



distanceMonitoraforestal: um pacote em R para estimar as densidades de populações de aves e mamíferos do Programa Monitora

Vitor Nelson Teixeira Borges-Júnior^{1,2,3}

 <https://orcid.org/0000-0002-7726-9871>

Luciana Ardenghi Fusinato^{3,4*}

 <https://orcid.org/0000-0001-9354-1698>

* Contato principal

Helga Correa Wiederhecker⁵

 <https://orcid.org/0000-0002-6454-0829>

Elildo Alves Ribeiro de Carvalho⁶

 <https://orcid.org/0000-0003-4356-2954>

¹ Universidade Federal do Rio de Janeiro, Departamento de Ecologia, Laboratório de Vertebrados, Rio de Janeiro/RJ, Brasil. <vntborgesjr@gmail.com>.

² Centro de Conhecimento em Biodiversidade, Belo Horizonte/MG, Brasil. <vntborgesjr@gmail.com>.

³ Piper 3D – Pesquisa, Educação e Consultoria Ambiental, Rio de Janeiro/RJ, Brasil. <vntborgesjr@gmail.com, lufusinato@gmail.com>.

⁴ Reserva Ecológica de Guapiaçu/REGUA, Guapiaçu – Cachoeiras de Macacu/RJ, Brasil. <lufusinato@gmail.com>.

⁵ WWF-Brasil, Brasília/DF, Brasil. <helgacorreia@wwf.org.br>.

⁶ Instituto Chico Mendes de Conservação da Biodiversidade/ICMBio, Centro Nacional de Pesquisa e Conservação de Mamíferos Carnívoros/CENAP, Atibaia/SP, Brasil. <elildo.carvalho-junior@icmbio.gov.br>.

Recebido em 05/03/2024 – Aceito em 02/10/2024

Como citar:

Borges-Júnior VNT, Fusinato LA, Wiederhecker HC, Carvalho EAR. distanceMonitoraforestal: um pacote em R para estimar as densidades de populações de aves e mamíferos do Programa Monitora. Biodivers. Bras. [Internet]. 2025; 15(2): 85-101. doi: 10.37002/biodiversidadebrasileira.v15i2.2536

Palavras-chave: Função de detecção; ajuste de modelos; ecologia de populações; conservação.

RESUMO – As ações humanas que levam à redução e à transformação de ambientes naturais são uma das principais causas da crescente perda global de biodiversidade. O aumento nas taxas de extinção de espécies e a redução das populações naturais são alguns indicadores dessa perda. No contexto brasileiro, o Programa Monitora do ICMBio é uma importante fonte de dados para o monitoramento da biodiversidade. Assim, uma parceria entre WWF e ICMBio foi feita em 2023 para buscar estratégias analíticas que impulsionassem a análise dos dados das espécies monitoradas. A partir dessa parceria, desenvolvemos o pacote distanceMonitoraforestal, criado para aplicar os modelos de análise por distância para estimar densidades populacionais para espécies de aves e mamíferos monitoradas pelo módulo Florestal Global. Neste artigo, apresentamos as funcionalidades do pacote a partir dos dados da população de *Sapajus apella* (macaco-prego) monitorada na REBIO do Uatumã e comentamos detalhes do método de análise de dados de amostragem por distância.



distanceMonitoraflorestal: an R package to estimate population density of birds and mammals of the Monitora Program

Keywords: Detection function; model adjustment; population ecology; conservation.

ABSTRACT – The anthropogenic activities that transform and reduce natural habitats are the major drivers of the global loss of biodiversity. The increase in species extinction rates and natural population declines are some indicators of this loss. In the Brazilian context, the ICMBio Monitora Program is an important source of data for monitoring biodiversity. Thus, a partnership between WWF and ICMBio was formed in 2023 to seek analytical strategies that would boost the analysis of data on monitored species. From this partnership, we developed the distanceMonitoraflorestal package, created to apply distance analysis models to estimate population densities for bird and mammal species monitored by the Global Forestry module. In this article, we present the package's functionalities based on data of the *Sapajus apella* population (tufted capuchin) at REBIO do Uatumã and comment on details of the distance sampling data analysis method.

distanceMonitoraflorestal: un paquete en R para estimar la densidad de poblaciones de aves y mamíferos del Programa Monitora

Palabras clave: Función de detección; ajuste del modelo; ecología de la población; conservación.

RESUMEN – Las acciones humanas que conducen a la reducción y transformación de los ambientes naturales son una de las principales causas de la creciente pérdida global de biodiversidad. El aumento de las tasas de extinción de especies y la reducción de las poblaciones naturales son algunos indicadores de esta pérdida. En el contexto brasileño, el Programa Monitora del ICMBio es una importante fuente de datos para el monitoreo de la biodiversidad. Así, en 2023 se formó una alianza entre WWF e ICMBio para buscar estrategias analíticas que impulsaran el análisis de datos sobre especies monitoreadas. A partir de esta asociación, desarrollamos el paquete distanceMonitoraflorestal, creado para aplicar modelos de análisis de distancia para estimar densidades de poblaciones de especies de aves y mamíferos monitoreadas por el módulo Forestal Global. En este artículo, presentamos las funcionalidades del paquete basadas en datos de la población de *Sapajus apella* (mono capuchino) monitoreada en la REBIO en Uatumã y comentamos detalles del método de análisis de datos de muestreo a distancia.

Introdução

Um dos principais efeitos das ações humanas sobre o planeta é a perda de biodiversidade, que pode ser observada, dentre outros indicadores, pela redução no tamanho populacional das espécies e aumento nas taxas de extinção [1][2][3]. Essa perda tem sido chamada por alguns autores de sexta extinção em massa e, considerando o viés sobre o atual conhecimento da biodiversidade, no qual vários grupos de seres vivos ainda são pouco conhecidos ou estudados, há um desafio de se estimar o quanto

de biodiversidade estamos perdendo [4]. Nesse contexto, iniciativas de monitoramento de dinâmicas que indicam perdas ou ganhos de biodiversidade, como variações nos tamanhos e densidades de populações naturais ou taxas de extinção/colonização de espécies, são essenciais e urgentes [5].

Em um país megadiverso como o Brasil, tanto os monitoramentos quanto a divulgação dos seus resultados são um desafio proporcional ao número de espécies e a sua área continental. Isso faz com que o número de espécies brasileiras monitoradas quanto ao tamanho populacional seja proporcionalmente

baixo diante da diversidade de espécies existente no país, como demonstrado para a fauna de vertebrados no Relatório Planeta Vivo [6][7]. Tal contexto demanda a atuação conjunta de diferentes setores, esferas e instituições para promover e ampliar o levantamento, a gestão e a divulgação dos dados sobre a biodiversidade. Na esfera Federal, o Instituto Chico Mendes de Conservação da Biodiversidade (ICMBio) implementou o Programa Nacional de Monitoramento da Biodiversidade (Monitora, Instruções Normativas ICMBio n. 3/2017 e n. 2/2022), com o intuito de impulsionar o monitoramento do estado da biodiversidade, de serviços ecossistêmicos e a sua dinâmica espacial e temporal em unidades de conservação federais. O programa foi delineado visando maior autonomia à gestão das UCs e ao órgão ambiental sobre a pesquisa da biodiversidade, desenvolvendo protocolos que pudessem ser executados com maior independência de especialistas e acadêmicos, com a participação das comunidades locais, e buscando agilidade e integração no levantamento e processamento dos dados sobre a dinâmica da biodiversidade brasileira.

Atualmente o Monitora é composto por três subprogramas: Terrestre, Aquático e Continental e o Marinho e Costeiro, que reúnem dados de mais de cem UCs federais dos diferentes biomas brasileiros. O subprograma Terrestre é composto pelos componentes Campestre/Savânica (15 UCs) e Florestal (64 UCs), sendo cada componente dividido em alvos globais e complementares. Fazem parte dos alvos globais do componente Florestal as borboletas frugívoras, as aves (com um alvo específico para as aves terrícolas cinegéticas), os mamíferos terrestres de médio e grande porte e as plantas arbóreas e arborescentes, para os quais foram elaborados protocolos de monitoramento implementados nas UCs federais [8][9][10].

Dentre os inúmeros desafios de um programa de monitoramento de grande escala está a análise dos dados obtidos. Em 2023, o *World Wildlife Fund for Nature* (WWF) somou esforços com o ICMBio através de uma parceria para que uma rotina de análise dos dados do Monitora fosse desenvolvida em R, uma linguagem de programação de código aberto e de livre acesso amplamente utilizada para condução de análises estatísticas [11]. O código precisaria ser o mais acessível possível para a aplicação por analistas ambientais e público interessado. Dessa parceria, desenvolvemos o pacote distanceMonitaraFlorestal. O pacote agrupa funcionalidades de diferentes

pacotes pré-existent, com o objetivo de facilitar a formatação, visualização e análise dos dados do Programa Monitora, tornando mais acessível o uso da linguagem R a partir de funções e fluxos de trabalho e de análise documentadas em língua portuguesa.

