



UNIVERSIDADE FEDERAL DE PERNAMBUCO
CENTRO DE INFORMÁTICA
PROGRAMA DE PÓS-GRADUAÇÃO EM CIÊNCIA DA COMPUTAÇÃO

PAULO MARTINS MONTEIRO

**Propostas de métodos baseados em Co-op training para aprendizado
semi-supervisionado em fluxos contínuos de dados**

Recife

2022

PAULO MARTINS MONTEIRO

**Propostas de métodos baseados em Co-op training para aprendizado
semi-supervisionado em fluxos contínuos de dados**

Trabalho apresentado ao Programa de Pós-graduação em Ciência da Computação do Centro de Informática da Universidade Federal de Pernambuco, como requisito parcial para obtenção do grau de Mestre em Ciência da Computação.

Área de Concentração: Inteligência Computacional

Orientador (a): Prof. Dr. Roberto Souto Maior de Barros

Recife

2022

Catálogo na fonte
Bibliotecária Monick Raquel Silvestre da S. Portes, CRB4-1217

M775p Monteiro, Paulo Martins
Propostas de métodos baseados em Co-op training para aprendizado semi-supervisionado em fluxos contínuos de dados / Paulo Martins Monteiro. – 2022.
84 f.: il., fig., tab.

Orientador: Roberto Souto Maior de Barros.
Dissertação (Mestrado) – Universidade Federal de Pernambuco. CIn, Ciência da Computação, Recife, 2022.
Inclui referências.

1. Inteligência computacional. 2. Fluxo contínuo de dados. I. Barros, Roberto Souto Maior de (orientador). II. Título.

006.31 CDD (23. ed.) UFPE - CCEN 2023-29

Paulo Martins Monteiro

“Propostas de métodos baseados em Co-op training para aprendizado semi-supervisionado em fluxos contínuos de dados”

Dissertação de Mestrado apresentada ao Programa de Pós-Graduação em Ciência da Computação da Universidade Federal de Pernambuco, como requisito parcial para a obtenção do título de Mestre em Ciência da Computação. Área de Concentração: Inteligência Computacional.

Aprovado em: 28/10/2022.

BANCA EXAMINADORA

Prof. Dr. Roberto Souto Maior de Barros
Centro de Informática / UFPE
(Orientador)

Prof. Dr. Paulo Maurício Gonçalves Júnior
Instituto Federal de Pernambuco, Campus Recife

Prof. Dr. Rodolfo Carneiro Cavalcante
Universidade Federal de Alagoas, Campus Arapiraca

Dedico esta dissertação a todos que apoiaram a minha trajetória até este momento.

AGRADECIMENTOS

Em primeiro lugar, gostaria de agradecer ao meu orientador Prof. Dr. Roberto Souto Maior de Barros, por me escolher como seu aluno e pacientemente me guiar durante este programa de pós-graduação.

A minha namorada Simone Gomes Zelaquett, que está sempre comigo me incentivando a evoluir.

Ao grupo de pesquisa de concept drift, especialmente a Silas Garrido Santos e José Luiz Pérez que sempre estiveram dispostos a tirar minhas duvidas.

Ao CIn-UFPE, pela oportunidade da realização do curso de pós-graduação de forma gratuita e por disponibilizar a sua infraestrutura.

O presente trabalho foi realizado com apoio da Coordenação de Aperfeiçoamento de Pessoal de Nível Superior – Brasil (CAPES) – Código de Financiamento 001.

RESUMO

No contexto de fluxo contínuo de dados, no qual os dados são gerados em tempo real, é comum a existência de dados sem rótulos, por exemplo, devido ao alto custo para rotulá-los. Para lidar com estes dados, estão sendo propostas estratégias de aprendizagem semi-supervisionada em que são utilizados dados rotulados e não rotulados ao mesmo tempo. Outro desafio típico dos fluxos contínuos de dados é a presença das chamadas mudanças de conceito (*concept drift*): neste cenário, a distribuição de probabilidade dos dados muda com o tempo, o que causa uma diminuição da precisão das classificações. Essa dissertação apresenta três novos métodos baseados na técnica de nossa autoria o *Co-op Training*, nos quais são utilizados dois classificadores que cooperam entre si para realizar predições em fluxo contínuos de dados. Estes algoritmos foram adaptados com o objetivo de obter uma melhor acurácia de classificação quando comparados ao método original e aos seus concorrentes. O primeiro método proposto é o *Co-op Training V2*, uma versão menos rigorosa do método original; o segundo é o *Co-op Training V3*, que utiliza apenas o nível de confiança de ambos os classificadores para rotular dados sem rótulo; e o último é o *Co-op Training V4*, que também utiliza apenas o nível de confiança na rotulação de dados, tendo o treinamento de ambos os classificadores como principal diferença para o V3. Os métodos propostos foram comparados aos algoritmos disponíveis no MOA-SS, a extensão do *Massive Online Analysis (MOA)* framework que foi utilizada para realizar os testes. Os experimentos utilizaram bases de dados artificiais e reais, tanto em conjuntos de dados sem mudanças de conceito quanto em cenários com mudanças de conceito. Finalmente, analisamos quais algoritmos se saíram melhor em cada um dos cenários testados utilizando como métrica a acurácia e o teste pos-hoc *Bonferroni-Dunn*, tendo o *Co-op Training* como a melhor opção para ser utilizado sem detector de mudança de conceito.

Palavras-chaves: fluxo contínuo de dados; classificação; aprendizagem semi-supervisionada; mudança de conceito; rotulação de dados.

ABSTRACT

In the context of continuous data flow, in which data is generated in real time, it is common to have unlabeled data, for example, due to the high cost of labeling it. To deal with such data, semi-supervised learning strategies are being proposed in which labeled and unlabeled data are used at the same time. Another typical challenge of continuous data stream is the presence of so-called concept drift: in this scenario, the probability distribution of the data changes over time, which causes a decrease in classification accuracy. This paper presents three new methods based on our technique, Co-op Training, in which two cooperating classifiers are used to make predictions on a continuous data stream. These algorithms have been adapted with the goal of obtaining better classification accuracy when compared to the original method and its competitors. The first proposed method is Co-op Training V2, a less rigorous version of the original method; the second is Co-op Training V3, which uses only the confidence threshold of both classifiers to label unlabeled data; and the last is Co-op Training V4, which also uses only the confidence threshold in labeling data, with the training of both classifiers as the main difference to V3. The proposed methods were compared to the algorithms available in MOA-SS, the extension of the Massive Online Analysis (MOA) framework that was used to perform the tests. The experiments used artificial and real databases, both in datasets without concept changes and in scenarios with concept changes. Finally, we analyze which algorithms did better in each of the tested scenarios using accuracy metrics and the Bonferroni-Dunn post hoc test, with Co-op Training as the best option to be used without a concept change detector.

Keywords: continuous data stream; classification; semi-supervised learning; concept drift; data labeling.

LISTA DE FIGURAS

Figura 1 – Fluxo de trabalho da aprendizagem supervisionada	23
Figura 2 – Demonstração do funcionamento do clustering	25
Figura 3 – a) Mudança abrupta, b) Mudança gradual e c) Mudança recorrente d) Mudança incremental	29
Figura 4 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases reais, sem detector e 3% dos dados rotulados.	43
Figura 5 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases reais, sem detector e 10% dos dados rotulados.	44
Figura 6 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases reais, sem detector e 30% dos dados rotulados.	44
Figura 7 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças abruptas, sem detector e 3% dos dados rotulados.	45
Figura 8 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças abruptas, sem detector e 10% dos dados rotulados.	46
Figura 9 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças abruptas, sem detector e 30% dos dados rotulados.	47
Figura 10 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças graduais, sem detector e 3% dos dados rotulados.	48
Figura 11 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças graduais, sem detector e 10% dos dados rotulados.	49
Figura 12 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças graduais, sem detector e 30% dos dados rotulados.	49
Figura 13 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases reais, sem detector e 3% dos dados rotulados.	50
Figura 14 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases reais, sem detector e 10% dos dados rotulados.	51
Figura 15 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases reais, sem detector e 30% dos dados rotulados.	52
Figura 16 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças abruptas, sem detector e 3% dos dados rotulados.	52

Figura 17 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças abruptas, sem detector e 10% dos dados rotulados.	53
Figura 18 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças abruptas, sem detector e 30% dos dados rotulados.	53
Figura 19 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças graduais, sem detector e 3% dos dados rotulados.	54
Figura 20 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças graduais, sem detector e 3% dos dados rotulados.	55
Figura 21 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças graduais, sem detector e 3% dos dados rotulados.	55
Figura 22 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases reais, com detector e 3% dos dados rotulados.	59
Figura 23 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases reais, com detector e 10% dos dados rotulados.	60
Figura 24 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases reais, com detector e 30% dos dados rotulados.	60
Figura 25 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças abruptas, com detector e 3% dos dados rotulados.	62
Figura 26 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças abruptas, com detector e 10% dos dados rotulados.	62
Figura 27 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças abruptas, com detector e 30% dos dados rotulados.	62
Figura 28 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças graduais, com detector e 3% dos dados rotulados.	62
Figura 29 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças graduais, com detector e 10% dos dados rotulados.	65
Figura 30 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças graduais, com detector e 30% dos dados rotulados.	65
Figura 31 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases reais, com detector e 3% dos dados rotulados.	65
Figura 32 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases reais, com detector e 10% dos dados rotulados.	67

Figura 33 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases reais, com detector e 30% dos dados rotulados.	67
Figura 34 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças abruptas, com detector e 3% dos dados rotulados.	67
Figura 35 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças abruptas, com detector e 10% dos dados rotulados.	67
Figura 36 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças abruptas, com detector e 30% dos dados rotulados.	67
Figura 37 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças graduais, com detector e 3% dos dados rotulados.	68
Figura 38 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças graduais, com detector e 10% dos dados rotulados.	68
Figura 39 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças graduais, com detector e 30% dos dados rotulados.	69

LISTA DE TABELAS

Tabela 1 – Característica dos conjuntos de dados utilizados.	37
Tabela 2 – Métodos propostos com bases reais, sem detector de mudança e 3% dos dados rotulados.	43
Tabela 3 – Métodos propostos com bases reais, sem detector de mudança e 10% dos dados rotulados.	44
Tabela 4 – Métodos propostos com bases reais, sem detector de mudança e 30% dos dados rotulados.	44
Tabela 5 – Métodos propostos com bases artificiais e mudanças abruptas, sem detector de mudança e 3% dos dados rotulados.	45
Tabela 6 – Métodos propostos com bases artificiais e mudanças abruptas, sem detector de mudança e 10% dos dados rotulados.	46
Tabela 7 – Métodos propostos com bases artificiais e mudanças abruptas, sem detector de mudança e 30% dos dados rotulados.	47
Tabela 8 – Métodos propostos com bases artificiais e mudanças graduais, sem detector de mudança e 3% dos dados rotulados.	48
Tabela 9 – Métodos propostos com bases artificiais e mudanças graduais, sem detector de mudança e 10% dos dados rotulados.	49
Tabela 10 – Métodos propostos com bases artificiais e mudanças graduais, sem detector de mudança e 30% dos dados rotulados.	50
Tabela 11 – Métodos concorrentes com bases reais, sem detector de mudança e 3% dos dados rotulados.	51
Tabela 12 – Métodos concorrentes com bases reais, sem detector de mudança e 10% dos dados rotulados.	51
Tabela 13 – Métodos concorrentes com bases reais, sem detector de mudança e 30% dos dados rotulados.	52
Tabela 14 – Métodos concorrentes com bases artificiais e mudanças abruptas, sem detector de mudança e 3% dos dados rotulados.	53
Tabela 15 – Métodos concorrentes com bases artificiais e mudanças abruptas, sem detector de mudança e 10% dos dados rotulados.	54

Tabela 16 – Métodos concorrentes com bases artificiais e mudanças abruptas, sem detector de mudança e 30% dos dados rotulados.	55
Tabela 17 – Métodos concorrentes com bases artificiais e mudanças graduais, sem detector de mudança e 3% dos dados rotulados.	56
Tabela 18 – Métodos concorrentes com bases artificiais e mudanças graduais, sem detector de mudança e 10% dos dados rotulados.	57
Tabela 19 – Métodos concorrentes com bases artificiais e mudanças graduais, sem detector de mudança e 30% dos dados rotulados.	58
Tabela 20 – Métodos propostos com bases reais, com detector de mudança e 3% dos dados rotulados.	59
Tabela 21 – Métodos propostos com bases reais, com detector de mudança e 10% dos dados rotulados.	60
Tabela 22 – Métodos propostos com bases reais, com detector de mudança e 30% dos dados rotulados.	60
Tabela 23 – Métodos propostos com bases artificiais e mudanças abruptas, com detector de mudança e 3% dos dados rotulados.	61
Tabela 24 – Métodos propostos com bases artificiais e mudanças abruptas, com detector de mudança e 10% dos dados rotulados.	61
Tabela 25 – Métodos propostos com bases artificiais e mudanças abruptas, com detector de mudança e 30% dos dados rotulados.	63
Tabela 26 – Métodos propostos com bases artificiais e mudanças graduais, com detector de mudança e 3% dos dados rotulados.	63
Tabela 27 – Métodos propostos com bases artificiais e mudanças graduais, com detector de mudança e 10% dos dados rotulados.	64
Tabela 28 – Métodos propostos com bases artificiais e mudanças graduais, com detector de mudança e 30% dos dados rotulados.	64
Tabela 29 – Métodos concorrentes com bases reais, com detector de mudança e 3% dos dados rotulados.	66
Tabela 30 – Métodos concorrentes com bases reais, com detector de mudança e 10% dos dados rotulados.	66
Tabela 31 – Métodos concorrentes com bases reais, com detector de mudança e 30% dos dados rotulados.	66

Tabela 32 – Métodos concorrentes com bases artificiais e mudanças abruptas, com detector de mudança e 3% dos dados rotulados.	68
Tabela 33 – Métodos concorrentes com bases artificiais e mudanças abruptas, com detector de mudança e 10% dos dados rotulados.	69
Tabela 34 – Métodos concorrentes com bases artificiais e mudanças abruptas, com detector de mudança e 30% dos dados rotulados.	70
Tabela 35 – Métodos concorrentes com bases artificiais e mudanças graduais, com detector de mudança e 3% dos dados rotulados.	71
Tabela 36 – Métodos concorrentes com bases artificiais e mudanças graduais, com detector de mudança e 10% dos dados rotulados.	72
Tabela 37 – Métodos concorrentes com bases artificiais e mudanças graduais, com detector de mudança e 30% dos dados rotulados.	73

LISTA DE ABREVIATURAS E SIGLAS

HDDM	Hoeffding Drift Detection Method
HT	Hoeffding Tree
MOA	Massive Online Analysis
MOA-SS	Massive Online Analysis-Semi Supervised
NB	Naive Bayes
RDDM	Reactive Drift Detection Method

SUMÁRIO

1	INTRODUÇÃO	17
1.1	OBJETIVOS	20
1.2	METODOLOGIA	20
1.3	ORGANIZAÇÃO DO TRABALHO	20
2	FUNDAMENTAÇÃO TEÓRICA	22
2.1	APRENDIZAGEM DE MÁQUINA	22
2.2	FLUXO DE DADOS	25
2.2.1	Principais classificadores	26
2.2.2	Mudança de conceito	27
2.2.2.1	<i>Definição</i>	27
2.2.2.2	<i>Tipos de mudança</i>	28
2.2.2.3	<i>Detectors de mudança de conceito</i>	29
2.3	FRAMEWORK	30
2.3.1	MOA	30
2.3.2	MOA-SS	31
2.4	MÉTODOS CONCORRENTES	31
3	MÉTODOS PROPOSTOS	32
3.1	CO-OP TRAINING V1	32
3.2	CO-OP TRAINING V2	33
3.3	CO-OP TRAINING V3	34
3.4	CO-OP TRAINING V4	35
4	EXPERIMENTOS E RESULTADOS	37
4.1	CONJUNTO DE DADOS	37
4.1.1	Bases de dados artificias	38
4.1.1.1	<i>SEA</i>	38
4.1.1.2	<i>AGRAWAL</i>	38
4.1.1.3	<i>RBF</i>	39
4.1.1.4	<i>LED</i>	39
4.1.1.5	<i>HEPATITIS</i>	39
4.1.2	Bases de dados reais	39

4.1.2.1	AIRLINES	40
4.1.2.2	COVERTYPE	40
4.1.2.3	ELECTRICITY	40
4.1.2.4	SENSOR	40
4.1.2.5	SPAM	41
4.1.2.6	CENSUS	41
4.1.2.7	POKERHAND	41
4.2	CONFIGURAÇÕES DOS EXPERIMENTOS	41
4.2.1	Métricas de avaliação	42
4.3	ANÁLISE DOS RESULTADOS SEM DETECTOR DE MUDANÇA DE CON- CEITO	43
4.3.1	Análise comparativa dos métodos propostos	43
4.3.2	Análise comparativa do melhor co-op vs métodos da ferramenta . . .	50
4.4	ANÁLISE DOS RESULTADOS COM DETECTOR DE MUDANÇA DE CONCEITO	59
4.4.1	Análise comparativa dos métodos propostos	59
4.4.2	Análise comparativa do melhor co-op vs métodos da ferramenta . . .	65
4.5	RESPOSTA A PERGUNTAS DE PESQUISA	74
5	CONCLUSÕES	76
5.1	PUBLICAÇÃO	78
5.2	TRABALHOS FUTUROS	78
	REFERÊNCIAS	79

1 INTRODUÇÃO

Atualmente, com o aumento do uso de smartphones, redes sociais, dispositivos de monitoramento de saúde, Internet das Coisas (IoT), a chegada de novas tecnologias como o 5G, e muitas outras aplicações, fazem com que tenhamos diversas fontes de informações de fluxo contínuos de dados em tempo real de forma constante ou infinita. Segundo o relatório do STATISTA (DEPARTMENT, 2022), está prevista a geração de 181 zettabytes de dados em 2025 e, devido ao alto volume de dados, torna-se humanamente impossível extrair manualmente as informações em tempo real. Para atender essa quantidade impressionante de dados, necessitamos de métodos que sejam rápidos e eficientes, que trabalhem em tempo real e isso requer uma quantidade razoável de recursos (BIFET et al., 2018).

Para resolver este desafio, são utilizadas técnicas de aprendizagem de máquina que buscam desenvolver sistemas que melhoram automaticamente o seu desempenho através de experiências (MITCHELL et al., 1990), tornando o algoritmo com a capacidade de fazer previsões, e.g. classificar radiografia para auxiliar no diagnóstico de pneumonia (FERREIRA; CAVALCANTE, 2019), traduzir linguagem de sinais (NEIVA, 2015) e detecção de fraudes em cartão de crédito (AWOYEMI; ADETUNMBI; OLUWADARE, 2017). As principais abordagens são a aprendizagem supervisionada, semi-supervisionada e não supervisionada.

Na aprendizagem supervisionada, é apresentado um conjunto de dados onde todos os exemplos contêm rótulo e consistem em uma ou mais entradas X e um valor de saída correspondente y , e visa construir um classificador ou regressor que consiga estimar o valor de saída para as entradas não vistas anteriormente.

Por outro lado, temos a aprendizagem não supervisionada, onde o conjunto de dados não contém rótulos, ou seja, são fornecidos apenas os valores de entrada. Esse tipo de algoritmo busca padrões a partir de similaridades. Neles são utilizadas técnicas de clustering, onde cada exemplo contido em um mesmo cluster são considerados similares, para isso são utilizados medidas de similaridade ou dissimilaridade, como, por exemplo, a distância euclidiana (WU; REHG, 2009), distância de Manhattan (MOHIBULLAH; HOSSAIN; HASAN, 2015) e distância de Mahalanobis (NASCIMENTO; TAVARES; SOUZA, 2015).

Por fim, temos a aprendizagem semi-supervisionada. Esta abordagem visa a combinação das duas aprendizagens apresentadas anteriormente, visto que são utilizados tanto dados rotulados quanto dados não rotulados. Os problemas da aprendizagem semi-supervisionada são

desafiadores e relevantes para aplicações de fluxo de dados, devido à escassez nos rótulos dos dados. Esta situação é comum ocorrer em dados reais devido à dificuldade de obtê-los e o processo de rotular os dados pode ser muito demorado e caro em processamento e memória, além de requerer um especialista humano. Atualmente diversas técnicas para aprendizagem semi-supervisionada são utilizadas, dentre elas o self-training (ALSAFARI; SADAQUI, 2021; TANHA; SOMEREN; AFSARMANESH, 2017), co-training (NING et al., 2021; ZHOU; LI et al., 2005), baseado em gráficos (CAMPS-VALLS; MARSHEVA; ZHOU, 2007; CHONG et al., 2020), co-regularization (SINDHWANI; NIYOGI; BELKIN, 2005; NIU et al., 2019), boosting (ZHANG et al., 2021; GRABNER; LEISTNER; BISCHOF, 2008), dentre outros.

Outro desafio encontrado na aprendizagem de fluxo contínuo de dados é a mudança de conceito. No mundo real, os conceitos, em geral, não são estáveis, causando mudanças com o tempo. Muitas vezes essas alterações tornam o modelo construído inconsistente ao receber novos dados, sendo necessária uma atualização regular do modelo (TSYMBAL, 2004). Diversos tipos de mudanças de conceito podem ocorrer, dentre eles a abrupta que ocorre quando há uma mudança de conceito de forma instantânea, a gradual tem uma transição mais lenta ao ser comparado ao abrupto, a mudança incremental ocorre quando há diversos conceitos intermediários entre dois conceitos, e enfim a mudança recorrente na qual o conceito antigo substituído reaparece depois de um período. Atualmente existem diversos detectores para lidar com as mudanças de conceito (Gonçalves Jr. et al., 2014). Neste trabalho utilizaremos dois detectores, Hoeffding Drift Detection Method (HDDM) e o Reactive Drift Detection Method (RDDM), considerados os melhores detectores para melhorar a precisão de classificação (BARROS; SANTOS, 2018).

Tendo em conta o que foi apresentado, neste trabalho demonstramos o método *Co-op Training* que foi desenvolvido para a dissertação e suas três variações. O *Co-op Training* foi inspirado na técnica de Co-Training sendo sua principal característica a cooperatividade de dois classificadores na previsão de rótulos. Para isto, ao chegar um exemplo não rotulado, o *Co-op Training* utiliza a combinação da previsão dos dois classificadores e o nível de confiança para treinar o *classificador base*. Na sua primeira variação denominada de *Co-op Training V2* criamos uma versão um pouco menos restrita, onde a construção é a mesma, tendo como diferencial a utilização da combinação da previsão dos dois classificadores ou o nível de confiança, a segunda variação intitulada de *Co-op Training V3* utiliza apenas o nível de confiança do *classificador base* ou o classificador secundário (*subclassificador*) com o intuito de treinar o *classificador base*, e por fim, temos o *Co-op Training V4* na qual assim como o *Co-op*

Training V3 utiliza o nível de confiança de ambos os classificadores, tendo como diferença a inversão de treinamento, ou seja, se o nível de confiança do *classificador base* atingir o limite definido pelo usuário, o classificador secundário será treinado utilizando o rótulo previsto pelo *classificador base*, e vice-versa.

Para a avaliação dos métodos propostos, realizamos os experimentos utilizando o ambiente de testes Massive Online Analysis (MOA) com a extensão Massive Online Analysis-Semi Supervised (MOA-SS), obtemos a média e o intervalo de confiança da acurácia após executar 30 repetições em cada cenário e utilizamos a média para medir o desempenho de todos os algoritmos. Além disso, foram realizadas diversas comparações entre os algoritmos propostos e os algoritmos contidos no MOA-SS. Estes experimentos foram realizados em diversos cenários, utilizando 3, 10 e 30% dos dados rotulados, 7 bases de dados reais e 5 artificiais (com variações de tamanho) contendo mudança de conceito, com e sem detectores de mudança de conceito. Utilizamos o HT e o NB como *classificador base* e classificador secundário e respectivamente o detector HDDM na sua versão A-test denominado HDDM_A e o RDDM.

Com o intuito de desenvolver um novo método para classificar fluxo de dados em ambiente semi-supervisionado, pretendemos investigar as seguintes questões de pesquisa:

- **(QP1):** Qual método dentre os propostos têm a melhor acurácia em ambiente com e sem detector?
- **(QP2):** Qual método dentre os concorrentes têm a melhor acurácia em ambiente com e sem detector?
- **(QP3):** Ocorreram diferenças significativas entre 3, 10 e 30% dos dados rotulados?
- **(QP4):** Os métodos propostos tiveram grande diferença ao lidar com mudança gradual e mudança abrupta?
- **(QP5):** As diferentes quantidades de instâncias nas bases artificiais interferiram no resultado?
- **(QP6):** Ocorreram diferenças ao utilizar detectores de mudança de conceito?

1.1 OBJETIVOS

O objetivo geral deste trabalho é propor algoritmos capazes de realizar a classificação em fluxos contínuos de dados no ambiente semi-supervisionado, podendo se comportar de forma positiva na presença das mudanças de conceitos tanto abrupta quanto gradual.

Como objetivos específicos, são propostos:

- Analisar abordagens propostas na literatura.
- Propor técnicas de classificação que lidem bem com fluxo de dados semi-supervisionado baseado na combinação de classificadores.
- Avaliar os métodos propostos em relação a outras abordagens.
- Analisar os resultados em diversos cenários.

1.2 METODOLOGIA

Para realizar estas experimentações foram desenvolvidos quatro algoritmos para lidar com fluxos contínuos de dados em ambiente semi-supervisionado baseado na técnica denominada Co-Training. Foram comparados os quatro algoritmos desenvolvidos para este trabalho e o que tiver o melhor desempenho será utilizado na comparação com outros 4 métodos disponíveis no MOA-SS. Todos os métodos foram comparados com 3, 10 e 30% dos dados rotulados, com e sem detector de mudança, e para simular os fluxos contínuos de dados foram utilizados bases reais e artificiais (com mudança abrupta e gradual), estas últimas contêm o tamanho de 100.000, 500.000, 750.000 e 1.000.000 instâncias. Foram utilizados os parâmetros padrões predefinidos na ferramenta e para avaliar executamos cada cenário 30 vezes para obter uma média da acurácia e o intervalo de confiança.

