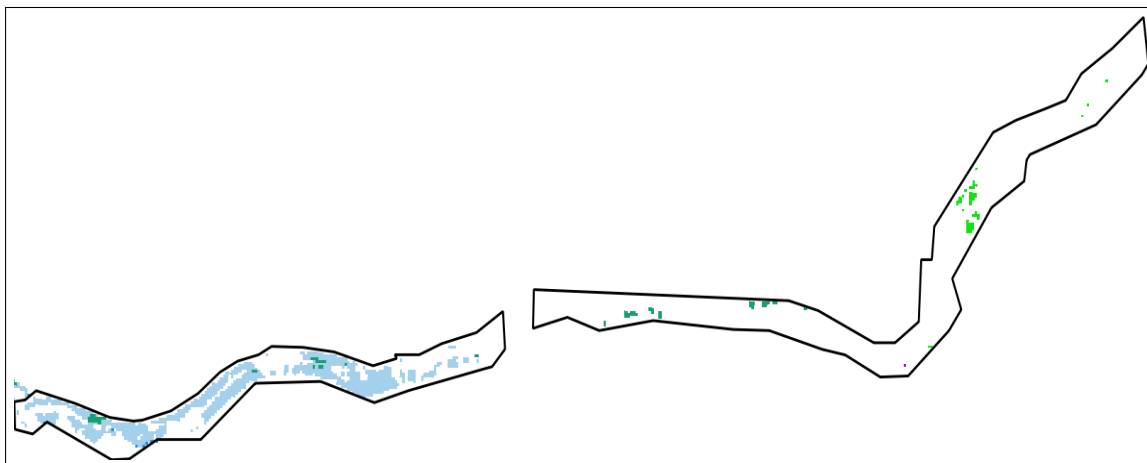


(b) Janela Temporal = 50 dias.

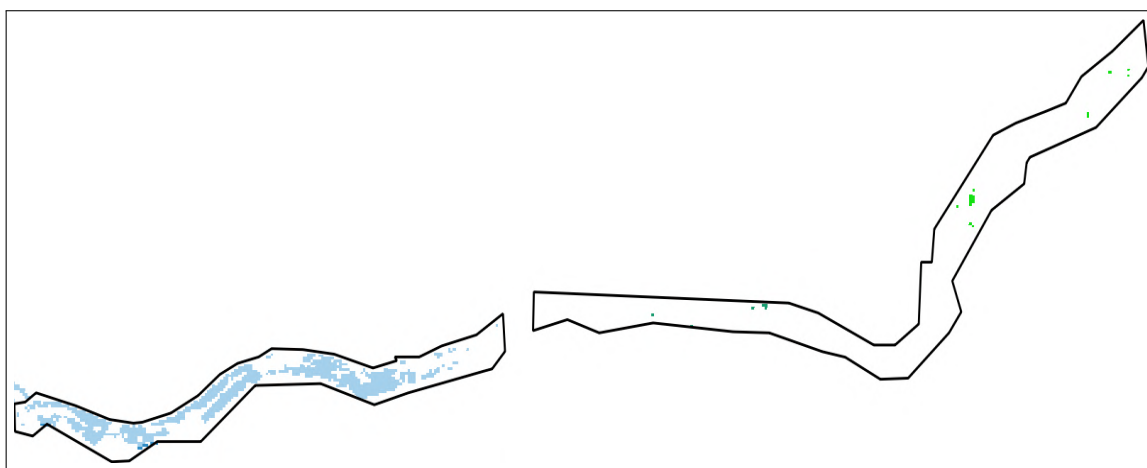


(c) Janela Temporal = 40 dias.

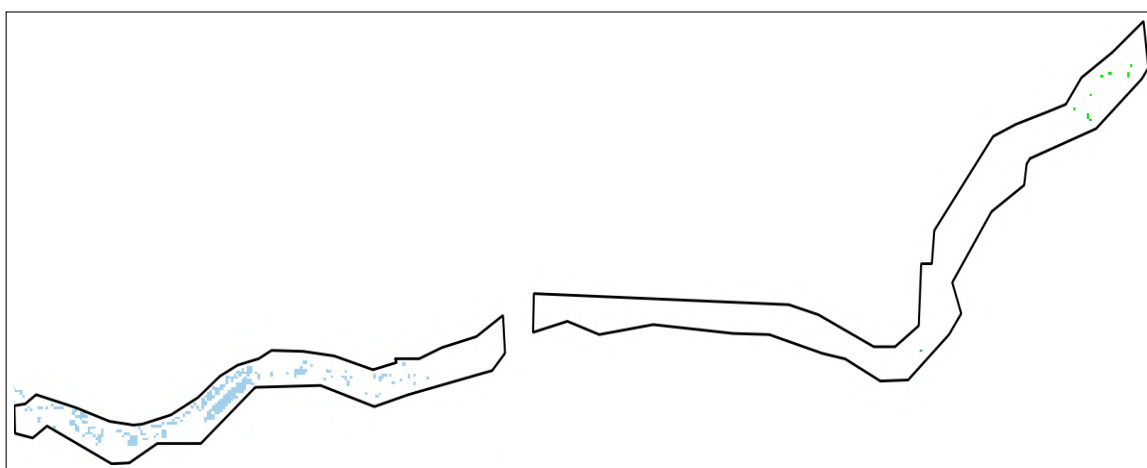
Figura 4.10: Resultados obtidos para a zona do Fundão para $\alpha = 0,05$ variando a janela temporal.



(a) Janela Temporal = 60 dias.



(b) Janela Temporal = 50 dias.



(c) Janela Temporal = 40 dias.

Figura 4.11: Resultados obtidos para a zona do Fundão para $\alpha = 0,005$ variando a janela temporal.

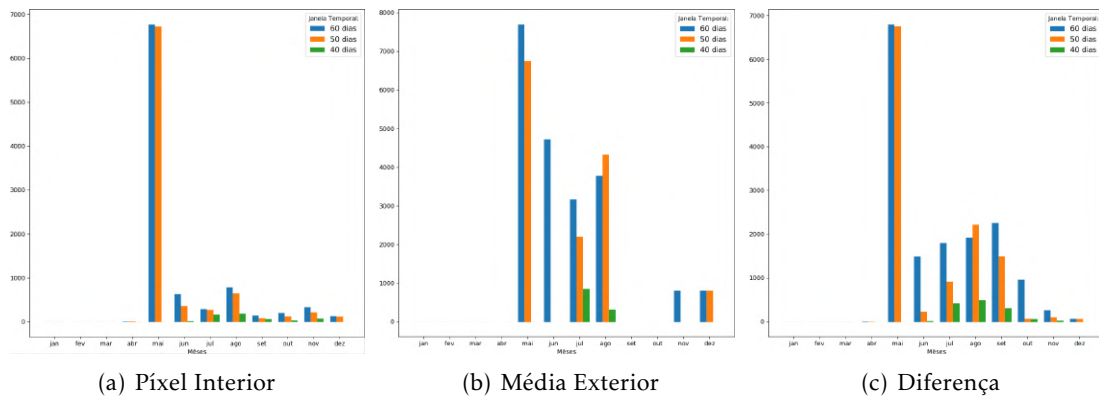


Figura 4.12: Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,005$.

Tabela 4.3: Exatidão para a zona do Fundão variando a janela temporal.

α	60 dias	50 dias	40 dias
0,05	59.4%	65.1%	67.4%
0,005	62.8%	63.3%	59.1%
0,0005	59.0%	57.5%	56.2%

Após analisar os resultados para o valor de $\alpha = 0,05$, foi possível observar que a redução da janela temporal resultou num menor número de ocorrências de falhas de normalidade nos conjuntos de dados. Isso levou a uma melhoria de 8% na exatidão dos resultados.

