

Apreçamento de debêntures ilíquidas combinando técnicas de aprendizado não supervisionado e supervisionado

Pricing of illiquid debentures using supervised and unsupervised machine learning techniques

Marcela Sousa Zuppini[†]
Afonso de Campos Pinto[‡]
Élia Yathie Matsumoto^{*}

Resumo A marcação a mercado de ativos ilíquidos é um desafio, dada a escassez de informações e negociações no mercado que possam indicar qual deve ser o seu preço justo. No caso de debêntures ilíquidas, uma das dificuldades na determinação do *spread* apropriado para esses ativos é a falta de valores de negociações anteriores. Este trabalho busca determinar o *spread* de debêntures ilíquidas com base nas suas características, nas informações sobre a saúde financeira dos emissores e na situação do mercado. O método apresentado utiliza técnicas de aprendizado não supervisionado e supervisionado, aliando suas respectivas características para a solução do problema de apreçamento de debêntures ilíquidas, podendo ser aplicada para outros tipos de ativos ilíquidos. Nos experimentos, a metodologia proposta se mostrou efetiva para apreçar debêntures ilíquidas sem valor anterior de *spread*.

Palavras-chave: Debêntures; Ativos ilíquidos; Aprendizado não supervisionado; Aprendizado supervisionado.

Código JEL: C02, C45, G12.

Abstract Marking illiquid assets to market is challenging due to the scarcity of information that could indicate their fair price. In the case of illiquid debentures, one of the hindrances to figuring out the proper spread for these kinds of assets is the lack of values from prior trading. This paper proposes to assess the spread of illiquid debentures based on their characteristics, information on the issuers' financial health, and market conditions. The presented method uses unsupervised and supervised learning techniques, combining their respective characteristics to solve the illiquid debentures pricing problem, with applications extending to other types of illiquid assets. In experiments, the proposed methodology proved effective for pricing illiquid debentures with no previous spread value.

Keywords: Debentures; Illiquid assets; Unsupervised learning; Supervised learning.

JEL Code: C02, C45, G12.

Submitted on July 13, 2020. Revised on May 26, 2021. Accepted on September 14, 2021. Published online in December 2021. Editor in charge: Marcelo Fernandes.

[†]Sao Paulo School of Economics – FGV, Brazil: marcela.zuppini@gmail.com.

[‡]Sao Paulo School of Economics – FGV, Brazil: afonso.pinto@fgv.br.

^{*}Sao Paulo School of Economics – FGV, Brazil: elia.matsumoto@fgv.br.

1. Introdução

A determinação do preço justo de ativos é de extrema importância em diversos contextos. Investidores buscam determinar o preço justo para encontrar oportunidades de maximizar seus resultados em transações no mercado. Para uma instituição financeira, apreçar corretamente seus fluxos de caixa para mensuração de obrigações é primordial para evitar problemas de liquidez e planejar investimentos futuros. Já um administrador de recursos de terceiros tem obrigação legal perante os órgãos reguladores de determinar o preço justo dos ativos que compõem sua carteira, de maneira a refletir para o investidor a quantia que seria recebida caso ele vendesse seus ativos, ou o valor que seria pago para transferir seu passivo (CPC, 2012).

Convencionalmente, o preço justo de ativos é dado pelo MtM (*mark-to-market*) ou marcação a mercado. Considerando-se ativos de renda fixa, a marcação a mercado consiste em descontar o valor esperado dos fluxos futuros de um ativo por uma taxa de mercado. Essa taxa reflete o custo de oportunidade dado pela taxa livre de risco e mais um *spread* de mercado, que traz embutido em si o prêmio pelo risco de crédito do emissor, o risco sistêmico inerente ao mercado, os impostos que serão cobrados sobre o lucro (Elton et al., 2001) e o risco de liquidez (Elton e Green, 1998). Esse *spread* de mercado pode ser calculado a partir de preços observados em um mercado líquido. Há ainda a possibilidade de obter junto à corretoras *spreads* para compra (*bid*) e venda (*ask*) de ativos. Existem também instituições que divulgam diariamente as taxas operadas no mercado, como a ANBIMA (Associação Brasileira das Entidades dos Mercados Financeiro e de Capitais) e a bolsa de valores B3. A ANBIMA também divulga taxas indicativas para debêntures. Para o cálculo dessa taxa indicativa ela considera negócios registrados no sistema REUNE, *calls* de corretoras e coleta de taxas junto aos *players* do mercado. As informações são filtradas por meio de critérios de seleção e relevância antes de serem consolidadas e divulgadas (ANBIMA, 2017). Isso significa que, mesmo que não haja negociações de uma determinada debênture no dia, ainda haverá uma taxa indicativa que será construída de acordo com a metodologia estabelecida pela ANBIMA e divulgada para o mercado. Porém, quando se tratam de ativos ilíquidos, não existem negócios no mercado de onde se possam capturar os *spreads*. As corretoras não divulgam ofertas de compra e venda. A ANBIMA não divulga taxa indicativa para diversas debêntures. Com isso, realizar a marcação a mercado de ativos ilíquidos se torna uma tarefa um pouco mais difícil.

Existem algumas maneiras de contornar esse problema. Damodaran (2005) destaca algumas maneiras de abordar essa tema, como por exemplo, determi-

nar o desconto ou *spread* de liquidez relacionando esse valor aos custos de transação e ao *holding period* do investidor. Outra maneira é ajustar os fatores de desconto de liquidez usando *proxies* como *bid-ask spread*, taxas de retorno e por meio da iliquidez sistêmica, ou, em outras palavras, pela demanda geral do mercado por liquidez, o que quer dizer que os descontos serão proporcionais à valorização da liquidez. Uma terceira maneira é considerar observações empíricas no mercado, ou seja, analisar para diversos vencimentos e diversos tipos de emissores, qual a diferença entre os preços dos ativos líquidos e ilíquidos. Estas abordagens exigem um mercado evoluído e líquido de onde se possam obter informações como o custo médio das transações de determinado tipo de ativo, a demanda geral por liquidez e os preços negociados para ativos líquidos e ilíquidos.

Uma outra estratégia é considerar o valor da iliquidez como uma opção, onde se adota a visão de que a iliquidez pode ser medida pelo tempo necessário até que se consiga operar um ativo (Koziol e Sauerbier, 2002). Nesta abordagem, compara-se o preço de um ativo que pode ser liquidado a qualquer momento no intervalo $[0, T]$, com o preço de um ativo ilíquido que só pode ser operado em um conjunto limitado de momentos, i.e., $0, t_1, t_2, \dots, t_N = T$. Assume-se que no vencimento T , o valor de face será pago em ambos os casos, ou seja, não há risco de crédito. Considera-se um investidor com *market timing* perfeito, ou seja, que sempre irá liquidar o ativo no momento ótimo, tanto no caso do ativo líquido quanto no ilíquido. A diferença entre o preço do ativo líquido e o ilíquido replica o *payoff* de uma *put lookback* e resulta no valor da iliquidez. Neste caso, o valor da iliquidez é superestimado devido à suposição do *market timing* perfeito. Outra desvantagem é que o risco de crédito é desconsiderado. Além disso, o modelo é desenvolvido para um *bond* que não paga cupons. Considerar ativos com pagamentos intermediários de juros torna a modelagem mais complicada.