Para cumprir com tal objetivo, apresentamos aqui o pacote e algumas de suas funcionalidades. O artigo traz o resultado de três meses de trabalho de elaboração, testes, refinamento e documentação dos códigos que foram aperfeiçoados durante e após o *workshop* de capacitação de dois dias promovido pela WWF e que contou com a participação de 13 analistas do ICMBio. Aqui ilustramos as etapas analíticas utilizando os dados de uma das populações de primata monitoradas no programa: a espécie *Sapajus apella* (macaco-prego-das-guianas) na REBIO do Uatumã (lat.-1,276, long.-59,494), município de Presidente Figueiredo, no estado do Amazonas. Escolhemos como modelo um primata para ilustrar também a possibilidade do uso de tamanho dos grupos como covariável no cálculo das estimativas de abundância e densidade populacionais. Ao final, apresentamos as variações nas estimativas de abundância e densidade ao longo dos anos, comentamos sua aplicabilidade, e discutimos as perspectivas futuras de uso e continuidade de desenvolvimento do pacote.

Pacote distanceMonitaraFlorestal

Base de dados

A base de dados que compõe o pacote engloba uma pequena fração da totalidade dos dados do programa Monitora. Ela reúne os dados obtidos através do protocolo de monitoramento dos alvos globais de aves cinegéticas e mamíferos terrestres de médio e grande porte do componente Florestal [12]. O delineamento amostral desse protocolo compreende unidades amostrais no formato de trilhas (transecções lineares), com extensão de até 5 km e 1 m de largura, demarcadas a cada 50 m, e estão sendo implantadas desde 2014. Esse delineamento possibilita a modelagem da detectabilidade dos animais, através do ajuste de uma função de detecção, sendo útil para estimativas de abundância e densidade populacional [13]. Para obter uma descrição detalhada do delineamento amostral e protocolo de amostragem de campo do Programa Monitora Componente Florestal o leitor pode consultar a documentação do ICMBio [8][9][10].

O *data frame* de dados que integra o pacote traz em cada uma das suas 27.887 linhas observações de 41 espécies de aves terrícolas cinegéticas e 156 espécies de mamíferos de médio e grande porte, realizadas entre os anos de 2014 e 2022. Os registros estiveram distribuídos entre quarenta UCs e 106 estações amostrais, em uma distância total de 11.876.350 km percorridos. Cada uma das 22 colunas corresponde às variáveis anotadas durante aplicação do protocolo de amostragem em campo como, por exemplo, a data e hora da observação, a identidade das espécies e sua distância perpendicular em relação a trilha percorrida. Os usuários do pacote podem acessar os nomes e descrições das informações contidas em cada coluna utilizando o comando *help* (*monitora_aves_masto_florestal*).

Descrição geral

O pacote *distanceMonitoraFlorestal* foi desenvolvido em linguagem R e reúne as funções para importação, transformação e exploração dos dados, assim como para a análise e visualização dos resultados. O pacote leva em seu nome a junção de palavras que fazem referência a i) “*distance*” abordagem analítica adotada para obter as estimativas de densidade das populações; ii) “*Monitora*” o nome do programa de monitoramento e iii) “*florestal*” o nome de sua componente. O termo “amostragem por distância” (*distance sampling*) foi usado pela primeira vez por Buckland e colaboradores em 1993 no primeiro livro a fazer uma síntese sobre o método [13]. No mesmo ano foi criado o programa *Distance* para o sistema operacional Windows [14][15], cuja última versão foi atualizada em 2023. Recentemente, foram desenvolvidos também alguns conjuntos de pacotes para análise de dados obtidos por amostragem por distância em linguagem R, como: *unmarked* [16]; *mrds* [17]; *Rdistance* [18]; e *Distance* [19]; dentre outros pacotes disponíveis no website do programa *Distance* (<https://distancesampling.org/>). Dentre esses pacotes, utilizamos algumas das funções de análise e exploração dos dados do pacote *Distance*, por considerá-lo um dos mais completos.

Os fluxos de análise foram divididos nas seguintes etapas, que estão disponíveis na seção *Articles* da página da internet do pacote: carregando uma base de dados atualizada; checagem, seleção e transformação dos dados; exploração-e-selecao-de-dados-para-analises; ajuste dos modelos: Fluxo 1 – Distâncias exatas com repetições; ajuste dos modelos: Fluxo 2 – Distâncias exatas sem repetições; ajuste

dos modelos: Fluxo 3 – Distâncias agrupadas com repetições; ajuste dos modelos: Fluxo 4 – Amostragem por distância com múltiplas covariáveis; ajuste dos modelos: Fluxo 5 – Amostragem por distância estratificada por unidade de conservação; ajuste dos modelos: Fluxo 6 – Amostragem por distância estratificada por ano. Ilustramos abaixo as principais etapas para a exploração e análise dos dados de nossa espécie-modelo. As recomendações ao longo do texto são baseadas nas principais referências para *distance sampling* [13][20][21][22].

Explorando os dados

Para que o método de análise por distância possa ser utilizado para estimativas baseadas em modelos, são recomendadas quantidades mínimas de observações e de transectos. Segundo Buckland e colaboradores [21], o número mínimo de réplicas para os transectos deve ser de 10-20, o que deve aumentar para espécies cujas populações são distribuídas em manchas. O número mínimo sugerido de animais ou grupos é de 60-80 animais (ou grupos) quando a amostragem é feita pelo método dos transectos lineares. É possível utilizar números menores que esses para realizar as análises, porém deve-se ter o cuidado de verificar se as funções de detecção estão bem modeladas. Os números recomendados se aplicam a cada função de detecção a ser modelada. Assim, quando se pretende estratificar os dados, dividindo-os em subconjuntos, seja por região geográfica ou por período amostral, é necessário ter um cuidado para que a suficiência amostral se mantenha dentro dos subconjuntos.

Algumas situações amostrais podem levar a efeitos indesejados na distribuição dos dados, dificultando o ajuste dos modelos. Os efeitos são: i) empilhamento das observações (*heaping*, Figura 1A), ocorre quando valores das distâncias perpendiculares são arredondados. Tal efeito é um acúmulo de observações sobre a mesma distância, que corresponde a um valor redondo. Um tipo específico de empilhamento é quando ele se dá sobre a distância zero, também chamado de pico próximo a distância zero (*spike near zero distance*, Figura 1B). Um segundo efeito: ii) movimento de resposta ao observador (Figura 1C), ocorre quando os animais podem ser repelidos ou atraídos pelo observador. Para animais que apresentam resposta de fuga em relação ao observador, o efeito no histograma de frequências é um aumento nas observações em distâncias

intermediárias. Para animais que são atraídos pelo observador, a tendência é de um aumento nas observações em distâncias próximas a zero, o que também pode gerar o padrão de pico próximo a zero. O efeito desse viés sobre as observações é enviesar também o ajuste dos modelos (formato da curva), as probabilidades de detecção e as abundâncias estimadas em cada faixa de distância. O terceiro efeito:

iii) superdispersão (Figura 1D), é resultante da falta de independência entre observações, como no caso das repetições amostrais em um mesmo transecto. Se os animais apresentarem um padrão de distribuição espacial que se mantém ao longo do tempo, o efeito das amostragens repetidas no histograma será o de picos em algumas distâncias e de reduções abruptas em outras.

