

Package ‘modeldata’

August 22, 2025

Title Data Sets Useful for Modeling Examples

Version 1.5.1

Description Data sets used for demonstrating or testing model-related packages are contained in this package.

License MIT + file LICENSE

URL <https://modeldata.tidymodels.org>,
<https://github.com/tidymodels/modeldata>

BugReports <https://github.com/tidymodels/modeldata/issues>

Depends R (>= 4.1)

Imports cli, dplyr, MASS, purrr, rlang (>= 1.1.0), tibble

Suggests covr, ggplot2, testthat (>= 3.0.0)

Config/Needs/website tidyverse/tidytemplate, tidymodels/tidymodels

Config/testthat/edition 3

Config/usethis/last-upkeep 2025-04-27

Encoding UTF-8

LazyData true

LazyDataCompression xz

RoxygenNote 7.3.2

NeedsCompilation no

Author Max Kuhn [aut, cre],
Posit Software, PBC [cph, fnd] (ROR: <<https://ror.org/03wc8by49>>)

Maintainer Max Kuhn <max@posit.co>

Repository CRAN

Date/Publication 2025-08-22 06:20:02 UTC

Contents

ad_data	3
ames	4
attrition	5
biomass	5
bivariate	6
car_prices	6
cat_adoption	7
cells	8
check_times	8
chem_proc_yield	10
Chicago	11
concrete	11
covers	12
credit_data	12
crickets	13
deliveries	13
drinks	14
grants	15
hepatic_injury_qsar	16
hotel_rates	17
hpc_cv	17
hpc_data	18
ischemic_stroke	19
leaf_id_flavia	20
lending_club	23
meats	23
mlc_churn	24
oils	25
parabolic	25
pathology	26
pd_speech	26
penguins	27
permeability_qsar	28
Sacramento	29
scat	29
sim_classification	30
small_fine_foods	35
Smithsonian	36
solubility_test	36
stackoverflow	37
steroidogenic_toxicity	37
tate_text	39
taxi	39
two_class_dat	40
two_class_example	40
wa_churn	41

ad_data	<i>Alzheimer's disease data</i>
---------	---------------------------------

Description

Alzheimer's disease data

Details

Craig-Schapiro et al. (2011) describe a clinical study of 333 patients, including some with mild (but well-characterized) cognitive impairment as well as healthy individuals. CSF samples were taken from all subjects. The goal of the study was to determine if subjects in the early states of impairment could be differentiated from cognitively healthy individuals. Data collected on each subject included:

- Demographic characteristics such as age and gender
- Apolipoprotein E genotype
- Protein measurements of Abeta, Tau, and a phosphorylated version of Tau (called pTau)
- Protein measurements of 124 exploratory biomarkers, and
- Clinical dementia scores

For these analyses, we have converted the scores to two classes: impaired and healthy. The goal of this analysis is to create classification models using the demographic and assay data to predict which patients have early stages of disease.

Value

ad_data a tibble

Source

Kuhn, M., Johnson, K. (2013) *Applied Predictive Modeling*, Springer.

Craig-Schapiro R, Kuhn M, Xiong C, Pickering EH, Liu J, Misko TP, et al. (2011) Multiplexed Immunoassay Panel Identifies Novel CSF Biomarkers for Alzheimer's Disease Diagnosis and Prognosis. PLoS ONE 6(4): e18850.

Examples

```
data(ad_data)
str(ad_data)
```

ames

Ames Housing Data

Description

A data set from De Cock (2011) has 82 fields were recorded for 2,930 properties in Ames IA. This version is copies from the AmesHousing package but does not include a few quality columns that appear to be outcomes rather than predictors.

Details

See this links for the sources below for more information as well as `?AmesHousing::make_ames`.

For these data, the training materials typically use:

```
library(tidymodels)

set.seed(4595)
data_split <- initial_split(ames, strata = "Sale_Price")
ames_train <- training(data_split)
ames_test  <- testing(data_split)

set.seed(2453)
ames_folds<- vfold_cv(ames_train)
```

Value

ames a tibble

Source

De Cock, D. (2011). "Ames, Iowa: Alternative to the Boston Housing Data as an End of Semester Regression Project," *Journal of Statistics Education*, Volume 19, Number 3.

<https://jse.amstat.org/v19n3/decock/DataDocumentation.txt>

<https://jse.amstat.org/v19n3/decock.pdf>

Examples

```
data(ames)
str(ames)
```

attrition	<i>Job attrition</i>
-----------	----------------------

Description

Job attrition

Details

These data are from the IBM Watson Analytics Lab. The website describes the data with “Uncover the factors that lead to employee attrition and explore important questions such as ‘show me a breakdown of distance from home by job role and attrition’ or ‘compare average monthly income by education and attrition’. This is a fictional data set created by IBM data scientists.”. There are 1470 rows.

Value

attrition a data frame

Source

The IBM Watson Analytics Lab website <https://www.ibm.com/communities/analytics/watson-analytics-blog/hr-employee-attrition/>

Examples

```
data(attrition)
str(attrition)
```

biomass	<i>Biomass data</i>
---------	---------------------

Description

Ghugare et al (2014) contains a data set where different biomass fuels are characterized by the amount of certain molecules (carbon, hydrogen, oxygen, nitrogen, and sulfur) and the corresponding higher heating value (HHV). These data are from their Table S.2 of the Supplementary Materials

Value

biomass a data frame

Source

Ghugare, S. B., Tiwary, S., Elangovan, V., and Tambe, S. S. (2013). Prediction of Higher Heating Value of Solid Biomass Fuels Using Artificial Intelligence Formalisms. *BioEnergy Research*, 1-12.

Examples

```
data(biomass)
str(biomass)
```

bivariate	<i>Example bivariate classification data</i>
-----------	--

Description

Example bivariate classification data

Details

These data are a simplified version of the segmentation data contained in caret. There are three columns: A and B are predictors and the column Class is a factor with levels "One" and "Two". There are three data sets: one for training (n = 1009), validation (n = 300), and testing (n = 710).

Value

bivariate_train, bivariate_test, bivariate_val
tibbles

Examples

```
data(bivariate)
str(bivariate_train)
str(bivariate_val)
str(bivariate_test)
```

car_prices	<i>Kelly Blue Book resale data for 2005 model year GM cars</i>
------------	--

Description

Kuiper (2008) collected data on Kelly Blue Book resale data for 804 GM cars (2005 model year).

Value

car_prices	data frame of the suggested retail price (column Price) and various characteristics of each car (columns Mileage, Cylinder, Doors, Cruise, Sound, Leather, Buick, Cadillac, Chevy, Pontiac, Saab, Saturn, convertible, coupe, hatchback, sedan and wagon)
------------	---

Source

Kuiper, S. (2008). Introduction to Multiple Regression: How Much Is Your Car Worth?, *Journal of Statistics Education*, Vol. 16 https://jse.amstat.org/jse_archive.htm#2008.