No caso do valor de $\alpha = 0,005$, também foi possível observar uma diminuição no número de ocorrências de falhas de normalidade ao diminuir a janela temporal. No entanto, é importante salientar que a redução da janela temporal para valores muito baixos pode resultar em perda de exatidão na classificação, assim como ocorre com a variação do nível de significância. Em outras palavras, a diminuição da janela temporal torna a metodologia mais rigorosa, acabando por detetar menos píxeis.

Para o valor de $\alpha = 0,0005$, observa-se a mesma tendência, sendo que a diminuição da janela temporal levou a uma queda na qualidade da classificação. Por isso, os resultados para esse valor não são apresentados, apenas a exatidão calculada.

4.2.2.2 Serra dos Candeeiros

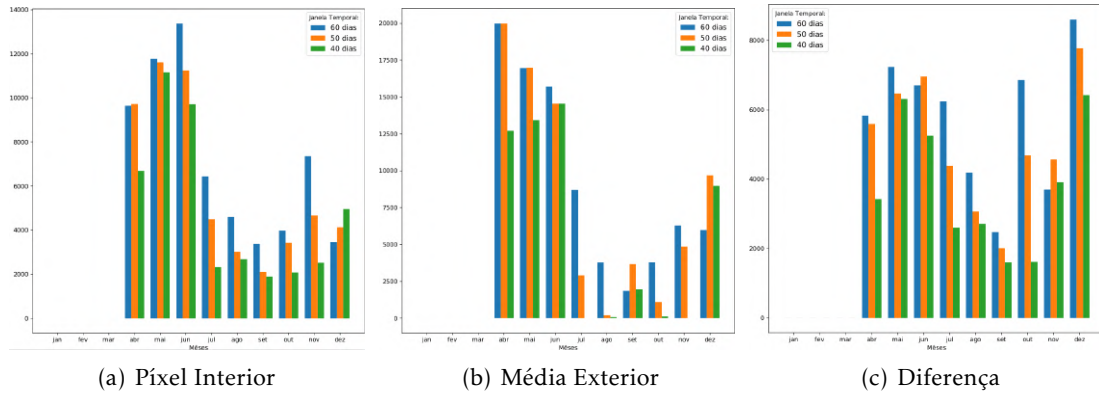


Figura 4.13: Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,05$.

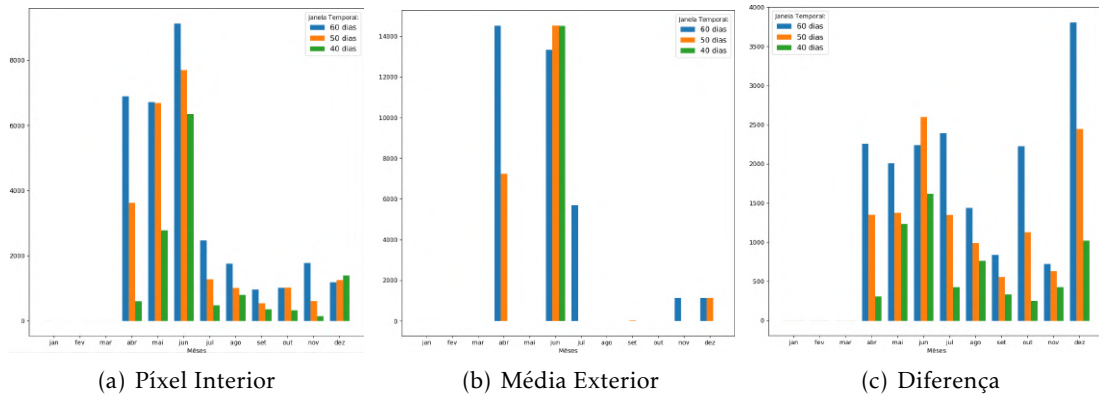


Figura 4.14: Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,005$.

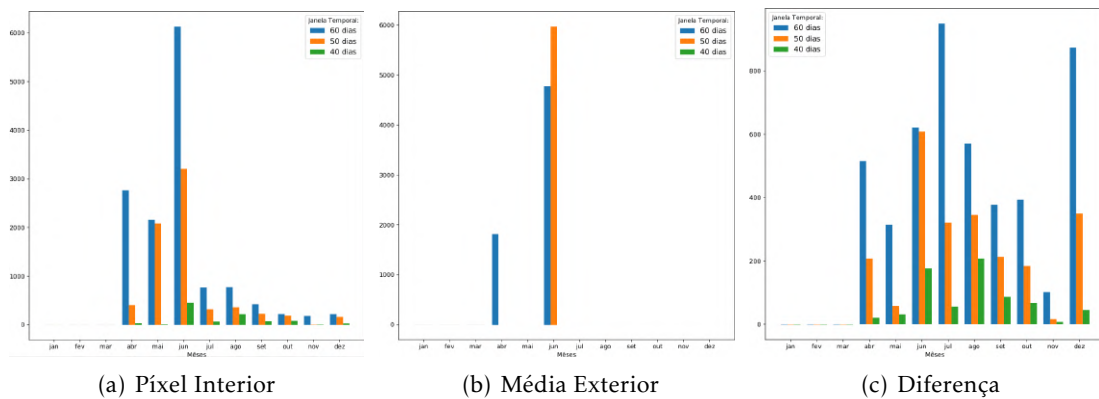
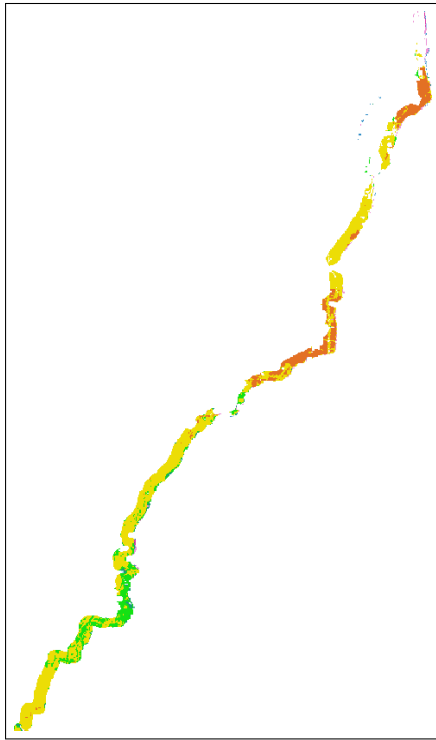
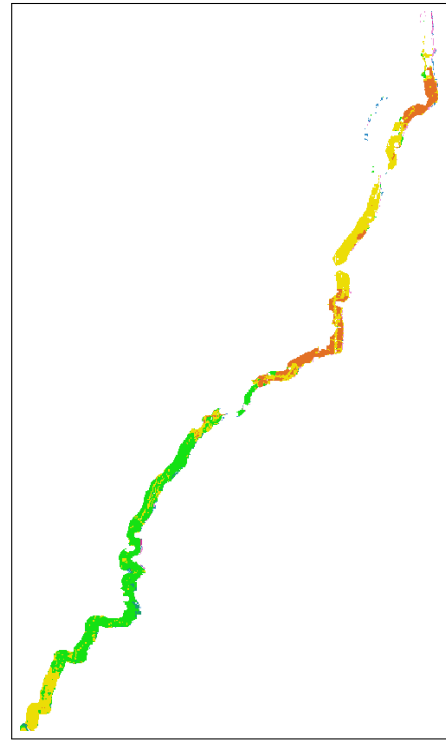


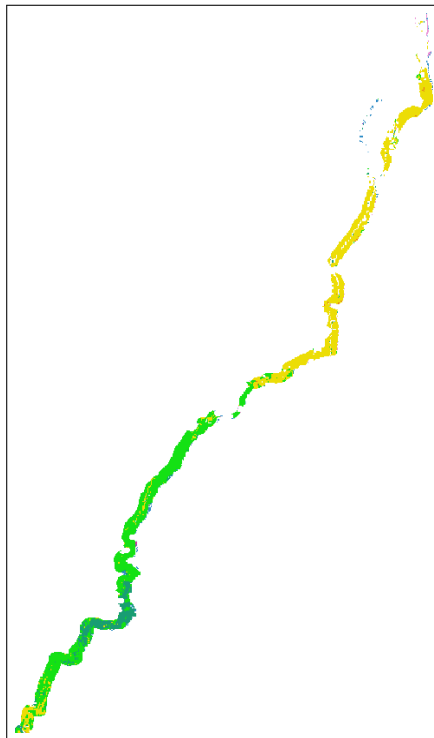
Figura 4.15: Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,0005$.