Uma forma mais simples para determinar o *spread* de ativos ilíquidos é considerar o *spread* de ativos líquidos do mesmo emissor ou as taxas negociadas em *credit default swaps* do emissor (Wernert, 2009). Uma outra alternativa seria considerar o *spread* de ativos semelhantes ao ilíquido, selecionados de acordo com alguns critérios como *rating* semelhante, *duration* próxima e mesmo setor econômico (Itau, 2017). Essas abordagens estão limitadas à existência de tais ativos.

É possível ainda elaborar modelos proprietários baseados em *ratings* (Mercantil do Brasil, 2016) ou discussões conduzidas por comitês (Patria, 2016). Nestes casos, o estudo se torna subjetivo.

Para apreçamento de ativos corporativos no Brasil, foi proposto utilizar o

modelo de [Merton \(1974\)](#) incluindo o risco de liquidez ([Secches, 2006](#)). O modelo não se mostrou eficaz, mas comprovou que o risco de liquidez é um fator importante na definição dos preços de debêntures.

Em geral, as abordagens citadas até aqui exigem um estudo caso a caso, dificultando assim sua utilização em larga escala. Portanto, é desejável buscar-se algoritmos computacionais que permitam a análise em massa de diferentes ativos, como ferramentas de inteligência computacional e estatística, entre elas, redes neurais e *clustering*, que foram os métodos escolhidos para serem explorados neste trabalho.

Há uma quantidade grande de estudos envolvendo redes neurais em finanças, descrevendo suas aplicações na solução de um amplo espectro de problemas de previsão ([Aldridge e Avellaneda, 2019](#)). A ferramenta já foi utilizada para previsão de *ratings* de títulos corporativos, apresentando bons resultados ([Dutta e Shekhar, 1988](#)). Utilizou-se redes neurais para estimar a saúde financeira de empresas ([Coats e Fant, 1993](#)). Cinco tipos de redes neurais foram testadas na construção de modelos de *credit score*, e mostrou-se que o uso das redes pode apresentar melhorias na precisão do *credit score* na ordem de 0,5% a 3% ([West, 2000](#)). Na previsão de taxas de câmbio, foram testados o passeio aleatório, modelos lineares e modelos não lineares, dentre eles redes neurais, obtendo resultados que indicam que os modelos lineares e não lineares, dependendo da natureza das séries, são melhores do que o passeio aleatório ([Medeiros et al., 2001](#)).

Mais especificamente no apreçamento de títulos, as redes neurais também têm sido utilizadas. Um conjunto de redes neurais foi utilizado para prever a curva *yield* de *bonds* do governo indiano, obtendo resultados satisfatórios, com erros abaixo de 1% entre o valor esperado e o observado ([Parekh et al., 2013](#)). Vários modelos, dentre eles redes neurais, foram utilizados para prever o preço futuro de *bonds* americanos, partindo de preços históricos, características dos títulos, como cupom pago e prazo até o vencimento, além de preços e volumes dos últimos negócios ([Ganguli e Dunnmon, 2017](#)). Os autores concluem que redes neurais e o método de mínimos quadrados generalizados fornecem os melhores resultados.

No mercado brasileiro, de acordo com o levantamento bibliográfico realizado por [Confessor e dos Santos \(2020\)](#), os estudos científicos apontam dois grandes grupos de fatores que afetam a valorização de debêntures: o primeiro composto por fatores relacionados às características da empresa emissora tais como *rating*, setor e desempenho financeiro, e o segundo relacionado aos aspectos da emissão do instrumento, como volume, vencimento, banco coordenador, garantias, entre outros.

O trabalho de [Curi \(2008\)](#) trata do uso de redes neurais para o apreçamento de debêntures, no qual são utilizados o volume de emissão de cada debênture e algumas características das empresas emissoras como *inputs* da rede. As taxas indicativas divulgadas pela ANBIMA são usadas como *benchmark*. O modelo proposto apresentou um erro grande, mas ainda assim resultados superiores ao modelo linear de mínimos quadrados ordinários. No ano em que esse trabalho foi feito, a quantidade de debêntures disponíveis e negociadas era bem pequena, o que resultou numa base de dados insuficiente para execução do treinamento ideal da rede neural. Assim, há espaço para evoluções nessa linha de estudo.

Este trabalho busca obter previsões dos *spreads* de debêntures, dado que atualmente há mais dados disponíveis para realização do estudo. Após realizar algumas verificações que serão detalhadas adiante, observou-se que as redes neurais não poderiam ser usadas para previsão direta da *spread*, mas sim para prever a primeira diferença dos *spreads*. Dessa forma, tendo apenas a diferença da variável desejada, foi necessário buscar outra forma de determinar o seu nível para, então, somar ao nível a primeira diferença e, assim, obter uma previsão da *spread*. A ferramenta escolhida foi a técnica de *clustering*. Assim, o objetivo deste trabalho é utilizar ferramentas de inteligência computacional, redes neurais e *clustering*, para a determinação dos *spreads* de ativos e analisar o quão próximos os resultados produzidos ficam dos valores esperados. O estudo foi feito com debêntures indexadas ao IPCA que possuem taxa indicativa divulgada pela ANBIMA, e cujas empresas emissoras divulgam informações de balanço.

Além da introdução, este artigo é composto por mais 4 seções. A metodologia proposta está descrita na Seção 2. Na Seção 3, detalhamos os experimentos realizados, desde o tratamento dos dados, passando pela escolha das variáveis, até a construção dos modelos. Na Seção 4, são apresentados os resultados dos experimentos, comparando-se os *spreads* obtidos nas simulações com aqueles negociados de fato no mercado. Por fim, são apresentadas as conclusões com comentários sobre elementos a serem explorados em trabalhos futuros.

2. Metodologia proposta

Como séries temporais de *spreads* podem não ser estacionárias e visto que, de acordo com a teoria de econometria de séries de tempo ([Enders, 2012](#)), a condição de estacionariedade da série dos dados é fundamental para a definição adequada de modelos de previsão, trabalhamos com a previsão da variação do *spread* (primeira diferença) que, neste caso, é uma série estacio-

nária.

Desta forma, nas definições a seguir, consideraremos uma série temporal do *spread* de debênture como não-estacionária, e a série de sua 1ª diferença, ou seja, a série de variação do *spread* de debênture como estacionária, e, a partir deste ponto do texto, considerando o arcabouço de um problema de estimação de valor futuro, denominaremos:

Valor futuro do *spread*, S_f : o valor a ser estimado no instante futuro, t ;

Valor anterior do *spread*, S_a : o valor do *spread* no período anterior ao instante futuro, t ;

Variação futura do valor do *spread*, $Y_f = S_f - S_a$: a diferença entre o valor futuro e o valor anterior do *spread*.

Debênture líquida: uma debênture que foi negociada no período anterior, ou seja, com valor do *spread* existente no período anterior ($\exists S_a$).

Debênture ilíquida: uma debênture não negociada no período anterior, ou seja, sem valor do *spread* existente no período anterior ($\nexists S_a$).

Além disso, a metodologia tem como premissa a disponibilidade de informações sobre os contratos de debêntures do período anterior, X_a , tais como taxa de emissão e volume, tanto no caso das debêntures líquidas quanto das ilíquidas.

No caso das debêntures líquidas, essas informações serão usadas como variáveis de entrada de modelos de aprendizado supervisionado, *Modelo_SP*, treinados para estimar as variações futuras do *spread*, Y_f .