Examples

```
data(car_prices)
str(car_prices)
```

cat_adoption	<i>Cat Adoption</i>
--------------	---------------------

Description

A subset of the cats at the animal shelter in Long Beach, California, USA.

Details

A data frame with 2257 rows and 19 columns:

- time** The time the cat spent at the shelter.
- event** The event of interest is the cat being homed or returned to its original location (i.e., owner or community). The non-event is the cat being transferred to another shelter or dying. Zero indicates a non-event (censored), and one corresponds to the event occurring.
- sex** The sex of the cat.
- neutered** Whether the cat is neutered.
- intake_condition** The intake condition of the cat.
- intake_type** The type of intake.
- latitude** Latitude of the intersection/cross street of intake or capture.
- longitude** Longitude of the intersection/cross street of intake or capture.
- black,brown,brown_tabby,calico,cream,gray,gray_tabby,orange,orange_tabby,tan,tortie,white** Indicators for the color/pattern of the cat's fur.

Value

tibble

Source

<https://data.longbeach.gov/explore/dataset/animal-shelter-intakes-and-outcomes/information/> on 2024-06-17

Examples

```
str(cat_adoption)
```

cells	<i>Cell body segmentation</i>
-------	-------------------------------

Description

Hill, LaPan, Li and Haney (2007) develop models to predict which cells in a high content screen were well segmented. The data consists of 119 imaging measurements on 2019. The original analysis used 1009 for training and 1010 as a test set (see the column called case).

Details

The outcome class is contained in a factor variable called class with levels "PS" for poorly segmented and "WS" for well segmented.

The raw data used in the paper can be found at the Biomedcentral website. The version contained in cells is modified. First, several discrete versions of some of the predictors (with the suffix "Status") were removed. Second, there are several skewed predictors with minimum values of zero (that would benefit from some transformation, such as the log). A constant value of 1 was added to these fields: avg_inten_ch_2, fiber_align_2_ch_3, fiber_align_2_ch_4, spot_fiber_count_ch_4 and total_inten_ch_2.

Value

cells a tibble

Source

Hill, LaPan, Li and Haney (2007). Impact of image segmentation on high-content screening data quality for SK-BR-3 cells, *BMC Bioinformatics*, Vol. 8, pg. 340, <https://bmcbioinformatics.biomedcentral.com/articles/10.1186/1471-2105-8-340>.

Examples

```
data(cells)
str(cells)
```

check_times	<i>Execution time data</i>
-------------	----------------------------

Description

These data were collected from the CRAN web page for 13,626 R packages. The time to complete the standard package checking routine was collected. In some cases, the package checking process is stopped due to errors and these data are treated as censored. It is less than 1 percent.

Details

As predictors, the associated package source code were downloaded and parsed to create predictors, including

- authors: The number of authors in the author field.
- imports: The number of imported packages.
- suggests: The number of packages suggested.
- depends: The number of hard dependencies.
- Roxygen: a binary indicator for whether Roxygen was used for documentation.
- gh: a binary indicator for whether the URL field contained a GitHub link.
- rforge: a binary indicator for whether the URL field contained a link to R-forge.
- descr: The number of characters (or, in some cases, bytes) in the description field.
- r_count: The number of R files in the R directory.
- r_size: The total disk size of the R files.
- ns_import: Estimated number of imported functions or methods.
- ns_export: Estimated number of exported functions or methods.
- s3_methods: Estimated number of S3 methods.
- s4_methods: Estimated number of S4 methods.
- doc_count: How many Rmd or Rnw files in the vignettes directory.
- doc_size: The disk size of the Rmd or Rnw files.
- src_count: The number of files in the src directory.
- src_size: The size on disk of files in the src directory.
- data_count: The number of files in the data directory.
- data_size: The size on disk of files in the data directory.
- testthat_count: The number of files in the testthat directory.
- testthat_size: The size on disk of files in the testthat directory.
- check_time: The time (in seconds) to run R CMD check using the "r-devel-windows-ix86+x86_64" flavor.
- status: An indicator for whether the tests completed.

Data were collected on 2019-01-20.

Value

check_times a data frame

Source

CRAN

Examples

```
data(check_times)
str(check_times)
```

chem_proc_yield

*Chemical manufacturing process data set***Description**

A data set that models yield as a function of biological material predictors and chemical structure predictors.

Details

This data set contains information about a chemical manufacturing process, in which the goal is to understand the relationship between the process and the resulting final product yield. Raw material in this process is put through a sequence of 27 steps to generate the final pharmaceutical product. The starting material is generated from a biological unit and has a range of quality and characteristics. The objective in this project was to develop a model to predict percent yield of the manufacturing process. The data set consisted of 177 samples of biological material for which 57 characteristics were measured. Of the 57 characteristics, there were 12 measurements of the biological starting material, and 45 measurements of the manufacturing process. The process variables included measurements such as temperature, drying time, washing time, and concentrations of by-products at various steps. Some of the process measurements can be controlled, while others are observed. Predictors are continuous, count, categorical; some are correlated, and some contain missing values. Samples are not independent because sets of samples come from the same batch of biological starting material.

Columns:

- yield: numeric
- bio_material_01 - bio_material_12: numeric
- man_proc_01 - man_proc_45: numeric

Value

chem_proc_yield
a tibble

Source

Kuhn, Max, and Kjell Johnson. *Applied predictive modeling*. New York: Springer, 2013.

Examples

```
data(chem_proc_yield)
str(chem_proc_yield)
```

Chicago	<i>Chicago ridership data</i>
---------	-------------------------------

Description

Chicago ridership data

Details

These data are from Kuhn and Johnson (2020) and contain an *abbreviated* training set for modeling the number of people (in thousands) who enter the Clark and Lake L station.

The date column corresponds to the current date. The columns with station names (Austin through California) are a *sample* of the columns used in the original analysis (for file size reasons). These are 14 day lag variables (i.e. date - 14 days). There are columns related to weather and sports team schedules.

The station at 35th and Archer is contained in the column Archer_35th to make it a valid R column name.

Value

Chicago	a tibble
stations	a vector of station names

Source

Kuhn and Johnson (2020), *Feature Engineering and Selection*, Chapman and Hall/CRC . <https://bookdown.org/max/FES/> and <https://github.com/topepo/FES>

Examples

```
data(Chicago)
str(Chicago)
stations
```

concrete	<i>Compressive strength of concrete mixtures</i>
----------	--

Description

Yeh (2006) describes an aggregated data set for experimental designs used to test the compressive strength of concrete mixtures. The data are used by Kuhn and Johnson (2013).

Value

concrete	a tibble
----------	----------

Source

Yeh I (2006). "Analysis of Strength of Concrete Using Design of Experiments and Neural Networks." *Journal of Materials in Civil Engineering*, 18, 597-604.

Kuhn, M., Johnson, K. (2013) *Applied Predictive Modeling*, Springer.

Examples

```
data(concrete)
str(concrete)
```

covers	<i>Raw cover type data</i>
--------	----------------------------

Description

These data are raw data describing different types of forest cover-types from the UCI Machine Learning Database (see link below). There is one column in the data that has a few difference pieces of textual information (of variable lengths).