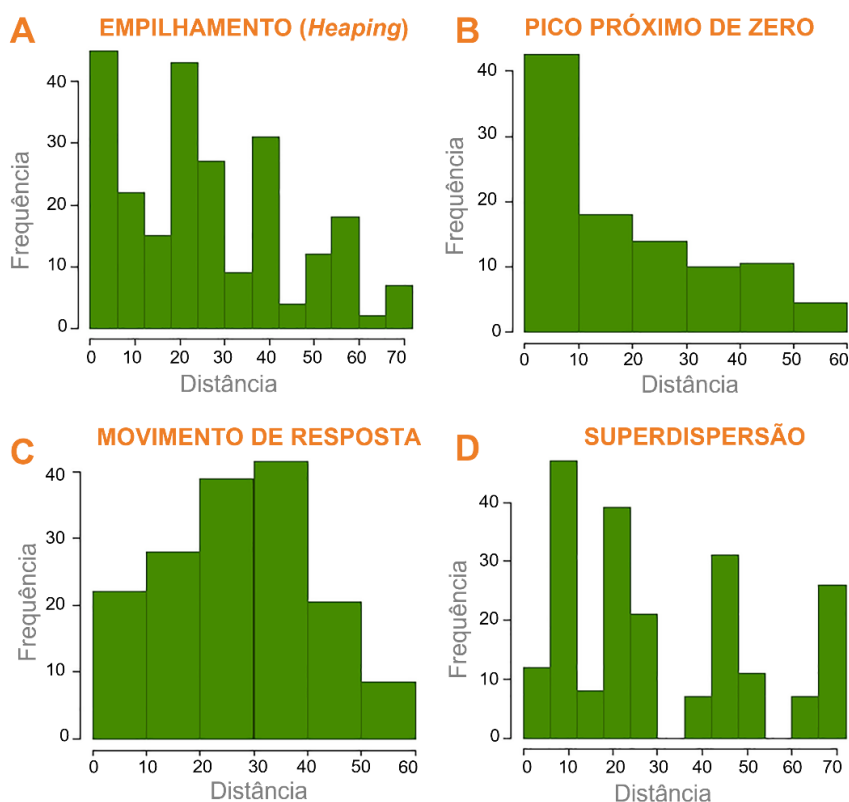


Figura 1 – Exemplos dos padrões indesejados de distribuição das frequências de observação em relação às distâncias do transecto: A = no empilhamento, observa-se picos a cada vinte unidades de distância, resultantes de arredondamento dos dados; B = no pico próximo a zero, observa-se um empilhamento das observações na distância zero; C = no movimento de resposta, a maior frequência de observações desloca-se para distâncias intermediárias; e D = na superdispersão, o viés de amostragens repetidas na mesma trilha pode levar a picos e quedas abruptas nas observações.

Para *Sapajus apella*, na REBIO do Uatumã, as amostragens foram realizadas entre julho de 2014 e dezembro de 2021 em três estações amostrais percorridas em 201 dias, totalizando um esforço de 1.045 km. O número total foi de 249 observações, com tamanhos de grupo que variaram de um a 22 indivíduos. Os anos de 2016 e 2018 foram os que apresentaram o maior número de observações (N = 40), enquanto 2020 teve o menor número (N = 18). O gráfico de distribuição das frequências de

observações pelas distâncias em relação à trilha, para os dados acumulados em todos os anos, demonstra um pico de observações em zero (Figura 2A). Através do *box plot* (Figura 2B) podemos ver que 25% das observações foram feitas até dois metros de distância da trilha, metade das observações até 13 m, e 75% até 25 m. A maior distância registrada foi 70 m, mas registros acima de 60 m foram considerados *outliers*.

Para gerar os gráficos como o da Figura 2, primeiro é necessário filtrar os dados que são

automaticamente carregados com o pacote, selecionando a unidade de conservação e a espécie foco da análise utilizando a função `filtrar_dados()`.

```
# filtrar dados e o esforço do dia para conter apenas
# observações com o mesmo esforço total
sapajus_apella_rebio_uatuma <- filtrar_dados(
  dados = monitora_aves_masto_florestal,
  esforco_dia == 5000,
  nome_uc == "rebio_do_uatuma",
  nome_sp == "sapajus_apella",
  validacao_obs = "especie"
)
```

Em seguida, os dados devem ser transformados para o formato processável pelas funcionalidades do pacote `Distance`, que estão inclusas no `distanceMointoraFlorestal`, utilizando a função `transformar_dados_formato_Distance()`. Nessa função, também deve ser indicado o nível de estratificação que determinará a forma como o esforço amostral (distância total percorrida em metros) será calculado

(exemplo: por ano; por ano e por estação do ano; por ano e por unidade de conservação). No código a seguir, a estratificação escolhida foi por ano. Logo, o esforço amostral foi calculado como a distância total percorrida em cada ano de amostragem:

```
# transformar os dados para o formato Distance
sapajus_apella_rebio_uatuma_distance <- transformar_
dados_formato_Distance(
  dados = sapajus_apella_rebio_uatuma,
  year
)
```

A partir dos dados transformados podemos gerar os gráficos da Figura 2, aplicando a função `plotar_distribuicao_distancia_estatico()`.

```
# gerar os gráficos de distribuição de distâncias
plotar_distribuicao_distancia_estatico(
  sapajus_apella_rebio_uatuma_distance,
  largura_caixa = 1
)
```

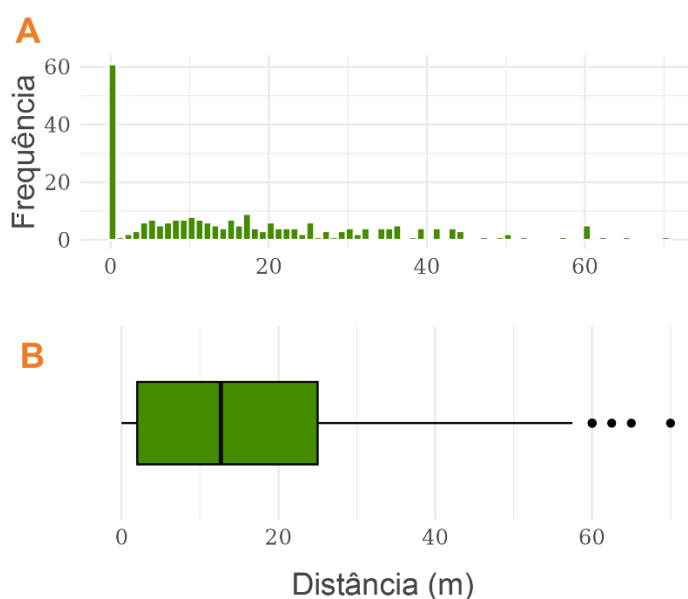


Figura 2 – Exploração dos dados de observações de *Sapajus apella* na REBIO do Uatuma para os anos de 2014 a 2021: A = distribuição de frequência das observações dos grupos em relação à distância da trilha e B = box plot das observações dos grupos em relação à distância da trilha.

Truncamento dos dados

A perda da detectabilidade com a distância perpendicular faz com que, a partir de determinada distância, os registros de observação contribuam pouco para a estimativa de probabilidade de detecção. Usar

todos os dados, nesses casos, pode resultar em maior variância nas taxas de encontro e em modelos piores em termos de ajuste aos dados. Assim, recomenda-se truncar os dados eliminando as observações feitas em distâncias onde a detectabilidade é muito baixa. A distância máxima que será utilizada para a análise

dos dados é chamada também de meia largura (*half-width*, w). Um dos critérios de corte sugerido seria eliminar entre 5 e 10% das observações mais distantes para o método de amostragem por transecto em linha, ou então a partir do ponto onde a probabilidade de detecção $g(w) = 0,15$ [22, pg. 16.]

A definição da distância de truncamento deve ser conduzida em duas etapas complementares. A primeira etapa consiste em avaliar estatisticamente o ajuste dos modelos pelo p-valor do teste de Cramér-von Mises. Nós incluímos nos códigos dos fluxos de análises uma função para implementar esta etapa do teste da distância de truncamento: `selecionar_distancia_truncamento()`. Abaixo, encontra-se o código utilizado para selecionar a distância de truncamento, tomando como exemplo os dados para o ano de 2014. Os parâmetros estimados estão listados na Tabela 1.

```
# filtrar dados para o ano de 2014
sapajus_apella_rebio_uatuma2014 <- filtrar_dados(
  dados = sapajus_apella_rebio_uatuma,
  ano == 2014
)
# transformar os dados para o formato Distance
sapajus_apella_rebio_uatuma_distance2014
<- transformar_dados_formato_Distance(
  dados = sapajus_apella_rebio_uatuma2014,
  year
)

# testar distancia de truncamento
sapajus_apella_rebio_uatuma_truncamento2014
<- selecionar_distancia_truncamento(sapajus_apella_
  rebio_uatuma_distance2014)

# retornar as abundâncias estimadas
sapajus_apella_rebio_uatuma_truncamento2014$modelos

# retornar a tabela comparativa do ajuste dos modelos
sapajus_apella_rebio_uatuma_truncamento2014$selecao
```

Tabela 1 – Estimativa da abundância e tabela comparativa do ajuste de modelos de função de detecção Half-normal para seleção da proporção superior a ser eliminada dos dados, para determinar a melhor distância de truncamento a partir do teste de Cramér-von Mises e dos valores de AIC. C-vM p-valor = valor de probabilidade do teste de Cramér-von Mises; P-hat = probabilidade de detecção estimada; ep P-hat = erro padrão da probabilidade de detecção.