1.3 ORGANIZAÇÃO DO TRABALHO

Esta dissertação está dividida em cinco capítulos organizados da seguinte forma:

- O Capítulo 2 apresenta a fundamentação teórica necessária para compreender os métodos propostos;

- No Capítulo 3, iremos nos aprofundar nos métodos propostos, descrevendo toda a estrutura dos algoritmos e a explicação do seu funcionamento;
- No Capítulo 4, são descritos os experimentos realizados como parte desta dissertação, sendo apresentados os resultados obtidos em diversos cenários;
- Por fim, no Capítulo 5 apresentaremos a conclusão do estudo e os trabalhos futuros.

2 FUNDAMENTAÇÃO TEÓRICA

Nesta seção iremos abordar o conteúdo base deste trabalho, para compreender os métodos propostos. A seção está dividida em 3 partes: aprendizagem de máquina, fluxo de dados e MOA-SS Framework.

Na seção 2.1 iremos apresentar o aprendizado de máquinas e as suas três principais subáreas, seção 2.2 contém informações sobre fluxo de dados, principais classificadores utilizados e mudança de conceito e, para finalizar, na última seção abordaremos o framework utilizado para realizar a implementação e experimentos dos métodos propostos.

2.1 APRENDIZAGEM DE MÁQUINA

Esta seção aborda três principais subáreas de aprendizagem de máquina, onde na literatura encontramos diversas definições sobre o tema. Em (FACELI et al., 2011), a aprendizagem de máquina é definida como:

Em AM, computadores são programados para aprender com a experiência. Para tal, empregam um princípio de inferência denominada indução, no qual se obtêm conclusões genéricas a partir de um conjunto particular de exemplos. Assim, algoritmos de AM aprendem a induzir uma função ou hipótese capaz de resolver um problema a partir de dados que representam instâncias do problema a ser resolvido (FACELI et al., 2011).

Ou seja, a partir de uma amostra de dados é estimada uma função ou hipótese, para que este algoritmo torne-se capaz de realizar as tarefas desejadas.

Existem diversas aplicações onde são utilizadas técnicas de AM para solucionar problemas, dentre elas, a análise de sentimento (AGARWAL; MITTAL, 2016), detecção e reconhecimento de face (SINGHAL et al., 2022), combate à lavagem de dinheiro (CHEN et al., 2018), detecção de intrusão na rede (SINCLAIR; PIERCE; MATZNER, 1999), entre outros.

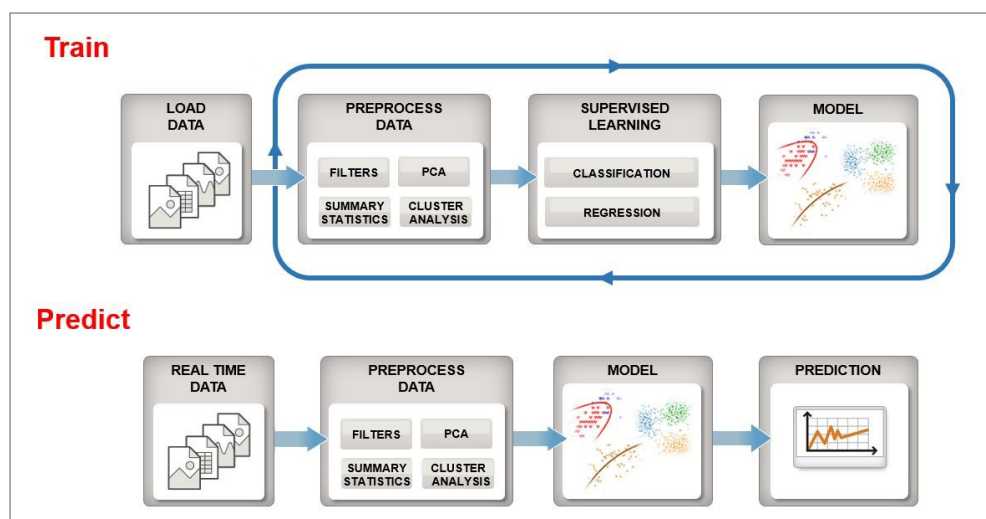
De acordo com (KANG; JAMESON, 2018), a aprendizagem de máquina pode ser dividida em quatro categorias: aprendizagem supervisionada, aprendizagem não supervisionada, aprendizagem semi-supervisionada e aprendizagem por reforço. Optamos por excluir desta seção a aprendizagem por reforço, dado que esta técnica não foi utilizada no desenvolvimento deste trabalho.

- **Aprendizagem Supervisionada:**

A aprendizagem supervisionada tem sido utilizada em diversas aplicações no mundo real

(KHATRI; ARORA; AGRAWAL, 2020)(SIVAKUMAR et al., 2018)(STAVENS; HOFFMANN; THRUN, 2007). Esta aprendizagem é semelhante ao aprendizado humano, porque devido às experiências passadas são obtidos novos conhecimentos com o intuito de melhorar a capacidade de realizar tarefas do mundo real (LIU, 2011). Essas experiências no mundo das máquinas são os dados coletados de alguma aplicação do mundo real, os dados consistem em um conjunto de exemplos que contém as variáveis de entrada X e uma variável de saída Y que são aprendidas e aplicadas um mapeamento para prever as saídas de dados nunca vistos anteriormente (CUNNINGHAM; CORD; DELANY, 2008). Na Fig. 1 podemos ver como funciona o processo de treinamento e predição. Na camada de treinamento os dados brutos são carregados e opcionalmente passa por uma etapa de pré-processamento, onde são aplicados diversas abordagens com objetivo de “limpar” os dados aumentando a precisão do modelo. Em seguida é aplicado o aprendizado supervisionado escolhendo qual técnica utilizar baseado no problema. Com isso, é gerado um modelo com capacidade de prever novos dados baseados nas experiências passadas. Ainda como pode ser visto na Fig.1, a etapa de previsão se inicia enviando para o modelo criado, dados para testes, ou seja, irá ser apresentado para o modelo, dados não rotulados. No meio deste processo deve ser aplicado as mesmas abordagens para pré-processamento de dados utilizado no momento do treinamento, após os dados chegarem no modelo, se inicia o processo de predição tendo como saída o valor previsto pelo modelo.

Figura 1 – Fluxo de trabalho da aprendizagem supervisionada



Fonte: (LOUSSAIEF; ABDELKRIM, 2016)

- **Aprendizagem Semi-Supervisionada:**

A aprendizagem semi-supervisionada está entre a aprendizagem supervisionada e não supervisionada, devido à possibilidade de ser treinado utilizando dados rotulados e não rotulados, estando apto para realizar tanto tarefa de classificação quanto de agrupamento (FILHO, 2009). O uso da abordagem semi-supervisionada tem chamado bastante atenção porque, na prática, para rotularmos os dados requer bastante esforço humano, tempo e dinheiro, sendo praticamente impossível em ambiente que contenha fluxo de dados. Ao utilizar esta abordagem reduzimos os custos, visto que não é necessário todos os dados estarem rotulados para realizar o treinamento (CALDAS, 2017).

Atualmente existem inúmeras aplicações para o aprendizado semi-supervisionado, como, por exemplo, análise de som pulmonar para diagnósticos de doenças pulmonares (CHAMBERLAIN et al., 2016), detecção de fraudes financeiras (WANG et al., 2019), classificar o tráfego de rede (ERMAN et al., 2007), e outros.

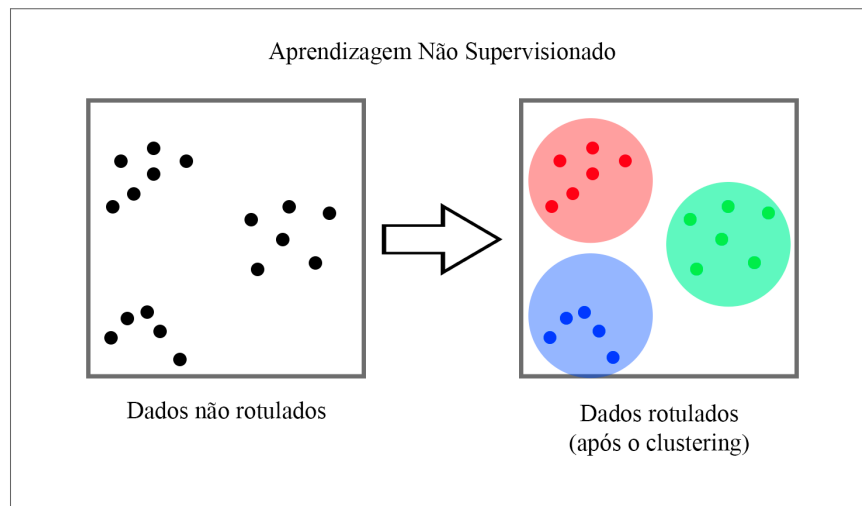
Atualmente encontram-se diversas abordagens de aprendizagem semi-supervisionada, dentre elas o Self-Training (YAROWSKY, 1995), na qual é um processo iterativo onde o classificador utiliza as suas próprias previsões para treinar a si mesmo. Ele armazena os dados em um lote e divide em duas janelas, L para os dados que contém rótulos e U para os dados que não contém rótulo. Em seguida é treinado com os dados da janela L e realiza as previsões dos rótulos da janela U, as previsões mais confiáveis são adicionados ao treinamento, repetindo todo o processo até a convergência. Outra abordagem bastante conhecida é a Co-Training (BLUM; MITCHELL, 1998) inicialmente utilizado para realizar classificação de páginas web, ele requer duas visões/hipóteses independentes, ou seja, pode haver bases de dados diferentes para um mesmo problema como, por exemplo: informações textuais e gráficas sobre um mesmo objeto. É necessário que os dados atendam a duas premissas para garantir um bom desempenho: premissa de independência (as visões devem ser condicionalmente independentes) e a premissa de suficiência (as visões devem ser suficientes para construir um bom classificador). Os dois algoritmos de aprendizado são treinados separadamente em cada visão e as previsões de cada classificador são usadas para ampliar o conjunto de treinamento do outro.

- **Aprendizagem Não Supervisionada:**

O Aprendizado não supervisionado, ao contrário do supervisionado, é utilizado quando não há rótulos disponíveis no conjunto de dados, ou seja, não se sabe quantas ou quais

classes existem. Uma das maneiras mais comuns para a realização das predições é a utilização da técnica de (*clustering*). O *clustering* é caracterizado pela maior similaridade dentro do mesmo *cluster* (grupo) e pela maior dissimilaridade entre diferentes *clusters* (SINAGA; YANG, 2020), para isso se utilizam medidas de distâncias, por exemplo: distância Euclidiana, distância de Mahalanobis, distância de Manhattan, ou seja, as instâncias que contêm características similares são agrupadas em um único cluster, enquanto as de características diferentes são agrupadas em clusters diferentes, como mostrado na Fig.2. Diversos algoritmos utilizam esta técnica, dentre elas o K-means (KAMPER; LIVESCU; GOLDWATER, 2017; BASAR et al., 2020), K-nearest neighbor (VAJDA; SANTOSH, 2016; KRAMER, 2011), Mean-shift Clustering (VIRUPAKSHAPPA; ORUKLU, 2019), etc.

Figura 2 – Demonstração do funcionamento do clustering



Fonte: O autor (2022)

2.2 FLUXO DE DADOS

Segundo (SANTOS; BARROS; Gonçalves Jr., 2015) “Os fluxos de dados são ambientes onde os dados fluem contínua e rapidamente, e em abundância.” Hoje em dia diversas fontes produzem dados continuamente como, por exemplo, e-mails, sensores, ligações telefônicas, etc. Essas fontes contêm uma geração contínua de uma indeterminada quantidade de dados, trazendo novos desafios devido ao armazenamento, manutenção e consulta do fluxo de dados (GAMA, 2010).

Os fluxos de dados são sequências de informações que chegam uma a uma em tempo real com tamanho infinito, tendo uma ordem temporal. Atualmente há alguns desafios ao encarar

um fluxo de dados, pois, ele é grande e rápido, ou seja, precisamos extrair as informações em tempo real. Com isso devemos obter uma solução de modo que utilize menos tempo e memória. Outro problema é que os dados podem evoluir, sendo necessário nossos modelos se adaptarem quando há alterações (BIFET et al., 2018).

As principais características do fluxo de dados são:

- A quantidade de dados é potencialmente infinita, o que impossibilita armazenar, sendo possível apenas armazenar um pequeno resumo dos dados.
- Devido à velocidade que os dados chegam, eles devem ser processados em tempo real e logo após descartados.
- Com o passar do tempo, a distribuição dos dados pode sofrer alterações, causando a mudança de conceito e tornando os dados irrelevantes ou prejudicar a qualidade/precisão do modelo.

2.2.1 Principais classificadores

Nesta seção discutimos os classificadores que selecionamos para utilizar como base nos métodos proposto e comparados, inicialmente esses classificadores foram propostos para tarefas de classificação estacionárias, mas também podem ser utilizados para classificar os fluxos de dados.

- **Hoeffding Tree:**

O Hoeffding Tree foi proposto em (DOMINGOS; HULTEN, 2000) e constitui um algoritmo baseado em árvore de decisão incremental, o que lhe dá a capacidade de aprender fluxo de dados, entretanto, com o pressuposto de que o contexto não muda com o tempo, ou seja, o modelo criado quando na presença de mudança de conceito ocorre uma redução de precisão.

O algoritmo utiliza o conceito de fronteira de Hoeffding (HOEFFDING, 1994), na qual quantifica um limiar de amostras fundamentais para realizar uma predição precisa (ABREU; CARVALHO; ABELÉM, 2021).

Uma vantagem em utilizar o Hoeffding Tree na aprendizagem semi-supervisionada é que o algoritmo explora o fato que uma pequena amostra pode ser suficiente para obter um atributo ideal para os nós das árvores de decisão (BIFET, 2009).

▪ Naive Bayes:

O classificador Naive Bayes (RISH, 2001), bastante utilizado por sua simplicidade e sua rapidez, é um classificador probabilístico que aplica o teorema de Bayes para realizar as predições, este classificador é considerado ingênuo porque assume que os atributos são independentes, ou seja, mesmo que as características dos dados dependam uma da outra ou da existência de outras características, o Naive Bayes considera que todas essas características sejam independentes.

Decidimos utilizar o Naive Bayes por ser um dos métodos de aprendizagem de máquina mais popular. Sua simplicidade combinada com a sua eficiência torna um método bastante atrativo (VIEGAS, 2015).

2.2.2 Mudança de conceito

No fluxo de dados em ambiente não estacionário um dos problemas mais desafiantes ocorre quando os dados sofrem mudanças ao longo do tempo, como, por exemplo, (DELANY et al., 2004) onde um classificador de spam verifica se os e-mails recebidos são spans ou não. Ao passar do tempo, os spammers irão modificar a sua maneira de construir as mensagens com o proposito de burlar o classificador, fazendo com que acredite que os e-mails não são spans. Uma das possíveis soluções para este problema seria a implementação de um detector de mudança de conceito.

Em 2.2.2.1 iremos definir o que seria a mudança de conceito, na seção 2.2.2.2 iremos apresentar os principais tipos de mudança de conceito e para finalizar em 2.2.2.3 compreendemos o que são os detectores e quais iremos utilizar nos métodos propostos.

2.2.2.1 Definição

A mudança de conceito ocorre quando há uma mudança na probabilidade posterior $P(c|X)$, onde a classe muda mantendo os valores de atributos da anterior. Devido a isto, o desempenho do classificador é diminuído.(AGRAHARI; SINGH, 2021)

Pode-se dizer que o conceito na qual os dados estão sendo coletados pode mudar de tempo a tempo. Estas mudanças que ocorrem ao longo do tempo se refletem de alguma forma nos exemplos utilizados para treinamento do modelo, fazendo com que as observações passadas

tornem-se irrelevantes para o estado atual. (GAMA, 2010)

Os algoritmos tradicionais de aprendizagem de máquina pressupõem que a distribuição dos dados é estática, mas a distribuição dos dados pode mudar com o tempo, tornando os algoritmos tradicionais de aprendizagem por lote, inadequados para modelos treinados com o fluxo de dados (WARES; ISAACS; ELYAN, 2019).

As mudanças de conceito podem ocorrer de diversas formas, na próxima subseção iremos detalhar cada uma.

2.2.2.2 Tipos de mudança

A mudança de conceito pode ser abrupta, gradual, incremental e recorrente. Na Fig.3 é demonstrado os diferentes tipos de mudança e sua respectiva velocidade. Em seguida detalharemos cada um deles:

- **Abrupto:**

Neste tipo, o antigo conceito do fluxo de dados é substituído de maneira repentina por um novo conceito (ver Fig. 3 (a)). Esta mudança brusca é conhecida por mudança abrupta. Devido a esta mudança, os algoritmos de aprendizagem tendem a ter sua precisão abruptamente degradada. Logo após o ocorrido, gradualmente o modelo tende a ter sua precisão melhorada conforme a chegada e treinamento dos novos dados.

- **Gradual:**

Ao contrário da mudança abrupta, com sua mudança de forma momentânea, a mudança gradual tem uma transição relativamente longa se comparado ao abrupto (ver Fig. 3 (b)). O novo conceito se sobrepõe ao poucos por um período até que este conceito se torne estável.

- **Incremental:**

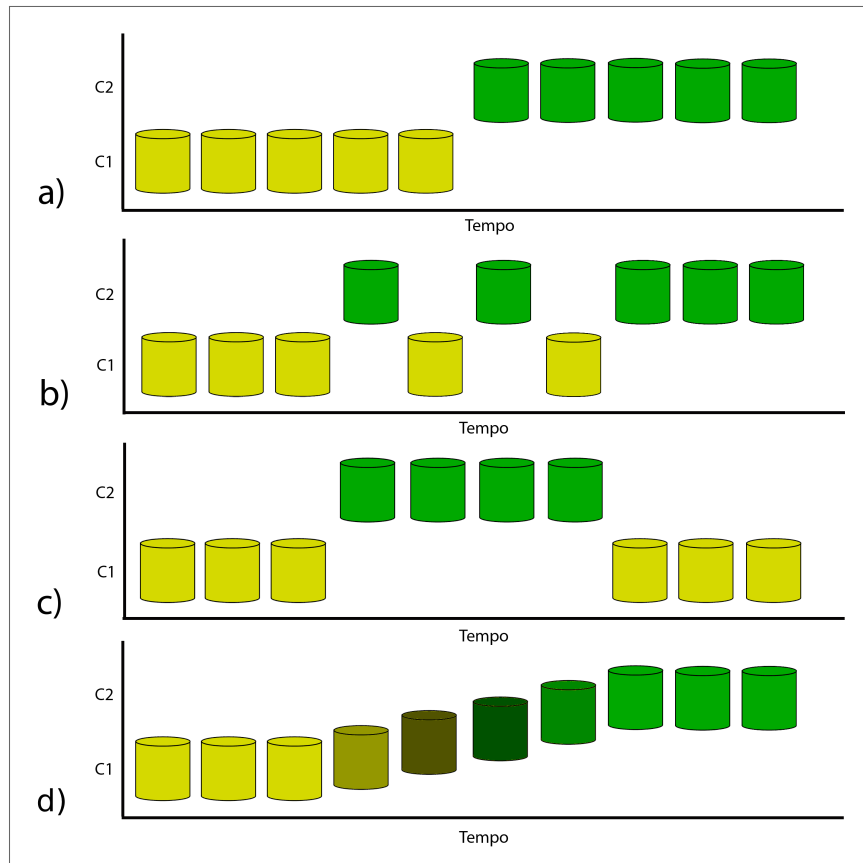
Neste tipo de mudança contém a presença de diversos conceitos intermediários entre os dois conceitos observados, indo de maneira gradativa do conceito C1 para o conceito C2. (ver Fig. 3 (d)).

- **Recorrente:**

Neste tipo de mudança, o conceito antigo que foi substituído pelo novo reaparece depois de um período de tempo (ver Fig. 3 (c)). A maioria dos detectores de mudança

são construídos para detectar apenas gradual e abrupto, pois, para detectar uma mudança recorrente existem alguns limites, como, por exemplo, memória para armazenar as informações antigas que podem ou não ser utilizadas no futuro.

Figura 3 – a) Mudança abrupta, b) Mudança gradual e c) Mudança recorrente d) Mudança incremental



Fonte: O autor (2022)

2.2.2.3 Detectores de mudança de conceito

Os detectores de mudanças de conceitos têm o papel de alertar ao classificador as alterações, no momento que o mesmo detecta que ocorreu uma mudança na distribuição dos dados (HIDALGO, 2017). São utilizadas diversas técnicas para detectar as mudanças, dentre eles o EDDM (BAENA-GARCIA et al., 2006) que utiliza o número de exemplo entre dois erros de classificação, ADWIN (BIFET; GAVALDA, 2007) utiliza a técnica de janela deslizante adaptável, O STED (NISHIDA; YAMAUCHI, 2007) compara a precisão do classificador em uma janela recente com o histórico, entre outros. Em seguida detalharemos os detectores HDDM e RDDM, estes dois detectores foram escolhidos para serem utilizados nestes experimentos porque de

acordo com (ANDERSON et al., 2019) o RDDM e o HDDM-A foram considerados os melhores detectores para melhorar a precisão de classificação.

- **HDDM:**

No HDDM (FRIAS-BLANCO et al., 2014) foram propostas duas abordagens, a primeira utiliza médias móveis, mais adequada para detectar mudanças abruptas (HDDM_A) e a segunda utiliza médias móveis ponderadas para lidar com mudanças graduais (HDDM_W). Os métodos monitoram ao longo do tempo a medida de desempenho (acertos e erros do classificador) para detectar mudanças significativas. Segundo os autores, estes métodos são de passos únicos, oferecendo limite sobre a taxa de falsos positivos e falsos negativos. Embora o autor tenha estudado o problema de mudança de conceitos assumindo que os dados são rotulados, ele afirma que estes métodos podem ser aplicados nos cenários mais realistas na qual os rótulos são mais difíceis de se obter.

- **RDDM:**

O RDDM (BARROS et al., 2017) tem como proposta original superar as deficiências e melhorar as detecções e a acurácia do DDM, tendo sido implementado no MOA, a sua principal ideia é reduzir o número de instâncias em conceitos muito longos e estáveis com o intuito de resolver um problema que causa uma perda de desempenho no DDM. É assumido que este problema é causado devido à baixa sensibilidade a mudança de conceitos em um número elevado de instâncias, porque é necessário um número bastante grande de erros de previsão para afetar suficientemente a taxa média de erro e detectar a mudança. Outras medidas também são tomadas para um melhor resultado.

2.3 FRAMEWORK

2.3.1 MOA

Apresentado em (BIFET et al., 2010), o MOA é um dos frameworks mais populares de código aberto para fluxo de dados, cujo objetivo é desenvolver uma estrutura experimental para realizar classificações, agrupamento, detectores de mudança de conceito, detector de outlier, entre outros. Tendo sua estrutura semelhante ao do WEKA (KHANALE; PATHAK, 2011). Este framework traz o benefício de comparar diversos cenários do mundo real ou artificial, tendo uma vasta gama de algoritmos, medidas de avaliação e geradores de dados.

2.3.2 MOA-SS

O MOA-SS (NGUYEN; GOMES; BIFET, 2019) foi desenvolvido para ser uma aba integrada dentro da API do MOA. Ele nos permite realizar experimentos semi-supervisionado em fluxo de dados dentro do ambiente MOA, além de permitir mascarar uma porcentagem dos dados, tanto do conjunto gerado automaticamente quanto definido pelo usuário, tornando dado supervisionado em semi-supervisionado ou não supervisionado.

2.4 MÉTODOS CONCORRENTES

- **Self-Training:**

Com o intuito de diminuir o ruído, este modelo utiliza lotes contendo uma certa quantidade de exemplos. Ao receber uma nova instância, ela é adicionada dentro do lote sendo separado em listas: U contém exemplos não rotulados e L exemplos rotulados que além de separados são utilizados para treinar. Ao completar o lote são classificados todas as instâncias em U, as instâncias que obtiverem o valor de confiança¹ acima do limite definido, serão utilizadas para treinar o modelo com o rótulo previsto.

- **Self-Training Incremental:**

Diferente do modelo anterior, este modelo é atualizado ao chegar uma nova instância (tem o princípio um por um). Ao receber a instância X, será verificado se ela tem rótulo, se sim este exemplo é utilizado para treinar, caso contrário a modelo irá prever o rótulo e se a confiança estiver acima do limite definido, o modelo assume que está correto e utiliza a instância com o rótulo previsto para treinar. Caso não atinja a confiança necessária, o exemplo é descartado.

- **Self-Training Weighting:**

Esta variação de self-training utiliza todas as instâncias para treinar o modelo, mas define um peso para cada exemplo conforme o nível de confiança da previsão.

¹ O valor de confiança é determinado pelo preditor com base na probabilidade de que um determinado rótulo previsto esteja correto

3 MÉTODOS PROPOSTOS

Este trabalho demonstra o método Co-op Training (MONTEIRO; SOARES; BARROS, 2021) desenvolvido durante o mestrado e suas três variações. Esta técnica foi inspirada no algoritmo de Co-Training (BLUM; MITCHELL, 1998), sendo uma principal característica a cooperatividade de dois classificadores (*classificador base* e *subclassificador*) na previsão de rótulos. Todos os métodos propostos funcionam de forma incremental, ou seja, têm o princípio de um por um. Ao chegar dados rotulados, todos os métodos treinam tanto o *classificador base* quanto o *subclassificador*, mas a diferença está na forma de tratar dados não rotulados. No Co-op Training, são necessários ambos classificadores terem o resultado igual e o nível de confiança do *classificador base* ser maior ou igual ao limite definido pelo usuário. No Co-op Training V2 a escolha é um pouco menos restrita, pois, o *classificador base* é treinado apenas quando ambos classificadores tenham o mesmo resultado ou o nível de confiança do *classificador base* é maior, ou igual ao limite. No Co-op Training V3 decidimos utilizar apenas o nível de confiança de ambos classificadores, que caso seja maior que o limite o *classificador base* é treinado. Para finalizar, o Co-op Training V4 assim como no Co-op Training V3 também utiliza apenas o nível de confiança, tendo como diferença a inversão de treinamento, ou seja, se o nível de confiança do *classificador base* é igual ou maior que o limite, o *subclassificador* será treinado utilizando o rótulo previsto pelo *classificador base*, e vice-versa. Todos os métodos foram implementados dentro do MOA-SS. Utilizamos analogias do mercado financeiro para diferenciar o comportamento de cada versão (Conservador, Moderado, Arrojado e Agressivo).