Value

covers a data frame

Source

<https://archive.ics.uci.edu/ml/machine-learning-databases/covtype/covtype.info>

Examples

```
data(covers)
str(covers)
```

credit_data	<i>Credit data</i>
-------------	--------------------

Description

These data are from the website of Dr. Lluís A. Belanche Muñoz by way of a github repository of Dr. Gaston Sanchez. One data point is a missing outcome was removed from the original data.

Value

credit_data a data frame

Source

<https://github.com/gastonstat/CreditScoring>, <http://bit.ly/2kkBFrk>

Examples

```
data(credit_data)
str(credit_data)
```

crickets	<i>Rates of Cricket Chirps</i>
----------	--------------------------------

Description

These data are from McDonald (2009), by way of Mangiafico (2015), on the relationship between the ambient temperature and the rate of cricket chirps per minute. Data were collected for two species of the genus *Oecanthus*: *O. exclamationis* and *O. niveus*. The data are contained in a data frame called `crickets` with a total of 31 data points.

Value

`crickets` a tibble

Source

Mangiafico, S. 2015. "An R Companion for the Handbook of Biological Statistics." <https://rcompanion.org/handbook/>.

McDonald, J. 2009. *Handbook of Biological Statistics*. Sparky House Publishing.

Examples

```
data(crickets)
str(crickets)
```

deliveries	<i>Food Delivery Time Data</i>
------------	--------------------------------

Description

Food Delivery Time Data

Details

These data are from a study of food delivery times in minutes (i.e., the time from the initial order to receiving the food) for a single restaurant. The data contains 10,012 orders from a specific restaurant. The predictors include:

- The time, in decimal hours, of the order.
- The day of the week for the order.
- The approximate distance in miles between the restaurant and the delivery location.
- A set of 27 predictors that count the number of distinct menu items in the order.

No times are censored.

Value

deliveries a tibble

Examples

```
data(deliveries)
str(deliveries)
```

drinks

Sample time series data

Description

Sample time series data

Details

Drink sales. The exact name of the series from FRED is: "Merchant Wholesalers, Except Manufacturers' Sales Branches and Offices Sales: Nondurable Goods: Beer, Wine, and Distilled Alcoholic Beverages Sales"

Value

drinks a tibble

Source

The Federal Reserve Bank of St. Louis website <https://fred.stlouisfed.org/series/S4248SM144NCEN>

Examples

```
data(drinks)
str(drinks)
```

grants*Grant acceptance data*

Description

A data set related to the success or failure of academic grants.

Details

The data are discussed in Kuhn and Johnson (2013):

"These data are from a 2011 Kaggle competition sponsored by the University of Melbourne where there was interest in predicting whether or not a grant application would be accepted. Since public funding of grants had decreased over time, triaging grant applications based on their likelihood of success could be important for estimating the amount of potential funding to the university. In addition to predicting grant success, the university sought to understand factors that were important in predicting success."

The data ranged from 2005 and 2008 and the data spending strategy was driven by the date of the grant. Kuhn and Johnson (2013) describe:

"The compromise taken here is to build models on the pre-2008 data and tune them by evaluating a random sample of 2,075 grants from 2008. Once the optimal parameters are determined, final model is built using these parameters and the entire training set (i.e., the data prior to 2008 and the additional 2,075 grants). A small holdout set of 518 grants from 2008 will be used to ensure that no gross methodology errors occur from repeatedly evaluating the 2008 data during model tuning. In the text, this set of samples is called the 2 0 0 8 holdout set. This small set of year 2008 grants will be referred to as the test set and will not be evaluated until set of candidate models are identified."

To emulate this, `grants_other` contains the training (pre-2008, $n = 6,633$) and holdout/validation data (2008, $n = 1,557$). `grants_test` has 518 grant samples from 2008. The object `grants_2008` is an integer vector that can be used to separate the modeling with the holdout/validation sets.

Value

`grants_other`, `grants_test`, `grants_2008`

two tibbles and an integer vector of data points used for training

Source

Kuhn and Johnson (2013). *Applied Predictive Modeling*. Springer.

Examples

```
data(grants)
str(grants_other)
str(grants_test)
str(grants_2008)
```

hepatic_injury_qsar *Predicting hepatic injury from chemical information*

Description

A quantitative structure-activity relationship (QSAR) data set to predict when a molecule has risk associated with liver function.

Details

This data set was used to develop a model for predicting compounds' probability of causing hepatic injury (i.e. liver damage). This data set consisted of 281 unique compounds; 376 predictors were measured or computed for each. The response was categorical (either "none", "mild", or "severe"), and was highly unbalanced.

This kind of response often occurs in pharmaceutical data because companies steer away from creating molecules that have undesirable characteristics. Therefore, well-behaved molecules often greatly outnumber undesirable molecules. The predictors consisted of measurements from 184 biological screens and 192 chemical feature predictors. The biological predictors represent activity for each screen and take values between 0 and 10 with a mode of 4. The chemical feature predictors represent counts of important sub-structures as well as measures of physical properties that are thought to be associated with hepatic injury.

Columns:

- class: ordered and factor (levels: 'none', 'mild', and 'severe')
- bio_assay_001 - bio_assay_184: numeric
- chem_fp_001 - chem_fp_192: numeric

Value

hepatic_injury_qsar
a tibble

Source

Kuhn, Max, and Kjell Johnson. *Applied predictive modeling*. New York: Springer, 2013.

Examples

```
data(hepatic_injury_qsar)
str(hepatic_injury_qsar)
```

hotel_rates*Daily Hotel Rate Data*

Description

A data set to predict the average daily rate for a hotel in Lisbon Portugal.

Details

Data are originally described in Antonio, de Almeida, and Nunes (2019). This version of the data is filtered for one hotel (the "Resort Hotel") and is intended as regression data set for predicting the average daily rate for a room. The data are post-2016; the 2016 data were used to have a predictor for the historical daily rates. See the `hotel_rates.R` file in the `data-raw` directory of the package to understand other filters used when creating this version of the data.

The agent and company fields were changed from random characters to use a set of random names. The outcome column is `avg_price_per_room`.

License:

No license was given for the data; See the reference below for source.

Source

<https://github.com/rfordatascience/tidytuesday/tree/master/data/2020/2020-02-11>

References

Antonio, N., de Almeida, A., and Nunes, L. (2019). Hotel booking demand datasets. *Data in Brief*, 22, 41-49.

Examples

```
## Not run:  
str(hotel_rates)  
  
## End(Not run)
```

hpc_cv*Class probability predictions*

Description

Class probability predictions

Details

This data frame contains the predicted classes and class probabilities for a linear discriminant analysis model fit to the HPC data set from Kuhn and Johnson (2013). These data are the assessment sets from a 10-fold cross-validation scheme. The data column columns for the true class (obs), the class prediction (pred) and columns for each class probability (columns VF, F, M, and L). Additionally, a column for the resample indicator is included.