Modelo	Função-chave	Abundância estimada	C-vM p-valor	P-hat	ep P-hat
25%	Half-normal	25,12	0,59	0,63	0,14
20%	Half-normal	26,72	0,65	0,63	0,13
15%	Half-normal	32,56	0,70	0,55	0,11
10%	Half-normal	30,23	0,67	0,63	0,11
5%	Half-normal	40,02	0,74	0,46	0,08

Como indicado, os códigos acima filtram os dados para o ano de 2014 e, em seguida, é ajustado um modelo simples do tipo half-normal (ver detalhes na próxima seção). Utilizando a configuração original da função, as proporções de 5%, 10%, 15%, 20% e 25% de truncamento serão testadas. Assim, podem ser removidas dos dados as observações mais distantes do transecto em proporções que variam de 5 a 25% do total das observações. Como resultado, obtivemos as abundâncias estimadas para cada modelo e uma tabela contendo: os resultados dos testes de bondade do ajuste de Cramér-von Mises; as probabilidades de detecção estimadas; e os respectivos valores de erro padrão. Quanto maior o p-valor do teste de Cramér-von Mises, melhor é o ajuste do modelo aos dados. De acordo com esse critério, a distância de

truncamento seria definida eliminando 5 ou 15% da porção superior das distâncias registradas. Contudo, tal resultado não deve ser avaliado de forma isolada, mas sim em conjunto com os resultados da segunda etapa que consiste em avaliar graficamente o ajuste da função de detecção aos dados.

Para auxiliar nessa tarefa, o pacote também traz uma função para gerar os histogramas de probabilidade de detecção empírica com as respectivas curvas da função de detecção ajustadas: `plotar_funcao_deteccao_selecao_distancia_truncamento()`. Para o exemplo ilustrado na Figura 3, uma boa alternativa seria selecionar o truncamento em 10 ou em 20% dos dados. Essas distâncias de truncamento apresentaram um bom ajuste da distribuição de

probabilidades da função de detecção, representado pelas curvas; à distribuição empírica de probabilidade de distâncias, representada pelos histogramas; além de apresentarem valores semelhantes de p-valor para o teste de Cramér-von Mises. Assim, a tabela de resultados dos testes de bondade não deve ser usada como critério único de escolha, mas sim de forma complementar aos resultados gráficos dos ajustes

das funções de detecção. Para as próximas etapas da análise, a distância de truncamento escolhida foi de 20%.

```
# histograma e função de detecção dos dados
truncados plotar_funcao_deteccao_selecao_distancia_
truncamento(sapajus_apella_rebio_uatuma_
truncamento2014)
```

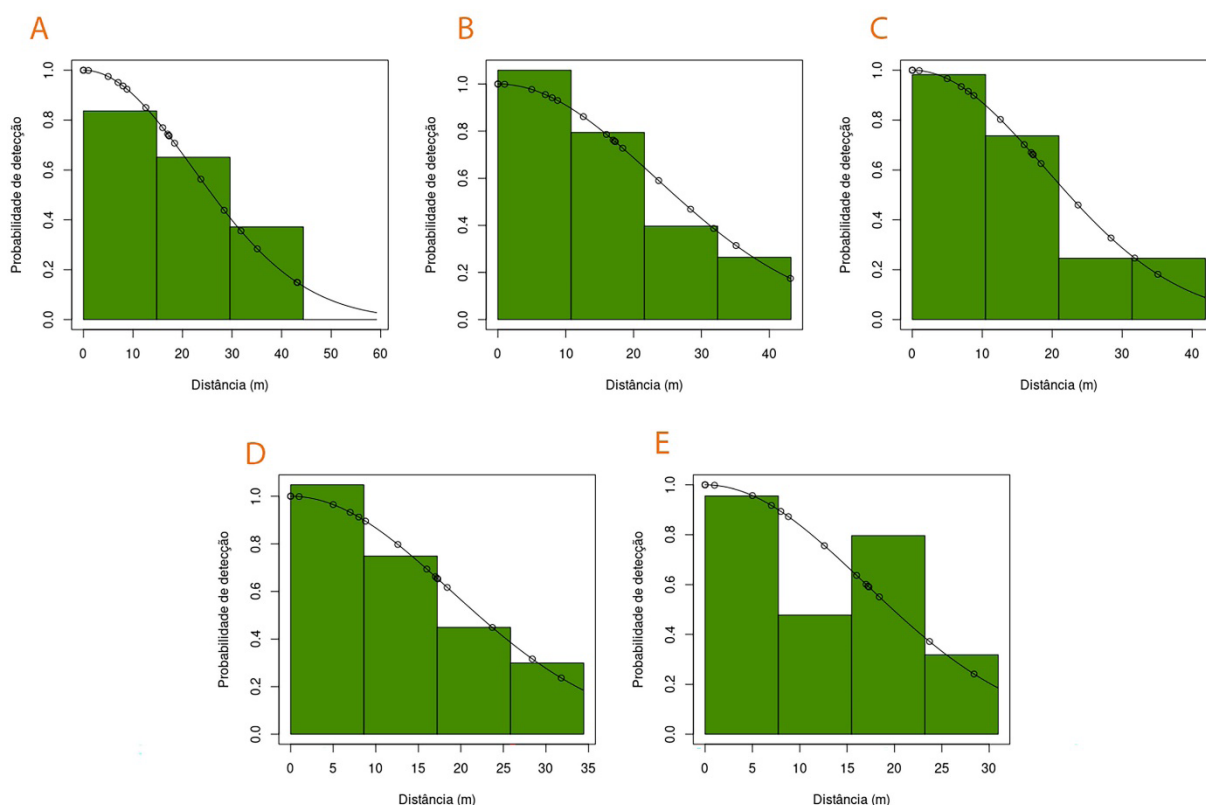


Figura 3 – Funções de detecção (curvas) do tipo half-normal ajustadas aos histogramas das probabilidades de detecção para as distâncias observadas, para truncamentos alternativos dos dados: A = truncamento de 5% das observações mais distantes; B = truncamento de 10% das observações mais distantes; C = truncamento de 15% das observações mais distantes; D = truncamento de 20% das observações mais distantes; e E = truncamento de 25% das observações mais distantes. Notar que o eixo da distância em metros vai diminuindo de valor à medida que os dados são truncados. Dados de *Sapajus apella* para a REBIO Uatumã para o ano de 2014.

Ajuste dos modelos de detectabilidade

A base da análise dos dados na amostragem por distância está na estimativa da detectabilidade, que é a probabilidade de se detectar os animais em função da sua distância em relação à trilha. A detectabilidade é estimada através do ajuste de uma função de detecção aos dados. A função de

detecção é formalmente representada como $g(x)$, que estima a probabilidade de se detectar um animal à distância x da linha. Assume-se, também, que na distância zero todos os animais são detectáveis, e a chance de detecção diminui com a distância em relação ao observador. A derivação das funções de detecção para que a área abaixo da curva seja igual a um leva ao modelo analítico expresso pela função

de densidade de probabilidade $f(x)$. Essa é análoga à função de detecção e possibilita a verificação da qualidade dos modelos ajustados pelos métodos de máxima verossimilhança, a partir do critério de Akaike.

As funções de detecção têm suas formas definidas pelas distribuições das funções chave. A função-chave é uma função que define a forma da curva de detectabilidade que será ajustada aos dados. As três distribuições de função-chave implementadas no pacote distanceMonitoraforestal são do tipo: uniforme (unif), half-normal (hn) e hazard-rate (hr). Dentre as funções disponíveis, há uma variação no número de parâmetros. Uma função uniforme possui dois parâmetros, uma largura e uma altura. Porém, no contexto da amostragem por distância, como a detectabilidade está sendo estimada dentro de uma área fixa, a probabilidade é igual em todas as distâncias e por isso, podemos considerar que ela não possui parâmetros. A função uniforme deve ser sempre acompanhada de um termo de ajuste. A função half-normal possui o parâmetro de escala, chamado de sigma. O parâmetro de escala é responsável pelo espalhamento da distribuição e, quanto maior o valor do parâmetro de escala, mais espalhada será a distribuição dos valores em torno da média. Já a função hazard-rate possui, além do parâmetro de escala, sigma, o parâmetro de forma, beta, que determina a forma da distribuição. As funções chave podem ainda ser otimizadas através dos termos de ajuste. Esses termos podem ser polinômios ou outras funções matemáticas que melhoraram o ajuste das funções chave para melhor refletir os dados observados. Os termos de ajuste implementados no pacote são do tipo cosseno, polinomial simples e polinomial de Hermite. Exceto para a função-chave uniforme, que deve ser utilizada sempre com um termo de ajuste, as demais funções chave não necessitam dessa transformação. Mas, na prática, qualquer termo de ajuste pode ser usado com qualquer função-chave.