3.1 CO-OP TRAINING V1

O Co-op Training V1 ou apenas Co-op Training é uma versão conservadora do algoritmo porque é necessário que ambos classificadores devem ter a mesma previsão além do nível de confiança do *classificador base* ser maior que o limiar definido pelo usuário, ou seja, menos instâncias serão rotuladas e treinadas. No algoritmo 1 é possível entender o seu funcionamento. Enquanto o fluxo de dados estiver ativo a variável X irá receber a instância atual, nas linhas 3 – 5, o método irá verificar se este exemplo contém o rótulo, caso contenha, ele irá utilizar esta instância para treinar tanto o *classificador base* quanto o *subclassificador*, caso este exemplo não contenha rótulo ele irá para a etapa de predição. Caso o algoritmo tenha obtido mais

de 100 instâncias, o mesmo irá seguir para a predição do rótulo ausente. Na linha 8 ele irá armazenar o nível de confiança que o classificador base tem de cada classe. Já na linha 9 irá armazenar o rótulo determinado pelo classificador base e na linha 10 irá armazenar o valor do nível de confiança do classificador e o mesmo processo ocorre nas linhas 11 – 13, tendo como diferença a utilização do subclassificador. Nas linhas 14 – 17, o método irá decidir se vai definir o rótulo deste algoritmo. Inicia-se na linha 14 com a verificação se a predição do classificador base é a mesma do subclassificador, caso seja a mesma e o nível de confiança do classificador seja maior que o limiar de confiança definido pelo usuário, o algoritmo irá rotular a instância com a classe prevista por ambos e treinar apenas o classificador base, a instância que não atender a expectativa da linha 14, não será rotulada.

Algoritmo 1: Algoritmo Co-op Training V1

input: fluxo de dados fd , classificadorBase F , subClassificador S , confiança T

```

1 while  $fd$  is active do
2    $X \leftarrow \text{proximalInstância}()$ ;
3   if  $y$  not missed then
4      $F.\text{trainOnInstance}(X,y)$ ;
5      $S.\text{trainOnInstance}(X,y)$ ;
6   else
7     if  $\text{contador} \geq 100$  then
8        $\text{predsBase} \leftarrow \text{predNormalizado}(F, X)$ ;
9        $\text{predBase} \leftarrow \text{predsBase.maxIndex}()$ ;
10       $\text{confBase} \leftarrow \text{predsBase.getValue}(\text{predBase})$ ;
11       $\text{predsSub} \leftarrow \text{predNormalizado}(S, X)$ ;
12       $\text{predSub} \leftarrow \text{predsSub.maxIndex}()$ ;
13       $\text{confSub} \leftarrow \text{predsSub.getValue}(\text{predSub})$ ;
14      if  $\text{predBase} == \text{predSub}$  and  $\text{confBase} \geq T$  then
15         $\text{instCopy} \leftarrow X.\text{copy}()$ ;
16         $\text{instCopy.setClass}(\text{predBase})$ ;
17         $F.\text{trainOnInstance}(\text{instCopy})$ ;
18      end
19    end
20  end
21   $\text{contador}++$ 
22 end

```

3.2 CO-OP TRAINING V2

O Co-op Training V2 consideramos ser uma versão moderada, pois é pouco restrito se comparado a V1, tendo a sua diferença na linha 14 do algoritmo (realçado) 2 onde, caso

a *predição do classificador base* for igual à *predição do subclassificador* ou, caso o nível de confiança do *classificador base* for maior ou igual ao limiar, a instância irá receber o rótulo previsto e o *classificador base* será treinado.

Algoritmo 2: Algoritmo Co-op Training V2

input: fluxo de dados fd , classificadorBase F , subClassificador S , confiança T

```

1 while  $fd$  is active do
2    $X \leftarrow \text{proximalInstância}()$ ;
3   if  $y$  not missed then
4      $F.\text{trainOnInstance}(X,y)$ ;
5      $S.\text{trainOnInstance}(X,y)$ ;
6   else
7     if  $\text{contador} \geq 100$  then
8        $\text{predsBase} \leftarrow \text{predNormalizado}(F, X)$ ;
9        $\text{predBase} \leftarrow \text{predsBase.maxIndex}()$ ;
10       $\text{confBase} \leftarrow \text{predsBase.getValue}(\text{predBase})$ ;
11       $\text{predsSub} \leftarrow \text{predNormalizado}(S, X)$ ;
12       $\text{predSub} \leftarrow \text{predsSub.maxIndex}()$ ;
13       $\text{confSub} \leftarrow \text{predsSub.getValue}(\text{predSub})$ ;
14      if  $\text{predBase} == \text{predSub}$  or  $\text{confBase} \geq T$  then
15         $\text{instCopy} \leftarrow X.\text{copy}()$ ;
16         $\text{instCopy.setClass}(\text{predBase})$ ;
17         $F.\text{trainOnInstance}(\text{instCopy})$ ;
18      end
19    end
20  end
21   $\text{contador}++$ 
22 end
  
```

3.3 CO-OP TRAINING V3

No Co-op training V3 consideramos ser uma versão arrojada, pelo fato de ter menos restrições na hora de escolher quais instâncias sem rótulo serão utilizadas. Assim como o Co-op training V2, tem a sua diferença na linha 14 (realçada), mas diferente das duas versões anteriores, esta versão utiliza apenas o nível de confiança para determinar se esta instância com o rótulo previsto deverá ser utilizada para treinar o *classificador base*, caso o nível de confiança do *classificador base* ou o nível de confiança do subclassificador *seja maior igual ao limiar definido pelo usuário*, o *classificador base* será treinado com a instância que teve o seu rótulo previsto.

Algoritmo 3: Algoritmo Co-op Training V3

input: fluxo de dados fd , classificadorBase F , subClassificador S , confiança T

```

1 while  $fd$  is active do
2    $X \leftarrow \text{proximalInstância}()$ ;
3   if  $y$  not missed then
4      $F.\text{trainOnInstance}(X,y)$ ;
5      $S.\text{trainOnInstance}(X,y)$ ;
6   else
7     if  $\text{contador} \geq 100$  then
8        $\text{predsBase} \leftarrow \text{predNormalizado}(F, X)$ ;
9        $\text{predBase} \leftarrow \text{predsBase.maxIndex}()$ ;
10       $\text{confBase} \leftarrow \text{predsBase.getValue}(\text{predBase})$ ;
11       $\text{predsSub} \leftarrow \text{predNormalizado}(S, X)$ ;
12       $\text{predSub} \leftarrow \text{predsSub.maxIndex}()$ ;
13       $\text{confSub} \leftarrow \text{predsSub.getValue}(\text{predSub})$ ;
14      if  $\text{confBase} \geq T$  or  $\text{confSub} \geq T$  then
15         $\text{instCopy} \leftarrow X.\text{copy}()$ ;
16         $\text{instCopy.setClass}(\text{predBase})$ ;
17         $F.\text{trainOnInstance}(\text{instCopy})$ ;
18      end
19    end
20  end
21   $\text{contador}++$ 
22 end

```

3.4 CO-OP TRAINING V4

Por fim, o Co-op Training V4 seria uma versão agressiva, nele são treinados os dois classificadores utilizando apenas o nível de confiança do classificador oposto, ou seja, caso o classificador oposto tenha classificado errado com um alto nível de confiança, o classificador será treinado com rótulo incorreto. Como pode ser visto no algoritmo 4, o Co-op Training V4 tem a sua estrutura similar as outras versões do Co-op training, se distinguindo na forma de treinar os dois classificadores, onde ele apenas verifica se o nível de confiança do classificador base é maior ou igual ao limiar de confiança definido pelo usuário, se sim, a instância irá ser rotulada com a predição e irá ser utilizada para treinar o subclassificador. O algoritmo teve sua linha 14 (realçada) modificada ao comparar com o Co-op Training V3 sendo adicionados a linha 19 – 23 (realçada) onde irá ocorrer o inverso, o algoritmo verificará se o nível de confiança do subclassificador é maior que o nível de confiança definido pelo usuário, caso seja maior ou igual à instância será rotulada e treinada pelo classificador base.

Algoritmo 4: Algoritmo Co-op Training V4

input: data stream ds , baseLearner F , subLearner S , thresholdBase T , thresholdSub R

```

1 while  $ds$  is active do
2    $X \leftarrow \text{nextInstance}()$ ;
3   if  $y$  not missed then
4      $F.\text{trainOnInstance}(X,y)$ ;
5      $S.\text{trainOnInstance}(X,y)$ ;
6   else
7     if  $\text{contador} \geq 100$  then
8        $\text{predsBase} \leftarrow \text{predNormalizado}(F, X)$ ;
9        $\text{predBase} \leftarrow \text{predsBase.maxIndex}()$ ;
10       $\text{confBase} \leftarrow \text{predsBase.getValue}(\text{predBase})$ ;
11       $\text{predsSub} \leftarrow \text{predNormalizado}(S, X)$ ;
12       $\text{predSub} \leftarrow \text{predsSub.maxIndex}()$ ;
13       $\text{confSub} \leftarrow \text{predsSub.getValue}(\text{predSub})$ ;
14      if  $\text{confBase} \geq T$  then
15         $\text{instCopy} \leftarrow X.\text{copy}()$ ;
16         $\text{instCopy.setClass}(\text{predBase})$ ;
17         $S.\text{trainOnInstance}(\text{instCopy})$ ;
18      end
19      if  $\text{confSub} \geq R$  then
20         $\text{instCopy} \leftarrow X.\text{copy}()$ ;
21         $\text{instCopy.setClass}(\text{predSub})$ ;
22         $F.\text{trainOnInstance}(\text{instCopy})$ ;
23      end
24    end
25  end
26   $\text{contador}++$ 
27 end
  
```

4 EXPERIMENTOS E RESULTADOS

Neste capítulo iremos descrever os experimentos realizados como parte desta dissertação e apresentar os resultados obtidos em diversos cenários: bases reais, artificiais, mudança de conceito abruptas e graduais, com e sem detectores. A seguir apresentaremos os bancos de dados utilizados nos experimentos.

4.1 CONJUNTO DE DADOS

Para validar os métodos desenvolvidos foram utilizadas 12 bases de dados (40 considerando as variações que podem conter mudanças graduais, abruptas e diferentes quantidades de instâncias), nas quais 7 são bases reais e 5 bases de dados artificiais, buscaram-se bases de dados diversos, desde relação à quantidade de classes quanto em quantidade de atributos, como pode ser visto na tabela 1. Nas seções abaixo detalharemos de forma resumida as bases utilizadas neste experimento.

Tabela 1 – Característica dos conjuntos de dados utilizados.

Dataset	Tipo	Exemplos	Atributos	Numérico	Categorico	Classes
Airlines	Real	539.383	7	3	4	2
Census	Real	299.284	41	13	28	2
Coverttype	Real	581.012	54	10	44	7
Electricity	Real	45.234	8	7	1	2
PokerHand	Real	1.025.009	10	10	0	10
Sensor	Real	2.219.803	5	5	0	58
Spam	Real	9.324	499	0	499	2
Agrawal	Artificial	100K, 500K, 750K e 1M	9	6	3	2
Hepatitis	Artificial	1M	19	6	13	2
LED	Artificial	100K, 500K, 750K e 1M	24	0	24	10
Random RBF	Artificial	100K, 500K, 750K e 1M	10	10	0	2
SEA	Artificial	100K, 500K, 750K e 1M	3	3	0	2

Fonte: O autor (2022)

4.1.1 Bases de dados artificias

Nesta seção são apresentados as bases de dados artificias geradas pelo MOA (BIFET et al., 2010) e utilizadas neste experimento.

4.1.1.1 SEA

Introduzido em (STREET; KIM, 2001), este conjunto de dados simula um fluxo de dados contendo mudança de conceito abrupta. O conjunto contém três atributos na qual apenas dois são relevantes para a classificação, é escolhido uma das funções possíveis:

- *Função 0: se $1(att1 + att2 \leq 8)$*
- *Função 1: se $1(att1 + att2 \leq 9)$*
- *Função 2: se $1(att1 + att2 \leq 7)$*
- *Função 3: se $1(att1 + att2 \leq 9.5)$*

Essas funções de classificação comparam a soma dos dois atributos relevantes com um valor limiar único para cada função de classificação, dependendo do resultado o gerador irá classificar com uma dos dois rótulos possíveis. Para simular a mudança de conceito a função de classificação é alterada.

4.1.1.2 AGRAWAL

Este conjunto de dados é baseado no artigo (AGRAWAL; IMIELINSKI; SWAMI, 1993), cujo objetivo é classificar se as pessoas que solicitaram o empréstimo, estão aptas a recebê-lo. Para isto ele verificará nove atributos nos quais três são categóricos (nível de educação, código postal e a marca do carro), e todos os outros são numéricos (salário, comissão, idade, valor da casa, quantidade de anos de casa própria, e o valor total do empréstimo solicitado). Todos os valores dos atributos são gerados de forma aleatória, sendo separados por uma classe binária: grupo A e grupo B. Existem dez funções para realizar a classificação a partir destes atributos, estas funções podem ser encontradas no artigo original (AGRAWAL; IMIELINSKI; SWAMI, 1993).

4.1.1.3 RBF

No *RandomRBF* (*Radial Basis Function*) (BIFET et al., 2010) é inicialmente gerado um conjunto aleatório de centros para cada classe. Para cada centro é definido de maneira aleatória um peso, um ponto central por atributo e um desvio padrão. Na criação de novas instâncias, um centro é escolhido de forma aleatória, considerando os pesos onde os centros com maiores pesos são mais propensos a serem escolhidos. Os valores de cada atributo são gerados de modo aleatório, e o vetor geral seja dimensionado para que seu comprimento, seja igual a um valor aleatório da distribuição gaussiana, para simular as mudanças de conceito são alteradas as posições dos centroides.

4.1.1.4 LED

Originalmente em (OLSHEN; STONE, 1984) o gerador produzia os dados para um display que contém 7 LEDs, envolvendo 7 atributos booleanos, mas o conjunto de dados utilizado além dos 7 atributos originais foram adicionados outros 17 atributos booleanos irrelevantes visando dificultar a predição, além de por padrão ter 10% dos dados contendo ruído (valores trocados). O intuito deste conjunto de dados é prever o valor entre 0 – 9 mostrado no display.

4.1.1.5 HEPATITIS

O conjunto foi gerado pelo BNG (RIJN et al., 2014), ou seja, os dados foram gerados artificialmente tendo como base um conjunto de dados real, ele é produzido com o intuito de ser utilizado como fluxo de dados, este conjunto de dados tenta prever se o paciente tem a doença hepatite, baseado em alguns atributos, dentre eles a idade, sexo, se utiliza esteroide, antivirais, se tem anorexia, entre outros.

4.1.2 Bases de dados reais

A seguir apresentaremos as bases de dados reais utilizados nos experimentos e coletados nos repositórios: *UCI Machine Learning Repository* (FRANK, 2010), *OpenML* (VANSCHOREN et al., 2013) e *Stream Data Mining Repository* (ZHU, 2010).

4.1.2.1 AIRLINES

Este conjunto tem um propósito de prever se o voo vai ser atrasado ou cancelado. Para isto os exemplos são compostos por 7 atributos, e os dados consistem em informações de partida e chegada de todos os voos comerciais, dia da semana, etc.

4.1.2.2 COVERTYPE

Covertime é um conjunto de cobertura florestal que abrange uma área de 30×30 metros, ele foi determinado pelo Sistema de Informações de Recursos (RIS) do Serviço Florestal dos EUA (USFS). As variáveis dos dados foram originalmente obtidas por Us Geological Survey (USGS) e da USFS. Este conjunto de dados contém 581.012 instâncias e 54 atributos categóricos e numéricos.

4.1.2.3 ELECTRICITY

Este conjunto de dados descrito por (HARRIES; WALES, 1999) e analisado por (GAMA et al., 2004). Consiste em dados coletados do Mercado Australiano de Eletricidade de Nova Gales do Sul. Neste local, os preços não são fixos e variam conforme a demanda e oferta do mercado. Este conjunto contém 45.312 instâncias no intervalo de 7 de maio de 1996 a 5 de dezembro de 1998, cada instância é referente a um período de 30 minutos, o intuito é prever a variação de preço da eletricidade em Nova Gales do Sul em relação a uma média móvel das últimas 24 horas.

4.1.2.4 SENSOR

O conjunto de dados SENSOR (ZHU, 2010) contém informações da temperatura, umidade, luz e voltagem que foram coletados de 54 sensores implementados no Intel Berkeley Research Lab, os registros são de um período de 2 meses (1 leitura a cada 1 – 3 minutos), cada rótulo é o ID do sensor, ou seja, é utilizado as informações para identificar corretamente qual é o ID do sensor. Neste conjunto de dados ocorre diversas mudanças de conceito, por exemplo, a temperatura dos sensores em alguns lugares podem aumentar, por exemplo, uma sala de reunião durante reuniões, ou um ambiente de trabalho em horário comercial.

4.1.2.5 SPAM

Spam são mensagens contendo propagandas de produtos/site, vírus, esquemas para ganhar dinheiro rápido, entre outros, na qual recebemos comumente por e-mail, enviado tanto por indivíduos quanto por mala direta, este conjunto de dados (FRANK, 2010) contém 500 atributos, na qual a maioria, possíveis palavras encontradas em e-mail que possa conter spam, na qual o último atributo, é a classe alvo onde, baseado nas palavras que contém no e-mail, é definido se é um e-mail spam ou um e-mail legítimo.

4.1.2.6 CENSUS

Os dados foram extraídos das Pesquisas Populacionais Atuais no período de 1994 e 1995, conduzidas por U.S. Census Bureau. O conjunto contém 42 atributos, dentre elas 1 atributo é a classe alvo e os outros dados demográficos que são relacionados ao emprego, o propósito é prever se a pessoa contém um rendimento acima ou abaixo de 50 mil dólares por ano.

4.1.2.7 POKERHAND

PokerHand foi apresentado em (CATTRAL; OPPACHER, 2007), nele cada instância representa uma mão composta por cinco cartas de baralhos, ele é composto por 10 atributos, cada carta é composto por dois atributos, o naipe (copas, espadas, paus, e ouros) e o valor da carta (ás, 2, 3, 4, ..., valete, rainha, rei), e tem como saída uma classe denominada "Poker hand" dez valores possíveis que representa a mão, foi utilizado o conjunto de dados original na qual contém mais de 1M de instâncias.

4.2 CONFIGURAÇÕES DOS EXPERIMENTOS

Nesta seção descrevemos todas as informações relacionadas à preparação e execução dos experimentos realizados para avaliar a acurácia dos métodos desenvolvidos em comparação aos métodos encontrados na ferramenta em diversos cenários.

Todos os experimentos tiveram uma frequência de 10 exemplos¹, sendo repetidos 30 vezes

¹ Quantidade de instâncias analisadas para obter uma amostra de desempenho (quanto menor, mais preciso o resultado)

para se obter uma média de acurácia com intervalo de confiança de 95%. Iniciamos comparando os métodos propostos em cenários sem detectores de mudança, com 3, 10 e 30% dos dados rotulados com bases reais e artificiais contendo mudança de conceito abrupto e gradual, respectivamente, com duração de 1 e 500 exemplos, localizados nas posições correspondentes a 25, 50 e 75% do total das instâncias, diferente dessas bases artificiais geradas pelo MOA, a base Hepatitis não há variações de tamanho e nem definições do tipo de mudança, por este motivo o resultado será repetido tanto na tabela de mudança abrupta quanto a tabela de mudança gradual. O método vencedor da maioria dos cenários foi comparado aos métodos já inclusos dentro da ferramenta, comparado nos mesmos cenários. Após o termino, iremos repetir os experimentos utilizando detectores de mudança de conceito. Os melhores resultados estão destacados em **negrito**.

Nos métodos que propomos utilizamos no classificador base o, Hoeffding Tree (HT) e no subclassificador Naive Bayes (NB), e como detectores de mudança de conceito o HDDM_A na sua versão e o RDDM respectivamente, nos métodos inclusos na ferramenta utilizamos o algoritmo HT e o detector o HDDM_A.

Os experimentos foram executados em uma máquina com a seguinte configuração: processador Intel Core i7-10750h, 16 GB de memória RAM e sistema operacional Windows 10, todos os testes foram executados no ambiente MOA-SS (NGUYEN; GOMES; BIFET, 2019), uma extensão do MOA (BIFET et al., 2010), permitindo a simulação do fluxo de dados semi-supervisionado e o controle da quantidade de dados rotulados, sendo possível alterar a porcentagem de rótulos visíveis.

4.2.1 Métricas de avaliação

Para estes experimentos foram utilizados as seguintes métricas:

- **Acurácia:** De modo geral, é medida a proporção de previsões corretas sobre o número total de instâncias avaliadas.

$$\text{Acurácia} = \frac{\text{Número correto de previsões}}{\text{Número total de previsões}}$$

- **Bonferroni-Dunn:** Utilizamos o diagrama de Bonferroni-Dunn cujo objetivo é mostrar que existem diferenças significativas entre um método e os seus concorrentes. É traçada uma linha horizontal entre os métodos, que indica não haver diferença significativa entre eles.

Tabela 2 – Métodos propostos com bases reais, sem detector de mudança e 3% dos dados rotulados.

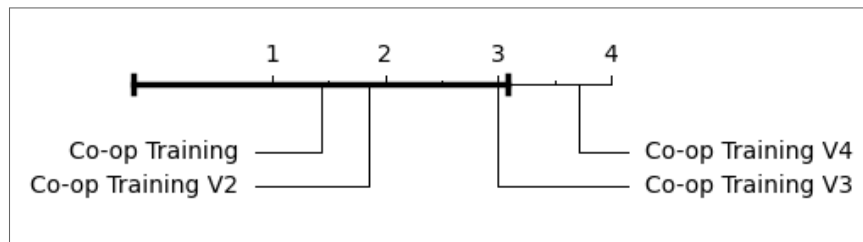
Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
AIRLINES	61,61 ± 3,98	61,13 ± 0,18	60,81 ± 0,20	57,27 ± 1,72
COVERTYPE	67,32 ± 4,38	56,88 ± 1,50	43,98 ± 1,26	38,48 ± 4,58
ELECTRICITY	67,97 ± 4,57	60,26 ± 2,33	49,20 ± 0,97	52,51 ± 1,46
SENSOR	17,58 ± 1,55	13,23 ± 0,49	13,03 ± 0,76	06,42 ± 0,42
SPAM	54,95 ± 5,44	66,93 ± 6,15	70,89 ± 5,88	49,00 ± 2,06
CENSUS	85,95 ± 5,60	80,74 ± 1,52	78,74 ± 1,31	68,70 ± 4,80
POKERHAND	49,37 ± 3,20	49,87 ± 0,25	46,89 ± 0,62	47,70 ± 1,30

Fonte: O autor (2022)

4.3 ANÁLISE DOS RESULTADOS SEM DETECTOR DE MUDANÇA DE CONCEITO

4.3.1 Análise comparativa dos métodos propostos

Figura 4 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases reais, sem detector e 3% dos dados rotulados.



Fonte: O autor (2022)

Iniciamos a análise com os experimentos da Tabela 2, na qual contém a comparação dos métodos propostos com 3% dos dados rotulados. Nela podemos perceber que o Co-op Training V1 se saiu melhor na maioria dos cenários, tendo o Co-op Training V4 com a pior performance, onde não obteve melhor desempenho em nenhuma das bases. Também observamos que ao utilizar uma pequena quantidade de dados rotulados o intervalo de confiança em algumas bases tendem a ficar com valor alto se comparado com as tabelas que contém uma maior quantia de rótulos. Ao analisar o diagrama de Bonferroni-Dunn (Ver Fig. 4) observamos que o Co-op Training V1 obteve o melhor desempenho, e uma diferença significativa se comparado com o Co-op Training V2 e V3.

Com 10% dos dados rotulados perceber-se que não ocorreram diferenças nas colocações tanto da Tabela (Ver Tab. 3) quanto no diagrama de Bonferroni (Ver Fig. 5).

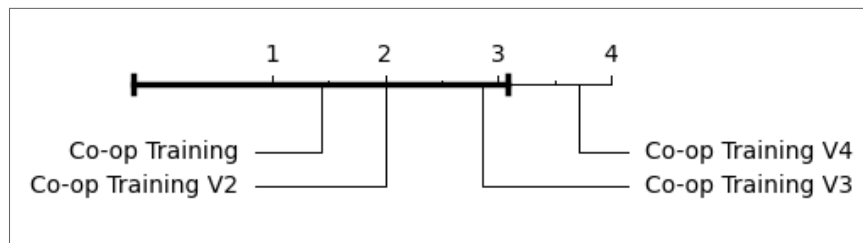
Na Tabela 4 aumentamos a quantidade de dados rotulados para 30%, nele podemos per-

Tabela 3 – Métodos propostos com bases reais, sem detector de mudança e 10% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
AIRLINES	62,78 ± 0,32	60,86 ± 0,29	60,21 ± 0,35	56,48 ± 1,41
COVERTYPE	70,71 ± 0,33	60,34 ± 1,43	53,86 ± 1,37	49,44 ± 1,50
ELECTRICITY	71,80 ± 0,79	62,11 ± 2,03	55,37 ± 1,53	55,68 ± 1,14
SENSOR	26,33 ± 0,86	20,21 ± 0,92	20,33 ± 0,74	10,24 ± 0,86
SPAM	51,44 ± 5,12	59,05 ± 3,77	62,96 ± 4,36	45,92 ± 1,65
CENSUS	87,84 ± 0,55	77,37 ± 2,13	74,60 ± 3,79	72,41 ± 5,82
POKERHAND 1M	49,63 ± 0,35	49,95 ± 0,20	48,01 ± 0,86	48,42 ± 1,09

Fonte: O autor (2022)

Figura 5 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases reais, sem detector e 10% dos dados rotulados.