Value

hpc_cv a data frame

Source

Kuhn, M., Johnson, K. (2013) *Applied Predictive Modeling*, Springer

Examples

```
data(hpc_cv)
str(hpc_cv)
```

hpc_data	<i>High-performance computing system data</i>
----------	---

Description

Kuhn and Johnson (2013) describe a data set where characteristics of unix jobs were used to classify there completion times as either very fast (1 min or less, VF), fast (1–50 min, F), moderate (5–30 min, M), or long (greater than 30 min, L).

Value

hpc_data a tibble

Source

Kuhn, M., Johnson, K. (2013) *Applied Predictive Modeling*, Springer.

Examples

```
data(hpc_data)
str(hpc_data)
```

ischemic_stroke

*Clinical data used to predict ischemic stroke***Description**

A data set to predict a binary outcome using imaging and patient data.

Details

These data were gathered to predict patient risk for ischemic stroke. A historical set of patients with a range of carotid artery blockages were selected. The data consisted of 126 patients, 44 of which had blockages greater than 70%. All patients had undergone Computed Tomography Angiography (CTA) to generate a detailed three-dimensional visualization and characterization of the blockage. These images were then analyzed in order to compute several features related to the disease, including: percent stenosis, arterial wall thickness, and tissue characteristics such as lipid-rich necrotic core and calcification.

The group of patients in this study also had follow-up information on whether or not a stroke occurred at a subsequent point in time. The data for each patient also included commonly collected clinical characteristics for risk of stroke such as whether or not the patient had atrial fibrillation, coronary artery disease, and a history of smoking. Demographics of gender and age were included as well. These readily available risk factors can be thought of as another potentially useful predictor set that can be evaluated. In fact, this set of predictors should be evaluated first to assess their ability to predict stroke since these predictors are easy to collect, are acquired at patient presentation, and do not require an expensive imaging technique.

Columns:

- stroke: factor (levels: 'yes' and 'no')
- nascet_scale: numeric
- calc_vol: numeric
- calc_vol_prop: numeric
- matx_vol: numeric
- matx_vol_prop: numeric
- lrnc_vol: numeric
- lrnc_vol_prop: numeric
- max_calc_area: numeric
- max_calc_area_prop: numeric
- max_dilation_by_area: numeric
- max_matx_area: numeric
- max_matx_area_prop: numeric
- max_lrnc_area: numeric
- max_lrnc_area_prop: numeric
- max_max_wall_thickness: numeric

- max_remodeling_ratio: numeric
- max_stenosis_by_area: numeric
- max_wall_area: numeric
- wall_vol: numeric
- max_stenosis_by_diameter: numeric
- age: integer
- male: integer
- smoking_history: integer
- atrial_fibrillation: integer
- coronary_artery_disease: integer
- diabetes_history: integer
- hypercholesterolemia_history: integer
- hypertension_history: integer

Value

ischemic_stroke
a tibble

Source

Kuhn, Max, and Kjell Johnson. *Feature Engineering and Selection: A Practical Approach for Predictive Models*. Chapman and Hall/CRC, 2019.

Examples

```
data(ischemic_stroke)
str(ischemic_stroke)
```

leaf_id_flavia	<i>Leaf identification data (Flavia)</i>
----------------	--

Description

Image analysis of leaves to predict species.

Details

From the original manuscript: "The Flavia dataset contains 1907 leaf images. There are 32 different species and each has 50-77 images. Scanners and digital cameras are used to acquire the leaf images on a plain background. The isolated leaf images contain blades only, without a petiole. These leaf images are collected from the most common plants in Yangtze, Delta, China. Those leaves were sampled on the campus of the Nanjing University and the Sun Yat-Sen arboretum, Nanking, China."

The reference below has details information on the features used for prediction.

Columns:

- species: factor (32 levels)
- apex: factor (9 levels)
- base: factor (6 levels)
- shape: factor (5 levels)
- denate_edge: factor (levels: 'no' and 'yes')
- lobed_edge: factor (levels: 'no' and 'yes')
- smooth_edge: factor (levels: 'no' and 'yes')
- toothed_edge: factor (levels: 'no' and 'yes')
- undulate_edge: factor (levels: 'no' and 'yes')
- outlying_polar: numeric
- skewed_polar: numeric
- clumpy_polar: numeric
- sparse_polar: numeric
- striated_polar: numeric
- convex_polar: numeric
- skinny_polar: numeric
- stringy_polar: numeric
- monotonic_polar: numeric
- outlying_contour: numeric
- skewed_contour: numeric
- clumpy_contour: numeric
- sparse_contour: numeric
- striated_contour: numeric
- convex_contour: numeric
- skinny_contour: numeric
- stringy_contour: numeric
- monotonic_contour: numeric
- num_max_points: numeric
- num_min_points: numeric
- diameter: numeric

- area: numeric
- perimeter: numeric
- physiological_length: numeric
- physiological_width: numeric
- aspect_ratio: numeric
- rectangularity: numeric
- circularity: numeric
- compactness: numeric
- narrow_factor: numeric
- perimeter_ratio_diameter: numeric
- perimeter_ratio_length: numeric
- perimeter_ratio_lw: numeric
- num_convex_points: numeric
- perimeter_convexity: numeric
- area_convexity: numeric
- area_ratio_convexity: numeric
- equivalent_diameter: numeric
- eccentricity: numeric
- contrast: numeric
- correlation_texture: numeric
- inverse_difference_moments: numeric
- entropy: numeric
- mean_red_val: numeric
- mean_green_val: numeric
- mean_blue_val: numeric
- std_red_val: numeric
- std_green_val: numeric
- std_blue_val: numeric
- correlation: numeric

Value

leaf_id_flavia a data frame

Source

Lakshika, Jayani PG, and Thiyanga S. Talagala. "Computer-aided interpretable features for leaf image classification." *arXiv preprint* arXiv:2106.08077 (2021).

https://github.com/SMART-Research/leaffeatures_paper

Examples

```
data(leaf_id_flavia)
str(leaf_id_flavia)
```

lending_club

Loan data

Description

Loan data

Details

These data were downloaded from the Lending Club access site (see below) and are from the first quarter of 2016. A subset of the rows and variables are included here. The outcome is in the variable Class and is either "good" (meaning that the loan was fully paid back or currently on-time) or "bad" (charged off, defaulted, of 21-120 days late). A data dictionary can be found on the source website.

Value

lending_club a data frame

Source

Lending Club Statistics <https://www.lendingclub.com/info/download-data.action>

Examples

```
data(lending_club)
str(lending_club)
```

meats

Fat, water and protein content of meat samples

Description

"These data are recorded on a Tecator Infratec Food and Feed Analyzer working in the wavelength range 850 - 1050 nm by the Near Infrared Transmission (NIT) principle. Each sample contains finely chopped pure meat with different moisture, fat and protein contents.

Details

If results from these data are used in a publication we want you to mention the instrument and company name (Tecator) in the publication. In addition, please send a preprint of your article to:

Karin Thente, Tecator AB, Box 70, S-263 21 Hoganas, Sweden

The data are available in the public domain with no responsibility from the original data source. The data can be redistributed as long as this permission note is attached."