(a) Janela Temporal = 60 dias.

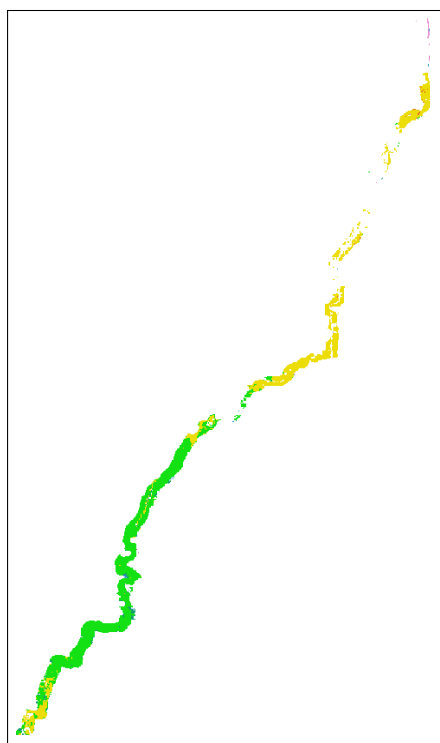


(b) Janela Temporal = 50 dias.

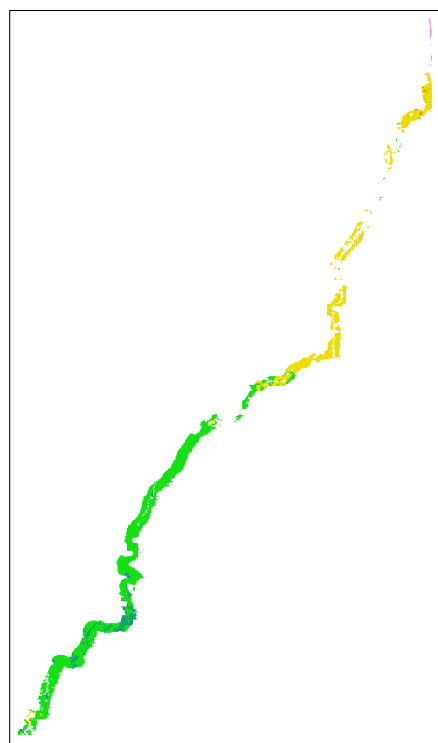


(c) Janela Temporal = 40 dias.

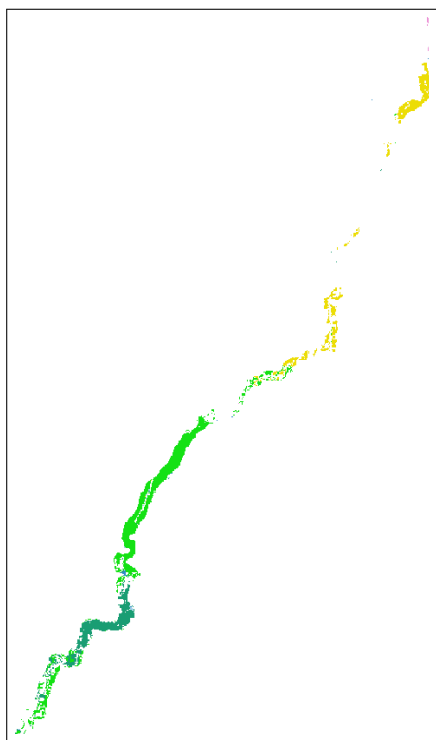
Figura 4.16: Resultados obtidos para a zona do Fundão para $\alpha = 0,05$ variando a janela temporal.



(a) Janela Temporal = 60 dias.

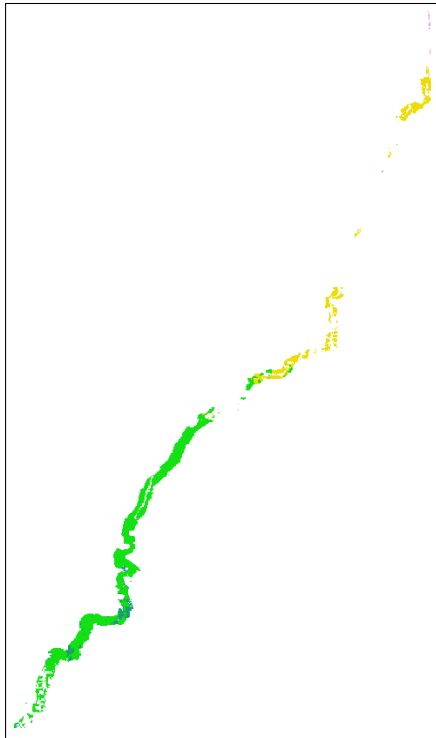


(b) Janela Temporal = 50 dias.

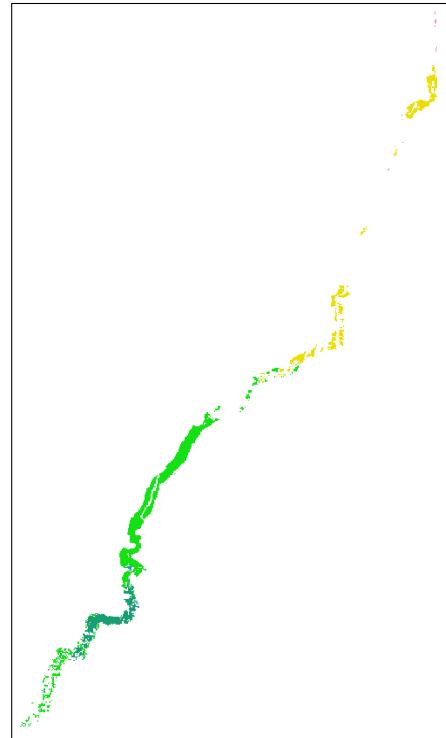


(c) Janela Temporal = 40 dias.

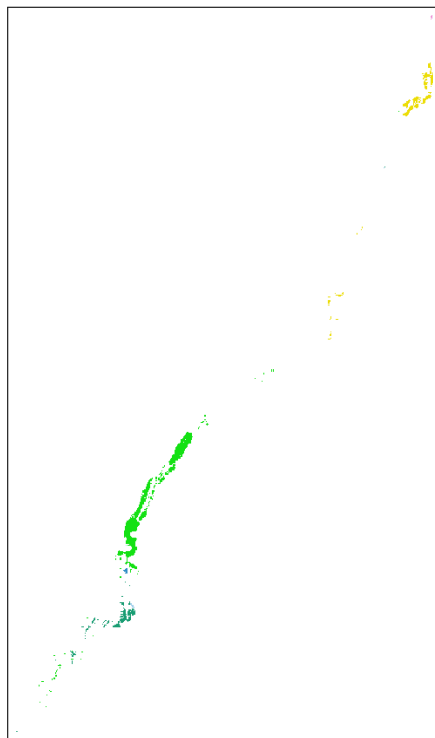
Figura 4.17: Resultados obtidos para a zona do Fundão para $\alpha = 0,005$ variando a janela temporal.