O ajuste dos modelos no pacote distanceMonitoraforestal é feito através da função `ajustar_modelos_Distance()`. Ela implementa o ajuste da função de detecção utilizando a função-chave selecionada pelo usuário com diferentes termos de ajuste, a depender da função-chave escolhida. Para função-chave half-normal, a função de detecção é ajustada sem termos de ajuste e com os termos cosseno e polinomial de Hermite; para um ajuste com a função-chave hazard-rate, a função de detecção é ajustada sem termos de ajuste e com os termos cosseno

e polinomial; já para um ajuste com a função-chave uniforme, os termos de ajuste cosseno e polinomial são utilizados. Ilustramos na Figura 4 os modelos ajustados para os dados de *Sapajus apella* em 2014, testando diferentes combinações entre função-chave e termos de ajuste. Utilizamos os seguintes códigos para os testes de modelos alternativos:

```
# ajustando diferentes modelos
# Uniforme com termos de ajuste Cosseno e polinomial
simples
sapajus_apella_distance_unif2014 <- ajustar_modelos_
Distance(
  dados = sapajus_apella_rebio_uatuma_distance2014,
  funcao_chave = "unif",
  truncamento = "20%"
)

# Half-Normal sem termos de ajuste e com termos de
ajuste Cosseno
sapajus_apella_distance_hn2014 <- ajustar_modelos_
Distance(
  dados = sapajus_apella_rebio_uatuma_distance2014,
  funcao_chave = "hn",
  truncamento = "20%"
)

# Hazard-rate sem termos de ajuste e com termos de
ajuste Cosseno
sapajus_apella_distance_hr2014 <- ajustar_modelos_
Distance(
  dados = sapajus_apella_rebio_uatuma_distance2014,
  funcao_chave = "hr",
  truncamento = "20%"
)
```

Escolhendo o melhor modelo

Para o ajuste da função de detecção, podem ser usados os seguintes critérios de avaliação: verossimilhança (L , *likelihood*), critério de informação de Akaike (*Akaike's information criterion* [AIC]), e medidas absolutas do ajuste dos modelos, através dos testes de bondade de ajuste. Esses critérios permitem avaliar a incerteza dos modelos de forma comparativa guiando a escolha do modelo que apresentou a menor incerteza. Logo, é necessário mais de um modelo de $f(x)$ para saber qual apresentou o maior valor de verossimilhança. Uma forma bastante usual de se comparar a verossimilhança entre modelos é através do critério de informação de Akaike, que leva em consideração também a simplicidade do modelo, pelo número de parâmetros (Tabela 2). A comparação entre os modelos pode ser feita no pacote através da função `selecionar_funcao_deteccao_termo_ajuste()`, que resultará em uma tabela de comparação entre os modelos (Tabela 2).

```
# Gerar tabela com o resumo comparativo dos modelos
melhor_modelo_sapajus_apella2014 <- selecionar_
funcao_deteccao_termo_ajuste(
  sapajus_apella_distance_unif2014$Cosseno,
  sapajus_apella_distance_unif2014$`Polinomial simples`,
  sapajus_apella_distance_hn2014$`Sem termo`,
  sapajus_apella_distance_hn2014$Cosseno,
```

```
  sapajus_apella_distance_hn2014$`Hermite polinomial`,
  sapajus_apella_distance_hr2014$`Sem termo`,
  sapajus_apella_distance_hr2014$Cosseno,
  sapajus_apella_distance_hr2014$`Polinomial simples`
)
```

```
melhor_modelo_sapajus_apella2014
```

Tabela 2 – Ajuste de diferentes modelos da função de detecção com base nos dados de observações de *Sapajus apella* na REBIO Uatumã no ano de 2014. Modelo indica a função-chave e o termo de ajuste, caso utilizado. Fórmula indica número de parâmetros do modelo. C-vM p-valor = valor de probabilidade do teste de Cramér-von Mises; P-hat = probabilidade de detecção estimada; se P-hat = erro padrão da probabilidade de detecção; AIC = critério de informação de Akaike.

Modelo	Fórmula	C-vM p-valor	P-hat	se P-hat	AIC	Delta AIC
Uniform with cosine adjustment term of order 1	-	0,65	0,62	0,11	119,03	0,00
Half-normal	~ 1	0,64	0,63	0,13	119,04	0,01
Uniform with simple polynomial adjustment term of order 2	-	0,59	0,70	0,09	119,26	0,23
Hazard-rate	~ 1	0,58	0,44	0,20	121,03	2,00

O critério de avaliação da bondade do ajuste compara a função de distribuição cumulativa de probabilidade da função-chave com a função de distribuição cumulativa empírica. Ou seja, diferente da seleção de modelos que utiliza o AIC para comparar um conjunto de modelos ajustados utilizando diferentes funções-chave, a avaliação de modelos pela bondade do ajuste compara o quão semelhantes são a distribuição cumulativa dos dados observados e a distribuição cumulativa teórica de uma função-chave, sendo realizada uma comparação par a par por vez. No contexto da amostragem por distância, a bondade do ajuste pode ser avaliada utilizando os seguintes métodos: o teste de *Goodness of Fit* e o Q-Q *Plot* com os testes correlacionados de Cramér-von Mises e Kolmogorov-Smirnov. Aqui apresentamos um exemplo de como conduzir o teste de Cramér-von Mises utilizando a função `testar_bondade_ajuste()`, que também resultará em uma tabela comparativa

dos resultados (Tabela 3). A implementação dessa função foi realizada de forma a permitir que o usuário realize mais de uma comparação par a par por vez e, para isso, primeiramente é necessário gerar uma lista com os modelos que serão avaliados.

```
# Gerar lista com os modelos
modelos_sapajus_apella2014 <- gerar_lista_modelos_
selecionados(
  sapajus_apella_distance_hn2014$Cosseno,
  sapajus_apella_distance_hr2014$`Sem termo`,
  sapajus_apella_distance_unif2014$Cosseno,
  sapajus_apella_distance_unif2014$`Polinomial simples`,
  sapajus_apella_distance_hn2014$`Sem termo`,
  nome_modelos_selecionados = melhor_modelo_
  sapajus_apella2014
)
```

```
# histograma e função de detecção dos modelos
plotar_funcao_deteccao_modelos_selecionados(modelos_
  sapajus_apella2014)
```

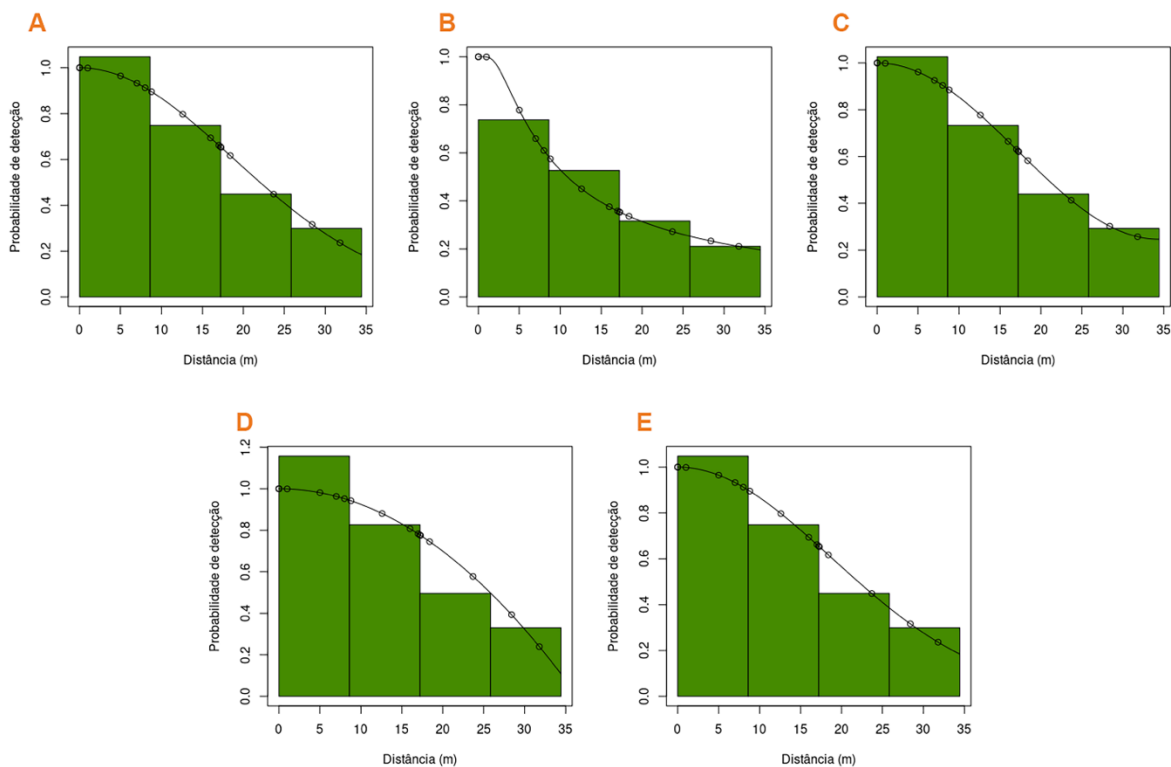


Figura 4 – Funções de detecção (curvas) ajustadas aos histogramas das probabilidades de detecção para as distâncias observadas: A = half-normal com termo de ajuste do tipo cosseno; B = hazard-rate sem termos de ajuste; C = uniforme com termo de ajuste do tipo cosseno; D = uniforme com termo de ajuste do tipo polinomial; e E = half-normal sem termos de ajuste. Dados de *Sapajus apella* para a REBIO Uatumã no ano de 2014.

```
# Teste de bondade de ajuste dos modelos e Q-Q plots
testar_bondade_ajuste(
  dados = modelos_sapajus_apella2014,
  plot = FALSE,
  chisq = FALSE,
)
```