"For each meat sample the data consists of a 100 channel spectrum of absorbances and the contents of moisture (water), fat and protein. The absorbance is -log10 of the transmittance measured by the spectrometer. The three contents, measured in percent, are determined by analytic chemistry."

Included here are the training, monitoring and test sets.

Value

meats a tibble

Examples

```
data(meats)
str(meats)
```

mlc_churn

Customer churn data

Description

A data set from the MLC++ machine learning software for modeling customer churn. There are 19 predictors, mostly numeric: state (categorical), account_length, area_code, international_plan (yes/no), voice_mail_plan (yes/no), number_vmail_messages, total_day_minutes, total_day_calls, total_day_charge, total_eve_minutes, total_eve_calls, total_eve_charge, total_night_minutes, total_night_calls, total_night_charge, total_intl_minutes, total_intl_calls, total_intl_charge, and number_customer_service_calls.

Details

The outcome is contained in a column called churn (also yes/no). A note in one of the source files states that the data are "artificial based on claims similar to real world".

Value

mlc_churn a tibble

Source

Originally at <http://www.sgi.com/tech/mlc/>

Examples

```
data(mlc_churn)
str(mlc_churn)
```

oils*Fatty acid composition of commercial oils*

Description

Fatty acid concentrations of commercial oils were measured using gas chromatography. The data is used to predict the type of oil. Note that only the known oils are in the data set. Also, the authors state that there are 95 samples of known oils. However, we count 96 in Table 1 (pgs. 33-35).

Value

oils a tibble

Source

Brodnjak-Voncina et al. (2005). Multivariate data analysis in classification of vegetable oils characterized by the content of fatty acids, *Chemometrics and Intelligent Laboratory Systems*, Vol. 75:31-45.

Examples

```
data(oils)
str(oils)
```

parabolic*Parabolic class boundary data*

Description

Parabolic class boundary data

Details

These data were simulated. There are two correlated predictors and two classes in the factor outcome.

Value

parabolic a data frame

Examples

```
data(parabolic)
str(parabolic)
```

pathology	<i>Liver pathology data</i>
-----------	-----------------------------

Description

Liver pathology data

Details

These data have the results of a x-ray examination to determine whether liver is abnormal or not (in the scan column) versus the more extensive pathology results that approximate the truth (in pathology).

Value

pathology a data frame

Source

Altman, D.G., Bland, J.M. (1994) "Diagnostic tests 1: sensitivity and specificity," *British Medical Journal*, vol 308, 1552.

Examples

```
data(pathology)
str(pathology)
```

pd_speech	<i>Parkinson's disease speech classification data set</i>
-----------	---

Description

Parkinson's disease speech classification data set

Details

From the UCI ML archive, the description is "The data used in this study were gathered from 188 patients with PD (107 men and 81 women) with ages ranging from 33 to 87 (65.1 p/m 10.9) at the Department of Neurology in Cerrahpaşa Faculty of Medicine, Istanbul University. The control group consists of 64 healthy individuals (23 men and 41 women) with ages varying between 41 and 82 (61.1 p/m 8.9). During the data collection process, the microphone is set to 44.1 KHz and following the physician's examination, the sustained phonation of the vowel /a/ was collected from each subject with three repetitions."

The data here are averaged over the replicates.

Value

pd_speech a data frame

Source

UCI ML repository (data) <https://archive.ics.uci.edu/ml/datasets/Parkinson%27s+Disease+Classification#>, Sakar et al (2019), "A comparative analysis of speech signal processing algorithms for Parkinson's disease classification and the use of the tunable Q-factor wavelet transform", *Applied Soft Computing*, V74, pg 255-263.

Examples

```
data(pd_speech)
str(pd_speech)
```

penguins	<i>Palmer Station penguin data</i>
----------	------------------------------------

Description

A data set from Gorman, Williams, and Fraser (2014) containing measurements from different types of penguins. This version of the data was retrieved from Allison Horst's palmerpenguins package on 2020-06-22.

Value

penguins a tibble

Source

Gorman KB, Williams TD, Fraser WR (2014) Ecological Sexual Dimorphism and Environmental Variability within a Community of Antarctic Penguins (*Genus Pygoscelis*). PLoS ONE 9(3): e90081. doi:10.1371/journal.pone.0090081

<https://github.com/allisonhorst/palmerpenguins>

Examples

```
data(penguins)
str(penguins)
```

permeability_qsar*Predicting permeability from chemical information*

Description

A quantitative structure-activity relationship (QSAR) data set to predict when a molecule can permeate cells.

Details

This pharmaceutical data set was used to develop a model for predicting compounds' permeability. In short, permeability is the measure of a molecule's ability to cross a membrane. The body, for example, has notable membranes between the body and brain, known as the blood-brain barrier, and between the gut and body in the intestines. These membranes help the body guard critical regions from receiving undesirable or detrimental substances. For an orally taken drug to be effective in the brain, it first must pass through the intestinal wall and then must pass through the blood-brain barrier in order to be present for the desired neurological target. Therefore, a compound's ability to permeate relevant biological membranes is critically important to understand early in the drug discovery process. Compounds that appear to be effective for a particular disease in research screening experiments, but appear to be poorly permeable may need to be altered in order improve permeability, and thus the compound's ability to reach the desired target. Identifying permeability problems can help guide chemists towards better molecules.

Permeability assays such as PAMPA and Caco-2 have been developed to help measure compounds' permeability (Kansy et al, 1998). These screens are effective at quantifying a compound's permeability, but the assay is expensive labor intensive. Given a sufficient number of compounds that have been screened, we could develop a predictive model for permeability in an attempt to potentially reduce the need for the assay. In this project there were 165 unique compounds; 1107 molecular fingerprints were determined for each. A molecular fingerprint is a binary sequence of numbers that represents the presence or absence of a specific molecular sub-structure. The response is highly skewed, the predictors are sparse (15.5% are present), and many predictors are strongly associated.

Columns:

- permeability: numeric
- chem_fp_0001 - chem_fp_1107: numeric

Value

permeability_qsar
a data frame

Source

Kuhn, Max, and Kjell Johnson. *Applied predictive modeling*. New York: Springer, 2013.

Examples

```
data(permeability_qsar)
str(permeability_qsar)
```

Sacramento*Sacramento CA home prices*

Description

This data frame contains house and sale price data for 932 homes in Sacramento CA. The original data were obtained from the website for the SpatialKey software. From their website: "The Sacramento real estate transactions file is a list of 985 real estate transactions in the Sacramento area reported over a five-day period, as reported by the Sacramento Bee." Google was used to fill in missing/incorrect data.

Value

Sacramento a tibble

Source

SpatialKey website: <https://support.spatialkey.com/spatialkey-sample-csv-data/>

Examples

```
data(Sacramento)
str(Sacramento)
```

scat*Morphometric data on scat*

Description

Reid (2015) collected data on animal feses in coastal California. The data consist of DNA verified species designations as well as fields related to the time and place of the collection and the scat itself. The data are on the three main species.