Essas estimativas, $\hat{Y}_f$ de Y_f , somadas aos valores anteriores do *spread*, S_a , produzirão as estimativas dos valores futuros do *spread*, $\hat{S}_f$ de S_f , conforme descrito nas equações 1 e 2:

$$\hat{Y}_f = \text{Modelo_SP}(X_a) \quad (1)$$

$$\hat{S}_f = \hat{Y}_f + S_a \quad (2)$$

Como, no caso de debêntures ilíquidas, os valores anteriores, S_a , não existem, este trabalho propõe um método para produzir uma *proxy* para essa informação, e iniciaremos com a definição de termos e nomenclaturas que utilizaremos para descrever os procedimentos.

2.1 Definição de termos e nomenclaturas

Neste item definimos os termos e nomes que utilizaremos para descrever a metodologia proposta e detalhar os experimentos executados neste trabalho.

Com relação à divisão de tempo, dividimos as observações em dois grupos segundo as datas de referência dos dados das debêntures e definimos dois períodos:

Período de treinamento: fazem parte deste conjunto, as observações de debêntures **líquidas** do segundo trimestre de 2013 até o segundo trimestre de 2017. Essas observações serão utilizadas para compor os conjuntos de dados de trabalho de **treinamento**, conforme descrito a seguir.

Período de teste (deve ser posterior ao Período de treinamento): fazem parte deste conjunto, as observações de debêntures **líquidas** e **ilíquidas** do terceiro trimestre de 2017. Essas observações serão usadas para compor os conjuntos de dados de trabalho de **teste** conforme o descrito a seguir.

Considerando a divisão das observações em período de treinamento e período de teste, definimos dois tipos de conjuntos de trabalho:

Conjunto de treinamento: é um sub-conjunto (amostra aleatória) do conjunto do período de **treinamento**, sendo, portanto, composto apenas por debêntures **líquidas**. Essas observações serão usadas no treinamento dos modelos.

Conjunto de teste: é um sub-conjunto (amostra aleatória) do conjunto do período de **teste**, composto por debêntures **líquidas** e **ilíquidas**. Essas observações serão usadas para verificar o desempenho dos modelos.

Nas Tabelas 1 e 2, listamos os nomes utilizados para identificar os conjuntos de **treinamento** e de **teste**, respectivamente:

Tabela 1
Identificação dos componentes do conjunto de treinamento

nome	descrição
STr_a	valores do <i>spread</i> de debêntures líquidas no instante anterior
XTr_a	variáveis explicativas de debêntures líquidas no instante anterior
YTr_a	variações do <i>spread</i> de debêntures líquidas no instante anterior

Na Tabela 3, listamos os nomes dos conjuntos para os quais produziremos estimativas. Para as debêntures ilíquidas, ou seja, no caso em que não existe valor anterior do *spread*, STe_a , propomos gerar as estimativas para essa informação, $i\hat{S}_a$, usando um modelo de aprendizado não supervisionado, *Modelo_NS*, conforme o descrito a seguir no item 2.2.

Tabela 2
Identificação dos componentes do conjunto de teste

nome	descrição
STe_a	valores do <i>spread</i> de debêntures líquidas no instante anterior
XTe_a	variáveis explicativas de debêntures líquidas no instante anterior
YTe_f	variações futuras do <i>spread</i> de debêntures líquidas
STe_f	valores futuros do <i>spread</i> de debêntures líquidas
$iSTe_a$	valores do <i>spread</i> de debêntures ilíquidas no instante anterior
$iXTe_a$	variáveis explicativas de debêntures ilíquidas no instante anterior
$iYTe_f$	variações futuras do <i>spread</i> de debêntures ilíquidas
$iSTe_f$	valores futuros do <i>spread</i> de debêntures ilíquidas

Tabela 3
Identificação dos nomes das
estimativas para o conjunto de teste

nome	descrição
$\hat{Y}_f$	estimativas de YTe_f
$\hat{S}_f$	estimativas de STe_f
$i\hat{S}_a$	estimativas de $iSTe_a$
$i\hat{Y}_f$	estimativas de $iYTe_f$
$i\hat{S}_f$	estimativas de $iSTe_f$

2.2 Proxy do valor anterior do *spread* de debêntures ilíquidas

Para produzir uma *proxy* do valor anterior do *spread* das debêntures ilíquidas do conjunto de teste, propomos criar um conjunto, $XInfo$, formado pela união das variáveis explicativas das debêntures **líquidas** do conjunto de **treinamento**, XTr_a , e das debêntures **ilíquidas** do conjunto de **teste**, $iXTe_a$, conforme o descrito na equação 3.

$$XInfo = \{ XTr_a \cup iXTe_a \} \quad (3)$$

Utilizamos esse conjunto, $XInfo$, para treinar um modelo, $Modelo_{NS}$, usando técnicas de método de aprendizado não supervisionado, para criar K agrupamentos: $(G^1, G^2, \dots, G^K) = Modelo_{NS}(XInfo)$.

Para produzir a estimativa do valor anterior do *spread* de cada debênture **ilíquida** k do conjunto de teste, seguimos os seguintes passos:

1. Identificamos j , o índice do grupo, G^j , no qual a debênture **ilíquida** k foi alocada.

2. Identificamos todas as debêntures líquidas alocadas nesse mesmo grupo G^j .
3. Calculamos a média dos valores anteriores dos *spreads* dessas debêntures **líquidas** alocadas no grupo G^j .

Propomos usar o valor dessa média como a estimativa do valor anterior do *spread* da debênture **ilíquida** k , conforme o descrito na equação 4

$$i\hat{S}_{k(a)} = \text{média}(STr_a \in G^j) \quad (4)$$

De acordo com Celebi e Aydin (2016), uma das principais aplicações de métodos de aprendizado não supervisionado é dividir uma amostra em grupos por meio da identificação de similaridades entre as observações que a compõem.

Essa é a ideia que fundamenta a escolha da média dos valores do *spread* das debêntures **líquidas** do conjunto de teste como estimativa do valor anterior do *spread* das debêntures **ilíquidas** do conjunto de teste alocadas no mesmo grupo, uma vez que, se alguma similaridade for identificada entre as informações dos dois tipos de debêntures, então é possível existir similaridade entre os valores de seus *spreads*.

2.3 Composição da previsão do valor futuro do *spread*

Conforme já mencionado, propomos utilizar técnicas de aprendizado supervisionado e treinar um modelo, *Modelo_SP*, usando XTr_a como entrada e YTr_f como saída.

Desta forma, o *Modelo_SP* servirá para estimar as variações do *spread* de debêntures líquidas e ilíquidas do conjunto de teste:

$$\begin{aligned}\hat{Y}_f &= \text{Modelo_SP}(XTe_a) \\ i\hat{Y}_f &= \text{Modelo_SP}(iXTe_a)\end{aligned}$$

Essas estimativas somadas às informações sobre os valores anteriores do *spread* formarão a previsão dos valores futuros do *spread*:

Para debêntures **líquidas**, o valor somado à variação é o valor anterior do *spread* existente, STe_a :

$$\hat{S}_f = STe_a + \hat{Y}_f$$

Para debêntures **ilíquidas**, o valor somado à variação é a estimativa produzida pelo procedimento de agrupamento, $i\hat{S}_a$, conforme o descrito anteriormente na equação 4:

$$i\hat{S}_f = i\hat{S}_a + i\hat{Y}_f$$

2.4 Comparação dos resultados

Os procedimentos descrito nos itens anteriores geram quatro tipos de estimativas para os valores futuros do *spread*, duas para debêntures líquidas e duas para ilíquidas:

Debêntures líquidas:

Anterior: Valor anterior do *spread* disponível (STe_a).