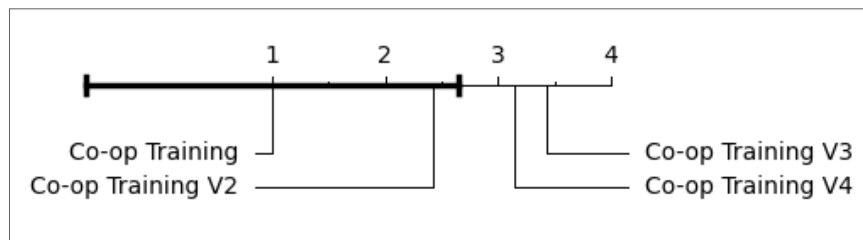
Fonte: O autor (2022)

Tabela 4 – Métodos propostos com bases reais, sem detector de mudança e 30% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
AIRLINES	64,08 ± 0,11	61,49 ± 0,17	60,51 ± 0,12	61,90 ± 0,46
COVERTYPE	74,11 ± 0,27	69,87 ± 0,71	63,98 ± 1,31	58,83 ± 0,60
ELECTRICITY	78,46 ± 0,39	74,81 ± 1,09	70,48 ± 1,68	66,39 ± 1,56
SENSOR	37,49 ± 1,23	29,12 ± 0,74	27,83 ± 1,00	31,40 ± 1,59
SPAM	63,48 ± 2,77	60,89 ± 1,70	59,01 ± 1,52	42,97 ± 2,45
CENSUS	90,24 ± 0,28	84,16 ± 1,10	80,08 ± 1,71	84,17 ± 1,10
POKERHAND 1M	50,43 ± 0,24	50,04 ± 0,06	49,40 ± 0,41	49,20 ± 0,40

Fonte: O autor (2022)

Figura 6 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases reais, sem detector e 30% dos dados rotulados.



Fonte: O autor (2022)

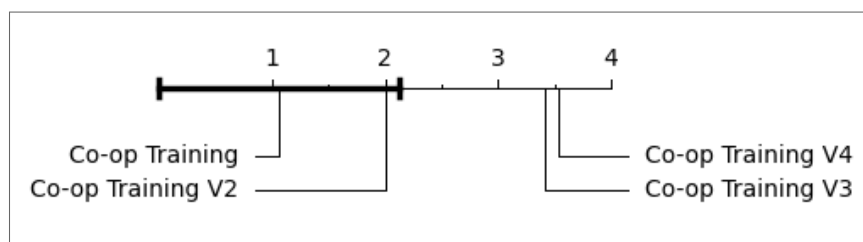
Tabela 5 – Métodos propostos com bases artificiais e mudanças abruptas, sem detector de mudança e 3% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
HEPATITIS	85,87 ± 5,52	67,58 ± 1,73	28,04 ± 0,59	79,67 ± 5,00
SEA 100K	76,86 ± 4,97	64,12 ± 3,61	49,83 ± 3,01	56,91 ± 3,36
AGRAWAL 100K	59,22 ± 4,47	53,83 ± 0,96	52,59 ± 1,02	51,04 ± 1,45
RBF 100K	65,13 ± 4,24	59,58 ± 0,68	57,89 ± 0,58	51,06 ± 0,76
LED 100K	58,27 ± 3,87	25,65 ± 2,75	18,41 ± 1,99	18,18 ± 2,42
SEA 500K	81,21 ± 5,23	75,50 ± 1,68	44,54 ± 3,78	60,03 ± 3,49
AGRAWAL 500K	67,27 ± 4,60	66,13 ± 0,98	53,75 ± 0,71	50,55 ± 1,46
RBF 500K	74,26 ± 4,80	69,05 ± 0,50	65,22 ± 0,42	50,88 ± 0,63
LED 500K	67,19 ± 4,34	29,87 ± 2,75	18,55 ± 2,20	15,86 ± 1,74
SEA 750K	81,97 ± 5,27	75,82 ± 1,54	44,14 ± 1,99	58,47 ± 3,41
AGRAWAL 750K	66,73 ± 4,64	56,67 ± 0,81	56,42 ± 0,87	50,91 ± 1,47
RBF 750K	75,84 ± 4,89	71,65 ± 0,39	67,05 ± 0,48	51,29 ± 0,94
LED 750K	68,74 ± 4,43	27,38 ± 2,94	17,68 ± 1,51	19,60 ± 2,55
SEA 1M	82,64 ± 5,31	77,18 ± 1,55	44,55 ± 2,19	54,76 ± 4,09
AGRAWAL 1M	67,00 ± 4,66	56,76 ± 0,86	56,14 ± 0,74	55,55 ± 0,83
RBF 1M	69,54 ± 4,51	72,57 ± 0,43	67,81 ± 0,31	51,51 ± 1,12
LED 1M	69,48 ± 4,47	26,45 ± 2,80	18,28 ± 1,98	18,58 ± 2,89

Fonte: O autor (2022)

ceber algumas diferenças, o Co-op Training V1 se saiu melhor em todas as bases, também podemos observar que a base SPAM com 3% dos dados rotulados obteve uma média de acurácia maior, mas ao observar o intervalo de confiança, podemos afirmar que, com uma maior quantidade de rótulos, temos uma acurácia mais estável, sem muitas variações. No diagrama de Bonferroni (Ver Fig. 6) observamos que o Co-op Training V1 se manteve bastante distante dos outros métodos, e que não houve uma diferença significativa em relação ao Co-op Training V2.

Figura 7 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças abruptas, sem detector e 3% dos dados rotulados.



Fonte: O autor (2022)

Analisando a tabela 5, na qual contém bases artificiais com as mudanças abruptas e 3% dos dados rotulados, pode-se constatar que o Co-op Training V1 demonstrou o melhor desem-

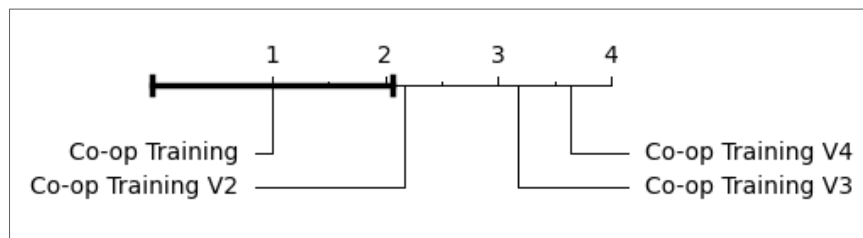
Tabela 6 – Métodos propostos com bases artificiais e mudanças abruptas, sem detector de mudança e 10% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
HEPATITIS	87,80 ± 0,07	83,99 ± 0,36	57,47 ± 1,47	84,11 ± 0,61
SEA 100K	79,98 ± 0,54	76,49 ± 1,31	67,60 ± 3,27	64,33 ± 2,44
AGRAWAL 100K	67,50 ± 1,87	55,84 ± 2,28	54,87 ± 1,83	49,26 ± 1,65
RBF 100K	69,97 ± 0,60	67,05 ± 0,96	62,94 ± 0,87	53,79 ± 1,91
LED 100K	68,64 ± 0,28	59,83 ± 1,37	36,54 ± 5,22	50,77 ± 1,93
SEA 500K	83,46 ± 0,25	81,39 ± 0,51	71,07 ± 2,52	66,57 ± 2,20
AGRAWAL 500K	72,98 ± 1,84	59,67 ± 2,18	58,30 ± 1,78	48,73 ± 1,55
RBF 500K	77,73 ± 0,36	75,83 ± 0,52	72,22 ± 0,49	55,31 ± 2,19
LED 500K	71,43 ± 0,18	64,06 ± 1,65	39,09 ± 3,38	53,82 ± 2,76
SEA 750K	83,82 ± 0,15	82,58 ± 0,48	68,23 ± 4,04	65,38 ± 2,73
AGRAWAL 750K	75,00 ± 1,33	60,27 ± 1,92	61,21 ± 1,91	47,37 ± 1,28
RBF 750K	79,67 ± 0,35	78,20 ± 0,41	73,74 ± 0,46	53,00 ± 1,86
LED 750K	71,84 ± 0,17	65,13 ± 1,60	36,77 ± 3,97	54,89 ± 2,54
SEA 1M	84,23 ± 0,12	82,67 ± 0,50	67,41 ± 3,66	66,23 ± 3,08
AGRAWAL 1M	76,57 ± 1,30	61,39 ± 2,10	68,58 ± 3,63	49,02 ± 1,68
RBF 1M	80,60 ± 0,35	79,19 ± 0,31	74,64 ± 0,44	54,65 ± 2,48
LED 1M	72,05 ± 0,13	64,14 ± 1,80	38,46 ± 3,51	56,28 ± 2,82

Fonte: O autor (2022)

penho, em seguida, o Co-op Training V2. Verificando o diagrama de Bonferroni (Ver Fig.7) percebemos que houve uma grande distância entre os dois melhores (Co-op Training V1 e V2) e os piores (Co-op Training V3 e V4). Este baixo desempenho dos dois métodos, se dar devido a utilização dos níveis de confiança dos classificadores para treinar com o rótulo previsto, ao ter apenas 3% dos dados rotulados, ocorreram diversos erros de classificações com um nível de confiança maior que o limiar, causando este baixo desempenho.

Figura 8 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças abruptas, sem detector e 10% dos dados rotulados.



Fonte: O autor (2022)

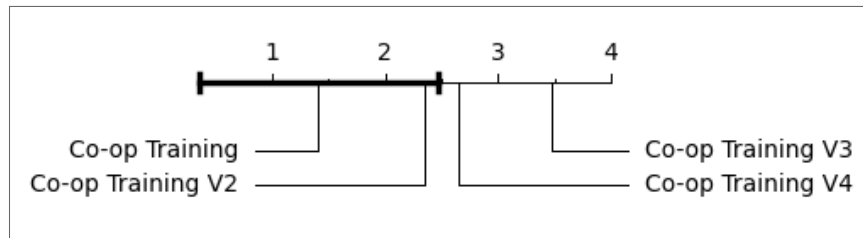
Na tabela 6, ao aumentar para 10% dos dados rotulados, o Co-op Training V1 obteve um desempenho melhor que anteriormente, tendo uma diferença significativa dos outros métodos como podemos ver na Fig.8.

Tabela 7 – Métodos propostos com bases artificiais e mudanças abruptas, sem detector de mudança e 30% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
HEPATITIS	89,43 ± 0,07	88,43 ± 0,12	86,52 ± 0,24	87,68 ± 0,06
SEA 100K	83,16 ± 0,10	82,40 ± 0,42	81,30 ± 0,49	82,41 ± 0,39
AGRAWAL 100K	72,13 ± 1,72	70,24 ± 2,42	67,85 ± 3,08	54,94 ± 1,20
RBF 100K	74,26 ± 0,43	73,54 ± 0,52	72,64 ± 0,62	70,52 ± 0,47
LED 100K	71,94 ± 0,09	71,67 ± 0,14	70,45 ± 0,69	72,14 ± 0,15
SEA 500K	85,06 ± 0,08	84,59 ± 0,17	83,79 ± 0,34	84,60 ± 0,37
AGRAWAL 500K	77,69 ± 2,01	73,47 ± 2,20	71,27 ± 3,51	55,40 ± 1,94
RBF 500K	81,96 ± 0,27	80,64 ± 0,31	79,02 ± 0,44	73,37 ± 0,23
LED 500K	73,01 ± 0,03	72,37 ± 0,19	68,51 ± 1,59	73,29 ± 0,04
SEA 750K	85,46 ± 0,08	85,04 ± 0,15	84,45 ± 0,28	85,61 ± 0,09
AGRAWAL 750K	80,40 ± 1,66	73,27 ± 2,69	73,10 ± 3,10	56,60 ± 1,88
RBF 750K	83,32 ± 0,19	82,35 ± 0,28	80,57 ± 0,44	73,83 ± 0,15
LED 750K	73,28 ± 0,03	77,52 ± 0,19	69,90 ± 0,99	73,50 ± 0,03
SEA 1M	85,74 ± 0,06	85,33 ± 0,12	84,75 ± 0,19	85,77 ± 0,14
AGRAWAL 1M	81,11 ± 1,37	76,69 ± 2,19	76,47 ± 2,57	56,48 ± 1,48
RBF 1M	84,46 ± 0,23	83,40 ± 0,27	81,62 ± 0,37	74,11 ± 0,16
LED 1M	73,35 ± 0,03	72,96 ± 0,23	69,50 ± 0,99	73,60 ± 0,02

Fonte: O autor (2022)

Figura 9 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças abruptas, sem detector e 30% dos dados rotulados.



Fonte: O autor (2022)

Aumentando a quantidade de dados rotulados para 30% (Ver Tab.7), presenciamos um trade off de desempenho entre o Co-op Training V1 e V4, afirmando o que foi dito anteriormente. Com o incremento da quantidade de dados rotulados, o classificador tende de ter mais sucesso nas classificações, melhorando o desempenho do classificador que usa apenas o nível de confiança. Ao analisarmos o Co-op Training V3 que também utiliza este método, observamos que ele não obteve um aumento de desempenho comparado ao V4, isso se deu devido à lógica utilizada na hora de treinamento, onde independente se o V3 ou V4 atingiu o limiar de confiança ele utiliza o rótulo do classificador base para treinar, ou seja, caso classifique errado com um baixo nível de confiança, mas o subclassificador realize a classificação

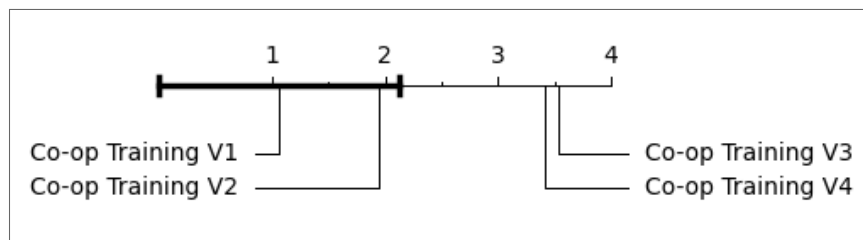
Tabela 8 – Métodos propostos com bases artificiais e mudanças graduais, sem detector de mudança e 3% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
HEPATITIS	85,87 ± 5,52	67,58 ± 1,73	28,04 ± 0,59	79,67 ± 5,00
SEA 100K	76,12 ± 4,95	64,77 ± 2,91	49,01 ± 3,74	54,70 ± 4,11
AGRAWAL 100K	59,34 ± 4,13	53,77 ± 1,05	53,50 ± 0,79	50,05 ± 1,42
RBF 100K	66,46 ± 4,72	58,98 ± 0,54	57,86 ± 0,49	50,78 ± 0,65
LED 100K	59,28 ± 3,92	24,09 ± 3,32	16,76 ± 2,07	18,31 ± 2,30
SEA 500K	67,28 ± 4,35	74,70 ± 1,58	45,65 ± 2,55	57,46 ± 3,59
AGRAWAL 500K	66,83 ± 4,68	55,85 ± 0,83	55,28 ± 0,66	49,47 ± 1,49
RBF 500K	73,90 ± 4,77	69,32 ± 0,40	65,71 ± 0,56	51,05 ± 0,74
LED 500K	67,25 ± 4,35	29,38 ± 3,48	18,03 ± 2,24	18,73 ± 2,60
SEA 750K	82,21 ± 5,29	75,28 ± 1,55	46,66 ± 2,43	58,55 ± 3,45
AGRAWAL 750K	67,24 ± 4,68	56,37 ± 0,93	54,92 ± 0,74	50,30 ± 1,54
RBF 750K	75,94 ± 4,90	71,13 ± 0,40	66,86 ± 0,31	50,51 ± 0,78
LED 750K	68,63 ± 4,42	22,69 ± 3,26	17,28 ± 2,07	18,97 ± 2,53
SEA 1M	82,51 ± 5,31	76,30 ± 1,33	43,71 ± 1,87	58,57 ± 3,11
AGRAWAL 1M	66,93 ± 4,62	56,60 ± 0,71	56,18 ± 0,68	49,87 ± 1,53
RBF 1M	77,22 ± 4,98	72,50 ± 0,46	67,67 ± 0,31	51,20 ± 0,88
LED 1M	69,37 ± 4,47	26,52 ± 2,44	18,16 ± 2,32	20,37 ± 2,92

Fonte: O autor (2022)

correta com um alto nível de confiança, o método irá treinar com o rótulo do classificador base (rótulo incorreto). Mesmo com a melhora do desempenho do V4, podemos notar que continuou obtendo uma diferença significativa no diagrama de bonferroni (Ver Fig. 9).

Figura 10 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças graduais, sem detector e 3% dos dados rotulados.



Fonte: O autor (2022)

Já nas bases com mudanças graduais com 3% dos dados rotulados (Ver Tab.8), o Co-op Training V1 se manteve com o melhor desempenho dentro dos métodos comparados. No diagrama de Bonferroni (Ver Fig.10) se manteve parecido comparando a base contendo mudanças abruptas, havendo apenas uma inversão entre o Co-op Training V3 e V4.

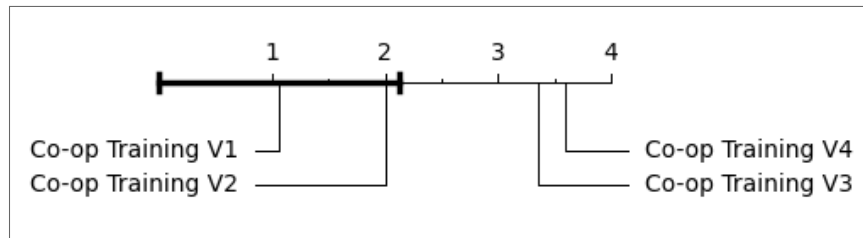
Aumentando para 10% dos dados rotulados (Ver Tab. 9), mantivemos o Co-op Training V1 e V2 com o melhor desempenho, tanto na acurácia (na maioria das bases), quanto no

Tabela 9 – Métodos propostos com bases artificiais e mudanças graduais, sem detector de mudança e 10% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
HEPATITIS	87,80 ± 0,07	83,99 ± 0,36	57,47 ± 1,47	84,11 ± 0,61
SEA 100K	80,48 ± 0,59	76,71 ± 1,42	68,34 ± 3,76	61,94 ± 4,73
AGRAWAL 100K	69,07 ± 1,19	54,98 ± 1,82	54,31 ± 1,67	50,22 ± 1,67
RBF 100K	70,18 ± 0,52	67,03 ± 0,96	63,16 ± 0,83	56,29 ± 2,24
LED 100K	68,82 ± 0,41	59,65 ± 1,49	40,76 ± 4,45	48,90 ± 2,82
SEA 500K	83,41 ± 0,25	81,50 ± 0,60	68,82 ± 3,96	66,77 ± 2,82
AGRAWAL 500K	72,94 ± 1,19	61,91 ± 2,26	58,77 ± 2,17	49,23 ± 1,58
RBF 500K	77,62 ± 0,38	76,23 ± 0,61	71,72 ± 0,58	52,40 ± 2,25
LED 500K	71,19 ± 0,20	63,46 ± 1,69	40,35 ± 2,56	55,51 ± 2,48
SEA 750K	75,24 ± 1,42	82,17 ± 0,58	69,69 ± 3,48	68,91 ± 2,47
AGRAWAL 750K	75,17 ± 1,51	59,97 ± 2,10	59,40 ± 2,18	48,23 ± 1,60
RBF 750K	79,16 ± 0,31	78,51 ± 0,36	73,22 ± 0,52	54,83 ± 2,17
LED 750K	71,76 ± 0,16	63,02 ± 1,97	39,22 ± 4,16	54,44 ± 2,47
SEA 1M	84,29 ± 0,16	83,07 ± 0,37	70,65 ± 3,82	67,25 ± 2,08
AGRAWAL 1M	76,50 ± 1,37	64,20 ± 2,22	62,08 ± 2,45	49,34 ± 1,58
RBF 1M	80,62 ± 0,30	79,13 ± 0,38	74,73 ± 0,44	56,17 ± 2,21
LED 1M	72,22 ± 0,13	62,79 ± 1,81	38,76 ± 3,46	55,93 ± 3,20

Fonte: O autor (2022)

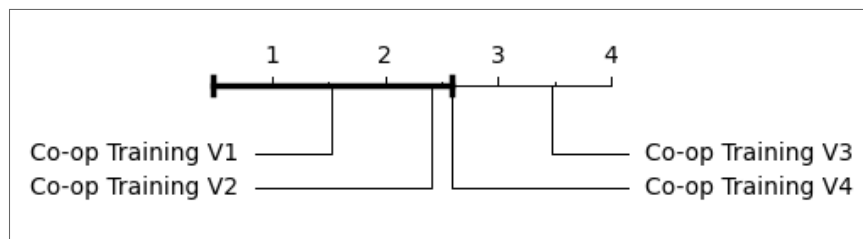
Figura 11 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças graduais, sem detector e 10% dos dados rotulados.



Fonte: O autor (2022)

diagrama de Bonferroni (11).

Figura 12 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças graduais, sem detector e 30% dos dados rotulados.



Fonte: O autor (2022)

Tabela 10 – Métodos propostos com bases artificiais e mudanças graduais, sem detector de mudança e 30% dos dados rotulados.

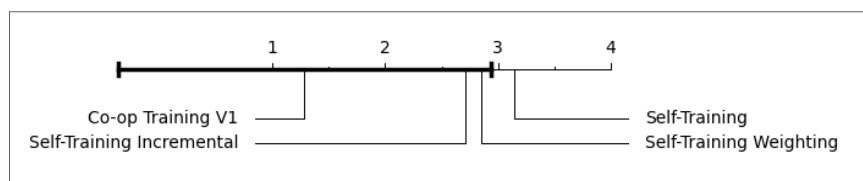
Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
HEPATITIS	89,43 ± 0,07	88,43 ± 0,12	86,52 ± 0,24	87,68 ± 0,06
SEA 100K	83,11 ± 0,20	81,85 ± 0,49	80,94 ± 0,80	82,44 ± 0,46
AGRAWAL 100K	71,96 ± 1,55	67,47 ± 2,78	64,40 ± 3,54	55,50 ± 0,70
RBF 100K	74,50 ± 0,39	73,49 ± 0,61	72,29 ± 0,84	69,81 ± 0,58
LED 100K	71,89 ± 0,08	71,68 ± 0,12	68,42 ± 3,47	72,13 ± 0,09
SEA 500K	85,12 ± 0,08	84,75 ± 0,16	83,81 ± 0,35	85,13 ± 0,16
AGRAWAL 500K	77,56 ± 1,77	73,74 ± 2,49	70,90 ± 3,23	55,59 ± 1,21
RBF 500K	81,80 ± 0,31	80,96 ± 0,37	78,88 ± 0,38	73,28 ± 0,22
LED 500K	72,99 ± 0,04	72,69 ± 0,16	68,38 ± 1,45	73,28 ± 0,04
SEA 750K	79,54 ± 1,49	85,05 ± 0,14	84,48 ± 0,26	85,36 ± 0,20
AGRAWAL 750K	80,17 ± 1,07	73,57 ± 2,92	72,27 ± 3,35	56,37 ± 2,61
RBF 750K	83,20 ± 0,22	82,25 ± 0,27	80,42 ± 0,34	73,88 ± 0,21
LED 750K	73,25 ± 0,02	72,60 ± 0,18	67,96 ± 2,04	73,48 ± 0,02
SEA 1M	85,71 ± 0,06	85,29 ± 0,12	84,92 ± 0,21	85,77 ± 0,18
AGRAWAL 1M	81,23 ± 1,33	76,77 ± 2,12	75,37 ± 2,64	55,48 ± 1,87
RBF 1M	84,54 ± 0,26	83,31 ± 0,28	81,29 ± 0,29	73,91 ± 0,19
LED 1M	73,38 ± 0,02	72,87 ± 0,17	69,35 ± 1,40	73,62 ± 0,02

Fonte: O autor (2022)

Diferente da tabela de mudanças abruptas (Ver Tab.10), na tabela de mudanças graduais com 30% dos dados rotulados (Ver Tab.10), o Co-op Training V4 venceu em mais bases, se confirmando ao verificar o diagrama de Bonferroni (Ver 12), onde o v4 está dentro do limite da linha horizontal, ou seja, podemos afirmar não haver diferenças significativas ao comparar com o V1.

4.3.2 Análise comparativa do melhor co-op vs métodos da ferramenta

Figura 13 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases reais, sem detector e 3% dos dados rotulados.