(a) Janela Temporal = 60 dias.



(b) Janela Temporal = 50 dias.



(c) Janela Temporal = 40 dias.

Figura 4.18: Resultados obtidos para a zona do Fundão para $\alpha = 0,0005$ variando a janela temporal.

Tabela 4.4: Exatidão para a zona da Serra dos Candeeiros variando a janela temporal.

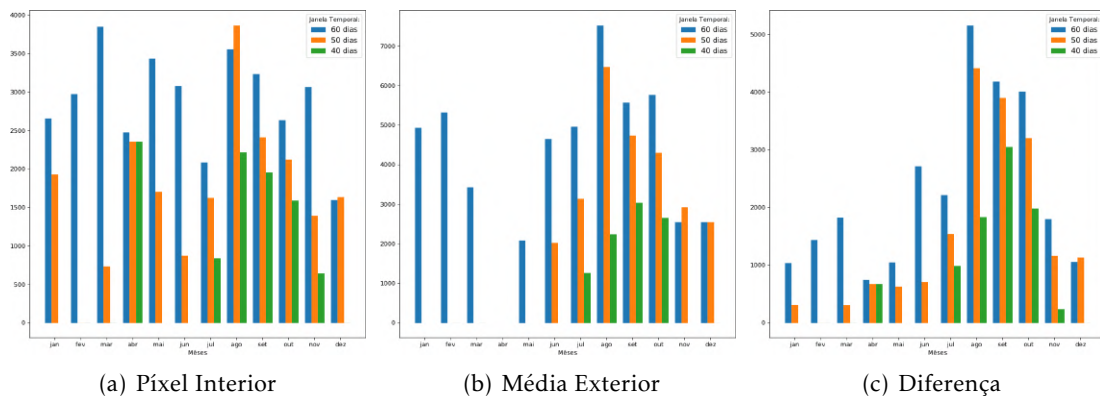
α	60 dias	50 dias	40 dias
0,05	21.5%	35.3%	55.2%
0,005	47.3%	53.5%	54.9%
0,0005	49.6%	55.6%	44.5%

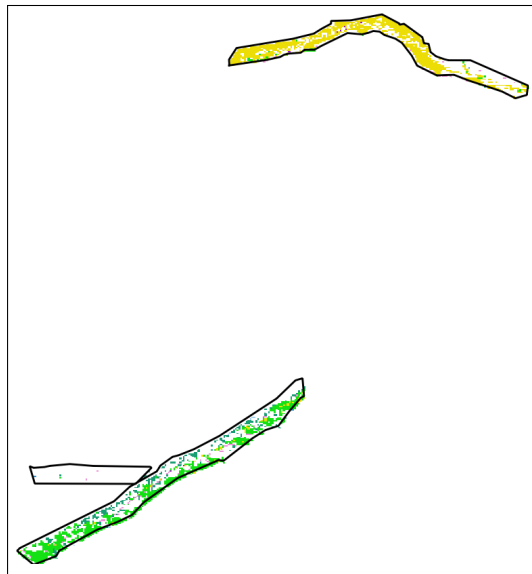
Ao reduzir a janela temporal para o valor de $\alpha = 0,05$, observou-se uma diminuição no número de falhas de normalidade nos conjuntos de dados, o que resultou numa melhoria de 33.7% na exatidão dos resultados obtidos.

No caso do valor de $\alpha = 0,005$, também se verificou uma diminuição no número de ocorrências de falhas de normalidade ao reduzir a janela temporal, obtendo-se uma melhoria de 7.6%.

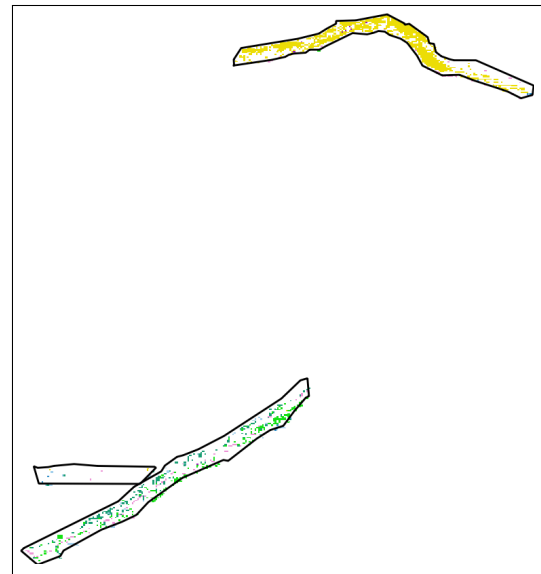
Para os resultados do valor de $\alpha = 0,0005$, observa-se uma tendência semelhante à verificada na zona do Fundão para o valor de $\alpha = 0,005$. Neste caso, também se verificou que a exatidão melhorou para a janela temporal de 50 dias, mas piorou para a janela temporal de 40 dias, onde se registou o mínimo número de ocorrências. Isso reforça a teoria de que, quando o número de ocorrências se aproxima de valores muito baixos, a qualidade da classificação diminui.

4.2.2.3 Sertão

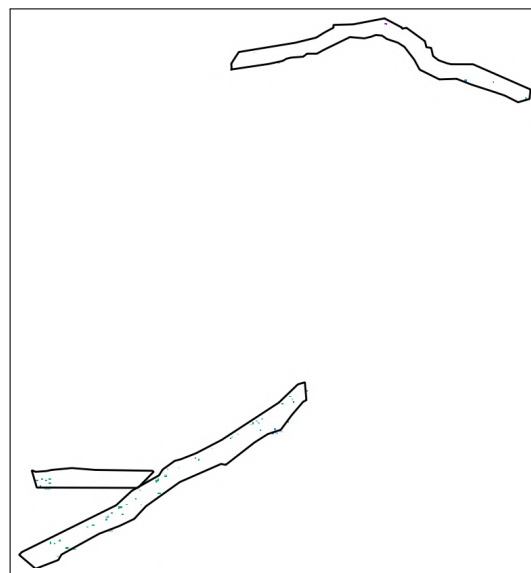
Figura 4.19: Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,05$.



(a) Janela Temporal = 60 dias.



(b) Janela Temporal = 50 dias.



(c) Janela Temporal = 40 dias.

Figura 4.20: Resultados obtidos para a zona do Fundão para $\alpha = 0,05$ variando a janela temporal.