Tabela 3 – Resultados do teste de bondade do ajuste de Cramér-von Mises para modelos de detecção com diferentes funções-chave e termos de ajuste com base nas observações de *Sapajus apella* na REBIO Uatumã no ano de 2014.

Modelo	W	P
Uniform with cosine adjustment term of order 1	0,08	0,64
Half-normal	0,10	0,58
Uniform with simple polynomial adjustment term of order 2	0,08	0,65
Hazard-rate	0,09	0,59

Ajuste dos modelos de detectabilidade com covariável

A detectabilidade do objeto foco da amostragem por distância pode apresentar viés causado por uma variedade de fatores. A hora do dia, a duração

da amostragem, a temperatura ou umidade do ar, a intensidade de chuva e a experiência do observador são alguns exemplos de fatores que podem causar vieses na detectabilidade. Esses fatores podem influenciar a probabilidade de detectar um objeto a uma certa distância, dando origem ao viés.

Uma forma de corrigir o viés da detectabilidade e, consequentemente, obter estimativas de abundância e densidade mais acuradas é o uso desses fatores como covariáveis nos modelos de detecção [23][24].

Para primatas que são animais que formam grupos, um efeito esperado é o chamado viés de tamanho [25]. A tendência é que as pequenas distâncias todos os grupos sejam detectáveis, mas somente os grupos maiores serão detectados a distâncias maiores. Nesse caso, deve-se considerar o tamanho como uma covariável para obtenção de estimativas de abundância e densidade mais acuradas. Esta é uma covariável especial, pois será usada tanto para corrigir efeitos de viés na detectabilidade quanto para corrigir a abundância. Para espécies que não formam grupos, cada observação é considerada um indivíduo. Para as que formam grupos, as informações de tamanho entrarão no cálculo da abundância. Nesses casos, o estimador de abundância Horvitz-Thompson irá incorporar tanto as informações do tamanho do grupo quanto da probabilidade de detecção condicionada às covariáveis pelo método da amostragem por distância com múltiplas covariáveis (Multiple-Covariate Distance Sampling [MCDS]) [24]. Para ajustar uma função de detecção utilizamos a função `ajustar_modelos_Distance()` informando a covariável no argumento `formula = ~ variável` e usando a distribuição do tipo half-normal, sem termos de ajuste.

```
# Ajuste do modelo half-normal sem termos de ajuste e
# com covariável sapajus_apella_distance_hn_size2014 <-
ajustar_modelos_Distance(
  dados = sapajus_apella_rebio_uatuma_distance2014,
  funcao_chave = "hn",
  truncamento = "20%",
  formula = ~ size
)
```

Ajuste do modelo global de detectabilidade, com efeito de tamanho de grupo, estratificado por ano

Além do uso de covariáveis, é possível estratificar os dados em subconjuntos para o ajuste dos modelos. No caso dos dados do Programa Monitora, a estratificação geográfica pode ser utilizada para dados de uma mesma espécie para diferentes UCs. Nesse caso, cada uma das UCs pode ser tratada como um estrato geográfico (`Region.Label`). Maiores detalhes para esse tipo de estratificação podem ser encontrados na *vignette* Ajuste dos modelos: Fluxo 5 – Amostragem por distância estratificada por unidade de conservação da documentação do pacote. Para avaliar a variação dos tamanhos populacionais ao longo do tempo, é interessante estratificar os dados por ano, conforme apresentamos no exemplo a seguir. A estratificação é uma etapa importante para o cálculo das estimativas de abundância e densidade, apresentadas na próxima seção. O gráfico do modelo global (2014-2021) com covariável ajustado às probabilidades de detecção está ilustrado na Figura 5.

```
# Ajuste do modelo half-normal sem termos de ajuste e
# com covariável estratificado por ano
sapajus_apella_distance_hn_year_size <- sapajus_apella_
rebio_uatuma_distance |>
  dplyr::mutate(Region.Label = year) |> # indica a
# estratificação por ano
ajustar_modelos_Distance(
  funcao_chave = "hn",
  truncamento = "20%",
  formula = ~ size
)

# Gráfico de ajuste da função de detecção às
# probabilidades de detecção
plotar_funcao_deteccao_modelos_selecionados(
  dados = sapajus_apella_distance_hn_year_size
)
```

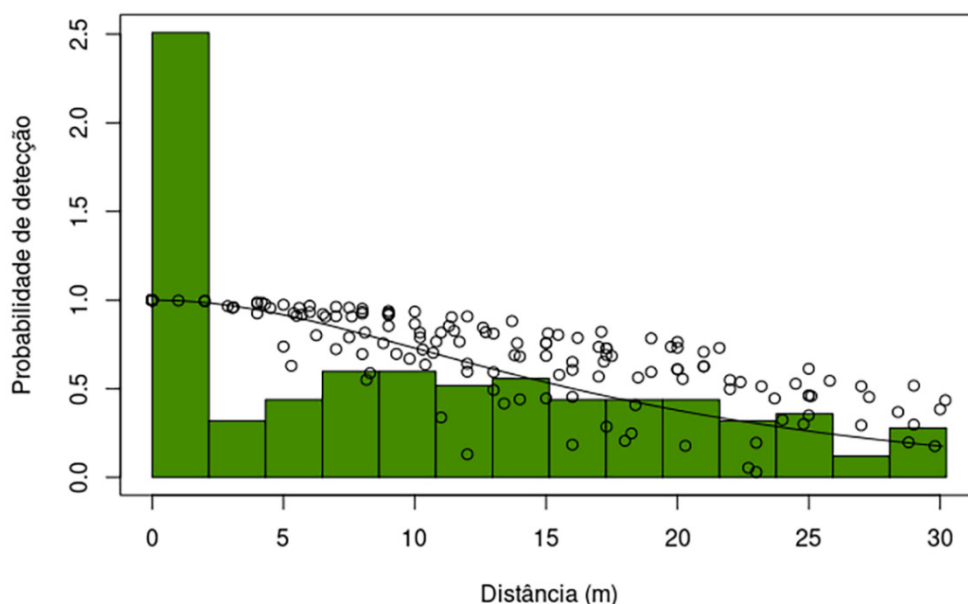


Figura 5 – Função de detecção do tipo half-normal com covariável de tamanho de grupo ajustada ao histograma das probabilidades de detecção para as distâncias observadas. Os pontos representam as probabilidades de detecção, dados os valores da covariável “tamanho de grupo” e a linha representa a função de detecção média considerando os tamanhos dos grupos. Dados de *Sapajus apella* para a REBIO Uatumã para os anos de 2014 a 2021.