Anterior + Modelo_SP: Valor anterior do *spread* mais a estimativa da variação futura produzida pelo modelo de **aprendizado supervisionado**, ($STe_a + \hat{Y}_f$).

Debêntures ilíquidas:

Modelo_NS: Estimativa do valor anterior do *spread* produzida pelo modelo de **aprendizado não supervisionado** ($i\hat{S}_a$).

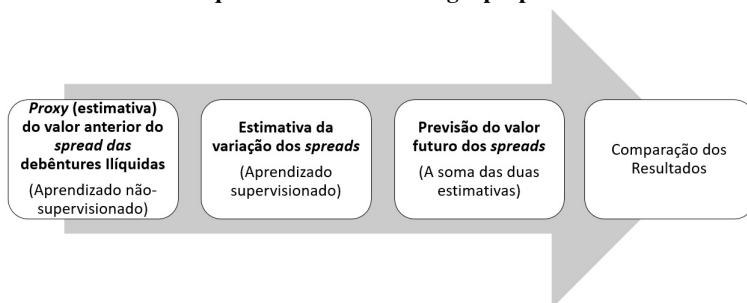
Modelo_NS + Modelo_SP: Estimativa do valor anterior do *spread* produzida pelo modelo de **aprendizado não supervisionado** mais a estimativa da variação futura produzida pelo modelo de **aprendizado supervisionado**, ($i\hat{S}_a + i\hat{Y}_f$).

De acordo com pesquisas sobre previsão de séries de tempo financeiras, o valor anterior da série, quando disponível, é tido como um bom estimador para seu próximo valor (Lahmiri, 2018; Arce et al., 2019). Por este motivo, propomos considerar o valor anterior do *spread* como principal referência de comparação para a metodologia proposta.

Esses quatro tipos de estimativas devem ser comparados com os valores reais do *spread* por meio do cálculo de métricas como raiz do erro quadrático médio (REQM) e coeficiente de determinação (R^2).

Os principais passos da metodologia estão esquematizados na Figura 1.

Figura 1
Esquemático da metodologia proposta



3. Descrição dos experimentos

Nesta seção, detalhamos os experimentos realizados neste trabalho para verificar o desempenho e a aplicabilidade da metodologia descrita na seção anterior. Começamos, na primeira subseção, com o detalhamento do processo de escolha e preparação dos dados. Nas subseções seguintes, detalhamos a implementação da metodologia:

- (i) Para produzir a *proxy* dos valores anteriores do *spread* das debêntures ilíquidas, implementamos o modelo de aprendizado não supervisionado *k-means*.
- (ii) Para estimar os valores futuros do *spread* de debêntures, tanto líquidas quanto ilíquidas, implementamos dois tipos de modelos de aprendizado supervisionado:
 1. Regressão Linear (RL) resolvido pelo método de mínimos quadrados.
 2. Rede Neural Artificial (NN) do tipo MLP (*Multi-Layered Perceptron*).

Destacamos que a metodologia pode ser implementada por meio de outras técnicas de aprendizado não supervisionado e supervisionado, e não somente as utilizadas neste experimento.

3.1 Escolha e preparação dos dados iniciais

Com o objetivo de determinar os valores dos *spreads* no mercado de debêntures, tomamos como base as variáveis sugeridas nos trabalhos de Curi (2008) e Mohanty e Dash (2019), e analisamos dados tanto sobre os ativos,

quanto sobre suas empresas emissoras. Como fonte de dados, usamos informações divulgadas pela B3 e informações de balanço das empresas disponibilizadas pela plataforma Economatica.

Além disso, para tentar capturar as condições do mercado brasileiro, consideramos incluir a taxa livre de risco e o índice EMBI¹ + Risco Brasil, divulgado pelo JP Morgan e disponível no site do Instituto de Pesquisas Econômica Aplicada (IPEA²).

Desta forma, para formar o conjunto de dados de trabalho, além da variável de interesse, S_f , que é o valor do *spread* das debêntures, analisamos a lista de variáveis apresentada a seguir:

1. Volume de emissão;
2. Taxa de emissão;
3. *Duration*: prazo médio remanescente até o vencimento;
4. Garantia (variável lógica): emissões com garantias reais ou fidejussórias recebem valor 1, as demais recebem valor 0;
5. Incentivo (variável lógica): ativos que possuem incentivo fiscal recebem o valor 1, as demais recebem o valor 0;
6. ROIC (*return on invested capital*): lucro líquido sobre o ativo total;
7. ROE (*return on equity*): lucro líquido sobre o patrimônio líquido;
8. Margem líquida: lucro líquido sobre a receita líquida;
9. Liquidez corrente: ativo circulante sobre passivo circulante;
10. Giro do ativo: receita líquida sobre ativo médio total;
11. Composição do endividamento: dívida de curto prazo sobre dívida total;
12. Alavancagem: lucro antes de juros e impostos sobre lucro antes de impostos;
13. Capital de giro sobre ativo: ativo circulante menos passivo circulante sobre ativo total;
14. Capital de giro sobre receita líquida: ativo circulante menos passivo circulante sobre receita líquida;
15. Volatilidade: volatilidade do *spread* observada nos 12 meses anteriores, obtida por meio do desvio-padrão;
16. Taxa livre de risco: curva construída com base nos contratos futuros de DI;
17. EMBI+Risco Brasil: índice que representa o risco soberano de países emergentes.

¹Emerging Markets Bond Index

²<http://www.ipeadata.gov.br/ExibeSerie.aspx?serid=40940&module=M>

Inicialmente, coletamos dados de 117 debêntures de 66 emissores, compreendendo o período de janeiro de 2013 a setembro de 2017, totalizando 2223 observações (117 debêntures para cada um dos 19 trimestres). Desse conjunto inicial, consideramos apenas as debêntures com contratos válidos durante todo o período, ou seja, descartamos 1155 observações referentes a debêntures com vencimento anterior a setembro de 2017 ou com data de emissão posterior a janeiro de 2013, formando uma base final de 1068 observações.

Para preparar os dados desse conjunto de observações executamos os procedimentos de tratamento e análise de dados descritos a seguir.

3.1.1 Verificação de estacionariedade

Para uniformizar as séries de tempos, calculamos as médias trimestrais dos valores do *spread* das 17 variáveis listadas no item anterior e o utilizamos o teste de KPSS (Baum, 2018) com nível de significância de 0,05 para verificar a estacionariedade de cada uma das séries.

No caso da série de valores do *spread*, S , o resultado do teste de KPSS para os valores em nível não rejeitou a hipótese de existência de raiz unitária, mas a série de primeiras diferenças resultou em estacionária.

Com isto, definimos a série de primeiras diferenças, ou seja, de taxas de variação trimestral dos valores futuros do *spread*, $Y_f = S_f - S_a$, como a variável dependente para os modelos de aprendizado supervisionado.

Para as séries das variáveis informativas, aplicamos o teste de KPSS e nos casos de resultados indicando não-estacionariedade, testamos as séries de primeiras diferenças e verificamos que a condição de estacionariedade foi atendida para as seguintes variáveis: 3. *Duration*; 7. ROE; 8. Margem líquida; 11. Composição de endividamento; 15. Volatilidade; 16. Taxa livre de risco.

Por este motivo, para essas 6 variáveis, substituímos suas séries originais em nível pelas séries de primeiras diferenças (de variação). Como consequência, descartamos a primeira observação do trimestre mais antigo, e construímos as observações com o seguinte conjunto de variáveis:

Variável dependente Y_f : Valor da variação do *spread* no trimestre futuro.