Fonte: O autor (2022)

Dentre os métodos concorrentes com 3% dos dados rotulados (Ver Tab.11), mais uma vez o Co-op Training obteve uma melhor acurácia na maioria das bases, tendo o Self-Training

Tabela 11 – Métodos concorrentes com bases reais, sem detector de mudança e 3% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
AIRLINES	61,61 ± 3,98	59,43 ± 0,97	58,71 ± 1,03	59,08 ± 0,61
COVERTYPE	67,32 ± 4,38	57,39 ± 1,48	59,11 ± 1,43	44,79 ± 2,84
ELECTRICITY	67,97 ± 4,57	66,99 ± 1,34	63,11 ± 2,28	49,85 ± 1,52
SENSOR	17,58 ± 1,55	7,55 ± 0,73	10,07 ± 0,46	15,98 ± 0,51
SPAM	54,95 ± 5,44	50,01 ± 3,61	65,93 ± 5,40	64,14 ± 5,77
CENSUS	85,95 ± 5,60	81,09 ± 3,62	83,47 ± 1,28	82,83 ± 2,34
POKERHAND	49,37 ± 320	47,01 ± 0,93	46,88 ± 0,55	47,81 ± 0,47

Fonte: O autor (2022)

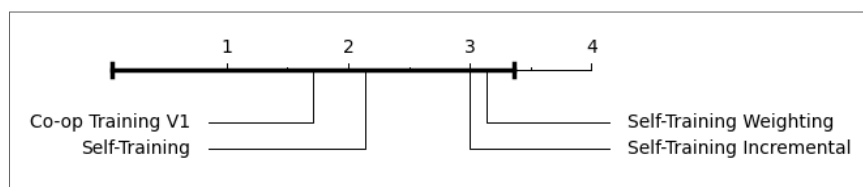
Tabela 12 – Métodos concorrentes com bases reais, sem detector de mudança e 10% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
AIRLINES	62,78 ± 0,32	60,17 ± 0,26	59,51 ± 1,31	59,46 ± 0,69
COVERTYPE	70,71 ± 0,33	61,34 ± 1,12	58,94 ± 1,31	48,57 ± 1,77
ELECTRICITY	71,80 ± 0,79	74,56 ± 0,86	62,53 ± 1,78	54,77 ± 1,40
SENSOR	26,33 ± 0,86	16,01 ± 0,57	12,84 ± 0,75	24,23 ± 0,63
SPAM	51,44 ± 5,12	53,88 ± 1,72	65,67 ± 5,16	61,39 ± 4,76
CENSUS	87,84 ± 0,55	87,99 ± 0,94	78,69 ± 2,58	78,51 ± 1,83
POKERHAND	49,63 ± 0,35	48,07 ± 0,82	47,73 ± 0,70	48,59 ± 0,56

Fonte: O autor (2022)

incremental em segundo lugar. No Bonferroni (Ver Fig.13) podemos verificar que, apenas o Self-Training tradicional teve uma diferença significativa em relação aos outros métodos.

Figura 14 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases reais, sem detector e 10% dos dados rotulados.



Fonte: O autor (2022)

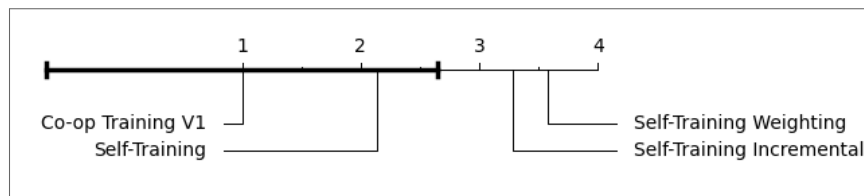
Já na Tabela 12, observamos que apenas o Self-Training Weighting não obteve o melhor desempenho em alguma base, mas mesmo assim analisando o diagrama (Ver Fig.14) observará que não houveram diferenças significativas entre as bases.

Tabela 13 – Métodos concorrentes com bases reais, sem detector de mudança e 30% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
AIRLINES	64,08 ± 0,11	60,22 ± 0,14	60,07 ± 0,53	59,29 ± 0,48
COVERTYPE	74,11 ± 0,27	69,74 ± 1,01	61,84 ± 1,50	53,12 ± 1,87
ELECTRICITY	78,46 ± 0,39	77,79 ± 0,50	74,68 ± 1,13	59,83 ± 1,97
SENSOR	37,49 ± 1,23	28,07 ± 0,82	19,16 ± 0,95	34,36 ± 0,80
SPAM	63,48 ± 2,77	62,69 ± 3,77	60,23 ± 2,03	61,73 ± 2,49
CENSUS	90,24 ± 0,28	89,28 ± 0,65	82,58 ± 1,66	82,11 ± 1,34
POKERHAND	50,43 ± 0,24	50,16 ± 0,17	49,56 ± 0,73	49,00 ± 0,59

Fonte: O autor (2022)

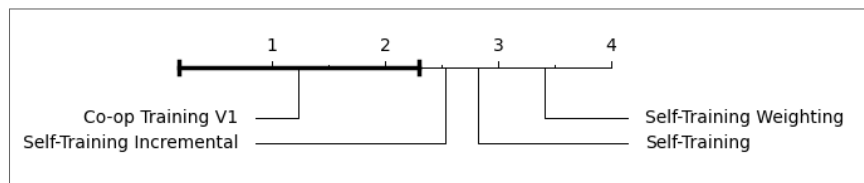
Figura 15 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases reais, sem detector e 30% dos dados rotulados.



Fonte: O autor (2022)

Aumentando para 30% dos dados rotulados (Ver Tab. 13), o Co-op Training obteve um desempenho invicto, tendo uma diferença significativa do Self-Training Incremental e o Self-Training Weighting.

Figura 16 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças abruptas, sem detector e 3% dos dados rotulados.



Fonte: O autor (2022)

Em relação à Tabela 14. Mesmo perdendo na base RBF 750K o Co-op Training obteve o melhor desempenho, tendo diferenças significativas dos outros métodos, como pode ser observado na Fig.16.

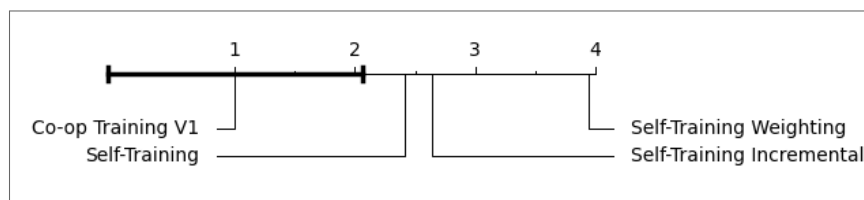
Tanto na tabela de mudanças abruptas com 10% dos dados rotulados (Ver Tab.15), quanto na de 30% (Ver Tab.16), o Co-op Training continuou se mostrando superior, nos diagramas (Ver Fig. 8 e Fig.9) podemos afirmar ao verificar que todos os algoritmos concorrentes ao

Tabela 14 – Métodos concorrentes com bases artificiais e mudanças abruptas, sem detector de mudança e 3% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
HEPATITIS	85,87 ± 5,52	81,44 ± 0,72	83,38 ± 0,28	74,50 ± 3,26
SEA 100K	76,86 ± 4,97	71,78 ± 2,30	69,92 ± 1,42	64,05 ± 1,90
AGRAWAL 100K	59,22 ± 4,47	55,86 ± 1,93	53,09 ± 1,30	54,01 ± 1,12
RBF 100K	65,13 ± 4,24	60,86 ± 2,26	61,71 ± 1,05	61,02 ± 1,25
LED 100K	58,27 ± 3,87	35,18 ± 4,72	21,09 ± 1,59	15,50 ± 1,34
SEA 500K	81,21 ± 5,23	73,68 ± 2,10	78,46 ± 0,70	72,79 ± 1,26
AGRAWAL 500K	67,27 ± 4,60	55,54 ± 2,51	55,47 ± 1,66	57,04 ± 1,01
RBF 500K	74,26 ± 4,80	68,36 ± 2,02	72,88 ± 0,70	69,62 ± 1,45
LED 500K	67,19 ± 4,34	36,42 ± 5,30	24,54 ± 2,15	16,20 ± 1,27
SEA 750K	81,97 ± 5,27	73,77 ± 2,14	79,46 ± 0,55	73,07 ± 2,24
AGRAWAL 750K	66,73 ± 4,64	57,42 ± 3,09	57,32 ± 1,72	56,93 ± 1,06
RBF 750K	75,84 ± 4,89	70,57 ± 1,07	76,52 ± 0,55	70,52 ± 2,11
LED 750K	68,74 ± 4,43	38,41 ± 4,02	25,10 ± 2,88	14,75 ± 1,55
SEA 1M	82,64 ± 5,31	73,89 ± 1,85	80,97 ± 0,39	73,99 ± 2,81
AGRAWAL 1M	67,00 ± 4,66	55,44 ± 2,12	55,80 ± 1,42	55,84 ± 1,23
RBF 1M	69,54 ± 4,51	70,35 ± 2,10	77,32 ± 0,36	71,27 ± 2,46
LED 1M	69,48 ± 4,47	36,09 ± 4,94	23,80 ± 2,40	17,11 ± 1,48

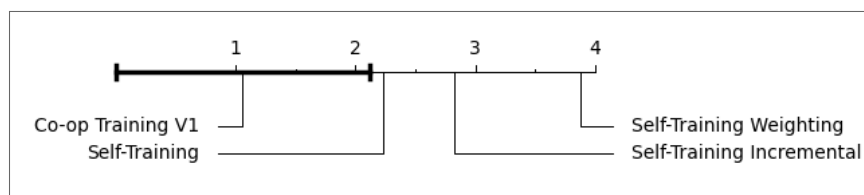
Fonte: O autor (2022)

Figura 17 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças abruptas, sem detector e 10% dos dados rotulados.



Fonte: O autor (2022)

Figura 18 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças abruptas, sem detector e 30% dos dados rotulados.



Fonte: O autor (2022)

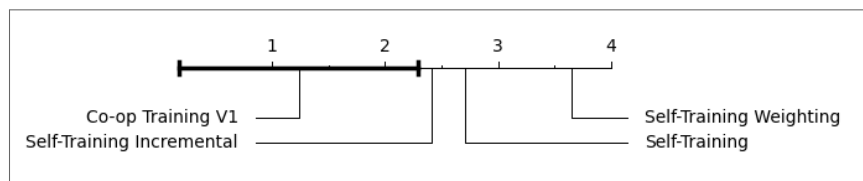
Tabela 15 – Métodos concorrentes com bases artificiais e mudanças abruptas, sem detector de mudança e 10% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
HEPATITIS	87,80 ± 0,07	85,26 ± 0,33	86,01 ± 0,18	70,38 ± 4,08
SEA 100K	79,98 ± 0,54	78,18 ± 1,11	75,63 ± 0,92	66,26 ± 1,40
AGRAWAL 100K	67,50 ± 1,87	64,22 ± 2,07	54,30 ± 2,73	56,46 ± 1,82
RBF 100K	69,97 ± 0,60	69,95 ± 1,02	64,77 ± 0,66	65,50 ± 1,19
LED 100K	68,64 ± 0,28	67,42 ± 1,18	61,86 ± 3,24	32,99 ± 1,83
SEA 500K	83,46 ± 0,25	79,60 ± 0,36	81,49 ± 0,51	71,80 ± 2,21
AGRAWAL 500K	72,98 ± 1,84	69,26 ± 3,07	58,93 ± 2,36	56,31 ± 1,27
RBF 500K	77,73 ± 0,36	75,48 ± 0,73	76,66 ± 0,68	72,98 ± 1,20
LED 500K	71,43 ± 0,18	66,52 ± 2,13	65,09 ± 2,60	33,82 ± 2,99
SEA 750K	83,82 ± 0,15	79,48 ± 0,70	82,40 ± 0,32	70,07 ± 2,72
AGRAWAL 750K	75,00 ± 1,33	66,23 ± 3,40	60,63 ± 3,07	57,78 ± 1,81
RBF 750K	79,67 ± 0,35	76,55 ± 0,79	79,01 ± 0,48	73,10 ± 2,00
LED 750K	71,84 ± 0,17	68,18 ± 0,99	67,07 ± 1,82	33,15 ± 2,06
SEA 1M	84,23 ± 0,12	79,76 ± 0,65	83,21 ± 0,28	74,61 ± 1,22
AGRAWAL 1M	76,57 ± 1,30	66,62 ± 3,24	60,52 ± 3,07	58,87 ± 1,66
RBF 1M	80,60 ± 0,35	76,80 ± 1,06	80,50 ± 0,38	75,30 ± 1,01
LED 1M	72,05 ± 0,13	67,83 ± 1,91	66,74 ± 1,78	34,41 ± 2,48

Fonte: O autor (2022)

Co-op Training estão fora da linha horizontal, ou seja, há diferenças significativas.

Figura 19 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças graduais, sem detector e 3% dos dados rotulados.



Fonte: O autor (2022)

Nas tabelas de bases artificiais com mudanças graduais com 3% (Ver Tab. 17), 10% (Ver Tab. 18), e 30% dos rótulos (Ver Tab. 19), o Co-op Training se manteve como a melhor opção para este tipo de base, como podemos afirmar nos diagramas (Ver Fig.19, Fig.20 e Fig.21), há uma diferença significativa entre o Co-op Training e os métodos concorrentes.

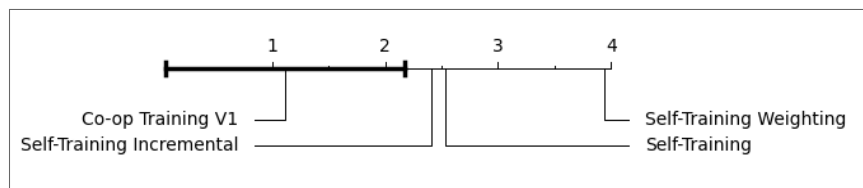
Nestes experimentos sem o uso de detectores, concluímos que o Co-op Training se tornou o melhor método dentre os testados.

Tabela 16 – Métodos concorrentes com bases artificiais e mudanças abruptas, sem detector de mudança e 30% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
HEPATITIS	89,43 ± 0,07	87,85 ± 0,18	88,07 ± 0,20	63,00 ± 3,27
SEA 100K	83,16 ± 0,10	81,60 ± 0,37	80,02 ± 0,77	67,41 ± 1,34
AGRAWAL 100K	72,13 ± 1,72	72,55 ± 1,60	62,04 ± 3,63	54,80 ± 1,31
RBF 100K	74,26 ± 0,43	74,00 ± 0,59	71,23 ± 0,66	70,44 ± 0,66
LED 100K	71,94 ± 0,09	71,53 ± 0,38	67,22 ± 3,31	67,55 ± 1,41
SEA 500K	85,06 ± 0,08	83,77 ± 0,35	84,05 ± 0,31	69,74 ± 2,17
AGRAWAL 500K	77,69 ± 2,01	76,08 ± 2,01	69,12 ± 3,39	58,08 ± 1,61
RBF 500K	81,96 ± 0,27	80,81 ± 0,35	80,18 ± 0,33	73,54 ± 1,11
LED 500K	73,01 ± 0,03	70,72 ± 0,47	67,96 ± 2,15	66,83 ± 1,64
SEA 750K	85,46 ± 0,08	84,14 ± 0,21	84,41 ± 0,26	71,01 ± 1,41
AGRAWAL 750K	80,40 ± 1,66	76,75 ± 2,09	72,80 ± 3,50	58,99 ± 1,51
RBF 750K	83,32 ± 0,19	80,48 ± 1,00	81,94 ± 0,39	75,02 ± 0,85
LED 750K	73,28 ± 0,03	70,47 ± 0,98	68,34 ± 1,72	68,90 ± 1,32
SEA 1M	85,74 ± 0,06	84,45 ± 0,27	85,07 ± 0,17	70,70 ± 1,39
AGRAWAL 1M	81,11 ± 1,37	79,43 ± 1,38	70,65 ± 3,15	54,43 ± 1,16
RBF 1M	84,46 ± 0,23	83,26 ± 0,25	82,78 ± 0,33	75,94 ± 0,68
LED 1M	73,35 ± 0,03	70,70 ± 0,79	69,55 ± 0,89	68,98 ± 1,20

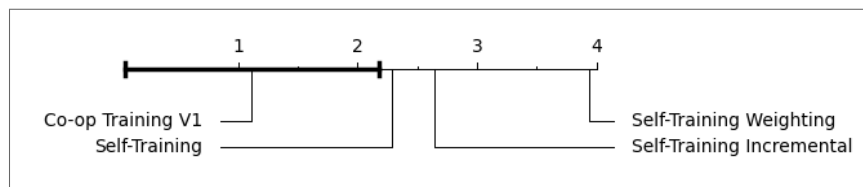
Fonte: O autor (2022)

Figura 20 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças graduais, sem detector e 3% dos dados rotulados.



Fonte: O autor (2022)

Figura 21 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças graduais, sem detector e 3% dos dados rotulados.



Fonte: O autor (2022)

Tabela 17 – Métodos concorrentes com bases artificiais e mudanças graduais, sem detector de mudança e 3% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
HEPATITIS	85,87 ± 5,52	81,44 ± 0,72	83,38 ± 0,28	74,50 ± 3,26
SEA 100K	76,12 ± 4,95	72,59 ± 1,96	71,84 ± 1,46	64,11 ± 2,30
AGRAWAL 100K	59,34 ± 4,13	57,91 ± 2,08	54,01 ± 1,03	53,44 ± 1,02
RBF 100K	66,46 ± 4,72	61,07 ± 1,83	61,42 ± 0,84	60,40 ± 1,46
LED 100K	59,28 ± 3,92	39,95 ± 3,68	21,27 ± 1,79	17,03 ± 1,48
SEA 500K	67,28 ± 4,35	74,26 ± 1,94	78,33 ± 0,56	71,33 ± 2,13
AGRAWAL 500K	66,83 ± 4,68	55,91 ± 2,89	55,72 ± 2,01	54,73 ± 1,12
RBF 500K	73,90 ± 4,77	67,46 ± 1,70	73,04 ± 0,64	67,72 ± 2,24
LED 500K	67,25 ± 4,35	33,39 ± 4,77	24,86 ± 2,35	15,15 ± 1,18
SEA 750K	82,21 ± 5,29	74,76 ± 1,94	79,60 ± 0,58	72,27 ± 2,16
AGRAWAL 750K	67,24 ± 4,68	57,75 ± 2,70	56,09 ± 1,82	57,26 ± 1,63
RBF 750K	75,94 ± 4,90	70,25 ± 1,79	75,76 ± 0,65	72,15 ± 1,96
LED 750K	68,63 ± 4,42	36,08 ± 4,66	25,75 ± 2,46	16,57 ± 1,57
SEA 1M	82,51 ± 5,31	75,30 ± 1,85	80,34 ± 0,45	72,88 ± 2,32
AGRAWAL 1M	66,93 ± 4,62	55,07 ± 1,91	56,45 ± 1,46	55,76 ± 1,01
RBF 1M	77,22 ± 4,98	71,06 ± 1,43	77,40 ± 0,52	71,22 ± 2,19
LED 1M	69,37 ± 4,47	34,65 ± 4,78	24,08 ± 2,37	16,56 ± 1,74

Fonte: O autor (2022)

Tabela 18 – Métodos concorrentes com bases artificiais e mudanças graduais, sem detector de mudança e 10% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
HEPATITIS	87,90 ± 0,07	85,26 ± 0,33	86,01 ± 0,18	70,38 ± 4,08
SEA 100K	80,48 ± 0,59	78,34 ± 0,09	75,15 ± 1,16	65,76 ± 2,27
AGRAWAL 100K	69,07 ± 1,19	64,44 ± 2,55	54,30 ± 2,74	53,64 ± 0,68
RBF 100K	70,18 ± 0,52	69,95 ± 0,82	64,87 ± 0,82	65,26 ± 1,21
LED 100K	68,82 ± 0,41	68,07 ± 0,71	63,43 ± 1,18	30,69 ± 2,97
SEA 500K	83,41 ± 0,25	80,06 ± 0,67	81,15 ± 0,41	70,40 ± 2,79
AGRAWAL 500K	72,94 ± 1,19	68,37 ± 2,47	58,77 ± 2,82	55,62 ± 3,78
RBF 500K	77,62 ± 0,38	75,14 ± 0,71	76,39 ± 0,59	73,51 ± 1,03
LED 500K	71,19 ± 0,20	67,80 ± 1,44	67,96 ± 0,75	34,12 ± 2,01
SEA 750K	75,24 ± 1,42	79,28 ± 0,91	82,47 ± 0,29	68,89 ± 3,34
AGRAWAL 750K	75,17 ± 1,51	67,27 ± 3,05	60,97 ± 3,37	59,85 ± 1,92
RBF 750K	79,16 ± 0,31	76,60 ± 0,77	79,05 ± 0,44	75,94 ± 0,97
LED 750K	71,76 ± 0,16	66,80 ± 1,66	67,69 ± 1,83	34,04 ± 2,60
SEA 1M	84,29 ± 0,16	80,19 ± 0,47	83,00 ± 0,32	74,04 ± 1,70
AGRAWAL 1M	76,50 ± 1,37	67,83 ± 2,96	64,34 ± 3,60	59,75 ± 1,60
RBF 1M	80,62 ± 0,30	77,91 ± 0,84	80,55 ± 0,39	74,51 ± 1,40
LED 1M	72,22 ± 0,13	66,30 ± 2,75	68,96 ± 1,45	33,82 ± 3,39

Fonte: O autor (2022)

Tabela 19 – Métodos concorrentes com bases artificiais e mudanças graduais, sem detector de mudança e 30% dos dados rotulados.

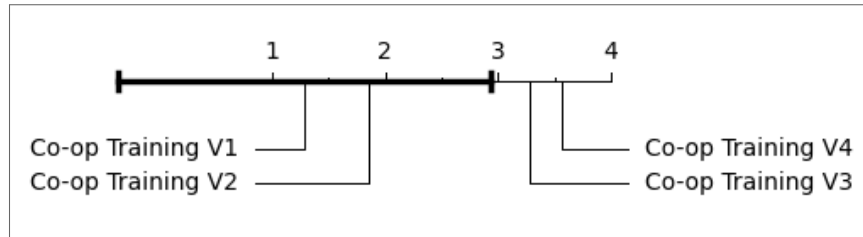
Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
HEPATITIS	89,43 ± 0,07	87,85 ± 0,18	88,07 ± 0,20	63,00 ± 3,27
SEA 100K	83,11 ± 0,20	81,94 ± 0,29	79,03 ± 0,95	68,11 ± 1,55
AGRAWAL 100K	71,96 ± 1,55	70,27 ± 2,26	58,41 ± 2,40	55,74 ± 1,59
RBF 100K	74,50 ± 0,39	73,99 ± 0,49	71,11 ± 0,80	71,01 ± 0,56
LED 100K	71,89 ± 0,08	71,25 ± 0,56	68,04 ± 1,61	67,53 ± 1,31
SEA 500K	85,12 ± 0,08	83,32 ± 0,32	83,63 ± 0,43	72,11 ± 1,49
AGRAWAL 500K	77,56 ± 1,77	75,08 ± 2,23	66,01 ± 3,49	57,54 ± 1,29
RBF 500K	81,80 ± 0,31	80,69 ± 0,37	79,52 ± 0,42	74,00 ± 0,94
LED 500K	72,99 ± 0,04	70,87 ± 0,70	68,75 ± 1,47	65,69 ± 2,02
SEA 750K	79,54 ± 1,49	83,82 ± 0,43	84,49 ± 0,27	72,35 ± 1,14
AGRAWAL 750K	80,17 ± 1,07	78,18 ± 1,38	72,22 ± 3,51	60,24 ± 1,56
RBF 750K	83,20 ± 0,22	82,18 ± 0,40	81,59 ± 0,36	75,69 ± 0,92
LED 750K	73,25 ± 0,02	70,48 ± 0,75	67,65 ± 1,59	67,96 ± 1,22
SEA 1M	85,71 ± 0,06	84,54 ± 0,21	84,86 ± 0,28	70,42 ± 1,43
AGRAWAL 1M	81,23 ± 1,33	76,14 ± 2,62	75,12 ± 2,96	59,88 ± 2,03
RBF 1M	84,54 ± 0,26	82,94 ± 0,32	82,98 ± 0,27	75,55 ± 0,85
LED 1M	73,38 ± 0,02	70,82 ± 0,64	69,33 ± 1,08	68,48 ± 1,22

Fonte: O autor (2022)

4.4 ANÁLISE DOS RESULTADOS COM DETECTOR DE MUDANÇA DE CONCEITO

4.4.1 Análise comparativa dos métodos propostos

Figura 22 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases reais, com detector e 3% dos dados rotulados.



Fonte: O autor (2022)

Iniciamos os experimentos com o uso de detector de mudanças de conceitos com a Tabela 20, nela podemos perceber que mesmo com o uso de detector o Co-op Training V1 continua tendo um melhor desempenho diante as outras versões, mesmo perdendo em duas bases (SPAM e POKERHAND), é possível verificar que a diferença foi pouca, principalmente na base POKERHAND, ao analisar o diagrama (Ver 22) afirmamos que o Co-op Training V1 e V2 estatisticamente tiveram desempenho parecido.

Já nas tabelas com 10% (Ver Tab.21) e 30% (Ver Tab.22), o Co-op Training se tornou invicto, vencendo todas as bases, e tendo respectivamente uma diferença significativa do Co-op Training V3/V4 (Ver Fig. 23) e do Co-op Training V3 (Ver Fig. 24).

Nas bases artificiais com mudanças abruptas com 3% (Ver Tab.23), 10% (Ver Tab.24) e 30% (Ver Tab.25), o Co-op Training venceu em todas ou na maioria das bases, nas bases na qual ele não obteve o melhor desempenho, ele alcançou uma acurácia próxima ao método que

Tabela 20 – Métodos propostos com bases reais, com detector de mudança e 3% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
AIRLINES	62,09 ± 0,49	60,76 ± 0,43	60,66 ± 0,07	48,98 ± 2,79
COVERTYPE	69,29 ± 0,28	56,20 ± 1,02	44,53 ± 1,29	50,08 ± 2,64
ELECTRICITY	67,97 ± 1,40	56,96 ± 2,24	50,65 ± 2,60	55,39 ± 1,20
SENSOR	26,71 ± 0,24	14,11 ± 0,54	14,13 ± 0,54	13,97 ± 0,62
SPAM	66,72 ± 2,27	69,66 ± 5,07	63,10 ± 5,04	50,25 ± 3,48
CENSUS	86,58 ± 1,31	80,29 ± 1,23	78,84 ± 1,73	68,02 ± 2,79
POKERHAND	49,37 ± 0,34	49,76 ± 0,27	47,70 ± 0,53	48,69 ± 1,06

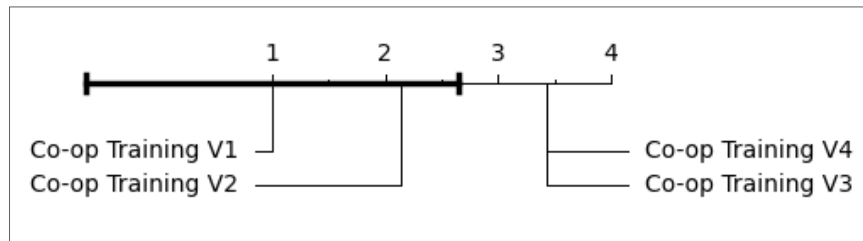
Fonte: O autor (2022)

Tabela 21 – Métodos propostos com bases reais, com detector de mudança e 10% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
AIRLINES	63,56 ± 0,32	60,65 ± 0,25	60,16 ± 0,27	56,58 ± 1,54
COVERTYPE	73,73 ± 0,33	61,07 ± 1,19	55,17 ± 1,18	53,91 ± 1,71
ELECTRICITY	72,67 ± 0,95	60,30 ± 2,31	54,95 ± 1,04	57,75 ± 1,03
SENSOR	49,22 ± 0,26	21,83 ± 0,89	19,43 ± 0,88	32,44 ± 0,65
SPAM	76,74 ± 1,94	61,93 ± 3,59	59,40 ± 3,70	47,84 ± 4,92
CENSUS	88,83 ± 0,36	75,52 ± 1,37	74,57 ± 1,97	67,74 ± 3,74
POKERHAND	49,69 ± 0,25	49,63 ± 0,34	48,35 ± 0,73	48,68 ± 1,01

Fonte: O autor (2022)

Figura 23 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases reais, com detector e 10% dos dados rotulados.