Taxas de encontro e estimativas de abundância e densidade

Esta é a etapa final de análise dos dados. Após o tratamento dos dados, ajuste e escolha dos modelos, chegamos aos resultados que realmente interessam para o monitoramento das espécies. As etapas anteriores de inclusão da covariável de tamanho de grupo e estratificação dos dados influenciarão nas estimativas de abundância e densidade, além de permitir que essas sejam calculadas ano a ano. Nesse ponto dos resultados, é importante avaliar também a incerteza associada a essas estimativas. A incerteza é medida pela variância da abundância e da densidade consideradas de modo que uma estimativa será menos precisa quanto maior for a variância associada. Ela serve de base para o cálculo de outras medidas de incerteza, como o coeficiente de variação (CV). O coeficiente de variação é útil para comparar variâncias quando as escalas e/ou unidades de medidas diferem. Assim, ele pode ser usado para comparar

a incerteza de valores estimados (ex. abundância, probabilidade de detecção) para diferentes conjuntos de dados ou diferentes arranjos de análise. Abaixo apresentamos as etapas para o cálculo das taxas de encontro e estimativas de abundância e densidade. Os resultados estão exemplificados nas Tabelas 4, 5 e 6, onde apresentamos somente as seis primeiras linhas dos dados para evitar tabelas muito extensas no corpo do manuscrito. Ao final, apresentamos o gráfico com a variação da densidade de *Sapajus apella* para os anos de 2014 a 2021 na REBIO Uatumã, sem e com os limites inferior e superior do intervalo de confiança de 95% (Figura 6).

```
# Área coberta pela Amostragem
resultado_area_coberta_amostragem <- gerar_resultados_
Distance(
  dados = sapajus_apella_distance_hn_year_size,
  resultado_selecao_modelos = "Half-normal",
  tipo_de_resultado = "area_estudo",
  estratificacao = TRUE
)
```

Tabela 4 – Seis primeiras linhas dos resultados indicando a função de detecção ajustada aos dados (Modelo), o ano, a área coberta (em m²), o esforço empregado (m), o número de indivíduos observados (n), o número de transectos (k), a taxa de encontro, o erro padrão (ep) da taxa de encontro e o coeficiente de variação (cv) da taxa de encontro para as observações de *Sapajus apella* na REBIO Uatumã entre os anos de 2014 e 2020.

Modelo	Ano	Área coberta (m ²)	Esforço (m)	n	k	Taxa de encontro	ep da taxa de encontro	cv da taxa de encontro
Half-normal	2014	19.958.400	330.000	121	3	3,66e-04	1,12e-04	0,30
Half-normal	2015	30.844.800	510.000	432	3	8,47e-04	2,16e-04	0,25
Half-normal	2016	36.288.000	600.000	352	3	5,86e-04	1,03e-04	0,17
Half-normal	2017	29.937.600	495.000	142	3	2,86e-04	4,05e-05	0,14
Half-normal	2018	36.288.000	600.000	210	3	3,50e-04	5,67e-05	0,16
Half-normal	2019	31.752.000	52.5000	169	3	3,21e-04	2,35e-04	0,73

```
# Abundância
resultado_abundancia <-
  gerar_resultados_Distance(
    dados = sapajus_apella_distance_hn_year_size,
```

```
resultado_selecao_modelos = "Half-normal",
tipo_de_resultado = "abundancia",
estratificacao = TRUE
)
```

Tabela 5 – Seis primeiras linhas dos resultados das estimativas de abundância de *Sapajus apella* em cada ano (2014-2021) e estação amostral da REBIO Uatumã. Em cada linha estão indicados o Modelo (função de detecção ajustada aos dados), o ano, o nome da estação amostral, a área total amostrada (em m²), o esforço empregado (m), a abundância estimada e o número de detecções de *Sapajus apella* em cada ano e estação amostral da REBIO Uatumã.

Modelo	Ano	Estação amostral	Área (m ²)	Esforço (m)	Abundância estimada	N de detecções
Half-normal	2014	Cachoeira	19.958.400	110.000	27,47	17
Half-normal	2014	Grid	19.958.400	110.000	103,26	59
Half-normal	2014	Tucumari	19.958.400	110.000	81,90	45
Half-normal	2015	Cachoeira	30.844.800	170.000	412,21	154
Half-normal	2015	Grid	30.844.800	170.000	489,60	202
Half-normal	2015	Tucumari	30.844.800	170.000	202,87	76

```
# Densidade
resultados_densidade <-
  gerar_resultados_Distance(
    dados = sapajus_apella_distance_hn_year_size,
```

```
resultado_selecao_modelos = "Half-normal",
tipo_de_resultado = "densidade",
estratificacao = TRUE
)
```

Tabela 6 – Seis primeiras linhas dos resultados das estimativas de densidade de *Sapajus apella* para cada ano (2014-2021) na REBIO Uatumã. Em cada linha estão indicados o Modelo (função de detecção ajustada aos dados), o ano, a densidade (número de indivíduos por hectare), o erro padrão, o coeficiente de variação, os limites inferiores e superiores dos intervalos de confiança (95%), e os graus de liberdade.

Modelo	Ano	Densidade (ind/ha)	Erro padrão	Coeficiente de variação	Limites inferiores e superiores dos intervalos de confiança (95 %)	Graus de liberdade
Half-normal	2014	0,10	3,45e-06	0,32	0,02 – 0,38	2,15
Half-normal	2015	0,35	8,83e-06	0,24	0,15 – 0,85	2,53
Half-normal	2016	0,21	2,55e-06	0,12	0,15 – 0,29	4,37
Half-normal	2017	0,09	8,29e-07	0,092	0,07 – 0,11	5,58
Half-normal	2018	0,10	1,59e-06	0,16	0,05 – 0,17	2,76
Half-normal	2019	0,10	8,62e-06	0,81	0,005 – 2,10	2,02

Visualização da densidade ao longo do período

```
plotar_resultados_Distance(
  dados = resultados_densidade,
  tipo_de_resultado = "densidade"
)
```

```
plotar_resultados_Distance(
  dados = resultados_densidade,
  tipo_de_resultado = "densidade",
  intervalo_confianca = TRUE
)
```

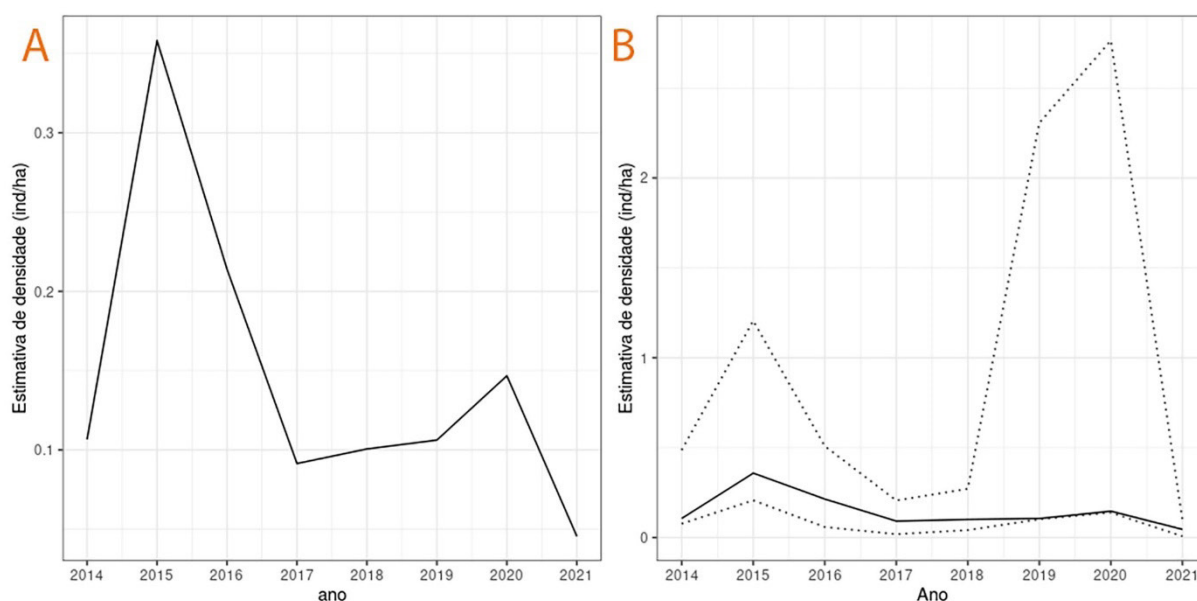


Figura 6 – Estimativas de densidade (indivíduos/ha): A = sem e B = com os limites inferior e superior do intervalo de confiança de 95%, considerando tamanho do grupo, de *Sapajus apella* para a REBIO Uatumã para os anos de 2014 a 2021.

Discussão

Apresentamos aqui a primeira versão do pacote distanceMonitorafflorestal. Acreditamos que o conjunto de funcionalidades atuais são capazes de gerar estimativas de densidade populacional de

dezenas de espécies de aves e mamíferos amostradas pelo Programa Monitora. Entre 2014 e 2022, o protocolo global do componente Florestal reuniu dados de quarenta UCs, em um total de 27.887 registros referentes a 197 espécies (156 de mamíferos e 41 de aves).