Conjunto de variáveis explicativas, $X_a = (X_a^1, \dots, X_a^{17})$ no trimestre anterior, composto por:

X_a^1 : Volume de emissão;	X_a^{10} : Giro do ativo;
X_a^2 : Taxa de emissão;	X_a^{11} : Variação da composição do endividamento;
X_a^3 : Variação da <i>duration</i> ;	X_a^{12} : Alavancagem;
X_a^4 : Garantia;	X_a^{13} : Capital de giro sobre ativo;
X_a^5 : Incentivo;	X_a^{14} : Capital de giro sobre receita líquida;
X_a^6 : ROIC;	X_a^{15} : Variação da volatilidade;
X_a^7 : Variação do ROE;	X_a^{16} : Variação da taxa livre de risco;
X_a^8 : Variação da margem líquida;	X_a^{17} : EMBI+Risco Brasil.
X_a^9 : Liquidez corrente;	

3.1.2 Definição dos períodos de treinamento e de teste

Seguindo o procedimento descrito na Seção 2 (na qual detalhamos a metodologia), definimos:

O período de abril de 2013 a junho de 2017 como **período de treinamento**, composto por 956 debêntures líquidas.

O último trimestre de coleta de dados (julho, agosto e setembro de 2017) como **período de teste**, composto por 112 debêntures, das quais, 104 são debêntures líquidas e 8 ilíquidas.

A Tabela 4 contém informações estatísticas dos valores do *spread* em nível, S_f , e a Tabela 5, dos valores das variações, $Y_f = S_f - S_a$. Nas duas tabelas, as estatísticas foram calculadas por período (treinamento e teste) e por tipo de debênture (líquida e ilíquida).

Tabela 4
Estatísticas por período dos valores do *spread* (S_f)

estatística	treinamento (líquidas)	teste (líquidas)	teste (ilíquidas)
média	7.58%	6.12%	5.28%
desvio padrão	2.22%	1.95%	0.42%
mínimo	4.62%	3.68%	4.69%
máximo	31.39%	13.54%	16.30%

Lembramos que, no caso de debêntures **ilíquidas**, o valor do *spread* anterior não existe, e, consequentemente, o valor de variação do *spread* também não. Por este motivo a Tabela 5 só exibe informações de debêntures **líquidas**.

Com relação às variáveis explicativas, na Tabela 6 exibimos a porcentagem de ocorrência de valor 1 nas variáveis lógicas X_a^4 e X_a^5 , e na Tabela 7, exibimos as médias das variáveis contínuas.

Tabela 5
Estatísticas por período dos
valores de variação do *spread* (Y_f)

estatística	treinamento (líquidas)	teste (líquidas)
média	0.11%	-0.53%
desvio padrão	1.06%	0.72%
mínimo	-1.68%	0.72%
máximo	5.50%	5.84%

Tabela 6
Frequência por período
das variáveis lógicas

variável	treinamento	teste
X_a^4 : Garantia	18.52%	22.12%
X_a^5 : Incentivo	49.19%	62.83%

Tabela 7
Médias por período das variáveis numéricas

variável	treinamento	teste
X_a^1 : Volume de emissão em milhões de reais	3.279	2.754
X_a^2 : Taxa de emissão	0.065	0.067
X_a^3 : Variação da <i>Duration</i>	-51.532	-43.825
X_a^6 : ROIC	0.00086	0.00091
X_a^7 : Variação da ROE	0.166	-1.503
X_a^8 : Variação da Margem líquida	-0.160	0.285
X_a^9 : Liquidez corrente	1.339	1.562
X_a^{10} : Giro do ativo	0.418	0.424
X_a^{11} : Variação da Composição do endividamento	0.773	-0.232
X_a^{12} : Alavancagem	9.666	2.179
X_a^{13} : Capital de giro sobre ativo	0.008	0.014
X_a^{14} : Capital de giro sobre receita líquida	0.095	0.060
X_a^{15} : Variação da Volatilidade	-0.108	0.014
X_a^{16} : Variação da Taxa livre de risco	-0.002	-0.012
X_a^{17} : EMBI+Risco Brasil	324	246

3.1.3 Definição dos conjuntos de treinamento e de teste

Novamente, seguindo o procedimento descrito na Seção 2, a partir dos conjuntos de período de treinamento e de teste, produzimos 100 conjuntos de trabalho de treinamento e de teste por reamostragem:

- Conjunto de **treinamento**: conjunto formado pela reamostragem de 80% das observações pertencentes ao **período de treinamento**. Este conjunto contém apenas debêntures líquidas que serão usadas para o treinamento dos modelos, (YTr_f, XTr_a) .
- Conjunto de **teste**: conjunto formado pela reamostragem de 80% das observações pertencentes ao **período de teste**. Este conjunto contém debêntures líquidas e ilíquidas que serão usadas para verificar o desempenho dos modelos e da metodologia: (YTe_f, XTe_a) e $(iYTe_f, iXTe_a)$.

3.1.4 Tratamento dos dados

Para os 100 conjuntos de treinamento e teste, aplicamos dois procedimentos de tratamento de dados:

1. Limpeza dos dados: com relação *missing data*, adotamos o procedimento de descarte das observações incompletas.
2. Normalização dos dados: com relação a amplitude dos dados, utilizamos a função *Z-Score* para normalizar as variáveis do conjunto de treinamento e utilizamos os valores de média e desvio padrão obtidos nesse processo para normalizar as variáveis do conjunto de teste.

Após a execução dos procedimentos citados, a partir das 956 observações do período de treinamento e das 112 do período de teste, geramos 100 conjuntos de trabalho de treinamento e de teste que serão usados para o treinamento dos modelos. A média e o desvio padrão (DP) do tamanho dos 100 conjuntos gerados estão listados na Tabela 8.

Tabela 8
Número de observações dos conjuntos de trabalho

estatística	treinamento(líquidas)	teste(líquidas)	teste(ilíquidas)
média ± dp	417 ± 6	48 ± 2	7 ± 1

3.2 Estimação da *proxy* do valor anterior do *spread* de debêntures ilíquidas

Para implementar a estratégia proposta de estimação de uma *proxy* do valor anterior do *spread* de debêntures ilíquidas, adotamos o modelo *k-means* com distância euclidiana, um dos modelos de aprendizado não supervisionado mais populares na área de *Machine Learning*.

A popularidade do *k-means* pode estar associada a existência de métodos heurísticos por meio dos quais é possível definir o número ótimo de agrupamentos tendo como base as informações das observações que estão sendo agrupadas.

Nos experimentos, aplicamos o método *Elbow* de definição de número de agrupamentos. Este método está baseado na ideia de aumentar o número de agrupamentos até que o aumento do número de grupos não produza agrupamentos significativamente distintos (Syakur et al., 2018).

Nas 100 execuções realizadas para os 100 conjuntos de trabalho, e para cada conjunto de trabalho:

1. Treinamos um modelo *k-means* (*Modelo_KMeans*), juntamente com o método *Elbow* de definição do número ótimo de agrupamentos, K , para agrupar o conjunto $XInfo = \{XTr_a \cup iXTe_a\}$, formado pela união das informações disponíveis das debêntures **líquidas** do conjunto de **treinamento** e das debêntures **ilíquidas** do conjunto de **teste**.
2. Para cada agrupamento estimado,

$$(G^1, \dots, G^K) = \text{Modelo_KMeans}(XInfo),$$

calculamos a média dos valores dos *spreads* das debêntures líquidas do grupo, e usamos o valor dessa média como *proxy* do valor do *spread* das debêntures ilíquidas do grupo, $i\hat{S}_a$.