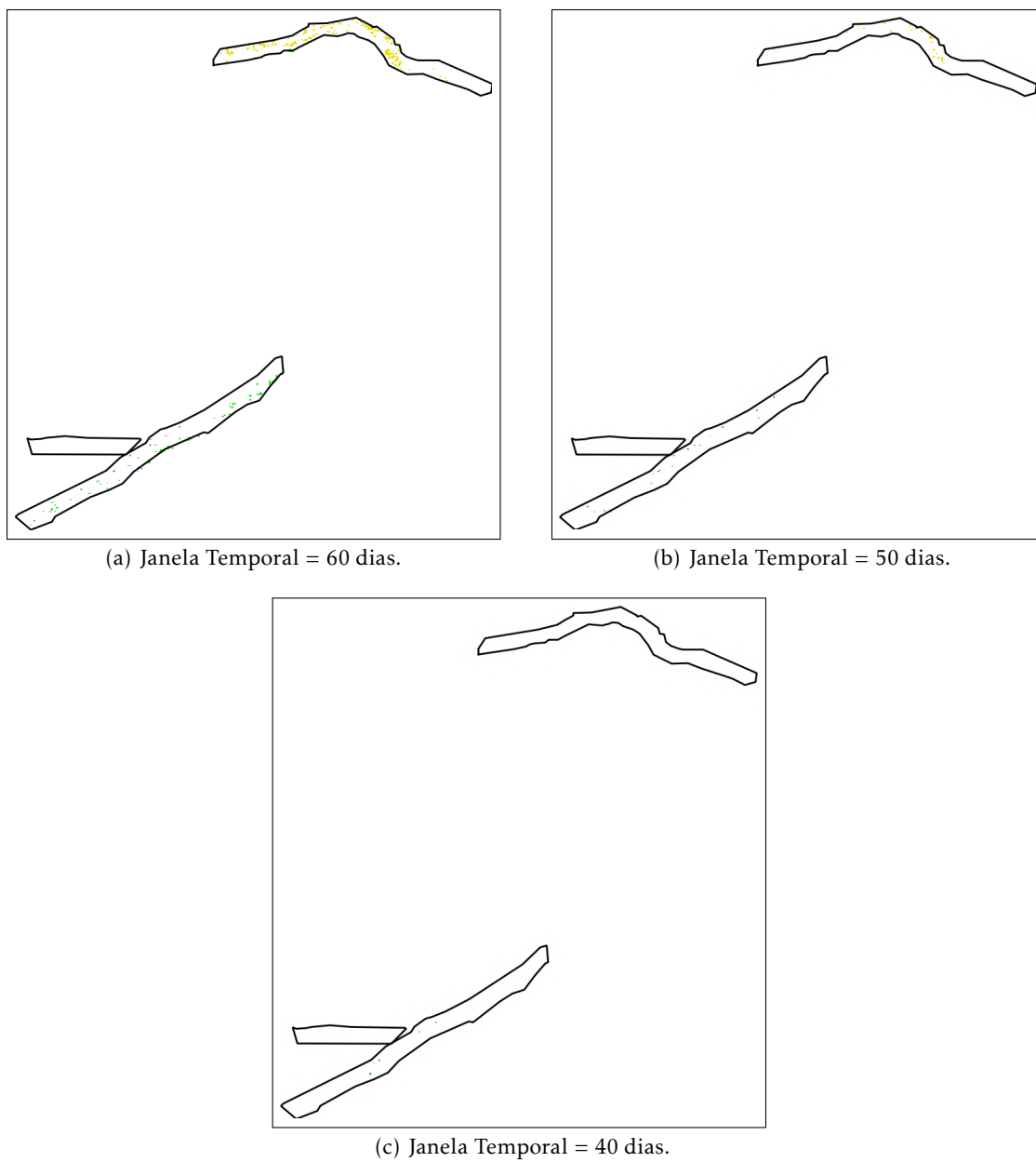


Figura 4.21: Resultados obtidos para a zona do Fundão para $\alpha = 0,005$ variando a janela temporal.

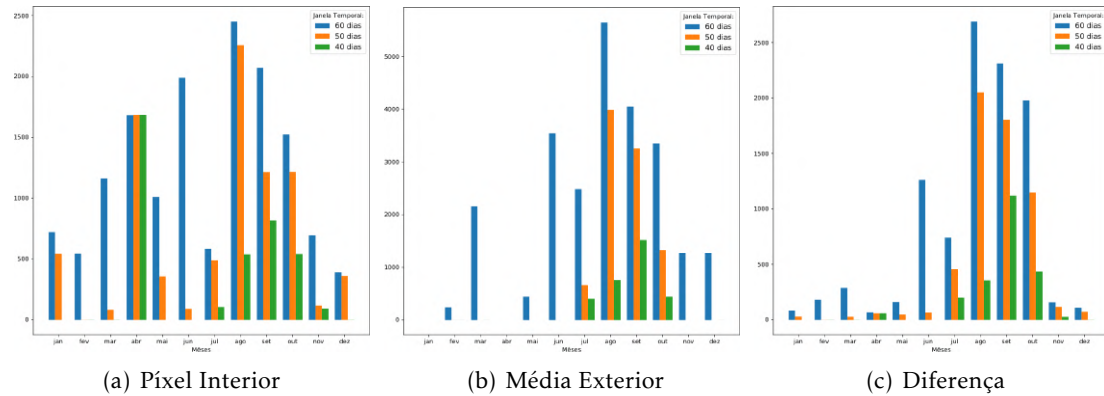


Figura 4.22: Ocorrências de falhas de normalidade nos conjuntos de dados para $\alpha = 0,005$.

Tabela 4.5: Exatidão para a zona da Sertã variando a janela temporal.

α	60 dias	50 dias	40 dias
0,05	57.0%	71.4%	70.6%
0,005	74.0%	72.0%	71.8%
0,0005	72.2%	71.9%	71.9%

Ao reduzir a janela temporal para o valor de $\alpha = 0,05$, verificou-se uma diminuição no número de ocorrências de falhas de normalidade nos conjuntos de dados, o que resultou numa melhoria de 14.4% na exatidão dos resultados obtidos.

Para o valor de $\alpha = 0,005$, também se observa que diminuindo a janela temporal o número de ocorrências de falhas de normalidade nos conjuntos de dados diminui. Contudo, neste caso não foi possível obter resultados melhores. Isto pode acontecer pelo facto de se usar um número baixo de observações, e diminuir ainda mais a janela temporal poderá levar à redução ou mesmo à não criação de alguns conjuntos de dados para a classificação.

Quanto ao valor de $\alpha = 0,0005$, verificou-se a mesma tendência. No entanto, os resultados para este caso não são apresentados, uma vez que se verificou uma diminuição da qualidade da classificação ao diminuir a janela temporal. Observou-se que, para janelas temporais de 50 e 40 dias, não foi classificado nenhum píxel, o que reforça a importância do número de observações.

Apesar de a exatidão ter um valor alto para este caso, isso deve-se ao facto de a faixa inferior apresentar um grande número de negativos verdadeiros (NV).

4.2.3 Sumário

Tabela 4.6: Exatidão obtida variando a janela temporal.

α	Janela temporal	Fundão	Serra dos Candeeiros	Sertão
0,05	60 dias	59.4%	21.5%	57.0%
	50 dias	65.1%	35.3%	71.4%
	40 dias	67.4%	55.2%	70.6%
0,005	60 dias	62.8%	47.3%	74.0%
	50 dias	63.3%	53.5%	72.0%
	40 dias	59.1%	54.9%	71.8%
0,0005	60 dias	59.0%	49.6%	72.2%
	50 dias	57.5%	55.6%	71.9%
	40 dias	56.2%	44.5%	71.9%

Ao analisar os valores apresentados na Tabela 4.6, verificou-se uma melhoria em todas as zonas quando se utilizou o valor de $\alpha = 0,05$. Além disso, para a maioria das zonas, a janela temporal de 40 dias apresentou resultados melhores às outras janelas temporais testadas.

Também se verificou que, para os casos do Fundão e Sertão, onde foram usadas menos observações, a exatidão diminuiu a partir do valor de $\alpha = 0,005$. No entanto, para a zona da Serra dos Candeeiros, essa relação não foi tão linear, sendo possível observar bons resultados para o valor de $\alpha = 0,0005$ com uma janela temporal de 50 dias.

Se for utilizado um número limitado de observações, por exemplo, duas por mês, é recomendável utilizar um valor de α entre 0,05 e 0,005, sendo recomendável uma janela temporal inferior a 60 dias para o valor de α de 0,05. No caso de se ter um bom número de observações, todos os valores de α podem ser utilizados, com janelas temporais menores que 60 dias para valores de α entre 0,05 e 0,0005 e com uma janela temporal entre os 50 e 60 dias para o valor de α de 0,0005. Deve-se ter em conta que a escolha do valor de α tem implicação no quanto esta metodologia deteta, para um valor alto deteta mais do que num valor baixo. Portanto, caso se pretenda uma classificação mais rigorosa, é recomendável utilizar um valor de α baixo.

É importante mencionar que, tendo em conta o número de casos analisados, não é correto generalizar esta conclusão, mesmo que seja verdadeira para os casos estudados. Para consolidar esta conclusão seria necessário realizar uma análise para um maior número de casos.