O uso do pacote Distance [19] como base para a elaboração das funções do pacote distanceMonitoraflorestal foi vantajoso por se tratar de um pacote relativamente completo dentro das abordagens mais convencionais de amostragem por distância. Entretanto, melhorias ainda podem ser feitas no distanceMonitoraflorestal a partir de funcionalidades de outros pacotes e envolvendo estratégias analíticas que ainda não foram implementadas. Um dos ajustes que pretendemos fazer é a incorporação da função-chave do tipo exponencial negativa. Até o momento, a função-chave do tipo half-normal, com ou sem termos de ajustes, tem demonstrado os melhores ajustes para uma maior quantidade de espécies testadas dentro do conjunto de dados do Programa Monitora. Foram bastante comuns, no conjunto de dados explorados, distribuições de observações com pico em zero. Para esse tipo de distribuição, a função-chave do tipo hazard-rate não é recomendada, pois infla as estimativas de probabilidade de detecção nas distâncias próximas a zero [21, p. 57]. Assim, recomendamos aos usuários do pacote que olhem a distribuição dos dados de observação em relação à distância e, caso tenham pico em zero, evitem usar modelos com a função-chave do tipo hazard-rate para estimativa de abundância.

Um outro aspecto que pretendemos aperfeiçoar, em termos analíticos, é acomodar melhor o excesso de repetições amostrais nas mesmas trilhas. Muitas repetições em uma mesma trilha pode gerar a superdispersão nos dados e prejudicar o ajuste dos modelos. Uma estratégia analítica para usar as repetições na amostragem por distância já foi implementada no pacote unmarked em modelos dinâmicos para múltiplas estações de amostragem, através da função distsampOpen(). A incorporação dessa função no distanceMonitoraflorestal deve melhorar o ajuste dos modelos e levar a estimativas de abundância e densidades mais acuradas.

Conclusão

Apresentamos uma breve introdução sobre como realizar as estimativas de abundância e densidade de aves e mamíferos, a partir de dados do Programa Monitora obtidos por amostragem por distância, utilizando o pacote distanceMonitoraflorestal no R. A partir dos dados de uma espécie de primata, demonstramos como o pacote oferece funcionalidades, todas em português, para operacionalizar as

etapas necessárias para a estimativa de abundância e densidade, como a seleção da proporção mais adequada para truncamento dos dados, o ajuste e a seleção da função de detecção e dos termos de ajuste. Também demonstramos como utilizar o pacote para avaliar diferentes modelos de forma comparativa e apresentar graficamente os resultados obtidos.

Existem outros pacotes disponíveis que podem ser utilizados para gerar estimativas de abundância e densidade, como o pacote unmarked e o próprio pacote Distance, utilizado como base para compilação deste pacote. Contudo, essa é a única ferramenta analítica desenvolvida para todas as etapas de análise dos dados do Programa Monitora, desde a exploração dos dados até a obtenção de séries históricas das dinâmicas de densidade populacional, documentada em língua portuguesa e pensada para um público menos familiarizado com o R. Sua documentação encontra-se disponível no site <https://vntborgesjr.github.io/distanceMonitoraflorestal/> e traz tutoriais que podem ser utilizados como referência para análise de diferentes conjuntos de dados, além da própria base de dados do Programa Monitora. A facilidade de acesso e aplicação do pacote faz dele uma excelente ferramenta para uso na geração de resultados e esperamos que ele seja amplamente utilizado para produção de relatórios técnicos, pesquisas acadêmicas e cursos sobre estimativas de abundância e densidade a partir da amostragem por distância.

Referências

1. Pimm SL, Jenkins CN, Abell R, Brooks TM, Gittleman JL, Joppa LN et al. The biodiversity of species and their rates of extinction, distribution, and protection. *Science*. 2014; 344. doi:10.1126/science.1246752.
2. Dirzo R, Young HS, Galetti M, Ceballos G, Isaac NJB, Collen B. Defaunation in the Anthropocene. *Science*. 2014; 345: 401-406.
3. Ceballos G, Ehrlich PR, Barnosky AD, García A, Pringle RM, Palmer TM. Accelerated modern human-induced species losses: Entering the sixth mass extinction. *Sci Adv*. 2015; 1. doi:10.1126/sciadv.1400253.
4. Cowie RH, Bouchet P, Fontaine B. The Sixth Mass Extinction: fact, fiction or speculation? *Biological Reviews*. 2022; 97: 640-663.
5. O'Connor B, Bojinski S, Röösli C, Schaepman ME. Monitoring global changes in biodiversity and climate essential as ecological crisis intensifies. *Ecol Inform*. 2020; 55. doi:10.1016/j.ecoinf.2019.101033.

6. WWF. Relatório Planeta Vivo 2022: Construindo uma sociedade positiva para a natureza. WWF: Gland, Suíça.; 2022.
7. Westveer J, Freeman R, McRae L, Marconi V, Almond REA, Grooten M. A deep dive into the Living Planet Index: a technical report. WWF: Gland, Switzerland.; 2022.
8. ICMBio/MMA. Estratégia do Programa Nacional de Monitoramento da Biodiversidade. Programa Monitora. Estrutura, articulações, perspectivas. Brasília; 2018.
9. ICMBio/MMA. Estrutura do Programa Monitora. Brasília; 2022.
10. ICMBio/MMA. Programa Monitora. Relatório Anual de Implementação. Brasília; 2023.
11. R-Core-Team. R: a language and environment for statistical computing. R Foundation for Statistical Computing: Vienna, Austria; 2024. Disponível em: <http://www.r-project.org>.
12. ICMBio/MMA. Monitoramento da biodiversidade: roteiro metodológico de aplicação. Brasília; 2014. Disponível em: www.icmbio.gov.br.
13. Buckland ST, Anderson DR, Burnham KP, Laake JL. Distance sampling: estimating abundance of biological populations. Springer-Science+Business Media, B. V.: London; 1993.
14. Laake JL, Buckland ST, Anderson DR, Burnham KP. DISTANCE User's Guide; 1993.
15. Thomas L, Buckland ST, Rexstad EA, Laake JL, Strindberg S, Hedley SL et al. Distance software: design and analysis of distance sampling surveys for estimating population size. *Journal of Applied Ecology*. 2010; 47: 5-14.
16. Kellner KF, Smith AD, Royle JA, Kéry M, Belant JL, Chandler RB. The unmarked R package: Twelve years of advances in occurrence and abundance modelling in ecology. *Methods Ecol Evol*. 2023; 14: 1408-1415.
17. Laake JL, Borchers DL, Thomas L, Miller DL, Bishop JRB. mrds: Mark-Recapture Distance Sampling. R package version 228; 2022.
18. McDonald T, Carlisle J, McDonald A, Nielson R, Augustine B, Griswald J et al. Rdistance: Distance Sampling Analyses. R package version 213 2019. [acesso em: 20 mar 2023]. Disponível em: <https://cran.r-project.org/src/contrib/Archive/Rdistance/>.
19. Miller DL, Rexstad E, Thomas L, Laake JL, Marshall L. Distance sampling in R. *J Stat Softw*. 2019; 89: 1-28.
20. Buckland S, Anderson D, Burnham K, Laake J, Borchers D, Thomas L. Advanced distance sampling: estimating abundance of biological populations. Oxford University Press: Oxford; 2004.
21. Buckland ST, Rexstad EA, Marques TA, Oedekoven CS. Methods in statistical ecology distance sampling: methods and applications. Springer International Publishing Switzezrland: New York; 2015. Disponível em: <http://www.springer.com/series/10235>.
22. Buckland ST, Anderson DR, Burnham KP, Laake JL, Borchers DL, Thomas L. Introduction to distance sampling: estimating abundance of biological populations. Oxford University Press: United Kingdom; 2001.
23. Marques FFC, Buckland ST. Incorporating covariates into standard line transect analyses. *Biometrics*. 2003; 59: 924-935.
24. Marques TA, Thomas L, Fancy SG, Buckland ST. Improving estimates of bird density using multiple-covariate distance sampling. *Auk*. 2007; 124: 1229-1243.
25. Buckland ST, Plumptre AJ, Thomas L, Rexstad EA. Line transect sampling of primates: can animal-to-observer distance methods work? *Int J Primatol*. 2010; 31: 485-499.

Biodiversidade Brasileira – BioBrasil.

Fluxo Contínuo e Edições Temáticas:

- Sustentabilidade da Araucária
- Programa Nacional de Monitoramento da Biodiversidade – Programa Monitora n.2, 2025

<http://www.icmbio.gov.br/revistaelectronica/index.php/BioBR>

Biodiversidade Brasileira é uma publicação eletrônica científica do Instituto Chico Mendes de Conservação da Biodiversidade (ICMBio) que tem como objetivo fomentar a discussão e a disseminação de experiências em conservação e manejo, com foco em unidades de conservação e espécies ameaçadas.

ISSN: 2236-2886