Os valores de *proxy* produzidos neste procedimento, $i\hat{S}_a$, foram usados para compor a estimativa do valor futuro do *spread* para as debêntures ilíquidas dos conjuntos de teste conforme o descrito no próximo item.

3.3 Composição da previsão do valor futuro do *spread*

Seguindo o procedimento descrito na metodologia, usamos as informações sobre debêntures líquidas dos conjuntos de treinamento, (YTr_f, XTr_a) , para treinar dois tipos de modelos de aprendizado supervisionado, e gerar as estimativas dos valores de variação do *spread* das debêntures (líquidas e ilíquidas) dos conjuntos de teste:

1. Regressão Linear (RL) com resolução pelo método dos mínimos quadrados.

$$RL_{\widehat{Y}Te_f} = Modelo_RL(XTe_a)$$

$$RL_{i\widehat{Y}Te_f} = Modelo_RL(iXTe_a)$$

2. Redes Neurais Artificiais (*Neural Network*, NN) do tipo MLP (*Multi-layered Perceptron*).

$$NN_{\widehat{Y}Te_f} = Modelo_NN(XTe_a)$$

$$NN_{i\widehat{Y}Te_f} = Modelo_NN(iXTe_a)$$

O modelo RL foi resolvido pelo método dos mínimos quadrados. Escolhemos os modelos RL porque esse tipo de modelo fornece instrumental para análise de correlação do tipo *ceteris paribus* entre as variáveis explicativas e a variável estimada. Além disso, os resultados dos modelos RL serviram como base de comparação para os resultados gerados pelos modelos NN.

No caso dos modelos NN, treinamos os modelos segundo os procedimentos de melhores práticas recomendados pela literatura técnica (Samarasinghe, 2007), usando a seguinte configuração básica:

1. Três camadas: entrada, oculta e saída.
2. Número de neurônios por camada:
 - 2.1 Nas camadas de entrada e oculta: com número de neurônios igual ao número de variáveis de entrada.
 - 2.2 Na última camada: um neurônio.
3. Algoritmo de treinamento: *Backpropagation Levenberg-Marquardt* (Sapna et al., 2012).

Na próxima seção, apresentamos os resultados numéricos produzidos pela execução dos procedimentos descritos.

4. Resultados numéricos e discussão

Os procedimentos descritos na seção anterior foram repetidos 100 vezes usando os 100 conjuntos de trabalhos gerados por meio de reamostragem das observações contidas nos períodos de **treinamento** e de **teste**. Os números apresentados nesta seção sintetizam esses resultados.

4.1 Variáveis selecionadas

Usamos o conjunto inicial de informações sobre as debêntures composto pelas 17 variáveis explicativas apresentadas na Seção 3.1 como base para o procedimento de seleção de variáveis

4.1.1 Modelos *k*-means

Seguindo os procedimentos descritos na metodologia, para definição do conjunto, *XInfo*, escolhemos as variáveis explicativas que estavam presentes em todas as observações das debêntures ilíquidas do conjunto de testes.

Nas 100 execuções, das 17 variáveis iniciais, 5 foram escolhidas 100 vezes e 6 foram escolhidas 2. Como as debêntures ilíquidas não possuem valor anterior, as variáveis não selecionadas foram as 6 para as quais foi necessário calcular a primeira diferença.

Considerando as 100 execuções, a Tabela 9 mostra os valores de média e desvio padrão (DP) do número de: (i) Grupos produzidos (#grupos); (ii) Variáveis de entrada por grupo (#var por grupo); (iii) Observações (debêntures) em cada grupo (#obs por grupo).

Tabela 9
***k*-means: características dos agrupamentos**

estatística	#grupos	#var por grupo	#obs por grupo
média $\pm$ dp	26 $\pm$ 6	5 $\pm$ 1	17 $\pm$ 6

A Tabela 10 lista a frequência das variáveis escolhidas nas 100 execuções.

Os conjuntos composto por essas 11 variáveis de entrada selecionadas pelos procedimentos de agrupamento definiram os conjuntos de variáveis iniciais para os modelos supervisionados.

4.1.2 Modelos de regressão linear (RL)

No caso dos modelos RL, utilizamos a análise de coeficientes usando significância de 5% para selecionar as variáveis de entrada.

Como resultado, nas 100 execuções, os modelos RL resultaram em modelos compostos por 3 variáveis de entrada, em média, e com desvio padrão de 1. No total, das 11 variáveis iniciais, 7 foram selecionadas. Na Tabela 11 listamos a frequência das variáveis selecionadas, juntamente com os valores das médias de seus coeficientes e de seus *p*-valores.

Tabela 10
***k*-means: frequência das variáveis escolhidas**

variável	frequência
X_a^1 : Volume de emissão	100
X_a^2 : Taxa de emissão	100
X_a^4 : Garantia	100
X_a^5 : Incentivo	100
X_a^6 : ROIC	2
X_a^9 : Liquidez corrente	2
X_a^{10} : Giro do ativo	2
X_a^{12} : Alavancagem	2
X_a^{13} : Capital de giro sobre ativo	2
X_a^{14} : Capital de giro sobre receita líquida	2
X_a^{17} : EMBI+Risco Brasil	100

Tabela 11
RL: frequência das variáveis escolhidas

variável	freq.	coef. (média)	<i>p</i> -valor (média)
X_a^1 : Volume de emissão	17	0,047	0,016
X_a^2 : Taxa de emissão	2	-0,066	0,008
X_a^4 : Garantia	94	0,150	0,004
X_a^5 : Incentivo	95	-0,060	0,010
X_a^{13} : Cap.giro sobre ativo	1	0,067	0,086
X_a^{14} : Cap.giro sobre receita líq.	1	-0,066	0,031
X_a^{17} : EMBI+Risco Brasil	100	0,130	0,000

Notamos uma relação positiva entre o volume de emissão e o *spread*, o que faz sentido, pois quanto maior o volume da dívida, maior é a chance de a empresa emissora não honrar o compromisso, o que leva a um *spread* mais alto. A mesma lógica se aplica ao EMBI+Risco Brasil, pois quanto maior a percepção de risco no país, maior a chance de instabilidade financeira, o que causa a elevação dos *spreads*. Notamos ainda relação positiva entre a garantia e o *spread*, o que é contra intuitivo, pois espera-se que um papel que possua garantias represente menos risco para o seu detentor, o que deveria ser refletido por um *spread* menor. No entanto, [John et al. \(2003\)](#) fazem um estudo do tema e encontram *spreads* maiores em emissões com garantias. Conforme discutido por [Sheng e Saito \(2005\)](#), no caso de insolvência da empresa, é pouco provável que os credores recebem essas garantias devido à hierarquia de pagamento prevista na Lei de Falência. Além disso, podemos conjecturar também que empresas que precisam dar garantias para conseguir emitir papéis, tendem a ser mais arriscadas.

Há uma relação negativa entre o incentivo fiscal e o *spread*, ou seja, papéis incentivados apresentam *spreads* menores, pois o detentor não paga imposto sobre o rendimento e, para ter essa vantagem, aceita papéis com taxas menores. A relação negativa também se observa no capital de giro sobre a receita líquida, pois quanto maior o valor desse indicador, maior a liquidez da empresa e menor deve ser o *spread*.