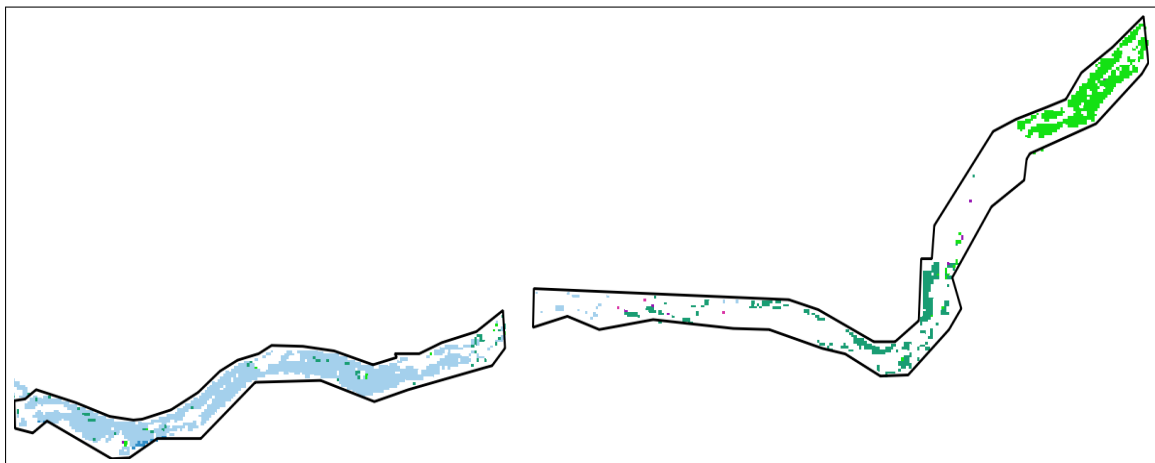
4.2.4 Pós-processamento

Para realizar os testes do pós-processamento, foram selecionados os casos com melhor exatidão apresentada nos testes anteriores. Para cada zona foram utilizados os seguintes parâmetros:

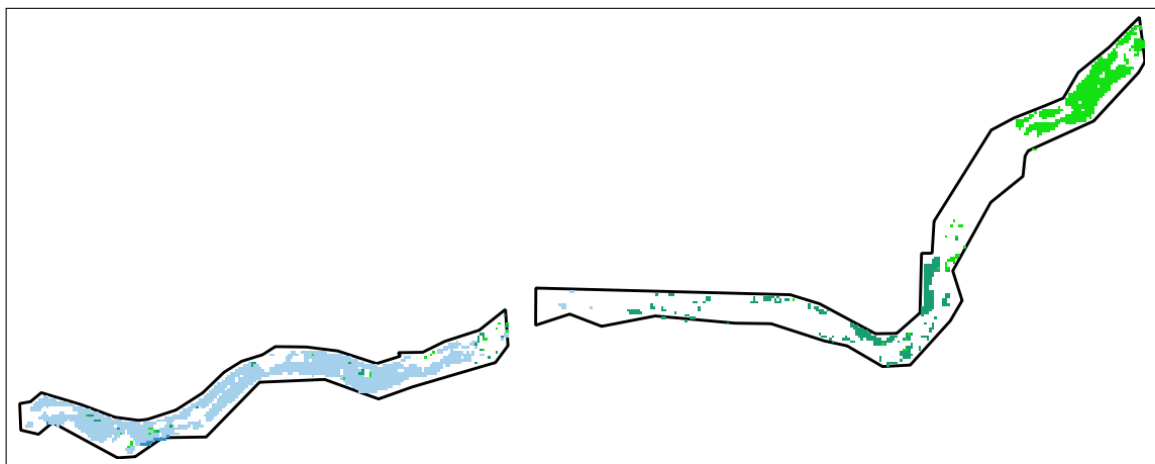
- **Fundão:** valor de $\alpha = 0,05$ e janela temporal de 40 dias;
- **Serra dos Candeeiros:** valor de $\alpha = 0,0005$ e janela temporal de 50 dias;
- **Sertão:** valor de $\alpha = 0,005$ e janela temporal de 60 dias.

De modo geral, ao observar as imagens obtidas, percebe-se que após o pós-processamento há uma redução das detecções dispersas, ou seja, dos píxeis isolados. As áreas detetadas ficam melhor delimitadas e preenchidas.

4.2.4.1 Fundão



(a) Resultados antes do pós-processamento



(b) Resultados após o pós-processamento

Figura 4.23: Resultados obtidos para a zona do Fundão aplicando um filtro de mediana.

4.2.4.2 Serra dos Candeeiros

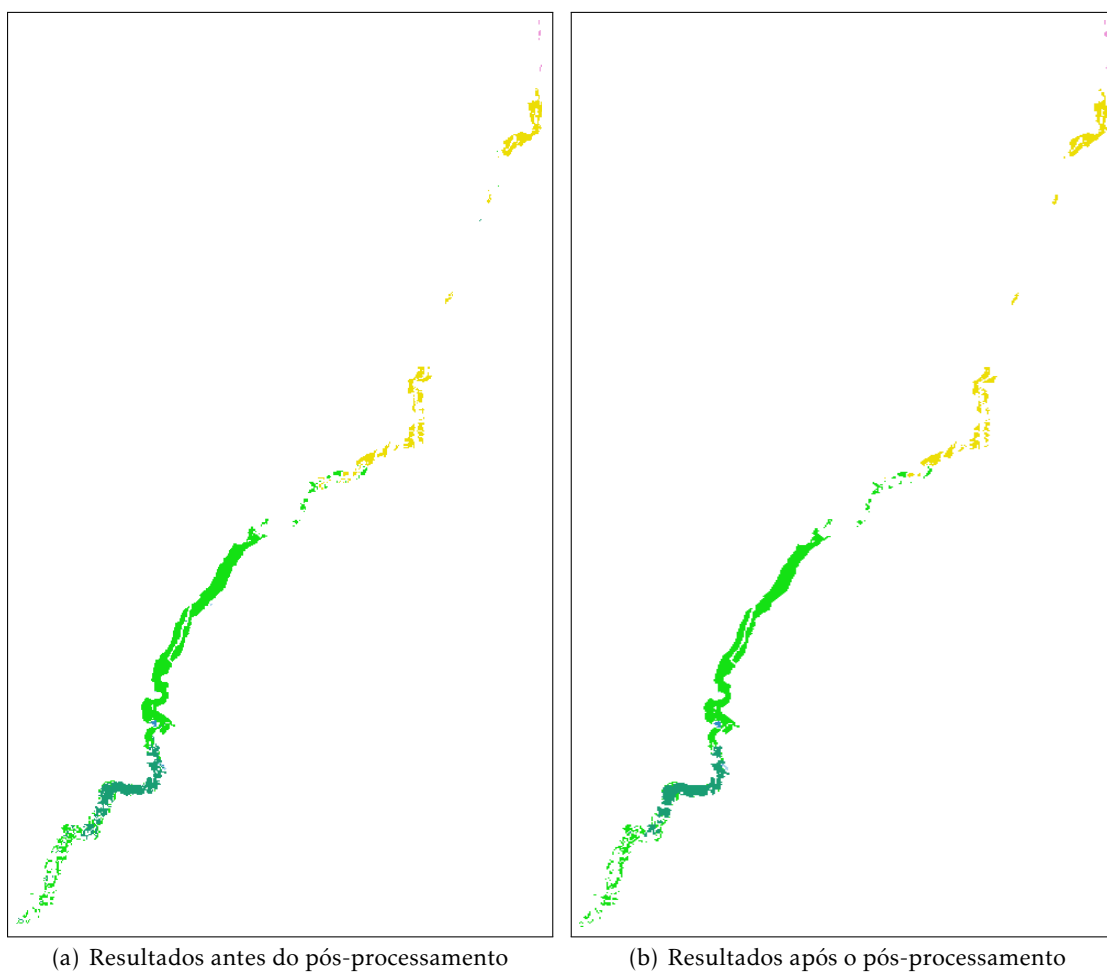
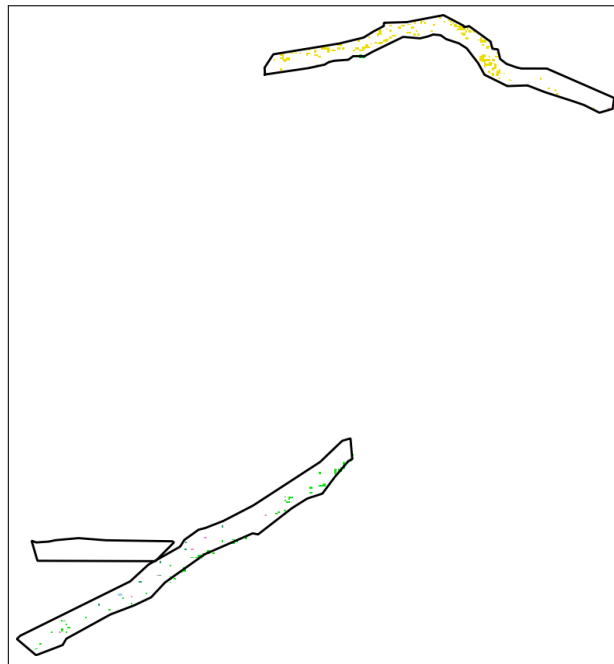
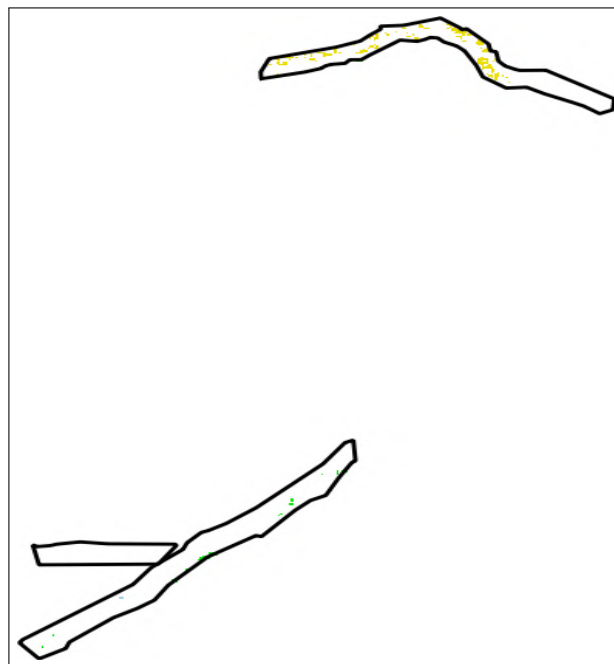