Value

scat a tibble

Source

Reid, R. E. B. (2015). A morphometric modeling approach to distinguishing among bobcat, coyote and gray fox scats. *Wildlife Biology*, 21(5), 254-262

Examples

```
data(scats)  
str(scats)
```

sim_classification	<i>Simulate datasets</i>
--------------------	--------------------------

Description

These functions can be used to generate simulated data for supervised (classification and regression) and unsupervised modeling applications.

Usage

```
sim_classification(  
  num_samples = 100,  
  method = "caret",  
  intercept = -5,  
  num_linear = 10,  
  keep_truth = FALSE  
)
```

```
sim_regression(  
  num_samples = 100,  
  method = "sapp_2014_1",  
  std_dev = NULL,  
  factors = FALSE,  
  keep_truth = FALSE  
)
```

```
sim_noise(  
  num_samples,  
  num_vars,  
  cov_type = "exchangeable",  
  outcome = "none",  
  num_classes = 2,  
  cov_param = 0  
)
```

```
sim_logistic(num_samples, eqn, correlation = 0, keep_truth = FALSE)
```

```
sim_multinomial(  
  num_samples,  
  eqn_1,  
  eqn_2,  
  eqn_3,  
  correlation = 0,
```

```
    keep_truth = FALSE
  )
```

Arguments

num_samples	Number of data points to simulate.
method	A character string for the simulation method. For classification, the single current option is "caret". For regression, values can be "sapp_2014_1", "sapp_2014_2", "van_der_laam_2007_1", "van_der_laam_2007_2", "hooker_2004", or "worley_1987". See Details below.
intercept	The intercept for the linear predictor.
num_linear	Number of diminishing linear effects.
keep_truth	A logical: should the true outcome value be retained for the data? If so, the column name is .truth.
std_dev	Gaussian distribution standard deviation for residuals. Default values are shown below in Details.
factors	A single logical for whether the binary indicators should be encoded as factors or not.
num_vars	Number of noise predictors to create.
cov_type	The multivariate normal correlation structure of the predictors. Possible values are "exchangeable" and "toeplitz".
outcome	A single character string for what type of independent outcome should be simulated (if any). The default value of "none" produces no extra columns. Using "classification" will generate a class column with num_classes values, equally distributed. A value of "regression" results in an outcome column that contains independent standard normal values.
num_classes	When outcome = "classification", the number of classes to simulate.
cov_param	A single numeric value for the exchangeable correlation value or the base of the Toeplitz structure. See Details below.
eqn, eqn_1, eqn_2, eqn_3	An R expression or (one sided) formula that only involves variables A and B that is used to compute the linear predictor. External objects should not be used as symbols; see the examples below on how to use external objects in the equations.
correlation	A single numeric value for the correlation between variables A and B.

Details

Specific Regression and Classification methods:

These functions provide several supervised simulation methods (and one unsupervised). Learn more by method:

method = "caret":

This is a simulated classification problem with two classes, originally implemented in `caret::twoClassSim()` with all numeric predictors. The predictors are simulated in different sets. First, two multivariate normal predictors (denoted here as `two_factor_1` and `two_factor_2`) are created with a correlation of about 0.65. They change the log-odds using main effects and an interaction:

intercept - 4 * two_factor_1 + 4 * two_factor_2 + 2 * two_factor_1 * two_factor_2

The intercept is a parameter for the simulation and can be used to control the amount of class imbalance.

The effects in the second set are linear with coefficients that alternate signs and have a sequence of values between 2.5 and 0.25. For example, if there were four predictors in this set, their contribution to the log-odds would be

-2.5 * linear_1 + 1.75 * linear_2 -1.00 * linear_3 + 0.25 * linear_4

(Note that these column names may change based on the value of num_linear).

The third set is a nonlinear function of a single predictor ranging between [0, 1] called non_linear_1 here:

(non_linear_1^3) + 2 * exp(-6 * (non_linear_1 - 0.3)^2)

The fourth set of informative predictors is copied from one of Friedman's systems and use two more predictors (non_linear_2 and non_linear_3):

2 * sin(non_linear_2 * non_linear_3)

All of these effects are added up to model the log-odds.

method = "sapp_2014_1":

This regression simulation is from Sapp et al. (2014). There are 20 independent Gaussian random predictors with mean zero and a variance of 9. The prediction equation is:

predictor_01 + sin(predictor_02) + log(abs(predictor_03)) +
predictor_04^2 + predictor_05 * predictor_06 +
ifelse(predictor_07 * predictor_08 * predictor_09 < 0, 1, 0) +
ifelse(predictor_10 > 0, 1, 0) + predictor_11 * ifelse(predictor_11 > 0, 1, 0) +
sqrt(abs(predictor_12)) + cos(predictor_13) + 2 * predictor_14 + abs(predictor_15) +
ifelse(predictor_16 < -1, 1, 0) + predictor_17 * ifelse(predictor_17 < -1, 1, 0) -
2 * predictor_18 - predictor_19 * predictor_20

The error is Gaussian with mean zero and variance 9.

method = "sapp_2014_2":

This regression simulation is also from Sapp et al. (2014). There are 200 independent Gaussian predictors with mean zero and variance 16. The prediction equation has an intercept of one and identical linear effects of log(abs(predictor)).

The error is Gaussian with mean zero and variance 25.

method = "van_der_laen_2007_1":

This is a regression simulation from van der Laan et al. (2007) with ten random Bernoulli variables that have a 40% probability of being a value of one. The true regression equation is:

2 * predictor_01 * predictor_10 + 4 * predictor_02 * predictor_07 +
3 * predictor_04 * predictor_05 - 5 * predictor_06 * predictor_10 +
3 * predictor_08 * predictor_09 + predictor_01 * predictor_02 * predictor_04 -
2 * predictor_07 * (1 - predictor_06) * predictor_02 * predictor_09 -
4 * (1 - predictor_10) * predictor_01 * (1 - predictor_04)

The error term is standard normal.

method = "van_der_laen_2007_2":

This is another regression simulation from van der Laan et al. (2007) with twenty Gaussians with mean zero and variance 16. The prediction equation is:

predictor_01 * predictor_02 + predictor_10^2 - predictor_03 * predictor_17 -
predictor_15 * predictor_04 + predictor_09 * predictor_05 + predictor_19 -
predictor_20^2 + predictor_09 * predictor_08

The error term is also Gaussian with mean zero and variance 16.

method = "hooker_2004":

Hooker (2004) and Sorokina *et al* (2008) used the following:

$$\begin{aligned} & \pi^{\sqrt{(predictor_01 * predictor_02) * \sqrt{(2 * predictor_03)} -}} \\ & \text{asin}(predictor_04) + \log(predictor_03 + predictor_05) - \\ & (predictor_09 / predictor_10) * \sqrt{(predictor_07 / predictor_08) -} \\ & predictor_02 * predictor_07 \end{aligned}$$

Predictors 1, 2, 3, 6, 7, and 9 are standard uniform while the others are uniform on $[0.6, 1.0]$. The errors are normal with mean zero and default standard deviation of 0.25.

method = "worley_1987":

The simulation system from Worley (1987) is based on a mechanistic model for the flow rate of liquids from two aquifers positioned vertically (i.e., the "upper" and "lower" aquifers). There are two sets of predictors:

- the borehole radius (radius_borehole from 0.05 to 0.15) and length (length_borehole from 1,120 to 1,680) .
- The radius of effect that the system has on collecting water (radius_influence from 100 to 50,000)

and physical properties:

- transmissibility_upper_aq
- potentiometric_upper_aq
- transmissibility_lower_aq
- potentiometric_lower_aq
- conductivity_borehole

A multiplicative error structure is used; the mechanistic equation is multiplied by an exponentiated Gaussian random error.