A relação negativa entre taxa de emissão e *spread* não tem sentido econômico, dado que, em geral, espera-se que quanto maior a taxa de emissão, maior seja o *spread*. Também não há sentido na relação positiva entre capital de giro sobre ativo e *spread*, pois quanto maior o indicador, menor é a possibilidade de problemas de solvência, o que deveria resultar em *spreads* menores. Embora seja difícil encontrar explicações que façam sentido do ponto de vista econômico-financeiro, acatamos a seleção dessas variáveis como relevantes na estimação dos *spreads*. Quatro variáveis não foram selecionadas: (i) X_a^6 : ROIC; (ii) X_a^9 : Liquidez corrente; (iii) X_a^{10} : Giro do ativo; (iv) X_a^{12} : Alavancagem.

4.1.3 Modelos de redes neurais artificiais (NN)

No caso dos modelos NN, para selecionar as variáveis de entrada, aplicamos o método “*Neighborhood Component Analysis*” (Yang et al., 2012).

Como resultado, nas 100 execuções, os modelos NN resultaram em modelos compostos por 5 variáveis de entrada, em média, e com desvio padrão de 1. Todas as 11 variáveis iniciais foram selecionadas. Na Tabela 12 listamos a frequência das variáveis selecionadas.

Tabela 12
NN: frequência das variáveis escolhidas

variável	frequência
X_a^1 : Volume de emissão	100
X_a^2 : Taxa de emissão	100
X_a^4 : Status de garantia	100
X_a^5 : Status de incentivo	100
X_a^6 : ROIC	2
X_a^9 : Liquidez corrente	2
X_a^{10} : Giro do ativo	2
X_a^{12} : Alavancagem	2
X_a^{13} : Capital de giro sobre ativo	2
X_a^{14} : Capital de giro sobre receita líquida	2
X_a^{17} : EMBI+Risco Brasil	100

Destacamos que o fato de os números nas Tabelas 12 e 10 resultarem

exatamente iguais foi apenas uma coincidência.

4.2 Comparação das estimativas produzidas

Para avaliar as estimativas produzidas pelos experimentos, utilizamos duas métricas de comparação: 1. Coeficiente de determinação (R^2) e; 2. Raiz da média do erro quadrado ($RMEQ$). A Tabela 13 mostra a média e desvio padrão das métricas produzidas pelas estimações geradas pelos modelos RL (RL_YTe_f), e NN (NN_YTe_f), para as variações do *spread* das debêntures líquidas dos conjuntos de teste.

Tabela 13
Variação do *spread* debêntures
líquidas do conjunto de teste

métricas	modelos RL	modelos NN
R^2	$10,6\% \pm 9,3\%$	$13,1\% \pm 11,0\%$
$RMEQ$	$0,087 \pm 0,024$	$0,236 \pm 0,077$

No caso de debêntures líquidas, para obter as estimativas dos valores futuros do *spread*, somamos os valores estimados pelos modelos RL e NN aos valores anteriores do *spread* (STe_a).

Lembramos que, para debêntures ilíquidas, os valores anteriores do *spread* não existem, por este motivo: (i) As métricas de comparação R^2 e $RMEQ$ não existem para as estimativas de variações do *spread*; (ii) Para obter as estimativas dos valores futuros do *spread*, somamos os valores estimados pelos modelos RL e NN aos valores estimados pelo método de aprendizado não supervisionado ($i\hat{S}_a$).

Conforme já mencionado, no caso de séries de tempo financeiras, o último valor disponível da série é tido como um preditor razoável do próximo valor. Assim, também incluímos as métricas produzidas pelas séries de valor anterior do *spread* nas Tabelas 14 e 15, que mostram o resumo dos resultados obtidos para as debêntures líquidas dos conjuntos de teste.

Tabela 14
Debêntures líquidas do conjunto de teste – descrição

série	descrição
STe_a	Valor anterior de <i>spread</i>
$STe_a + RL_YTe_f$	Valor anterior de <i>spread</i> + Estimativa da RL
$STe_a + NN_YTe_f$	Valor anterior de <i>spread</i> + Estimativa da NN

Tabela 15
Debêntures líquidas do conjunto de teste – métricas

série	R^2	$RMEQ$
STe_a	$85,7\% \pm 3,5\%$	$0,398 \pm 0,058$
$STe_a + RL_YTe_f$	$85,1\% \pm 3,2\%$	$0,430 \pm 0,053$
$STe_a + NN_YTe_f$	$83,9\% \pm 4,1\%$	$0,486 \pm 0,076$

De acordo com as métricas obtidas, para debêntures líquidas, a melhor estimativa para os valores futuros do *spread* é dada pela série dos valores anteriores do *spread*.

No caso das debêntures **ilíquidas**, usamos as estimativas geradas pelo método de aprendizado não supervisionado no lugar do valor anterior do *spread*. As Tabelas 16 e 17 mostram o resumo dos resultados obtidos para as debêntures **ilíquidas** dos conjuntos de **teste**.

Tabela 16
Debêntures ilíquidas do conjunto de teste – descrição

série	descrição
$i\hat{S}_a$	Estimativa do método não supervisionado (MNS)
$i\hat{S}_a + RL_YTe_f$	Estimativa do MNS + Estimativa da RL
$i\hat{S}_a + NN_YTe_f$	Estimativa do MNS + Estimativa da NN

Tabela 17
Debêntures ilíquidas do conjunto de teste – métricas

série	R^2	$RMEQ$
$i\hat{S}_a$	$58,6\% \pm 25,3\%$	$0,392 \pm 0,161$
$i\hat{S}_a + RL_YTe_f$	$72,1\% \pm 15,9\%$	$0,360 \pm 0,126$
$i\hat{S}_a + NN_YTe_f$	$80,0\% \pm 11,1\%$	$0,377 \pm 0,113$

Nos experimentos, o método de aprendizado não supervisionado (*k*-means), isoladamente, produziu estimativas com R^2 de 58,6%. E, quando somado às estimativas geradas pelos métodos de aprendizado supervisionado, RL e NN, os valores da métrica R^2 subiram, respectivamente, para 72,1% e 80,0%.

Os experimentos evidenciaram que, apesar de a melhor estimativa para o valor futuro do *spread* ser o valor anterior do *spread*, na falta dessa informação, a metodologia proposta, de aliar um modelo de aprendizado não supervisionado a um modelo de aprendizado supervisionado, se mostrou uma boa alternativa, com alto potencial para ser aplicada na prática.

5. Conclusões

O objetivo desse trabalho foi investigar alternativas para estimar o valor de *spread* de debêntures ilíquidas, um tema bastante relevante e um problema real no mercado brasileiro pois, considerando o período de julho de 2017 a julho de 2018, de 1329 debêntures registradas, apenas 592 haviam tido negócios, ou seja, menos de 45% do total de debêntures em circulação.

Para atingir esse objetivo, propomos uma metodologia que combina técnicas de aprendizado não supervisionado e supervisionado para gerar essas estimativas.

Mais especificamente, propomos a aplicação de aprendizado não supervisionado para criar agrupamentos considerando papéis líquidos e ilíquidos, de maneira que possamos utilizar a média dos *spreads* dos papéis líquidos como *proxy* dos *spreads* de papéis ilíquidos (nível). Propomos ainda estimar a variação dos *spreads* entre um período e outro por meio de aprendizado supervisionado. Desta forma, para debêntures ilíquidas, a estimativa do valor do *spread* para o período seguinte é dada pela soma das duas estimativas.