Figura 4.24: Resultados obtidos para a zona da Serra dos Candeeiros aplicando um filtro de mediana.

4.2.4.3 Sertã



(a) Resultados antes do pós-processamento



(b) Resultados após o pós-processamento

Figura 4.25: Resultados obtidos para a zona da Sertã aplicando um filtro de mediana.

CONCLUSÕES E TRABALHO FUTURO

5.1 Conclusão

O trabalho realizado nesta dissertação permitiu o desenvolvimento de três *scripts* em *Python*: um para o pré-processamento das observações de satélite, um para a metodologia de aprendizagem supervisionada com o auxílio de uma RNA e um para a metodologia de classificação com o teste estatístico *Welch's t-Test*. Utilizou-se a linguagem de programação *Python*, uma vez que o *software* QGIS permite executar *scripts* em *Python*.

Após uma primeira análise, constatou-se que a primeira metodologia não era fiável para uma abordagem a píxel, enquanto que a segunda metodologia apresentou resultados promissores. Assim sendo, o resto do trabalho focou-se em torno da segunda metodologia. Foram realizados testes à metodologia com o objetivo de apurar os parâmetros que proporcionam uma boa classificação.

No teste de variação do nível de significância (α) do *Welch's t-Test*, verificou-se que a metodologia não classifica bem para valores altos de α (ex: 0,05) e que, com a diminuição do valor de α , o número de píxeis detetados como alteração diminui. Ou seja, quanto mais baixo for o valor de α , mais rigorosa é a classificação. Também se observou que o número de observações do conjunto de dados influencia a classificação: em conjuntos de dados pequenos (cerca de duas observações por mês), valores baixos de α (ex: 0,0005) são demasiados rigorosos, pelo que é preferível a utilização de valores α superiores a esse. Em conjuntos de dados de dimensão razoável (cerca de três ou mais observações por mês), observa-se uma melhor classificação para valores baixos de α (ex: 0,0005). Ou seja, a classificação também é rigorosa, mas como o número de dados é maior, a rigidez da classificação compensa.

Outro teste realizado à metodologia consistiu na verificação da normalidade dos conjuntos de dados, uma vez que o *Welch's t-Test* tem um pressuposto estatístico em que os conjuntos de dados devem verificar uma distribuição normal. Verificou-se uma diminuição das ocorrências de falhas de normalidade com a diminuição da janela temporal usada para a criação dos conjuntos de dados. A nível dos resultados, constatou-se que para valores altos de α (ex: 0,05), os resultados melhoram, permitindo concluir que a metodologia

deteta melhor quando a maioria dos conjuntos de dados verificam a distribuição normal. Contudo, para valores baixos de α (ex: 0,0005), em que a classificação já é rigorosa, a diminuição da janela temporal torna-a ainda mais rigorosa. Por fim, verificou-se que o pós-processamento desenvolvido permite delimitar melhor as áreas de píxeis detetados e remover deteções isoladas.

5.2 Trabalho Futuro

Após a conclusão deste estudo, é possível identificar tarefas para trabalho futuro na perspectiva de melhoramento do mesmo e de criação de novos conhecimentos. Numa primeira fase, poderia ser implementado um *plugin* para QGIS que executasse os *scripts* do pré-processamento e da segunda metodologia, permitindo ao utilizador escolher os parâmetros, como por exemplo, o nível de significância (α) e a janela temporal.

Seria importante, realizar uma análise da metodologia para várias classes da carta de ocupação do solo, variando principalmente o nível de significância (α) e/ou outros parâmetros, de forma a perceber se para certas classes, alguns parâmetros poderiam melhorar a classificação e caso assim se verifique ajustar a implementação para que se utilizem os parâmetros favoráveis para cada classe. Para isso, também é necessário criar um conjunto novo de áreas de estudo, mais recentes que 2017, que incluam áreas de várias classes da COS, não se limitando apenas à classe matos e floresta.

Embora não tenha sido utilizado nesta dissertação, seria interessante incluir superpíxeis em ambas as metodologias, para analisar e comparar os resultados obtidos. Eventualmente, isso poderia ser relevante para abordagens a objeto, uma vez que os superpíxeis são objetos que simulam os píxeis.