The references give feasible ranges for each of these variables. See also Morris *et al* (1993).

sim_noise():

This function simulates a number of random normal variables with mean zero. The values can be independent if cov_param = 0. Otherwise the values are multivariate normal with non-diagonal covariance matrices. For cov_type = "exchangeable", the structure has unit variances and covariances of cov_param. With cov_type = "toeplitz", the covariances have an exponential pattern (see example below).

Logistic simulation:

sim_logistic() provides a flexible interface to simulating a logistic regression model with two multivariate normal variables A and B (with zero mean, unit variances and correlation determined by the correlation argument).

For example, using eqn = A + B would specify that the true probability of the event was

$$\text{prob} = 1 / (1 + \exp(A + B))$$

The class levels for the outcome column are "one" and "two".

Multinomial simulation:

sim_multinomial() can generate data with classes "one", "two", and "three" based on the values in arguments eqn_1, eqn_2, and eqn_3, respectively. Like [sim_logistic\(\)](#) these equations use predictors A and B.

The individual equations are evaluated and exponentiated. After this, their values are, for each row of data, normalized to add up to one. These probabilities are then passed to `stats::rmultinom()` to generate the outcome values.

References

- Hooker, G. (2004, August). Discovering additive structure in black box functions. In *Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining* (pp. 575-580). DOI: 10.1145/1014052.1014122
- Morris, M. D., Mitchell, T. J., and Ylvisaker, D. (1993). Bayesian design and analysis of computer experiments: use of derivatives in surface prediction. *Technometrics*, 35(3), 243-255.
- Sapp, S., van der Laan, M. J., and Canny, J. (2014). Subsemble: an ensemble method for combining subset-specific algorithm fits. *Journal of applied statistics*, 41(6), 1247-1259. DOI: 10.1080/02664763.2013.864263
- Sorokina, D., Caruana, R., Riedewald, M., and Fink, D. (2008, July). Detecting statistical interactions with additive groves of trees. In *Proceedings of the 25th international conference on Machine learning* (pp. 1000-1007). DOI: 10.1145/1390156.1390282
- Van der Laan, M. J., Polley, E. C., and Hubbard, A. E. (2007). Super learner. *Statistical applications in genetics and molecular biology*, 6(1). DOI: 10.2202/1544-6115.1309.
- Worley, B. A. (1987). Deterministic uncertainty analysis (No. ORNL-6428). Oak Ridge National Lab.(ORNL), Oak Ridge, TN/

Examples

```
set.seed(1)
sim_regression(100)
sim_classification(100)

# Flexible logistic regression simulation
if (rlang::is_installed("ggplot2")) {
  library(dplyr)
  library(ggplot2)

  sim_logistic(1000, ~ .1 + 2 * A - 3 * B + 1 * A * B, corr = .7) |>
    ggplot(aes(A, B, col = class)) +
    geom_point(alpha = 1/2) +
    coord_equal()

  f_xor <- ~ 10 * xor(A > 0, B < 0)
  # or
  f_xor <- rlang::expr(10 * xor(A > 0, B < 0))

  sim_logistic(1000, f_xor, keep_truth = TRUE) |>
    ggplot(aes(A, B, col = class)) +
    geom_point(alpha = 1/2) +
    coord_equal() +
    theme_bw()
}

## How to use external symbols:
```

```
a_coef <- 2
# splice the value in using rlang's !! operator
lp_eqn <- rlang::expr(!!a_coef * A+B)
lp_eqn
sim_logistic(5, lp_eqn)

# Flexible multinomial regression simulation
if (rlang::is_installed("ggplot2")) {

}
```

small_fine_foods	<i>Fine foods example data</i>
------------------	--------------------------------

Description

Fine foods example data

Details

These data are from Amazon, who describe it as "This dataset consists of reviews of fine foods from amazon. The data span a period of more than 10 years, including all ~500,000 reviews up to October 2012. Reviews include product and user information, ratings, and a plaintext review."

A subset of the data are contained here and are split into a training and test set. The training set sampled 10 products and retained all of their individual reviews. Since the reviews within these products are correlated, we recommend resampling the data using a leave-one-product-out approach. The test set sampled 500 products that were not included in the training set and selected a single review at random for each.

There is a column for the product, a column for the text of the review, and a factor column for a class variable. The outcome is whether the reviewer gave the product a 5-star rating or not.

Value

training_data, testing_data
tibbles

Source

<https://snap.stanford.edu/data/web-FineFoods.html>

Examples

```
data(small_fine_foods)
str(training_data)
str(testing_data)
```

Smithsonian	<i>Smithsonian museums</i>
-------------	----------------------------

Description

Geocodes for the Smithsonian museums (circa 2018).

Value

Smithsonian a tibble

Source

https://en.wikipedia.org/wiki/List_of_Smithsonian_museums

Examples

```
data(Smithsonian)
str(Smithsonian)
```

solubility_test	<i>Solubility predictions from MARS model</i>
-----------------	---

Description

Solubility predictions from MARS model

Details

For the solubility data in Kuhn and Johnson (2013), these data are the test set results for the MARS model. The observed solubility (in column solubility) and the model results (prediction) are contained in the data.

Value

solubility_test
a data frame

Source

Kuhn, M., Johnson, K. (2013) *Applied Predictive Modeling*, Springer

Examples

```
data(solubility_test)
str(solubility_test)
```

stackoverflow*Annual Stack Overflow Developer Survey Data*

Description

Annual Stack Overflow Developer Survey Data

Details

These data are a collection of 5,594 data points collected on developers. These data could be used to try to predict who works remotely (as used in the source listed below).

Value

stackoverflow a tibble

Source

Julia Silge, *Supervised Machine Learning Case Studies in R*

<https://supervised-ml-course.netlify.com/chapter2>

Raw data: <https://insights.stackoverflow.com/survey/>

Examples

```
data(stackoverflow)
str(stackoverflow)
```

steroidogenic_toxicity*Predicting steroidogenic toxicity with assay data*

Description

A set of *in vitro* assays are used to quantify the risk of reproductive toxicity via the disruption of steroidogenic pathways.