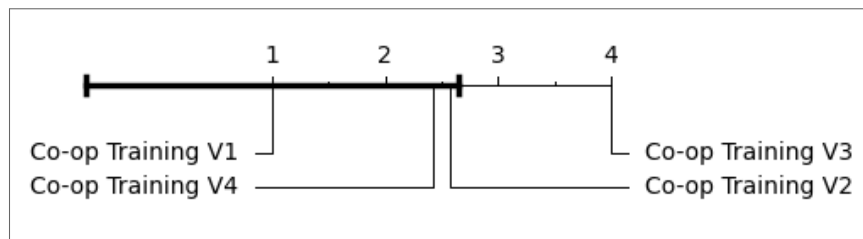
Fonte: O autor (2022)

Tabela 22 – Métodos propostos com bases reais, com detector de mudança e 30% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
AIRLINES	65,93 ± 0,15	61,35 ± 0,25	60,18 ± 0,22	61,81 ± 0,55
COVERTYPE	77,24 ± 0,16	67,38 ± 0,94	60,53 ± 1,37	69,76 ± 0,34
ELECTRICITY	78,64 ± 0,37	74,35 ± 0,99	70,42 ± 1,58	74,01 ± 0,68
SENSOR	70,97 ± 0,21	61,26 ± 0,44	50,89 ± 1,03	63,60 ± 0,15
SPAM	83,15 ± 1,38	73,42 ± 3,95	61,62 ± 4,57	64,90 ± 6,34
CENSUS	90,30 ± 0,20	80,21 ± 0,86	75,09 ± 1,26	80,44 ± 0,66
POKERHAND	49,98 ± 0,11	49,78 ± 0,27	48,32 ± 0,65	49,09 ± 0,45

Fonte: O autor (2022)

Figura 24 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases reais, com detector e 30% dos dados rotulados.



Fonte: O autor (2022)

Tabela 23 – Métodos propostos com bases artificiais e mudanças abruptas, com detector de mudança e 3% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
HEPATITIS	85,80 ± 0,12	67,11 ± 1,61	27,92 ± 0,56	75,09 ± 6,37
SEA 100K	76,03 ± 0,60	67,12 ± 3,35	50,35 ± 3,54	47,59 ± 4,15
AGRAWAL 100K	60,78 ± 2,00	54,18 ± 0,78	53,74 ± 0,85	50,07 ± 1,55
RBF 100K	65,56 ± 0,57	59,24 ± 0,64	63,20 ± 0,91	50,80 ± 0,55
LED 100K	58,53 ± 1,18	24,02 ± 2,85	17,69 ± 1,32	16,27 ± 2,26
SEA 500K	81,15 ± 0,26	73,27 ± 2,03	46,22 ± 2,47	59,50 ± 3,15
AGRAWAL 500K	68,58 ± 1,58	56,17 ± 0,49	53,81 ± 0,84	50,23 ± 1,54
RBF 500K	74,25 ± 0,47	69,13 ± 0,45	66,05 ± 0,35	51,07 ± 0,89
LED 500K	67,68 ± 0,38	24,32 ± 3,74	14,59 ± 1,45	22,34 ± 2,43
SEA 750K	82,05 ± 0,26	73,99 ± 1,33	42,70 ± 1,78	43,88 ± 2,40
AGRAWAL 750K	71,09 ± 1,53	56,88 ± 0,93	52,10 ± 4,66	56,27 ± 1,35
RBF 750K	75,62 ± 0,35	73,59 ± 1,37	66,66 ± 0,34	67,18 ± 0,48
LED 750K	68,67 ± 0,27	27,18 ± 3,21	18,45 ± 1,74	20,50 ± 2,25
SEA 1M	82,69 ± 0,15	76,79 ± 1,28	42,63 ± 1,54	58,95 ± 2,99
AGRAWAL 1M	71,05 ± 1,65	59,11 ± 1,17	54,65 ± 0,79	48,94 ± 1,50
RBF 1M	77,23 ± 0,49	72,48 ± 0,47	67,51 ± 0,39	50,68 ± 0,47
LED 1M	69,66 ± 0,23	28,23 ± 4,08	19,29 ± 2,33	15,89 ± 2,29

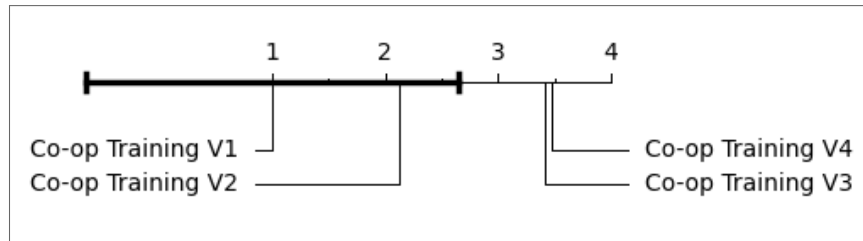
Fonte: O autor (2022)

Tabela 24 – Métodos propostos com bases artificiais e mudanças abruptas, com detector de mudança e 10% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
HEPATITIS	87,74 ± 0,10	84,13 ± 0,29	57,14 ± 1,78	83,85 ± 0,55
SEA 100K	80,00 ± 0,57	76,33 ± 1,24	63,13 ± 4,54	64,51 ± 1,62
AGRAWAL 100K	71,20 ± 1,16	58,54 ± 2,36	55,08 ± 1,70	49,69 ± 1,67
RBF 100K	70,10 ± 0,60	67,33 ± 0,84	62,79 ± 0,72	55,85 ± 1,77
LED 100K	69,13 ± 0,38	63,15 ± 1,37	41,29 ± 4,91	53,10 ± 2,96
SEA 500K	83,75 ± 0,22	81,20 ± 0,82	67,96 ± 3,67	66,19 ± 2,99
AGRAWAL 500K	77,94 ± 1,27	63,90 ± 1,68	61,64 ± 1,95	50,14 ± 1,62
RBF 500K	77,32 ± 0,53	76,01 ± 0,39	71,90 ± 0,56	54,48 ± 2,18
LED 500K	71,39 ± 0,17	63,37 ± 1,42	38,92 ± 3,95	56,44 ± 1,73
SEA 750K	84,13 ± 0,16	82,36 ± 0,52	72,21 ± 3,33	71,49 ± 2,87
AGRAWAL 750K	78,16 ± 1,22	64,14 ± 1,47	62,95 ± 1,94	60,21 ± 1,37
RBF 750K	79,11 ± 0,35	77,63 ± 0,36	73,57 ± 0,24	73,95 ± 0,39
LED 750K	71,83 ± 0,16	62,18 ± 1,74	36,57 ± 4,31	39,01 ± 2,61
SEA 1M	84,62 ± 0,17	83,09 ± 0,47	70,46 ± 3,79	64,05 ± 1,60
AGRAWAL 1M	79,22 ± 1,11	65,94 ± 1,24	63,02 ± 2,55	48,19 ± 1,55
RBF 1M	79,51 ± 0,60	79,17 ± 0,42	74,66 ± 0,40	54,25 ± 1,29
LED 1M	73,39 ± 0,03	64,06 ± 1,78	41,35 ± 2,35	56,56 ± 1,34

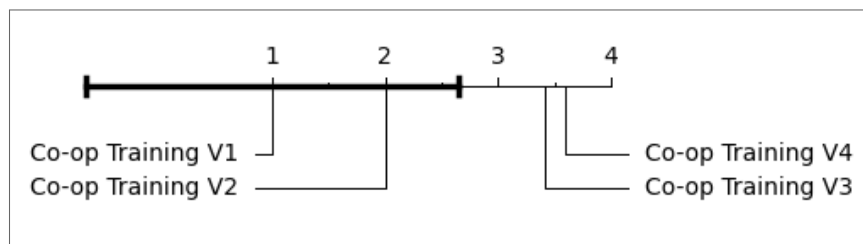
Fonte: O autor (2022)

Figura 25 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças abruptas, com detector e 3% dos dados rotulados.



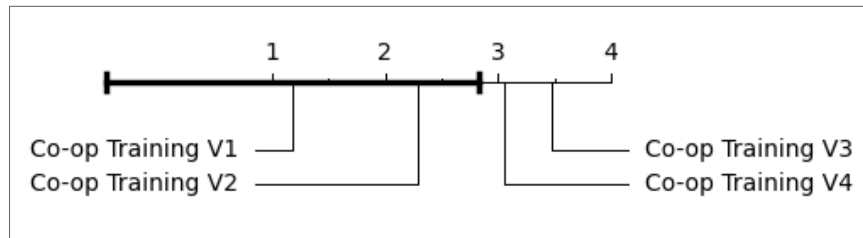
Fonte: O autor (2022)

Figura 26 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças abruptas, com detector e 10% dos dados rotulados.



Fonte: O autor (2022)

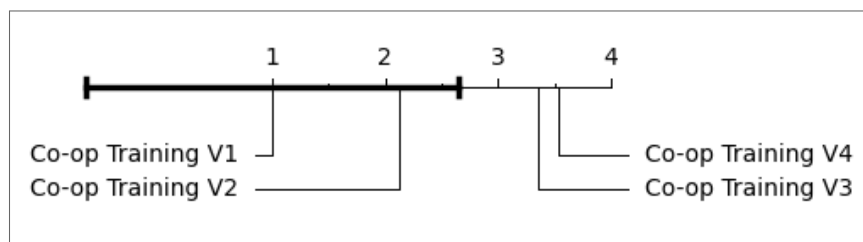
Figura 27 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças abruptas, com detector e 30% dos dados rotulados.



Fonte: O autor (2022)

atingiu o melhor desempenho, tendo sua maior diferença na base SEA 500K, ficando a 0,03 de acurácia atrás do maior. Analisando os diagramas de Bonferroni, apenas o Co-op Training V1 e V2 obtiveram desempenhos aproximados.

Figura 28 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças graduais, com detector e 3% dos dados rotulados.



Fonte: O autor (2022)

Tabela 25 – Métodos propostos com bases artificiais e mudanças abruptas, com detector de mudança e 30% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
HEPATITIS	89,34 ± 0,07	88,48 ± 0,17	86,48 ± 0,25	87,54 ± 0,13
SEA 100K	83,56 ± 0,12	82,82 ± 0,36	81,09 ± 0,76	82,39 ± 0,37
AGRAWAL 100K	74,37 ± 1,27	70,13 ± 1,55	61,98 ± 3,16	54,26 ± 1,28
RBF 100K	74,22 ± 0,36	73,59 ± 0,36	72,39 ± 0,41	69,30 ± 0,37
LED 100K	72,26 ± 0,08	72,09 ± 0,08	66,50 ± 3,77	72,22 ± 0,11
SEA 500K	86,12 ± 0,11	85,25 ± 0,22	84,11 ± 0,34	86,15 ± 0,12
AGRAWAL 500K	81,37 ± 1,41	75,05 ± 2,54	70,86 ± 3,09	54,69 ± 1,93
RBF 500K	81,83 ± 0,37	80,83 ± 0,32	79,23 ± 0,37	73,17 ± 0,22
LED 500K	73,41 ± 0,03	72,81 ± 0,11	66,52 ± 2,55	73,37 ± 0,03
SEA 750K	86,73 ± 0,06	85,95 ± 0,11	85,01 ± 0,35	84,89 ± 0,37
AGRAWAL 750K	82,47 ± 1,35	75,96 ± 2,22	72,99 ± 3,13	74,13 ± 2,90
RBF 750K	83,39 ± 0,26	82,07 ± 0,46	80,76 ± 6,07	80,37 ± 0,37
LED 750K	73,31 ± 0,02	72,94 ± 0,11	66,07 ± 1,92	65,21 ± 1,86
SEA 1M	87,10 ± 0,05	86,48 ± 0,08	85,76 ± 0,19	86,73 ± 0,13
AGRAWAL 1M	85,36 ± 1,00	79,08 ± 2,27	75,93 ± 3,32	55,15 ± 1,21
RBF 1M	84,13 ± 0,46	83,05 ± 0,30	80,79 ± 0,41	73,69 ± 0,21
LED 1M	73,59 ± 0,02	73,03 ± 0,08	68,07 ± 1,70	73,60 ± 0,03

Fonte: O autor (2022)

Tabela 26 – Métodos propostos com bases artificiais e mudanças graduais, com detector de mudança e 3% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
HEPATITIS	85,80 ± 0,12	67,11 ± 1,61	27,92 ± 0,56	75,09 ± 6,37
SEA 100K	75,86 ± 0,42	63,23 ± 4,06	47,15 ± 3,15	58,38 ± 2,87
AGRAWAL 100K	60,47 ± 1,88	54,37 ± 0,80	52,52 ± 0,72	49,72 ± 1,52
RBF 100K	65,56 ± 0,63	58,82 ± 0,72	57,17 ± 0,56	51,54 ± 0,83
LED 100K	59,28 ± 0,82	23,68 ± 2,76	16,84 ± 1,62	16,14 ± 3,30
SEA 500K	81,26 ± 0,28	72,88 ± 1,59	45,46 ± 2,70	56,02 ± 3,58
AGRAWAL 500K	68,20 ± 1,51	56,30 ± 0,60	55,93 ± 0,92	50,55 ± 1,47
RBF 500K	74,27 ± 0,48	68,84 ± 0,46	66,09 ± 1,89	51,11 ± 0,93
LED 500K	67,08 ± 0,39	31,21 ± 3,59	19,73 ± 1,60	16,25 ± 2,02
SEA 750K	82,04 ± 0,18	77,20 ± 0,87	43,71 ± 2,06	54,58 ± 4,33
AGRAWAL 750K	70,30 ± 1,39	50,15 ± 1,61	55,50 ± 0,85	47,73 ± 1,41
RBF 750K	75,52 ± 0,48	71,17 ± 0,36	66,95 ± 0,36	52,62 ± 1,49
LED 750K	68,59 ± 0,29	25,95 ± 2,52	17,40 ± 1,89	21,51 ± 3,58
SEA 1M	82,70 ± 0,19	76,97 ± 1,35	43,35 ± 2,06	53,94 ± 4,08
AGRAWAL 1M	70,02 ± 1,28	58,40 ± 0,77	56,48 ± 0,80	50,17 ± 1,61
RBF 1M	77,01 ± 0,44	72,23 ± 0,38	67,56 ± 0,37	51,89 ± 0,81
LED 1M	69,48 ± 0,26	24,49 ± 1,97	16,66 ± 1,69	17,20 ± 2,03

Fonte: O autor (2022)

Tabela 27 – Métodos propostos com bases artificiais e mudanças graduais, com detector de mudança e 10% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
HEPATITIS	87,74 ± 0,10	84,13 ± 0,29	57,14 ± 1,78	83,85 ± 0,55
SEA 100K	80,00 ± 0,57	76,72 ± 1,41	68,35 ± 3,20	65,12 ± 1,76
AGRAWAL 100K	69,16 ± 1,35	56,04 ± 2,27	53,85 ± 1,39	48,42 ± 1,50
RBF 100K	70,00 ± 0,44	65,88 ± 1,18	63,18 ± 0,99	54,36 ± 2,13
LED 100K	68,70 ± 0,38	60,67 ± 1,22	36,43 ± 4,96	51,31 ± 3,60
SEA 500K	83,43 ± 0,28	81,93 ± 0,46	67,34 ± 3,75	62,44 ± 2,95
AGRAWAL 500K	77,08 ± 1,24	63,47 ± 2,19	58,95 ± 1,82	47,42 ± 1,38
RBF 500K	77,73 ± 0,04	76,27 ± 0,48	71,60 ± 0,50	55,15 ± 2,14
LED 500K	71,24 ± 0,17	61,30 ± 2,18	38,05 ± 3,50	56,95 ± 1,50
SEA 750K	84,10 ± 0,17	81,96 ± 0,60	68,62 ± 3,61	66,73 ± 2,18
AGRAWAL 750K	77,81 ± 1,52	51,62 ± 1,84	62,40 ± 1,76	47,25 ± 1,39
RBF 750K	79,41 ± 0,33	78,27 ± 0,20	72,96 ± 0,45	54,34 ± 1,98
LED 750K	71,81 ± 0,14	61,91 ± 1,85	42,31 ± 3,68	58,59 ± 3,98
SEA 1M	84,03 ± 0,17	83,28 ± 0,34	69,68 ± 3,08	66,98 ± 2,35
AGRAWAL 1M	80,90 ± 0,76	65,89 ± 1,56	61,75 ± 1,38	50,02 ± 1,68
RBF 1M	80,56 ± 0,35	78,38 ± 1,35	74,81 ± 0,48	51,95 ± 2,06
LED 1M	72,19 ± 0,13	63,26 ± 2,05	40,95 ± 2,81	62,02 ± 2,02

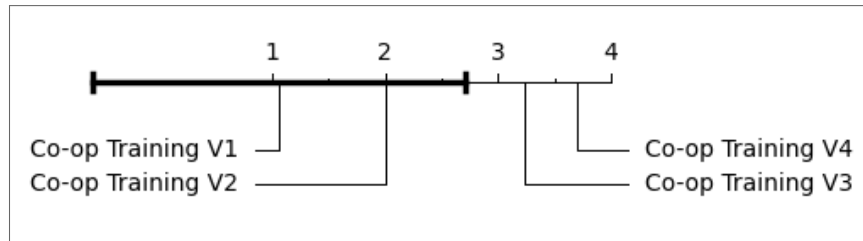
Fonte: O autor (2022)

Tabela 28 – Métodos propostos com bases artificiais e mudanças graduais, com detector de mudança e 30% dos dados rotulados.

Dataset	Co-op Training V1	Co-op Training V2	Co-op Training V3	Co-op Training V4
HEPATITIS	89,34 ± 0,07	88,48 ± 0,17	86,48 ± 0,25	87,54 ± 0,13
SEA 100K	83,50 ± 0,17	82,10 ± 0,44	81,20 ± 0,56	83,04 ± 0,33
AGRAWAL 100K	73,38 ± 1,38	65,88 ± 2,70	66,00 ± 2,57	55,25 ± 0,98
RBF 100K	74,25 ± 0,43	73,53 ± 0,61	72,89 ± 0,64	70,14 ± 0,37
LED 100K	72,23 ± 0,08	71,89 ± 0,15	69,46 ± 1,58	70,79 ± 1,95
SEA 500K	86,10 ± 0,08	85,21 ± 0,17	84,22 ± 0,31	86,02 ± 0,21
AGRAWAL 500K	81,77 ± 1,61	76,36 ± 2,19	73,90 ± 3,07	54,39 ± 1,43
RBF 500K	81,22 ± 0,69	80,75 ± 0,37	78,82 ± 0,44	73,05 ± 0,25
LED 500K	73,19 ± 0,03	72,86 ± 0,24	66,30 ± 1,93	73,42 ± 0,03
SEA 750K	86,76 ± 0,07	86,21 ± 0,12	85,17 ± 0,21	86,19 ± 0,49
AGRAWAL 750K	82,91 ± 1,60	56,21 ± 0,62	72,53 ± 7,43	54,49 ± 1,69
RBF 750K	83,44 ± 0,29	82,26 ± 0,36	80,05 ± 0,49	73,42 ± 0,18
LED 750K	73,29 ± 0,03	72,85 ± 0,16	65,91 ± 2,18	73,62 ± 0,02
SEA 1M	86,72 ± 0,08	86,37 ± 0,14	85,81 ± 0,23	86,92 ± 0,08
AGRAWAL 1M	84,74 ± 1,04	79,52 ± 2,27	76,04 ± 2,85	53,42 ± 1,43
RBF 1M	84,29 ± 0,25	82,86 ± 0,21	81,04 ± 0,46	73,90 ± 0,22
LED 1M	73,35 ± 0,03	73,06 ± 0,06	67,58 ± 1,62	72,17 ± 2,00

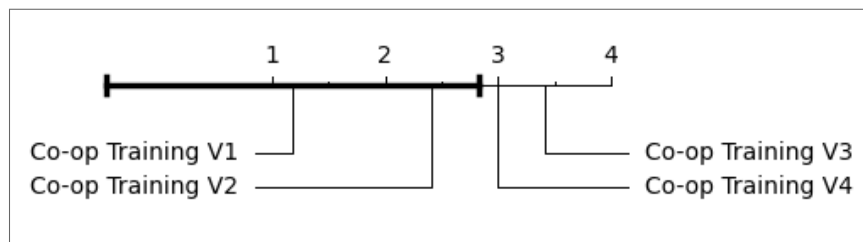
Fonte: O autor (2022)

Figura 29 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças graduais, com detector e 10% dos dados rotulados.



Fonte: O autor (2022)

Figura 30 – Diagrama de Bonferroni-Dunn - Métodos propostos com bases artificiais e mudanças graduais, com detector e 30% dos dados rotulados.

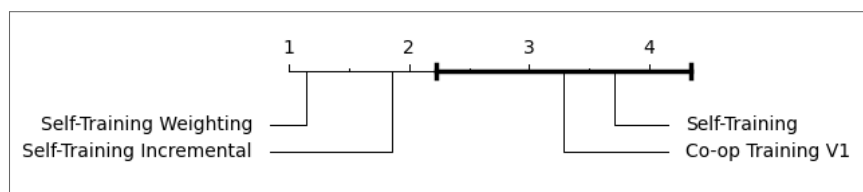


Fonte: O autor (2022)

Já nas bases artificiais com mudanças graduais (Ver Tab. 26, 27 e 28). O Co-op Training continua obtendo o melhor desempenho, não tendo uma diferença significativa na acurácia nas bases na qual ele não ficou em primeiro. Assim como nas bases com mudanças abruptas, nas de mudanças graduais o diagrama de Bonferroni continuou mantendo o Co-op Training V3 e V4 fora da linha horizontal, ou seja, há diferenças significativas entre o Co-op Training V1/V2 e V3/V4.

4.4.2 Análise comparativa do melhor co-op vs métodos da ferramenta

Figura 31 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases reais, com detector e 3% dos dados rotulados.



Fonte: O autor (2022)

Ao utilizar detectores de mudança de conceito ocorreu uma grande diferença nos resultados, observando as tabelas, podemos perceber que o Self-Training Incremental e o Self-Training

Tabela 29 – Métodos concorrentes com bases reais, com detector de mudança e 3% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
AIRLINES	62,09 $\pm$ 0,49	59,90 $\pm$ 0,17	64,58 $\pm$ 0,10	64,69 $\pm$ 0,16
COVERTYPE	69,29 $\pm$ 0,28	68,70 $\pm$ 0,19	72,14 $\pm$ 0,29	72,64 $\pm$ 0,33
ELECTRICITY	67,97 $\pm$ 1,40	69,60 $\pm$ 0,79	75,08 $\pm$ 0,36	75,25 $\pm$ 0,47
SENSOR	26,71 $\pm$ 0,24	08,83 $\pm$ 4,43	30,13 $\pm$ 0,97	30,73 $\pm$ 0,71
SPAM	66,72 $\pm$ 2,27	79,09 $\pm$ 1,59	84,11 $\pm$ 0,98	83,18 $\pm$ 1,73
CENSUS	86,58 $\pm$ 1,31	85,87 $\pm$ 0,85	90,84 $\pm$ 0,08	90,89 $\pm$ 0,04
POKERHAND	49,37 $\pm$ 0,34	46,74 $\pm$ 0,04	49,93 $\pm$ 0,06	49,96 $\pm$ 0,06

Fonte: O autor (2022)

Tabela 30 – Métodos concorrentes com bases reais, com detector de mudança e 10% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
AIRLINES	63,56 $\pm$ 0,32	62,33 $\pm$ 0,06	65,62 $\pm$ 0,12	65,56 $\pm$ 0,15
COVERTYPE	73,73 $\pm$ 0,33	74,99 $\pm$ 0,44	76,87 $\pm$ 0,19	77,01 $\pm$ 0,22
ELECTRICITY	72,67 $\pm$ 0,95	76,83 $\pm$ 0,32	78,11 $\pm$ 0,25	78,16 $\pm$ 0,30
SENSOR	49,22 $\pm$ 0,26	21,68 $\pm$ 0,05	55,70 $\pm$ 0,56	56,00 $\pm$ 0,40
SPAM	76,74 $\pm$ 1,94	84,48 $\pm$ 0,64	85,59 $\pm$ 0,67	85,48 $\pm$ 0,68
CENSUS	88,83 $\pm$ 0,36	88,42 $\pm$ 0,50	91,24 $\pm$ 0,10	91,16 $\pm$ 0,08
POKERHAND	49,69 $\pm$ 0,25	46,96 $\pm$ 0,04	50,46 $\pm$ 0,13	50,24 $\pm$ 0,10

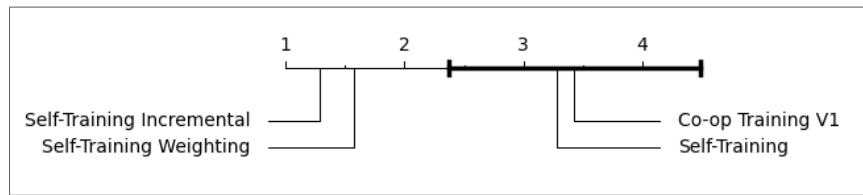
Fonte: O autor (2022)

Tabela 31 – Métodos concorrentes com bases reais, com detector de mudança e 30% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
AIRLINES	65,93 $\pm$ 0,15	64,40 $\pm$ 0,03	66,09 $\pm$ 0,07	65,98 $\pm$ 0,08
COVERTYPE	77,24 $\pm$ 0,16	79,54 $\pm$ 0,07	80,05 $\pm$ 0,08	79,99 $\pm$ 0,07
ELECTRICITY	78,64 $\pm$ 0,37	81,04 $\pm$ 0,17	81,38 $\pm$ 0,20	81,34 $\pm$ 0,20
SENSOR	70,97 $\pm$ 0,21	64,29 $\pm$ 0,08	75,45 $\pm$ 0,24	75,42 $\pm$ 0,18
SPAM	83,15 $\pm$ 1,38	88,38 $\pm$ 0,21	88,99 $\pm$ 0,30	88,87 $\pm$ 0,33
CENSUS	90,30 $\pm$ 0,20	90,40 $\pm$ 0,06	90,83 $\pm$ 0,06	90,83 $\pm$ 0,06
POKERHAND	49,98 $\pm$ 0,11	47,84 $\pm$ 0,04	50,61 $\pm$ 0,20	50,58 $\pm$ 0,21

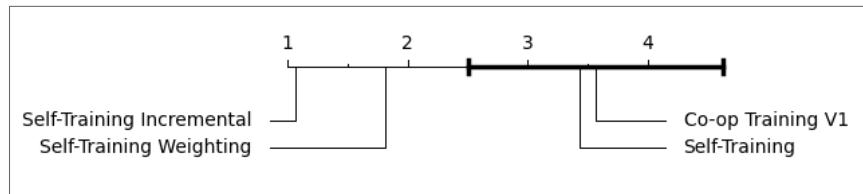
Fonte: O autor (2022)

Figura 32 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases reais, com detector e 10% dos dados rotulados.