BIBLIOGRAFIA

- Achanta, R., Shaji, A., Smith, K., Lucchi, A., Fua, P., & Ssstrunk, S. (2012). SLIC Superpixels Compared to State-of-the-Art Superpixel Methods. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 34(11), 2274–2282. <https://doi.org/10.1109/TPAMI.2012.120> (ver p. 17)
- Achanta, R., Shaji, A., Smith, K., Lucchi, A., Fua, P., & Ssstrunk, S. (2010). SLIC Superpixels. *EPFL Technical report* (ver p. 17).
- Aubard, V., Pereira-Pires, J. E., Campagnolo, M. L., Pereira, J. M. C., Mora, A., & Silva, J. M. N. (2020). Fully Automated Countrywide Monitoring of Fuel Break Maintenance Operations. *Remote Sensing*, 12(18). <https://doi.org/10.3390/rs12182879> (ver pp. 26, 28)
- Baatz, M., & Schape, A. (2000). Multiresolution Segmentation: an optimization approach for high quality multi-scale image segmentation. *Angewandte Geographische Informationsverarbeitung. Wichmann-Verlag, Heidelberg*, 12–23 (ver p. 19).
- Benz, U. C., Hofmann, P., Willhauck, G., Lingenfelder, I., & Heynen, M. (2004). Multi-resolution, object-oriented fuzzy analysis of remote sensing data for GIS-ready information. *ISPRS Journal of Photogrammetry & Remote Sensing*, 58(6), 239–258. <https://doi.org/http://dx.doi.org/10.1016/j.isprsjprs.2003.10.002> (ver p. 19)
- Breiman, L. (2001). Random Forests. *Machine Learning*, 45, 5–32. <https://doi.org/10.1023/A:1010933404324> (ver p. 21)
- Caetano, M., Igreja, C., & Marcelino, F. (2019). Especificaes tcnicas da Carta de Uso e Ocupao do Solo (COS) de Portugal Continental para 2018. Relatrio Tcnico. Direo-Geral do Territrio. (Ver p. 13).
- Campbell, J. B., & Wynne, R. H. (2011). Introduction to Remote Sensing. The Guilford Press. (Ver p. 6).
- Campbell, J., & Shin, M. (2011). Essentials of Geographic Information Systems. Saylor Foundation. (Ver p. 12).
- Chang, K.-T. (2006). Introduction to Geographic Information Systems. McGraw-Hill College. (Ver p. 12).
- DGT. (2022). *Mapas DG-Territrio*. <http://mapas.dgterritorio.gov.pt/>. (Ver p. 13)

- DPFVAP. (2014). Manual de rede primária. (Ver pp. 2, 3).
- Drăguț, L., Tiede, D., & Levick, S. R. (2010). ESP: a tool to estimate scale parameter for multiresolution image segmentation of remotely sensed data. *International Journal of Geographical Information Science*, 24(6), 859–871. <https://doi.org/https://doi.org/10.1080/13658810903174803> (ver p. 19)
- ESA. (2013). Sentinel-2 Data Sheet. https://esamultimedia.esa.int/docs/S2-Data_Sheet.pdf. (Ver p. 11)
- ESA. (s.d.-a). *Level-2A Algorithm Overview*. <https://sentinels.copernicus.eu/web/sentinel/technical-guides/sentinel-2-msi/level-2a/algorithm>. (Ver p. 30)
- ESA. (s.d.-b). *The Sentinel missions*. https://www.esa.int/Applications/Observing_the_Earth/Copernicus/The_Sentinel_missions. (Ver p. 10)
- ESA. (s.d.-c). *Sentinel Overview*. <https://sentinel.esa.int/web/sentinel/missions>. (Ver p. 9)
- ESA. (s.d.-d). *Spatial Resolution*. <https://sentinels.copernicus.eu/web/sentinel/user-guides/sentinel-2-msi/resolutions/spatial>. (Ver p. 11)
- Espectro Eletromagnético. (2017). <https://estudodacor.wordpress.com/aspectos-fisicos/espectro-da-luz/espectro-2/espectro-eletromagnetico-2/#main>. (Ver p. 7)
- Felzenszwalb, P. F., & Huttenlocher, D. P. (2004). Efficient Graph-Based Image Segmentation. *International Journal of Computer Vision*, 59, 167–181. <https://doi.org/https://doi.org/10.1023/B:VISI.0000022288.19776.77> (ver p. 16)
- Jinlin, C., Chunzhi, Y., Guangkui, X., & Ning, L. (2018). Image Segmentation Method Using Fuzzy C Mean Clustering Based on Multi-Objective Optimization. *Journal of Physics: Conference Series*, 1004, 87–94. <https://doi.org/10.1088/1742-6596/1004/1/012035> (ver p. 18)
- Juneja, P., & Kashyap, R. (2016). Energy based Methods for Medical Image Segmentation. *International Journal of Computer Applications*, 146(6), 22–27. <https://doi.org/10.5120/ijca2016910808> (ver p. 16)
- Kornilov, A. S., & Safonov, I. V. (2018). An Overview of Watershed Algorithm Implementations in Open Source Libraries. *Journal of Imaging*, 4(10). <https://doi.org/10.3390/jimaging4100123> (ver p. 15)
- Liu, D., & Xia, F. (2010). Assessing object-based classification: advantages and limitations. *Remote Sensing Letters*, 1, 187–194. <https://doi.org/10.1080/01431161003743173> (ver p. 14)
- Najar, C., Hebel, F., Speich, S., Canepa, S., & Agnolotti, V. (2010). What is GIS? http://www.gitta.info/what_gis/en/text/what_gis.pdf (ver p. 12)
- NASA. (2009). *First Photo of Earth From a Weather Satellite, TIROS-1*. https://www.nasa.gov/topics/earth/earthday/gall_tiros.html. (Ver p. 9)
- NASA. (2016). *The Television Infrared Observation Satellite Program (TIROS)*. <https://science.nasa.gov/missions/tiros>. (Ver p. 8)
- Neural Network. (2015). <http://www.extremetech.com/wp-content/uploads/2015/07/NeuralNetwork.png>. (Ver p. 22)

- Pereira-Pires, J. E., Aubard, V., Ribeiro, R. A., Pereira, J. M., Silva, J. M. N., & Mora, A. (2020). Semi-Automatic Methodology for Fire Break Maintenance Operations Detection with Sentinel-2 Imagery and Artificial Neural Network. *Remote Sensing*, 12(6), 909–927. <https://doi.org/https://doi.org/10.3390/rs12060909> (ver p. 21)
- Pereira-Pires, J. E., Aubard, V., Silva, J. M. N., Ribeiro, R. A., Pereira, J. M., Fonseca, J. M., Campagnolo, M. L., & Mora, A. (2020). Pixel-based and object-based change detection methods for assessing fuel break maintenance. *2020 International Young Engineers Forum (YEF-ECE)*, 49–54. <https://doi.org/10.1109/YEF-ECE49388.2020.9171818> (ver p. 14)
- Rohde, R. A., & Costa, J. (2020). *Espectro da radiação solar na Terra*. <https://commons.wikimedia.org/w/index.php?curid=94300841>. (Ver p. 7)
- Scheffler, D. (2017). *AROSICS: An Automated and Robust Open-Source Image Co-Registration Software for Multi-Sensor Satellite Data* (Versão 0.2.1) [This is the version as used in Scheffler et al. (2017): <https://www.mdpi.com/2072-4292/9/7/676>]. Zenodo. <https://doi.org/10.5281/zenodo.3743085>. (Ver p. 27)
- Stevens, J., Smith, J. M., & Bianchetti, R. A. (2012). *Mapping Our Changing World* (A. M. MacEachren & D. J. Peuquet, Eds.). Department of Geography. (Ver p. 8).
- Stutz, D., Hermans, A., & Leibe, B. (2018). Superpixels: An evaluation of the state-of-the-art. *Computer Vision and Image Understanding*, 166, 1–27. <https://doi.org/https://doi.org/10.1016/j.cviu.2017.03.007> (ver pp. 15–17)
- Tonbul, T. K.
bibinitperiod H. (2018). An experimental comparison of multi-resolution segmentation, SLIC and K-means clustering for object-based classification of VHR imagery. *International Journal of Remote Sensing*, 39(12), 1–17. <https://doi.org/http://dx.doi.org/10.1080/01431161.2018.1506592> (ver p. 21)