Details

H295R cells were used to measure the effect with two sets of assay results. The first includes a set of protein measurements on: cytochrome P450 enzymes ("cyp"s), STAR, and 3BHSD2. The second include hormone measurements for DHEA, progesterone, testosterone, and cortisol.

Columns:

- class: factor (levels: 'toxic' and 'nontoxic')
- cyp_11a1: numeric

- cyp_11b1: numeric
- cyp_11b2: numeric
- cyp_17a1: numeric
- cyp_19a1: numeric
- cyp_21a1: numeric
- hsd3b2: numeric
- star: numeric
- progesterone: numeric
- testosterone: numeric
- dhea: numeric
- cortisol: numeric

Value

A tibble with columns

- class: factor(levels: toxic and nontoxic)
- cyp_11a1: numeric
- cyp_11b1: numeric
- cyp_11b2: numeric
- cyp_17a1: numeric
- cyp_19a1: numeric
- cyp_21a1: numeric
- hsd3b2: numeric
- star: numeric
- progesterone: numeric
- testosterone: numeric
- dhea: numeric
- cortisol: numeric

Source

Maglich, J. M., Kuhn, M., Chapin, R. E., & Pletcher, M. T. (2014). More than just hormones: H295R cells as predictors of reproductive toxicity. *Reproductive Toxicology*, 45, 77-86.

Examples

```
data(steroidogenic_toxicity)
str(steroidogenic_toxicity)
```

tate_text	<i>Tate Gallery modern artwork metadata</i>
-----------	---

Description

Metadata such as artist, title, and year created for recent artworks owned by the Tate Gallery. Only artworks created during or after 1990 are included, and the metadata source was last updated in 2014. The Tate Gallery provides these data but requests users to be respectful of their [guidelines for use](#).

Value

tate_text a tibble

Source

- <https://github.com/tategallery/collection>
- <https://www.tate.org.uk/>

Examples

```
data(tate_text)
str(tate_text)
```

taxi	<i>Chicago taxi data set</i>
------	------------------------------

Description

A data set containing information on a subset of taxi trips in the city of Chicago in 2022.

Details

The source data are originally described on the linked City of Chicago data portal. The data exported here are a pre-processed subset motivated by the modeling problem of predicting whether a rider will tip or not.

tip Whether the rider left a tip. A factor with levels "yes" and "no".

distance The trip distance, in odometer miles.

company The taxi company, as a factor. Companies that occurred few times were binned as "other".

local Whether the trip's starting and ending locations are in the same community. See the source data for community area values.

dow The day of the week in which the trip began, as a factor.

month The month in which the trip began, as a factor.

hour The hour of the day in which the trip began, as a numeric.

Value

tibble

Source

<https://data.cityofchicago.org/Transportation/Taxi-Trips/wrvz-psew>

Examples

taxi

two_class_dat	<i>Two class data</i>
---------------	-----------------------

Description

Two class data

Details

There are artificial data with two predictors (A and B) and a factor outcome variable (Class).

Value

two_class_dat a data frame

Examples

```
data(two_class_dat)
str(two_class_dat)
```

two_class_example	<i>Two class predictions</i>
-------------------	------------------------------

Description

Two class predictions

Details

These data are a test set form a model built for two classes ("Class1" and "Class2"). There are columns for the true and predicted classes and column for the probabilities for each class.

Value

two_class_example
a data frame

Examples

```
data(two_class_example)
str(two_class_example)
```

wa_churn	<i>Watson churn data</i>
----------	--------------------------

Description

Watson churn data

Details

These data were downloaded from the IBM Watson site (see below) in September 2018. The data contain a factor for whether a customer churned or not. Alternatively, the tenure column presumably contains information on how long the customer has had an account. A survival analysis can be done on this column using the churn outcome as the censoring information. A data dictionary can be found on the source website.

Value

wa_churn a data frame

Source

IBM Watson Analytics <https://ibm.co/2sOvyvy>

Examples

```
data(wa_churn)
str(wa_churn)
```

Index

* datasets

- ad_data, [3](#)
- ames, [4](#)
- attrition, [5](#)
- biomass, [5](#)
- bivariate, [6](#)
- car_prices, [6](#)
- cat_adoption, [7](#)
- cells, [8](#)
- check_times, [8](#)
- Chicago, [11](#)
- concrete, [11](#)
- covers, [12](#)
- credit_data, [12](#)
- crickets, [13](#)
- deliveries, [13](#)
- drinks, [14](#)
- grants, [15](#)
- hotel_rates, [17](#)
- hpc_cv, [17](#)
- hpc_data, [18](#)
- lending_club, [23](#)
- meats, [23](#)
- mlc_churn, [24](#)
- oils, [25](#)
- parabolic, [25](#)
- pathology, [26](#)
- pd_speech, [26](#)
- penguins, [27](#)
- Sacramento, [29](#)
- scat, [29](#)
- small_fine_foods, [35](#)
- Smithsonian, [36](#)
- solubility_test, [36](#)
- stackoverflow, [37](#)
- tate_text, [39](#)
- two_class_dat, [40](#)
- two_class_example, [40](#)
- wa_churn, [41](#)
- ad_data, [3](#)
- ames, [4](#)
- attrition, [5](#)
- biomass, [5](#)
- bivariate, [6](#)
- bivariate_test (bivariate), [6](#)
- bivariate_train (bivariate), [6](#)
- bivariate_val (bivariate), [6](#)
- car_prices, [6](#)
- caret::twoClassSim(), [37](#)
- cat_adoption, [7](#)
- cells, [8](#)
- check_times, [8](#)
- chem_proc_yield, [10](#)
- Chicago, [11](#)
- concrete, [11](#)
- covers, [12](#)
- credit_data, [12](#)
- crickets, [13](#)
- deliveries, [13](#)
- drinks, [14](#)
- grants, [15](#)
- grants_2008 (grants), [15](#)
- grants_other (grants), [15](#)
- grants_test (grants), [15](#)
- hepatic_injury_qsar, [16](#)
- hotel_rates, [17](#)
- hpc_cv, [17](#)
- hpc_data, [18](#)
- ischemic_stroke, [19](#)
- leaf_id_flavia, [20](#)
- lending_club, [23](#)
- meats, [23](#)

mlc_churn, 24

oils, 25

parabolic, 25

pathology, 26

pd_speech, 26

penguins, 27

permeability_qsar, 28

Sacramento, 29

scat, 29

sim_classification, 30

sim_logistic(sim_classification), 30

sim_logistic(), 33

sim_multinomial(sim_classification), 30

sim_noise(sim_classification), 30

sim_regression(sim_classification), 30

small_fine_foods, 35

Smithsonian, 36

solubility_test, 36

stackoverflow, 37

stations(Chicago), 11

stats::rmultinom(), 34

steroidogenic_toxicity, 37

tate_text, 39

taxi, 39

testing_data(small_fine_foods), 35

training_data(small_fine_foods), 35

two_class_dat, 40

two_class_example, 40

wa_churn, 41