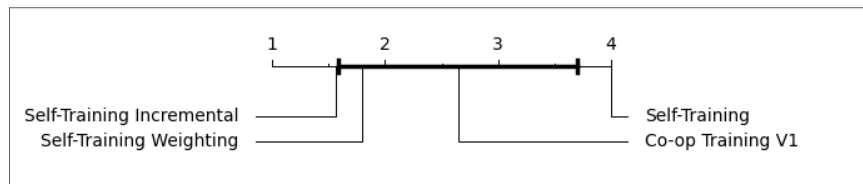
Fonte: O autor (2022)

Figura 33 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases reais, com detector e 30% dos dados rotulados.



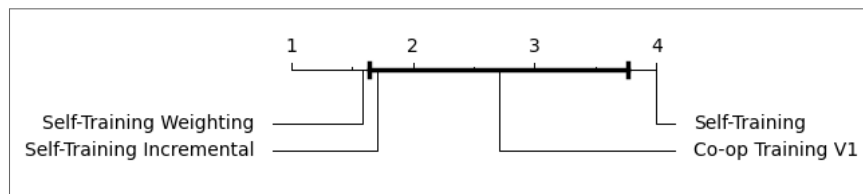
Fonte: O autor (2022)

Figura 34 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças abruptas, com detector e 3% dos dados rotulados.



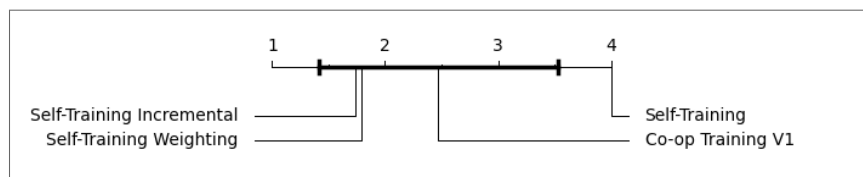
Fonte: O autor (2022)

Figura 35 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças abruptas, com detector e 10% dos dados rotulados.



Fonte: O autor (2022)

Figura 36 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças abruptas, com detector e 30% dos dados rotulados.



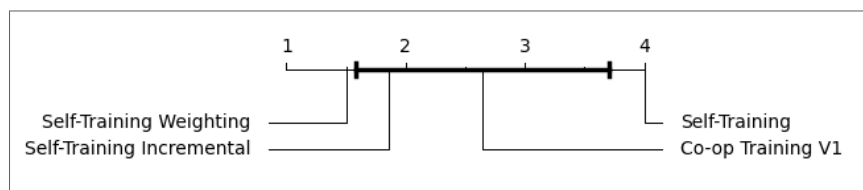
Fonte: O autor (2022)

Tabela 32 – Métodos concorrentes com bases artificiais e mudanças abruptas, com detector de mudança e 3% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
HEPATITIS	85,80 $\pm$ 0,12	82,33 $\pm$ 0,34	88,08 $\pm$ 0,07	88,12 $\pm$ 0,07
SEA 100K	76,03 $\pm$ 0,60	73,28 $\pm$ 0,88	82,55 $\pm$ 1,86	83,69 $\pm$ 0,29
AGRAWAL 100K	60,78 $\pm$ 2,00	60,06 $\pm$ 0,63	72,17 $\pm$ 0,54	71,92 $\pm$ 0,57
RBF 100K	65,56 $\pm$ 0,57	59,09 $\pm$ 0,55	69,52 $\pm$ 0,35	69,39 $\pm$ 0,40
LED 100K	58,53 $\pm$ 1,18	23,46 $\pm$ 0,23	66,06 $\pm$ 0,20	66,11 $\pm$ 0,33
SEA 500K	81,15 $\pm$ 0,26	74,24 $\pm$ 0,59	85,92 $\pm$ 0,15	85,92 $\pm$ 0,14
AGRAWAL 500K	68,58 $\pm$ 1,58	59,92 $\pm$ 1,10	77,51 $\pm$ 0,45	77,83 $\pm$ 0,37
RBF 500K	74,25 $\pm$ 0,47	59,10 $\pm$ 0,19	72,72 $\pm$ 0,35	72,55 $\pm$ 0,27
LED 500K	67,68 $\pm$ 0,38	25,08 $\pm$ 1,59	71,74 $\pm$ 0,14	71,60 $\pm$ 0,13
SEA 750K	82,05 $\pm$ 0,26	74,73 $\pm$ 0,42	85,99 $\pm$ 0,06	85,92 $\pm$ 0,16
AGRAWAL 750K	71,09 $\pm$ 1,53	61,78 $\pm$ 1,57	79,24 $\pm$ 0,22	79,16 $\pm$ 0,37
RBF 750K	75,62 $\pm$ 0,35	59,51 $\pm$ 0,20	73,33 $\pm$ 0,37	73,51 $\pm$ 0,33
LED 750K	68,67 $\pm$ 0,27	23,88 $\pm$ 0,12	72,42 $\pm$ 0,08	72,44 $\pm$ 0,06
SEA 1M	82,69 $\pm$ 0,15	75,15 $\pm$ 0,37	86,35 $\pm$ 0,22	86,17 $\pm$ 0,07
AGRAWAL 1M	71,05 $\pm$ 1,65	61,33 $\pm$ 0,69	80,41 $\pm$ 0,41	80,28 $\pm$ 0,53
RBF 1M	77,23 $\pm$ 0,49	59,04 $\pm$ 0,16	74,84 $\pm$ 0,48	74,29 $\pm$ 0,32
LED 1M	69,66 $\pm$ 0,23	24,00 $\pm$ 0,09	72,69 $\pm$ 0,07	72,63 $\pm$ 0,10

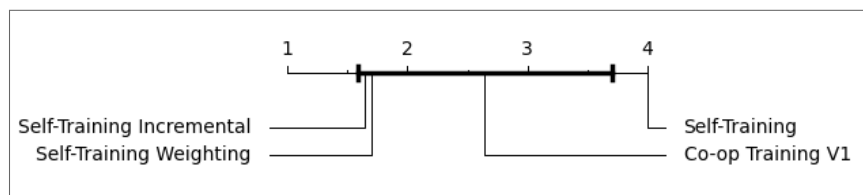
Fonte: O autor (2022)

Figura 37 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças graduais, com detector e 3% dos dados rotulados.



Fonte: O autor (2022)

Figura 38 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças graduais, com detector e 10% dos dados rotulados.



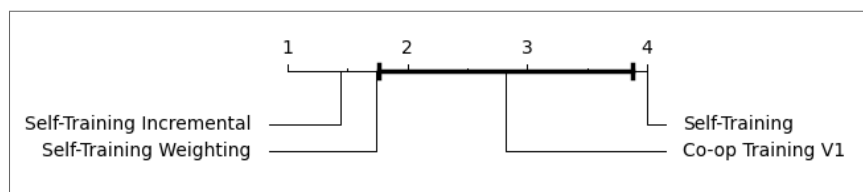
Fonte: O autor (2022)

Tabela 33 – Métodos concorrentes com bases artificiais e mudanças abruptas, com detector de mudança e 10% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
HEPATITIS	87,74 $\pm$ 0,10	84,83 $\pm$ 0,12	89,32 $\pm$ 0,07	89,20 $\pm$ 0,08
SEA 100K	80,00 $\pm$ 0,57	79,11 $\pm$ 0,27	85,38 $\pm$ 0,17	85,44 $\pm$ 0,21
AGRAWAL 100K	71,20 $\pm$ 1,16	66,39 $\pm$ 0,36	75,83 $\pm$ 0,40	75,52 $\pm$ 0,49
RBF 100K	70,10 $\pm$ 0,60	63,15 $\pm$ 0,17	71,33 $\pm$ 0,34	71,12 $\pm$ 0,32
LED 100K	69,13 $\pm$ 0,38	41,40 $\pm$ 0,25	70,60 $\pm$ 0,19	70,70 $\pm$ 0,14
SEA 500K	83,75 $\pm$ 0,22	79,08 $\pm$ 0,15	86,36 $\pm$ 0,07	86,25 $\pm$ 0,09
AGRAWAL 500K	77,94 $\pm$ 1,27	65,98 $\pm$ 0,25	81,78 $\pm$ 0,33	81,95 $\pm$ 0,49
RBF 500K	77,32 $\pm$ 0,53	63,43 $\pm$ 0,07	76,77 $\pm$ 0,45	77,07 $\pm$ 0,49
LED 500K	71,39 $\pm$ 0,17	41,61 $\pm$ 0,06	73,08 $\pm$ 0,05	73,04 $\pm$ 0,06
SEA 750K	84,13 $\pm$ 0,16	79,15 $\pm$ 0,12	86,73 $\pm$ 0,08	86,71 $\pm$ 0,08
AGRAWAL 750K	78,16 $\pm$ 1,22	66,73 $\pm$ 0,99	83,63 $\pm$ 0,45	83,08 $\pm$ 0,54
RBF 750K	79,11 $\pm$ 0,35	63,40 $\pm$ 0,05	78,12 $\pm$ 0,54	78,45 $\pm$ 0,52
LED 750K	71,83 $\pm$ 0,16	41,75 $\pm$ 0,07	73,27 $\pm$ 0,05	73,29 $\pm$ 0,07
SEA 1M	84,62 $\pm$ 0,17	79,00 $\pm$ 0,11	86,89 $\pm$ 0,08	86,83 $\pm$ 0,06
AGRAWAL 1M	79,22 $\pm$ 1,11	66,98 $\pm$ 0,80	83,94 $\pm$ 0,49	84,11 $\pm$ 0,37
RBF 1M	79,51 $\pm$ 0,60	63,29 $\pm$ 0,06	79,21 $\pm$ 0,54	80,25 $\pm$ 0,50
LED 1M	73,39 $\pm$ 0,03	41,73 $\pm$ 0,04	73,42 $\pm$ 0,06	73,45 $\pm$ 0,04

Fonte: O autor (2022)

Figura 39 – Diagrama de Bonferroni-Dunn - Métodos concorrentes com bases artificiais e mudanças graduais, com detector e 30% dos dados rotulados.



Fonte: O autor (2022)

obtiveram os melhores resultados, enquanto com o Co-op Training mesmo com o pequeno aumento não foi possível ultrapassar, o motivo de não ter sido o melhor algoritmo como no teste sem detector se dar devido aos falsos positivos, detectado pelos detectores de mudança, causando uma perda de acurácia. Mesmo perdendo na maioria das bases de dados, podemos observar nos diagramas que nas bases artificiais (abrupta e gradual) o Co-op Training não obteve um resultado com uma diferença significativa em alguns momentos do Self-Training Incremental ou Self-Training Weighting. Já o self-training, além de não ter tido o melhor

Tabela 34 – Métodos concorrentes com bases artificiais e mudanças abruptas, com detector de mudança e 30% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
HEPATITIS	89,34 $\pm$ 0,07	86,38 $\pm$ 0,08	90,21 $\pm$ 0,06	90,24 $\pm$ 0,08
SEA 100K	83,56 $\pm$ 0,12	83,00 $\pm$ 0,14	85,93 $\pm$ 0,13	86,15 $\pm$ 0,13
AGRAWAL 100K	74,37 $\pm$ 1,27	71,25 $\pm$ 0,91	80,18 $\pm$ 0,26	80,14 $\pm$ 0,47
RBF 100K	74,22 $\pm$ 0,36	67,42 $\pm$ 0,09	74,14 $\pm$ 0,33	74,72 $\pm$ 0,31
LED 100K	72,26 $\pm$ 0,08	56,84 $\pm$ 0,10	72,51 $\pm$ 0,07	72,44 $\pm$ 0,09
SEA 500K	86,12 $\pm$ 0,11	83,07 $\pm$ 0,04	87,12 $\pm$ 0,06	87,22 $\pm$ 0,04
AGRAWAL 500K	81,37 $\pm$ 1,41	70,73 $\pm$ 0,15	84,86 $\pm$ 0,58	84,51 $\pm$ 0,43
RBF 500K	81,83 $\pm$ 0,37	67,57 $\pm$ 0,21	81,43 $\pm$ 0,59	81,04 $\pm$ 0,51
LED 500K	73,41 $\pm$ 0,03	57,36 $\pm$ 0,06	73,44 $\pm$ 0,03	73,44 $\pm$ 0,03
SEA 750K	86,73 $\pm$ 0,06	83,07 $\pm$ 0,05	87,59 $\pm$ 0,06	87,58 $\pm$ 0,06
AGRAWAL 750K	82,47 $\pm$ 1,35	70,58 $\pm$ 0,09	85,29 $\pm$ 0,26	85,72 $\pm$ 0,63
RBF 750K	83,39 $\pm$ 0,26	67,55 $\pm$ 0,03	82,24 $\pm$ 0,54	82,45 $\pm$ 0,62
LED 750K	73,31 $\pm$ 0,02	57,43 $\pm$ 0,04	73,55 $\pm$ 0,04	73,54 $\pm$ 0,02
SEA 1M	87,10 $\pm$ 0,05	83,12 $\pm$ 0,03	87,88 $\pm$ 0,03	87,85 $\pm$ 0,04
AGRAWAL 1M	85,36 $\pm$ 1,00	71,11 $\pm$ 0,65	86,77 $\pm$ 0,71	85,94 $\pm$ 0,65
RBF 1M	84,13 $\pm$ 0,46	67,44 $\pm$ 0,03	83,68 $\pm$ 0,73	83,82 $\pm$ 0,60
LED 1M	73,59 $\pm$ 0,02	57,40 $\pm$ 0,05	73,56 $\pm$ 0,03	73,55 $\pm$ 0,03

Fonte: O autor (2022)

desempenho em alguma base, conforme o diagrama de Bonferroni nas bases artificiais, foi o único a ter uma diferença significativa de todos os métodos. Nas bases, independentes da quantidade de rótulos o Co-op Training se mostrou inferior tanto em acurácia quanto ao diagrama de Bonferroni.

Tabela 35 – Métodos concorrentes com bases artificiais e mudanças graduais, com detector de mudança e 3% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
HEPATITIS	85,80 $\pm$ 0,12	82,33 $\pm$ 0,34	88,08 $\pm$ 0,07	88,12 $\pm$ 0,07
SEA 100K	75,86 $\pm$ 0,42	73,83 $\pm$ 0,73	83,20 $\pm$ 0,26	83,96 $\pm$ 0,27
AGRAWAL 100K	60,47 $\pm$ 1,88	59,96 $\pm$ 0,55	72,86 $\pm$ 0,49	72,17 $\pm$ 0,46
RBF 100K	65,56 $\pm$ 0,63	59,23 $\pm$ 0,71	70,00 $\pm$ 0,41	70,10 $\pm$ 0,27
LED 100K	59,28 $\pm$ 0,82	23,68 $\pm$ 0,32	65,49 $\pm$ 0,36	65,88 $\pm$ 0,26
SEA 500K	81,26 $\pm$ 0,28	74,74 $\pm$ 0,64	86,06 $\pm$ 0,15	86,06 $\pm$ 0,13
AGRAWAL 500K	68,20 $\pm$ 1,51	61,17 $\pm$ 0,84	77,78 $\pm$ 0,29	77,21 $\pm$ 0,39
RBF 500K	74,27 $\pm$ 0,48	59,40 $\pm$ 0,13	72,68 $\pm$ 0,25	72,79 $\pm$ 0,39
LED 500K	67,08 $\pm$ 0,39	24,03 $\pm$ 0,14	71,60 $\pm$ 0,13	71,89 $\pm$ 0,08
SEA 750K	82,04 $\pm$ 0,18	74,41 $\pm$ 0,38	85,93 $\pm$ 0,17	86,07 $\pm$ 0,10
AGRAWAL 750K	70,30 $\pm$ 1,39	60,98 $\pm$ 0,87	79,07 $\pm$ 0,40	79,02 $\pm$ 0,44
RBF 750K	75,52 $\pm$ 0,48	59,59 $\pm$ 0,22	74,08 $\pm$ 0,34	73,68 $\pm$ 0,43
LED 750K	68,59 $\pm$ 0,29	23,79 $\pm$ 0,12	72,34 $\pm$ 0,10	72,38 $\pm$ 0,08
SEA 1M	82,70 $\pm$ 1,28	74,82 $\pm$ 0,36	86,06 $\pm$ 0,11	86,19 $\pm$ 0,11
AGRAWAL 1M	70,02 $\pm$ 1,28	59,83 $\pm$ 0,19	80,25 $\pm$ 0,34	79,99 $\pm$ 0,41
RBF 1M	77,01 $\pm$ 0,44	59,15 $\pm$ 0,16	74,69 $\pm$ 0,29	75,07 $\pm$ 0,36
LED 1M	69,48 $\pm$ 0,26	24,00 $\pm$ 0,09	72,63 $\pm$ 0,13	72,73 $\pm$ 0,11

Fonte: O autor (2022)

Tabela 36 – Métodos concorrentes com bases artificiais e mudanças graduais, com detector de mudança e 10% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
HEPATITIS	87,74 $\pm$ 0,10	84,83 $\pm$ 0,12	89,32 $\pm$ 0,07	89,20 $\pm$ 0,08
SEA 100K	80,00 $\pm$ 0,57	78,69 $\pm$ 0,31	85,37 $\pm$ 0,16	85,44 $\pm$ 0,22
AGRAWAL 100K	69,16 $\pm$ 1,35	65,79 $\pm$ 0,49	75,88 $\pm$ 0,51	75,58 $\pm$ 0,25
RBF 100K	70,00 $\pm$ 0,44	62,91 $\pm$ 0,33	71,69 $\pm$ 0,30	71,64 $\pm$ 0,23
LED 100K	68,70 $\pm$ 0,38	41,11 $\pm$ 0,20	70,52 $\pm$ 0,16	70,63 $\pm$ 0,19
SEA 500K	83,43 $\pm$ 0,28	79,05 $\pm$ 0,16	86,33 $\pm$ 0,11	86,36 $\pm$ 0,09
AGRAWAL 500K	77,08 $\pm$ 1,24	66,25 $\pm$ 0,48	81,78 $\pm$ 0,46	81,74 $\pm$ 0,41
RBF 500K	77,73 $\pm$ 0,04	63,44 $\pm$ 0,09	77,14 $\pm$ 0,45	76,88 $\pm$ 0,52
LED 500K	71,24 $\pm$ 0,17	41,64 $\pm$ 0,11	73,06 $\pm$ 0,07	73,02 $\pm$ 0,08
SEA 750K	84,10 $\pm$ 0,17	79,13 $\pm$ 0,15	86,63 $\pm$ 0,08	86,68 $\pm$ 0,07
AGRAWAL 750K	77,81 $\pm$ 1,52	66,12 $\pm$ 0,46	83,74 $\pm$ 0,65	83,98 $\pm$ 0,71
RBF 750K	79,41 $\pm$ 0,33	63,39 $\pm$ 0,07	79,12 $\pm$ 0,39	77,69 $\pm$ 0,50
LED 750K	71,81 $\pm$ 0,14	43,73 $\pm$ 2,73	73,27 $\pm$ 0,08	73,30 $\pm$ 0,08
SEA 1M	84,03 $\pm$ 0,17	79,07 $\pm$ 0,11	86,95 $\pm$ 0,04	86,88 $\pm$ 0,04
AGRAWAL 1M	80,90 $\pm$ 0,76	66,59 $\pm$ 0,56	83,82 $\pm$ 0,58	83,86 $\pm$ 0,43
RBF 1M	80,56 $\pm$ 0,35	63,31 $\pm$ 0,06	79,47 $\pm$ 0,68	79,67 $\pm$ 0,68
LED 1M	72,19 $\pm$ 0,13	43,75 $\pm$ 2,70	73,41 $\pm$ 0,02	73,40 $\pm$ 0,03

Fonte: O autor (2022)

Tabela 37 – Métodos concorrentes com bases artificiais e mudanças graduais, com detector de mudança e 30% dos dados rotulados.

Dataset	Co-op Training V1	Self-Training	Self-Training Incremental	Self-Training Weighting
HEPATITIS	89,34 $\pm$ 0,07	86,38 $\pm$ 0,08	90,21 $\pm$ 0,06	90,24 $\pm$ 0,08
SEA 100K	83,50 $\pm$ 0,17	83,20 $\pm$ 0,12	86,02 $\pm$ 0,10	83,55 $\pm$ 0,33
AGRAWAL 100K	73,38 $\pm$ 1,38	70,53 $\pm$ 0,25	79,95 $\pm$ 0,36	79,97 $\pm$ 0,30
RBF 100K	74,25 $\pm$ 0,43	67,38 $\pm$ 0,09	74,66 $\pm$ 0,40	74,47 $\pm$ 0,37
LED 100K	72,23 $\pm$ 0,08	56,96 $\pm$ 0,11	72,54 $\pm$ 0,10	72,49 $\pm$ 0,09
SEA 500K	86,10 $\pm$ 0,08	83,12 $\pm$ 0,05	87,16 $\pm$ 0,07	87,16 $\pm$ 0,07
AGRAWAL 500K	81,77 $\pm$ 1,61	70,88 $\pm$ 0,50	84,85 $\pm$ 0,62	84,85 $\pm$ 0,54
RBF 500K	81,22 $\pm$ 0,69	67,64 $\pm$ 0,05	81,50 $\pm$ 0,46	81,48 $\pm$ 0,48
LED 500K	73,19 $\pm$ 0,03	57,33 $\pm$ 0,06	73,41 $\pm$ 0,04	73,42 $\pm$ 0,03
SEA 750K	86,76 $\pm$ 0,07	83,12 $\pm$ 0,04	87,61 $\pm$ 0,07	87,60 $\pm$ 0,04
AGRAWAL 750K	82,91 $\pm$ 1,60	71,05 $\pm$ 0,51	85,85 $\pm$ 0,53	85,35 $\pm$ 0,31
RBF 750K	83,44 $\pm$ 0,29	67,62 $\pm$ 0,03	83,60 $\pm$ 0,35	82,20 $\pm$ 0,54
LED 750K	73,29 $\pm$ 0,39	57,39 $\pm$ 0,04	73,54 $\pm$ 0,03	73,55 $\pm$ 0,02
SEA 1M	86,72 $\pm$ 0,08	83,06 $\pm$ 0,03	87,82 $\pm$ 0,06	87,83 $\pm$ 0,06
AGRAWAL 1M	84,74 $\pm$ 1,04	71,20 $\pm$ 0,66	86,40 $\pm$ 0,90	85,65 $\pm$ 0,35
RBF 1M	84,29 $\pm$ 0,25	67,48 $\pm$ 0,03	84,16 $\pm$ 0,39	83,61 $\pm$ 0,42
LED 1M	73,35 $\pm$ 0,03	57,52 $\pm$ 0,04	73,60 $\pm$ 0,02	73,60 $\pm$ 0,02

Fonte: O autor (2022)

4.5 RESPOSTA A PERGUNTAS DE PESQUISA

Nesta seção iremos avaliar e responder as seis questões de pesquisas que este trabalho se propôs a investigar.

- **QP1: Qual método dentre os propostos têm a melhor acurácia em ambiente com e sem detector?** Analisando os diagramas de Bonferroni e a tabela de acurácia, concluímos que o melhor método foi o Co-op Training V1, seguido por Co-op Training V2 que mostrou ter um desempenho estatisticamente similar. Por fim, o Co-op Training V3 e V4 foram os piores.
- **QP2: Qual método dentre os concorrentes têm a melhor acurácia em ambiente com e sem detector?** Sem o uso de detectores, o Self-Training se mostrou melhor na maioria dos cenários, seguido pelo Self-Training incremental. Quando adicionado o detector de mudança cenário muda, pois, o Self-Training se torna o pior método, porque não se sai melhor em nenhum cenário, então, nos cenários que utilizam detectores de mudanças o Self-Training Incremental se sai melhor, com um desempenho bastante similar ao do Self-Training Weighting.
- **QP3: Ocorreram diferenças significativas entre 3% 10% e 30% dos dados rotulados?** Foi observado um aumento de acurácia significativo na maioria das bases e a diminuição do intervalo de confiança, mas em poucos casos foram vistos uma perda de desempenho. Com mais rótulos disponíveis, os classificadores tendem a ter mais desempenho, pois, tem mais dados corretos disponíveis para serem utilizado para treinamento.
- **QP4: Os métodos propostos tiveram grande diferença ao lidar com mudanças graduais e mudanças abruptas?** Tanto nos cenários com o uso de detectores, quanto nos cenários que não foram utilizados os detectores, a maioria se mantiveram estáveis, mas foram observados alguns casos na qual tiveram perda ou ganho significativo de desempenho.
- **QP5: As diferentes quantidades de instâncias nas bases artificiais interferiram no resultado?** O que observamos foram pequenas oscilações, mas geralmente houve pequenos incrementos nas médias das acurácias.