Nos experimentos, para as debêntures ilíquidas, o método de aprendizado não supervisionado produziu estimativas com R^2 superiores a 55%, um resultado, por si só, razoável. E a adição das estimativas geradas pelos métodos de aprendizado supervisionado, conseguiram elevar o R^2 para valores entre 70% a 80%.

Assim, acreditamos que a principal contribuição deste artigo é a proposta de combinação de técnicas de aprendizado não supervisionado (*k*-means) com supervisionado (modelos de previsão RL e NN), aliando as suas respectivas características para a solução do problema de apreçamento de debêntures ilíquidas no mercado brasileiro, produzindo resultados satisfatórios para a gestão destes instrumentos.

Os resultados obtidos nos experimentos também indicam que há mais elementos a serem explorados, tanto com relação à incrementos na coleta de dados, quanto à melhoria na configuração dos métodos computacionais utilizados.

À medida que o mercado de debêntures for se aquecendo e novas emissões forem feitas, maior e mais rica será a base de dados disponível para pesquisa. Um aumento expressivo nos dados de treinamento com quantidade maior e mais diversa de observações ajudaria a melhorar o desempenho dos modelos.

Trabalhos futuros podem, ainda, explorar as aplicações da metodologia proposta com a implementação de outras combinações de modelos não supervi-

sionado e supervisionados, como por exemplo, agrupamento hierárquico com **Random Forest** ou **Support Vector Regression**.

Há espaço ainda para estudar o quanto o desempenho da metodologia é afetado pela mudança de corte nas janelas temporais dos períodos de treinamento e de teste, bem como pela mudança da série temporal objeto de estudo, por meio da aplicação da metodologia para séries de *spreads* de outros indicadores como CDI e IGPM.

Referências

- Aldridge, I. e Avellaneda, M. (2019). Neural networks in finance: Design and performance, *Journal of Financial Data Science* **Vol. 1**(No. 4): pp. 39–62.
- ANBIMA (2017). Conselho de regulação e melhores práticas de negociação de instrumentos financeiros.
URL: <http://www.anbima.com.br/data/files/DA/C7/CB/00/23BF95104FEB5B9568A80AC2/Deliberacao-n20.pdf>
- Arce, P., Antognini, J., Kristjanpoller, W. e Salinas, L. (2019). Fast and adaptive cointegration based model for forecasting high frequency financial time series, *Computational Economics* **54**(1): 99–112.
- Baum, C. (2018). KPSS: Stata module to compute kwiatkowski-phillips-schmidt-shin test for stationarity.
- Celebi, M. E. e Aydin, K. (2016). *Unsupervised learning algorithms*, Springer.
- Coats, P. K. e Fant, L. F. (1993). Recognizing financial distress patterns using a neural network tool, *Financial Management* **Vol. 22**(No. 3): pp. 142–155.
- Confessor, K. L. A. e dos Santos, J. F. (2020). A valorização das debêntures: Fatores evidenciados na literatura, *Revista Brasileira de Administração Científica* **Vol. 11**(No. 1).
- CPC (2012). Mensuração do valor justo. Comitê de Pronunciamentos Contábeis.
URL: http://static.cpc.aatb.com.br/Documentos/395_CPC_46_rev%2012.pdf
- Curi, L. Z. (2008). *Aplicação de redes neurais na precificação de debêntures*, Master's thesis, Escola de Economia de São Paulo. Fundação Getúlio Vargas.

- Damodaran, A. (2005). Marketability and value: Measuring the illiquidity discount, *Stern School of Business* .
- Dutta, S. e Shekhar, S. (1988). Bond rating: A nonconservative application of neural networks, *IEEE* .
- Elton, E. J. e Green, C. T. (1998). Tax and liquidity effects in pricing government bonds, *The Journal of Finance* **LIII**(5).
- Elton, E. J., Gruber, M. J., Agrawal, D. e Mann, C. (2001). Explaining the rate spread on corporate bonds, *The Journal of Finance* **LVI**(1): 247–278.
- Enders, W. (2012). Applied econometric time series, *Privredna kretanja i ekonomska politika* **132**: 93.
- Ganguli, S. e Dunnmon, J. (2017). Better models for prediction of bond prices, *Stanford University* .
- Itau, S. S. (2017). Manual de marcação a mercado.
URL: https://www.itau.com.br/_arquivosstaticos/SecuritiesServices/defaultTheme/PDF/ManualPrecificacao.pdf
- John, K., Lynch, A. W. e Puri, M. (2003). Credit rating, collateral and loan characteristics: implication for yield, *Journal of Business* **76**: 371–410.
- Koziol, C. e Sauerbier, P. (2002). Valuation of bond illiquidity: An option-theoretical approach, *University of Mannheim* .
- Lahmiri, S. (2018). Minute-ahead stock price forecasting based on singular spectrum analysis and support vector regression, *Applied Mathematics and Computation* **320**: 444–451.
- Medeiros, M. C., Veiga, A. e Pedreira, C. E. (2001). Modeling exchange rates: smooth transitions, neural networks, and linear models, *IEEE* **Vol. 12**: 755–764.
- Mercantil do Brasil, C. S. C. (2016). Manual de precificação de ativos.
URL: https://mercantildobrasil.com.br/Anexos/PDFs/Geral/FundosInvestimentos/Manual_de_Precificacao_MBCorretora.pdf
- Merton, R. C. (1974). On the pricing of corporate debt: The risk structure of interest rates, *Journal of Finance* **Vol. 29**(No. 2): pp. 449–470.

- Mohanty, S. e Dash, R. (2019). A support vector regression framework for indian bond price prediction, *2019 International Conference on Applied Machine Learning (ICAML)*, IEEE, pp. 69–72.
- Parekh, B., Shah, N., Mehta, R. e Shah, H. (2013). Bond market prediction using an ensemble of neural networks, *International Journal of Computer Applications* **Vol. 82**(No. 4).
- Patria, I. (2016). Manual de precificação de ativos.
URL: <http://www.patria.com/Content/Upload/859a90e5-40dd-4983-8409-2960f103cace-manual-de-precificacao-patria.pdf>
- Samarasinghe, S. (2007). *Neural Networks for Applied Sciences and Engineering*, first edition ed., Auerbach Publications, Boca Raton, NY, USA.
- Sapna, S., Tamilarasi, A., Kumar, M. P. et al. (2012). Backpropagation learning algorithm based on levenberg marquardt algorithm, *Comp Sci Inform Technol (CS and IT)* **2**: 393–398.
- Secches, P. (2006). A influência do risco de liquidez no apreçamento de debêntures, *Escola de Economia de São Paulo. Fundação Getulio Vargas* .
- Sheng, H. H. e Saito, R. (2005). Determinantes de spread das debêntures no mercado brasileiro, *Revista de Administração de Empresas* **Vol. 40**(No. 2): pp. 193–205.
- Syakur, M., Khotimah, B., Rochman, E. e Satoto, B. (2018). Integration *k*-means clustering method and elbow method for identification of the best customer profile cluster, *IOP Conference Series: Materials Science and Engineering* **336**(1): 012017.
- Wernert, P. (2009). Illiquid instruments - how to price illiquid or otc instruments at market value. Zanders Treasury & Finance Solutions.
URL: <http://zanders.eu/en/latest-insights/illiquid-instruments/>
- West, D. (2000). Neural network credit scoring models, *Elsevier Science* .
- Yang, W., Wang, K. e Zuo, W. (2012). Neighborhood component feature selection for high-dimensional data., *JCP* **7**(1): 161–168.