- **QP6: Ocorreram diferenças ao utilizar detectores de mudança de conceito?** Nos experimentos sem o uso de detectores de mudanças o Co-op Training foi o melhor método na maioria dos cenários, ao ser adicionado o detector, o desempenho se manteve superior dentre os métodos propostos, mas diante o Self-Traning Incremental e Weighting o resultado foi o inverso, se manteve dentre os piores, sendo o melhor apenas se comparado ao Self-Training. O Co-op Training não se saiu melhor também no cenário com o uso de detectores, provavelmente devido aos falsos positivos encontrados em excesso pelos detectores, causando um desempenho abaixo do esperado.

5 CONCLUSÕES

Nesta dissertação foi realizada uma análise comparativa em ambiente com fluxo contínuo de dados, mudança abrupta e gradual, bases reais e artificiais, com 3%, 10% e 30% dos dados rotulados. Os métodos avaliados foram: Co-op Training V1 e suas variações (Co-op Training V2, Co-op Training V3 e Co-op Training V4) e Self-Training e suas variações (Self-Training Incremental e Self-Training Weighting). Foi decidido utilizar os parâmetros padrão dos métodos, nessa avaliação utilizamos a acurácia como medida de desempenho.

O detalhe de funcionamento do Co-op Training e suas variações além da sua principal diferença foi apresentado no Capítulo 3.

No Capítulo 4 foi realizada uma análise comparativa entre os métodos Co-op Training sendo o melhor comparado com os Self-Training. Utilizamos 3%, 10% e 30% dos dados rotulados para comparar se há uma diferença significativa nesse aumento de dados contendo rótulos. Foram escolhidas 7 bases reais: AIRLINES, COVERTYPE, ELETRICITY, SENSOR, SPAM, CENSUS E POKERHAND. Também foram selecionadas 5 bases artificiais: HEPATITIS, SEA, AGRAWAL, RandomRBF e LED, tendo esses as suas versões 500K, 750K e 1M de instâncias, além de mudança abrupta e gradual, com exceção do HEPATITIS, pois, só tem uma única versão.

Cada experimento foi executado 30 vezes, sendo obtido a média da acurácia junto ao intervalo de confiança com 95%. Primeiramente analisamos o desempenho dos métodos sem detectores de mudança, inicialmente na comparação entre os métodos propostos, observamos que o Co-op Training nas bases reais foi o melhor com 3% dos dados rotulados mesmo perdendo em algumas bases, mas obteve um desempenho superior ao aumentar para 10% e 30%, já nas bases artificiais com mudança abrupta, o Co-op Training se saiu melhor com 3% e 10% dos dados enquanto com 30% perdeu em várias bases se comparado ao Co-op Training V4. Nas bases com mudança gradual novamente o Co-op Training obteve o melhor desempenho, com 10% enquanto nos 30% ocorreu um empate nos números de bases que obtiveram maior desempenho com o Co-op Training V4, chegamos à conclusão que dentre os métodos propostos sem o uso de detectores o Co-op Training foi o melhor método, por este motivo a seguir detalharemos o desempenho diante o Self-Training e as suas variações. Nas bases reais com 3% e 10% assim como com os métodos propostos o Co-op Training perdeu em algumas bases, mas mesmo assim obteve o melhor desempenho na maioria, enquanto nos

30% venceu em todas, nas mudanças abruptas com 10% o Co-op Training venceu de forma invicta, nas graduais perdeu apenas na base SEA para o Self-Training Incremental e em um RandomRBF. Analisando o desempenho ao utilizar detectores de mudança de conceito foi notado um grande aumento de desempenho em algumas bases como, por exemplo: SENSOR. Iniciamos o experimento novamente comparando os métodos propostos, na base real o Co-op Training venceu a maioria das bases nas três quantidades de dados rotulados, nas bases contendo mudança abrupta tanto o Co-op Training V4 venceu em várias bases com 30% dos rótulos, mas o Co-op Training venceu a maioria, e para finalizar, nas graduais novamente o Co-op Training se saiu melhor, por este motivo novamente iremos utilizar o Co-op Training na comparação com o Self-Training e as suas variações. Ao utilizar os detectores, observamos um cenário totalmente diferente ao comparar nosso método com os encontrados na ferramenta. Nas bases reais o Co-op Training perdeu em todas tendo em primeiro Self-Training Incremental, seguido do Self-Training Weighting, nas bases com mudança abrupta o Co-op Training nas três opções de quantidade de dados rotulados venceu apenas em algumas bases RBF e uma LED, perdendo novamente para o Self-Training Incremental que obteve o melhor desempenho comparando com o Self-Training Weighting, e para finalizar, nas bases graduais o Co-op Training novamente perdeu na maioria dos cenários.

Conforme os resultados obtidos nos experimentos, nos cenários que não foram utilizados o detector de mudança de conceito o Co-op Training obteve o melhor desempenho de todos os métodos utilizados para comparar, tendo o seu desempenho aumentado ao utilizar uma maior quantidade de dados rotulados, enquanto ao utilizar detector manteve um valor aproximado de acurácia, mas o Self-Training Incremental e Self-Training Weighting obteve um aumento significativo causando esta derrota no Co-op Training. Consequentemente, baseado nos resultados podemos concluir que o Co-op Training obtém melhores desempenhos ao não utilizar detectores de mudança de conceito. Para entender melhor o motivo em um trabalho futuro iremos analisar detalhadamente o comportamento dos detectores e realizar outra análise para encontrar qual a melhor combinação de classificadores e detectores, ao utilizar detectores obtivemos um melhor desempenho no Self-Training Incremental, seguido do Self-Training Weighting. Devemos optar preferencialmente pela utilização do Co-op Training em cenários onde não há mudança de conceito, ou que seja necessário o baixo uso de poder computacional.

5.1 PUBLICAÇÃO

O método Co-op Training desenvolvido para a elaboração desta dissertação, foi submetido, aceito e apresentado na conferência IEEE SMC 2021.

MONTEIRO, P. M.; SOARES, E.; BARROS, R. S. M. Co-op training: a semi-supervised learning method for data streams. In: 2021 IEEE International Conference on Systems, Man, and Cybernetics (SMC). [S.l.], 2021. p. 933–938.

5.2 TRABALHOS FUTUROS

- Realizar um estudo para encontrar a melhor combinação de detectores de mudança para elevar a acurácia do Co-op Training;
- Utilizar outras métricas para avaliar o desempenho.
- Produzir experimentos utilizando como classificador outros algoritmos ensembles e avaliar o desempenho;
- Adicionar métodos que estão no estado da arte para uma melhor comparação.

REFERÊNCIAS

- ABREU, D.; CARVALHO, I.; ABELÉM, A. Seleção de características em stream para a detecção de ataques de rede. In: SBC. *Anais do XXI Simpósio Brasileiro em Segurança da Informação e de Sistemas Computacionais*. [S.l.], 2021. p. 197–210.
- AGARWAL, B.; MITTAL, N. Machine learning approach for sentiment analysis. In: *Prominent feature extraction for sentiment analysis*. [S.l.]: Springer, 2016. p. 21–45.
- AGRAHARI, S.; SINGH, A. K. Concept drift detection in data stream mining: A literature review. *Journal of King Saud University-Computer and Information Sciences*, Elsevier, 2021.
- AGRAWAL, R.; IMIELINSKI, T.; SWAMI, A. Database mining: A performance perspective. *IEEE Transactions on Knowledge and Data Engineering*, v. 5, n. 6, p. 914–925, 1993. Special issue on Learning and Discovery in Knowledge-Based Databases. Disponível em: <<http://www.almaden.ibm.com/software/quest/Publications/ByDate.html>>.
- ALSAFARI, S.; SADAQUI, S. Semi-supervised self-training of hate and offensive speech from social media. *Applied Artificial Intelligence*, Taylor & Francis, v. 35, n. 15, p. 1621–1645, 2021.
- ANDERSON, R.; KOH, Y. S.; DOBBIE, G.; BIFET, A. Recurring concept meta-learning for evolving data streams. *Expert Systems with Applications*, Elsevier, v. 138, p. 112832, 2019.
- AWOYEMI, J. O.; ADETUNMBI, A. O.; OLUWADARE, S. A. Credit card fraud detection using machine learning techniques: A comparative analysis. In: IEEE. *2017 international conference on computing networking and informatics (ICCNi)*. [S.l.], 2017. p. 1–9.
- BAENA-GARCIA, M.; CAMPO-ÁVILA, J. del; FIDALGO, R.; BIFET, A.; GAVALDA, R.; MORALES-BUENO, R. Early drift detection method. In: *Fourth international workshop on knowledge discovery from data streams*. [S.l.: s.n.], 2006. v. 6, p. 77–86.
- BARROS, R. S.; CABRAL, D. R.; Gonçalves Jr., P. M.; SANTOS, S. G. T. C. RDDM: Reactive drift detection method. *Expert Systems with Applications*, Elsevier, v. 90, p. 344–355, 2017.
- BARROS, R. S. M.; SANTOS, S. G. T. C. A large-scale comparison of concept drift detectors. *Information Sciences*, Elsevier, v. 451, p. 348–370, 2018.
- BASAR, S.; ALI, M.; OCHOA-RUIZ, G.; ZAREEI, M.; WAHEED, A.; ADNAN, A. Unsupervised color image segmentation: A case of rgb histogram based k-means clustering initialization. *Plos one*, Public Library of Science San Francisco, CA USA, v. 15, n. 10, p. e0240015, 2020.
- BIFET, A. Adaptive learning and mining for data streams and frequent patterns. *ACM SIGKDD Explorations Newsletter*, ACM New York, NY, USA, v. 11, n. 1, p. 55–56, 2009.
- BIFET, A.; GAVALDA, R. Learning from time-changing data with adaptive windowing. In: SIAM. *Proceedings of the 2007 SIAM international conference on data mining*. [S.l.], 2007. p. 443–448.
- BIFET, A.; GAVALDA, R.; HOLMES, G.; PFAHRINGER, B. *Machine learning for data streams: with practical examples in MOA*. [S.l.]: MIT press, 2018.

- BIFET, A.; HOLMES, G.; PFAHRINGER, B.; KRANEN, P.; KREMER, H.; JANSEN, T.; SEIDL, T. Moa: Massive online analysis, a framework for stream classification and clustering. In: PMLR. *Proceedings of the First Workshop on Applications of Pattern Analysis*. [S.l.], 2010. p. 44–50.
- BLUM, A.; MITCHELL, T. Combining labeled and unlabeled data with co-training. In: *Proceedings of 11th annual conference on Computational learning theory*. [S.l.: s.n.], 1998. p. 92–100.
- CALDAS, W. L. Proposta de dois métodos semi-supervisionados baseados na máquina de aprendizagem mínima utilizando co-training. 2017.
- CAMPS-VALLS, G.; MARSHEVA, T. V. B.; ZHOU, D. Semi-supervised graph-based hyperspectral image classification. *IEEE transactions on Geoscience and Remote Sensing*, IEEE, v. 45, n. 10, p. 3044–3054, 2007.
- CATTRAL, R.; OPPACHER, F. Discovering rules in the poker hand dataset. In: *Proceedings of the 9th annual conference on Genetic and evolutionary computation*. [S.l.: s.n.], 2007. p. 1870–1870.
- CHAMBERLAIN, D.; KODGULE, R.; GANELIN, D.; MIGLANI, V.; FLETCHER, R. R. Application of semi-supervised deep learning to lung sound analysis. In: IEEE. *2016 38th annual international conference of the IEEE engineering in medicine and biology society (EMBC)*. [S.l.], 2016. p. 804–807.
- CHEN, Z.; KHOA, L. D. V.; TEOH, E. N.; NAZIR, A.; KARUPPIAH, E. K.; LAM, K. S. Machine learning techniques for anti-money laundering (aml) solutions in suspicious transaction detection: a review. *Knowledge and Information Systems*, Springer, v. 57, n. 2, p. 245–285, 2018.
- CHONG, Y.; DING, Y.; YAN, Q.; PAN, S. Graph-based semi-supervised learning: A review. *Neurocomputing*, Elsevier, v. 408, p. 216–230, 2020.
- CUNNINGHAM, P.; CORD, M.; DELANY, S. J. Supervised learning. In: *Machine learning techniques for multimedia*. [S.l.]: Springer, 2008. p. 21–49.
- DELANY, S. J.; CUNNINGHAM, P.; TSYMBAL, A.; COYLE, L. A case-based technique for tracking concept drift in spam filtering. In: SPRINGER. *International Conference on Innovative Techniques and Applications of Artificial Intelligence*. [S.l.], 2004. p. 3–16.
- DEPARTMENT, S. Published by S. R. *Total Data Volume Worldwide 2010-2025*. 2022. Disponível em: <<https://www.statista.com/statistics/871513/worldwide-data-created/>>.
- DOMINGOS, P.; HULTEN, G. Mining high-speed data streams. In: *Proceedings of the sixth ACM SIGKDD international conference on Knowledge discovery and data mining*. [S.l.: s.n.], 2000. p. 71–80.
- ERMAN, J.; MAHANTI, A.; ARLITT, M.; COHEN, I.; WILLIAMSON, C. Semi-supervised network traffic classification. In: *Proceedings of the 2007 ACM SIGMETRICS international conference on Measurement and modeling of computer systems*. [S.l.: s.n.], 2007. p. 369–370.
- FACELI, K.; LORENA, A. C.; GAMA, J.; CARVALHO, A. C. P. d. L. F. d. *Inteligência artificial: uma abordagem de aprendizado de máquina*. [S.l.]: LTC, 2011.

- FERREIRA, A.; CAVALCANTE, R. Análise comparativa entre algoritmos de aprendizagem de máquina em classificação de imagens de radiografia no auxílio ao diagnóstico de pneumonia. In: SBC. *Anais da XIX Escola Regional de Computação Bahia, Alagoas e Sergipe*. [S.l.], 2019. p. 245–254.
- FILHO, V. M. *Um novo algoritmo de agrupamento semisupervisionado baseado no Fuzzy C-Means*. Dissertação (Mestrado) — Universidade Federal de Pernambuco, 2009.
- FRANK, A. Uci machine learning repository. irvine, ca: University of california, school of information and computer science. <http://archive.ics.uci.edu/ml>, 2010.
- FRIAS-BLANCO, I.; CAMPO-ÁVILA, J. del; RAMOS-JIMENEZ, G.; MORALES-BUENO, R.; ORTIZ-DIAZ, A.; CABALLERO-MOTA, Y. Online and non-parametric drift detection methods based on hoeffding's bounds. *IEEE Transactions on Knowledge and Data Engineering*, IEEE, v. 27, n. 3, p. 810–823, 2014.
- GAMA, J. *Knowledge discovery from data streams*. [S.l.]: CRC Press, 2010.
- GAMA, J.; MEDAS, P.; CASTILLO, G.; RODRIGUES, P. Learning with drift detection. In: SPRINGER. *Brazilian symposium on artificial intelligence*. [S.l.], 2004. p. 286–295.
- Gonçalves Jr., P. M.; SANTOS, S. G. T. C.; BARROS, R. S. M.; VIEIRA, D. C. L. A comparative study on concept drift detectors. *Expert Systems with Applications*, Elsevier, v. 41, n. 18, p. 8144–8156, 2014.
- GRABNER, H.; LEISTNER, C.; BISCHOF, H. Semi-supervised on-line boosting for robust tracking. In: SPRINGER. *European conference on computer vision*. [S.l.], 2008. p. 234–247.
- HARRIES, M.; WALES, N. S. Splice-2 comparative evaluation: Electricity pricing. Citeseer, 1999.
- HIDALGO, J. I. G. *Experiências com variações prequential para avaliação da aprendizagem em fluxo de dados*. Dissertação (Mestrado) — Universidade Federal de Pernambuco, 2017.
- HOEFFDING, W. Probability inequalities for sums of bounded random variables. In: *The collected works of Wassily Hoeffding*. [S.l.]: Springer, 1994. p. 409–426.
- KAMPER, H.; LIVESCU, K.; GOLDWATER, S. An embedded segmental k-means model for unsupervised segmentation and clustering of speech. In: IEEE. *2017 IEEE Automatic Speech Recognition and Understanding Workshop (ASRU)*. [S.l.], 2017. p. 719–726.
- KANG, M.; JAMESON, N. J. Machine learning: Fundamentals. *Prognostics and Health Management of Electronics: Fundamentals, Machine Learning, and the Internet of Things*, Wiley Online Library, p. 85–109, 2018.
- KHANALE, P. B.; PATHAK, V. M. Weka: A dynamic software suit for machine learning & exploratory data analysis. In: . [S.l.: s.n.], 2011.
- KHATRI, S.; ARORA, A.; AGRAWAL, A. P. Supervised machine learning algorithms for credit card fraud detection: a comparison. In: IEEE. *2020 10th International Conference on Cloud Computing, Data Science & Engineering (Confluence)*. [S.l.], 2020. p. 680–683.
- KRAMER, O. Dimensionality reduction by unsupervised k-nearest neighbor regression. In: IEEE. *2011 10th international conference on machine learning and applications and workshops*. [S.l.], 2011. v. 1, p. 275–278.

- LIU, B. Supervised learning. In: *Web data mining*. [S.l.]: Springer, 2011. p. 63–132.
- LOUSSAIEF, S.; ABDELKRIM, A. Machine learning framework for image classification. In: IEEE. *2016 7th International Conference on Sciences of Electronics, Technologies of Information and Telecommunications (SETIT)*. [S.l.], 2016. p. 58–61.
- MITCHELL, T.; BUCHANAN, B.; DEJONG, G.; DIETTERICH, T.; ROSENBLOOM, P.; WAIBEL, A. Machine learning. *Annual review of computer science*, Annual Reviews 4139 El Camino Way, PO Box 10139, Palo Alto, CA 94303-0139, USA, v. 4, n. 1, p. 417–433, 1990.
- MOHIBULLAH, M.; HOSSAIN, M. Z.; HASAN, M. Comparison of euclidean distance function and manhattan distance function using k-medoids. *International Journal of Computer Science and Information Security*, LJS Publishing, v. 13, n. 10, p. 61, 2015.
- MONTEIRO, P. M.; SOARES, E.; BARROS, R. S. M. de. Co-op training: a semi-supervised learning method for data streams. In: IEEE. *2021 IEEE International Conference on Systems, Man, and Cybernetics (SMC)*. [S.l.], 2021. p. 933–938.
- NASCIMENTO, E. G. S.; TAVARES, O. de L.; SOUZA, A. F. D. A cluster-based algorithm for anomaly detection in time series using mahalanobis distance. In: THE STEERING COMMITTEE OF THE WORLD CONGRESS IN COMPUTER SCIENCE, COMPUTER ... *Proceedings on the International Conference on Artificial Intelligence (ICAI)*. [S.l.], 2015. p. 622.
- NEIVA, D. H. *Uma proposta de sistema para tradução entre linguagens de sinais*. Dissertação (Mestrado) — Universidade Federal de Pernambuco, 2015.
- NGUYEN, M. H. L.; GOMES, H. M.; BIFET, A. Semi-supervised learning over streaming data using MOA. In: *IEEE International Conference on Big Data (Big Data)*. [S.l.: s.n.], 2019. p. 553–562.
- NING, X.; WANG, X.; XU, S.; CAI, W.; ZHANG, L.; YU, L.; LI, W. A review of research on co-training. *Concurrency and computation: practice and experience*, Wiley Online Library, p. e6276, 2021.
- NISHIDA, K.; YAMAUCHI, K. Detecting concept drift using statistical testing. In: SPRINGER. *International conference on discovery science*. [S.l.], 2007. p. 264–269.
- NIU, X.; HAN, H.; SHAN, S.; CHEN, X. Multi-label co-regularization for semi-supervised facial action unit recognition. *Advances in neural information processing systems*, v. 32, 2019.
- OLSHEN, L. B. J. F. R.; STONE, C. Classification and regression trees. In: _____. Belmont, California: Wadsworth International Group, 1984. p. 43–49. ISBN 0412048418. Disponível em: <<http://www.ics.uci.edu/~mlern/databases/led-display-creator/>>.
- RIJN, J. N. van; HOLMES, G.; PFAHRINGER, B.; VANSCHOREN, J. Algorithm selection on data streams. In: SPRINGER. *International Conference on Discovery Science*. [S.l.], 2014. p. 325–336.
- RISH, I. An empirical study of the Naive Bayes classifier. In: *IJCAI workshop empirical methods in artif. intellig.* [S.l.: s.n.], 2001. v. 3, p. 41–46.

- SANTOS, S. G. T. C.; BARROS, R. S. M.; Gonçalves Jr., P. M. Optimizing the parameters of drift detection methods using a genetic algorithm. In: IEEE. *2015 IEEE 27th International Conference on Tools with Artificial Intelligence (ICTAI)*. [S.l.], 2015. p. 1077–1084.
- SINAGA, K. P.; YANG, M.-S. Unsupervised k-means clustering algorithm. *IEEE access*, IEEE, v. 8, p. 80716–80727, 2020.
- SINCLAIR, C.; PIERCE, L.; MATZNER, S. An application of machine learning to network intrusion detection. In: IEEE. *Proceedings 15th Annual Computer Security Applications Conference (ACSAC'99)*. [S.l.], 1999. p. 371–377.
- SINDHWANI, V.; NIYOGI, P.; BELKIN, M. A co-regularization approach to semi-supervised learning with multiple views. In: CITESEER. *Proceedings of ICML workshop on learning with multiple views*. [S.l.], 2005. v. 2005, p. 74–79.
- SINGHAL, P.; SRIVASTAVA, P. K.; TIWARI, A. K.; SHUKLA, R. K. A survey: Approaches to facial detection and recognition with machine learning techniques. In: SPRINGER. *Proceedings of Second Doctoral Symposium on Computational Intelligence*. [S.l.], 2022. p. 103–125.
- SIVAKUMAR, S.; NAYAK, S. R.; VIDYANANDINI, S.; KUMAR, J. A.; PALAI, G. An empirical study of supervised learning methods for breast cancer diseases. *Optik*, Elsevier, v. 175, p. 105–114, 2018.
- STAVENS, D.; HOFFMANN, G.; THRUN, S. Online speed adaptation using supervised learning for high-speed, off-road autonomous driving. In: *IJCAI*. [S.l.: s.n.], 2007. p. 2218–2224.
- STREET, W. N.; KIM, Y. A streaming ensemble algorithm (sea) for large-scale classification. In: *Proceedings of the seventh ACM SIGKDD international conference on Knowledge discovery and data mining*. [S.l.: s.n.], 2001. p. 377–382.
- TANHA, J.; SOMEREN, M. V.; AFSARMANESH, H. Semi-supervised self-training for decision tree classifiers. *International Journal of Machine Learning and Cybernetics*, Springer, v. 8, n. 1, p. 355–370, 2017.
- TSYMBAL, A. The problem of concept drift: definitions and related work. *Computer Science Department, Trinity College Dublin*, Citeseer, v. 106, n. 2, p. 58, 2004.
- VAJDA, S.; SANTOSH, K. A fast k-nearest neighbor classifier using unsupervised clustering. In: SPRINGER. *International conference on recent trends in image processing and pattern recognition*. [S.l.], 2016. p. 185–193.
- VANSCHOREN, J.; RIJN, J. N. van; BISCHL, B.; TORGO, L. Openml: Networked science in machine learning. *SIGKDD Explorations*, ACM, New York, NY, USA, v. 15, n. 2, p. 49–60, 2013. Disponível em: <<http://doi.acm.org/10.1145/2641190.2641198>>.
- VIEGAS, F. A. R. Explorando estratégias bayesianas eficientes e eficazes para classificação de texto. Universidade Federal de Minas Gerais, 2015.
- VIRUPAKSHAPPA, K.; ORUKLU, E. Unsupervised machine learning for ultrasonic flaw detection using gaussian mixture modeling, k-means clustering and mean shift clustering. In: IEEE. *2019 IEEE International Ultrasonics Symposium (IUS)*. [S.l.], 2019. p. 647–649.

- WANG, D.; LIN, J.; CUI, P.; JIA, Q.; WANG, Z.; FANG, Y.; YU, Q.; ZHOU, J.; YANG, S.; QI, Y. A semi-supervised graph attentive network for financial fraud detection. In: IEEE. *2019 IEEE International Conference on Data Mining (ICDM)*. [S.l.], 2019. p. 598–607.
- WARES, S.; ISAACS, J.; ELYAN, E. Data stream mining: methods and challenges for handling concept drift. *SN Applied Sciences*, Springer, v. 1, n. 11, p. 1–19, 2019.
- WU, J.; REHG, J. M. Beyond the euclidean distance: Creating effective visual codebooks using the histogram intersection kernel. In: IEEE. *2009 IEEE 12th International Conference on Computer Vision*. [S.l.], 2009. p. 630–637.
- YAROWSKY, D. Unsupervised word sense disambiguation rivaling supervised methods. In: *33rd annual meeting of the association for computational linguistics*. [S.l.: s.n.], 1995. p. 189–196.
- ZHANG, B.; WANG, Y.; HOU, W.; WU, H.; WANG, J.; OKUMURA, M.; SHINOZAKI, T. Flexmatch: Boosting semi-supervised learning with curriculum pseudo labeling. *Advances in Neural Information Processing Systems*, v. 34, p. 18408–18419, 2021.
- ZHOU, Z.-H.; LI, M. et al. Semi-supervised regression with co-training. In: *IJCAI*. [S.l.: s.n.], 2005. v. 5, p. 908–913.
- ZHU, X. Stream data mining repository. URL: <http://www.cse.fau.edu/~xqzhu/stream.html>, 2010.